



Fuzzy Labeling for Abstract Argumentation: An Empirical Evaluation

Célia da Costa Pereira, Mauro Dragoni, Andrea G. B. Tettamanzi, Serena Villata

► To cite this version:

Célia da Costa Pereira, Mauro Dragoni, Andrea G. B. Tettamanzi, Serena Villata. Fuzzy Labeling for Abstract Argumentation: An Empirical Evaluation. Tenth International Conference on Scalable Uncertainty Management (SUM 2016) , Sep 2016, Nice, France. pp.126 - 139, 10.1007/978-3-319-45856-4_9 . hal-01377557

HAL Id: hal-01377557

<https://hal.science/hal-01377557>

Submitted on 7 Oct 2016

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

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

Fuzzy Labeling for Abstract Argumentation: An Empirical Evaluation

Célia da Costa Pereira¹, Mauro Dragoni², Andrea G. B. Tettamanzi¹, and
Serena Villata¹

¹ Université Côte d’Azur, CNRS, Inria, I3S, France
{celia.pereira, andrea.tettamanzi, serena.villata}@unice.fr

² Fondazione Bruno Kessler, Trento, Italy
dragoni@fbk.eu

Abstract. Argumentation frameworks have to be evaluated with respect to argumentation semantics to compute the set(s) of accepted arguments. In a previous approach, we proposed a fuzzy labeling algorithm for computing the (fuzzy) set of acceptable arguments, when the sources of the arguments in the argumentation framework are only partially trusted. The convergence of the algorithm was proved, and the convergence speed was estimated to be linear, as it is generally the case with iterative methods. In this paper, we provide an experimental validation of this algorithm with the aim of carrying out an empirical evaluation of its performance on a benchmark of argumentation graphs. Results show the satisfactory performance of our algorithm, even on complex graph structures as those present in our benchmark.

1 Introduction

In crisp argumentation, arguments are evaluated, following a specific semantics, as being acceptable or not acceptable, as shown by Dung [9]. Roughly, accepted arguments are those arguments which are not attacked by other accepted arguments, and rejected arguments are those attacked by accepted arguments. The set of accepted arguments, called *extensions*, represent consistent set(s) of arguments that can be accepted together. However, in many applications, such as decision making and agent-based recommendation systems, such a crisp evaluation of the arguments is not suitable to represent the complexity of the considered scenario. To address this issue, we [8] proposed to perform a *fuzzy* evaluation of the arguments of an argumentation framework. Such a fuzzy evaluation of the arguments is originated by the observation that some arguments may come from only partially trusted sources. To represent the degrees of trust, we rethink the usual crisp argument evaluation [9, 6] by evaluating arguments in terms of fuzzy degrees of acceptability. In our previous contribution [8], we proved that the fuzzy labeling algorithm used to assign to the arguments fuzzy degrees of acceptability converges, and we discussed its convergence speed. However, no empirical evaluation was addressed to support this discussion.

In this paper, we face this issue by providing an extensive evaluation of the performance and scalability of the fuzzy labeling algorithm with respect to a benchmark of abstract argumentation frameworks. We first select three existing datasets for abstract argumentation tasks used in the literature, namely the datasets created by Bistarelli *et al.* [3], by Cerutti *et al.* [7], and by Vallati *et al.* [12]. Moreover, we generate our own dataset of abstract argumentation frameworks by randomly combining some well known *graph patterns* in argumentation theory into 20,000 bigger argumentation frameworks. Second, we study the behaviour of the algorithm with respect to the frameworks in the benchmark, to check whether its performance are satisfiable even considering huge and complex networks as those represented in the datasets, e.g., presenting an increasing number of strongly connected components.

The reminder of the paper is as follows. In Section 2, we provide the basics of abstract argumentation theory and fuzzy set theory. Section 3 firstly recalls the main concepts behind the definition of the fuzzy labeling algorithm presented in [8], and secondly, describes its current implementation. In Section 4, we describe the four datasets which compose the benchmark used to evaluate our algorithm, and we report about the obtained results. Conclusions end the paper.

2 Preliminaries

In this section, we provide some insights about abstract argumentation theory and fuzzy sets.

2.1 Abstract argumentation theory

We provide the basics of Dung's abstract argumentation [9].

Definition 1. (*Abstract argumentation framework*) An abstract argumentation framework is a pair $\langle \mathcal{A}, \rightarrow \rangle$ where $\mathcal{A}$ is a set of elements called arguments and $\rightarrow \subseteq \mathcal{A} \times \mathcal{A}$ is a binary relation called attack. We say that an argument A_i attacks an argument A_j if and only if $(A_i, A_j) \in \rightarrow$.

Dung [9] presents several acceptability semantics which produce zero, one, or several sets of accepted arguments. These semantics are grounded on two main concepts, called conflict-freeness and defence.

Definition 2. (*Conflict-free, Defence*) Let $C \subseteq \mathcal{A}$. A set C is conflict-free if and only if there exist no $A_i, A_j \in C$ such that $A_i \rightarrow A_j$. A set C defends an argument A_i if and only if for each argument $A_j \in \mathcal{A}$ if A_j attacks A_i then there exists $A_k \in C$ such that A_k attacks A_j .

Definition 3. (*Acceptability semantics*) Let C be a conflict-free set of arguments, and let $\mathcal{D} : 2^{\mathcal{A}} \mapsto 2^{\mathcal{A}}$ be a function such that $\mathcal{D}(C) = \{A | C \text{ defends } A\}$.

- C is admissible if and only if $C \subseteq \mathcal{D}(C)$.
- C is a complete extension if and only if $C = \mathcal{D}(C)$.

- C is a grounded extension if and only if it is the smallest (w.r.t. set inclusion) complete extension.
- C is a preferred extension if and only if it is a maximal (w.r.t. set inclusion) complete extension.
- C is a stable extension if and only if it is a preferred extension that attacks all arguments in $\mathcal{A} \setminus C$.

The concepts of Dung’s semantics are originally stated in terms of sets of arguments. It is equal to express these concepts using argument *labeling* [10, 13, 6] In a reinstatement labeling [6], an argument is labeled “in” if all its attackers are labeled “out” and it is labeled “out” if it has at least an attacker which is labeled “in”.

Definition 4. (*AF-labeling* [6]) Let $\langle \mathcal{A}, \rightarrow \rangle$ be an abstract argumentation framework. An AF-labeling is a total function $lab : \mathcal{A} \rightarrow \{in, out, undec\}$. We define $in(lab) = \{A_i \in \mathcal{A} | lab(A_i) = in\}$, $out(lab) = \{A_i \in \mathcal{A} | lab(A_i) = out\}$, $undec(lab) = \{A_i \in \mathcal{A} | lab(A_i) = undec\}$.

Definition 5. (*Reinstatement labeling* [6]) Let lab be an AF-labeling.

- a in-labeled argument is said to be legally in iff all its attackers are labeled out.
- a out-labeled argument is said to be legally out iff it has at least one attacker that is labeled in.
- an undec-labelled argument is said to be legally undec iff not all its attackers are labelled out and it does not have an attacker that is labelled in.

Definition 6. An admissible labelling is a labelling lab where each in-labelled argument is legally in and each out-labelled argument is legally out. lab is a complete labeling if there are no arguments illegally in and illegally out and illegally undec. We say that lab is a

- grounded, iff $in(lab)$ is minimal (w.r.t. set inclusion);
- preferred, iff $in(lab)$ is maximal (w.r.t. set inclusion);
- stable, iff $undec(lab) = \emptyset$

2.2 Fuzzy Sets

Fuzzy sets [14] are a generalization of classical (crisp) sets obtained by replacing the characteristic function of a set A , χ_A , which takes up values in $\{0, 1\}$ ($\chi_A(x) = 1$ iff $x \in A$, $\chi_A(x) = 0$ otherwise) with a *membership function* μ_A , which can take up any value in $[0, 1]$. The value $\mu_A(x)$ or, more simply, $A(x)$ is the membership degree of element x in A , i.e., the degree to which x belongs in A .

A fuzzy set is completely defined by its membership function. Therefore, it is useful to define a few terms describing various features of this function. Given a fuzzy set A , its *core* is the (conventional) set of all elements x such that $A(x) = 1$;

its *support*, $\text{supp}(A)$, is the set of all x such that $A(x) > 0$. A fuzzy set is *normal* if its core is nonempty. The set of all elements x of A such that $A(x) \geq \alpha$, for a given $\alpha \in (0, 1]$, is called the α -cut of A , denoted A_α .

The usual set-theoretic operations of union, intersection, and complement can be defined as a generalization of their counterparts on classical sets by introducing two families of operators, called triangular norms and triangular co-norms. In practice, it is usual to employ the min norm for intersection and the max co-norm for union. Given two fuzzy sets A and B , and an element x ,

$$(A \cup B)(x) = \max\{A(x), B(x)\}; \quad (1)$$

$$(A \cap B)(x) = \min\{A(x), B(x)\}; \quad (2)$$

$$\bar{A}(x) = 1 - A(x). \quad (3)$$

Finally, given two fuzzy sets A and B , $A \subseteq B$ if and only if, for every element x , $A(x) \leq B(x)$.

3 Fuzzy Labeling for Abstract Argumentation

In this section, we recall the fuzzy labeling algorithm for abstract argumentation [8] we want to empirically evaluate over the available datasets for abstract argumentation to study its performance. For a complete description of the algorithm and its convergence theorem as well as the comparison with the related approaches we remind the reader to [8]. Moreover, we report about the implementation we develop to test the performances of the algorithm.

3.1 Algorithm

In order to account for the fact that arguments may originate from sources that are trusted only to a certain degree, the (crisp) abstract argumentation structure described in Section 2 may be extended by allowing gradual membership of arguments in the set of arguments $\mathcal{A}$. We have that $\mathcal{A}$ is a fuzzy set of trustworthy arguments, and $\mathcal{A}(A)$, the membership degree of argument A in $\mathcal{A}$, is given by the trust degree of the most reliable (i.e., trusted) source that offers argument A ,³

$$\mathcal{A}(A) = \max_{s \in \text{src}(A)} \tau_s, \quad (4)$$

where $\text{src}(A)$ is the set of sources proposing argument A and τ_s is the degree to which source $s \in \text{src}(A)$ is trusted. We do not make any further assumptions on the trust model, as it is out of the scope of this paper. However, we refer the interested reader to [11], where a more detailed description of how the source trustworthiness degree can be computed starting from elements, like the source sincerity and expertise, is provided.

³ Here, we suppose that the agent is optimistic. To represent a pessimistic behaviour, we should use the min operator, for example.

Definition 7. (*Fuzzy AF-labeling*) Let $\langle \mathcal{A}, \rightarrow \rangle$ be an abstract argumentation framework. A fuzzy AF-labeling is a total function $\alpha : \mathcal{A} \rightarrow [0, 1]$.

Such an α may also be regarded as (the membership function of) the fuzzy set of acceptable arguments: $\alpha(A) = 0$ means the argument is outright unacceptable, $\alpha(A) = 1$ means the argument is fully acceptable, and all cases inbetween are provided for.

Definition 8. (*Fuzzy Reinstatement Labeling*) Let α be a fuzzy AF-labeling. We say that α is a fuzzy reinstatement labeling iff, for all arguments A ,

$$\alpha(A) = \min\{\mathcal{A}(A), 1 - \max_{B:B \rightarrow A} \alpha(B)\}. \quad (5)$$

The above definition combines two intuitive postulates of fuzzy labeling: (1) the acceptability of an argument should not be greater than the degree to which the arguments attacking it are unacceptable and (2) an argument cannot be more acceptable than the degree to which its sources are trusted: $\alpha(A) \leq \mathcal{A}(A)$.

We can verify that the fuzzy reinstatement labeling is a generalization of the crisp reinstatement labeling of Definition 5, whose *in* and *out* labels are particular cases corresponding, respectively, to $\alpha(A) = 1$ and $\alpha(A) = 0$. The intermediate cases, $0 < \alpha(A) < 1$ correspond to a continuum of degrees of “undecidedness”, of which 0.5 is but the most undecided representative.

We denote by $\alpha_0 = \mathcal{A}$ the initial labeling, and by α_t the labeling obtained after the t^{th} iteration of the labeling algorithm.

Definition 9. Let α_t be a fuzzy labeling. An iteration in α_t is carried out by computing a new labeling α_{t+1} for all arguments A as follows:

$$\alpha_{t+1}(A) = \frac{1}{2}\alpha_t(A) + \frac{1}{2}\min\{\mathcal{A}(A), 1 - \max_{B:B \rightarrow A} \alpha_t(B)\}. \quad (6)$$

This defines a sequence $\{\alpha_t\}_{t=0,1,\dots}$ of labelings which always converges to a limit fuzzy labeling, as proven in [8]. Moreover, the convergence speed is linear: in practice, a small number of iterations is enough to compute the limit up to the desired precision. The fuzzy labeling of a fuzzy argumentation framework is thus the limit of $\{\alpha_t\}_{t=0,1,\dots}$.

Definition 10. Let $\langle \mathcal{A}, \rightarrow \rangle$ be a fuzzy argumentation framework. A fuzzy reinstatement labeling for such argumentation framework is, for all arguments A ,

$$\alpha(A) = \lim_{t \rightarrow \infty} \alpha_t(A). \quad (7)$$

3.2 Implementation

The fuzzy labeling algorithm has been implemented by using the Java language without using any specific library. The first version of the algorithm was developed with the support of multi-threading where the update of each node was parallelized. However, preliminary tests run in multi-threading mode reported

some computational overhead that make the observation of algorithm performance not consistent. For this reason and for measuring performance values easing the comparison with the obtained results, we opted for running all tests in a single-thread mode.

Each run was performed on a server equipped with a Xeon E5-2609 v2 @ 2.50 Ghz. According to official user-based benchmarks,⁴ the single thread mark of the CPU is 1229 points. This value can be used as reference for normalizing results concerning the timing of each run obtained on other machines.

4 Evaluation

In this section, we study the behaviour and the performances of the fuzzy-labeling algorithm over a benchmark for abstract argumentation, and then we report about the obtained results.

The aim of our experimental analysis is to assess the scalability of the fuzzy-labeling algorithm concerning two perspectives:

- the number of iterations needed for convergence with respect to the number of the nodes in the graph, and
- the time needed for convergence with respect to the number of the nodes in the graph.

It must be stressed that the time needed for convergence depends on (i) the time needed for computing each iteration, (ii) the time needed to update the α of each single argument, and (iii) the number of iterations required for the labeling to converge.

4.1 Benchmark

The benchmark we used to evaluate the performances of the fuzzy labeling algorithm is composed of different datasets for abstract argumentation tasks used in the literature. More precisely, we have considered the following datasets:

- The Perugia dataset [2, 4, 3]:⁵ the dataset is composed of randomly generated directed-graphs. To generate random graphs, they adopted two different libraries. The first one is the Java Universal Network/Graph Framework (JUNG), a Java software library for the modeling, generation, analysis and visualization of graphs. The second library they used is NetworkX, a Python software package for the creation, manipulation, and study of the structure, dynamics, and functions of complex networks. Three kinds of networks are generated:

⁴ https://www.cpubenchmark.net/CPU_mega_page.html

⁵ The dataset is available at <http://www.dmi.unipg.it/conarg/dwl/networks.tgz>.

- In the Erdős-Rényi graph model, the graph is constructed by randomly connecting n nodes. Each edge is included in the graph with probability p independent from every other edge. For the generation of these argumentation graphs, they adopted $p = c \log n/n$ (with c empirically set to 2.5), which ensures the connectedness of such graphs.
- The Kleinberg graph model adds a number of directed long-range random links to an $n \times n$ lattice network, where vertices are the nodes of a grid with undirected edges between any two adjacent nodes. Links have a non-uniform distribution that favors edges to close nodes over more distant ones.
- In the Barabási-Albert graph model, at each time step, a new vertex is created and connected to existing vertices according to the principle of “preferential attachment”, such that vertices with higher degree have a higher probability of being selected for attachment.

For more details about the generation of these networks as well as the graph models, we refer the reader to [2, 4, 3].

- The dataset used by Cerutti *et al.* in their KR 2014 paper [7] (which we will call the KR dataset): the dataset has been generated to evaluate a meta-algorithm for the computation of preferred labelings, based on the general recursive schema for argumentation semantics called SCC-Recursiveness. The dataset is composed of three sets of argumentation frameworks, namely:
 - 790 randomly generated argumentation frameworks where the number of strongly connected components (SCC) is 1, varying the number of arguments between 25 and 250 with a step of 25.
 - 720 randomly generated argumentation frameworks where the number of strongly connected components varies between 5 and 45 with a step of 5. The size of the SCCs is determined by normal distributions with means between 20 and 40 with a step of 5, and with a fixed standard deviation of 5. They similarly varied the probability of having attacks between arguments among SCCs.
 - 2800 randomly generated argumentation frameworks where the number of strongly connected components is between 50 and 80 with a step of 5.
- The dataset presented by Vallati *et al.* at ECAI 2014 [12] (which we will call the ECAI dataset): the dataset was produced to study the features of argumentation frameworks. More precisely, it is composed of 10,000 argumentation frameworks generated using a parametric random approach allowing to select (probabilistically average, standard deviation) the density of attacks for each strongly connected component, and how many arguments (probabilistically) in each SCC attack how many arguments (probabilistically) in how many (probabilistically) other SCCs. The number of arguments ranges between 10 and 40,000, and they exploited a 10-fold cross-validation approach on a uniform random permutation of the instances.

The availability of real-world benchmarks for argumentation problems is quite limited, with some few exceptions like [5] or AIFdb.⁶ However, these bench-

⁶ <http://corpora.aifdb.org>


```

a → b, b → a
a → b, b → a, b → c
a → b, b → a, b → c, c → d, d → c
a → b, b → c, c → d, d → e, e → f, f → e
a → b, b → c, c → d, d → c
a → b, b → c, c → a
a → b, b → a, b → c, c → c
a → b, b → a, b → c, c → c, d → d
a → b, b → a, b → c, a → c, c → d, d → c
a → b, b → a, b → c, a → c, c → d
a → b, b → c, c → a, b → d, a → d, c → d, d → e, e → d
a → b, b → c, e → c, c → d
a → b, b → a, b → c, c → d, d → e, e → c
a → b, b → c, c → c
a → b, b → c, c → a, a → d, b → d, c → d
a → b, a → c, c → a, c → d, d → c, d → a, a → d, c → e, d → f

```

Fig. 1. The “patterns” used for constructing the Sophia Antipolis dataset.

marks are tailored towards problems of argument mining and their representation as abstract argumentation frameworks usually leads to topologically simple graphs, such as cycle-free graphs. These kinds of graphs are not suitable for evaluating the computational performance of solvers for abstract argumentation problems. For this reason, we decided to use artificially generated graphs as benchmarks, in line with the preliminary performance evaluation of Bistarelli *et al.* [2].

In order to ensure the consideration of all kinds of interesting “patterns” that could appear in argumentation frameworks (e.g., the abstract argumentation frameworks used to exemplify the behaviour of the semantics in [1]), we have generated further graphs by composing these basic well-known examples of *interesting* argumentation patterns (shown in Figure 1) into bigger frameworks.

Our generated dataset (which we will call the Sophia Antipolis dataset)⁷ consists of 20,000 argumentation graphs created through a random aggregation of the patterns shown above. This process has been executed with different settings in order to obtain complex graphs of specific sizes. In particular, a set of 1,000 argumentation graphs is generated for graph sizes from 5,000 to 100,000 nodes, with incremental steps of 5,000 nodes each. The aggregation of patterns has been done incrementally, and the connections (edges) between single patterns were generated randomly. The number of created graphs and their different sizes should support the evaluation of argumentation reasoning algorithms under a broad number of scenarios.

4.2 Results

Figures 2–11 summarize the behavior of the fuzzy-labeling algorithm on the four datasets we considered. For each dataset, we applied the algorithm to the argumentation graphs with all argument weights set to 1 (i.e., arguments coming from fully trusted sources) and with random weights (i.e., arguments coming

⁷ The Sophia Antipolis dataset is available at <https://goo.gl/pN1M9r>.

from a variety of more or less trusted sources as it may be the case in application scenarios like multiagent systems). From a first inspection of the figures, it is clear that certain graph types are harder than others: the Sophia Antipolis appears to be the hardest, followed by the Erdős-Rényi, the Barabási-Albert, and the KR + ECAI datasets. The Kleinberg dataset appears to be the easiest. Furthermore, for all datasets, the graphs with random weights never require a smaller number of iterations for convergence than their counterparts with all weights fixed to 1.

Figures 2 and 3 show the behaviour of the fuzzy labeling algorithm when applied to the Barabási-Albert dataset. In particular, Figure 2 (left-hand side) illustrates the evolution of the number of iterations needed to reach convergence when all the weights are equal to 1. We can notice that the curve follows a logarithmic rise with the increasing of the number of nodes. The figure illustrated through the (right-hand side) curve represents the evolution of the time needed to reach the convergence. It shows a behaviour rather linear. However, we can notice that the slope of the curve decreases with the increasing of the number of nodes. A similar behaviour is depicted in Figure 3 which illustrates the evolution of the quantity of time (in ms) needed to reach the convergence when the weights are assigned randomly. These two illustrations clearly show the capability of the fuzzy labeling algorithm to handle a growing amount of data.

In Figures 4 and 5, we present the behaviour of the fuzzy labeling algorithm when applied to the Erdős-Rényi dataset. We can notice that when all the weights are equal to 1, the convergence is reached very quickly both when considering the number of iterations, and the quantity of time needed for convergence. However, while such a quantity is quite similar with respect to the case in which the weights are randomly assigned, we can notice that the number of iterations needed for convergence is higher with respect to the behaviour illustrated in Figure 4. This can be due to the fact that the Erdős-Rényi dataset is constructed by randomly connecting the nodes. As we can see in Figures 6 and 7, the convergence with the Kleinberg dataset is even globally faster, either when all weights are equal to 1 or when the weights are randomly assigned. Instead, the behaviour on the Sophia Antipolis dataset, shown in Figures 10 and 11, is quite similar to the behaviors obtained with the Barabási-Albert dataset.

It is less evident, but the fuzzy-labeling algorithm behaves on the KR + ECAI dataset (illustrated in Figures 8 and 9) much like it does on the Barabási-Albert and Sophia Antipolis datasets, with the exception of a few small graphs which are outliers and which demand a relatively large number of iterations to converge. Nevertheless, the time behavior of Barabási-Albert, Sophia Antipolis and KR + ECAI is qualitatively identical.

Despite the differences among the various graph types, we have a rate of increase in time which is at most log-linear for all graph types and for all weight assignments.

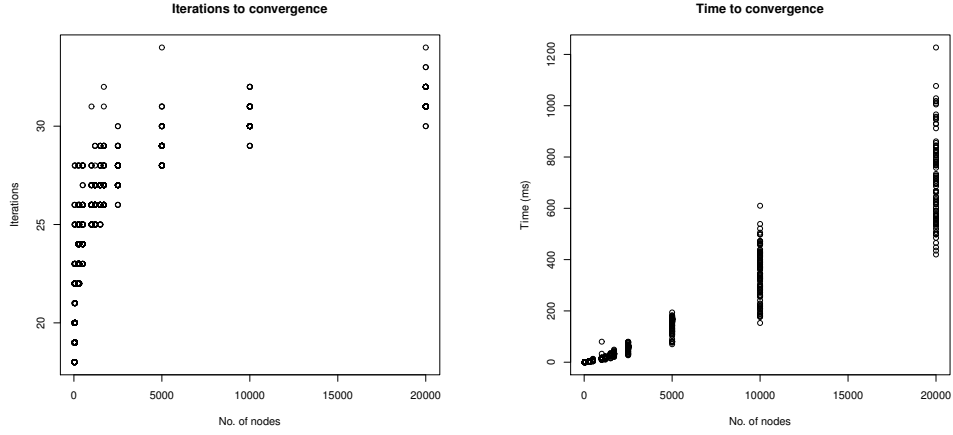


Fig. 2. Barabási-Albert dataset of the Perugia benchmark with all weights equal to 1.0.

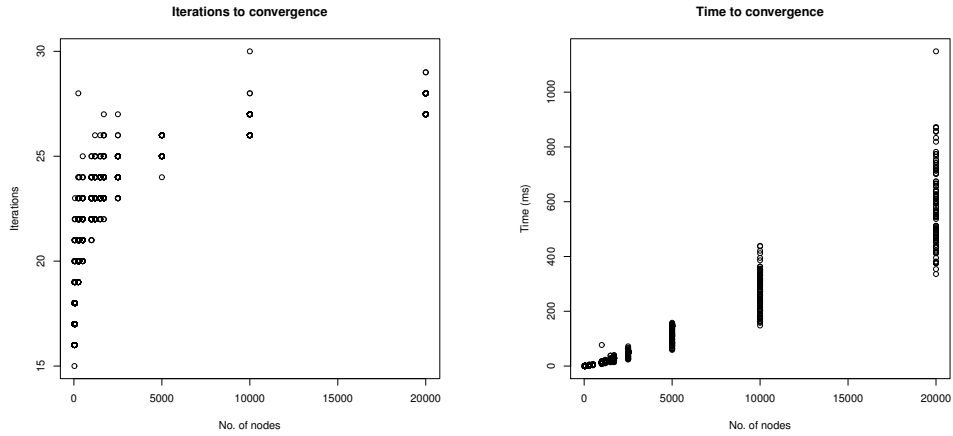


Fig. 3. Barabási-Albert dataset of the Perugia benchmark with random weights.

5 Conclusions

We have evaluated the performance of the fuzzy labeling algorithm proposed in [8] on a benchmark consisting of four datasets of argumentation graphs having widely different characteristics. The experimental results clearly indicate that the fuzzy labeling algorithm scales up nicely in all circumstances, and is thus a viable argumentation reasoning tool.

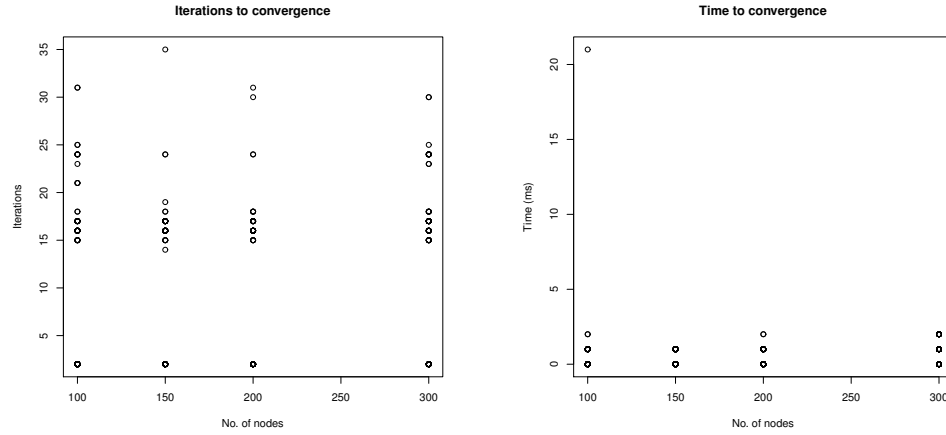


Fig. 4. The Erdős-Rényi dataset of the Perugia benchmark with all weights equal to 1.0.

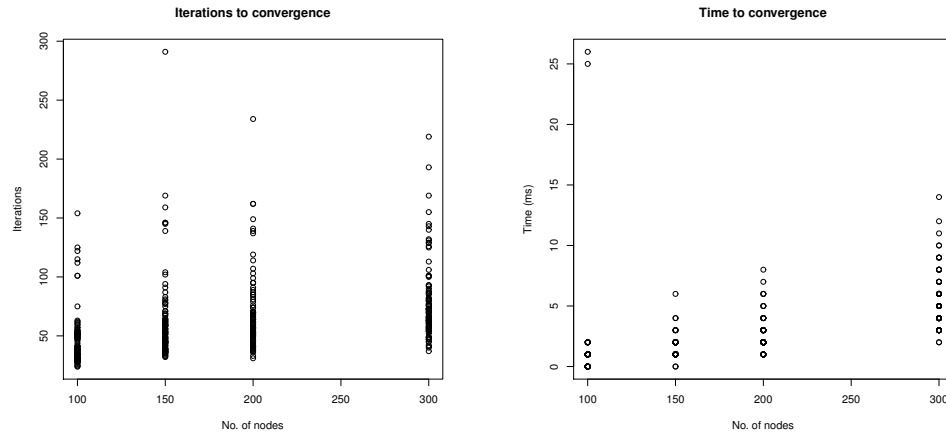


Fig. 5. The Erdős-Rényi dataset of the Perugia benchmark with random weights.

References

1. Baroni, P., Caminada, M., Giacomin, M.: An introduction to argumentation semantics. *Knowledge Eng. Review* 26(4), 365–410 (2011)
2. Bistarelli, S., Rossi, F., Santini, F.: A first comparison of abstract argumentation systems: A computational perspective. In: *Proceedings of the 28th Italian Conference on Computational Logic*. pp. 241–245 (2013)
3. Bistarelli, S., Rossi, F., Santini, F.: Benchmarking hard problems in random abstract afs: The stable semantics. In: *Computational Models of Argument - Proceedings of COMMA 2014*. pp. 153–160 (2014)

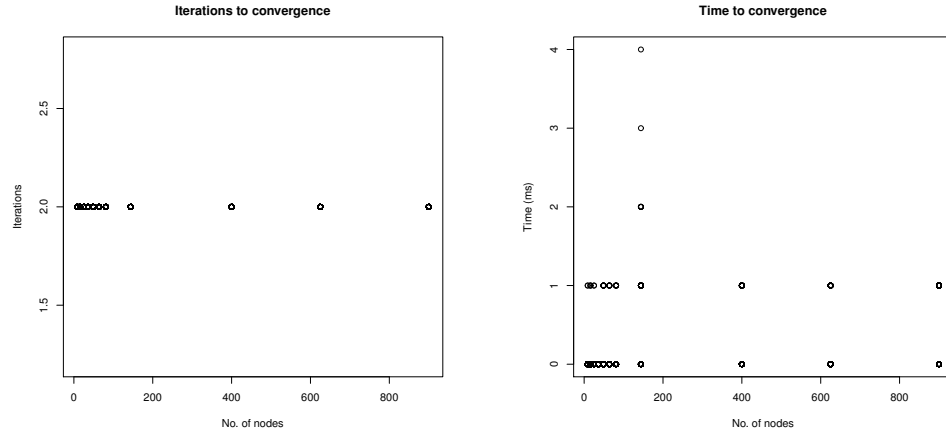


Fig. 6. The Kleinberg dataset of the Perugia benchmark with all weights equal to 1.0.

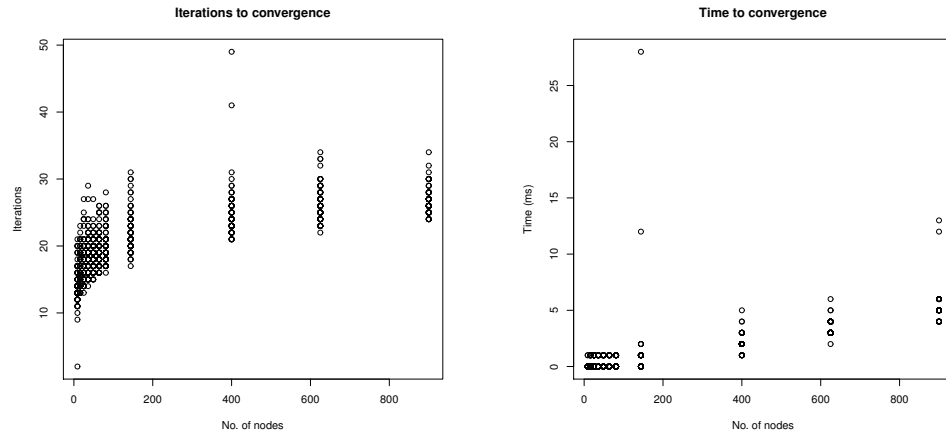


Fig. 7. The Kleinberg dataset of the Perugia benchmark with random weights.

4. Bistarelli, S., Rossi, F., Santini, F.: A first comparison of abstract argumentation reasoning-tools. In: ECAI 2014 - 21st European Conference on Artificial Intelligence. pp. 969–970 (2014)
5. Cabrio, E., Villata, S.: NoDE: A benchmark of natural language arguments. In: Computational Models of Argument - Proceedings of COMMA 2014. pp. 449–450 (2014)
6. Caminada, M.: On the issue of reinstatement in argumentation. In: Proceedings of JELIA. LNCS, vol. 4160, pp. 111–123. Springer (2006)
7. Cerutti, F., Giacomini, M., Vallati, M., Zanella, M.: An SCC recursive meta-algorithm for computing preferred labellings in abstract argumentation. In: Principles of Knowledge Representation and Reasoning: Proceedings of the Fourteenth

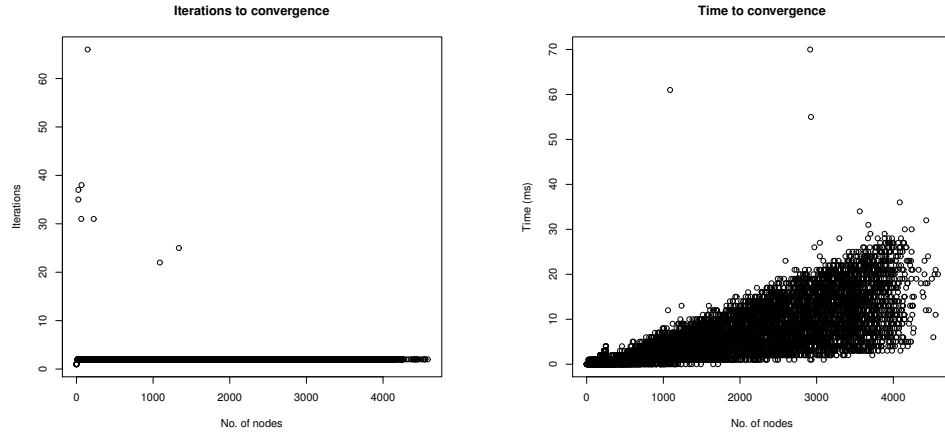


Fig. 8. The benchmark consisting of the KR + ECAI dataset with all weights equal to 1.0.

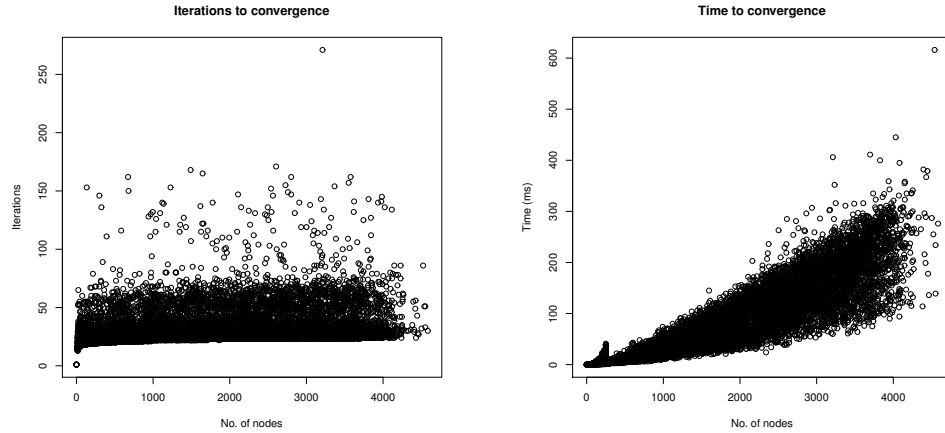


Fig. 9. The benchmark consisting of the KR + ECAI dataset with random weights.

International Conference, KR 2014. (2014)

8. da Costa Pereira, C., Tettamanzi, A., Villata, S.: Changing one's mind: Erase or rewind? In: IJCAI 2011, Proceedings of the 22nd International Joint Conference on Artificial Intelligence. pp. 164–171 (2011)
9. Dung, P.M.: On the acceptability of arguments and its fundamental role in non-monotonic reasoning, logic programming and n-person games. *Artif. Intell.* 77(2), 321–358 (1995)
10. Jakobovits, H., Vermeir, D.: Robust semantics for argumentation frameworks. *J. Log. Comput.* 9(2), 215–261 (1999)

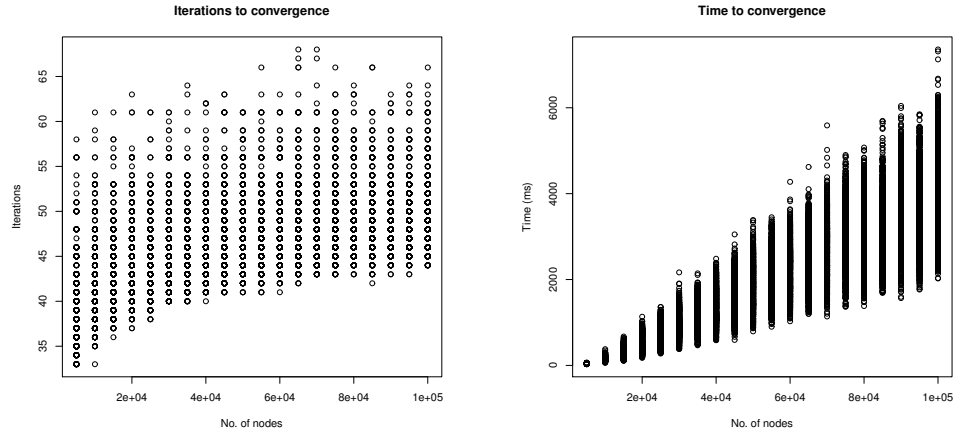


Fig. 10. The Sophia Antipolis benchmark with all weights equal to 1.0.

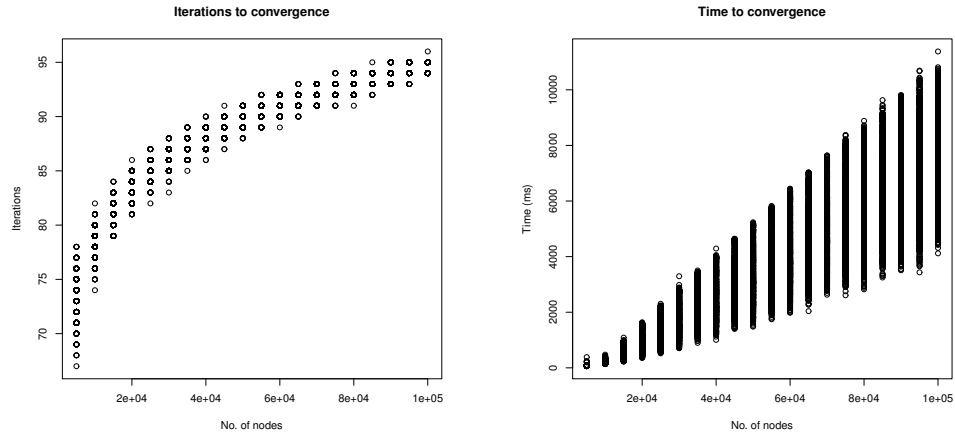


Fig. 11. The Sophia Antipolis benchmark with random weights.

11. Paglieri, F., Castelfranchi, C., da Costa Pereira, C., Falcone, R., Tettamanzi, A., Villata, S.: Trusting the messenger because of the message: feedback dynamics from information quality to source evaluation. *Computational & Mathematical Organization Theory* 20(2), 176–194 (2014)
12. Vallati, M., Cerutti, F., Giacomini, M.: Argumentation frameworks features: an initial study. In: *ECAI 2014 - 21st European Conference on Artificial Intelligence*. pp. 1117–1118 (2014)
13. Verheij, B.: Artificial argument assistants for defeasible argumentation. *Artif. Intell.* 150(1-2), 291–324 (2003)
14. Zadeh, L.A.: Fuzzy sets. *Information and Control* 8, 338–353 (1965)