

Assumption-based Argumentation Through the Lens of Approximation Fixpoint Theory

Iosif APOSTOLAKIS^a, Hannes STRASS^b and Johannes P. WALLNER^a

^a*Graz University of Technology*

^b*TUD Dresden University of Technology*

ORCID ID: Iosif Apostolakis <https://orcid.org/0009-0003-2124-3773>, Hannes Strass

<https://orcid.org/0000-0001-6180-6452>, Johannes P. Wallner

<https://orcid.org/0000-0002-3051-1966>

Abstract. Approximation fixpoint theory (AFT) was successfully employed as a uniform approach to define the semantics of major formalisms in abstract argumentation, such as argumentation frameworks and abstract dialectical frameworks. So far AFT was not directly applied for capturing formalisms in structured argumentation. We take up this opportunity and place flat assumption-based argumentation (ABA) in AFT.

1. Introduction

Computational argumentation [1] constitutes the branch of Artificial Intelligence (AI) concerned with studying representations of argument structures, and subsequent methods of argumentative reasoning. Application avenues for approaches to computational argumentation can be found, e.g., in medical or legal reasoning [2].

Key for argumentative reasoning in this field is a formally defined foundation of structure and semantics of arguments. Commonly, one distinguishes on the one hand abstract argumentation [3] and, on the other hand, structured argumentation [4]. The former encompasses formal approaches that represent arguments as abstract entities and which are mainly concerned with relations between arguments to govern reasoning. Argumentative reasoning is then based on argumentation semantics that specify, e.g., which arguments can be deemed acceptable together. Structured argumentation approaches specify workflows [5] from (rule-based) knowledge bases to instantiations of arguments, and subsequent reasoning, often following reasoning as defined in abstract argumentation.

Mirroring heterogeneity in application domains and varying roots such as legal reasoning, logic programming (LP), or non-monotonic reasoning, multiple structured argumentation formalisms have been developed, such as assumption-based argumentation (ABA) [6], ASPIC⁺ [7], defeasible logic programming (DeLP) [8], deductive argumentation [9], Carneades [10], and Gorgias [11].

While several of these formalisms can be connected, each structured argumentation formalism is an expressive framework with several parameters and language fragments, and properties of each formalism typically have to be studied in the specific formalism.

Looking at, arguably, a thematic “neighbour” [12] of structured argumentation, namely LP, we find some general frameworks studying, e.g., whole families of non-monotonic semantics of LPs, such as equilibrium logic [13] and approximation fixpoint theory (AFT) [14,15]. Specifically, the latter one already found its way into foundational investigations of semantics in computational argumentation: main families of semantics of abstract dialectical frameworks (ADFs) [16] have been studied and proposed by placing ADFs into the AFT framework [17,18], and later also the most prominent abstract argumentation formalism, argumentation frameworks (AFs) [3], was placed within AFT.

In addition to laying foundations of semantics and a clean definition, AFT also allows to study properties in the general framework of AFT and then “transfer” them to specific formalisms, e.g., fuzzy-style logics [19], complexity analysis in AFT [20], and non-determinism [21,22].

Nevertheless, AFT was, to the best of our knowledge, not applied directly to formalisms in structured argumentation. In this paper, we present a first step towards “AFT for structured argumentation” by placing the prominent fragment of flat ABA in AFT. Specifically, our main contributions are as follows.

- In order to situate ABA in AFT, a main ingredient is a consequence operator and an approximation thereof. We first present such operators, motivated by similar operators known from logic programming [23,14].
- Utilizing a main contribution of AFT – extending approximators to *stable* approximators – we place major ABA semantics in AFT and formally prove semantic correspondences.
- In a similar vein, we show that a variant operator related to an ABA version of Dung’s [3] characteristic function leads to another way of placing ABA in AFT.

In contrast to earlier work on ABA and AFT [24] (going the reverse way) and utilizing translations [12] of ABA to LP and utilizing existing AFT for LP, we provide a direct instantiation of ABA in AFT via providing operators directly suitable for ABA.

2. Background

We introduce main preliminaries in this section.

Assumption-based Argumentation. Our work builds on the Assumption-Based Argumentation (ABA) frameworks originally proposed by Bondarenko et al. [6]. Central to ABA is a deductive system, defined as a pair $(\mathcal{L}, \mathcal{R})$, where $\mathcal{L}$ denotes a formal language and $\mathcal{R}$ a finite set of inference rules defined over $\mathcal{L}$. Each rule $r \in \mathcal{R}$ has the form $a_0 \leftarrow a_1, \dots, a_n$, where $a_i \in \mathcal{L}$. We denote the head of a rule by $head(r) = a_0$ and its (possibly empty) body by $body(r) = \{a_1, \dots, a_n\}$.

An ABA framework extends this deductive system by designating a subset of $\mathcal{L}$ as assumptions and assigning to each assumption a corresponding contrary.

Definition 1. An ABA framework is a tuple $D = (\mathcal{L}, \mathcal{R}, \mathcal{A}, \bar{\cdot})$, with $(\mathcal{L}, \mathcal{R})$ a deductive system, $\mathcal{A} \subseteq \mathcal{L}$ a non-empty set of assumptions, and $\bar{\cdot}$ a total function mapping assumptions $a \in \mathcal{A}$ to their contrary $\bar{a} \in \mathcal{L}$.

Subsequently we will tacitly assume a given ABA framework $D = (\mathcal{L}, \mathcal{R}, \mathcal{A}, \bar{\cdot})$, unless specified otherwise.

We extend the contrary notation to sets as follows: for $S \subseteq \mathcal{A}$, $\bar{S} = \{\bar{a} \mid a \in S\}$. In this work, we focus on finite and flat ABA frameworks, which means that $\mathcal{L}$ and $\mathcal{R}$ are assumed to be finite, and that for every rule $r \in \mathcal{R}$, it holds that $\text{head}(r) \notin \mathcal{A}$.

For a set A of assumptions we often employ the deductive closure of A , that is, the set of atoms that are derivable from A in D . Formally, for $S \subseteq \mathcal{L}$, we define the function

$$\mathbf{R}_D(S) := \{x \in \mathcal{L} \mid \text{there is } x \leftarrow B \in \mathcal{R} \text{ with } B \subseteq S\}.$$

Consequently, $\mathbf{R}_D^*$ then denotes the “reflexive transitive closure” of $\mathbf{R}_D$, that is, formally $\mathbf{R}_D^*(S) := \bigcup_{i \geq 0} \mathbf{R}_D^i(S)$ with $\mathbf{R}_D^0(S) := S$ and $\mathbf{R}_D^{i+1}(S) := \mathbf{R}_D(\mathbf{R}_D^i(S))$. Intuitively, $\mathbf{R}_D^*(A)$ is then the deductive closure of A . We omit the subscript D from $\mathbf{R}$ when the context is clear. By definition, for any $S \subseteq \mathcal{L}$ and $A \subseteq \mathcal{A}$, we find $A \subseteq \mathbf{R}^*(A)$, and $\mathbf{R}(S) \setminus \mathcal{A} = \mathbf{R}(S)$ (as we assume flat ABA frameworks). It is also clear that $\mathbf{R}$ and $\mathbf{R}^*$ are $\subseteq$ -monotone, that is, $S_1 \subseteq S_2 \subseteq \mathcal{L}$ implies both $\mathbf{R}(S_1) \subseteq \mathbf{R}(S_2)$ and $\mathbf{R}^*(S_1) \subseteq \mathbf{R}^*(S_2)$.

In assumption-based argumentation, the attack relation represents conflicts arising when one set of assumptions derives the contrary of an assumption in another set. This relation can be extended between sets of assumptions, as follows.

Definition 2. Let $A, B \subseteq \mathcal{A}$. We say that A attacks B in D iff there exists some $b \in B$ such that $\bar{b} \in \mathbf{R}^*(A)$. We denote $A^+ := \{a \in \mathcal{A} \mid \bar{a} \in \mathbf{R}^*(A)\}$.

Attacks are used to define semantic concepts such as conflict-freeness and defense.

Definition 3. A set $A \subseteq \mathcal{A}$ is conflict-free in D , written $A \in cf(D)$, iff A does not attack itself. A defends a set $B \subseteq \mathcal{A}$ iff, for any $C \subseteq \mathcal{A}$ that attacks B , we have that A attacks C .

We can then define standard ABA semantics as follows.

Definition 4. Let $A \in cf(D)$.

- A is admissible in D , $A \in adm(D)$, iff A defends itself.
- A is complete in D iff $A \in adm(D)$ and A defends no assumption in $\mathcal{A} \setminus A$.
- A is grounded in D iff it is subset-minimal complete.
- A is preferred in D iff it is subset-maximal admissible.
- A is stable in D , $A \in stb(D)$, iff A attacks every $a \in \mathcal{A} \setminus A$.

Example 1. Consider the ABA D with $\mathcal{A} = \{a, b, c, d\}$, $\mathcal{L} = \mathcal{A} \cup \bar{\mathcal{A}}$, and deduction rules $\mathcal{R} = \{\bar{b} \leftarrow a, \bar{c} \leftarrow b, \bar{d} \leftarrow c, \bar{c} \leftarrow d\}$. Here, the set $\{a\}$ is a conflict-free set, as $\bar{a}$ cannot be derived by assumption a . Additionally, this set is admissible and complete, as it defends just itself. The set is also grounded, as it is the $\subseteq$ -least complete set. However, this set is neither stable nor preferred. The sets $\{a, c\}$ and $\{a, d\}$ on the other hand are both stable and preferred.

Approximation Fixpoint Theory. An operator is a (total) function $O: V \rightarrow V$ on a set V . We are typically interested in operators on sets V with an internal structure: At the minimum, we require $(V, \leq)$ to be a partially ordered set. More often, $(V, \leq)$ will be a complete lattice, that is, such that every subset $S \subseteq V$ has a least upper bound (*lub*) $\bigvee S \in V$ and greatest lower bound (*glb*) $\bigwedge S \in V$; this entails that $(V, \leq)$ has both a least

element $0 = \bigwedge V = \bigvee \emptyset$ and a greatest element $1 = \bigvee V = \bigwedge \emptyset$. An operator $O: V \rightarrow V$ is *monotone* iff $x \leq y$ implies $O(x) \leq O(y)$, and *anti-monotone* iff $x \leq y$ implies $O(y) \leq O(x)$. An element $x \in V$ with $O(x) = x$ is a *fixpoint* of O ; if $O(x) \leq x$ then x is a *pre-fixpoint* of O ; if $x \leq O(x)$ then x is a *post-fixpoint* of O . A fundamental result in lattice theory, Tarski's fixpoint theorem [25], states that whenever O is a monotone operator on a complete lattice $(V, \leq)$, the set $\{x \in V \mid O(x) = x\}$ of its fixpoints forms a complete lattice itself, and so has a least element $\text{lfp}(O)$. While this result is very useful, it is not constructive; however, constructive versions exist and tell us how to actually construct least fixpoints [26,27]: To this end, one defines, for any $x \in V$, the iterated (transfinite) application of O to x via $O^0(x) := x$, $O^{\alpha+1}(x) := O(O^\alpha(x))$ for successor ordinals, and $O^\beta(x) := \bigvee \{O^\alpha(x) \mid \alpha < \beta\}$ for limit ordinals. It is then known that (for monotone O) there exists an ordinal α such that $O^\alpha(0) = \text{lfp}(O)$ [28].

In order to apply Tarski's fixpoint theorem [25] (or its constructive versions), it is thus necessary to have an operator that is monotone, a condition that is no longer fulfilled when extending the syntax of logic programs from *definite* to *normal* logic programs, i.e., when allowing "negation as failure" to occur in bodies of rules. To still have a useful fixpoint theory in the case of non-monotone operators, Denecker, Marek, and Truszczyński (henceforth DMT) [14] fundamentally generalized the theory underlying the fixpoint-based approach to semantics, thereby founding what is now known as *approximation fixpoint theory*. Its underlying idea is that when an operator of interest does not have properties that guarantee the existence of fixpoints, to then *approximate* this operator in a more fine-grained algebraic structure where fixpoint existence is guaranteed.

Formally – following ideas of Belnap [29], Ginsberg [30], and Fitting [31] –, DMT [14] moved from a complete lattice $(V, \leq)$ to its associated *bilattice* on the set $V^2 := V \times V$. So where elements of V correspond to candidates for the semantics (interpretations), the pairs contained in V^2 correspond to *approximations* of such candidates. More technically, a pair $(x, y) \in V^2$ approximates all $z \in V$ with $x \leq z \leq y$. A pair $(x, y) \in V^2$ is called *consistent* iff $x \leq y$, in other words, if the interval $[x, y] := \{z \in V \mid x \leq z \leq y\}$ is non-empty; a pair of the form (x, x) is *exact*. There are two natural orderings on V^2 :

- $(x_1, y_1) \leq_t (x_2, y_2) :\iff x_1 \leq x_2 \text{ and } y_1 \leq y_2$, the *truth ordering* extending the (truth) lattice ordering $\leq$;
- $(x_1, y_1) \leq_p (x_2, y_2) :\iff x_1 \leq x_2 \text{ and } y_2 \leq y_1$, the *precision ordering* comparing intervals by precision of approximation.

Just as pairs of V^2 approximate elements of V , a foundational idea of Denecker et al. [14] was that operators $O: V \rightarrow V$ can be approximated by operators on V^2 : An *approximator* is an operator $\mathcal{A}: V^2 \rightarrow V^2$ that is $\leq_p$ -monotone and that maps exact pairs to exact pairs. By virtue of the latter condition, we say that $\mathcal{A}$ *approximates* the operator $O: V \rightarrow V$ with $O(x) = \mathcal{A}(x, x)_1$ (where $(x, y)_1 = x$ and $(x, y)_2 = y$); by virtue of the first condition, $\mathcal{A}$ is guaranteed to possess ($\leq_p$ -least) fixpoints. An approximator $\mathcal{A}$ is *symmetric* iff $\mathcal{A}(x, y)_2 = \mathcal{A}(y, x)_1$ for all $(x, y) \in V^2$; this allows to specify $\mathcal{A}$ by giving only $\mathcal{A}(\cdot, \cdot)_1$.

The main contribution of DMT [14] was the association of the *stable approximator* $\mathcal{A}^{st}$ to an approximator $\mathcal{A}$: For a complete lattice $(V, \leq)$ and an approximator $\mathcal{A}: V^2 \rightarrow V^2$, define the *stable approximator for* $\mathcal{A}$ as $\mathcal{A}^{st}: V^2 \rightarrow V^2$ by $\mathcal{A}^{st}(x, y) := (\text{lfp}(\mathcal{A}(\cdot, y)_1), \text{lfp}(\mathcal{A}(x, \cdot)_2))$, where $\mathcal{A}(\cdot, y)_1$ (resp. $\mathcal{A}(x, \cdot)_2$) denotes the operator that maps any $z \in V$ to $\mathcal{A}(z, y)_1$ (resp. to $\mathcal{A}(x, z)_2$). Among other results, DMT [14] also showed that this construction is well-defined because both relevant oper-

ators $\mathcal{A}(\cdot, y)_1 : V \rightarrow V$ and $\mathcal{A}(x, \cdot)_2 : V \rightarrow V$ are $\leq$ -monotone (as $\mathcal{A}$ is an approximator) and thus possess least fixpoints. Furthermore, any approximator maps consistent pairs to consistent pairs, and the stable approximator $\mathcal{A}^{st}$ maps consistent post-fixpoints of $\mathcal{A}$ to consistent pairs. Consequently, not only is $\text{lfp}(\mathcal{A})$ consistent (and called the *Kripke-Kleene fixpoint* of $\mathcal{A}$), but so is $\text{lfp}(\mathcal{A}^{st})$, the *well-founded fixpoint* of $\mathcal{A}$. Moreover, a fixpoint (x, y) of $\mathcal{A}^{st}$, called a *stable fixpoint* of $\mathcal{A}$, is always a $\leq$ -minimal fixpoint of O (when $\mathcal{A}$ approximates O). The term “stable” is not coincidental: When using the logic program consequence operator T_P by van Emden and Kowalski [23] as O and defining a suitable approximator $\mathcal{A}$, then the exact fixpoints of $\mathcal{A}^{st}$ correspond one-to-one to the stable models of the logic program P [14].

3. An Operator for ABA and Its Approximator

In this section we first reformulate standard ABA semantics by means of two basic operators, namely deductive closure $\mathbf{R}^*$ and an operator $\mathbf{A}$ mapping a set to the assumptions it does not contradict. We then introduce an operator on $2^{\mathcal{L}}$ and an associated approximator on the bilattice, which will allow us to reconstruct ABA semantics in AFT.

In the previous section, we defined the operator $\mathbf{R}^*$ mapping sets of atoms to their deductive closure. Now we also define the operator $\mathbf{A}$ that maps a set S of atoms to the set of assumptions whose contraries do not lie in S . Formally:

$$\mathbf{A}(S) := \{a \in \mathcal{A} \mid \bar{a} \notin S\}$$

By construction, $\mathbf{A}$ is $\subseteq$ -antimonotone, so by $A \subseteq \mathbf{R}^*(A)$ we get $\mathbf{A}(\mathbf{R}^*(A)) \subseteq \mathbf{A}(A)$ for any $A \subseteq \mathcal{A}$. With operators $\mathbf{A}(\cdot)$ and $\mathbf{R}^*(\cdot)$, we are able to equivalently reformulate semantical notions of ABA as follows. We will use the following observation to characterize the set of assumptions being attacked by a set $S \subseteq \mathcal{L}$:

$$S^+ = \{a \in \mathcal{A} \mid \bar{a} \in \mathbf{R}^*(S)\} = \mathcal{A} \setminus \{a \in \mathcal{A} \mid \bar{a} \notin \mathbf{R}^*(S)\} = \mathcal{A} \setminus \mathbf{A}(\mathbf{R}^*(S))$$

Consequently, the set of assumptions $\mathbf{A}(\mathbf{R}^*(S))$ contains exactly those assumptions that are *not* attacked by S . This observation yields the following reformulation of the standard ABA semantics.

Proposition 1. *Given a set $A \subseteq \mathcal{A}$, the following hold:*

- *A is conflict-free iff $A \subseteq \mathbf{A}(\mathbf{R}^*(A))$,*
- *A defends b iff $b \in \mathbf{A}(\mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(A))))$,*
- *A is admissible iff $A \subseteq \mathbf{A}(\mathbf{R}^*(A))$ and $A \subseteq \mathbf{A}(\mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(A))))$,*
- *A is complete iff $A \subseteq \mathbf{A}(\mathbf{R}^*(A))$ and $A = \mathbf{A}(\mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(A))))$,*
- *A is stable iff $A = \mathbf{A}(\mathbf{R}^*(A))$.*

Proof. For the first item, we have that A is conflict-free iff $a \in A \implies a \notin A^+$ iff $a \in A \implies \bar{a} \notin \mathbf{R}^*(A)$ iff $a \in A \implies a \in \mathbf{A}(\mathbf{R}^*(A))$. For the second item, A defends b iff $\forall C : b \in C^+ \implies C \cap A^+ \neq \emptyset$ iff $\forall C : b \in C^+ \implies C \cap (\mathcal{A} \setminus \mathbf{A}(\mathbf{R}^*(A))) \neq \emptyset$ iff $\forall C : b \notin \mathbf{A}(\mathbf{R}^*(C)) \implies C \not\subseteq \mathbf{A}(\mathbf{R}^*(A))$ iff $\forall C : C \subseteq \mathbf{A}(\mathbf{R}^*(A)) \implies b \in \mathbf{A}(\mathbf{R}^*(C))$ iff $b \in \mathbf{A}(\mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(A))))$ where the last “if” direction holds because $\mathbf{A}(\mathbf{R}^*(\cdot))$ is $\subseteq$ -antimonotone whence $C \subseteq \mathbf{A}(\mathbf{R}^*(A))$ implies $b \in \mathbf{A}(\mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(A)))) \subseteq \mathbf{A}(\mathbf{R}^*(C))$.

Using these two, the remaining items follow. $\square$

Example 2. Consider the ABA given in Example 1. Then $\mathbf{A}(\mathbf{R}^*({a})) = \mathbf{A}({a, \bar{b}}) = \{a, c, d\}$ and therefore $\{a\} \in cf(D)$. Further, $\mathbf{A}(\mathbf{R}^*(\mathbf{A}(\mathbf{R}^*({a})))) = \mathbf{A}(\mathbf{R}^*({a, c, d})) = \{a\}$, which implies that $\{a\}$ is also admissible and complete. However, $\{a\}$ is not stable since $\{a\} \neq \{a, c, d\}$. The set $\{a, c\}$ though, is stable, as $\mathbf{A}(\mathbf{R}^*({a, c})) = \{a, c\}$.

The results of Proposition 1 could be stated even more succinctly using an operator that is inspired by the work of Pollock [32] and that we will introduce in the next section; we did not want to anticipate this definition, and also are operating in a different lattice here ($2^{\mathcal{L}}$ instead of $2^{\mathcal{A}}$). In order to motivate the ABA “one-step consequence” operator we are going to introduce, we first define some intermediate semantical notions.

Definition 5. Let D be an ABA framework.

1. A set $S \subseteq \mathcal{L}$ is non-contradictory iff for all $a \in \mathcal{A}$, we have $a \notin S$ or $\bar{a} \notin S$.
2. A set $S \subseteq \mathcal{L}$ is supported for D iff $\mathbf{R}(S) = S \setminus \mathcal{A}$ and S is $\subseteq$ -maximally non-contradictory.

In other words, a non-contradictory set S never contains an assumption *and* its contrary. A set $S \subseteq \mathcal{L}$ being *maximally* non-contradictory thus essentially means that for every assumption $a \in \mathcal{A}$, the set S *either* contains a *or* its contrary $\bar{a}$. So “supported” here means something different for assumptions and non-assumptions: An assumption a being supported in S refers to the absence of its contrary $\bar{a}$ from S , with S also containing *all* assumptions whose contraries are absent from it; a non-assumption x having support in S refers to x being derivable by a rule $x \leftarrow B$ from S , and S conversely also containing *all* atoms that are derivable from it. The name “supported” is inspired by a similar notion from normal logic programs [33].

Example 3. Let $\mathcal{A} = \{a, b, c, d\}$ and $\mathcal{R} = \{\bar{b} \leftarrow a, \bar{c} \leftarrow b, \bar{c} \leftarrow d\}$. Then the set $S = \{a, c, \bar{b}\}$, is non-contradictory as for all assumptions it contains, their contraries do not lie in the set. However, it is not $\subseteq$ -maximally non-contradictory, as the set $\{a, c, d, \bar{b}\}$ is also non-contradictory. Now the latter set is not supported as it does not contain $\bar{c}$ which can be derived by rule $\bar{c} \leftarrow d$. Finally, the set $\{a, d, \bar{b}, \bar{c}\}$ is supported.

We consider the complete lattice $(2^{\mathcal{L}}, \subseteq)$ of all sets of atoms ordered by set inclusion. The operator $\mathbf{B}: 2^{\mathcal{L}} \rightarrow 2^{\mathcal{L}}$ we introduce assigns to any set $S \subseteq \mathcal{L}$ a new set $\mathbf{B}(S)$ that represents one step of consequence operation and assumption update with respect to D . Intuitively, the non-assumptions in $\mathbf{B}(S)$ are those atoms that are derivable from S via rules in $\mathcal{R}$ (in one step of rule application), while the assumptions in $\mathbf{B}(S)$ are determined by whether their contraries are present in S .

Definition 6. Let D be an ABA framework. Define the operator $\mathbf{B}_D: 2^{\mathcal{L}} \rightarrow 2^{\mathcal{L}}$ via

$$\mathbf{B}_D(S) := \mathbf{R}(S) \cup \mathbf{A}(S)$$

As usual, we drop the subscript and write $\mathbf{B}_D$ as $\mathbf{B}$ when D is clear from context. In our first reconstruction using $\mathbf{B}$, we show that its fixpoints characterize supported sets.

Proposition 2. A set $S \subseteq \mathcal{L}$ is a fixpoint of $\mathbf{B}$ iff S is supported for D .

Proof. S is a fixpoint of $\mathbf{B}$ iff $\mathbf{B}(S) = S$ iff $\mathbf{R}(S) \cup \mathbf{A}(S) = S$ iff $\mathbf{R}(S) \setminus \mathcal{A} = S \setminus \mathcal{A}$ and $\mathbf{A}(S) \cap \mathcal{A} = S \cap \mathcal{A}$ iff $\mathbf{R}(S) = S \setminus \mathcal{A}$ and $\mathbf{A}(S) = S \cap \mathcal{A}$ iff $\mathbf{R}(S) = S \setminus \mathcal{A}$ and S is maximally non-contradictory iff S is supported for D .

It remains to show that $\mathbf{A}(S) = S \cap \mathcal{A}$ iff S is maximally non-contradictory.

To this end, consider any $a \in \mathcal{A}$. If $\mathbf{A}(S) = S \cap \mathcal{A}$ and $a \in S$, then $a \in \mathbf{A}(S)$ and thus $\bar{a} \notin S$, so S is non-contradictory. If $a \notin S$, then $a \notin \mathbf{A}(S)$ whence $\bar{a} \in S$; if $\bar{a} \notin S$, then $a \in \mathbf{A}(S) = S$, thus S is maximally non-contradictory.

Conversely, let S be maximally non-contradictory. If $a \in S$, then $\bar{a} \notin S$ whence $a \in \mathbf{A}(S)$. If $a \in \mathbf{A}(S)$, then $\bar{a} \notin S$ and since S is maximally non-contradictory we get $a \in S$. Thus $\mathbf{A}(S) = S \cap \mathcal{A}$. $\square$

In particular, for any fixpoint S of $\mathbf{B}$, every non-assumption in S is justified by a rule with body contained in S , and every assumption that can be consistently accepted given S is contained in S . Conversely, any supported and maximally non-contradictory set S satisfies $\mathbf{B}(S) = S$. This vindicates our choice of $\mathbf{B}$ as an ABA consequence operator from the point of view of logic programming, where the characterization of supported model semantics via the consequence operator's fixpoints is well-known [31].

To connect with approximation fixpoint theory, we move from $(2^{\mathcal{L}}, \subseteq)$ to the bilattice $(2^{\mathcal{L}} \times 2^{\mathcal{L}}, \leq_p)$. Recall that for $I, U \subseteq \mathcal{L}$, we put $(I, U) \leq_p (I', U')$ iff $I \subseteq I'$ and $U' \subseteq U$. Pairs (I, U) with $I \subseteq U$ are called consistent and approximate all sets S such that $I \subseteq S \subseteq U$. There, I collects atoms that are considered to be *in*, while U contains those that are *in* or *undecided*, similarly as in abstract argumentation [34].

An approximator for $\mathbf{B}$ is an operator $A: 2^{\mathcal{L}} \times 2^{\mathcal{L}} \rightarrow 2^{\mathcal{L}} \times 2^{\mathcal{L}}$ that is $\leq_p$ -monotone and maps exact pairs (S, S) to $(\mathbf{B}(S), \mathbf{B}(S))$, in the standard sense of approximation fixpoint theory. We will now define a concrete approximator for $\mathbf{B}$ that we will use later on to reconstruct ABA semantics and define new ones.

Definition 7. Let D be an ABA framework. We define a symmetric operator

$$\mathbf{C}_D: (2^{\mathcal{L}} \times 2^{\mathcal{L}}) \rightarrow (2^{\mathcal{L}} \times 2^{\mathcal{L}}) \quad \text{with} \quad \mathbf{C}_D(I, U)_1 = \mathbf{R}(I) \cup \mathbf{A}(U).$$

The first component thus collects all atoms derivable by $\mathcal{R}$ in one step from the current lower bound I , together with those assumptions that are compatible with the upper bound U , while the second component is obtained symmetrically. (In other words, $\mathbf{C}_D(I, U)_2 := \mathbf{R}(U) \cup \mathbf{A}(I)$.) We will typically drop the subscript D when it is clear from context. On exact pairs (S, S) this operator coincides with $\mathbf{B}$, since both the derivability and assumption parts are evaluated with respect to the same set S . It is straightforward to verify that $\mathbf{C}$ is $\leq_p$ -monotone.

Lemma 3. $\mathbf{C}$ is $\leq_p$ -monotone.

Proof. Let $(L_1, U_1) \leq_p (L_2, U_2)$. Then $L_1 \subseteq L_2$ and $U_2 \subseteq U_1$. Consider first $x \in \mathcal{L} \setminus \mathcal{A}$. If $x \in \mathbf{C}(L_1, U_1)_1$, then $x \in \mathbf{R}(L_1) \subseteq \mathbf{R}(L_2) \subseteq \mathbf{C}(L_2, U_2)_1$ by monotonicity of $\mathbf{R}(\cdot)$ directly. Consider now $a \in \mathcal{A}$. If $a \in \mathbf{C}(L_1, U_1)_1$, that is, $a \in \mathbf{A}(U_1)$, then $\bar{a} \notin U_1$, whence by $U_2 \subseteq U_1$ we have $\bar{a} \notin U_2$, thus $a \in \mathbf{A}(U_2) \subseteq \mathbf{C}(L_2, U_2)_1$. $\square$

Consequently, $\mathbf{C}$ is an approximator of $\mathbf{B}$ on the bilattice $(2^{\mathcal{L}} \times 2^{\mathcal{L}}, \leq_p)$. This allows us to apply the stable operator construction to $\mathbf{C}$ and relate its various fixpoints to the standard ABA semantics. Approximation fixpoint theory associates with $\mathbf{C}$ the stable

April 2022

operator $\mathbf{C}^{st}: 2^{\mathcal{L}} \times 2^{\mathcal{L}} \rightarrow 2^{\mathcal{L}} \times 2^{\mathcal{L}}$ defined by $\mathbf{C}^{st}(I, U) = (\text{lfp}(\mathbf{C}(\cdot, U)_1), \text{lfp}(\mathbf{C}(\cdot, I)_1))$. We prepare for subsequent reconstructions by taking a closer look at this $\mathbf{C}^{st}$.

Lemma 4. *For any $S \subseteq \mathcal{L}$, we have $\text{lfp}(\mathbf{C}(\cdot, S)_1) = \mathbf{R}^*(\mathbf{A}(S))$.*

Proof. “ $\subseteq$ ”: By definition, $\mathbf{A}(S) \subseteq \mathbf{R}^*(\mathbf{A}(S))$ and $\mathbf{R}(\mathbf{R}^*(\mathbf{A}(S))) = \mathbf{R}^*(\mathbf{A}(S))$. Thus we get that $\mathbf{C}(\mathbf{R}^*(\mathbf{A}(S)), S)_1 = \mathbf{R}(\mathbf{R}^*(\mathbf{A}(S))) \cup \mathbf{A}(S) = \mathbf{R}^*(\mathbf{A}(S))$. Hence $\mathbf{R}^*(\mathbf{A}(S))$ is a fixpoint of $\mathbf{C}(\cdot, S)_1$ and therefore $\text{lfp}(\mathbf{C}(\cdot, S)_1) \subseteq \mathbf{R}^*(\mathbf{A}(S))$.

“ $\supseteq$ ”: We will show that for all $j \geq 0$, $\mathbf{R}^j(\mathbf{A}(S)) \subseteq \text{lfp}(\mathbf{C}(\cdot, S)_1)$ by induction on j , where we denote $C_S^* := \text{lfp}(\mathbf{C}(\cdot, S)_1)$.

For $j = 0$ we get $\mathbf{R}^0(\mathbf{A}(S)) = \mathbf{A}(S) \subseteq \mathbf{R}(\emptyset) \cup \mathbf{A}(S) = \mathbf{C}(\emptyset, S)_1 \subseteq C_S^*$.

In the induction step, consider $\mathbf{R}^{j+1}(\mathbf{A}(S)) = \mathbf{R}(\mathbf{R}^j(\mathbf{A}(S)))$. By the induction hypothesis, $\mathbf{R}^j(\mathbf{A}(S)) \subseteq C_S^*$. By $\subseteq$ -monotonicity of $\mathbf{R}$, we obtain directly that $\mathbf{R}^{j+1}(\mathbf{A}(S)) = \mathbf{R}(\mathbf{R}^j(\mathbf{A}(S))) \subseteq \mathbf{R}(C_S^*) \subseteq \mathbf{R}(C_S^*) \cup \mathbf{A}(S) = \mathbf{C}(C_S^*, S)_1 = C_S^*$. $\square$

Thus it follows that $\mathbf{C}^{st}(I, U) = (\mathbf{R}^*(\mathbf{A}(U)), \mathbf{R}^*(\mathbf{A}(I)))$. We next relate fixpoints of $\mathbf{C}^{st}$ to complete sets. Intuitively, a consistent fixpoint (I, U) of $\mathbf{C}^{st}$ represents a three-valued view of the ABA framework.

Proposition 5. *If a consistent pair (I, U) is a fixpoint of $\mathbf{C}^{st}$, then $I \cap \mathcal{A}$ is a complete set. Conversely, if a set $S \subseteq \mathcal{A}$ is a complete set, then the pair $(\mathbf{R}^*(S), \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(S))))$ is a consistent fixpoint of $\mathbf{C}^{st}$.*

Proof. Let $S \in \text{com}(D)$. We will show that the pair $(\mathbf{R}^*(S), \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(S))))$ constitutes a fixpoint of $\mathbf{C}^{st}$, i.e. $\mathbf{C}^{st}(\mathbf{R}^*(S), \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(S)))) = (\mathbf{R}^*(S), \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(S))))$.

By Lemma 4 we get $\text{lfp}(\mathbf{C}(\cdot, \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(S))))_1 = \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(S))))) = \mathbf{R}^*(S)$, by Proposition 1 because S is complete. Additionally, by Lemma 4, $\text{lfp}(\mathbf{C}(\cdot, \mathbf{R}^*(S))_1) = \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(S)))$. Then overall for a complete set S it holds that $(\mathbf{R}^*(S), \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(S))))$ constitutes a fixpoint of $\mathbf{C}^{st}$. It remains to show that the pair $(\mathbf{R}^*(S), \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(S))))$ is consistent, that is, that $\mathbf{R}^*(S) \subseteq \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(S)))$. To this end, S being complete by Proposition 1 yields that $S \subseteq \mathbf{A}(\mathbf{R}^*(S))$ which by $\subseteq$ -monotonicity yields the claim.

For the other direction, assume $(I, U) = \mathbf{C}^{st}(I, U)$ with $I \subseteq U$. We will prove that $S := I \cap \mathcal{A}$ is a complete set in D . To start, $S = \mathbf{A}(U)$ since $I = \mathbf{R}(I) \cup \mathbf{A}(U)$. Also, $I = \text{lfp}(\mathbf{C}(\cdot, U)_1) = \mathbf{R}^*(\mathbf{A}(U))$ by Lemma 4. Symmetrically, we obtain $U = \mathbf{R}^*(\mathbf{A}(I))$. Since by consistency $I \subseteq U$ and $\mathbf{A}$ is antimonotone, we get $\mathbf{A}(U) \subseteq \mathbf{A}(I) = \mathbf{A}(\mathbf{R}^*(\mathbf{A}(U)))$, which shows that S is conflict-free by Proposition 1. Additionally it holds that $\mathbf{A}(U) = \mathbf{A}(\mathbf{R}^*(\mathbf{A}(I))) = \mathbf{A}(\mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(\mathbf{A}(U)))))$. Therefore, $\mathbf{A}(U) = I \cap \mathcal{A}$ is complete. $\square$

Stable sets are those complete sets that also attack every assumption they do not contain. On the approximation side, this corresponds exactly to exact fixpoints of the stable operator.

Proposition 6. *If (I, I) is a fixpoint of $\mathbf{C}^{st}$ then $I \cap \mathcal{A}$ is a stable set of D ; conversely, if $A \subseteq \mathcal{A}$ is a stable set of D , then $(\mathbf{R}^*(A), \mathbf{R}^*(A))$ is a fixpoint of $\mathbf{C}^{st}$.*

Proof. If (I, I) is a fixpoint of $\mathbf{C}^{st}$ then $I = \mathbf{R}^*(\mathbf{A}(I))$ by Lemma 4. Thus $\mathbf{R}^*(I) = \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(I)))$ by monotonicity of $\mathbf{R}^*$. By Proposition 1, $\mathbf{R}^*(I) \cap \mathcal{A} = I \cap \mathcal{A}$ is stable.

Conversely, if $A \subseteq \mathcal{A}$ is stable, Proposition 1 yields that $A = \mathbf{A}(\mathbf{R}^*(A))$. Thus $\mathbf{R}^*(A) = \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(A)))$ and by Lemma 4, we get $\mathbf{R}^*(A) = \text{lfp}(\mathbf{C}^{st}(\cdot, \mathbf{R}^*(A))_1)$, whence the pair $(\mathbf{R}^*(A), \mathbf{R}^*(A))$ is a fixpoint of $\mathbf{C}^{st}$. $\square$

Admissible sets can be characterized by the stable operator, as pairs whose refinement by the approximator constitute an $\leq_p$ -improvement of the input pair itself.

Proposition 7. *If a consistent pair (I, U) is a post-fixpoint of $\mathbf{C}^{st}$, i.e. $(I, U) \leq_p \mathbf{C}^{st}(I, U)$, then $\mathbf{A}(U)$ is admissible in D . Conversely, if $A \subseteq \mathcal{A}$ is admissible in D , then $(\mathbf{R}^*(A), \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(A))))$ is a consistent post-fixpoint of $\mathbf{C}^{st}$.*

Proof. If $(I, U) \leq_p \mathbf{C}^{st}(I, U)$, then $I \subseteq \text{lfp}(\mathbf{C}(\cdot, U)_1) = \mathbf{R}^*(\mathbf{A}(U))$ and $\text{lfp}(\mathbf{C}(\cdot, I)_1) = \mathbf{R}^*(\mathbf{A}(I)) \subseteq U$ by Lemma 4, respectively. Thus $\mathbf{R}^*(\mathbf{A}(U)) \subseteq \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(\mathbf{A}(I))))$ by antimonotonicity, and in combination, $I \subseteq \mathbf{R}^*(\mathbf{A}(U)) \subseteq \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(\mathbf{A}(I))))$. Additionally, $I \subseteq U$ yields $I \subseteq \mathbf{R}^*(\mathbf{A}(U)) \subseteq \mathbf{R}^*(\mathbf{A}(I))$. Altogether,

$$I \subseteq \mathbf{R}^*(\mathbf{A}(U)) \subseteq \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(\mathbf{A}(I)))) \subseteq \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(\mathbf{A}(U)))) \subseteq \mathbf{R}^*(\mathbf{A}(I)) \subseteq U.$$

Thus from $\mathbf{R}^*(\mathbf{A}(U)) \subseteq U$ we get $\mathbf{A}(U) \subseteq \mathbf{A}(\mathbf{R}^*(\mathbf{A}(U)))$ whence $\mathbf{A}(U)$ is conflict-free by Proposition 1. Moreover, from $\mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(\mathbf{A}(U)))) \subseteq U$ we can infer $\mathbf{A}(U) \subseteq \mathbf{A}(\mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(\mathbf{A}(U)))))$ whence $\mathbf{A}(U)$ is admissible by Proposition 1.

Conversely, if $A \subseteq \mathcal{A}$ is admissible, then Proposition 1 yields that $A \subseteq \mathbf{A}(\mathbf{R}^*(A))$ and $A \subseteq \mathbf{A}(\mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(A))))$. Thus by monotonicity of $\mathbf{R}^*$ we get $\mathbf{R}^*(A) \subseteq \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(A)))$ (whence the pair is consistent) and $\mathbf{R}^*(A) \subseteq \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(A)))))$. By Lemma 4, $\mathbf{R}^*(A) \subseteq \text{lfp}(\mathbf{C}(\cdot, \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(A))))_1) = \mathbf{C}^{st}(\mathbf{R}^*(A), \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(A))))_1$. On the other hand, $\mathbf{C}^{st}(\mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(A))), \mathbf{R}^*(A))_1 = \text{lfp}(\mathbf{C}(\cdot, \mathbf{R}^*(A))_1) = \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(A)))$ by Lemma 4. Thus $(\mathbf{R}^*(A), \mathbf{R}^*(\mathbf{A}(\mathbf{R}^*(A))))$ is a consistent post-fixpoint of $\mathbf{C}^{st}$. $\square$

Among complete sets, the grounded set is the least one w.r.t. $\subseteq$. On the approximation side, this corresponds to the least precise fixpoint of the stable operator.

Proposition 8. *The consistent pair $(I, \mathbf{R}^*(\mathbf{A}(I)))$ is the least fixpoint of $\mathbf{C}^{st}$ iff $I \cap \mathcal{A}$ is the grounded set of D .*

In terms of AFT, preferred sets correspond to those fixpoints of $\mathbf{C}^{st}$ whose accepted assumptions cannot be strictly increased while maintaining the fixpoint property.

Proposition 9. *A consistent pair $(I, \mathbf{R}^*(\mathbf{A}(I)))$ is a $\leq_p$ -maximal fixpoint of $\mathbf{C}^{st}$ iff $I \cap \mathcal{A}$ is a preferred set.*

The proofs of the last two propositions are similar in technique to the previous ones, but had to be left out due to a lack of space.

4. The Characteristic ABA Operator and Its Canonical Approximator

In this section, we will consider a slightly different way to reconstruct ABA semantics within AFT. Mainly, we will use another underlying complete lattice: instead of considering all sets of atoms $2^{\mathcal{L}}$, we will only consider sets of assumptions $2^{\mathcal{A}}$.

Using the well-known connection between assumption-based argumentation and abstract argumentation frameworks, our plan is then to employ the known reconstruction of semantics for abstract argumentation frameworks within AFT [17] for ABA as well. We will observe that this introduces a number of operators that are useful for ABA in their own right, and have seemingly not been considered in the literature so far. Firstly, we introduce an operator that (in the abstract argumentation case) goes back to Pollock [32].

April 2022

Definition 8. For any ABA framework D , define $\mathbf{U}: 2^{\mathcal{A}} \rightarrow 2^{\mathcal{A}}$ with $\mathbf{U}(A) := \mathbf{A}(\mathbf{R}^*(A))$.

As $\mathbf{R}$ is $\subseteq$ -monotone, also $\mathbf{R}^*$ is $\subseteq$ -monotone; thus with $\mathbf{A}$ being $\subseteq$ -antimonotone, we get that $\mathbf{U}$ is $\subseteq$ -antimonotone.

Next, we introduce the ABA counterpart of Dung's [3] characteristic operator on abstract argumentation frameworks.

Definition 9. For any D , define $\mathbf{F}: 2^{\mathcal{A}} \rightarrow 2^{\mathcal{A}}$ with $\mathbf{F}(A) := \{a \in \mathcal{A} \mid A \text{ defends } a\}$.

As observed by Dung [3] for AFs, the two operators are in a close relationship.

Proposition 10. For any ABA framework D and $A \subseteq \mathcal{A}$, we have $\mathbf{U}(\mathbf{U}(A)) = \mathbf{F}(A)$.

The approximator we are going to use will approximate the antimonotone operator $\mathbf{U}$. Making use of Proposition 10 and a standard result of AFT [14, Proposition 8], the approximator will be the *canonical* approximator of $\mathbf{U}$.

Definition 10. For any ABA framework D , the operator $\mathbf{G}: (2^{\mathcal{A}} \times 2^{\mathcal{A}}) \rightarrow (2^{\mathcal{A}} \times 2^{\mathcal{A}})$ is symmetric with its first component given by $\mathbf{G}(I, U)_1 := \mathbf{U}(U)$.

That is, overall we obtain $\mathbf{G}(I, U) = (\mathbf{U}(U), \mathbf{U}(I))$ and clearly $\mathbf{G}$ approximates $\mathbf{U}$ as $\mathbf{G}(S, S) = (\mathbf{U}(S), \mathbf{U}(S))$. In a typical AFT setting, we would now consider the stable approximator of $\mathbf{G}$ that is given (through AFT) by $\mathbf{G}^{st}(S, T) = (\text{lfp}(\mathbf{G}(\cdot, T)_1), \text{lfp}(\mathbf{G}(\cdot, S)_1))$. However, as the following result shows, in this case this approximator is identical to $\mathbf{G}$.

Proposition 11. The stable approximator of $\mathbf{G}$ is the approximator $\mathbf{G}$, that is, $\mathbf{G}^{st} = \mathbf{G}$.

Proof. Let $S, T \subseteq \mathcal{A}$ be arbitrary. By the definition of $\mathbf{G}$, we have $\mathbf{G}(S, T)_1 = \mathbf{U}(T)$ which does not depend on S . Hence, $\text{lfp}(\mathbf{G}(\cdot, T)_1) = \mathbf{U}(T)$. Therefore we get that $\mathbf{G}^{st}(S, T) = (\mathbf{U}(T), \mathbf{U}(S)) = \mathbf{G}(S, T)$. $\square$

This shows that, as for AFs, the canonical approximator is already stable, so we do not need an additional stable-approximator “iteration” to obtain the intended semantics. In what follows, we will use $\mathbf{G}$ to reconstruct ABA semantics. The first result shows that complete sets are precisely the first components of consistent fixpoints of $\mathbf{G}$.

Proposition 12. Let $X \subseteq \mathcal{A}$. Then $X \in \text{com}(D)$ iff there exists a set $Y \supseteq X$ such that the pair (X, Y) is a consistent fixpoint of $\mathbf{G}$.

Proof. Let (X, Y) be a consistent fixpoint of $\mathbf{G}$. Then $(X, Y) = \mathbf{G}(X, Y) \implies X = \mathbf{U}(Y), Y = \mathbf{U}(X)$. Hence, $X = \mathbf{U}(\mathbf{U}(X))$ and $X \subseteq \mathbf{U}(X) = Y$ (due to consistency), which means it is complete. Conversely, let $X \in \text{com}(D)$. Then the pair $(X, Y = \mathbf{U}(X))$ is consistent since $X \subseteq \mathbf{U}(X)$ (by definition of complete). Also it is a fixpoint of $\mathbf{G}$ as $Y = \mathbf{U}(X)$, and $X = \mathbf{U}(\mathbf{U}(X)) = \mathbf{U}(Y)$. $\square$

Likewise, stable sets are characterized by exact fixpoints of $\mathbf{G}$.

Proposition 13. Let $X \subseteq \mathcal{A}$. Then $X \in \text{stb}(D)$ iff the pair (X, X) is a fixpoint of $\mathbf{G}$.

Proof. By definition of stable sets, we get that $X = \mathbf{U}(X)$. Consider the pair (X, X) , and its image through $\mathbf{G}$. For a pair (X, Y) to be a fixpoint, it must hold that $X = \mathbf{U}(Y)$, and $Y = \mathbf{U}(X)$. Here both these hold as $X = Y$. Conversely, the pair (X, X) is a fixpoint of $\mathbf{G}$ which implies that $X = \mathbf{U}(X)$. Hence X is stable. $\square$

Admissible and preferred sets arise when we relax or strengthen the exact fixpoint requirement and instead look at post-fixpoints and maximal fixpoints, respectively.

Proposition 14. *Let $X \subseteq \mathcal{A}$. Then X is an admissible set iff (X, Y) is post-fixpoint of $\mathbf{G}$ for some $X \subseteq Y$.*

Proposition 15. *Let $X \subseteq \mathcal{A}$. Then X is a preferred set iff (X, Y) is a $\leq_p$ -maximal fixpoint of $\mathbf{G}$ for some $X \subseteq Y$.*

The grounded set is characterized as the least precise fixpoint of $\mathbf{G}$.

Proposition 16. *Let $X \subseteq \mathcal{A}$. Then X is the grounded set iff (X, Y) is the least fixpoint of $\mathbf{G}$ for some $Y \supseteq X$.*

5. Conclusions

In this work we reconstructed the major semantics of assumption-based argumentation within the general framework of approximation fixpoint theory. In particular, we showed two ways of defining operators and their approximators that both capture the standard semantics of ABA. The first operator/approximator pair $(\mathbf{B}, \mathbf{C})$ is inspired by ABA's close connections to normal logic programs and is similar in spirit to the van Emden/Kowalski one-step consequence operator for logic programming [23], and its respective approximator [14]. The second operator/approximator pair $(\mathbf{U}, \mathbf{G})$ is inspired by ABA's connections to Dung's argumentation frameworks and is similar in spirit to operators proposed by Pollock [32], Dung [3], and our previous reconstruction of AF semantics in AFT [17].

By placing ABA in AFT, we open up further research towards general frameworks capturing main aspects of various structured argumentation approaches, e.g., complementing other general investigations such as properties of attack relations in structured argumentation [35]. As another direct consequence, taking into account the ability of AFT to also reconstruct *fuzzy* logic programming [19], our reconstruction of ABA in AFT paves the way for weighted generalizations of ABA, e.g. in the spirit of weighted abstract dialectical frameworks [36,37].

References

- [1] Baroni P, Gabbay D, Giacomin M, van der Torre L, editors. Handbook of Formal Argumentation. College Publications; 2018.
- [2] Atkinson K, Baroni P, Giacomin M, Hunter A, Prakken H, Reed C, et al. Towards Artificial Argumentation. AI Mag. 2017;38(3):25-36.
- [3] Dung PM. On the Acceptability of Arguments and its Fundamental Role in Nonmonotonic Reasoning, Logic Programming and n-Person Games. Artif Intell. 1995;77(2):321-58.
- [4] Besnard P, García AJ, Hunter A, Modgil S, Prakken H, Simari GR, et al. Introduction to structured argumentation. Argument Comput. 2014;5(1):1-4.
- [5] Caminada M, Amgoud L. On the evaluation of argumentation formalisms. Artif Intell. 2007;171(5-6):286-310.
- [6] Bondarenko A, Dung PM, Kowalski RA, Toni F. An Abstract, Argumentation-Theoretic Approach to Default Reasoning. Artif Intell. 1997;93:63-101.
- [7] Modgil S, Prakken H. A general account of argumentation with preferences. Artif Intell. 2013;195:361-97.

- [8] García AJ, Simari GR. Defeasible Logic Programming: An Argumentative Approach. *Theory Pract Log Program*. 2004;4(1-2):95-138.
- [9] Besnard P, Hunter A. *Elements of Argumentation*. MIT Press; 2008.
- [10] Gordon TF, Prakken H, Walton D. The Carneades model of argument and burden of proof. *Artif Intell*. 2007;171(10-15):875-96.
- [11] Kakas AC, Moraitis P, Spanoudakis NI. *GORGias*: Applying argumentation. *Argument Comput*. 2019;10(1):55-81.
- [12] Caminada M, Schulz C. On the Equivalence between Assumption-Based Argumentation and Logic Programming. *J Artif Intell Res*. 2017;60:779-825.
- [13] Pearce D. A New Logical Characterisation of Stable Models and Answer Sets. In: *Proc. NMELP*. Lecture Notes in Computer Science. Springer; 1996. p. 57-70.
- [14] Denecker M, Marek V, Truszczyński M. Approximations, Stable Operators, Well-Founded Fixpoints and Applications in Nonmonotonic Reasoning. In: *Logic-Based Artificial Intelligence*; 2000. p. 127-44.
- [15] Denecker M, Marek VW, Truszczyński M. Ultimate approximation and its application in nonmonotonic knowledge representation systems. *Information and Computation*. 2004;192(1):84-121.
- [16] Brewka G, Woltran S. Abstract Dialectical Frameworks. In: *Proc. KR. AAAI Press*; 2010. p. 102-11.
- [17] Strass H. Approximating operators and semantics for abstract dialectical frameworks. *Artif Intell*. 2013;205:39-70.
- [18] Brewka G, Ellmauthaler S, Strass H, Wallner JP, Woltran S. Abstract Dialectical Frameworks. An Overview. *FLAP*. 2017;4(8).
- [19] Kettmann P, Heynink J, Strass H. Approximation Fixpoint Theory as a Unifying Framework for Fuzzy Logic Programming Semantics. In: *Proc. IJCAI*. ijcai.org; 2025. p. 4544-52.
- [20] Strass H, Wallner JP. Analyzing the computational complexity of abstract dialectical frameworks via approximation fixpoint theory. *Artif Intell*. 2015;226:34-74.
- [21] Killen S, You J, Heynink J. An Alternative Theory of Stable Revision for Nondeterministic Approximation Fixpoint Theory and the Relationships. In: *Proc. AAAI*. AAAI Press; 2025. p. 15033-40.
- [22] Heynink J, Arieli O, Bogaerts B. Non-deterministic approximation fixpoint theory and its application in disjunctive logic programming. *Artif Intell*. 2024;331:104110.
- [23] van Emden MH, Kowalski RA. The Semantics of Predicate Logic as a Programming Language. *J ACM*. 1976;23(4):733-42.
- [24] Heynink J, Arieli O. Argumentative Reflections of Approximation Fixpoint Theory. In: *Proc. COMMA*. Frontiers in Artificial Intelligence and Applications. IOS Press; 2020. p. 215-26.
- [25] Tarski A. A Lattice-Theoretical Fixpoint Theorem and Its Applications. *Pacific Journal of Mathematics*. 1955;5(2):285-309.
- [26] Markowsky G. Chain-complete posets and directed sets with applications. *Algebra Universalis*. 1976;6:53-68.
- [27] Cousot P, Cousot R. Constructive versions of Tarski's fixed point theorems. *Pacific Journal of Mathematics*. 1979;82(1):43-57.
- [28] Bourbaki N. Sur le théorème de Zorn. *Archiv der Mathematik*. 1949;2(6):434-7.
- [29] Belnap ND. A useful four-valued logic. In: *Modern uses of multiple-valued logic*; 1977. p. 5-37.
- [30] Ginsberg ML. Multivalued logics: a uniform approach to reasoning in artificial intelligence. *Comput Intell*. 1988;4:265-316.
- [31] Fitting M. Fixpoint Semantics for Logic Programming: A Survey. *Theor Comput Sci*. 2002;278(1-2):25-51.
- [32] Pollock JL. Defeasible Reasoning. *Cogn Sci*. 1987;11(4):481-518.
- [33] Clark KL. Negation as Failure. In: *Logic and Data Bases, Symposium on Logic and Data Bases, Centre d'études et de recherches de Toulouse, France, 1977*. Advances in Data Base Theory. New York: Plenum Press; 1977. p. 293-322.
- [34] Caminada MWA, Gabbay DM. A Logical Account of Formal Argumentation. *Stud Logica*. 2009;93(2-3):109-45.
- [35] Dung PM, Thang PM. Fundamental properties of attack relations in structured argumentation with priorities. *Artif Intell*. 2018;255:1-42.
- [36] Brewka G, Strass H, Wallner JP, Woltran S. Weighted Abstract Dialectical Frameworks. In: *Proc. AAAI*. AAAI Press; 2018. p. 1779-86.
- [37] Bogaerts B. Weighted Abstract Dialectical Frameworks through the Lens of Approximation Fixpoint Theory. In: *Proc. AAAI*. AAAI Press; 2019. p. 2686-93.