



Bipolar argumentative semantics for t -norm based fuzzy logics

Esther Anna Corsi¹

Received: 13 May 2023 / Accepted: 28 August 2025
© The Author(s) 2025

Abstract

We investigate the relations between bipolar logical argumentation frames and the three main t -norm-based fuzzy logics: *Lukasiewicz* ($\mathbf{L}$), *Gödel* ($\mathbf{G}$) and *product* logic ($\mathbf{II}$). The arguments that we consider are complex entities, and the inference relation between the support and the claim is instantiated with the standard consequence relation of one specific fuzzy logic. We introduce several *argumentative principles* defined in terms of the attack and the support relation that refine the existence of such relations whenever the involved arguments share some propositional formulas. Through the notion of argumentative immunity introduced in Corsi and Fermüller (2018a) we finally recover complete semantics for $\mathbf{L}$, $\mathbf{G}$ and $\mathbf{II}$.

Keywords Argumentative Semantics · Bipolar Argumentation Frameworks · Fuzzy Logic · Logical Argumentation Theory · T -norm · T -norm-based Argumentation Theory

1 Introduction

In his seminal work (Dung 1995), Dung introduces abstract argumentation frameworks. Dung conceives the arguments as completely abstract entities related to each other by a binary relation intuitively understood as an attack relation. Therefore, abstract argumentation frameworks can be represented as directed graphs in which the nodes represent the arguments and the edges represent the attack relations between them. In the present work, we consider a variation of Dung's original argumentative setting that differs in the definition of the arguments and the relations that subsist among them.

Many are the works aimed at deductively formalising Dung's work (García and Simari 2004; Amgoud and Besnard 2009; Rahwan and Simari 2009; Amgoud and Besnard 2010; Amgoud et al. 2011; Besnard and Hunter 2001; Goriannis and Hunter 2011). Following the same approach, we understand the arguments as complex entities made of three parts: the *support*, the *claim* and the *method of inference* that makes explicit how the claim follows from the support. In

addition, in the argumentative frameworks we will consider, the arguments are related not only with a “negative” relation, the attack relation but also with a “positive” one, the support relation. Argumentation frameworks with both the attack and the support relations are called *bipolar* (Cayrol and Lagasque-Schiex 2005b; Amgoud et al. 2008; Boella et al. 2010; Cayrol and Lagasque-Schiex 2013; Proietti et al. 2019; Prakken et al. 2020; Yu and Van der Torre 2020; Gonzalez et al. 2021).

Suppose argument A is attacked by a set of arguments and supported by a different set of arguments. In that case, we might still be deciding whether to consider A a reasonable point of view, assigning it value 1 as per *acceptance of A* , or not, assigning it value 0 as per *not acceptance of A* . Thus, intermediate acceptance levels might be necessary to formalise a position of indecision regarding a given argument's acceptance status.

Since fuzzy logic relies on the idea that truth comes in degrees (Cintula et al. 2011), we might think that the connection between argumentation frameworks, in our specific case bipolar argumentation frameworks, and fuzzy logic is straightforward. However, even if we consider argumentative frameworks with only the attack relation, their relation with classical and some non-classical propositional logic is only sometimes clear (see Besnard and Hunter (2008)). Sequent-based argumentation frameworks represent one

✉ Esther Anna Corsi
esther.corsi@unimi.it

¹ LUCI Lab, Department of Philosophy, University of Milan,
Via Festa del Perdono 7, Milan 20122, Italy

possible way to establish such a connection (see Arieli and Straßer (2015)).

In the present work, as the method of inference between the support set and the claims of arguments, we consider the consequence relations relative to the main three t-norm based fuzzy logics (Cintula et al. 2011): Łukasiewicz (?), Gödel (G) and product logic (Π). In these new frameworks, we introduce *argumentative principles* defined using both attack and support relations. As the attack principles in Corsi and Fermüller (2017, 2018b, 2018a), the argumentative principles refine the existence (or non-existence) of the attack and support relations once the arguments involved share, in their claims, some or even all propositional formulas. Following the same approach as Corsi and Fermüller (2018a), we define *argumentative immunity* in bipolar t-norm based argumentative frameworks. Thus, through the notion of argumentative immunity jointly with specific sets of argumentative principles, we recover sound and complete semantics for L, G and Π.

It is important to point out that we are not claiming that fuzzy logic is the logic of bipolar argumentation frameworks. Instead, we investigate under which (argumentative) conditions it is possible to recover new semantics for the above-mentioned logics.

With the concept of immunity, as pointed out in Corsi and Fermüller (2018a), we have an alternative notion of logical validity not based on Tarskian-style semantics but that only refers to t-norm based arguments and specific restrictions on the definition of the attack and support relations. Thus, the present work aims to contribute to the literature on alternative semantics for fuzzy logics like, e.g. Giles's game-based semantics for Łukasiewicz logic (Giles 1982, 1977), the voting semantics by Lawry (1998) and the acceptability semantics by Paris (1997).

From the point of view of argumentation theory, possible applications of the argumentative frameworks and principles introduced remain an open line of research for future investigations.

The paper is organised as follows. In Sect. 2, we recall some notions of abstract and logical argumentation theory. In Sect. 3, we give a glimpse of the three main t-norm based fuzzy logics. The first contribution of the paper can be found in Sect. 4 with the introduction of bipolar t-norm based logical argumentation frameworks and bipolar principles. In Sect. 5, we prove completeness theorems for L, G and Π. Finally, in Sect. 6, we discuss related work, and in Sect. 7, we conclude.

2 Abstract and logical argumentation theory

Dung's abstract argumentation frameworks are defined as a set of arguments and a binary relation, the attack relation, defined over it Dung (1995). In Dung's setting arguments have no inner structure, they are abstract entities and also the attack relation is not instantiated in any specific way. These very general frameworks can be used to model various scenarios in Artificial Intelligence and related fields.

Definition 1 (Argumentation framework (AF) (Dung 1995)) An argumentation framework is a pair $\mathcal{A} = \langle Ar, \mathcal{R}_{\rightarrow} \rangle$ where:

1. Ar is a set of arguments, and
2. $\mathcal{R}_{\rightarrow}$ is a binary relation over Ar ($\mathcal{R}_{\rightarrow} \subseteq Ar \times Ar$).

Given an argumentation framework $\mathcal{A} = \langle Ar, \mathcal{R}_{\rightarrow} \rangle$, we say that $S \subseteq Ar$ *attacks* an argument $b \in Ar$ if there exists $a \in S$ such that $(a, b) \in \mathcal{R}_{\rightarrow}$. In the sequel, whenever an argument a attacks an argument b , instead of writing $(a, b) \in \mathcal{R}_{\rightarrow}$ we will use a more intuitive notation and write $a \rightarrow b$. We will also use $\rightarrow$ to denote $\mathcal{R}_{\rightarrow}$ itself. $S \subseteq Ar$ *defends* an argument a whenever S attacks b for every argument b that attacks a .

To enlarge the expressive power and consequently possible applications of abstract argumentation theory, frameworks with also a positive relation among arguments, called *support*, have been investigated (Karacapilidis and Papadias 2001; Verheij 2002; Amgoud et al. 2008; Cayrol and Lagasque-Schiex 2013); these are the *bipolar* argumentative frameworks.

Definition 2 (Bipolar Argumentation Framework (BAF)) A bipolar argumentation framework is a triplet $\mathcal{A} = \langle Ar, \mathcal{R}_{\rightarrow}, \mathcal{R}_{\Rightarrow} \rangle$ such that:

- Ar is a set of arguments,
- $\mathcal{R}_{\rightarrow}$ is the *attack* relation and it is a binary relation over Ar ($\mathcal{R}_{\rightarrow} \subseteq Ar \times Ar$),
- $\mathcal{R}_{\Rightarrow}$ is the *support* relation and it is a binary relation over Ar ($\mathcal{R}_{\Rightarrow} \subseteq Ar \times Ar$).

As for AFs, instead of writing $(a, b) \in \mathcal{R}_{\Rightarrow}$, we use the notation $a \Rightarrow b$. The symbols $\nrightarrow$ and $\nRightarrow$ stand for “not attacking” and “not supporting”.

In some specific situations, support and attack are dependent notions. For example, if in a given argumentation framework there is an argument b that defends and argument a , then we infer that b *supports* a . However, whenever the arguments are instantiated with actual propositions, this

way of reasoning does not always apply, as the following example shows.

Example 1 (Cayrol and Lagasque-Schiex (2005a)) There is a family that would like to go hiking. They prefer sunny weather to a cloudy one and cloudy weather to a rainy one. If it is rainy, then they will not go hiking. Otherwise, they will go. However, clouds could be a sign of rain. They look at the sky early in the morning, and it is cloudy.

The following exchange of informal arguments occurs:

- a: Today it is a holiday, and we go hiking.
- b: The weather will be cloudy; clouds are a sign of rain, so we should cancel the hiking.
- c: These clouds are early patches of mist; the day will be sunny without clouds, and the weather will be not cloudy.
- d: Clouds will not grow. Thus, the weather will be cloudy but not rainy.

In Example 1, we have that $d \longrightarrow c \longrightarrow b \longrightarrow a$. Even just by looking at the graphical representation of the example in Fig. 1, we see that the argument d defends b and c defends a . However, both c and d support the idea of going on a hike. Therefore, considering chains of arguments and counter-arguments, counting the links between them and assessing the odd ones as attacks and the even ones as supports is an oversimplification of the definition of support. Thus, in bipolar argumentation frameworks, the definition of the attack relation should be independent of the definition of the support relation. However, through the argumentative principles, as shown in Sect. 5, it is possible to enforce some constraints on the existence of both relations in a given framework.

With the introduction of argumentative principles, we do not pretend to contribute to the theory of bipolar argumentation frameworks directly, but we rather use some concepts of these frameworks to provide a new type of semantic framework for fuzzy logic.

In logical argumentation theory, both the arguments and the relations among them are instantiated. The following definition of *logical argument* has been broadly used in

the literature (Besnard and Hunter 2001; Gorogiannis and Hunter 2011; Amgoud and Cayrol 2002).

Definition 3 (Deductive Argument Besnard and Hunter (2008)) Let $\mathcal{L} = \langle \mathcal{L}, \vdash \rangle$ be a propositional logic ($\mathcal{L}$ stands for the propositional language and $\vdash$ for a Tarskian consequence relation over $\mathcal{L}$). Let Δ be a finite set of propositional formulas in $\mathcal{L}$ (the database). An *argument* based on $\mathcal{L}$ and Δ is a pair $\langle \Gamma, \varphi \rangle$ s.t. $\Gamma \subseteq \Delta$ and:

1. $\Gamma \vdash \varphi$
2. Γ is consistent
3. Γ is a minimal subset of Δ satisfying (1). $S(\langle \Gamma, \varphi \rangle)$ denotes the support of the argument and $C(\langle \Gamma, \varphi \rangle)$ its claim. For all $S \subseteq \Delta$, $Arg(S)$ denotes the set of all arguments that can be built from S .

Although this definition seems to be a good starting point for constructing and identifying arguments, we will now point out some limitations of this approach. As stated in D'Agostino and Modgil (2018), formalisations should be compatible with real-world modes of reasoning and consider that real-world computational abilities are resource-bounded. Deciding whether the support set of a given argument is consistent or not is, in general, and depending on the underlying logic, an NP-complete decision problem (Besnard and Hunter 2006) and determining whether it is minimal is a Π_2^P -complete problem for CL and at least as hard for many other logics (Parsons et al. 2003). These two assumptions are generally computationally unfeasible but feasible in many situations. In fact, in the generating process of the arguments from a given set Δ of $\mathcal{L}$ -formulas, $\subseteq$ -minimality and consistency of the support sets can be checked while constructing the arguments.

An alternative definition of logical argument has been introduced in Arieli (2013) and widely investigated by Arieli and Straßer in Arieli and Straßer (2016, 2015, 2019, 2014, 2020). In Arieli and Straßer's approach, arguments are understood as *sequents*, as introduced by Gentzen (1935), and in their definition, the assumptions about the minimality and consistency of the support are omitted.

In logical argumentation theory, the attack relation can be instantiated in several ways.

Definition 4 (Attack Relations Gorogiannis and Hunter (2011); Besnard and Hunter (2018)) Let Δ be a database on the underlying logic $\mathcal{L} = \langle \mathcal{L}, \vdash \rangle$ and $\langle \Gamma_1, \psi_1 \rangle$ and $\langle \Gamma_2, \psi_2 \rangle$ two deductive (or sequent-based) arguments in $Arg(\Delta)$. We define the following attack relations by listing the conditions

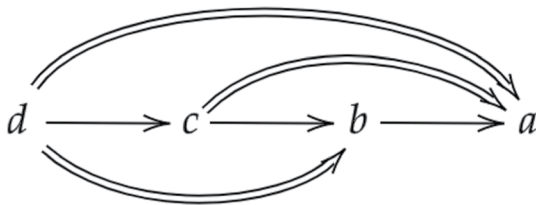


Fig. 1 Graphical representation of Example 1

under which $\langle \Gamma_1, \psi_1 \rangle$ attacks $\langle \Gamma_2, \psi_2 \rangle$. On the left side, we list the symbol for each attack relation.

- [Def] $\langle \Gamma_1, \psi_1 \rangle$ is a *defeater* of $\langle \Gamma_2, \psi_2 \rangle$ if $\psi_1 \vdash \neg \bigwedge \Gamma_2$.
- [Ucut] $\langle \Gamma_1, \psi_1 \rangle$ is a *direct undercut* of $\langle \Gamma_2, \psi_2 \rangle$ if $\psi_1 \vdash \neg \bigwedge \Gamma'_2$ and $\neg \bigwedge \Gamma'_2 \vdash \psi_1$ for some $\Gamma'_2 \subseteq \Gamma_2$.
- [Reb] $\langle \Gamma_1, \psi_1 \rangle$ is a *rebuttal* of $\langle \Gamma_2, \psi_2 \rangle$ if $\psi_1 \vdash \neg \psi_2$ and $\neg \psi_2 \vdash \psi_1$.
- [C-Reb] $\langle \Gamma_1, \psi_1 \rangle$ is a *compact rebuttal* of $\langle \Gamma_2, \psi_2 \rangle$ if $\Gamma_1 \vdash \neg \psi_2$.
- [D-Reb] $\langle \Gamma_1, \psi_1 \rangle$ is a *defeat rebuttal* of $\langle \Gamma_2, \psi_2 \rangle$ if $\psi_1 \vdash \neg \psi_2$.
- [I-Reb] $\langle \Gamma_1, \psi_1 \rangle$ is an *indirect rebuttal* of $\langle \Gamma_2, \psi_2 \rangle$ if there exists $\varphi \in \mathcal{L}$ such that $\psi_1 \vdash \varphi$ and $\psi_2 \vdash \neg \varphi$.

In the existing literature, additional attack relations are defined, but in the present paper, we only consider some of them. A logical argumentation framework is defined as follows.

Definition 5 (Logical Argumentation Framework (LAF)) Let $\mathfrak{L} = \langle \mathcal{L}, \vdash \rangle$ be a propositional logic, Δ a database over $\mathcal{L}$ -formulas and $\mathcal{A}$ a set of attack relations. A *logical argumentation framework* over Δ is a pair $A = \langle \text{Arg}(\Delta), \text{Attack}(\mathcal{A}) \rangle$ where $\text{Arg}(\Delta)$ is the set of all arguments generated by Δ and $\text{Attack}(\mathcal{A}) \subseteq \text{Arg}(\Delta) \times \text{Arg}(\Delta)$ is an attack relation such that $(A_1, A_2) \in \text{Attack}(\mathcal{A})$ iff there is some $\mathcal{R} \in \mathcal{A}$ such that $A_1 \mathcal{R}$ -attacks A_2 .

In logical argumentation theory, the support relation can be instantiated in several ways. The many definitions reflect the different kinds of support among arguments. In the following, only the definition of *direct support* is given, which can be seen as a positive counterpart to *defeat rebuttal*.

Definition 6 (Direct Support - [D-Sup]) Let $\langle \Gamma_1; \psi_1 \rangle$ and $\langle \Gamma_2; \psi_2 \rangle$ be two deductive

arguments in $\text{Arg}(\Delta)$. We say that $\langle \Gamma_1, \psi_1 \rangle$ *directly supports* $\langle \Gamma_2, \psi_2 \rangle$ ($\langle \Gamma_1, \psi_1 \rangle \xrightarrow{[D-Sup]} \langle \Gamma_2, \psi_2 \rangle$) if $\psi_1 \vdash \psi_2$.

We can find many intermediary levels of abstraction between Dung's style argumentation frameworks and logical argumentation frameworks. In Corsi and Fermüller (2017, 2018a), the authors have introduced semi-abstract argumentation frameworks (SAFs). SAFs are logical argumentation frameworks in which the arguments are still understood as complex entities, but the claims are the only part instantiated. In SAFs, *attack principles* (see Corsi and Fermüller (2017)) have been introduced and they can be intuitively seen as rules that refine the existence of the attack relations whenever the arguments involved share some propositional formula. See Fig. 2.

Definition 7 (Semi-Abstract Argumentation framework (SAF))

A *semi-abstract argumentation framework (SAF)* is a pair $S = \langle C(\mathcal{AF}), \mathcal{R}_{\rightarrow} \rangle$, where $C(\mathcal{AF})$ is a set formulas representing claims of arguments of a logical argumentation framework $A = \langle \text{Arg}(\Delta), \text{Attack}(\mathcal{A}) \rangle$, i.e., $C(\mathcal{AF}) = \{ \varphi \mid \langle \Gamma, \varphi \rangle \in \text{Arg}(\Delta) \}$, and $\mathcal{R}_{\rightarrow}$ is an attack relation defined over $C(\mathcal{AF}) \times C(\mathcal{AF})$ such that $(\varphi_1, \varphi_2) \in \mathcal{R}_{\rightarrow}$ iff there is some $\langle \Gamma_1, \varphi_1 \rangle, \langle \Gamma_2, \varphi_2 \rangle$ in $\text{Arg}(\Delta)$ and $\mathcal{R} \in \mathcal{A}$ such that $\langle \Gamma_1, \varphi_1 \rangle \mathcal{R}$ -attacks $\langle \Gamma_2, \varphi_2 \rangle$.

In a *semi-abstract* setting, $X \rightarrow A$ can be interpreted as “there exists an argument with claim X that attacks an argument with claim A .” Let us recall the attack principles.

- (A. $\wedge$) If $X \rightarrow A$ or $X \rightarrow B$ then $X \rightarrow A \wedge B$.
- (A. $\vee$) If $X \rightarrow A \vee B$ then $X \rightarrow A$ and $X \rightarrow B$.
- (A. $\supset$) If $X \rightarrow B$ and $X \not\rightarrow A$ then $X \rightarrow A \supset B$.
- (A. $\neg$) If $X \rightarrow A$ then $X \not\rightarrow \neg A$.

- (C. $\wedge$) If $X \rightarrow A \wedge B$ then $X \rightarrow A$ or $X \rightarrow B$.
- (C. $\vee$) If $X \rightarrow A$ and $X \rightarrow B$ then $X \rightarrow A \vee B$.
- (C. $\supset$) If $X \rightarrow A \supset B$ then $X \rightarrow B$ and $X \not\rightarrow A$.
- (C. $\neg$) If $X \not\rightarrow A$ then $X \rightarrow \neg A$.

In a fully instantiated logical argumentation framework the attack principle (A. $\wedge$) can be read as follows. In a given framework, if there exists an argument with claim X such that it attacks (using a specific attack relation) an argument with claim A , then there are arguments with claim $A \wedge B$ that are also attacked by the argument with claim X . Thus, the attack principles can also be defined in fully instantiated frameworks.

As shown in Corsi and Fermüller (2017), all of the above attack principles can be used, together with the definition of an argumentative entailment relation, to recover a complete semantics for CL (classical logic).

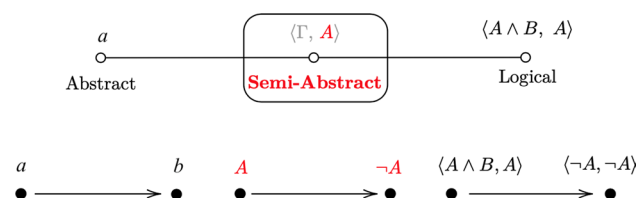


Fig. 2 Levels of abstraction relative to the different argumentation frameworks introduced (Esther Anna 2025)

It is not surprising that some of these attack principles are intuitively very demanding and also, considering the following modal interpretation of the attack relation, not formally justified.

$$\iota(X \longrightarrow A) = \Box(X \wedge \neg A).$$

The attack principles justified by the modal interpretation ι of the attack relations are:

$$\text{MAP} = \{(A.\wedge), (A.\vee), (C.\vee), (C.\supset)*, (A.\neg)\}.$$

Neglecting the problematic attack principles, it is possible to characterise a fragment of the classical sequent calculus in which some of the logical rules are missing. This calculus is referred as LM and consists of the rules for classical logic LK except for $(\neg, r)$, $(\wedge, r)$ and $(\supset, l)$.

3 T-norm based fuzzy logics

We will now recall the basic definitions relative to Łukasiewicz (L), Gödel (G) and product logic (Π), the three fundamental t-norm based fuzzy logics (Hájek 1998).

Definition 8 (t-norm) A *t-norm* (or *triangular-norm*) is a binary function $*$ on $[0, 1]$ such that

1. $*$ is commutative and associative.
2. $*$ is non-decreasing in both arguments.
3. $1 * x = x$ and $0 * x = 0$ for all $x \in [0, 1]$.

A t-norm $*$ is a *continuous t-norm* if it is continuous on $[0, 1]^2$. For an extensive study of t-norms and continuous t-norms see Klement et al. (2013).

The Łukasiewicz, Gödel and product t-norms are, respectively,

$$\begin{aligned} x *_{\mathbf{L}} y &= \max\{0, x + y - 1\}, \\ x *_{\mathbf{G}} y &= \min\{x, y\} \text{ and} \\ x *_{\mathbf{\Pi}} y &= x \cdot y. \end{aligned}$$

Considering the continuous t-norms as truth functions for conjunction, the corresponding truth function for implication can be defined uniquely. The truth function for implication is called *residuum*.

Let $*$ be a continuous t-norm. Then, for any $x, y, z \in [0, 1]$, the operation

$$x \supset_* y = \max\{z \mid x * z \leq y\}$$

is the unique operation satisfying the condition

$$(x * z) \leq y \text{ iff } z \leq (x \supset_* y).$$

Here below, we recall the residuum of $*_{\mathbf{L}}$, $*_{\mathbf{G}}$ and $*_{\mathbf{\Pi}}$.

$$\begin{aligned} x \supset_{*_{\mathbf{L}}} y &= \min\{1, 1 - x + y\} \\ x \supset_{*_{\mathbf{G}}} y &= \begin{cases} 1 & \text{if } x \leq y \\ y & \text{otherwise.} \end{cases} \\ x \supset_{*_{\mathbf{\Pi}}} y &= \begin{cases} 1 & \text{if } x \leq y \\ \frac{y}{x} & \text{otherwise.} \end{cases} \end{aligned}$$

Since for any continuous t-norm $*$ the function $\min$ and $\max$ can be defined in terms of $*$ and $\supset_*^1$, and they extend the bivalent truth-tables for classical conjunction and disjunction on $\{0, 1\}$, respectively, we can consider them the truth-functions of *weak conjunction* ($\wedge$) and *weak disjunction* ($\vee$) while the one interpreted by the t-norm is referred as *strong conjunction* ($\&$). In the case of Gödel logic the two conjunctions coincide, but they differ for all other t-norms. Negation can be defined as $\neg_* x := x \supset_* 0$, the fuzzy interpretation of *reductio ad absurdum*.

The (residual) negation of the three continuous t-norms considered is defined as follows:

$$\neg_{*_{\mathbf{G}}} x = \neg_{*_{\mathbf{\Pi}}} x = \begin{cases} 1 & \text{if } x = 0 \\ 0 & \text{otherwise.} \end{cases}$$

$$\neg_{*_{\mathbf{L}}} x = 1 - x.$$

The functions $*$, $\supset_*$, $\min$ and $\max$ equip the interval $[0, 1]$ with an algebraic structure that can be used for a standard definition of algebraic semantics for fuzzy logic. These algebras are called *t-algebras* (see Cintula et al. (2011)) and we denote them by $[0, 1]^*$. The *t-algebra* of $*$ is the algebra

$$[0, 1]^* = \langle [0, 1], *, \supset_*, \min, \max, 0, 1 \rangle.$$

For any set of K of continuous t-norms we denote the corresponding set of t-algebras by $\mathbb{K}$, and vice versa. Now we can define the syntax and standard semantics of logics based on continuous t-norms as follows.

Definition 9 (Syntax and Standard Semantics of Logics of Continuous T-norms Cintula et al. (2011)) The language $\mathcal{L}$ of the propositional fuzzy logic L_K of continuous t-norms consists of the propositional variables $p, q, r, \dots$, the binary propositional connectives $\&$ (strong conjunctions), $\supset$ (implication), $\wedge$ (weak conjunction), $\vee$ (weak disjunction), the

¹ $\min\{x, y\} = x * (x \supset_* y)$ and $\max\{x, y\} = \min\{(x \supset_* y) \supset_* y, (y \supset_* x) \supset_* x\}.$

unary propositional connective $\neg$ (negation) and the propositional constants $\bar{1}$ (truth) and $\bar{0}$ (falsity).

The set of propositional variables of $\mathcal{L}$ will be denoted by Var and the set of all formulas of $\mathcal{L}$ by $Fm_{\mathcal{L}}$. Capital letters denote formulas and upper case greek letters denote sets of formulas.

A $[0, 1]$ -evaluation of propositional variables is a mapping $e_* : Var \rightarrow [0, 1]$. The evaluation of propositional variables extends uniquely the $*$ -evaluation to all formulas by the following recursive definition. For any proposition variable p and any formula A, B :

$$\begin{aligned} e_*(p) &= e(p) \quad e_*(\bar{0}) = 0 \\ e_*(\bar{1}) &= 1 \quad e_*(\neg A) = \neg_*(e_*(A)) \\ e_*(A \& B) &= e_*(A) * e_*(B) \\ e_*(A \supset B) &= e_*(A) \supset_* e_*(B) \\ e_*(A \wedge B) &= \min\{e_*(A), e_*(B)\} \\ e_*(A \vee B) &= \max\{e_*(A), e_*(B)\} \end{aligned}$$

The *standard consequence relation* of L_K is defined as follows.

Definition 10 (Standard Consequence Relation of L_K Cintula et al. (2011)) A $*$ -evaluation e_* is a $*$ -model of a set Γ of formulas if $e(A) = 1$ for all $A \in \Gamma$. A formula A is a standard consequence of Γ in L_K ($\Gamma \models_K A$) if for each $* \in K$, all $*$ -models e_* of Γ are $*$ -models of A .

A formula A is L_K -valid, $\models_K A$, iff $e_*(A) = 1$ for all $*$ -evaluations e_* .

In general it is not possible to axiomatize the consequence relation $\models_K$ by rules with finitely many premises, but its finitary version² $\models_K^{fin}$ is finitely axiomatizable for any set $\mathbb{K}$ of t-algebras, with *modus ponens* as the only deduction rule.

Definition 11 The logic of a set K of continuous t-norms (or equivalently the logic of $\mathbb{K}$) is identified with the finitary consequence relation $\models_K^{fin}$ and denoted by L_K .

The logics of $*_L$, $*_G$ and $*_{\Pi}$ are, respectively, *Lukasiewicz* (L), *product* (Π) and *Gödel logic* (G). The logic of all continuous t-norms is the *basic logic* (BL).

In Definition 10, “1” is the only designated truth-value and the consequence relation is defined semantically by the truth-preserving paradigm. It is possible to define an alternative consequence relation, the *order-based consequence relation*, that preserves not only truth in its maximal degree,

but all available degrees. Refer to Font (2015) for a detailed introduction to the two consequence relations.

Definition 12 (Order-Based Consequence Relation) A formula A is a *order-based consequence* of Γ in L_K ($\Gamma \models_K^{\leq} A$) if for all evaluation $e_* \inf_{\gamma \in \Gamma} e_*(\gamma) \leq e_*(A)$.

For all three logics L , G and Π we can define a *standard consequence relation* and an *order-based consequence relation*. If we only consider validity, then the two consequence relations coincide, i.e. $\models_K A$ iff $\models_K^{\leq} A$. Note that, *modus ponens* is not sound with the order-based consequence relation.

Axiomatic systems for logics of continuous t-norms can be defined. A Hilbert-style axiomatisation sound and complete for Łukasiewicz Logic is the following.

$$\begin{aligned} [Tr] : & (F \supset G) \supset ((G \supset H) \supset (F \supset H)) \\ [We] : & F \supset (G \supset F) \\ [Ex] : & (F \supset (G \supset H)) \supset (G \supset (F \supset H)) \\ [\wedge-1] : & (F \wedge G) \supset F \\ [\wedge-2] : & (F \wedge G) \supset G \\ [\wedge-3] : & (H \supset F) \supset ((H \supset G) \supset (H \supset (F \wedge G))) \\ [\vee-1] : & F \supset (F \vee G) \\ [\vee-2] : & G \supset (F \vee G) \\ [\vee-3] : & (G \supset F) \supset ((H \supset F) \supset ((G \vee H) \supset F)) \\ [Lin] : & (F \supset G) \vee (G \supset F) \\ [\perp] : & \perp \supset F \\ [Waj] : & ((F \supset G) \supset G) \supset ((G \supset F) \supset F) \end{aligned}$$

together with the deduction rule of *modus ponens*.

$$[MP] : \text{ from } F \text{ and } F \supset G, \text{ infer } G.$$

Theorem 1 The above Hilbert-style system is sound and complete for Łukasiewicz logic. In other words: a formula F is derivable in the system iff F is L -valid.

As axiomatic system for Gödel logic we consider $[Tr]$, $[We]$, $[Ex]$, $[\wedge-1]$, $[\wedge-2]$, $[\wedge-3]$, $[\vee-1]$, $[\vee-2]$, $[\vee-3]$, $[Lin]$, $[\perp]$ plus $[Con]$ and *modus ponens* as inference rule.

$$[Con] : (F \supset (F \supset G)) \supset (F \supset G)$$

As axiomatic system for product logic we consider

$$\begin{aligned} [BL1] : & (F \supset B) \supset ((G \supset H) \supset (F \supset H)) \\ [BL4] : & F \& (F \supset G) \supset G \& (G \supset F) \\ [BL5a] : & (F \& G \supset H) \supset (F \supset (G \supset H)) \\ [BL5b] : & (F \supset (G \supset H)) \supset (F \& G \supset H) \\ [BL6] : & ((F \supset G) \supset H) \supset (((G \supset F) \supset H) \supset H) \\ [BL7] : & \perp \supset F \\ [P] : & \neg F \vee ((F \supset F \& G) \supset G) \end{aligned}$$

² $\Gamma \models_K^{fin} A$ iff there is a finite set $\Gamma' \subseteq \Gamma$ such that $\Gamma' \models_K A$.

together with *modus ponens*.

4 Bipolar T-norm based argumentative frameworks

The consequence relations recalled in Definition 10 and Definition 12 can be considered as the method of inference between a support set and a claim of an argument. Thus, we can recover two definitions of *t-norm based arguments*. In these definitions, as in Arieli and Straßer (2016, 2015, 2019), we do not make any assumption on the support set except for requiring it to be a set of formulas in $Fm_{\mathcal{L}}$.

Definition 13 (Argument based on $\models_K$) Let $\mathcal{L}$ be the language of the propositional fuzzy logic L_K and $Fm_{\mathcal{L}}$ the set its well formed formulas. Let Δ be a finite set of propositional formulas in $Fm_{\mathcal{L}}$. An argument based on $\models_K$ is a pair $\langle \Gamma_X, X \rangle$ such that

1. $\Gamma_X \cup \{X\} \subseteq Fm_{\mathcal{L}}$,
2. $\Gamma_X \subseteq \Delta$ and
3. $\Gamma_X \models_K X$.

Definition 14 (Argument based on $\models_K^{\leq}$)

Let $\mathcal{L}$ be the language of the propositional fuzzy logic L_K and $Fm_{\mathcal{L}}$ the set its well formed formulas. Let Δ be a finite set of propositional formulas in $Fm_{\mathcal{L}}$. An argument based on $\models_K$ is a pair $\langle \Gamma_X, X \rangle$ such that

1. $\Gamma_X \cup \{X\} \subseteq Fm_{\mathcal{L}}$,
2. $\Gamma_X \subseteq \Delta$ and
3. $\Gamma_X \models_K^{\leq} X$.

To better understand the argumentative characteristics of these two ways of defining a t-norm based argument, we consider several attack relations and explore which attack principles are justified, and which are not. Since the attack principles are defined in terms of the claims of the arguments, for this analysis, we consider attack relations defined, at least partially, in terms of the claims of the arguments as well. The attack relations we consider are the following: *defeat*, *compact rebuttal*, *defeating rebuttal* and *indirect rebuttal*.

Following Definition 13, we have that the argument $\langle \Gamma_A, A \rangle$ attacks the argument $\langle \Gamma_B, B \rangle$ using the *defeat* attack relation iff $A \models_K \neg \bigwedge \Gamma_B$, i.e. for every evaluation e such that $e(A) = 1$, then $e(\neg \bigwedge \Gamma_B) = 1$. By Definition 14, the argument $\langle \Gamma_A, A \rangle$ attacks the argument $\langle \Gamma_B, B \rangle$ using the *defeat* attack relation whenever $A \models_K^{\leq} \neg \bigwedge \Gamma_B$, i.e. for every evaluation e , $e(A) \leq e(\neg \bigwedge \Gamma_B)$.

The attack principles are implications between (disjunction or conjunction) of assertions of the form $X \longrightarrow A$ or $X \not\rightarrow A$. Thus, in each attack principle we can distinguish the *premise* and the *conclusion*. For example, the premise of $(A \wedge)$ is “ $X \longrightarrow A$ or $X \longrightarrow B$ ”, while its conclusion is “ $X \longrightarrow A \wedge B$ ”. We say that an attack principle is *justified* whenever the attacking condition of the conclusion of the principle *logically follows* (in L_K) from the attacking condition of the premise. We now analyse the attack principles considering *defeating rebuttal* as attacking relation and both definitions of L-based arguments. The other cases can be found in Appendix A.

4.1 Analysis of the attack principles considering $[D-Reb]$ and $\models_L$

We assume to work in fully instantiated argumentation frameworks such that the arguments mentioned in the attack principles belong to the set of arguments.

(A. $\wedge$) If $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_A, A \rangle$, then $X \models_L \neg A$

, i.e. for any evaluation e s.t. $e(X) = 1$, $e(\neg A) = 1 - e(A) = 1$. Since for any evaluation e $e(\neg(A \wedge B)) = 1 - e(A \wedge B) = 1 - \min\{e(A), e(B)\} \geq 1 - e(A)$, if $1 - e(A) = 1$, then also $1 - e(A \wedge B) = 1$, i.e. $X \models_L \neg(A \wedge B)$ and the principle is justified.

(C. $\wedge$) If $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_{A \wedge B}, A \wedge B \rangle$, then

for any evaluation e such that $e(X) = 1$, $e(\neg(A \wedge B)) = 1$ i.e. $1 - e(A \wedge B) = 1 - \min\{e(A), e(B)\} = 1$. Therefore, either $e(A) = 0$ or $e(B) = 0$, i.e.

either $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_A, A \rangle$, or

$\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_B, B \rangle$ and the principle is justified.

(A. $\vee$) If $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_{A \vee B}, A \vee B \rangle$, then

$X \models_L \neg(A \vee B)$, i.e. for any evaluation e s.t. $e(X) = 1$, $e(\neg(A \vee B)) = 1 - e(A \vee B) = 1 - \max\{e(A), e(B)\} = 1$. Therefore, $\max\{e(A), e(B)\} = 0$, which implies both $e(A) = 0$ and $e(B) = 0$, i.e.

$\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_A, A \rangle$ and

$\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_B, B \rangle$. The principle is justified.

(C. $\vee$) If $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_A, A \rangle$ and

$\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_B, B \rangle$, then for any evaluation e s.t. $e(X) = 1$, we have that $e(\neg A) = 1$ and $e(\neg B) = 1$, i.e. $e(A) = 0$ and $e(B) = 0$. Therefore, $e(\neg(A \vee B)) = 1 - \max\{e(A), e(B)\} = 1$,

, i.e. $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_{A \vee B}, A \vee B \rangle$ and the principle is justified.

(A. $\supset$)* If $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_B, B \rangle$ and $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_A, A \rangle$, then $X \models_{\mathbf{L}} \neg B$ and $X \not\models_{\mathbf{L}} \neg A$, i.e. for any evaluation e s.t. $e(X) = 1$, then $e(\neg B) = 1 - e(B) = 1$, which implies $e(B) = 0$. Moreover, there exists at least one evaluation e' s.t. $e'(X) = 1$ and $e'(A) > 0$. To have the principle hold we would need that for any evaluation e s.t. $e(X) = 1$, then $e(\neg(A \supset B)) = 1$, i.e. $1 - e(A \supset B) = 1$ from which it would follow that $e(A \supset B) = 0$. Since $e(A \supset B) = \min\{1, 1 - e(A) + e(B)\}$, $e(A \supset B) = 0$ only if $1 - e(A) + e(B) = 0$, i.e. $e(A) = 1$ and $e(B) = 0$, but this does not follow from the premises and the principle is not justified.

(C. $\supset$)* If $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_{A \supset B}, A \supset B \rangle$, then $X \models_{\mathbf{L}} \neg(A \supset B)$, i.e. whenever there is an evaluation e s.t. $e(X) = 1$, then $e(\neg(A \supset B)) = 1 - e(A \supset B) = 1 - \min\{1, 1 - e(A) + e(B)\} = 1$. Thus, $1 - e(A) + e(B) = 0$, i.e. $e(A) = 1$ and $e(B) = 0$. From $e(B) = 0$ and $e(A) = 1$ it follows that $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_B, B \rangle$ and $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_A, A \rangle$. Thus, the attack principle is justified. From the premise it follows also that $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_{\neg A}, \neg A \rangle$.

(A. $\neg$) If $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_A, A \rangle$, then $X \models_{\mathbf{L}} \neg A$, i.e. for any evaluation e s.t. $e(X) = 1$, $e(\neg A) = 1 - e(A) = 1$ i.e. $e(A) = 0$. We want to show that there exists an evaluation e' s.t. $e'(X) = 1$ and $e'(\neg(\neg A)) < 1$, i.e. $e'(A) < 1$. From the premise we have that $e(A) = 0$. Therefore, $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_{\neg A}, \neg A \rangle$ and the principle is justified.

(C. $\neg$) If $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_{\neg A}, \neg A \rangle$, then for some evaluation e' s.t. $e'(X) = 1$, $e'(\neg\neg A) = e'(A) < 1$. However, from this hypothesis it does not follow that for any evaluation e s.t. $e(X) = 1$, $e(\neg A) = 1 - e(A) = 1$, i.e. $e(A) = 0$. Thus, the principle is not justified.

4.2 Analysis of the attack principles considering $[D-Reb]$ and $\models_{\mathbf{L}}^{\leq}$

(A. $\wedge$) If $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_B, A \rangle$, then $X \models_{\mathbf{L}}^{\leq} \neg A$, i.e. for every evaluation e , $e(X) \leq e(\neg A) = 1 - e(A)$. Since for every evaluation e $e(A \wedge B) = \min\{e(A), e(B)\} \leq e(A)$, $1 - e(A) \leq 1 - e(A \wedge B)$. Therefore, $e(X) \leq 1 - e(A \wedge B)$, i.e. $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_{A \wedge B}, A \wedge B \rangle$ and the principle is justified.

(C. $\wedge$) If $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_{A \wedge B}, A \wedge B \rangle$, then $X \models_{\mathbf{L}}^{\leq} \neg(A \wedge B)$ i.e. for every evaluation e $e(X) \leq e(\neg(A \wedge B)) = 1 - \min\{e(A), e(B)\}$. It could be that there exists an evaluation e_i s.t. $e_i(A \wedge B) = e_i(A)$ and other evaluation e_j for which $e_j(A \wedge B) = e_j(B)$. Therefore, we cannot deduce that for every evaluation e either $e(X) \leq e(\neg A)$ or $e(X) \leq e(\neg B)$. The principle is not justified.

(A. $\vee$) If $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_{A \vee B}, A \vee B \rangle$, then $X \models_{\mathbf{L}}^{\leq} \neg(A \vee B)$ i.e. for every evaluation e , $e(X) \leq e(\neg(A \vee B)) = 1 - e(A \vee B) = 1 - \max\{e(A), e(B)\}$. Since for every evaluation e , $e(A \vee B) \geq e(A)$ and $e(A \vee B) \geq e(B)$, it follows that $1 - e(A \vee B) \leq 1 - e(A)$ and $1 - e(A \vee B) \leq 1 - e(B)$, i.e. $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_A, A \rangle$ and $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_B, B \rangle$. The principle is justified.

(C. $\vee$) If $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_A, A \rangle$ and $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_B, B \rangle$, then $X \models_{\mathbf{L}}^{\leq} \neg A$ and $X \models_{\mathbf{L}}^{\leq} \neg B$, i.e. for every evaluation e , $e(X) \leq 1 - e(A)$ and $e(X) \leq 1 - e(B)$. Since $e(A \wedge B) = \min\{e(A), e(B)\} = e(A)$ or $e(A \wedge B) = \min\{e(A), e(B)\} = e(B)$, $e(X) \leq 1 - \max\{e(A), e(B)\}$ i.e. $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_{A \wedge B}, A \wedge B \rangle$ and the principle is justified.

(A. $\supset$)* If $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_B, B \rangle$ and $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_A, A \rangle$, then $X \models_{\mathbf{L}}^{\leq} \neg B$ and $X \not\models_{\mathbf{L}}^{\leq} \neg A$, i.e. for every evaluation e , $e(X) \leq 1 - e(B)$ and for some evaluation e^* , $e^*(X) > 1 - e^*(A)$. We would need to show that for every evaluation e ,

$e(X) \leq 1 - \min\{1, 1 - e(A) + e(B)\}$, but this does not follow from the premises and the principle is not justified.

- (C. $\supset$)* If $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_{A \supset B}, A \supset B \rangle$, then for every evaluation e , $X \models_{\mathbf{L}}^{\leq} \neg(A \supset B)$, i.e. $e(X) \leq 1 - e(A \supset B) = 1 - \min\{1, 1 - e(A) + e(B)\}$. Since for every evaluation e , $\min\{1, 1 - e(A) + e(B)\} \geq e(B)$, it follows that $1 - \min\{1, 1 - e(A) + e(B)\} \leq 1 - e(B)$, i.e. $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_B, B \rangle$.
- . However, the other part of the claim, $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_A, A \rangle$, does not follow from the premise and the principle is not justified.
- (A. $\neg$) If $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_A, A \rangle$, then $X \models_{\mathbf{L}}^{\leq} \neg A$, i.e. for every evaluation e , $e(X) \leq 1 - e(A)$. To have the principle justified, we need to find some evaluation e' s.t. $e'(X) > e'(A)$, but this does not follow from the premise and the principle is not justified.
- (C. $\neg$) If $\langle \Gamma_X, X \rangle \xrightarrow{[D-Reb]} \langle \Gamma_{\neg A}, \neg A \rangle$, then $X \not\models_{\mathbf{L}}^{\leq} \neg \neg A$, i.e. there exists an evaluation e' s.t. $e'(X) > e'(\neg \neg A) = e'(A)$. However, from this premise it does not follow that for every evaluation e , $e(X) \leq e(\neg A) = 1 - e(A)$. Thus, the principle is not justified.

Whenever the attack relation is instantiated with *direct rebuttal* and the arguments are defined using Definition 13, the attack principles justified are $AP_{\models_{\mathbf{L}}}^{[D-Reb]} = \{(A.\wedge), (C.\wedge), (A.\vee), (C.\vee), (C.\supset), (A.\neg)\}$. If Definition 14 is the one used for the definition of t-norm based arguments, then the attack principles justified are $AP_{\models_{\mathbf{L}}}^{[D-Reb]} = \{(A.\wedge), (A.\vee), (C.\vee)\}$.

In Tables 1, 2 and 3, we summarise the attack principles justified by the *defeat*, *compact rebuttal* and *indirect rebuttal* attack relation considering the two definitions of t-norm based argument. For the case of *defeat*, we have indicated which additional conditions the arguments need to satisfy to have the attack principle hold. Interestingly, the attack principles justified by most of the attack relations under Definition 13, are the same satisfied by the modal interpretation of the attack relation previously recalled and introduced in Corsi and Fermüller (2017). This is not the case whenever the order-based definition is considered. In this case, none of the attack relations considered satisfy either (C. $\supset$)* or (A. $\neg$). The attack principle (C. $\neg$) is problematic with both definitions. However, this is not surprising since it is a very demanding principle. Again, we have a confirmation that some of the attack principles are more acceptable, i.e. easier

Table 1 Attack principles justified by the *defeat* attack relation

[Def]	$\models_{\mathbf{L}}$	$\models_{\mathbf{L}}^{\leq}$
(A. $\wedge$)	$\Gamma_A \subseteq \Gamma_{A \wedge B}$	$\Gamma_A \subseteq \Gamma_{A \wedge B}$
(A. $\vee$)	$\Gamma_{A \vee B} \subseteq \Gamma_A$ and $\Gamma_{A \vee B} \subseteq \Gamma_B$	$\Gamma_{A \vee B} \subseteq \Gamma_A$ and $\Gamma_{A \vee B} \subseteq \Gamma_B$
(A. $\supset$)*	$\Gamma_B \subseteq \Gamma_{A \supset B}$	$\Gamma_B \subseteq \Gamma_{A \supset B}$
(A. $\neg$)	$\mathbf{x}$	$\mathbf{x}$
(C. $\wedge$)	$\Gamma_{A \wedge B} \subseteq \Gamma_A$ or $\Gamma_{A \wedge B} \subseteq \Gamma_B$	$\Gamma_{A \wedge B} \subseteq \Gamma_A$ or $\Gamma_{A \wedge B} \subseteq \Gamma_B$
(C. $\vee$)	$\Gamma_A \subseteq \Gamma_{A \vee B}$ or $\Gamma_B \subseteq \Gamma_{A \vee B}$	$\Gamma_A \subseteq \Gamma_{A \vee B}$ or $\Gamma_B \subseteq \Gamma_{A \vee B}$
(C. $\supset$)*	$\Gamma_{A \supset B} \subseteq \Gamma_B, \Gamma_A \subseteq \Gamma_{A \supset B}$ and $\gamma_i^* \notin \Gamma_A$	$\mathbf{x}$
(C. $\neg$)	$\mathbf{x}$	$\mathbf{x}$

Table 2 Attack principles justified by the *compact rebuttal* attack relation

[C-Reb]	$\models_{\mathbf{L}}$	$\models_{\mathbf{L}}^{\leq}$
(A. $\wedge$)	$\checkmark$	$\checkmark$
(A. $\vee$)	$\checkmark$	$\checkmark$
(A. $\supset$)*	$\mathbf{x}$	$\mathbf{x}$
(A. $\neg$)	$\checkmark$	$\mathbf{x}$
(C. $\wedge$)	$\mathbf{x}$	$\mathbf{x}$
(C. $\vee$)	$\checkmark$	$\mathbf{x}$
(C. $\supset$)*	$\checkmark$	$\mathbf{x}$
(C. $\neg$)	$\mathbf{x}$	$\mathbf{x}$

Table 3 Attack principles justified by the *indirect rebuttal* attack relation

[I-Reb]	$\models_{\mathbf{L}}$	$\models_{\mathbf{L}}^{\leq}$
(A. $\wedge$)	$\checkmark$	$\checkmark$
(A. $\vee$)	$\checkmark$	$\checkmark$
(A. $\supset$)*	$\mathbf{x}$	$\mathbf{x}$
(A. $\neg$)	$\mathbf{x}$	$\mathbf{x}$
(C. $\wedge$)	$\mathbf{x}$	$\mathbf{x}$
(C. $\vee$)	$\checkmark$	$\checkmark$
(C. $\supset$)*	$\mathbf{x}$	$\mathbf{x}$
(C. $\neg$)	$\mathbf{x}$	$\mathbf{x}$

to justify, than others. In particular, those attack principles satisfied by the modal interpretation of the attack relation are justified in different scenarios: in sequent-based argumentation frameworks and now also in t-norm based ones.

As shown in Corsi and Fermüller (2018a), semi-abstract argumentation frameworks are not adequate to characterise fuzzy logics. In Corsi and Fermüller (2018a), a complete argumentative semantics for $\mathbf{L}$, $\mathbf{G}$ and $\mathbf{II}$ is recovered considering *weighted* semi-abstract argumentation frameworks, i.e. argumentative frameworks where the attack relation is graded. An alternative way to recover complete semantics for the above mentioned logics is to consider *bipolar* t-norm based argumentative frameworks. From now on,

we consider as t-norm-based arguments only those defined according to Definition 13.

Definition 15 (Bi-LAF) A *Bipolar L-based Argumentation framework* is a triplet $AF = \langle Ar, Attack(\mathcal{A}), Support(\mathcal{S}) \rangle$ such that:

Ar is a set of L-Lbased arguments, i.e. $\langle \Gamma, \psi \rangle \in Ar$ if only if $\Gamma \models_{\mathbf{L}} \psi$.

$\mathcal{A}$ is a set of attack rules and $\mathcal{S}$ is a set of support rules.

$(\langle \Gamma_A, A \rangle, \langle \Gamma_B, B \rangle) \in Attack(\mathcal{A})$ iff there is some $\mathcal{R} \in \mathcal{A}$ such that $\langle \Gamma_A, A \rangle \mathcal{R}$ -attacks $\langle \Gamma_B, B \rangle$.

$(\langle \Gamma_A, A \rangle, \langle \Gamma_B, B \rangle) \in Support(\mathcal{S})$ iff there is some $\mathcal{R} \in \mathcal{S}$ such that $\langle \Gamma_A, A \rangle \mathcal{R}$ -supports $\langle \Gamma_B, B \rangle$.

The definition of Bipolar G-based and Bipolar II-based Argumentation Frameworks is similar. In bipolar argumentation frameworks, we can define (bipolar) *principles* that are similar to the attack principles, but in these, both relations are used. For example, for implication we can introduce the following principle.

$(C_{Bi.\supset})$ If $\langle \Gamma_X, X \rangle \longrightarrow \langle \Gamma_{A \supset B}, A \supset B \rangle$, then

$\langle \Gamma_X, X \rangle \longrightarrow \langle \Gamma_B, B \rangle$ and

$\langle \Gamma_X, X \rangle \Longrightarrow \langle \Gamma_A, A \rangle$. In words: If an argu-

ment with claim X attacks an argument with claim $A \supset B$, then the argument with claim X also attacks an argument with claim B and supports an argument with claim A .

If we consider a Bi-LAF argumentation framework such that $\mathcal{A} = \{[D-Reb]\}$ and $\mathcal{S} = \{[D-Sup]\}$, the principle $(C_{Bi.\supset})$ is justified. In fact, if $X \models_{\mathbf{L}} \neg(A \supset B)$, then whenever $e(X) = 1$, we have that $e(\neg(A \supset B)) = 1 - e(A \supset B) = 1$. Thus, $e(A \supset B) = 0$ holds. Since $e(A \supset B) = \min\{1, 1 - e(A) + e(B)\}$, if $e(A \supset B) = 0$, then $1 - e(A) + e(B) = 0$, i.e. $e(A) = 1$ and $e(B) = 0$. Therefore, for any evaluation e such that $e(X) = 1$, $e(A) = 1$ and $e(B) = 0$, we have that $X \models_{\mathbf{L}} A$ and $X \models_{\mathbf{L}} \neg B$.

In the same setting, also the following principles hold.

$(A_{S.\vee})$ If $\langle \Gamma_X, X \rangle \Longrightarrow \langle \Gamma_{A \vee B}, A \vee B \rangle$, then

$\langle \Gamma_X, X \rangle \Longrightarrow \langle \Gamma_A, A \rangle$ or $\langle \Gamma_X, X \rangle \Longrightarrow \langle \Gamma_B, B \rangle$.

If $X \models_{\mathbf{L}} A \vee B$, then whenever $e(X) = 1$, we have that $e(A \vee B) = \min\{e(A), e(B)\} = 1$. Thus, either $e(A) = 1$ or $e(B) = 1$, i.e. either $X \models_{\mathbf{L}} A$ or $X \models_{\mathbf{L}} B$.

$(A_{S.\supset})$ If $\langle \Gamma_X, X \rangle \Longrightarrow \langle \Gamma_A, A \rangle$ and

$\langle \Gamma_X, X \rangle \Longrightarrow \langle \Gamma_{A \supset B}, A \supset B \rangle$, then

$\langle \Gamma_X, X \rangle \Longrightarrow \langle \Gamma_B, B \rangle$.

If $X \models_{\mathbf{L}} A$ and $X \models_{\mathbf{L}} A \supset B$, then whenever $e(X) = 1$ we have that $e(A) = 1$ and $e(A \supset B) = 1$. Since $e(A \supset B) = \min\{1, 1 - e(A) + e(B)\}$ and $e(A) = 1$, $e(A \supset B) = \min\{1, e(B)\} = e(B)$. Thus, $e(B) = 1$ and $X \models_{\mathbf{L}} B$.

$(A_{S.\neg})$ If $\langle \Gamma_X, X \rangle \Longrightarrow \langle \Gamma_A, A \rangle$, then $\langle \Gamma_X, X \rangle \not\Longrightarrow \langle \Gamma_{\neg A}, \neg A \rangle$.

If $X \models_{\mathbf{L}} A$, then whenever $e(X) = 1$ we have that $e(A) = 1$. Thus, $e(\neg A) = 0$ and $X \not\models_{\mathbf{L}} \neg A$.

$(C_{Bi.\neg})$ If $\langle \Gamma_X, X \rangle \longrightarrow \langle \Gamma_A, A \rangle$, then

$\langle \Gamma_X, X \rangle \Longrightarrow \langle \Gamma_{\neg A}, \neg A \rangle$.

If $X \models_{\mathbf{L}} \neg A$, then whenever $e(X) = 1$ we have that $e(\neg A) = 1$, i.e. $\langle \Gamma_X, X \rangle \Longrightarrow \langle \Gamma_{\neg A}, \neg A \rangle$.

$(A.\top)$ $\langle \Gamma_X, X \rangle \Longrightarrow \langle \Gamma_{\top}, \top \rangle$ for every argument $\langle \Gamma_X, X \rangle$.

$X \models_{\mathbf{L}} \top$ for every X .

5 Characterising L, G and II

To relate fuzzy logics to the realm of all possible t-norm based argumentation frameworks that satisfy certain principles like the ones discussed in the previous section, we define a closure operation on t-norm based argumentation frameworks similar to the one introduced in Corsi and Fermüller (2018a).

Definition 16 (Syntactic Closure of Bi-LAFs) Given Δ a finite set of propositional formulas, a Bi-LAF AF is *syntactically closed* with respect to Δ if all formulas and subformulas of formulas in Δ occur as claims of some argument in AF .

Definition 16 extends to Bi-GAFs and Bi-IIAFs in the obvious way and we will suppress the explicit reference to Δ whenever the context makes clear what formulas are expected to be available as claims of arguments.

As explained in Corsi and Fermüller (2018a), the following notion of *immunity*, provides a new view on logical validity that only refers to claims of arguments, to the attacks and to the supports between them.

Definition 17 Let P be a set of principles. A formula F is *P-argumentatively immune* (shortly: *P-immune*) if in all syntactically closed Bi-LAFs (with respect to F) that satisfy the principles in P , F is not attacked.

In addition, as validity for the three logics analysed in the paper is co-NP-complete (see Baaz et al. (2001) and Haniková (2002)), the respective forms of argumentative immunity can, presumably, be verified much more efficiently than semantic properties that appear in the context of non-monotonic reasoning, see the work of Gottlob in Gottlob (1992).

Let us denote with $\mathbf{L} P$ the following set of principles. Some of them are attack principles, the others are the principles just introduced.

$\mathbf{LP} = \{(\mathbf{A}.\wedge), (\mathbf{C}.\wedge), (\mathbf{A}.\vee), (\mathbf{C}.\vee), (\mathbf{A}.\top), (\mathbf{C}_{Bi}.\supset), (\mathbf{A}_S.\supset), (\mathbf{A}_S.\vee), (\mathbf{C}_{Bi}.\neg)\}$.

As proved in the following proposition, $\mathbf{LP}$ -immune arguments are closed over *modus ponens*.

Proposition 2 (Closure of $\mathbf{LP}$ -immune arguments over Modus Ponens) *If A and $A \supset B$ are argumentatively $\mathbf{LP}$ -immune, then also B is argumentatively $\mathbf{LP}$ -immune.*

Proof Since both A and $A \supset B$ are $\mathbf{LP}$ -immune, then for any X in a syntactically closed Bi- $\mathbf{LAF}$ framework $X \not\models_{\mathbf{L}} \neg A$ and $X \not\models_{\mathbf{L}} \neg(A \supset B)$. Thus, there is an evaluation e_1 s.t. $e_1(X) = 1$ and $e_1(\neg A) < 1$ and an evaluation e_2 s.t. $e_2(X) = 1$ and $e_2(\neg(A \supset B)) < 1$. Our claim is that $X \not\models_{\mathbf{L}} \neg B$ for any X in the framework, i.e. there exists an evaluation e_3 s.t. $e_3(X) = 1$ and $e_3(\neg B) < 1$. Since A belongs to the framework, also $X \wedge A$ is in the framework and from the hypothesis it must be that $X \wedge A \not\models_{\mathbf{L}} \neg A$ and $X \wedge A \not\models_{\mathbf{L}} \neg(A \supset B)$, i.e. there is an evaluation e_4 s.t. $e_4(X \wedge A) = 1$ and $e_4(\neg(A \supset B)) < 1$. Since $e_4(X \wedge A) = \min\{e_4(X), e_4(A)\} = 1$, then $e_4(A) = 1$. Moreover, since $e_4(\neg(A \supset B)) = e_4(A \supset B) = \min\{1, 1 - e_4(A) + e_4(B)\}$ and $e_4(A) = 1$, we have that $e_4(\neg(A \supset B)) = 1 - \min\{1, e_4(B)\} = 1 - e_4(B)$. Thus, from $e_4(\neg(A \supset B)) < 1$, it follows $e_4(B) > 0$ and $e_4(\neg B) < 1$. Therefore, for any argument X in the framework we can find an evaluation showing that $X \not\models_{\mathbf{L}} \neg B$. $\square$

Theorem 3 (Bipolar Argumentative Soundness of $\mathbf{L}$) *Every $\mathbf{L}$ -valid formula is argumentatively $\mathbf{LP}$ -immune.*

Proof By Theorem 1 and Proposition 2 remain to check that the axioms of the axiomatic system for Łukasiewicz logic recalled in Sect. 3 are $\mathbf{LP}$ -immune. In the following, we implicitly assume that all arguments occur in a Bi- $\mathbf{LAF}$ that is syntactically closed with respect to the axiom in question. In each case, we argue indirectly, and we derive a contradiction from the assumption that there is an argument X that attacks the axiom in question.

[Tr] If $X \rightarrow (F \supset G) \supset ((G \supset H) \supset (F \supset H))$, then by $(\mathbf{C}_{Bi}.\supset)$ we have that $X \Rightarrow F \supset (G \supset H)$ and $X \rightarrow ((F \supset G) \supset (F \supset H))$. Again by $(\mathbf{C}_{Bi}.\supset)$ we have $X \Rightarrow F \supset G$ and $X \rightarrow F \supset H$ from which it follows, by $(\mathbf{C}_{Bi}.\supset)$, $X \rightarrow H$ and $X \Rightarrow F$. Since $X \Rightarrow F$ and $X \Rightarrow F \supset (G \supset H)$, by $(\mathbf{A}_S.\supset)$, we have $X \Rightarrow G \supset H$ and from $X \Rightarrow F \supset G$ and $X \Rightarrow F$, by $(\mathbf{A}_S.\supset)$, it follows $X \Rightarrow G$. Again by $(\mathbf{A}_S.\supset)$, from $X \Rightarrow G \supset H$ and $X \Rightarrow G$ it follows $X \Rightarrow H$ which goes against $X \rightarrow H$.

[We] If $X \rightarrow F \supset (G \supset F)$, by $(\mathbf{C}_{Bi}.\supset)$ it follows $X \rightarrow F \supset G$ and $X \Rightarrow F$. From $X \rightarrow F \supset G$ it follows by $(\mathbf{C}_{Bi}.\supset)$ $X \Rightarrow F$ and $X \rightarrow F$ which is incompatible with $X \Rightarrow F$.

[Ex] If $X \rightarrow (F \supset (G \supset H)) \supset (G \supset (F \supset H))$ by $(\mathbf{C}_{Bi}.\supset)$ it follows that $X \rightarrow (G \supset (F \supset H))$ and $X \Rightarrow (F \supset (G \supset H))$. Again, by $(\mathbf{C}_{Bi}.\supset)$, we have that $X \rightarrow F \supset H$ and $X \Rightarrow G$. Since $X \rightarrow F \supset H$ it follows also $X \rightarrow H$ and $X \Rightarrow F$. From $X \Rightarrow (F \supset (G \supset H))$ and $X \Rightarrow F$ by $(\mathbf{A}_S.\supset)$ it follows $X \Rightarrow G \supset H$ and by $X \Rightarrow G \supset H$ and $X \Rightarrow G$ it follows $X \Rightarrow H$ which is incompatible with $X \rightarrow H$.

[$\wedge$ -1] If $X \rightarrow (F \wedge G) \supset F$ by $(\mathbf{C}_{Bi}.\supset)$ it follows $X \rightarrow F$, $X \Rightarrow F \wedge G$ and $X \not\rightarrow F \wedge G$. Moreover by $(\mathbf{A}.\wedge)$ we have $X \not\rightarrow F$ and $X \not\rightarrow G$ which contradicts $X \rightarrow F$.

[$\wedge$ -2] See the previous case.

[$\wedge$ -3] If $X \rightarrow (H \supset F) \supset ((H \supset G) \supset (H \supset (F \wedge G)))$ by $(\mathbf{C}_{Bi}.\supset)$ it follows $X \rightarrow (H \supset G) \supset (H \supset (F \wedge G))$ and $X \Rightarrow H \supset F$ and again, by the same principle we have $X \rightarrow H \supset (F \wedge G)$ and $X \Rightarrow H \supset G$. Since $X \rightarrow H \supset (F \wedge G)$ it follows $X \rightarrow F \wedge G$ and $H \Rightarrow H$. From $X \rightarrow F \wedge G$ by $(\mathbf{C}.\wedge)$ it follows that $X \rightarrow F$ or $X \rightarrow G$. Since $X \Rightarrow H \supset F$ and $X \Rightarrow H$, by $(\mathbf{A}_S.\supset)$ we have $X \Rightarrow F$ and by $X \Rightarrow H$ and $X \Rightarrow H \supset G$ it follows $X \Rightarrow G$, but this is in contradiction with $X \rightarrow F$ or $X \rightarrow G$.

[$\vee$ -1] If $X \rightarrow F \supset (F \vee G)$ by $(\mathbf{C}_{Bi}.\supset)$ we have $X \rightarrow F \vee G$ and $X \Rightarrow F$. From $X \rightarrow F \vee G$ by $(\mathbf{A}.\vee)$ it follows $X \rightarrow G$ and $X \rightarrow F$ which is in contradiction with $X \Rightarrow F$.

[$\vee$ -2] See the previous case.

[$\vee$ -3] If $X \rightarrow (G \supset F) \supset ((H \supset F) \supset ((G \vee H) \supset F))$ by $(\mathbf{C}_{Bi}.\supset)$ we have $X \rightarrow (H \supset F) \supset ((G \vee H) \supset F)$ and $X \Rightarrow G \supset F$. Again, by $(\mathbf{C}_{Bi}.\supset)$, we have also $X \rightarrow (G \vee H) \supset F$ and $X \Rightarrow H \supset F$. From $X \rightarrow (G \vee H) \supset F$ it follows also $X \rightarrow F$ and $X \Rightarrow G \vee H$. Since $X \Rightarrow G \vee H$, by $(\mathbf{A}_S.\vee)$ either

$$X \Rightarrow G \quad (1)$$

or

$$X \Rightarrow H. \quad (2)$$

In case (1), since $X \Rightarrow G$ and $X \Rightarrow G \supset F$, by $(A_S.\supset)$ we have that $X \Rightarrow F$. However from $X \rightarrow F$ and $(C_{Bi}.\neg)$ it follows $X \Rightarrow \neg F$ and this is in contradiction with $X \Rightarrow F$. In case (2), we reach the contradiction in the same way.

[Lin] If $X \rightarrow (F \supset G) \vee (G \supset F)$ from $(A.\vee)$ it follows $X \rightarrow F \supset G$ and $X \rightarrow G \supset F$. From $X \rightarrow F \supset G$, by $(C_{Bi}.\supset)$, it follows $X \rightarrow G$ and $X \Rightarrow F$ while from $X \rightarrow G \supset F$ it follows $X \rightarrow F$ and $X \Rightarrow G$ and we have reached a contradiction.

[$\perp$] If $X \rightarrow \perp \supset F$, then by $(C_{Bi}.\supset)$ it follows $X \rightarrow F$ and $X \Rightarrow \perp$. Thus by $(A_S.\neg)$ and $(A.\top)$ we reach a contradiction.

[Waj] If $X \rightarrow ((F \supset G) \supset G) \supset ((G \supset F) \supset F)$, then by $(C_{Bi}.\supset)$ we have $X \rightarrow (G \supset F) \supset F$ and $X \Rightarrow (F \supset G) \supset G$. From $X \rightarrow (G \supset F) \supset F$ again by $(C_{Bi}.\supset)$ we have $X \rightarrow F$ and $X \Rightarrow G \supset F$. Therefore,

$$X \models_{\mathbf{L}} \neg F, \quad (3)$$

$$X \models_{\mathbf{L}} G \supset F \quad (4)$$

and

$$X \models_{\mathbf{L}} (F \supset G) \supset G. \quad (5)$$

From (3) it follows that whenever $e(X) = 1$ $1 - e(F) = 1$, i.e. $e(F) = 0$. From (4) it follows that whenever $e(X) = 1$, $e(G \supset F) = \min\{1, 1 - e(G) + e(F)\} = 1$ and given $e(F) = 0$ this is satisfied only if $e(G) = 0$. From (5) we have that whenever $e(X) = 1$ for some evaluation e , $e((F \supset G) \supset G) = 1$, i.e. $\min\{1, 1 - e(F \supset G) + e(G)\} = 1$. Since $e(F \supset G) = \min\{1, 1 - e(F) + e(G)\}$, given $e(F) = 0$ and $e(G) = 0$, we have $e(F \supset G) = 1$. Therefore,

Table 4 Principles used in Theorem 5

Axiom	Principles used in the proof
[Tr]	$(C_{Bi}.\supset), (A_S.\supset)$
[We]	$(C_{Bi}.\supset)$
[Ex]	$(C_{Bi}.\supset), (A_S.\supset)$
[$\wedge$ -1]	$(C_{Bi}.\supset), (A.\wedge)$
[$\wedge$ -2]	$(C_{Bi}.\supset), (A.\wedge)$
[$\wedge$ -3]	$(C_{Bi}.\supset), (A_S.\supset), (C.\wedge)$
[$\vee$ -1]	$(C_{Bi}.\supset), (A.\vee)$
[$\vee$ -2]	$(C_{Bi}.\supset), (A.\vee)$
[$\vee$ -3]	$(C_{Bi}.\supset), (A_S.\supset), (A_S.\vee), (C_{Bi}.\neg)$
[Lin]	$(C_{Bi}.\supset), (A.\vee)$
[$\perp$]	$(C_{Bi}.\supset), (A_S.\neg), (A.\top)$
[Waj]	$(C_{Bi}.\supset)$

$\min\{1, 1 - e(F \supset G) + e(G)\} = 0$ while it should have been 1.

In Table 4 we summarise which principles have been used in the corresponding section of the proof. $\square$ **Theorem 4** (Bipolar Argumentative Completeness of $\mathbf{L}$) Every argumentatively $\mathbf{L}$ P-immune formula is $\mathbf{L}$ -valid.

Proof We have to show that if F is not an $\mathbf{L}$ -valid formula, then it is not $\mathbf{L}$ P-immune. If F is not an $\mathbf{L}$ -valid formula, then there is an evaluation e such that $e(F) < 1$. Since $e(\neg F) = 1 - e(F)$, any formula F is attacked by its negation, i.e. $\neg F \rightarrow F$. $\square$

Concerning the attack principles needed to prove a completeness theorem with G-based argumentation frameworks, we first need to verify that the interpretation of $(C.\wedge)$, $(A.\vee)$, and $(C.\supset)^*$ are justified in Bi-GAFs where $\mathcal{A} = \{[D-Reb]\}$ and $\mathcal{S} = \{[D-Sup]\}$.

$(C.\wedge)$ If $X \models_G \neg(A \wedge B)$, then for any evaluation e such that $e(X) = 1$, we have that $e(\neg(A \wedge B)) = 1$. Thus, $\min\{e(A), e(B)\} = 0$ which implies either $e(A) = 0$ or $e(B) = 0$. Therefore, either $X \models_G \neg A$, or $X \models_G \neg B$.

$(A.\vee)$ If $X \models_G \neg(A \vee B)$, then for any evaluation e such that $e(X) = 1$, then $e(\neg(A \vee B)) = 1$. Therefore, $\max\{e(A), e(B)\} = 0$, which implies $e(A) = 0$ and $e(B) = 0$, i.e. $X \models_G \neg B$ and $X \models_G \neg A$.

$(C.\supset)^*$ If $X \models_G \neg(A \supset B)$, then for any evaluation e such that $e(X) = 1$, then also $e(\neg(A \supset B)) = 1$ and this happens only if $e(A \supset B) = 0$. The only case in which the implication has value 0 is whenever $e(B) = 0$ and $e(A) > 0$. Therefore, $e(\neg B) = 1$ and $e(\neg A) = 0$, i.e. $X \models_G \neg B$ and $X \not\models_G \neg A$.

In addition, we also need the following principles.

$(C'_{Bi}.\supset)$ If $\langle \Gamma_X, X \rangle \rightarrow \langle \Gamma_{A \supset B}, A \supset B \rangle$, then $\langle \Gamma_X, X \rangle \rightarrow \langle \Gamma_B, B \rangle$ and $\langle \Gamma_X, X \rangle \Rightarrow \langle \Gamma_A, A \rangle$. If $X \models_G \neg(A \supset B)$, then whenever there is an evaluation e such that $e(X) = 1$, then $e(\neg(A \supset B)) = 1$. Therefore, $e(A \supset B) = 0$, from which it follows $e(B) = 0$ and $e(A) > 0$. From $e(B) = 0$ it follows $e(\neg B) = 1$ and from $e(A) > 0$ it follows $e(\neg A) = 0$ and $e(\neg \neg A) = 1$. Conclusively, $X \models_G \neg B$ and $X \models_G \neg \neg A$.

$(A.\perp)$ $\langle \Gamma_X, X \rangle \rightarrow \langle \Gamma_{\perp}, \perp \rangle$ for every argument $\langle \Gamma_X, X \rangle$. $X \models_G \neg \perp$ for every X .

Let us define the set GP of principles as follows.

$$GP = \{(C.\wedge), (A.\vee), (C.\supset)^*, (A.\perp), (C'_{Bi}.\supset)\}$$

Proposition 5 (Closure of GP-immune arguments over Modus Ponens) *If A and $A \supset B$ are argumentatively GP-immune, then also B is argumentatively GP-immune.*

Proof We have to show that if

$$X \not\models_G \neg A \quad (6)$$

and

$$X \not\models_G \neg(A \supset B), \quad (7)$$

then

$$X \not\models_G \neg B. \quad (8)$$

From (6) it follows that there exists an evaluation e_1 such that $e_1(X) = 1$ and $e_1(\neg A) < 1$, from which it follows $e_1(\neg A) = 0$ and $e_1(A) > 0$. From 7 it follows that there exists an evaluation e_2 such that $e_2(X) = 1$ and $e_2(\neg(A \supset B)) < 1$, i.e. $e_2(\neg(A \supset B)) = 0$ and $e_2(A \supset B) > 0$. If we consider an evaluation e^* such that $e^*(X) = 1$, $e^*(A) = 1$ and $e^*(A \supset B) = 1$, then this evaluation e^* is compatible with both hypothesis 6 and 7. Since $e^*(A \supset B) = 1$, we have that $e^*(A) \leq e^*(B)$. Since $e^*(A) = 1$, it follows $e^*(B) = 1$ and $e^*(\neg B) = 0$, i.e. $X \not\models_G \neg B$. $\square$

Theorem 6 (Bipolar Argumentative Soundness of G) *Any formula F is G-valid iff it is GP-immune.*

Proof ($\Rightarrow$) Given Proposition 5, it remain to show that every G-axiom is argumentatively GP-immune.

[Tr] If there is an argument X such that $X \longrightarrow (F \supset G) \supset ((G \supset H) \supset (F \supset H))$, then, by $(C'_{Bi} \cdot \supset)$ it follows

$$X \longrightarrow (G \supset H) \supset (F \supset H) \quad (9)$$

and

$$X \Longrightarrow \neg\neg(F \supset G). \quad (10)$$

From 9 and $(C'_{Bi} \cdot \supset)$ it follows

$$X \longrightarrow F \supset H \quad (11)$$

and

$$X \Longrightarrow \neg\neg G \supset H. \quad (12)$$

From 11 and $(C'_{Bi} \cdot \supset)$ it follows

$$X \longrightarrow H \quad (13)$$

and

$$X \Longrightarrow \neg\neg F. \quad (14)$$

From 14 it follows that whenever there is an evaluation e such that $e(X) = 1$, then $e(\neg\neg F) = 1$. Therefore, $e(\neg F) = 0$ and

$$e(F) > 0. \quad (15)$$

From 12 it follows that whenever there is an evaluation e such that $e(X) = 1$, then $e(\neg\neg(G \supset H)) = 1$, i.e. $e(\neg(G \supset H)) = 0$ and $e(G \supset H) > 0$. Since $e(H) = 0$, from 13 it follows, $e(G) = 0$. From 10 it follows that $e(\neg\neg(F \supset G)) = 1$, $e(\neg(F \supset G)) = 0$ and $e(F \supset G) > 0$. Since $e(G) = 0$, it follows that $e(F) = 0$, but this is in contradiction with 15.

[We] If there is some argument X such that $X \longrightarrow F \supset (G \supset F)$, then, by $(C'_{Bi} \cdot \supset)$ we have that

$$X \longrightarrow G \supset F \quad (16)$$

and

$$X \Longrightarrow \neg\neg F. \quad (17)$$

From 16 and $(C'_{Bi} \cdot \supset)$ it follows

$$X \longrightarrow F \quad (18)$$

and

$$X \Longrightarrow \neg\neg G. \quad (19)$$

From 17 it follows that whenever there is an evaluation e such that $e(X) = 1$, $e(\neg F) = 1$, i.e. $e(F) = 0$. From 18 it follows that whenever there is an evaluation e such that $e(X) = 1$, $e(\neg\neg F) = 1$, i.e. $e(\neg F) = 0$ and $e(F) > 0$, but this is in contradiction with $e(F) = 0$.

[Ex] If there is some argument X such that $X \longrightarrow (F \supset (G \supset H)) \supset (G \supset (F \supset H))$, then, by $(C'_{Bi} \cdot \supset)$ we have

$$X \longrightarrow G \supset (F \supset H) \quad (20)$$

and

$$X \Longrightarrow \neg\neg(F \supset (G \supset H)). \quad (21)$$

From 20 and $(C'_{Bi} \supset)$ it follows

$$X \longrightarrow F \supset H \quad (22)$$

and

$$X \Longrightarrow \neg\neg G \quad (23)$$

From 22 and $(C'_{Bi} \supset)$ we have

$$X \longrightarrow H \quad (24)$$

and

$$X \Longrightarrow \neg\neg F. \quad (25)$$

From 24 it follows that whenever there is an evaluation e such that $e(X) = 1$, it follows $e(H) = 0$. From 23 it follows that $e(F) > 0$ and from 23 it follows $e(G) > 0$. From 21 it follows $e(\neg\neg(F \supset (G \supset H))) = 1$. Therefore, $e(F \supset (G \supset H)) > 0$. This last inequality holds if either

$$e(G \supset H) > 0 \quad (26)$$

or

$$e(F) = e(G \supset H). \quad (27)$$

From 26 and $e(H) = 0$, it follows $e(G) = 0$, but this is in contradiction with $e(G) > 0$. From 27, $e(G) > 0$ and $e(H) = 0$ it follows $e(G \supset H) = 0$. Therefore, $e(F) = 0$, but this is in contradiction with $e(F) > 0$.

[$\wedge$ -1] If there is some argument X such that $X \longrightarrow (F \wedge G) \supset F$, then, by $(C'_{Bi} \supset)$ it follows $X \longrightarrow F$ and $X \Longrightarrow \neg\neg(F \wedge G)$. Therefore, whenever there is an evaluation e such that $e(X) = 1$, then $e(F) = 0$ and $e(\neg\neg(F \wedge G)) = 1$ from which it follows $\min\{e(F), e(G)\} > 0$. In particular, $e(F) > 0$ that is in contradiction with $e(F) = 0$.

[$\wedge$ -2] This case is similar to the previous one.

[$\wedge$ -3] If there is some argument X such that $X \longrightarrow (H \supset F) \supset ((H \supset G) \supset (H \supset (F \wedge G)))$, then, by $(C'_{Bi} \supset)$ it follows

$$X \longrightarrow (H \supset G) \supset (H \supset (F \wedge G)) \quad (28)$$

and

$$X \Longrightarrow \neg\neg(H \supset F). \quad (29)$$

From 28 and $(C'_{Bi} \supset)$ it follows

$$X \longrightarrow H \supset (F \wedge G) \quad (30)$$

and

$$X \Longrightarrow \neg\neg(H \supset G). \quad (31)$$

From 30 and $(C'_{Bi} \supset)$ it follows

$$X \longrightarrow (F \wedge G) \quad (32)$$

and

$$X \Longrightarrow \neg\neg H. \quad (33)$$

From 32 and $(C \wedge)$ it follows either

$$X \longrightarrow F \quad (34)$$

or

$$X \longrightarrow G. \quad (35)$$

In the case 34 we have that whenever there is an evaluation e such that $e(X) = 1$, then $e(F) = 0$. From 33 it follows that $e(H) > 0$ and from 29 we have that $e(\neg\neg(H \supset F)) = 1$. Therefore, $e(H \supset F) > 0$. However from $e(H) > 0 = e(F)$ it follows $e(H \supset F) = 0$ and we have found a contradiction. In the case 35 we can reach a contradiction in a very similar way.

[$\vee$ -1] If there is some argument X such that

$X \longrightarrow F \supset (F \vee G)$, then, by $(C'_{Bi} \supset)$ we have

$$X \longrightarrow F \vee G \quad (36)$$

and

$$X \Longrightarrow \neg\neg F. \quad (37)$$

From 36 and $(A \vee)$ it follows $X \longrightarrow F$ and $X \longrightarrow G$. Therefore, whenever there is an evaluation e such that $e(X) = 1$, $e(F) = 0$ and $e(G) = 0$, but this is in contradiction with what follows from 37: $e(F) > 0$.

[$\vee$ -2] This case is similar to the previous one.

[$\vee$ -3] If there is some argument X such that $X \longrightarrow (G \supset F) \supset ((H \supset F) \supset ((G \vee H) \supset F))$, then, by $(C'_{Bi} \supset)$ we have

$$X \longrightarrow (H \supset F) \supset ((G \vee H) \supset F) \quad (38)$$

and

$$X \Longrightarrow \neg\neg(G \supset F). \quad (39)$$

From 38 and $(C'_{Bi} \supset)$ it follows

$$X \longrightarrow (G \vee H) \supset F \quad (40)$$

and

$$X \Longrightarrow \neg\neg(H \supset F). \quad (41)$$

From 40 and $(C'_{Bi} \supset)$ it follows

$$X \longrightarrow F \quad (42)$$

and

$$X \Longrightarrow \neg\neg(G \vee H). \quad (43)$$

From 42 it follows that whenever there is an evaluation e such that $e(X) = 1$, $e(F) = 0$. From 43 it follows $e(G \vee H) > 0$, i.e. $\max\{e(G), e(H)\} > 0$. Therefore, either

$$e(G) > 0 \quad (44)$$

or

$$e(H) > 0. \quad (45)$$

In the first case we have that $e(G) > e(F) = 0$ from which it follows $e(G \supset F) = e(F) = 0$, but this is in contradiction with what follows from 39: $e(G \supset F) > 0$. In the case 45 we reach a contradiction from 41 as in the previous point.

[Lin] If there is some argument X such that $X \longrightarrow (F \supset G) \vee (G \supset F)$, by $(C'_{Bi} \supset)$ we have

$$X \longrightarrow F \supset G \quad (46)$$

and

$$X \longrightarrow G \supset F. \quad (47)$$

From 46 and $(C \supset)^*$ it follows $X \longrightarrow G$ and $X \not\rightarrow F$, but this is in contradiction with what follows from 47 and $(C \supset)^*$: $X \longrightarrow F$ and $X \not\rightarrow G$.

[$\perp$] If there is some argument X such that $X \longrightarrow \perp \supset F$, then, by $(C \supset)^*$ we have $X \longrightarrow F$ and $X \not\rightarrow \perp$. However, this it cannot be since every argument is supposed to attack *falsum* as $(A \perp)$ requires.

[Con] If there is some argument X such that $X \longrightarrow (F \supset (F \supset G)) \supset (F \supset G)$, by $(C'_{Bi} \supset)$ we have

Table 5 Principles used in Theorem 6

Axiom	Principles used in the proof
[Tr]	$(C'_{Bi} \supset)$
[We]	$(C'_{Bi} \supset)$
[Ex]	$(C'_{Bi} \supset)$
$[\wedge-1]$	$(C'_{Bi} \supset)$
$[\wedge-2]$	$(C'_{Bi} \supset)$
$[\wedge-3]$	$(C'_{Bi} \supset), (C \wedge)$
$[\vee-1]$	$(C'_{Bi} \supset), (A \vee)$
$[\vee-2]$	$(C'_{Bi} \supset)$
$[\vee-3]$	$(C'_{Bi} \supset)$
[Lin]	$(C \supset)^*$
$[\perp]$	$(C \supset)^*, (A \perp)$
[Con]	$(C'_{Bi} \supset)$

$$X \longrightarrow F \supset G \quad (48)$$

and

$$X \Longrightarrow \neg\neg(F \supset (F \supset G)). \quad (49)$$

From 48 and $(C'_{Bi} \supset)$ it follows

$$X \longrightarrow G \quad (50)$$

and

$$X \Longrightarrow \neg\neg F. \quad (51)$$

From 50 and 51 it follows that whenever there is an evaluation e such that $e(X) = 1$, then $e(G) = 0$ and $e(F) > 0$. Therefore, since $e(F) > e(G) = 0$, $e(F \supset G) = 0$ and $e(F \supset (F \supset G)) = 0$, but this is in contradiction with what follows from 49: $e(F \supset (F \supset G)) > 0$ $\square$.

Theorem 7 (*Bipolar Argumentative Completeness of G*)
Every argumentatively GP-immune formula is G-valid.

Proof We prove the theorem indirectly. Suppose there is a formula F that is not G-valid, we will show that F is not argumentatively GP-immune. Since F that is not G-valid, then there is an evaluation e such that $e(F) < 1$. We distinguish two possibilities:

$$e(F) = 0 \quad (52)$$

or

$$e(F) \in (0, 1). \quad (53)$$

If $e(F) = 0$ (52), then $e(\neg F) = 1$ and $\neg F \models_G \neg F$, i.e. $\neg F \rightarrow F$. If $e(F) \in (0, 1)$ (53), then $e(\neg F) = 0$. Therefore, $e(\neg F) \leq e(F)$ and $e(F \supset \neg F) = 0$, i.e. also the formula $F \supset \neg F$ is not G-valid and $\neg(F \supset \neg F) \rightarrow (F \supset \neg F)$. $\square$

In Table 5 we summarise which principle has been used in the corresponding sub-case of Theorem 6.

To recover a complete semantics for product logic, we proceed in a similar way. The attack principles needed are: (A. $\vee$), (A. $\neg$), (C. $\supset$)*, (A. $\supset$)* together with two additional attack principles regarding strong conjunction: (A. &) and (C. &).

- (A. &) If $\langle \Gamma_X, X \rangle \rightarrow \langle \Gamma_A, A \rangle$ or $\langle \Gamma_X, X \rangle \rightarrow \langle \Gamma_B, B \rangle$, then $\langle \Gamma_X, X \rangle \rightarrow \langle \Gamma_{A \& B}, A \& B \rangle$.
- (C. &) If $\langle \Gamma_X, X \rangle \rightarrow \langle \Gamma_{A \& B}, A \& B \rangle$, then $\langle \Gamma_X, X \rangle \rightarrow \langle \Gamma_A, A \rangle$ or $\langle \Gamma_X, X \rangle \rightarrow \langle \Gamma_B, B \rangle$.

Let us now verify that the above mentioned principles are justified in Bi- Π AFs where $\mathcal{A} = \{[D-Reb]\}$ and $\mathcal{S} = \{[D-Sup]\}$.

- (A. $\neg$) If $X \models_{\Pi} \neg \neg A$, then for any evaluation e such that $e(X) = 1$, $e(\neg \neg A) = 1$. Therefore, $e(\neg A) = 0$, i.e. $X \not\models_{\Pi} \neg A$.
- (A. $\vee$) If $X \models_{\Pi} \neg(A \vee B)$, then for any evaluation e such that $e(X) = 1$, then $e(\neg(A \vee B)) = 1$. Therefore, $e(A \vee B) = \max\{e(A), e(B)\} = 0$, from which it follows $e(A) = 0$ and $e(B) = 0$, i.e. $X \models_{\Pi} \neg A$ and $X \models_{\Pi} \neg B$.
- (A. &) If $X \models_{\Pi} \neg A$ or $X \models_{\Pi} \neg B$, then for any evaluation e such that $e(X) = 1$, then either $e(\neg A) = 1$, i.e. $e(A) = 0$, or $e(\neg B) = 1$, i.e. $e(B) = 0$. Therefore, $e(A) \cdot e(B) = 0$, i.e. $X \models_{\Pi} \neg(A \& B)$.
- (C. &) If $X \models_{\Pi} \neg(A \& B)$, then for any evaluation e such that $e(X) = 1$, then $e(\neg(A \& B)) = 1$. Therefore, $e(A \& B) = 0$ which implies either $e(A) = 0$, or $e(B) = 0$.
- (A. $\supset$)* If $X \models_{\Pi} \neg B$ and $X \not\models_{\Pi} \neg A$, then, whenever there is an evaluation e such that $e(X) = 1$, then $e(\neg B) = 1$ and $e(\neg A) < 1$. Therefore, $e(B) = 0$ and $e(A) > 0$ which implies $e(A \supset B) = 0$ and $e(\neg(A \supset B)) = 1$.
- (C. $\supset$)* If $X \models_{\Pi} \neg(A \supset B)$, then for any evaluation e such that $e(X) = 1$, then $e(\neg(A \supset B)) = 1$. Therefore, $e(A \supset B) = 0$ and this happens only if $e(B) = 0$ and $e(A) > 0$. From this it follows $e(\neg B) = 1$ and $e(\neg A) = 0$, i.e. $X \models_{\Pi} \neg B$ and $X \not\models_{\Pi} \neg A$.

Let us define the set ΠP of principles as follows.

$$\Pi P = \{(A.\neg), (A.\vee), (A.\perp), (A. \&), (C. \&), (A.\supset)^*\}$$

Proposition 8 (Closure of ΠP -immune arguments over Modus Ponens) *If A and $A \supset B$ are argumentatively ΠP -immune, then also B is argumentatively ΠP -immune.*

Proof We have to show that if, for any X in the framework,

$$X \not\models_{\Pi} \neg A \quad (54)$$

and

$$X \not\models_{\Pi} \neg(A \supset B), \quad (55)$$

then

$$X \not\models_{\Pi} \neg B. \quad (56)$$

From 54 it follows that there is an evaluation e_1 such that $e_1(X) = 1$ and $e_1(\neg A) < 1$, from which it follows $e_1(A) > 0$. From 55 it follows there is an evaluation e_2 such that $e_2(X) = 1$ and $e_2(\neg(A \supset B)) < 1$, i.e. $e_2(\neg(A \supset B)) = 0$ and $e_2(A \supset B) > 0$. In particular, we have that $e_2(A \supset B) > 0$ if either

$$e_2(A) \leq e_2(B), \quad (57)$$

from which it follows $e_2(A \supset B) = 0$, or

$$e_2(A) > e_2(B) \quad \text{and} \quad e_2(B) > 0, \quad (58)$$

from which it follows $e_2(A \supset B) = \frac{e_2(B)}{e_2(A)}$. If 58 is the case we have already identified an evaluation such that $e(X) = 1$ and $e(\neg B) = 0$ from which it follows $X \not\models_{\Pi} \neg B$. Otherwise, since both X and A belong to the framework, also $X \wedge A$ it does and from the hypothesis we have $X \wedge A \not\models_{\Pi} \neg A$ and $X \wedge A \not\models_{\Pi} \neg(A \supset B)$, i.e. there is an evaluation e_3 such that $e_3(X \wedge A) = 1$ (from which it follows $e_3(X) = 1$ and $e_3(A) = 1$) and $e_3(\neg(A \supset B)) < 1$. Therefore, $e_3(A \supset B) > 0$ and if $e_3(A \supset B) = 1$, then $1 = e_3(A) \leq e_3(B)$ from which it follows $e_3(B) = 1$; if $e_3(A \supset B) \in (0, 1)$, then $e_3(A) > e_3(B) > 0$. Conclusively in both cases $e_3(\neg B) = 0$. $\square$

Theorem 9 (Adequateness Theorem for G) *Any formula F is G-valid iff it is GP-immune.*

Proof Given Proposition 8, it remain to show that every axiom is argumentatively ΠP -immune.

- [BL1] Suppose there is an argument X such that $X \rightarrow (F \supset B) \supset ((G \supset H) \supset (F \supset H))$. By (C. $\supset$)* we have that

$$X \longrightarrow (G \supset H) \supset (F \supset H) \quad (59) \quad X \longrightarrow F \quad (71)$$

and

$$X \not\rightarrow F \supset B. \quad (60) \quad X \not\rightarrow G. \quad (72)$$

From 59 and (C. $\supset$)* it follows

$$X \longrightarrow F \supset H \quad (61) \quad X \not\rightarrow F \quad (73)$$

and

$$X \not\rightarrow G \supset H. \quad (62) \quad X \not\rightarrow F \supset G. \quad (74)$$

From 61 and (C. $\supset$)* it follows

$$X \longrightarrow H \quad (63) \quad X \longrightarrow F \quad (75)$$

and

$$X \not\rightarrow F. \quad (64) \quad X \not\rightarrow G. \quad (76)$$

From 60 and (A. $\supset$)* it follows

$$X \not\rightarrow G \quad (65)$$

or

$$X \longrightarrow F. \quad (66)$$

From 62 and (A. $\supset$)* it follows $X \not\rightarrow H$ or $X \longrightarrow G$. Since 63 holds, then also $X \longrightarrow G$ and $X \longrightarrow F$, but this is in contradiction with 64.

[BL4] Suppose there is an argument X such that $X \longrightarrow F \& (F \supset G) \supset G \& (G \supset F)$. By (C. $\supset$)* we have

$$X \longrightarrow G \& (G \supset F) \quad (67)$$

and

$$X \not\rightarrow F \& (F \supset G). \quad (68)$$

By 67 and (C. &) it follows either

$$X \longrightarrow G \quad (69)$$

or

$$X \longrightarrow G \supset F. \quad (70)$$

From 70 and (C. $\supset$)* it follows

and

$$X \not\rightarrow G. \quad (72)$$

From 68 and (A. &) it follows

$$X \not\rightarrow F \quad (73)$$

and

$$X \not\rightarrow F \supset G. \quad (74)$$

From 74 and (A. $\supset$)* it follows either

$$X \longrightarrow F \quad (75)$$

or

$$X \not\rightarrow G. \quad (76)$$

Since 72 holds, then 75 does not hold and $X \not\rightarrow G$. Therefore, 69 does not hold and 71 it does, but this is in contradiction with 73.

[BL5a] Suppose there is an argument X such that $X \longrightarrow (F \& G \supset H) \supset (F \supset (G \supset H))$. Therefore, by (C. $\supset$)* we have

$$X \longrightarrow F \supset (G \supset H) \quad (77)$$

and

$$X \not\rightarrow F \& G \supset H. \quad (78)$$

From 77 and (C. $\supset$)* we have that

$$X \longrightarrow G \supset H \quad (79)$$

and

$$X \not\rightarrow F. \quad (80)$$

From 79 and (C. $\supset$)* it follows

$$X \longrightarrow H \quad (81)$$

and

$$X \not\rightarrow G. \quad (82)$$

From 78 and (A. $\supset$)* it follows either

- $X \longrightarrow F \& G$ (83)
- or
- $X \longrightarrow G$. (84)
- Since 81 holds, 84 does not and from 83 and (C. &) we have either $X \longrightarrow F$, that cannot be because of 80, or $X \longrightarrow G$ that is in contradiction with 82.
- [BL5b] Suppose there is an argument X such that $X \longrightarrow (F \supset (G \supset H)) \supset (F \& G \supset H)$. By (C. $\supset$)* we have
- $X \longrightarrow (F \& G) \supset H$ (85)
- and
- $X \not\rightarrow F \supset (G \supset H)$. (86)
- From 85 and (C. $\supset$)* it follows
- $X \longrightarrow H$ (87)
- and
- $X \not\rightarrow F \& G$. (88)
- From 88 and (A. &) it follows
- $X \not\rightarrow F$ (89)
- and
- $X \not\rightarrow G$. (90)
- From 86 and (A. $\supset$)* it follows
- $X \not\rightarrow G \supset H$ (91)
- or
- $X \longrightarrow F$ (92)
- and from 81 and (A. $\supset$)* we have
- $X \not\rightarrow H$ (93)
- or
- $X \longrightarrow G$. (94)
- Since 89 holds, 92 does not. Since 87 holds, 93 does not. Therefore, should hold both $X \not\rightarrow G$ 90 and $X \longrightarrow G$ (94), but this cannot happen.
- [BL6] Suppose there is an argument X such that $X \longrightarrow ((F \supset G) \supset H) \supset (((G \supset F) \supset H) \supset H)$. Therefore, by (C. $\supset$)* we have
- $X \longrightarrow ((G \supset F) \supset H) \supset H$ (95)
- and
- $X \not\rightarrow (F \supset G) \supset H$. (96)
- From 95 and (C. $\supset$)* it follows
- $X \longrightarrow H$ (97)
- and
- $X \not\rightarrow (G \supset F) \supset H$. (98)
- From 96 and (A. $\supset$)* we have either
- $X \not\rightarrow H$ (99)
- or
- $X \longrightarrow F \supset G$. (100)
- From 100 and (C. $\supset$)* it follows
- $X \longrightarrow G$ (101)
- and
- $X \not\rightarrow F$. (102)
- From 98 and (A. $\supset$)* it follows either
- $X \not\rightarrow H$ (103)
- or
- $X \longrightarrow G \supset F$ (104)
- from which it follows, by (C. $\supset$)*, $X \longrightarrow F$ and $X \not\rightarrow G$. Since 97 holds, 99 and 103 do not and 100 it does. However, both 101 and 102 are in contradiction with what follows from 104, and this it cannot happen.
- [BL7] Suppose there is an argument X such that $X \longrightarrow \perp \supset F$. Therefore, by (C. $\supset$)* we have $X \longrightarrow F$ and $X \not\rightarrow \perp$, but this it cannot be

since every argument is supposed to attack $\perp$ as (A. $\perp$) requires.

[P] Suppose there is an argument X such that $X \longrightarrow \neg F \vee ((F \supset F \& G) \supset G)$. Therefore, by (A. $\vee$) we have

$$X \longrightarrow \neg F \quad (105)$$

and

$$X \longrightarrow (F \supset (F \& G)) \supset F. \quad (106)$$

From 105 and (A. $\neg$), it follows

$$X \not\rightarrow F. \quad (107)$$

From 106 and (C. $\supset$)* it follows $X \longrightarrow F$ and $X \not\rightarrow F \supset (F \& G)$. However $X \longrightarrow F$ is in contradiction with 105.

Theorem 10 (Bipolar Argumentative Completeness of P) *Every argumentatively ΠP -immune formula is P -valid.*

Proof We prove the theorem indirectly. Suppose there is a formula F that is not P -valid, we will show that F is not argumentatively ΠP -immune. Since F that is not P -valid, then there is an evaluation e such that $e(F) < 1$. We distinguish two possibilities: (a) $e(F) = 0$, and (b) $e(F) \in (0, 1)$.

- (a) If $e(F) = 0$, then $e(\neg F) = 1$ and $\neg F \models_{\Pi} \neg F$, i.e. $\neg F \longrightarrow F$.
- (b) If $e(F) \in (0, 1)$ then $e(\neg F) = 0$. Therefore, $e(\neg F) < e(F)$ and $e(F \supset \neg F) = e(\neg F) = 0$, i.e. also the formula $F \supset \neg F$ is not P -valid and $\neg(F \supset \neg F) \longrightarrow (F \supset \neg F)$. $\square$

6 Related work

In Sect. 2, we recalled the notion of attack principles; in Sect. 4, we introduced bipolar and support principles. The analysis of argumentative frameworks based on such principles can be seen as a postulate-based investigation, a standard mode of analysis in formal argumentation. Caminada and Amgoud did the first works on postulates in argumentation (see Caminada and Amgoud (2005, 2007)). The postulates introduced in these works are more general than the principles recalled and introduced in the previous sections and the correspondence with the two can be found only in specific cases. In addition, also the general perspective is different. In the work of Caminada and Amgoudit, the postulates serve as metrics for measuring the quality of a given

argumentative framework and they are desirable properties that any well-defined framework should satisfy. By contrast, we do not justify any principle a priori but rather analyse, under specific and given circumstances, which of these principles apply and which do not.

In Gorogiannis and Hunter (2011), Gorogiannis and Hunter introduce additional logic-based postulates about the attack relations. Also in this case, under specific circumstances, some of our attack principles coincide with such postulates. These postulates are useful for classifying the attack relations. In fact, the authors use them to characterise some of the attack relations and show that an arbitrary attack relation satisfies a specific list of postulates if and only if the attack relation $\mathcal{R}$ is a specific element of $\mathcal{A} = \{[\text{Def}], [\text{Ucut}], [\text{Reb}], [\text{C-Reb}], [\text{D-Reb}], [\text{I-Reb}]\}$. Let us now recall some of the postulates introduced Gorogiannis and Hunter (2011) and give the corresponding interpretation within the argumentative frameworks defined in Sect. 2. Let $\langle \Gamma_1, \psi_1 \rangle$, $\langle \Gamma_2, \psi_2 \rangle$ and $\langle \Gamma_3, \psi_3 \rangle$ be three deductive arguments (Definition 3).

- (D0) If $\langle \Gamma_1, \psi_1 \rangle \equiv \langle \Gamma'_1, \psi'_1 \rangle$ and $\langle \Gamma_2, \psi_2 \rangle \equiv \langle \Gamma'_2, \psi'_2 \rangle$ (the sign $\equiv$ stands for logical equivalence), then $\langle \Gamma_1, \psi_1 \rangle \mathcal{R}$ -attacks $\langle \Gamma_2, \psi_2 \rangle$ iff $\langle \Gamma'_1, \psi'_1 \rangle \mathcal{R}$ -attacks $\langle \Gamma'_2, \psi'_2 \rangle$.
- (D1) If $\langle \Gamma_1, \psi_1 \rangle \mathcal{R}$ -attacks $\langle \Gamma_2, \psi_2 \rangle$, then $\psi_1 \cup \Gamma_2 \vdash \perp$.
- (D2) If $\langle \Gamma_1, \psi_1 \rangle \mathcal{R}$ -attacks $\langle \Gamma_2, \psi_2 \rangle$ and $\psi_3 \equiv \psi_1$, then $\langle \Gamma_3, \psi_3 \rangle \mathcal{R}$ -attacks $\langle \Gamma_2, \psi_2 \rangle$.

Postulate (D0) is a syntax-independence requirement. The syntax of the components of two arguments should not play a role in defining the attack relations.

Postulate (D1) affirms that if an argument attacks another argument, then the claim of the attacking argument is inconsistent with the support of the attacked one.

Postulate (D2) states that arguments with equivalent claims should attack the same argument. As future work, it would be interesting to characterise different support relations in terms of support principles.

In the recent work by Arieli, Borg and Straßer Arieli et al. (2023), the authors provide an exhaustive postulate-based study of the properties of logical argumentation frameworks. In addition, they give a complete characterization of their semantics and inference relations. In this work, arguments are not defined as in Definition 3 and only minimal requirements are made on their structure. By not requiring that the support sets are neither minimal nor consistent, the authors capture a wide range of logic-based arguments and avoid some computational concerns. With the same perspective as the other cited works, the authors consider several desirable properties and then check in which setting they can be warranted. In the present work, we rather investigate which principles can be formally justified without prior desire for their level of acceptability.

A comprehensive analysis on the rationality postulates in the ASPIC⁺ and ABA systems can be found in D'Agostino and Modgil (2018); Modgil and Prakken (2013, 2018).

7 Conclusions and future work

In the present paper, we have introduced new semantics for the main three t-norm based fuzzy logics based on t-norm based arguments and bipolar frameworks. We have focused on a specific instantiation of both attack and support relation, namely *direct rebuttal* and *direct support*. This choice is related to the fact that the principles considered are defined in terms of the claims of the arguments. Considering argumentative principles defined using also the supports of the arguments, it would be interesting to investigate whether it is possible to recover additional complete semantics for the same logics based on different instantiations of the attack and support relation. The overall methodology used to prove our results is similar to the one used in Corsi and Fermüller (2018a). What changes is the underlying framework both on the level of definition of arguments, we have introduced t-norm based arguments, and on the possible relations among them, we consider instantiations of both the attack and the support relation.

Some of the principles introduced are easier to justify than others. This is not surprising and it has already been the case in the characterisation of classical and fuzzy logics in Corsi and Fermüller (2017) and Corsi and Fermüller (2018a). The specificity of some principles might be seen as a way to look at the characteristics of the various logics from a new perspective and answer to the question “what is needed on the level of argumentative frameworks in order to characterise, e.g., Łukasiewicz logic?”

As future work, we aim to investigate the potential benefits of argumentation-based semantics of fuzzy logics in various types of applications. It would be of particular interest to examine whether the argumentative principles outlined in Section 4 are applicable in real-world scenarios such as discussion analysis and conflict resolution. Thus, it would then be relevant to compare the present work with Gonzalez et al. (2021); Karamlou et al. (2019).

Appendix A: Justification of the attack principles considering $\mathbb{L}$ -based arguments

A.0.1 [Def] and $\models_{\mathbb{L}}$ -based arguments

- (A. $\wedge$) If $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_A, A \rangle$, then $X \models_{\mathbb{L}} \neg \bigwedge \Gamma_A$, i.e. for any evaluation e s.t. $e(X) = 1$, $e(\neg \bigwedge \Gamma_A) = 1$. For every evaluation e we have $e(\neg \bigwedge \Gamma_A) = 1 - e(\bigwedge \Gamma_A) = 1 - \min_{\gamma \in \Gamma_A} (e(\gamma))$. Therefore, considering the argument $\langle \Gamma_{A \wedge B}, A \wedge B \rangle$, if $\Gamma_A \subseteq \Gamma_{A \wedge B}$, we have $1 - \min_{\gamma \in \Gamma_A} (e(\gamma)) \leq 1 - \min_{\gamma \in \Gamma_{A \wedge B}} (e(\gamma))$ and if $1 - \min_{\gamma \in \Gamma_A} (e(\gamma)) = 1$ also $1 - \min_{\gamma \in \Gamma_{A \wedge B}} (e(\gamma)) = 1$ i.e. $X \models_{\mathbb{L}} \neg \bigwedge \Gamma_{A \wedge B}$. The attack principle is justified.
- (C. $\wedge$) If $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_{A \wedge B}, A \wedge B \rangle$, then $X \models_{\mathbb{L}} \neg \bigwedge \Gamma_{A \wedge B}$, i.e. for any evaluation e s.t. $e(X) = 1$, $e(\neg \bigwedge \Gamma_{A \wedge B}) = 1 - \min_{\gamma \in \Gamma_{A \wedge B}} (e(\gamma)) = 1$. Therefore, $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_A, A \rangle$ or $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_B, B \rangle$ only if either $\Gamma_{A \wedge B} \subseteq \Gamma_A$ or $\Gamma_{A \wedge B} \subseteq \Gamma_B$.
- (A. $\vee$) If $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_{A \vee B}, A \vee B \rangle$, then $X \models_{\mathbb{L}} \neg \bigwedge \Gamma_{A \vee B}$, i.e. for any evaluation e s.t. $e(X) = 1$, $e(\neg \bigwedge \Gamma_{A \vee B}) = 1 - \min_{\gamma \in \Gamma_{A \vee B}} (e(\gamma)) = 1$. Therefore, $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_A, A \rangle$ and $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_B, B \rangle$ only if $\Gamma_{A \vee B} \subseteq \Gamma_A$ and $\Gamma_{A \vee B} \subseteq \Gamma_B$.
- (C. $\vee$) If $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_A, A \rangle$ and $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_B, B \rangle$, then $X \models_{\mathbb{L}} \neg \bigwedge \Gamma_A$ and $X \models_{\mathbb{L}} \neg \bigwedge \Gamma_B$, i.e. for any evaluation e s.t. $e(X) = 1$ $e(\neg \bigwedge \Gamma_A) = 1 - \min_{\gamma \in \Gamma_A} (e(\gamma)) = e(\neg \bigwedge \Gamma_B) = 1 - \min_{\gamma \in \Gamma_B} (e(\gamma)) = 1$. Therefore, $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_{A \vee B}, A \vee B \rangle$ only if either $\Gamma_A \subseteq \Gamma_{A \vee B}$ or $\Gamma_B \subseteq \Gamma_{A \vee B}$.
- (A. $\supset$)* If $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_{A \supset B}, A \supset B \rangle$, then for any evaluation e s.t. $e(X) = 1$, then $e(\neg \bigwedge \Gamma_{A \supset B}) = 1 - \min_{\gamma \in \Gamma_{A \supset B}} (e(\gamma)) = 1$. If $\Gamma_{A \supset B} \subseteq \Gamma_B$, then $\min_{\gamma \in \Gamma_B} (e(\gamma)) \leq \min_{\gamma \in \Gamma_{A \supset B}} (e(\gamma))$ and $1 - \min_{\gamma \in \Gamma_{A \supset B}} (e(\gamma)) \leq 1 - \min_{\gamma \in \Gamma_B} (e(\gamma))$

. Therefore, $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_B, B \rangle$. To show that $\langle \Gamma_X, X \rangle \not\xrightarrow{[Def]} \langle \Gamma_A, A \rangle$ follows from the hypothesis, we need to find a specific evaluation e^* s.t. $e^*(X) = 1$, but $1 - \min_{\gamma \in \Gamma_A} (e^*(\gamma)) < 1$ while $1 - \min_{\gamma \in \Gamma_{A \supset B}} (e^*(\gamma)) = 1$

. Given an evaluation e s.t. $e(X) = 1$, if we indicate with γ_i^* the elements of $\Gamma_{A \supset B}$ s.t. $e(\gamma_i^*) = \min_{\gamma \in \Gamma_{A \supset B}} (e(\gamma))$, if $\Gamma_A \subset \Gamma_{A \supset B}$ and $\gamma_i^* \notin \Gamma_A$ for any i , then $\min_{\gamma \in \Gamma_{A \supset B}} (e(\gamma)) < \min_{\gamma \in \Gamma_A} (e(\gamma))$ from which it follows

$1 - \min_{\gamma \in \Gamma_A} (e(\gamma)) < 1 - \min_{\gamma \in \Gamma_{A \supset B}} (e(\gamma))$.

- (C. $\supset$)* If $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_B, B \rangle$ and $\langle \Gamma_X, X \rangle \not\xrightarrow{[Def]} \langle \Gamma_A, A \rangle$, then for any evaluation e s.t. $e(X) = 1$, $e(\neg \bigwedge \Gamma_B) = 1 - \min_{\gamma \in \Gamma_B} (e(\gamma)) = 1$ and there is an evaluation e^* s.t. $e^*(X) = 1$, but $e^*(\neg \bigwedge \Gamma_A) = 1 - \min_{\gamma \in \Gamma_A} (e^*(\gamma)) < 1$. Therefore, $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_{A \supset B}, A \supset B \rangle$ only if $\Gamma_B \subseteq \Gamma_{A \supset B}$.

- (A. $\neg$) The principle is not justified.
(C. $\neg$) The principle is not justified.

A.0.2 [Def] and $\models_{\mathbf{L}}$ -based arguments

- (A. $\wedge$) If $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_A, A \rangle$, then $X \models_{\mathbf{L}} \neg \bigwedge \Gamma_A$, i.e. for every evaluation e , $e(X) \leq e(\neg \bigwedge \Gamma_A) = 1 - \min_{\gamma \in \Gamma_A} (e(\gamma))$. If we consider the argument $\langle \Gamma_{A \wedge B}, A \wedge B \rangle$, if $\Gamma_A \subseteq \Gamma_{A \wedge B}$, we have $1 - \min_{\gamma \in \Gamma_A} (e(\gamma)) \leq 1 - \min_{\gamma \in \Gamma_{A \wedge B}} (e(\gamma))$. Therefore, for every e we have $e(X) \leq e(\neg \bigwedge \Gamma_{A \wedge B})$, i.e. $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_{A \wedge B}, A \wedge B \rangle$ and the principle is justified.
- (C. $\wedge$) If $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_{A \wedge B}, A \wedge B \rangle$, then it follows either $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_A, A \rangle$ or $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_B, B \rangle$ if either $\Gamma_{A \wedge B} \subseteq \Gamma_A$ or $\Gamma_{A \wedge B} \subseteq \Gamma_B$, respectively.
- (A. $\vee$) If $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_{A \vee B}, A \vee B \rangle$, then it follows $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_A, A \rangle$ and $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_B, B \rangle$ only if $\Gamma_{A \vee B} \subseteq \Gamma_A$ and $\Gamma_{A \vee B} \subseteq \Gamma_B$.
- (C. $\vee$) If $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_A, A \rangle$ and

$\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_B, B \rangle$, then it follows

$\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_{A \vee B}, A \vee B \rangle$ only if either $\Gamma_A \subseteq \Gamma_{A \vee B}$ or $\Gamma_B \subseteq \Gamma_{A \vee B}$.

- (A. $\supset$)* If $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_{A \supset B}, A \supset B \rangle$, then for any evaluation e $e(X) \leq e(\neg \bigwedge \Gamma_{A \supset B})$. Therefore, if $\Gamma_{A \supset B} \subseteq \Gamma_B$ $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_B, B \rangle$. However from the hypothesis we cannot deduce that $\langle \Gamma_X, X \rangle \not\xrightarrow{[Def]} \langle \Gamma_A, A \rangle$. The principle is not justified.
- (C. $\supset$)* If $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_B, B \rangle$ and $\langle \Gamma_X, X \rangle \not\xrightarrow{[Def]} \langle \Gamma_A, A \rangle$, then it follows $\langle \Gamma_X, X \rangle \xrightarrow{[Def]} \langle \Gamma_{A \supset B}, A \supset B \rangle$ only if $\Gamma_B \subseteq \Gamma_{A \supset B}$.
- (A. $\neg$) The principle is not justified.
(C. $\neg$) The principle is not justified.

A.0.3 [C-Reb-1] and $\models_{\mathbf{L}}$ -based arguments

- (A. $\wedge$) If $\langle \Gamma_X, X \rangle \xrightarrow{[C-Reb-1]} \langle \Gamma_A, A \rangle$, then $\Gamma_X \models_{\mathbf{L}} \neg A$, i.e. for any evaluation e s.t. $e(\bigwedge \Gamma_X) = 1$, $e(\neg A) = 1 - e(A) = 1$. Since for every evaluation e $e(A) \geq (A \wedge B)$, if $e(A) = 0$ also $e(A \wedge B) = 0$ and $e(\neg(A \wedge B)) = 1$, i.e. $\langle \Gamma_X, X \rangle \xrightarrow{[C-Reb-1]} \langle \Gamma_{A \wedge B}, A \wedge B \rangle$.
- (C. $\wedge$) If $\langle \Gamma_X, X \rangle \xrightarrow{[C-Reb-1]} \langle \Gamma_{A \wedge B}, A \wedge B \rangle$, then for any evaluation e s.t. $e(\bigwedge \Gamma_X) = 1$, $e(\neg(A \wedge B)) = 1 - e(A \wedge B) = 1$, i.e. $\min(e(A), e(B)) = 0$. However, this is different from having either $e(A) = 0$ or $e(B) = 0$. The principle is not justified.
- (A. $\vee$) If $\langle \Gamma_X, X \rangle \xrightarrow{[C-Reb-1]} \langle \Gamma_{A \vee B}, A \vee B \rangle$, then for any evaluation e s.t. $e(\bigwedge \Gamma_X) = 1$, we have $e(\neg(A \vee B)) = 1 - e(A \vee B) = 1$, i.e. $1 - e(A \vee B) = 1 - \max(e(A), e(B)) = 1$ from which it follows $\max(e(A), e(B)) = 0$. Since $e(A) \leq \max(e(A), e(B))$ and $e(B) \leq \max(e(A), e(B))$ for any evaluation e , if $\max(e(A), e(B)) = 0$ also $e(A) = 0$ and $e(B) = 0$, i.e. $\langle \Gamma_X, X \rangle \xrightarrow{[C-Reb-1]} \langle \Gamma_A, A \rangle$ and $\langle \Gamma_X, X \rangle \xrightarrow{[C-Reb-1]} \langle \Gamma_B, B \rangle$.
- (C. $\vee$) If $\langle \Gamma_X, X \rangle \xrightarrow{[C-Reb-1]} \langle \Gamma_A, A \rangle$ and $\langle \Gamma_X, X \rangle \xrightarrow{[C-Reb-1]} \langle \Gamma_B, B \rangle$, then for any evaluation e s.t. $e(\bigwedge \Gamma_X) = 1$ we have $e(\neg A) = 1 - e(A) = 1$ and

$e(\neg B) = 1 - e(B) = 1$, i.e. $e(A) = 0$ and $e(B) = 0$. Therefore, $\max(e(A), e(B)) = 0$, $1 - \max(e(A), e(B)) = 1$ and

$$\langle \Gamma_X, X \rangle \xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_{A \vee B}, A \vee B \rangle.$$

(A. $\supset$)* If $\langle \Gamma_X, X \rangle \xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_{A \supset B}, A \supset B \rangle$

, then for any evaluation e s.t. $e(\bigwedge \Gamma_X) = 1$, $e(\neg(A \supset B)) = 1 - e(A \supset B) = 1$, i.e. $e(A \supset B) = 0$. Since $e(A \supset B) = \min(1, 1 - e(A) + e(B))$, $1 - e(A) + e(B) = 0$, i.e. $1 + e(B) = e(A)$ which implies $e(B) = 0$ and $e(A) = 1$. Therefore, whenever an evaluation e is such that $e(\bigwedge \Gamma_X) = 1$, then $e(\neg B) = 1 - e(B) = 1$ and $e(\neg A) = 1 - e(A) = 0$

, i.e. $\langle \Gamma_X, X \rangle \xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_B, B \rangle$ and

$\langle \Gamma_X, X \rangle \not\xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_A, A \rangle$. The principle is justified.

(C. $\supset$)* If $\langle \Gamma_X, X \rangle \xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_B, B \rangle$

and $\langle \Gamma_X, X \rangle \not\xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_A, A \rangle$

, then whenever an evaluation e is s.t.

$e(\bigwedge \Gamma_X) = 1$, $e(\neg B) = 1 - e(B) = 1$

, i.e. $e(B) = 0$. Since $e(B) = 0$

$e(A \supset B) = \min(1, 1 - e(A) + e(B)) = \min(1, 1 - e(A)) = 1 - e(A)$. In order

to have $\langle \Gamma_X, X \rangle \xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_{A \supset B}, A \supset B \rangle$

we need to show that $e(A \supset B) \leq e(B)$

, i.e. $1 - e(A) \leq e(B)$. In fact, if this

last inequality holds we would have

$1 = 1 - e(B) \leq 1 - 1 + e(A) = 1 - e(A \supset B)$

, i.e. $\langle \Gamma_X, X \rangle \xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_{A \supset B}, A \supset B \rangle$

. However, we are under the hypothesis that $e(B) = 0$. Therefore, we would need $e(A) = 1$ for any e while from the hypothesis that

$\langle \Gamma_X, X \rangle \not\xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_A, A \rangle$ we can only have that for some e^* s.t. $e^*(\bigwedge \Gamma_X) = 1$, $e^*(A) > 0$. The principle is not justified.

(A. $\neg$) If $\langle \Gamma_X, X \rangle \xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_A, A \rangle$, then

for any evaluation e s.t. $e(\bigwedge \Gamma_X) = 1$

, $e(\neg A) = 1 - e(A) = 1$ i.e. $e(A) = 0$.

Therefore, $e(\neg\neg A) = 1 - 1 + e(A) = 0$, i.e.

$\langle \Gamma_X, X \rangle \xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_B, \neg A \rangle$. The principle is justified.

(C. $\neg$) If $\langle \Gamma_X, X \rangle \not\xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_B, \neg A \rangle$, then there is some evaluation e^* s.t. $e^*(\bigwedge \Gamma_X) = 1$ and

$e^*(\neg\neg A) = e^*(A) < 1$. In order to have

$\langle \Gamma_X, X \rangle \xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_A, A \rangle$ we would need

$e^*(\neg A) = 1 - e^*(A) = 1$ i.e. $e^*(A) = 0$, but from the hypothesis we can only infer $e^*(A) < 1$. Therefore, the principle is not justified.

A.0.4 [C-Reb-1] and $\models_L^{\leq}$ -based arguments

(A. $\wedge$) The principle is justified for the same reason (A. $\wedge$) holds using Definition 13.

(C. $\wedge$) The principle is not justified because from $e(\bigwedge \Gamma_X) \leq 1 - e(A \wedge B)$ we cannot deduce neither $e(\bigwedge \Gamma_X) \leq 1 - e(A)$ nor $e(\bigwedge \Gamma_X) \leq 1 - e(B)$ since we only know that $1 - e(A) \leq 1 - \min(e(A), e(B))$.

(A. $\vee$) The principle is justified for the same reason (A. $\vee$) holds using Definition 13.

(C. $\vee$) The principle is not justified.

(A. $\supset$)* If $\langle \Gamma_X, X \rangle \xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_{A \supset B}, A \supset B \rangle$

, then for every evaluation e $e(\bigwedge \Gamma_X) \leq 1 - e(A \supset B)$

. Since $e(A \supset B) = \min(1, 1 - e(A) + e(B))$, if

$1 - e(A) + e(B) < 1$, then $e(A \supset B) = 1 - e(A) + e(B)$

. Therefore, $e(B) \leq e(A \supset B)$. If $1 - e(A) + e(B) \geq 1$

$e(A \supset B) = 1$ and also in this case $e(B) \leq e(A \supset B)$

. In any case $\langle \Gamma_X, X \rangle \xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_B, B \rangle$. However, from the hypothesis we cannot deduce that

$\langle \Gamma_X, X \rangle \not\xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_A, A \rangle$. Thus, the principle is not justified, but its shorter version (A. $\supset$)* holds.

(C. $\supset$)* If $\langle \Gamma_X, X \rangle \xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_B, B \rangle$, then for every evaluation e $e(\bigwedge \Gamma_X) \leq 1 - e(B)$ and in order to show

that $\langle \Gamma_X, X \rangle \xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_{A \supset B}, A \supset B \rangle$ we would need

$1 - e(B) \leq 1 - e(A \supset B)$, but this it cannot be because for every evaluation e , $e(B) \leq e(A \supset B)$. The principle is not justified.

(A. $\neg$) If $\langle \Gamma_X, X \rangle \xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_A, A \rangle$, then for any evaluation e $e(\bigwedge \Gamma_X) \leq e(\neg A) = 1 - e(A)$. In order to show that $\langle \Gamma_X, X \rangle \xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_B, \neg A \rangle$ we would need $e^*(\bigwedge \Gamma_X) > 1 - e^*(\neg A) = e^*(A)$ for some evaluation e^* , but it does not follow from the hypothesis and the principle is not justified.

(C. $\neg$) If $\langle \Gamma_X, X \rangle \not\xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_B, \neg A \rangle$, then there exists some evaluation e^* s.t. $e^*(\bigwedge \Gamma_X) > e^*(A)$ and in order to have $\langle \Gamma_X, X \rangle \xrightarrow{[C\text{-}Reb-1]} \langle \Gamma_A, A \rangle$ we would need that for any evaluation e $e(\bigwedge \Gamma_X) \leq 1 - e(A)$, but this cannot be deduced from the hypothesis. Therefore, the principle is not justified.

A.0.5 [I-Reb] and $\models_{\mathbf{L}}$ -based arguments

(A. $\wedge$) If $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_A, A \rangle$, then there is a formula $\varphi \in Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}} \varphi$ and $A \models_{\mathbf{L}} \neg\varphi$. Therefore, whenever there is an evaluation e s.t. $e(X) = 1$ $e(\varphi) = 1$ and whenever $e(A) = 1$, $e(\neg\varphi) = 1$. Since $e(A \wedge B) = \min\{e(A), e(B)\}$, whenever $e(A \wedge B) = 1$, both $e(A) = 1$ and $e(B) = 1$. From the hypothesis we have that from $e(A) = 1$ it follows $e(\neg\varphi) = 1$. Conclusively $A \wedge B \models_{\mathbf{L}} \neg\varphi$, i.e.

$\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_{A \wedge B}, A \wedge B \rangle$ and the principle is justified.

(C. $\wedge$) If $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_{A \wedge B}, A \wedge B \rangle$, then there is a formula $\varphi \in Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}} \varphi$ and $A \wedge B \models_{\mathbf{L}} \neg\varphi$. Therefore, whenever there is an evaluation e s.t. $e(X) = 1$ $e(\varphi) = 1$ and whenever $e(A \wedge B) = 1$ we have $e(\neg\varphi) = 1$. If $e(A \wedge B) = 1$, then $\min\{e(A), e(B)\} = 1$, i.e. both $e(A) = 1$ and $e(B) = 1$ and from this we cannot deduce that having just, for example, $e(A) = 1$ is enough to conclude $e(\neg\varphi) = 1$. Therefore, the principle is not justified.

(A. $\vee$) If $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_{A \vee B}, A \vee B \rangle$, then there is a formula $\varphi \in Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}} \varphi$ and $A \vee B \models_{\mathbf{L}} \neg\varphi$. Therefore, whenever there is an evaluation e s.t. $e(X) = 1$, then $e(\varphi) = 1$ and whenever $e(A \vee B) = 1$, $e(\neg\varphi) = 1$. If $e(A \vee B) = 1$, then $\max\{e(A), e(B)\} = 1$. Conclusively whenever there is an evaluation e s.t. $e(A) = 1$, $e(A \vee B) = 1$ and $e(\neg\varphi) = 1$. The same holds with B . This implies that both $A \models_{\mathbf{L}} \neg\varphi$ and $B \models_{\mathbf{L}} \neg\varphi$, i.e. $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_A, A \rangle$ and $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_B, B \rangle$ and the principle is justified.

(C. $\vee$) If $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_A, A \rangle$ and $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_B, B \rangle$, then there are φ and φ' in $Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}} \varphi$, $X \models_{\mathbf{L}} \varphi'$, $A \models_{\mathbf{L}} \neg\varphi$ and $B \models_{\mathbf{L}} \neg\varphi'$. Therefore, whenever there is an evaluation e s.t. $e(X) = 1$, then $e(\varphi) = 1$ and $e(\varphi') = 1$, which implies $e(\varphi \wedge \varphi') = \min\{e(\varphi), e(\varphi')\} = 1$. Whenever there is an evaluation e s.t. $e(A \vee B) = 1$, then $\max\{e(A), e(B)\} = 1$. This implies that at least one between $e(A)$ and $e(B)$ is 1. From the hypothesis it follows either $e(\neg\varphi) = 1$ or $e(\neg\varphi') = 1$, i.e. $e(\varphi) = 0$ or $e(\varphi') = 0$. Conclusively we have that

whenever there is an evaluation e s.t. $e(X) = 1$, $e(\varphi \wedge \varphi') = 1$ and whenever $e(A \wedge B) = 1$, $e(\neg(\varphi \wedge \varphi')) = 1 - e(\varphi \wedge \varphi') = 1$, i.e.

$\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_{A \vee B}, A \vee B \rangle$.

(A. $\supset$)* If $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_{A \supset B}, A \supset B \rangle$, then there is a formula $\varphi \in Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}} \varphi$ and $A \supset B \models_{\mathbf{L}} \neg\varphi$, i.e. whenever there is an evaluation e s.t. $e(X) = 1$, $e(\varphi) = 1$ and whenever $e(A \supset B) = 1$, then $e(\neg\varphi) = 1$. Our first claim is that $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_B, B \rangle$

. Therefore, we have to show that there is

$\varphi' \in Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}} \varphi'$ and $B \models_{\mathbf{L}} \neg\varphi'$.

Whenever there is an evaluation e s.t. $e(B) = 1$, since $e(A \supset B) = \min\{1, 1 - e(A) + e(B)\}$, also $e(A \supset B) = 1$ and from the hypothesis we have that $e(\neg\varphi) = 1$. Therefore, $X \models_{\mathbf{L}} \varphi$

and $B \models_{\mathbf{L}} \neg\varphi$, i.e. $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_B, B \rangle$.

However from the hypothesis it does not follow that $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_A, A \rangle$ and the principle is not justified.

(C. $\supset$)* If $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_B, B \rangle$, then there is a formula $\varphi \in Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}} \varphi$ and $B \models_{\mathbf{L}} \neg\varphi$. If $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_A, A \rangle$, then for any formula $\varphi' \in Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}} \varphi'$ $A \not\models_{\mathbf{L}} \neg\varphi'$, i.e. there is at least an evaluation e^* s.t. $e^*(A) = 1$ and $e^*(\neg\varphi') = 1 - e^*(\varphi') < 1$, which implies $e^*(\varphi') > 0$. Therefore, whenever $e(A \supset B) = 1$, then $\min\{1, 1 - e(A) + e(B)\} = 1$ and this happens whenever $e(B) \geq e(A)$. However for some evaluations e^* , $e^*(A) = 1$ and $e^*(\varphi) > 0$. At the same time, since we are under the assumption that $e^*(A \supset B) = 1$, we have $e^*(B) \geq e^*(A) = 1$. Therefore, $e^*(B) = 1$ and from the first hypothesis $e^*(\varphi) = 0$, but this is a contradiction and the principle is not justified.

(A. $\neg$) If $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_A, A \rangle$, then there is a formula $\varphi \in Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}} \varphi$ and $A \models_{\mathbf{L}} \neg\varphi$, i.e. whenever there is an evaluation e s.t. $e(X) = 1$, $e(\varphi) = 1$ and whenever $e(A) = 1$, then $e(\neg\varphi) = 1$. We should then show that there is an evaluation e^* s.t. $e^*(\neg A) = 1$ and $e^*(\neg\varphi) < 1$, but this does not follow from the hypothesis and the principle is not justified.

(C. $\neg$) The attack principle is not justified.

A.0.6 [I-Reb] and $\models_{\mathbf{L}}^{\leq}$ -based arguments

- (A. $\wedge$) If $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_A, A \rangle$, then there is a formula $\varphi \in Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}}^{\leq} \varphi$ and $A \models_{\mathbf{L}}^{\leq} \neg\varphi$. Therefore, for any evaluation e $e(X) \leq e(\varphi)$ and $e(A) \leq e(\neg\varphi)$. Since $e(A \wedge B) = \min\{e(A), e(B)\}$, $e(A \wedge B) \leq e(A)$ from which it follows $e(A \wedge B) \leq e(A) \leq e(\neg\varphi)$, i.e. $A \wedge B \models_{\mathbf{L}}^{\leq} \neg\varphi$ and the principle is justified.
- (C. $\wedge$) If $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_{A \wedge B}, A \wedge B \rangle$, then there is a formula $\varphi \in Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}}^{\leq} \varphi$ and $A \wedge B \models_{\mathbf{L}}^{\leq} \neg\varphi$, i.e. for any evaluation e $e(X) \leq e(\varphi)$ and $e(A \wedge B) \leq e(\neg\varphi)$. However, since $e(A \wedge B) = \min\{e(A), e(B)\} \leq e(A)$ and $e(A \wedge B) = \min\{e(A), e(B)\} \leq e(B)$ from the hypothesis we cannot conclude neither $A \models_{\mathbf{L}}^{\leq} \neg\varphi$ or $B \models_{\mathbf{L}}^{\leq} \neg\varphi$ and the principle is not justified.
- (A. $\vee$) If $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_{A \vee B}, A \vee B \rangle$, then there is a formula $\varphi \in Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}}^{\leq} \varphi$ and $A \vee B \models_{\mathbf{L}}^{\leq} \neg\varphi$, i.e. for every evaluation e $e(X) \leq e(\varphi)$ and $e(A \vee B) \leq e(\neg\varphi)$. Since $e(A \vee B) = \max\{e(A), e(B)\}$ we have both $e(A) \leq e(A \vee B) \leq e(\neg\varphi)$ and $e(B) \leq e(A \vee B) \leq e(\neg\varphi)$, i.e. $A \models_{\mathbf{L}}^{\leq} \neg\varphi$ and $B \models_{\mathbf{L}}^{\leq} \neg\varphi$ and the principle is justified.
- (C. $\vee$) If $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_A, A \rangle$ and $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_B, B \rangle$, then there are φ and φ' in $Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}}^{\leq} \varphi$, $X \models_{\mathbf{L}}^{\leq} \varphi'$, $A \models_{\mathbf{L}}^{\leq} \neg\varphi$ and $B \models_{\mathbf{L}}^{\leq} \neg\varphi'$. Therefore, for every evaluation e $e(X) \leq e(\varphi)$, $e(A) \leq e(\neg\varphi)$, $e(X) \leq e(\varphi')$ and $e(B) \leq e(\neg\varphi')$. Since $e(X) \leq e(\varphi)$ and $e(X) \leq e(\varphi')$, $e(x) \leq e(\varphi \wedge \varphi')$. In fact, whenever $e(\varphi \wedge \varphi') = \min\{e(\varphi), e(\varphi')\} = e(\varphi)$ from $e(X) \leq e(\varphi)$ we have $e(X) \leq e(\varphi \wedge \varphi')$ and the same holds whenever $e(\varphi \wedge \varphi') = e(\varphi')$. Since $e(A \vee B) = \max\{e(A), e(B)\}$, whenever $e(A \vee B) = e(A)$, $e(A \vee B) = e(A) \leq 1 - e(\varphi) \leq 1 - e(\varphi \wedge \varphi')$ and the same holds if $e(A \vee B) = e(B)$. Therefore, in both cases we have that $e(A \wedge B) \leq 1 - e(\varphi \wedge \varphi')$ and the principle is justified.
- (A. $\supset$)* If $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_{A \supset B}, A \supset B \rangle$, then there is a formula $\varphi \in Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}}^{\leq} \varphi$ and $A \supset B \models_{\mathbf{L}}^{\leq} \neg\varphi$, i.e. for every evaluation e

$e(X) \leq e(\varphi)$ and $e(A \supset B) \leq e(\neg\varphi)$. Since for every evaluation e $e(B) \leq 1 - e(A) + e(B)$ and $e(A \supset B) = \min\{1, 1 - e(A) + e(B)\}$, $e(B) \leq e(A \supset B)$ and from the hypothesis $e(B) \leq e(A \supset B) \leq e(\neg\varphi)$, i.e. $B \models_{\mathbf{L}}^{\leq} \neg\varphi$ and $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_B, B \rangle$. However from the hypothesis we cannot deduce $\langle \Gamma_X, X \rangle$

$\xrightarrow{[I-Reb]} \langle \Gamma_A, A \rangle$ and the principle is not justified.

- (C. $\supset$)* If $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_B, B \rangle$, then there is a formula $\varphi \in Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}}^{\leq} \varphi$ and $B \models_{\mathbf{L}}^{\leq} \neg\varphi$, i.e. for every evaluation e $e(X) \leq e(\varphi)$ and $e(B) \leq e(\neg\varphi)$. If $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_A, A \rangle$, then $\varphi' \in Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}}^{\leq} \varphi'$ and $A \models_{\mathbf{L}}^{\leq} \neg\varphi'$, i.e. for some evaluation e^* $e^*(X) \leq e^*(\varphi)$ and $e^*(A) > e^*(\neg\varphi)$. We would need to show that there is a formula $\varphi'' \in Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}}^{\leq} \varphi''$ and $A \supset B \models_{\mathbf{L}}^{\leq} \neg\varphi''$, i.e. for every evaluation e $e(X) \leq e(\varphi)$ and $e(A \supset B) \leq e(\neg\varphi'')$, but this does not follow from the hypothesis and the principle is not justified.
- (A. $\neg$) If $\langle \Gamma_X, X \rangle \xrightarrow{[I-Reb]} \langle \Gamma_A, A \rangle$, then there is a formula $\varphi \in Fm_{\mathbf{L}}$ s.t. $X \models_{\mathbf{L}}^{\leq} \varphi$ and $A \models_{\mathbf{L}}^{\leq} \neg\varphi$, i.e. for every evaluation e $e(X) \leq e(\varphi)$ and $e(A) \leq e(\neg\varphi)$. We would need to show that for some evaluation e^* , $e^*(\neg A) > e^*(\neg\varphi)$ and the principle is not justified.
- (C. $\neg$) The principle is not justified.

Acknowledgements The author acknowledges funding by the Department of Philosophy “Piero Martinetti” of the University of Milan under the Project “Departments of Excellence 2023-2027” awarded by the Italian Ministry of Education, University and Research (MIUR).

Data availability Not applicable.

Declarations

Compliance with Ethical Standards This article does not contain any studies with human participants or animals performed by any of the authors. Furthermore the authors declare that there are no conflicts of interest.

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted

use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

References

- Amgoud L, Besnard P (2009) Bridging the gap between abstract argumentation systems and logic. In *International Conference on Scalable Uncertainty Management*, pages 12–27. Springer
- Amgoud L, Besnard P (2010) A formal analysis of logic-based argumentation systems. In *International Conference on Scalable Uncertainty Management*, pages 42–55. Springer
- Amgoud L, Besnard P, Vesic S (2011) Identifying the core of logic-based argumentation systems. In *2011 IEEE 23rd International Conference on Tools with Artificial Intelligence*, pages 633–636. IEEE
- Amgoud L, Cayrol C (2002) A reasoning model based on the production of acceptable arguments. *Ann Math Artif Intell* 34(1–3):197–215
- Amgoud L, Cayrol C, Lagasque-Schieux M-C, Livet P (2008) On bipolarity in argumentation frameworks. *Int J Intell Syst* 23(10):1062–1093
- Arieli O (2013) A sequent-based representation of logical argumentation. In *International Workshop on Computational Logic in Multi-Agent Systems*, pages 69–85. Springer
- Arieli O, Borg A, Straßer C (2023) A postulate-driven study of logical argumentation. *Artif Intell* 322:103966
- Arieli O, Straßer C (2014) Dynamic derivations for sequent-based logical argumentation. In *Proc. 5th International Conference on Computational Models of Argument (COMMA 2014)*, pages 89–100
- Arieli O, Straßer C (2015) Sequent-based logical argumentation. *Argum Comput* 6(1):73–99
- Arieli O, Straßer C (2016) Deductive argumentation by enhanced sequent calculi and dynamic derivations. *Electron Notes Theor Comput Sci* 323:21–37
- Arieli O, Straßer C (2019) Logical argumentation by dynamic proof systems. *Theor Comput Sci*. <https://doi.org/10.1016/j.tcs.2019.02.019>
- Arieli O, Straßer C (2020) On minimality and consistency tolerance in logical argumentation frameworks. In *Proc. 8th International Conference on Computational Models of Argument (COMMA 2020)*, To appear in *Frontiers in Artificial Intelligence and Applications*. IOS press
- Baaz M, Hájek P, Montagna F, Veith H (2001) Complexity of t-tautologies. *Ann Pure Appl Logic* 113(1–3):3–11
- Besnard P, Hunter A (2001) A logic-based theory of deductive arguments. *Artif Intell* 128(1–2):203–235
- Besnard P, Hunter A (2006) Knowledgebase compilation for efficient logical argumentation. *Proceedings of the Tenth International Conference on Principles of Knowledge Representation and Reasoning* 6:123–133
- Besnard P, Hunter A (2008) *Elements of Argumentation*. MIT Press, Cambridge
- Besnard P, Hunter A (2018) A review of argumentation based on deductive arguments. *Handbook of Formal Argumentation*, pages 437–484
- Boella G, Gabbay DM, van der Torre L, Villata S (2010) Support in abstract argumentation. In *Proceedings of the Third International Conference on Computational Models of Argument (COMMA'10)*, pages 40–51. *Frontiers in Artificial Intelligence and Applications*, IOS Press
- Caminada M, Amgoud L (2005) An axiomatic account of formal argumentation. In *AAAI* 6:608–613
- Caminada M, Amgoud L (2007) On the evaluation of argumentation formalisms. *Artif Intell* 171(5–6):286–310
- Cayrol C, Lagasque-Schieux M-C (2005a) Gradual valuation for bipolar argumentation frameworks. In *European Conference on Symbolic and Quantitative Approaches to Reasoning and Uncertainty*, pages 366–377. Springer
- Cayrol C, Lagasque-Schieux M-C (2005b) On the acceptability of arguments in bipolar argumentation frameworks. In *Symbolic and Quantitative Approaches to Reasoning with Uncertainty: 8th European Conference, ECSQARU 2005, Barcelona, Spain, July 6–8, 2005. Proceedings 8*, pages 378–389. Springer
- Cayrol C, Lagasque-Schieux M-C (2013) Bipolarity in argumentation graphs: towards a better understanding. *Int J Approx Reason* 54(7):876–899
- Cintula P, Fermüller CG, Hájek P, Noguera C (2011, 2015) *Handbook of Mathematical Fuzzy Logic (in 3 volumes)*, volume 37/38/58 of *Studies in Logic, Mathematical Logic and Foundations*. College Publications
- Corsi EA, Fermüller CG (2017) Logical argumentation principles, sequents, and nondeterministic matrices. In *Logic, Rationality and Interaction (LORI 2017)*, volume 10455 of *LNCS*, pages 422–437. Springer
- Corsi EA, Fermüller CG (2018a) Connecting fuzzy logic and argumentation frames via logical attack principles. *Soft Comput*. <https://doi.org/10.1007/s00500-018-3513-2>
- Corsi EA, Fermüller CG (2018b) From semi-abstract argumentation to logical consequence. In *Oswald, S. and Maillat, D., editors, Argumentation and Inference: Proceedings of the 2nd European Conference on Argumentation, Fribourg 2017*, volume 2, pages 151–164. London: College Publications
- D'Agostino M, Modgil S (2018) Classical logic, argument and dialectic. *Artif Intell* 262:15–51
- Dung PM (1995) On the acceptability of arguments and its fundamental role in nonmonotonic reasoning, logic programming and n-person games. *Artif Intell* 77(2):321–357
- Esther Anna C (2025) Attack principles in sequent-based argumentation theory. *J Logic Comput* 35(3):exad080
- Font JM (2015) Consequence and degrees of truth in many-valued logic. In *Petr Hájek on Mathematical Fuzzy Logic*, pages 117–141. Springer
- García AJ, Simari GR (2004) Defeasible logic programming: an argumentative approach. *Theory Pract Log Program* 4(1–2):95–138
- Gentzen G (1935) Untersuchungen über das Logische Schließen I & II. *Mathematische Zeitschrift*, 39(1):176–210, 405–431
- Giles R (1977) A non-classical logic for physics. In: *Wojeicki R, Malinkowski G (eds) Selected Papers on Łukasiewicz Sentential Calculi*. Polish Academy of Sciences, pp 13–51
- Giles R (1982) Semantics for fuzzy reasoning. *Int J Man-Mach Stud* 17(4):401–415
- Gonzalez MGE, Budán MC, Simari GI, Simari GR (2021) Labeled bipolar argumentation frameworks. *J Artif Intell Res* 70:1557–1636
- Gorogiannis N, Hunter A (2011) Instantiating abstract argumentation with classical logic arguments: postulates and properties. *Artif Intell* 175(9–10):1479–1497
- Gottlob G (1992) Complexity results for nonmonotonic logics. *J Log Comput* 2(3):397–425
- Hájek P (1998) *Metamathematics of Fuzzy Logic*. Kluwer Academic Publishers
- Haniková Z (2002) A note on the complexity of propositional tautologies of individual t-algebras. *Neural Netw World* 12(5):453–460
- Karacapilidis N, Papadias D (2001) Computer supported argumentation and collaborative decision making: the hermes system. *Inf Syst* 26(4):259–277
- Karamlou A, Cyra K, Toni F (2019) Deciding the winner of a debate using bipolar argumentation. In *ACM, I. ., editor, AAMAS '19*

- Proceedings of the 18th International Conference on Autonomous Agents and MultiAgent Systems
- Klement EP, Mesiar R, Pap E (2013) Triangular norms, volume 8. Springer Science & Business Media
- Lawry J (1998) A voting mechanism for fuzzy logic. *Int J Approx Reason* 19(3–4):315–333
- Modgil S, Prakken H (2013) A general account of argumentation with preferences. *Artif Intell* 195:361–397
- Modgil S, Prakken H (2018) Abstract rule-based argumentation. In: Baroni P, Gabbay D, Giacomin M, van der Torre L (eds) *Handbook of Formal Argumentation*. College Publications
- Paris JB (1997) A semantics for fuzzy logic. *Soft Comput* 1(3):143–147
- Parsons S, Wooldridge M, Amgoud L (2003) Properties and complexity of some formal inter-agent dialogues. *J Log Comput* 13(3):347–376
- Prakken H, Grasso F, Green NL, Schneider J, Wells S et al (2020) On validating theories of abstract argumentation frameworks: the case of bipolar argumentation frameworks. *Computational Models of Natural Argument CEUR Workshop Proceedings* 2669:21–30
- Proietti C, Grossi D, Smets S, Velázquez-Quesada FR (2019) Bipolar argumentation frameworks, modal logic and semantic paradoxes. *Logic, Rationality, and Interaction: 7th International Workshop, LORI 2019, Chongqing, China, October 18–21, 2019, Proceedings* 7, pages 214–229
- Rahwan I, Simari GR (2009) *Argumentation in artificial intelligence*, vol 47. Springer
- Verheij B (2002) On the existence and multiplicity of extensions in dialectical argumentation. In *Proceedings of the 9th International Workshop on Non-Monotonic Reasoning (NMR 2002)*, pages 416–425
- Yu L, Van der Torre L (2020) A principle-based approach to bipolar argumentation. In *NMR 2020 Workshop Notes*, volume 227

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.