

A Formal Framework for Combining Legal Reasoning Methods

Henry Prakken

Department of Information and Computing Sciences,
Utrecht University and Faculty of Law, University of
Groningen
The Netherlands
h.prakken@uu.nl

Giovanni Sartor

University of Bologna, Law Faculty-CIRSFID
European University Institute, Florence
Italy
giovanni.sartor@unibo.it

ABSTRACT

This paper proposes a novel argumentation-based approach to combine legal-reasoning methods that each solve a subproblem of an overall legal problem. The methods can be of any nature (for instance, logical, case-based or probabilistic), as long as their input-output behaviour can be described at the metalevel with deductive or defeasible rules. The model is formulated in the *ASPIC*⁺ framework, to profit from its metatheory and explanation methods, and to allow for disagreement about how to solve a subproblem. The model is not meant to be directly implementable but to serve as a semantics for architectures and implementations.

CCS CONCEPTS

• Computing methodologies → Artificial intelligence;

KEYWORDS

Legal argumentation, legal problem solving methods

ACM Reference Format:

Henry Prakken and Giovanni Sartor. 2023. A Formal Framework for Combining Legal Reasoning Methods. In *Nineteenth International Conference for Artificial Intelligence and Law (ICAIL 2023)*, June 19–23, 2023, Braga, Portugal. ACM, New York, NY, USA, 10 pages. <https://doi.org/10.1145/3594536.3595129>

1 INTRODUCTION

In legal reasoning different issues have to be addressed, which may require different reasoning methods. Typically, to decide a case a judge has to assess the facts at stake as well as to identify and apply the relevant legal sources. Consequently, also a lawyer to advice a client on a case, has to address both factual and legal issues, considering what arguments may be relevant to either. Thus both factual and normative reasoning is needed. Further distinctions are possible within the domains of factual and normative reasoning. For instance, in addressing factual issues different reasoning methods can be used: common sense defeasible generalisations may be relied on; scientific theories may be deductively applied; statistical inference may be used; trained intuition by experts or judges may support certain assessment. This plurality of human approaches

is reflected in different methods for automated inference, such as defeasible reasoners, deductive reasoners, statistical tools and neural networks. Similarly, with regard to the law, different reasoning methods can be deployed depending on whether regulations are applied, precedents are referred to, or impacts of relevant interests or values are assessed. These tasks may themselves be analysed into different inference steps. For instance, defeasible legal reasoning may be applied to the facts of a case only if the facts are linked to the predicates that occur in legal rules. The matching of the descriptions of concrete facts and the abstract predicates occurring in legal rules is called ‘subsumption’ by legal theorists. Unless the knowledge representation is supplemented by rules or an ontology that bridges facts and rules, the modelling of subsumption may require case-based reasoning.

In this paper we propose a novel argumentation-based approach to combine the application of legal-reasoning methods to different parts of a legal problem. The model allows for disagreement about what is a suitable way to solve a subproblem. It is inspired by [12], in which an ad-hoc metalevel formalism was proposed. We replace it with *ASPIC*⁺, for two reasons. First, the metatheory of *ASPIC*⁺ automatically applies, such as all results on satisfaction of the rationality postulates and, second, existing explanation methods for *ASPIC*⁺ can be used to explain an outcome, for instance, the argument game for grounded semantics [10], which is arguably very intuitive. Moreover, we show how a problem in the modelling of burden of proof first discussed in [11] can be solved as an application of our approach. Our model is not meant to be directly implementable but to serve as a semantics for implementations or more concrete formal models, with special emphasis on explanations of outcomes of a combination of reasoning methods.

This paper is organised as follows. In Section 2 we summarise the theory of abstract argumentation frameworks and the *ASPIC*⁺ framework on which we will build. Then in Section 3 we present our formal model for combining reasoning methods, which we apply to several examples in Section 4. We conclude in Section 5.

2 FORMAL PRELIMINARIES

In this section we present our formal preliminaries, being the theory of abstract argumentation frameworks [6] and the *ASPIC*⁺ framework for structured approaches to argumentation [9].

2.1 Abstract Argumentation Frameworks

An *abstract argumentation framework* (*AF*) is a pair $(\mathcal{A}, \mathcal{D})$, where $\mathcal{A}$ is a set of arguments and $\mathcal{D} \subseteq \mathcal{A} \times \mathcal{A}$ is a relation of defeat.¹

¹Dung used the term ‘attack’ but since we will interpret it as the *ASPIC*⁺ defeat relation, we will use ‘defeat’.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.
ICAIL 2023, June 19–23, 2023, Braga, Portugal

© 2023 Copyright held by the owner/author(s). Publication rights licensed to ACM.
ACM ISBN 979-8-4007-0197-9/23/06...\$15.00
<https://doi.org/10.1145/3594536.3595129>

The theory of AFs [6] identifies sets of arguments (called *extensions*) which are internally coherent and defend themselves against attack. An argument $A \in \mathcal{A}$ is *defended* by a set by $S \subseteq \mathcal{A}$ if for all $B \in \mathcal{A}$: if B attacks A , then some $C \in S$ attacks B . Then given an AF,

- E is *admissible* if E is conflict-free and defends all its members;
- E is a *complete extension* if E is admissible and $A \in E$ iff A is defended by E ;
- E is a *preferred extension* if E is a $\subseteq$ -maximal admissible set;
- E is a *stable extension* if E is admissible and attacks all arguments outside it;
- $E \subseteq \mathcal{A}$ is the *grounded extension* if E is the least fixpoint of operator F , where $F(S)$ returns all arguments defended by S .

It holds that any preferred, stable or grounded extension is a complete extension. For $T \in \{\text{complete, preferred, grounded, stable}\}$, X is *sceptically* or *credulously* justified under the T semantics if X belongs to all, respectively at least one, T extension.

2.2 The ASPIC⁺ Framework

The ASPIC⁺ framework [9] defines abstract argumentation systems as structures consisting of a logical language $\mathcal{L}$ and two sets $\mathcal{R}_s$ and $\mathcal{R}_d$ of strict and defeasible inference rules defined over $\mathcal{L}$. In this paper we for simplicity assume that $\mathcal{L}$ contains ordinary negation $\neg$ but all new definitions proposed in this paper can be easily adapted to versions of ASPIC⁺ with asymmetric negation, such as negation as failure. Arguments are constructed from a knowledge base (a subset of $\mathcal{L}$) by chaining inferences over $\mathcal{L}$ into acyclic graphs. Formally,

Definition 2.1. [Argumentation System] an argumentation system (AS) is a triple $AS = (\mathcal{L}, \mathcal{R}, n)$ where:

- $\mathcal{L}$ is a logical language with a negation symbol $\neg$;
- $\mathcal{R} = \mathcal{R}_s \cup \mathcal{R}_d$ is a finite set of strict ($\mathcal{R}_s$) and defeasible ($\mathcal{R}_d$) inference rules of the form $\{\varphi_1, \dots, \varphi_n\} \rightarrow \varphi$ and $\{\varphi_1, \dots, \varphi_n\} \Rightarrow \varphi$ respectively (where φ_i, φ are meta-variables ranging over wff in $\mathcal{L}$), such that $\mathcal{R}_s \cap \mathcal{R}_d = \emptyset$. Here, $\varphi_1, \dots, \varphi_n$ are called the *antecedents* and φ the *consequent* of the rule.
- n is a partial function such that $n : \mathcal{R}_d \rightarrow \mathcal{L}$.

Informally, $n(r)$ is a well-formed formula (wff) in $\mathcal{L}$ which says that the defeasible rule $r \in \mathcal{R}$ is applicable, so that an argument claiming $\neg n(r)$ attacks an inference step in the argument using r . We write $\psi = -\varphi$ just in case $\psi = \neg\varphi$ or $\varphi = \neg\psi$. We use $\leadsto$ as a variable ranging over $\{\rightarrow, \Rightarrow\}$. Since the order of antecedents of a rule does not matter, we sometimes write $S \leadsto \varphi$ where S is the set of all antecedents of the rule.

Definition 2.2. [Knowledge bases] A knowledge base in an AS $= (\mathcal{L}, \mathcal{R}, n)$ is a set $\mathcal{K} \subseteq \mathcal{L}$ consisting of two disjoint subsets $\mathcal{K}_n$ (the *axioms*) and $\mathcal{K}_p$ (the *ordinary premises*).

Definition 2.3. [Argumentation theories] An argumentation theory is a pair $(AS, \mathcal{K})$ where AS is an argumentation system and $\mathcal{K}$ a knowledge base in AS .

Definition 2.4. [Arguments] A *argument* A on the basis of an argumentation theory AT is a structure obtainable by applying one or more of the following steps finitely many times:

- (1) φ if $\varphi \in \mathcal{K}$ with: $\text{Prem}(A) = \{\varphi\}$; $\text{Conc}(A) = \varphi$; $\text{Prop}(A) = \{\varphi\}$, $\text{Sub}(A) = \{\varphi\}$; $\text{Rules}(A) = \emptyset$; $\text{DefRules}(A) = \emptyset$; $\text{TopRule}(A) = \text{undefined}$.
- (2) $A_1, \dots, A_n \leadsto \psi$ if $A_1, \dots, A_n$ are arguments such that $\psi \notin \text{Conc}(\{A_1, \dots, A_n\})$ and $\text{Conc}(A_1), \dots, \text{Conc}(A_n) \leadsto \psi \in \mathcal{R}$ with:
 $\text{Prem}(A) = \text{Prem}(A_1) \cup \dots \cup \text{Prem}(A_n)$;
 $\text{Conc}(A) = \psi$;
 $\text{Prop}(A) = \text{Prop}(A_1) \cup \dots \cup \text{Prop}(A_n) \cup \{\psi\}$,
 $\text{Sub}(A) = \text{Sub}(A_1) \cup \dots \cup \text{Sub}(A_n) \cup \{A\}$;
 $\text{Rules}(A) = \text{Rules}(A_1) \cup \dots \cup \text{Rules}(A_n) \cup \{\text{Conc}(A_1), \dots, \text{Conc}(A_n) \leadsto \psi\}$;
 $\text{DefRules}(A) = \text{Rules}(A) \cap \mathcal{R}_d$;
 $\text{TopRule}(A) = \text{Conc}(A_1), \dots, \text{Conc}(A_n) \leadsto \psi$.

$\text{Prem}_n(A) = \text{Prem}(A) \cap \mathcal{K}_n$ and $\text{Prem}_p(A) = \text{Prem}(A) \cap \mathcal{K}_p$. Furthermore, argument A is *strict* if $\text{DefRules}(A) = \emptyset$ and *defeasible* otherwise, and A is *firm* if $\text{Prem}_p(A) = \emptyset$, otherwise A is *plausible*.

The set of all arguments on the basis of AT is denoted by $\mathcal{A}_{AT}$.

Each of the functions Func in this definition is also defined on sets of arguments $S = \{A_1, \dots, A_n\}$ as follows: $\text{Func}(S) = \text{Func}(A_1) \cup \dots \cup \text{Func}(A_n)$. Note that the $\rightarrow$ and $\Rightarrow$ symbols are overloaded to denote both inference rules and arguments. In this paper we do for simplicity not discuss variants of ASPIC⁺ in which the premises of an argument must be consistent (see [9]). All new definitions proposed in this paper directly apply to these versions.

Definition 2.5. [Attack] Argument A attacks argument B iff A *undercuts* or *rebuts* or *undermines* B , where:

- A *undercuts* B (on B') iff $\text{Conc}(A) = -n(r)$ and $B' \in \text{Sub}(B)$ such that B' 's top rule r is defeasible.
- A *rebuts* B (on B') iff $\text{Conc}(A) = -\varphi$ for some $B' \in \text{Sub}(B)$ of the form $B'_1, \dots, B'_n \Rightarrow \varphi$.
- A *undermines* B (on φ) iff $\text{Conc}(A) = -\varphi$ for some $\varphi \in \text{Prem}(B) \cap \mathcal{K}_p$.

Definition 2.6. [Structured Argumentation Frameworks] A structured argumentation framework (SAF) defined by an argumentation theory AT is a triple $(\mathcal{A}, C, \leq)$ where $\mathcal{A}$ is the set of all arguments on the basis of AT , $\leq$ is an ordering on $\mathcal{A}$ and $(X, Y) \in C$ iff X attacks Y .

In this paper we assume that the argument ordering $\leq$ is determined by two preorders $\leq$ on $\mathcal{R}_d$ and $\leq'$ on $\mathcal{K}_p$. The notion of *defeat* is defined as follows. Undercutting attacks succeed as *defeats* independently of preferences over arguments, while rebutting and undermining attacks succeed only if the attacked argument is not stronger than the attacking argument. Below $A < B$ is defined as usual as $A \leq B$ and $B \not\leq A$ and $A \approx B$ as $A \leq B$ and $B \leq A$. In this paper we assume that for no arguments A and B both $A < B$ and $B < A$ hold.

Definition 2.7. [Defeat] Argument A *defeats* argument B iff either A undercuts B ; or A rebuts or undermines B on B' and $A \not\leq B'$.

Abstract argumentation frameworks are then generated from SAFs as follows:

Definition 2.8 (Argumentation frameworks). An abstract argumentation framework (AF) corresponding to a SAF $= (\mathcal{A}, C, \leq)$ is

a pair $(\mathcal{A}, \mathcal{D})$ such that $\mathcal{D}$ is the defeat relation on $\mathcal{A}$ determined by *SAF*.

We can then define nonmonotonic consequence notions for well-formed formulas (wff). A wff $\varphi \in \mathcal{L}$ is *sceptically justified* on the basis of a *SAF* under semantics T if φ is the conclusion of a sceptically justified argument on the basis of the *AF* corresponding to the *SAF* under semantics T , and *credulously justified* on the basis of a *SAF* under semantics T if φ is not sceptically justified and is the conclusion of a credulously justified argument on the basis of the *AF* corresponding to the *SAF* under semantics T .

3 MAIN IDEAS AND FORMALISM

The idea is that a reasoning problem is solved by a connected series of problem solving modules. Each module solves a subproblem of the overall problem. It receives input from and can provide output to other modules. These connections between multiple modules are defined by *output-input metalevel rules*, which are *ASPIC⁺* rules that have as antecedents outputs of one or more modules and have as consequent a single input to another module. A crucial element of our approach is that the antecedents and consequent of each metalevel rule are expressed in the metalanguages of the various modules. Together the sets of modules and output-input metalevel rules specify a problem decomposition. For example, a legal-reasoning module reasoning about civil liability in medical surgery cases could apply legal rules to proven facts with rule-based argumentation, where these facts are provided by a combination of two other modules, a Bayesian network computing posterior probabilities of factual propositions on the basis of the evidence in the case combined with another rule-based argumentation module that determines the relevant burdens and standards of proof.

Another important element of our approach is that while each module has its own problem solving method, its input-output behaviour is represented by a set of *input-output metalevel rules*, which are *ASPIC⁺* rules of which both the antecedents and the consequent are described in the same module's metalanguage: the antecedents describe the inputs to the module while each consequent describes an output of the module according to the module's problem-solving method. For instance, for the evidential Bayesian-network module the antecedents of the input-output metarules would specify a Bayesian network while their consequents would state that particular probability values can be derived from the Bayesian network specified by the metarule's antecedents. Thus while each module reasons with its own method, the application of its method can still be described at the metalevel in *ASPIC⁺*. When these internal input-output metalevel rules of a module are combined with the output-input metalevel rules *between* modules, this allows an explanation at the metalevel in *ASPIC⁺* of how the overall problem is solved by the problem decomposition. Moreover, since these two kinds of metarules (the input-output rules describing a single module and the output-input rules connecting multiple modules) are defeasible, our approach allows to represent conflicts about what is the correct input for a module by way of the existence of conflicting metalevel arguments in *ASPIC⁺*. For instance, in our medical-liability example there could be an alternative argumentation- or scenario-based evidential module that outputs an alternative to the Bayesian-network view on which facts can be proven. Then

alternative input-output metarules of the legal-reasoning method could be triggered to provide arguments for alternative solutions to the medical liability problem.

Our approach is not meant to be directly implementable but it is meant to serve as a semantics for implementations or more concrete formal models. While the output-input metalevel rules have to be specified in advance when specifying the problem decomposition, the input-output metalevel rules can be generated dynamically during an explanation of the overall problem solution, where (as further explained below) only the relevant rules need to be generated.

The first definition formalises an abstract view on problem-solving methods, abstracting from its method and just focussing on its input-output behaviour.

Definition 3.1. A *problem-solving module* M is a tuple of four elements $(\mathcal{L}_M^I, \mathcal{L}_M^O, R_M, \mathcal{R}_M^{io})$ where

- $\mathcal{L}_M^I$ is M 's input language and $\mathcal{L}_M^O$ is M 's output language;
- $R_M: \text{Pow}(\mathcal{L}_M^I) \rightarrow \text{Pow}(\mathcal{L}_M^O)$ is M 's problem-solving mechanism, assigning sets of outputs to sets of inputs;
- $\mathcal{R}_M^{io}$ is the set of all rules of the form $S \Rightarrow \varphi$ where
 - S is a finite subset of $\mathcal{L}_M^I$ and $\varphi \in \mathcal{L}_M^O$; and
 - $\varphi \in R_M(S)$.

Note that $\mathcal{R}_M^{io}$ describes the behaviour of M in a set of *ASPIC⁺* rules defined over the input- and output languages of M .

Definition 3.2. Given a finite set $\mathcal{M}$ of problem-solving modules, a set $\mathcal{R}_M^{oi}$ of *output-input rules* for $\mathcal{M}$ is a set of rules of the form $S \Rightarrow \varphi$ where each element of S (which is finite) is from the output language of some $M \in \mathcal{M}$ and φ is from the input language of some $M \in \mathcal{M}$ that has no element from $\mathcal{L}_M^O$ in S .

In practical applications it makes sense to make sets of output-input rules non-circular but since *ASPIC⁺* can formally handle cyclic arguments, we will not formally impose this requirement here.

Definition 3.3. Given a pair $PS = (\mathcal{M}, \mathcal{R}_M^{oi})$, where $\mathcal{M}$ is a set finite of problem-solving modules and $\mathcal{R}_M^{oi}$ is a finite of input-output rules for $\mathcal{M}$, a *problem specification* is an *ASPIC⁺* structured argumentation framework $(\mathcal{A}^{PS}, \mathcal{C}^{PS}, \leq^{PS})$ defined by an argumentation theory $AT^{PS} = (\mathcal{A}^{PS}, \mathcal{K}^{PS})$ such that

- $\mathcal{L}^{PS}$ is the union of all input and output languages of any $M \in \mathcal{M}$
- $\mathcal{R}_s^{PS} = \{S \rightarrow \varphi \mid S \text{ is finite and } S \vdash_{\text{FOL}} \varphi\}$
- $\mathcal{R}_d^{PS}$ is the union of $\mathcal{R}_M^{oi}$ and all $\mathcal{R}_M^{io}$ of any $M \in \mathcal{M}$
- $\mathcal{K}^{PS} = \mathcal{K}_n^{PS}$ is a subset of $\{\varphi \mid \varphi \in \mathcal{L}_M^I \text{ of some } M \in \mathcal{M}\}$
- $\leq^{PS}$ includes specificity orderings for each $\mathcal{R}_M^{io}$

The metalevel strict rules consist of all classically valid first-order inferences over $\mathcal{L}^{PS}$ since $\mathcal{L}^{PS}$ is a first-order language by construction. The last clause of this definition allows that given inputs of a module can be represented in the global metalevel $\mathcal{K}_n^{PS}$. For example, a module reasoning with an *ASPIC⁺* instantiation will specify an argumentation theory and an argument ordering, and a Bayesian-network module will specify a Bayesian network. In addition, $\mathcal{K}_n^m$ can contain all relevant axioms of the problem-solving modules. For instance, for Bayesian-network modules it

can contain the axioms of probability theory, or for modules based on $ASPIC^+$ it can contain an assumption that $\mathcal{K}_n$ is consistent. As further explained in Section 4, these two design choices are to make sure that inconsistent input specifications (for example, two different unconditional probabilities for the same probabilistic variable in a BN or inputs that make $\mathcal{K}_n$ inconsistent in $ASPIC^+$) generate conflicting arguments at the metalevel.

We now explain the idea of a module's internal input-output rules $\mathcal{R}_M^{io}$ in more detail. The idea is to describe the input-output behaviour of a module as defined by R_M with a set of $ASPIC^+$ rules $S \Rightarrow \varphi$ where S is a set of statements in the input metalanguage of the module and φ is a solution of the module on the assumption that S completely describes the module's input. There are such rules for each possible S and for every solution of that S . Below metarules of this kind will be denoted with subscripted names io . Furthermore, the idea is to only accept the solutions of the most specific rule, that is, the rule that gathers all inputs. To this end, a specificity preference mechanism can be applied to the problem specification that guarantees that in case of conflict only the rule gathering all inputs is applied. It is not guaranteed that there is always a single most specific rule. Situations in which this is not the case reflect genuine conflicts about what should be the input of the module.

To illustrate these ideas with a small formal example, consider a module M with a classical-logic instance of $ASPIC^+$ (a propositional language, building arguments as classical-logic inferences from consistent subsets of the knowledge base, only undermining attack, no preferences). Consider p, q and $\neg q$ as possible elements of the knowledge base. Three relevant arguments as regards these propositions then are p, q and $\neg q$, where q and $\neg q$ defeat each other. Then we have (at least) the following rules in $\mathcal{R}_M^{io}$:

- $io_1: p \in \mathcal{K}_p \Rightarrow p$ is sceptically justified under the grounded semantics on the basis of $\mathcal{K}_p$
- $io_2: q \in \mathcal{K}_p \Rightarrow q$ is sceptically justified under the grounded semantics on the basis of $\mathcal{K}_p$
- $io_3: p \in \mathcal{K}_p, q \in \mathcal{K}_p \Rightarrow q$ is sceptically justified under the grounded semantics on the basis of $\mathcal{K}_p$
- $io_4: q \in \mathcal{K}_p, \neg q \in \mathcal{K}_p \Rightarrow q$ is not sceptically justified under the grounded semantics on the basis of $\mathcal{K}_p$
- $io_5: p \in \mathcal{K}_p, q \in \mathcal{K}_p, \neg q \in \mathcal{K}_p \Rightarrow p$ is sceptically justified under the grounded semantics on the basis of $\mathcal{K}_p$
- $io_6: p \in \mathcal{K}_p, q \in \mathcal{K}_p, \neg q \in \mathcal{K}_p \Rightarrow q$ is not sceptically justified under the grounded semantics on the basis of $\mathcal{K}_p$

Let us suppose that the inputs received from other modules are $q \in \mathcal{K}_p, q \in \mathcal{K}_p, \neg q \in \mathcal{K}_p$. Then all rules are triggered in that their antecedents are in $\mathcal{K}_p$. As regards the status of q there is a conflict between rules io_2 and io_3 on the one hand and io_4 and io_6 on the other. Rule io_4 is more specific than rule io_2 while rule io_6 is more specific than both rules io_2 and io_3 . As a consequence, both rules io_4 and io_6 will be applied at the cost of rules io_2 and io_3 . In addition, since there is no conflict about the status of p , both rule io_1 and rule io_5 apply. ('applied and apply' here mean that they give rise to justified arguments for their consequents.)

Of course, in an implementation it would not be a good idea to generate all of these input-output metarules. Instead, only the input-output behaviour of the fullest specification of $\mathcal{K}_n$ should be observed; this is equivalent to only generating rules io_5 and io_6 .

4 APPLICATIONS

In this section we apply the formalism and ideas from the previous section to several realistic application scenario's in AI & law.

4.1 Combining Rule-based and Evidential Reasoning under Burden of Proof

First we model an example from [5] concerning reasoning under burden of proof.

Example 4.1 (Civil law example). Let us consider a case in which a doctor caused harm to a patient by misdiagnosing his case. Assume that there is no doubt that the doctor harmed the patient: she failed to diagnose cancer, which consequently spread and became incurable. However, it is uncertain whether or not the doctor followed the guidelines governing this case: it is unclear whether she prescribed all the tests that were required by the guidelines, or whether she failed to prescribe some tests that would have enabled cancer to be detected. Assume that, under the applicable law, doctors are liable for any harm suffered by their patients, but they can avoid liability if they show that they were diligent (not negligent) in treating the patient, i.e., that they exercised due care. Thus, rather than the patient having the burden of proving that doctors have been negligent (as it should be the case according to the general principles), doctors have the burden of providing their diligence. Let us assume that the law also says that doctors are considered to be diligent if they followed the medical guidelines that govern the case. In this case, given that the doctor has the burden of persuasion on her diligence, and that she failed to provide a convincing argument for it, the legally correct solution is that she should be ordered to compensate the patient.

In [5] this example was formalised as a single $ASPIC^+$ argumentation theory as follows.

Example 4.2 (Civil law example in $ASPIC^+$).

$\mathcal{K}_n = \{e_1, e_2, e_3\};$

$\mathcal{R} =$

- $r_1: e_1 \Rightarrow \neg \text{guidelines}$
- $r_2: e_2 \Rightarrow \text{guidelines}$
- $r_3: e_3 \Rightarrow \text{harm}$
- $r_4: \neg \text{guidelines} \Rightarrow \neg \text{dueDiligence}$
- $r_5: \text{guidelines} \Rightarrow \text{dueDiligence}$
- $r_6: \text{harm}, \neg \text{dueDiligence} \Rightarrow \text{liable}\}$

We assume no rule preferences. Then the following arguments can be built as displayed in Figure 1.

Without preferences the (direct) defeat relations are the following:

- arguments A4 and A5 defeat each other,
- A7 and A8 defeat each other.

Any of the standard semantics for $ASPIC^+$ yields that A9 is only defensible. However, this does not agree with the fact that the burden of proof is on *dueDiligence* so that, if both *dueDiligence* and $\neg \text{dueDiligence}$ are defensible, as is here the case, A9 should be justified since A6 for *harm* is justified and A7 for $\neg \text{dueDiligence}$ is not overruled.

In [5] this outcome is achieved by applying a nonstandard labelling semantics in which an argument that is burdened is out if it

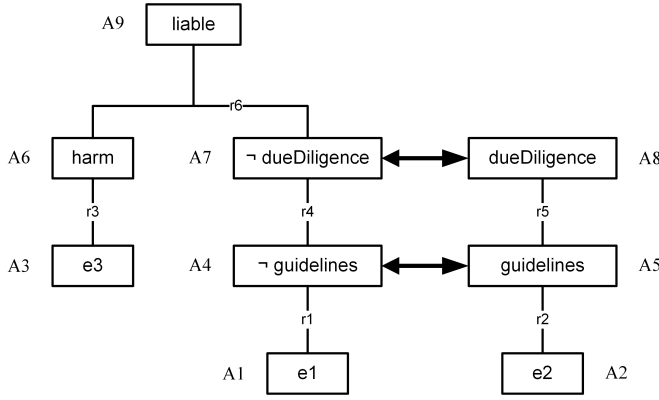


Figure 1: ‘Standard’ arguments in Example 4.1.

is defeated by an argument that is in or undecided (instead of the usual ‘out if defeated by an argument that is in’). In this example this gives the desired result that in any labelling A8 is out so A7 is in, so A7 is justified in any semantics. However, a theoretical drawback of this proposal is that several desirable properties of the various semantics cannot be proven. For example, the grounded extension is not guaranteed to be unique. We now show how our alternative approach avoids this problem while retaining the standard argumentation semantics and treating the example in the desired way.

We divide the above specification into two separate modules: the Legal Rule module LR contains the legal rules of the example, while the Evidential module Ev contains all the evidential rules of the example. In addition we have an $ASPIC^+$ module B for determining the burdens of persuasion that provides the burdens of proof. All three modules apply $ASPIC^+$ but Ev uses the weakest-link argument ordering while LR and B use the last-link ordering. This is since arguably weakest link fits better with epistemic reasoning while last link fits better with normative reasoning [9, Section 3.2]. Although in our example we assume no preferences, this still illustrates the flexibility of our approach.

The overall metalevel $\mathcal{K}_n^m$ of our problem specification contains

$$\begin{aligned} e_1 &\in \mathcal{K}_n(Ev), e_2 \in \mathcal{K}_n(Ev), e_3 \in \mathcal{K}_n(Ev) \\ r_1 &\in \mathcal{R}_d(Ev), r_2 \in \mathcal{R}_d(Ev), r_3 \in \mathcal{R}_d(Ev) \\ r_4 &\in \mathcal{R}_d(LR), r_5 \in \mathcal{R}_d(LR), r_6 \in \mathcal{R}_d(LR) \\ \text{Burden}(\text{dueDiligence}) &\in \mathcal{K}_n(B), \text{Burden}(\text{Harm}) \in \mathcal{K}_n(B) \end{aligned}$$

Next we specify output-input metalevel rules (more precisely, rule schemes for all their ground instances) that take output of the Ev and B modules and provide input for the LR module. (The need for ‘defensible or’ will be explained below).

oi_1 : $\text{Burden}(\varphi)$ is defensible or justified on the basis of $\text{SAF}(B)$, φ is justified on the basis of $\text{SAF}(Ev) \Rightarrow \varphi \in \mathcal{K}_n(LR)$

oi_2 : $\text{Burden}(\varphi)$ is defensible or justified on the basis of $\text{SAF}(B)$, φ is not justified on the basis of $\text{SAF}(Ev) \Rightarrow \neg\varphi \in \mathcal{K}_n(LR)$

So the idea is that only formulas or their negations for which a

burden is defined are outputted from Ev into LR . This yields the following instantiations of rule schemes oi_1 and oi_2 :

$oi_1(\text{harm})$: $\text{Burden}(\text{harm})$ is defensible or justified on the basis of $\text{SAF}(B)$, harm is justified on the basis of $\text{SAF}(Ev) \Rightarrow \text{harm} \in \mathcal{K}_n(LR)$

$oi_2(\text{dueDiligence})$: $\text{Burden}(\text{dueDiligence})$ is defensible or justified on the basis of $\text{SAF}(B)$, dueDiligence is not justified on the basis of $\text{SAF}(Ev) \Rightarrow \neg\text{dueDiligence} \in \mathcal{K}_n(LR)$

We then have an unattacked argument on the basis of our problem specification for the conclusion that *liable* is justified on the basis of $\text{SAF}(LR)$ as displayed in Figure 2 (with some abbreviations). Moreover, the content of and flow between the various modules in the problem specification is displayed in Figure 3 (in which elements that are empty, such as $\mathcal{K}_p$ in all modules, are left implicit).

The meta-argument uses the following input-output rule of LR :

io_1 : $\text{Harm} \in \mathcal{K}_n(LR), \neg\text{dueDiligence} \in \mathcal{K}_n(LR), r_3 \in \mathcal{R}_d(LR) \Rightarrow \text{liable}$ is justified on the basis of $\text{SAF}(LR)$

Furthermore, its subargument for ‘*Burden(harm)* is justified on the basis of $\text{SAF}(B)$ ’ uses the following input-output rule of B :

io_2 : $\text{Burden}(\text{harm}) \in \mathcal{K}_n(B) \Rightarrow \text{Burden}(\text{harm})$ is justified on the basis of $\text{SAF}(B)$

And the subargument for ‘*harm* is justified on the basis of $\text{SAF}(Ev)$ ’ uses the following input-output rule of Ev :

io_3 : $e_3 \in \mathcal{K}_n(Ev), r_3 \in \mathcal{R}(Ev) \Rightarrow \text{harm}$ is justified on the basis of $\text{SAF}(Ev)$

The subargument for ‘*Burden(dueDiligence)* is justified on the basis of $\text{SAF}(B)$ ’ uses the following input-output rule of Ev :

io_4 : $\text{Burden}(\text{dueDiligence}) \in \mathcal{K}_n(B) \Rightarrow \text{Burden}(\text{dueDiligence})$ is justified on the basis of $\text{SAF}(B)$

And the subargument for ‘*dueDiligence* is not justified on the basis of $\text{SAF}(Ev)$ ’ uses:

io_5 : $e_1, e_2 \in \mathcal{K}_n(Ev), r_1, r_2 \in \mathcal{R}(Ev), \leq(Ev) = \approx, \Rightarrow \text{dueDiligence}$ is not justified on the basis of $\text{SAF}(Ev)$

The argument that *liable* is justified can be paraphrased as follows. *liable* is justified in LR since LR contains *harm* and $\neg\text{dueDiligence}$ as necessary facts. For *harm* this is so since it is justified in Ev since e_3 is a necessary fact in Ev and r_3 a rule in Ev , which yields an unattacked argument for *harm* in LR . For $\neg\text{dueDiligence}$ this is so since $\text{Burden}(\text{dueDiligence})$ is justified in B (since it is a necessary fact in B) and *dueDiligence* is not justified in Ev . The latter is so since Ev not only contains e_2 as a fact and r_2 and r_5 as defeasible rules but also contains e_1 as a fact and r_1 and r_4 as defeasible rules. Then [here the local explanation method of Ev can be inserted].

We now modify the example to illustrate the treatment of conflicting inputs for a module. Instead of having $\text{Burden}(\text{dueDiligence}) \in$

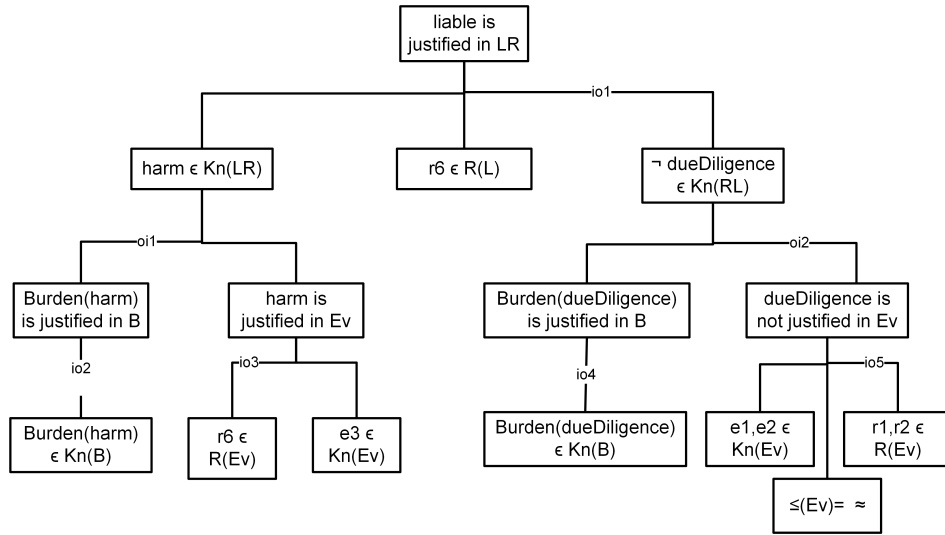


Figure 2: Metalevel arguments in Example 4.1 (1).

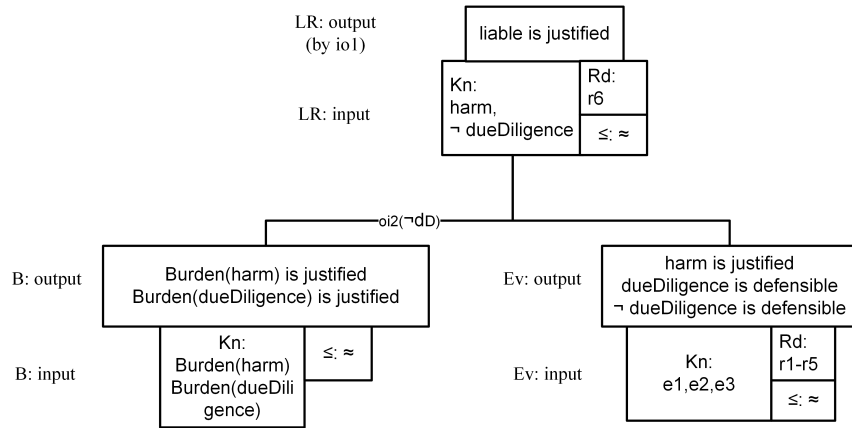


Figure 3: Problem solving modules in Example 4.1 (1).

$\mathcal{K}_n(B)$ we now have both $\text{Burden}(\text{dueDiligence}) \in \mathcal{K}_p(B)$ and $\text{Burden}(\neg \text{dueDiligence}) \in \mathcal{K}_p(B)$. We also add to $\mathcal{K}_n(\mathcal{M})$ that $\mathcal{K}_n(LR)$ is consistent. Then $\text{Burden}(\text{dueDiligence})$ is defensible in B , which still triggers oi_1 . However, now oi_2 is triggered for both burden statements:

$oi_2(\text{dueDiligence})$: $\text{Burden}(\text{dueDiligence})$ is defensible or justified on the basis of $\text{SAF}(B)$, dueDiligence is not justified on the basis of $\text{SAF}(\text{Ev}) \Rightarrow \neg \text{dueDiligence} \in \mathcal{K}_n(LR)$

$oi_2(\neg \text{dueDiligence})$: $\text{Burden}(\neg \text{dueDiligence})$ is defensible or justified on the basis of $\text{SAF}(B)$, $\neg \text{dueDiligence}$ is not justified on the basis of $\text{SAF}(\text{Ev}) \Rightarrow \text{dueDiligence} \in \mathcal{K}_n(LR)$

So for both dueDiligence and $\neg \text{dueDiligence}$ we obtain arguments that their contradictories are in $\mathcal{K}_n(LR)$. For the conclusion that $\neg \text{dueDiligence}$ is in $\mathcal{K}_n(LR)$ it is the same argument as in Figure 2 except that the intermediate conclusion ' $\text{Burden}(\text{dueDiligence})$ is

justified in B ' must now be replaced with ' $\text{Burden}(\text{dueDiligence})$ is defensible in B ' and the premise ' $\text{Burden}(\text{dueDiligence}) \in \mathcal{K}_n(B)$ ' must be replaced with ' $\text{Burden}(\text{dueDiligence}) \in \mathcal{K}_p(B)$ '. The argument for the conclusion that dueDiligence is in $\mathcal{K}_n(LR)$ is of a similar form with the obvious alternative replacements.

Together with the axiom in $\mathcal{K}_M$ that $\mathcal{K}_{LR}$ is consistent, this then also yields arguments for both dueDiligence and $\neg \text{dueDiligence}$ that they are not in $\mathcal{K}_n(LR)$ (we leave it to the reader to verify the construction of these arguments). So both arguments for the conclusions that they are in $\mathcal{K}_n(LR)$ are defeated. Actually, this also holds for both arguments for the conclusions that they are not in $\mathcal{K}_n(LR)$, since we assume no preferences between o-i rules at the metalevel. The upshot of all this is that there now also is an argument that liable is not justified in LR , since no argument for it can be constructed. Or, more precisely, we now obtain an additional version of LR with a different knowledge base in which no such argument can be constructed. The new situation as regards the

modules is displayed in Figure 4. Note that we now we have the following equally specific applicable io-rules in LR :

io_1 : $harm \in \mathcal{K}_n(LR), \neg dueDiligence \in \mathcal{K}_n(LR), r_3 \in \mathcal{R}_d(LR) \Rightarrow liable$ is justified on the basis of $SAF(LR)$
 io_6 : $harm \in \mathcal{K}_n(LR), \neg dueDiligence \notin \mathcal{K}_n(LR), r_3 \in \mathcal{R}_d(LR) \Rightarrow liable$ is not justified on the basis of $SAF(M(R))$.

The result is that both conclusions on whether *liable* is justified in LR are defensible on the basis of the overall metalevel SAF that specifies the problem decomposition. Note that here it does not matter that these conclusions do not contradict each other (since formally they refer to different versions of LR with different knowledge bases): to resolve the conflict, a choice has to be made in module B between the conflicting arguments on who has the burden of proof concerning *dueDiligence*. This conflict is modelled at the metalevel as two defensible arguments on the basis of the metalevel SAF .

Concluding this subsection, note that we have provided an alternative solution to a theoretical problem first noted by [11] and also discussed by [5, 14]. The problem is how to reconcile the concept of shifts in the burden of proof with ‘standard’ semantics for abstract argumentation frameworks. While in [11] alternative semantics were proposed as a solution, in [14] a solution within the standard semantics was proposed. In [5] a problem with that solution was identified and another nonstandard solution was proposed. However, as noted above, a theoretical drawback of this proposal is that several desirable properties of the various semantics cannot be proven, such as uniqueness of the grounded extension. We have instead suggested that the solution lies in decomposing a reasoning problem into multiple components that each can have their own definitions and semantics, while shifts in the burden of proof are modelled in the connections between these components, the semantics of which is also standard, namely the theory of Section 2. Our approach thus also agrees with [15]’s suggestion that legal proof debates are essentially meta-theoretic.

4.2 Reasoning with and about Bayesian Networks

We next show how a Bayesian-network (BN) application can be combined with rule-based argumentation. We first illustrate how a BN can generate facts for LR . We limit ourselves to BNs with only boolean variables.

4.2.1 From a BN to rule-based argumentation. A Bayesian network is a pair (G, Pr) where G is a directed acyclic graph (V, D) , where V is a set of boolean probabilistic variables and $D \subseteq V \times V$ is a set of probabilistic dependencies, and Pr is probability function which for every $v \in V$ specifies $Pr(v \mid parents(v))$ (the conditional probability tables for each $v \in V$). As for notation, for any $v \in V$ we often write that v is true as v (letting the context disambiguate) and that v is false as $\neg v$. Finally, a subset E of V , the evidence, can be declared as certainly true.

We replace the $ASPIC^+$ -style Ev module of the previous subsection with a possible BN of the same evidential problem. The BN is visualised in Figure 5, where the probabilities are the posterior ones after entering the evidence $(e_1, e_2$ and $e_3)$. The probability that a variable is false (true) is listed at the top (bottom). We make no

claims as to the adequacy of this BN for modelling the example; we only use it to illustrate how a logical and probabilistic reasoning method can be combined in our formal framework.

First, the structure of the BN is specified in $\mathcal{K}^n(Ev)$. We need to specify which variables are in V and which probabilistic dependencies are in D . So we have the following statements in $\mathcal{K}_n(Ev)$.

$dueDiligence \in V(Ev), Harm \in V(Ev), Guidelines \in V(Ev), e_1 \in V(Ev), e_2 \in V(Ev), e_3 \in V(Ev)$.
 $(Harm, dueDiligence) \in D(Ev), (dueDiligence, Guidelines) \in D(Ev)$,
 $(Harm, e_3) \in D(Ev), (Guidelines, e_1) \in D(Ev), (Guidelines, e_2) \in D(Ev)$.

We also assume that the contents of the conditional probability tables (CPT) are given, so $\mathcal{K}_n(Ev)$ also contains the following statements (note that $Pr(\neg v \mid v') = 1 - Pr(v \mid v')$, so it can be left implicit).

$Pr(Harm) = 0.1 \in Pr(Ev)$ (This specifies the prior probability of harm).
 $Pr(e_3 \mid Harm) = 0.92 \in Pr(Ev), Pr(e_3 \mid \neg Harm) = 0.1 \in Pr(Ev)$ (These probabilities specify the CPT for e_3).
 $Pr(dueDiligence \mid Harm) = 0.3 \in Pr(Ev), Pr(dueDiligence \mid \neg Harm) = 0.75 \in Pr(Ev)$ (the CPT for *dueDiligence*).
 $Pr(Guidelines \mid dueDiligence) = 0.7 \in Pr(Ev), Pr(Guidelines \mid \neg dueDiligence) = 0.2 \in Pr(Ev)$ (the CPT for *Guidelines*) (the CPT for *Guidelines*).
 $Pr(e_1 \mid Guidelines) = 0.05 \in Pr(Ev), Pr(e_1 \mid \neg Guidelines) = 0.9 \in Pr(Ev)$ (the CPT for e_1)
 $Pr(e_2 \mid Guidelines) = 0.9 \in Pr(Ev), Pr(e_2 \mid \neg Guidelines) = 0.05 \in Pr(Ev)$ (the CPT for e_2)

Finally, we need to specify the evidence in the BN module:

$e1 \in E(Ev), e2 \in E(Ev), e3 \in E(Ev)$

Next we specify the metarules that connect the output of the Ev and B modules to the input of the LR module. The following rules mimic the format of rules oi_1 and oi_2 above and attempt to formalise the proof standard of *on the balance of probabilities*. (If one wants to model debates about what is the correct proof standard, then these rules can be refined, but we omit a discussion of this for simplicity.)

oi_4 : Burden(φ) is defensible or justified on the basis of $SAF(B)$, E is all the evidence in Ev , $Pr(\varphi \mid E) > 0.5$ on the basis of the BN in Ev , $\Rightarrow \varphi \in \mathcal{K}_n(LR)$

oi_5 : Burden(φ) is defensible or justified on the basis of $SAF(B)$, E is all the evidence in Ev , $Pr(\varphi \mid E) \leq 0.5$ on the basis of the BN in Ev $\Rightarrow \neg \varphi \in \mathcal{K}_n(LR)$

Then metalevel arguments can be constructed just as in Figure 3 by replacing the subarguments using input-output rules io_3 and io_5 of Ev with the following rules (assuming that $Pr(dueDiligence \mid E) < 0.5$, which is the case if $Pr(dueDiligence \mid Harm)$ is changed to 0.1 as in the next subsection).

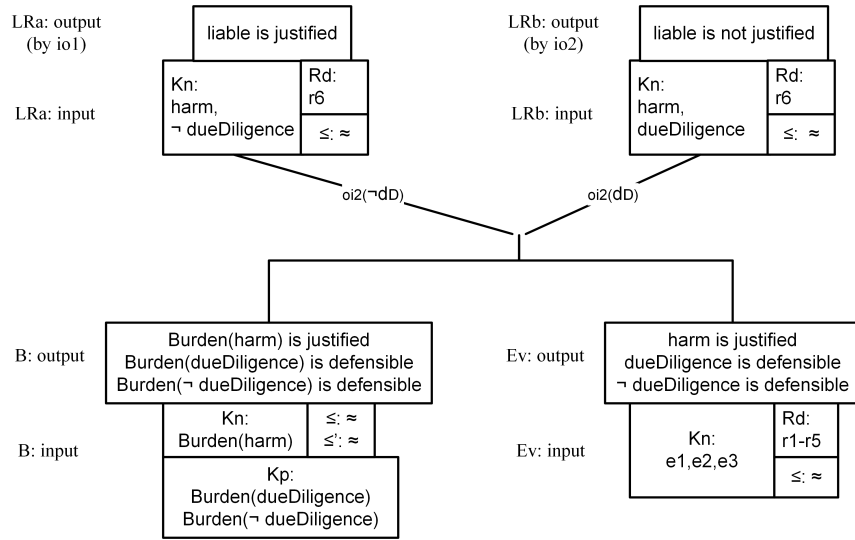


Figure 4: Problem solving modules in Example 4.1 (2).

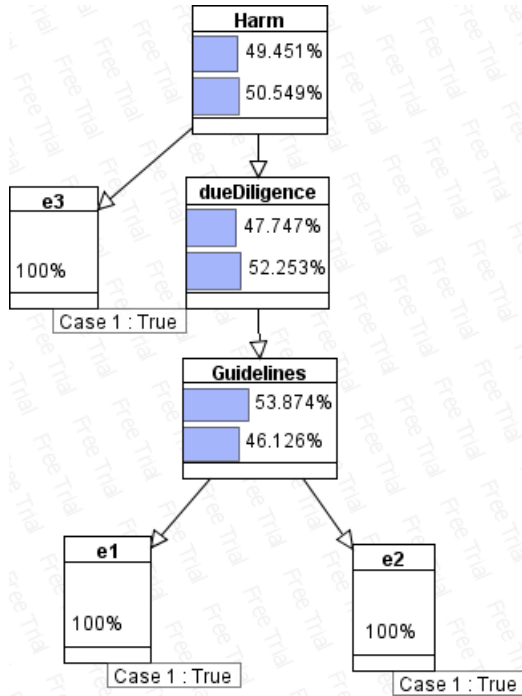


Figure 5: A Bayesian network module.

io_5' : (specification of BN) $\Rightarrow Pr(harm | E) > 0.5$ on the basis of the BN in Ev.

io_5'' : (specification of BN) $\Rightarrow Pr(dueDiligence | E) < 0.5$ on the basis of the BN in Ev.

4.2.2 From argument-scheme-based argumentation to a BN. Next we illustrate how the input probabilities of the Ev module can be argued about in another module, following the approach of [13].

We assume an $ASPIC^+$ -style module ArS with argument schemes for arguing about probabilities for a BN. Among these schemes is the expert testimony scheme, which is a defeasible rule scheme in $\mathcal{R}_d(ArS)$:

$exp(x)$: x is an expert as regards as regards $Pr(H|E)$, x says that $Pr(H|E) = n \Rightarrow Pr(H|E) = n$

For our case we have the following facts in $\mathcal{K}_n(ArS)$:

- (1) Tony is an expert as regards $Pr(dueDiligence|Harm)$, (2) Tony says that $Pr(dueDiligence|Harm) = 0.3$
- (3) Lucy is an expert as regards $Pr(dueDiligence|Harm)$, (4) Lucy says that $Pr(dueDiligence|Harm) = 0.1$

Furthermore, we want that disagreeing experts can result in multiple conflicting inputs for the Ev module, so we write the following output-input metarule:

oi_6 : $Pr(H|E) = n$ is defensible or justified on the basis of $ArS \Rightarrow Pr(H|E) = n \in Pr(BN)$

Then if two contradictory probability statements are both defensible, then oi_6 applies to both of them and then the axioms of probability, which are in $\mathcal{K}_n(\mathcal{M})$, will similarly to in Section 4.1 give rise to conflicting arguments about whether these statements are or are not in $\mathcal{K}_n(BN)$. In consequence, there will be defensible metalevel arguments for, respectively, the conclusions that $(\neg)dueDiligence$ is in $\mathcal{K}_n(LR)$ (even though they do not defeat each other). That in turn gives rise to two versions of the LR module as in Figure 4, which then again gives two defensible metalevel arguments that *liable* is, respectively, is not justified. All this is visualised in Figure 7.

In this example the metalevel conflict arose because of defensible arguments within an (argumentation-based) module. Another

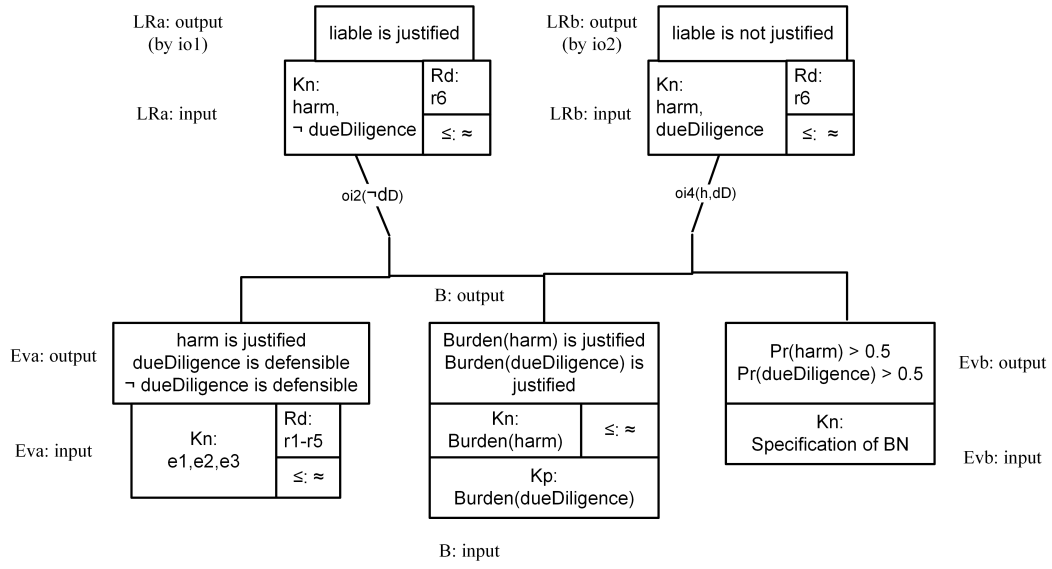


Figure 6: Conflicting evidential modules.

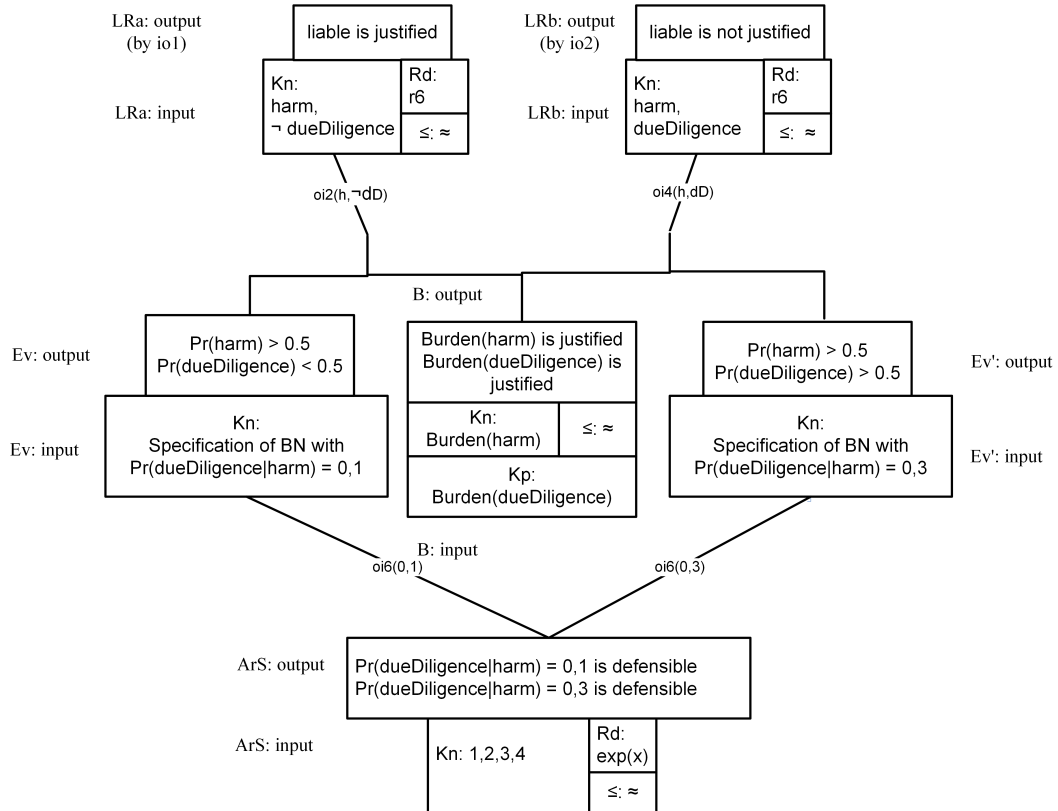


Figure 7: Conflicting output from argumentation about probabilities.

way in which conflicting inputs can arise is when two alternative modules provide conflicting inputs for a model. To illustrate this,

assume that two different versions of *Ev* are both proposed as the right way to solve the evidential problem of whether due diligence

can be proven: *Eva*, which is the above *ASPIC*⁺-based *Ev* module, and *Evb*, which is the first of the above BN-based *Ev* modules, with $Pr(\text{dueDiligence}|\text{harm}) = 0.3$. Suppose also that *dueDiligence* is not justified in *Eva* but that according to *Evb* we have that $Pr(\text{dueDiligence} \mid E) > 0.5$ (as in Figure 5). Then *Eva* wants to put $\neg\text{dueDiligence}$ into $\mathcal{K}_n(LR)$ while *Evb* instead proposes to put *dueDiligence* into $\mathcal{K}_n(LR)$ (namely, on the basis of oi_4). Together with a metalevel axiom in $\mathcal{K}_n^m$ that $\mathcal{K}_n(LR)$ is consistent, this again yields two rebutting metalevel arguments about what should be put into $\mathcal{K}_n(LR)$. Then rules io_1 and io_2 describing two versions of *LR* can both be applied and again two defensible arguments for, respectively, ‘*liable* is justified on the basis of *SAF(LR)*’ and ‘*liable* is not justified on the basis of *SAF(LR)*’ can be created. The situation as regards the modules is displayed in Figure 6.

4.3 Other possible applications

We end this section with a brief discussion of some other possible applications of our approach. First, it could be used for combining rule- and rule-based reasoning. For instance, the IBP system [1] combines a deductive model of rule-based reasoning with HYPO/CATO-style case-based reasoning on whether the conditions of a rule are satisfied. A remodelling in our approach would make it possible to replace the deductive rule model with a defeasible one, such as *ASPIC*⁺ or Carneades [8]. Moreover, at the ‘input’ side the case-based module could be connected with another reasoning model for determining which factors apply to a case. As illustrated in Section 4, we could then model conflicts about which factors apply in a case, leading to alternative copies of the case-based and rule-based modules.

Another application is combining rule- or logic-based reasoning with machine-learning applications for establishing the truth of factual predicates. For example, rules on early release from prison may have conditions about the probability of recidivism, which may be verified by machine-learning approaches.

5 CONCLUSION

In this paper we proposed a novel argumentation-based approach to combine the application of legal-reasoning methods to different parts of a legal problem. The model was formulated in the *ASPIC*⁺ framework, to profit from its metatheory and explanation methods. An interesting feature of our approach is that it can model disagreement about what is a suitable method to model a legal problem. The model is not meant to be directly implementable but to serve as a semantics for architectures and implementations, with special emphasis on explanations of outcomes of a combination of reasoning methods. We illustrated the power of our approach with applications to evidential reasoning and reasoning under burden of proof. Among other things, we proposed a novel way to reconcile ‘standard’ argumentation semantics with the concept of shifts in the burden of proof, based on the idea that a standard semantics can be applied in individual reasoning modules while shifts in the burden of proof are modelled in the connections between various modules, which have an *ASPIC*⁺ semantics.

The idea to combine reasoning methods in a modular way is not new in AI. In [4] a modular approach to logic programming is proposed which inspired much further work. For instance, in

[7] a modular version of assumption-based argumentation [16] was proposed for representing and reasoning with alternative legal theories. Modules can attack each other but the logic of each of them is assumption-based argumentation. Other related work is [3], in which a framework for multi-context systems is proposed where each context can have its own logic (monotonic or not) while bridge rules specify the information flow between contexts. These bridge rules are nonmonotonic since in their body they can refer to non-consequences of a module. Their semantics is, unlike in our approach, not given by an existing nonmonotonic logic but by a formalism especially developed for multi-context systems. While [3]’s approach thus has some similarities with ours, there are also differences and [3] do not discuss legal applications. Finally, in [2] modularity is studied in the context of abstract argumentation frameworks [6]. In future work it would be interesting to investigate the formal relations between our approach and all this related work.

REFERENCES

- [1] K.D. Ashley and S. Brüninghaus. 2009. Automatically classifying case texts and predicting outcomes. *Artificial Intelligence and Law* 17 (2009), 125–165.
- [2] P. Baroni, G. Boella, F. Cerutti, M. Giacomin, L. van der Torre, and S. Villata. 2014. On the Input/Output behavior of argumentation frameworks. *Artificial Intelligence* 217 (2014), 144–197.
- [3] G. Brewka and T. Eiter. 2007. Equilibria in heterogeneous nonmonotonic multi-context systems. In *Proceedings of the 22nd National Conference on Artificial Intelligence (AAAI-07)*, 385–390.
- [4] A. Brogi, P. Mancarella, D. Pedreschi, and F. Turini. 1994. Modular Logic Programming. *ACM Transactions on Programming Languages and Systems* 16 (1994), 1361–1398.
- [5] R. Calegari and G. Sartor. 2021. Burdens of persuasion and standards of proof in structured argumentation. In *Logic and Argumentation. 4th International Conference, CLAR 2021 Hangzhou, China, October 20–22, 2021, Proceedings*, P. Baroni C. Benz Müller and Y.N. Wang (Eds.). Number 13040 in Springer Lecture Notes in AI. Springer, 40–59.
- [6] P.M. Dung. 1995. On the acceptability of arguments and its fundamental role in nonmonotonic reasoning, logic programming, and *n*-person games. *Artificial Intelligence* 77 (1995), 321–357.
- [7] P.M. Dung and G. Sartor. 2011. The modular logic of private international law. *Artificial Intelligence and Law* 19 (2011), 233–261.
- [8] T.F. Gordon, H. Prakken, and D.N. Walton. 2007. The Carneades model of argument and burden of proof. *Artificial Intelligence* 171 (2007), 875–896.
- [9] S. Modgil and H. Prakken. 2018. Abstract rule-based argumentation. In *Handbook of Formal Argumentation*, P. Baroni, D. Gabbay, M. Giacomin, and L. van der Torre (Eds.). Vol. 1. College Publications, London, 286–361.
- [10] H. Prakken. 1999. Dialectical proof theory for defeasible argumentation with defeasible priorities (preliminary report). In *Formal Models of Agents (Springer Lecture Notes in AI, 1760)*, J.-J.Ch. Meyer and P.-Y. Schobbens (Eds.). Springer Verlag, Berlin, 202–215.
- [11] H. Prakken. 2001. Modelling defeasibility in law: logic or procedure? *Fundamenta Informaticae* 48 (2001), 253–271.
- [12] H. Prakken. 2008. Combining modes of reasoning: an application of abstract argumentation. In *Proceedings of the 11th European Conference on Logics in Artificial Intelligence (JELIA 2008) (Springer Lecture Notes in AI, 5293)*, S. Hoelldobler, C. Lutz, and H. Wansing (Eds.). Springer Verlag, Berlin, 349–361.
- [13] H. Prakken. 2020. A new use case for argumentation support tools: supporting discussions of Bayesian analyses of complex criminal cases. *Artificial Intelligence and Law* 28 (2020), 27–49.
- [14] H. Prakken and G. Sartor. 2011. On modelling burdens and standards of proof in structured argumentation. In *Legal Knowledge and Information Systems. JURIX 2011: The Twenty-fourth Annual Conference*, K. Atkinson (Ed.). IOS Press, Amsterdam etc., 83–92.
- [15] B. Schafer. 2009. Twelve angry men or one good woman? Asymmetric relations in evidentiary reasoning. In *Legal Evidence and Proof: Statistics, Stories, Logic*, H. Kaptein, H. Prakken, and B. Verheij (Eds.). Ashgate Publishing, Farnham, 255–282.
- [16] F. Toni. 2014. A tutorial on assumption-based argumentation. *Argument and Computation* 5 (2014), 89–117.