

Integrated Specification and Verification of Security Protocols and Policies

Report**Author(s):**

Frau, Simone; Torabi Dashti, Muhammad

Publication date:

2011

Permanent link:

<https://doi.org/10.3929/ethz-a-006804400>

Rights / license:

In Copyright - Non-Commercial Use Permitted

Integrated Specification and Verification of Security Protocols and Policies

Simone Frau
ETH Zürich

Mohammad Torabi-Dashti
ETH Zürich

Abstract—We propose a language for formal specification of service-oriented architectures. The language supports the integrated specification of communication level events, policy level decisions, and the interaction between the two. We show that the reachability problem is decidable for a fragment of service-oriented architectures. The decidable fragment is well suited for specifying, and reasoning about, security-sensitive architectures. In the decidable fragment, the attacker controls the communication media. The policies of services are centered around the *trust application* and *trust delegation* rules, and can also express RBAC systems with role hierarchy. The fragment is of immediate practical relevance: We report on the specification and verification of two security-sensitive architectures, stemming from the e-government and e-health domains.

I. INTRODUCTION

Context. Security protocols and authorization logics are two major techniques used in securing software systems. A central role of any (security) protocol is to give meaning to the messages that are exchanged in the course of the protocol [1]. For example, a signed X.509 certificate sent by a certificate authority is in many security protocols *meant* to imply that the authority endorses the public key and its owner, mentioned in the certificate. There are several ways to make the meanings of messages, and in general actions, of a protocol explicit, e.g. by associating epistemic effects to the actions [2]. In this paper, we propose a formal language for specifying service-oriented architectures, in which

- the messages received by a service are interpreted in terms of policy statements of the service, and
- the authorization policies of the service constrain the actions the service can perform.

The proposed language is well suited for integrated specification of security protocols and authorization policies in service-oriented architectures. We see a service-oriented architecture as a collection of finitely many services which communicate over insecure media. Each service consists of a number of processes that run in parallel and share a *policy engine*. Processes communicate by sending and receiving messages, as it is usual in asynchronous message passing environments. Each send event is constrained by a *guard*, and each receive event leads to an *update*. Guards and updates belong to the *policy level*, as opposed to send and receive events which constitute the *communication level*. In anthropomorphic terms, services “think” at the policy level, and “talk” at the communication level.

From an operational point of view, guards are predicates which, if derivable by the policy engine of a service, allow the service to perform a corresponding send action, cf. Dijkstra’s guarded command language. Updates are also predicates at the policy level. When a service receives a message in one of its processes, it adds the corresponding update predicates to its policy engine. Intuitively, updates associate meanings to the messages a service receives in terms of predicates in the policy level. The notion of updates is similar to the *assumptions* which are relied upon after receiving a message, in the trust management model of Guttman et al. [3].

Motivations. The separation between the communication and policy levels is a useful abstraction for better understanding each of these levels. Indeed, distributed authorization logics, such as [4]–[6], typically abstract away the communication level events by assuming that all the policy statements exchanged among the participants are signed certificates. This frees the modelers from specifying the exact routes through which the statements travel, etc. The abstraction however obscures how each of the policy statements are represented (as messages) in a given application, how messages are interpreted as policy statements, whether there is a place for misinterpretation, etc. For instance, one would expect that the statement “Ann says employee(Piet)” is added to the policy engine of a service, only after the service receives a message which is *meant* to indicate that Piet is Ann’s employee. However, the meanings of messages (determined by their format, who has signed them, etc.) is often not specified in the policy level. Therefore, in a concrete environment, it is unclear whether or not the attacker can fake a message which would mean that Piet is Ann’s employee, even though he is not. Similarly, formal specifications of security protocols, e.g. as in [7], [8], fully detail the format of the exchanged messages, while the meanings of messages in terms of policy statements are left unspecified.

While maintaining the separation between the communication and policy levels, we believe that, for a thorough security analysis, the interaction between the two levels must also be defined precisely. A typical specification in our proposed language thus consists of three components: communication level events, policy level decisions, and the interface between the two. As the interface between the levels is explicitly present in the specifications, a more precise security analysis of service-oriented architectures becomes possible. This singles out our specification language from the formalisms which focus on

either the communication or the policy level, and hence neglect their interactions.

Decision procedure. We assume that the attacker is in direct control of the communication media, i.e. messages are passed through the attacker. The message composition capabilities of the attacker may, for example, reflect the Dolev-Yao threat model [9]. The attacker can indirectly affect the policies of the participating services, by sending tampered messages which in turn affect the update predicates.

A generic *reachability* problem is defined for service-oriented architectures specified in the language. The reachability problem subsumes the secrecy problem for security protocols and the safety problem for authorization policies (these notions are defined in the paper). The reachability problem turns out to be undecidable in general, even when assuming a finite bound on the number of participating services. We give a decision algorithm for the reachability problem under the following two conditions: (1) the message composition and decomposition capabilities of the attacker reflect the Dolev-Yao threat model, and (2) policy engines of (honest) services are centered around the *trust application* and *trust delegation* rules à la DKAL [6], besides *type-1* theories (formally defined in section IV). Type-1 theories are sufficiently expressive for modelling, e.g., RBAC systems with role hierarchy. The trust application and trust delegation rules, which are the core of many distributed authorization logics [4]–[6], intuitively state that

- (Trust application) If Ann trusts Mike on statement f , and Mike says f , then Ann believes f holds.
- (Trust delegation) If Ann trusts Mike on statement f , and Mike delegates the right to state f to, e.g., Piet, then Ann trusts Piet on statement f .

Trust delegation often contributes to the resilience and flexibility of access control systems. In practice, however, for a given application, trust delegation may, or may not, be allowed. Our formalism and decision algorithm can be adapted to exclude (transitive) trust delegation, if desired.

The decidable fragment is of practical interest: several industrial service-oriented architectures studied in the context of AVANTSSAR [10] (The EU Project on Automated Validation of Trust and Security of Service-oriented Architectures) fall into this fragment. As a comprehensive example, a case study on specifying and verifying an on-line car registration service [10], stemming from the European initiative for *points of single contact*, is reported in this paper.

To prove our decidability result, we encode the derivation of authorization predicates in the policy engine of a service into message inference trees induced by the Dolev-Yao deduction rules. The encoding benefits us in two ways: (1) it simplifies the decidability proof, and (2) it allows us to build upon existing tools which have been originally developed for verifying security protocols. In particular, we have extended the constraint solver of Millen and Shmatikov [7] in Prolog to validate service-oriented architectures.

Note that verification algorithms for correctness of security protocols and authorization logics have been mostly developed

in isolation. For instance, it has been shown that the secrecy problem is decidable for security protocols with a bounded number of sessions [7], [8]. For these results the local computational power of the processes is limited to pattern matching, hence not fully accounting for authorization policies of the participants. Likewise, (un)decidability results for the safety problem in the HRU access control matrix model [11], and authorization logics such as [4]–[6] abstract away communication level events and their effects on policy level decisions. In contrast, our decision algorithm for reachability takes the communication and policy levels into account, and also covers the interface between them.

Related work. Our proposed language can be used to specify security protocols as it is common in the literature, e.g. see [7]. Authorization policies are modeled as logic programs in the language. Logic programs have been extensively used for specifying and reasoning about policies, e.g. see BINDER [4], SECPAL [5], and DKAL [6].

Recent progress in analyzing business processes, augmented with authorization policies, is related to our work, e.g. see [12], [13]. These studies focus on using specific formalisms and techniques for selected case studies, and thus do not consider decidability issues in general. A notable exception to this is [14], where workflows, policy level predicates and their interfaces are all formalized in first-order logic. Reachability is not considered in [14].

Road map. Section II describes an on-line car registration service: this serves as our running example. Section III introduces the syntax and semantics of the language we propose for specifying service-oriented architectures. There, we also formally define the reachability problem. A decidable fragment of architectures is identified in section IV. Section V gives a formal specification of the running example in the decidable fragment. A constraint solving algorithm for deciding reachability in the fragment is given in section VI. Section VII concludes the paper. Some of the proofs are relegated to appendix A in order not to disrupt the flow of the paper.

II. A RUNNING EXAMPLE: CRP

We give an informal description of a service-oriented architecture for online car registration procedure, originated from the European initiative for *points of single contact*, see [10]. We refer to the case study as CRP. A formal model of CRP and its verification results are given in subsequent sections.

CRP involves a number of parties: Mike the new owner of a car, Piet the employee of the car registration office, Ann the head of the car registration office, the human resources department of the office, called *hr*, and the central repository server, referred to as *cr*.

Mike buys a car, but before he is allowed to drive it, he must register the car at the car registration office. The office provides an online registration service. After producing a document containing all the necessary data for registering the car, Mike sends the document to one of the employees of the office, Piet. If the document is valid, Piet sends it to the central repository server to be permanently stored. The *cr* allows only

the employees of the car registration office to write on the server, and Piet is not initially known to the *cr* as an employee. The *cr* trusts Ann, the head of the office, on who is an employee of the office. Ann, however, has decided to delegate this task (or, right) to the human resources department *hr*, and communicates this decision to the *cr* with a certificate. Consequently, the *cr* inquiries the *hr* on the status of Piet, to which the *hr* replies with a certificate stating that Piet is an employee of the office. Finally, the *cr* accepts Piet's request to store the document. After storing the document, the *cr* sends back an acknowledgement to Piet. Piet, in turn, sends a token of successful completion of the registration to Mike.

Below, we present the message exchange pattern for the CRP. The usual primitives for security protocols are employed: asymmetric encryptions $\{\cdot\}$, signatures $\text{sig}(\cdot, \cdot)$, public key constructors $\text{pk}(\cdot)$, hash functions $h(\cdot)$ and pairing $(\cdot, \cdot)$. When confusion is unlikely, we simply write x, y for the pair (x, y) , and write $[x]_a$ for $a, x, \text{sig}(a, x)$. We assume these cryptographic operators are ideal, à la Dolev and Yao [9]. Below, terms in sans-serif are flags, i.e. unique constants which denote the purpose of their accompanying messages.

$$\begin{aligned} \text{mike} \rightarrow \text{piet} &: \{\text{mike}, \text{doc}\}_{\text{pk}(\text{piet})}, [h(\text{doc})]_{\text{mike}} \\ \text{piet} \rightarrow \text{cr} &: \{\text{mike}, \text{doc}\}_{\text{pk}(\text{cr})}, [\text{ann}, h(\text{mike}, \text{doc})]_{\text{piet}} \\ \text{cr} \rightarrow \text{piet} &: [h(\text{mike}, \text{doc})]_{\text{cr}} \\ \text{piet} \rightarrow \text{mike} &: [h(\text{mike}, \text{doc}), \text{success_token}]_{\text{piet}} \end{aligned}$$

The delegation of Ann's right to the *hr*, and the *hr*'s attestation that Piet is an employee, go over a separate exchange.

$$\begin{aligned} \text{cr} \rightarrow \text{ann} &: \text{piet}, \text{empl_status} \\ \text{ann} \rightarrow \text{cr} &: [\text{piet}, \text{is_empl}, \text{delegated_to}, \text{hr}]_{\text{ann}} \\ \text{cr} \rightarrow \text{hr} &: \text{piet}, \text{empl_status} \\ \text{hr} \rightarrow \text{cr} &: [\text{piet}, \text{is_empl}]_{\text{hr}} \end{aligned}$$

This specification falls short of capturing the internal reasonings of the participants involved, as described informally above. For instance, it is not clear how the participants must interpret the messages they receive, and how the *cr* ascertains Piet's right to store. These relations are often formalized in authorization logics, such as DKAL. However, specifying the internal reasoning in, e.g. SECPAL and DKAL, would fall short of determining the actual messages exchanged, or how the messages are related to policy statements.

In section V, we specify CRP in our formal language. Then, we can answer questions such as: whether the attacker can take the CRP system into a state in which, Eve, who is not an employee of the office, can write into the repository. Such questions cannot be meaningfully posed when considering the communication or the policy levels of CRP in isolation.

III. A LANGUAGE FOR SPECIFYING SERVICE ORIENTED ARCHITECTURES

The syntax and semantics of the language are defined below. A typical architecture specified in the language consists of communication events of services, how received messages are interpreted, how services make logical decisions, and how these decisions affect the communication events.

A. Syntax

A *signature* is a tuple $(\Sigma, \mathcal{V}, \mathcal{P})$, where Σ is a countable set of functions, $\mathcal{V}$ is a countable set of variables, $\mathcal{P}$ is a nonempty finite set of predicates, and these sets are pairwise disjoint. We use the capital letters $A, B, \dots$ to denote the elements of $\mathcal{V}$. The free term algebra induced by Σ , with variables $\mathcal{V}$, is denoted $\mathcal{T}_{\Sigma(\mathcal{V})}$. A *message* is an element of $\mathcal{T}_{\Sigma(\emptyset)}$, i.e. a ground term. The set of *atoms* $\mathcal{A}_{\Sigma(\mathcal{V})}$ is defined as $\{p(t_1, \dots, t_n) \mid p \in \mathcal{P}, t_i \in \mathcal{T}_{\Sigma(\mathcal{V})}, \text{arity of } p \text{ is } n\}$. The total function $\text{var} : \mathcal{T}_{\Sigma(\mathcal{V})} \cup \mathcal{A}_{\Sigma(\mathcal{V})} \rightarrow 2^{\mathcal{V}}$ gives the set of variables appearing in terms and atoms. A *fact* is an atom with no variables.

Fix a signature. An *event* is of either of the following forms:

- $g_1 \vee \dots \vee g_k \blacktriangleright s(t)$,
- $r(t) \triangleright u$

where $t \in \mathcal{T}_{\Sigma(\mathcal{V})}$, each g_i , with $1 \leq i \leq k$ and $k \geq 0$, is a finite set of atoms, and u is a finite set of atoms. Intuitively, an event of the form $g_1 \vee \dots \vee g_k \blacktriangleright s(t)$ denotes *guarded send*, where term t is sent to the network only if the *guard* $g_1 \vee \dots \vee g_k$ is evaluated to "true" (for the exact definition, see section III-B). An event of the form $r(t) \triangleright u$ denotes *receive* followed by *update*, where receiving term t results in adding the update u to the policy level. The function var extends to events as $\text{var}(g_1 \vee \dots \vee g_k \blacktriangleright s(t)) = \bigcup_{1 \leq i \leq k} \text{var}(g_i) \cup \text{var}(t)$ and $\text{var}(r(t) \triangleright u) = \text{var}(t) \cup \text{var}(u)$.

We write E for the set of all events, and E^* for the set of all finite sequences of events. A *process* is a finite sequence of events $\epsilon = e_1 \dots e_n$ where

- If $e_i = g_1 \vee \dots \vee g_k \blacktriangleright s(t)$, then $\text{var}(e_i) \subseteq \bigcup_{1 \leq j < i} \text{var}(e_j)$, for all $1 \leq i \leq n$.
- If $e_i = r(t) \triangleright u$, then $\text{var}(u) \subseteq \text{var}(t) \cup \bigcup_{1 \leq j < i} \text{var}(e_j)$, for all $1 \leq i \leq n$.

These conditions intuitively state that the behavior of any process depends deterministically on its input.

A *service* π is a tuple $(\eta_\pi, \Omega_\pi, \mathbf{I}_\pi)$, where η_π is a finite set of processes, Ω_π is a finite set of facts, called the *knowledge* of π , and $\mathbf{I}_\pi$ is a finite set of Horn clauses, called the *intensional knowledge* of π . A Horn clause is of the form $a \leftarrow a_1, \dots, a_n$, with $n \geq 0$, and $a, a_1, \dots, a_n$ being atoms.

An *attacker* model $\mathbb{A}$ is a service with no processes, i.e. it is a pair $(\Omega_{\mathbb{A}}, \mathbf{I}_{\mathbb{A}})$, where $\Omega_{\mathbb{A}}$ is a finite set of facts, called the knowledge of the attacker, and $\mathbf{I}_{\mathbb{A}}$ is a finite set of Horn clauses, referred to as the intensional knowledge of the attacker. The attacker is also able to send and receive messages; these capabilities are reflected in the execution model described in section III-B.

A (service-oriented) *architecture* is a tuple $((\Sigma, \mathcal{V}, \mathcal{P}), \Pi, \mathbb{A})$, where $(\Sigma, \mathcal{V}, \mathcal{P})$ is a signature, Π is a finite nonempty set of services and $\mathbb{A}$ is an attacker model, where Π and $\mathbb{A}$ are defined using the signature $(\Sigma, \mathcal{V}, \mathcal{P})$. In order to avoid trivial name clashes, it is assumed that variables that appear in different processes of an architecture are distinct.

We assume that for any architecture $((\Sigma, \mathcal{V}, \mathcal{P}), \Pi, \mathbb{A})$, the predicate $\mathcal{K}$, of arity 1, belongs to $\mathcal{P}$. This predicate is in particular used to model the knowledge of the attacker $\mathbb{A}$.

B. Semantics

Let s be a finite set of facts, and $\mathbf{I}$ be a finite set of Horn clauses. The *closure* of s under $\mathbf{I}$, denoted $[s]^\mathbf{I}$, is the smallest set that contains s and moreover $\forall(a \leftarrow a_1, \dots, a_n) \in \mathbf{I}. \forall \sigma. a_1\sigma, \dots, a_n\sigma \in [s]^\mathbf{I} \implies a\sigma \in [s]^\mathbf{I}$, where σ is a total (grounding) substitution function for the Horn clause $a \leftarrow a_1, \dots, a_n$; that is $\sigma : (\text{var}(a) \cup \bigcup_{1 \leq i \leq n} \text{var}(a_i)) \rightarrow \mathcal{T}_{\Sigma(\emptyset)}$. The existence of $[s]^\mathbf{I}$ follows immediately from the monotonicity of Horn theories [15]. The *ground deduction* problem for a finite set of Horn clauses $\mathbf{I}$, asks whether $a \in [s]^\mathbf{I}$, for an arbitrary fact a and a finite set of facts s .

Let $((\Sigma, \mathcal{V}, \mathcal{P}), \Pi, \mathbb{A})$ be an architecture, consisting of services $\Pi = \{1, \dots, \ell\}$, with $\pi = (\eta_\pi^0, \Omega_\pi^0, \mathbf{I}_\pi)$ for $\pi \in \Pi$, and $\mathbb{A} = (\Omega_\mathbb{A}^0, \mathbf{I}_\mathbb{A})$. A *configuration* of the architecture is a tuple $((\eta_1, \Omega_1), \dots, (\eta_\ell, \Omega_\ell), \Omega_\mathbb{A})$, where η_i is a finite set of processes, for $1 \leq i \leq \ell$, and $\Omega_\mathbb{A}$ and Ω_i are finite sets of facts. The *initial* configuration of the architecture is $z^0 = ((\eta_1^0, \Omega_1^0), \dots, (\eta_\ell^0, \Omega_\ell^0), \Omega_\mathbb{A}^0)$.

Each architecture $((\Sigma, \mathcal{V}, \mathcal{P}), \Pi, \mathcal{A})$ is attributed with a Kripke structure which represents all the executions of the architecture. This is explained informally in the following. Suppose a process belonging to service $\pi \in \Pi$ can perform a guarded send event $g_1 \vee \dots \vee g_k \blacktriangleright s(t)$. Then, the guard is evaluated against the policy engine of π . The guard is interpreted as the “disjunction” of the “conjunctions” of atoms in each g_i , with $1 \leq i \leq k$. If the guard can be derived in the policy engine (i.e. the guard evaluates to “true”), then the process sends the term t to the network; that is, t is immediately added to the knowledge of the attacker. We remark that variables appearing in a guarded send event originate in previous receive events in the process. Therefore, the variables in term t have already been instantiated with ground values. Therefore, only messages (i.e. terms with no variables) are sent to the network. Now, suppose a process belonging to service $\pi \in \Pi$ can perform a receive event $r(t) \triangleright u$. If there exists a grounding substitution σ such that $t\sigma$ can be derived from the attacker’s knowledge, then the process receives $t\sigma$ and the predicates $u\sigma$ are added to the knowledge set of service π . The formal definition is given below.

An architecture $((\Sigma, \mathcal{V}, \mathcal{P}), \Pi, \mathcal{A})$, with $\Pi = \{1, \dots, \ell\}$, induces a Kripke structure (S, S^0, T) , where S and T are the smallest sets satisfying

- $S^0 = z^0$, $z^0 \in S$.
- If $z = ((\eta_1, \Omega_1), \dots, (\eta_i, \Omega_i), \dots, (\eta_\ell, \Omega_\ell), \Omega_\mathbb{A}) \in S$, then
 - If $\mathbf{e} \in \eta_i$ with $\mathbf{e} = (g_1 \vee \dots \vee g_k \blacktriangleright s(t))\mathbf{e}'$, $\mathbf{e}' \in E^*$, and $\exists j. 1 \leq j \leq k \wedge g_j \subseteq [\Omega_i]^\mathbf{I_i}$, then $z' \in S$ and $(z, z') \in T$, with $z' = ((\eta_1, \Omega_1), \dots, (\eta'_i, \Omega_i), \dots, (\eta_\ell, \Omega_\ell), \Omega_\mathbb{A} \cup \{\mathcal{K}(t)\})$, and $\eta'_i = \eta_i \setminus \{\mathbf{e}\} \cup \{\mathbf{e}'\}$.
 - If $\mathbf{e} \in \eta_i$ with $\mathbf{e} = (r(t) \triangleright u)\mathbf{e}'$, $\mathbf{e}' \in E^*$, and there exists a substitution σ where $\mathcal{K}(t\sigma) \in [\Omega_\mathbb{A}]^\mathbf{I_\mathbb{A}}$, then $z' \in S$ and $(z, z') \in T$, where the configuration z' is defined as $z' = ((\eta_1, \Omega_1), \dots, (\eta'_i, \Omega_i \cup u\sigma), \dots, (\eta_\ell, \Omega_\ell), \Omega_\mathbb{A})$ and $\eta'_i = \eta_i \setminus \{\mathbf{e}\} \cup \{\mathbf{e}'\sigma\}$.

Given a configuration Z and Kripke structure (S, S^0, T) , we

say Z is *reachable* in (S, S^0, T) iff $(S^0, Z) \in T^*$, where T^* is the reflexive transitive closure of T .

C. The reachability decision problem

Given an architecture $\text{arch} = ((\Sigma, \mathcal{V}, \mathcal{P}), \Pi, \mathbb{A})$, with $\Pi = \{1, \dots, \ell\}$, and a fact $f \in \mathcal{A}_{\Sigma(\emptyset)}$, the *reachability* problem $\text{REACH}(\text{arch}, a, f)$, with $a \in \Pi \cup \{\mathbb{A}\}$, asks whether there exists a reachable configuration $((\eta_1, \Omega_1), \dots, (\eta_\ell, \Omega_\ell), \Omega_\mathbb{A})$ in the Kripke structure induced by the architecture such that $f \in [\Omega_a]^\mathbf{I_a}$, or not. These cases are respectively denoted by $\text{REACH}(\text{arch}, a, f) = \text{T}$ and $\text{REACH}(\text{arch}, a, f) = \text{F}$.

The decision problem REACH subsumes the secrecy problem for security protocols and the safety problem for authorization logics. The secrecy problem asks whether the attacker can learn a (supposedly secret) message m via interacting with other services. This can be represented as $\text{REACH}(\text{arch}, \mathbb{A}, \mathcal{K}(m))$. The safety problem asks whether an authorization predicate f (e.g. $\text{knows}(\text{can_read}(\text{ann}, \text{file12}))$) can be derived by a service, say π , with $\pi \in \Pi$. This corresponds to $\text{REACH}(\text{arch}, \pi, f)$.

Due to the computational power of Horn clauses, it is in general undecidable whether a guard evaluates to “true” or not. Therefore, REACH is in general undecidable (even though any architecture consists of finitely many services). However, if the ground deduction problem is decidable for the intensional knowledge sets of the participants and the attacker, then a semi-decision algorithm can be constructed for REACH using dovetailing. This is due to the fact that any architecture induces a finite number of interleavings of events (although an architecture in general induces an infinitely-branching Kripke structure); cf. *symbolic traces* defined in section VI-B.

IV. A DECIDABLE FRAGMENT

In this section, we identify a fragment of intensional knowledge (for the participants and the attacker) that admits decision algorithms for the reachability problem. In this decidable fragment, referred to as $\mathbf{A}_1$, the intensional knowledge of the attacker is fixed to the standard Dolev-Yao (DY) deduction rules, as formalized in, e.g., [7]. The policies (i.e. intensional knowledge) of honest services are centered around *trust application* TA , and (transitive) *trust delegation* TD rules, adopted and adapted from DKAL , and can also express typical RBAC systems with role hierarchy. Next, we introduce the notion of *infons*, cf. [6], [16].

Infons are pieces of information; for example, $\text{can_read}(\text{piet}, \text{file12})$ stipulating that Piet can read a certain *file12*. An infon does not admit a truth value, i.e. it is never false or true. Instead, infons are the interfaces between the communication level and policy level for honest services. That is, if the policy engine of a service, say Ann, derives the predicate $\text{knows}(\text{can_read}(\text{piet}, \text{file12}))$, then Ann “knows” that Piet may read this file, and may thus grant him read access to *file12*. Note that “knows” in this context, and also in DKAL , is a predicate symbol and not a modality as in logics of knowledge. In fact, “knows” here is closer to the notion of *belief* rather than *knowledge*, in epistemic terms.

Infons are different from predicates (i.e. policy statements) in that they can be *nested*, i.e. infons are constructed by applying *infor constructors* to message terms and other infons. We assume that in any signature $(\Sigma, \mathcal{V}, \mathcal{P})$, the set Σ can be partitioned into Σ_{msg} and Σ_{infor} , so that Σ_{infor} is the set of infor constructor functions, and Σ_{msg} is the set of message constructor functions. The set of infons $Infons$ is formally defined as the smallest set satisfying the property: if $t_1, \dots, t_n \in \mathcal{T}_{\Sigma_{msg}(\mathcal{V})} \cup Infons$ and $f \in \Sigma_{infor}$ with arity of f being n , then $f(t_1, \dots, t_n) \in Infons$. Note that $Infons \cap \mathcal{T}_{\Sigma_{msg}(\mathcal{V})} = \emptyset$; in particular messages are not infons.

We assume that for honest services the policy statements (i.e. the inhabitants of the policy level) are predicates over *Infons*. That is, for these services, the knowledge ranges over pieces of information. In contrast, the knowledge of the attacker ranges directly over message terms. Below, disjoint union of two sets is denoted by $\sqcup$.

Definition 1. *Fragment A_1 consists of all architectures $arch = ((\Sigma, \mathcal{V}, \mathcal{P}), \Pi, \mathbb{A})$, which satisfy the following syntactical conditions:*

- $(\Sigma, \mathcal{V}, \mathcal{P})$ is a signature, with $\Sigma = \Sigma_{msg} \sqcup \Sigma_{infor}$, and:
 - A finite subset of constants in Σ_{msg} , denoted by *Agents*, represents the set of the names of the services in Π .¹
 - Apart from nullary functions (i.e. constants), Σ_{msg} only contains the functions $\{\cdot\}$, $\{\cdot\}_\cdot$, $sig(\cdot, \cdot)$, $pk(\cdot)$, $h(\cdot)$, $(\cdot, \cdot)$. These represent respectively asymmetric and symmetric encryption, digital signature, public key constructor, hash and pairing functions, interpreted as usual.
 - Σ_{infor} contains in particular the functions $\theta(\cdot, \cdot)$ and $\sigma(\cdot, \cdot)$. These intuitively stand for trusted on and said, respectively, with $\theta, \sigma : Agents \times Infons \rightarrow Infons$.
 - $\mathcal{P} = \{\mathcal{K}\}$, with $\mathcal{K}$ being a unary predicate. Intuitively, $\mathcal{K}$ stands for “knows”.
- Any service $\pi = (\eta_\pi^0, \Omega_\pi^0, \mathbf{I}_\pi)$ in Π meets the conditions:
 - For all processes in η_π^0 , all the terms sent and received are elements of $\mathcal{T}_{\Sigma_{msg}(\mathcal{V})}$.
 - Ω_π^0 is a finite set of ground atoms of the form $\mathcal{K}(i)$, with $i \in Infons$.
 - $\mathbf{I}_\pi$ includes the *TA* and *TD* rules, respectively represented by $\mathcal{K}(X) \leftarrow \mathcal{K}(\theta(A, X)), \mathcal{K}(\sigma(A, X))$, and $\mathcal{K}(\theta(A, \theta(B, X))) \leftarrow \mathcal{K}(\theta(A, X))$. The set of all the other rules in $\mathbf{I}_\pi$ constitutes a type-1 theory, as defined below, and σ and θ do not appear in this set.
- $\mathbb{A} = (\Omega_{\mathbb{A}}^0, \mathbf{I}_{\mathbb{A}})$, where $\Omega_{\mathbb{A}}^0$ is a finite subset of $\{\mathcal{K}(t) \mid t \in \mathcal{T}_{\Sigma_{msg}(\emptyset)}\}$, and $\mathbf{I}_{\mathbb{A}}$ consists of the rules that reflect the capabilities of the standard *DY* attacker as formalized in

¹Intuitively, one (or more) public key is attributed to each element of *Agents*. The public keys are known to everyone, and to the attacker in particular. Using the public key of $a \in Agents$ one can encrypt messages for a and can verify the authenticity of the messages signed by a .

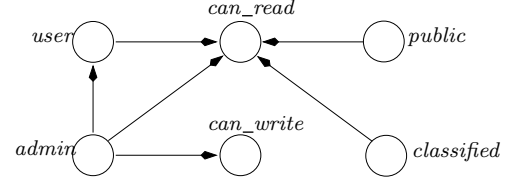


Fig. 1. Dependency graph, example 1.

appendix A or [17]; for instance, $\mathcal{K}(X) \leftarrow \mathcal{K}((X, Y))$, and $\mathcal{K}(\{X\}_Y) \leftarrow \mathcal{K}(X), \mathcal{K}(Y)$.

We define type-1 theories in order to extend the policies of (honest) services beyond *TA* and *TD*, e.g. to express typical RBAC systems with role hierarchy.

Definition 2. *A finite set of Horn clauses T , defined over signature $(\Sigma, \mathcal{V}, \mathcal{P})$, with $\Sigma = \Sigma_{msg} \sqcup \Sigma_{infor}$, is called a type-1 theory, iff*

- (a) All clauses in T have the form $p(t) \leftarrow p_1(t_1), \dots, p_\ell(t_\ell)$, where $p, p_1, \dots, p_\ell \in \mathcal{P}$, and $t, t_1, \dots, t_\ell \in Infons$.
- (b) For all $a \leftarrow a_1, \dots, a_\ell$ in T , $\bigcup_{i \in \{1, \dots, \ell\}} \text{var}(a_i) \subseteq \text{var}(a)$.
- (c) The infor dependency graph of T is acyclic. The infor dependency graph of T is a directed graph defined by the pair $(\Sigma_{infor}, \text{Edges})$ with $(f, g) \in \text{Edges}$ iff there exists a Horn clause in T using g in its head, and f in its body.

We remark that neither *TA* nor *TD* fall into type-1 theories, due to conditions (b) and (c) in definition 2, respectively.

Example 1. Consider a file server which implements an RBAC system with two roles, *user* and *admin*. Users may read any *public* file, admins may read any *classified* file, and admins may also write to any file. Admins inherit all the rights attributed to users. Below, we give a type-1 theory which describes this RBAC system.

$$\begin{aligned}
 \mathcal{K}(\text{user}(A)) &\leftarrow \mathcal{K}(\text{admin}(A)) \\
 \mathcal{K}(\text{can_read}(A, F)) &\leftarrow \mathcal{K}(\text{user}(A)), \mathcal{K}(\text{public}(F)) \\
 \mathcal{K}(\text{can_read}(A, F)) &\leftarrow \mathcal{K}(\text{admin}(A)), \mathcal{K}(\text{classified}(F)) \\
 \mathcal{K}(\text{can_write}(A, F)) &\leftarrow \mathcal{K}(\text{admin}(A))
 \end{aligned}$$

Here Σ_{msg} contains the set of identities of involved services, and names of files, while $\mathcal{P} = \{\mathcal{K}\}$, and $\Sigma_{infor} = \{\text{user}, \text{admin}, \text{public}, \text{classified}, \text{can_read}, \text{can_write}\}$ with obvious arities. The infor dependency graph for this theory, shown in figure 1, is indeed acyclic. •

The type-1 policy of example 1 is *centralized* in the sense that it describe the policies of a single entity, i.e. the file server. In the context of an architecture, this type-1 theory may represent the policies of a single service, cf. section V.

In section VI, we give a decision algorithm for the reachability problem for architectures in fragment A_1 . The decision algorithm is based upon encoding policy level computations of services into message derivation trees of the standard *DY* model. Let us now continue with a formal specification of *CRP* (section II) as an A_1 architecture.

V. A FORMAL SPECIFICATION OF CRP

We give a formal specification of CRP by defining the architecture $\text{crp} = ((\Sigma, \mathcal{V}, \mathcal{P}), \Pi, \mathbb{A})$. The architecture, as we will see, indeed falls into the $\mathbf{A}_1$ fragment. The signature Σ contains the standard cryptographic primitives $\{\cdot\cdot, \{\cdot\cdot\}, pk(\cdot), h(\cdot), sig(\cdot, \cdot), (\cdot, \cdot) \in \Sigma_{msg}$, and the infon constructors, namely $can_store(\cdot), empl(\cdot), head(\cdot), \sigma(\cdot, \cdot), \theta(\cdot, \cdot) \in \Sigma_{infn}$. The set of constants $\Sigma_C \subset \Sigma_{msg}$ used in crp is defined as $\Sigma_C = \{mike, piet, ann, cr, hr, eve, doc, empl_status, is_empl, delegated_to, success_token\}$ with eve being the identity of the attacker.

The set $\mathcal{P}$ contains only one unary predicate, $\mathcal{K}$, used for modelling the knowledge of the services and the attacker. Note that the knowledge of each of the services and the attacker are stored separately. In order to avoid unnecessary cluttering, we suppress the predicate symbol $\mathcal{K}$.

The attacker $\mathbb{A} = (\Omega_{\mathbb{A}}^0, \mathbf{I}_{\mathbb{A}})$ has the initial knowledge $\Omega_{\mathbb{A}}^0 = \Sigma_C \setminus \{doc\} \cup \{pk(A) \mid A \in \{mike, piet, ann, cr, hr, eve\}\}$, and her intensional knowledge reflects the usual DY message derivation rules. Despite the fact that the attacker's knowledge set is always finite, she can generate an infinite set of terms by pairing, hashing, signing and encrypting the terms that are known to her (cf. appendix A).

The intensional knowledge for $mike$, $piet$, and ann consists of the TA and TD rules only. The intensional knowledge of the hr contains the type-1 theory $\{empl(X) \leftarrow head(X)\}$, besides TA and TD . The intensional knowledge of the cr , in addition to TA and TD contains the type-1 theory $\{empl(X) \leftarrow head(X), can_store(X) \leftarrow empl(X)\}$, which constitute a simple hierarchical RBAC system, where employees have the right to store documents in the cr , and heads of the office inherit all rights of employees.

Below, we describe the processes executed by the participating services, and their initial knowledge. Recall that capital letters denote variables. To avoid name clashes, variables appearing in different processes should be tagged with process names. The tagging is however omitted in the following to ease the presentation.

Citizen (mike). Mike's initial knowledge is empty.

$$\begin{aligned} \emptyset &\triangleright s(\{mike, doc\}_{pk(piet)}, [h(doc)]_{mike}) \\ r([h(mike, doc), success_token]_{piet}) &\triangleright \emptyset \end{aligned}$$

Mike sends document doc to Piet, in order to be stored in the cr . He then waits to receive a "success" token from Piet. The names $mike$ and $piet$, and also doc are constants in Mike's specification.

Employee (piet). Piet's initial knowledge is empty.

$$\begin{aligned} r(\{C, D\}_{pk(piet)}, [h(D)]_C) &\triangleright \emptyset \\ \emptyset &\triangleright s(\{C, D\}_{pk(cr)}, [ann, h(C, D)]_{piet}) \\ r([h(C, D)]_{cr}) &\triangleright \emptyset \\ \emptyset &\triangleright s([h(C, D), success_token]_{piet}) \end{aligned}$$

After Piet receives a request, from a citizen C , to store a document D , he sends a corresponding request to the cr . Then he waits for confirmation (of successful storing) from

the cr , after which he notifies the citizen on the completed transaction.

Head (ann). Ann's initial knowledge is empty.

$$\begin{aligned} r(E, empl_status) &\triangleright \emptyset \\ \emptyset &\triangleright s([E, is_empl, delegated_to, hr]_{ann}) \end{aligned}$$

When Ann receives a request for information on the status of a principal E , she replies that the task of providing such information has been delegated to the hr . Ann would typically execute a few instances of this process in parallel.

Human Resources (hr). The initial knowledge of hr contains all employees and heads of the office. That is, the hr 's initial knowledge is $\{empl(piet), head(ann)\}$.

$$\begin{aligned} r(E, empl_status) &\triangleright \emptyset \\ \{empl(E)\} &\triangleright s([E, is_empl]_{hr}) \end{aligned}$$

After receiving a request for the status of a principal E , the hr confirms that E is an employee of the office by sending the message $[E, is_empl]_{hr}$, only if $empl(E)$ can be derived in the policy engine of the hr .

Central Repository (cr). The cr service consists of two processes, executed in parallel. The initial knowledge of the cr is $\{head(ann)\}$.

Central Repository's main process.

$$\begin{aligned} r(\{C, D\}_{pk(cr)}, [H, h(C, D)]_E) &\triangleright \{\theta(H, empl(E))\} \\ \{head(H)\} &\triangleright s(E, empl_status) \\ r([E, is_empl]_F) &\triangleright \{\sigma(F, empl(E))\} \\ \{can_store(E)\} &\triangleright s([h(C, D)]_{cr}) \end{aligned}$$

After receiving the first message from E , requesting to store a document D as an employee of the office headed by H , the cr asserts that H is trusted on whether E is an employee of the office, or not. Next, if H is indeed the head of the office, then the cr asks H for clarifying the employment status of E . If the cr , after receiving the third message, ascertains that E has the right to store, then the document is stored (not formalized here), and E is notified.

Central Repository's delegation handler.

$$\begin{aligned} r([E, is_empl, delegated_to, HR]_H) &\triangleright \\ &\{\sigma(H, \theta(HR, empl(E)))\} \end{aligned}$$

This process receives, independently of the other process, messages from an office head H to delegate to the HR the right to declare the employment status of E .

The formalization given above models the inquiries which the cr conducts via a "broadcast" send. Namely, when the cr sends the message $(E, empl_status)$, the message is intended for H (i.e. ann), but it is received also by F (i.e. the hr). This is because ann replies to the message with a delegation certificate, which is consumed by the delegation handler process of the cr . It is thus F (i.e. the hr) who actually responds to the message $(E, empl_status)$ in the cr 's main process.

The derivation tree of figure 2 shows how the cr ascertains Piet's right to write into the repository; the correspondence between derivation trees and computing closures under Horn

$$\begin{array}{c}
\frac{\theta(ann, empl(piet))}{\theta(ann, \theta(hr, empl(piet)))} \text{TD} \quad \frac{\sigma(ann, \theta(hr, empl(piet)))}{\theta(hr, empl(piet))} \text{TA} \quad \frac{\sigma(hr, empl(piet))}{\frac{empl(piet)}{can_store(piet)}} \text{type-1} \\
\hline
\end{array}$$

Fig. 2. A derivation tree for the CRP case study

clauses is immediate. The rules (e.g. *TA* or type-1 theory of the *cr*) used in each derivation step are also shown in the figure.

Remark 1. *The formalization can be further extended by adding the following process to the service Ann:*

$$\begin{array}{l}
r(E, empl_status) \triangleright \emptyset \\
\{empl(E)\} \blacktriangleright s([E, is_empl]_{ann})
\end{array}$$

We also add the fact $empl(piet)$ to the initial knowledge of Ann. In this extension Ann can decide whether to relegate the task to ascertain that Peter is an employee of the office to the Human Resources department or to carry out the task herself (or both). It is immediate that the main process of *cr* remains the same in both these scenarios: it need not be a priori “aware” of whether the delegation takes place or not.

VI. DECIDING REACHABILITY

In this section, we give an algorithm for deciding REACH in $\mathbf{A}_1$ architectures. We start with presenting an encoding from policy statements into (message) terms. Then, we extend the constraint solving algorithm of Millen and Shmatikov for deciding reachability in the “encoded” $\mathbf{A}_1$ architectures.

The purpose of the encoding is to replace the logic programs of the participating services in an $\mathbf{A}_1$ architecture with the derivation rules of the DY model. Then, intuitively, the attacker and all the services would be equipped with the reasoning power of the DY model, which is well understood and comes with decision algorithms for reachability.

A. Encoding policy level computations

Below, to simplify the presentation, we suppress the predicate symbol $\mathcal{K}$ from facts, and work directly with infons; indeed $\mathcal{K}$ is the only predicate symbol in $\mathbf{A}_1$ architectures. The encoding consists of two functions: ζ which maps infons to $\mathcal{T}_{\Sigma(V)}$, and $\mathcal{E}$ which maps infons to guards. We start with an initial encoding for trust application only. Then, we extend the encoding to trust delegation and type-1 theories.

TA only. We recursively define the encoding for infon i :

$$\zeta(i) = \begin{cases} \llbracket \zeta(X), sig(\bar{A}, \zeta(X)) \rrbracket_{\zeta(\theta(A, X))} & \text{if } i = \sigma(A, X) \\ \theta(\bar{A}, \zeta(X)) & \text{if } i = \theta(A, X) \\ i & \text{otherwise} \end{cases}$$

where $\bar{\cdot} : Agents \rightarrow \overline{Agents}$ is a bijection which associates a unique name to each element of $Agents$. Elements of $\overline{Agents}$ belong to $\mathcal{T}_{\Sigma_{msg}(\emptyset)}$, and are defined solely for the encoding function ζ , i.e. they do not appear in the specifications of

architectures. In particular, the attacker does not have the private keys associated to the members of $\overline{Agents}$.

Here, the encoding of the infon $\sigma(a, x)$ is the ciphertext $\llbracket x, sig(\bar{a}, x) \rrbracket_{\theta(\bar{a}, x)}$ from which the infon x (i.e. what service a said) can be obtained using the *symmetric decryption* (*Sdec*) rule of DY only if the key $\theta(\bar{a}, x)$ (i.e. a is trusted on x) is obtained first. This indicates that if *TA* is applicable on a set of facts at the policy level, then the DY rule

$$\mathcal{K}(X) \leftarrow \mathcal{K}(\llbracket X \rrbracket_K), \mathcal{K}(K) \quad Sdec$$

is applicable to the terms resulting from the encoding. Intuitively, the role of *sig* is to ensure that terms of the form $\llbracket x, sig(\bar{a}, x) \rrbracket_{\theta(\bar{a}, x)}$ can be constructed using the DY rules only if a corresponding $\sigma(a, x)$ can be derived in the policy level. Recall that the attacker does not know the private key of $\bar{a}$.

TA, TD and type-1 theories. In order to include *TD* and type-1 theories in the encoding, we define an *expansion* function $\mathcal{E}$ that for any atom returns a guard. We motivate the expansion function via a simple example. Suppose the fact $\theta(a, i)$ is present in the policy engine of a service, and the guard $\theta(a, \theta(b, i))$ is to be evaluated. The *TD* rule implies that the guard can be derived, while there is no corresponding inference tree (in the DY model) for $\zeta(\theta(a, \theta(b, i)))$, given $\zeta(\theta(a, i))$. The set of infons which yield $\theta(a, \theta(b, i))$ via applying only the *TD* rule is however finite. This finite set of infons can be seen as a guard, namely $\{\theta(a, \theta(b, i))\} \vee \{\theta(a, i)\}$. The fact that $\theta(a, i)$ yields $\theta(a, \theta(b, i))$ in the policy engine is reflected in the DY model by: either $\zeta(\theta(a, \theta(b, i)))$ or $\zeta(\theta(a, i))$, or both, are obtained from $\zeta(\theta(a, i))$.

The expansion function $\mathcal{E}_P(Q, i)$ is defined for finite sets of Horn clauses P and Q , and infon i :

$$\begin{aligned}
\mathcal{E}_P(\emptyset, i) &= \{i\} \\
\mathcal{E}_P(\{r \leftarrow r_1, \dots, r_\ell\} \sqcup Q', i) &= \\
&\begin{cases} \mathcal{E}_P(Q', i) \vee \\ (\mathcal{E}_P(P, r_1\rho) \cup \dots \cup \mathcal{E}_P(P, r_\ell\rho)) & \text{if } i = r\rho \\ \mathcal{E}_P(Q', i) & \text{if } \neg\exists\rho. i = r\rho \end{cases}
\end{aligned}$$

where $\cup$ distributes over $\vee$, i.e. $S \cup (S_1 \vee S_2) = (S_1 \vee S_2) \cup S = (S_1 \cup S) \vee (S_2 \cup S)$. Here Q is a support theory used only to ensure that the expansion of any infon results in a finite set; see lemma 1 below.

Lemma 1. *Let $P = P^1 \cup \{TD\}$, with P^1 being the type-1 theory in a service of an $\mathbf{A}_1$ architecture. Then, $\mathcal{E}_P(P, i)$ is a finite set for any infon i .*

Proof: Immediate, since the dependency graph of P^1 is acyclic, P^1 does not contain θ , and the Horn clause which encodes TD strictly decreases the number of θ functions. ■

We write $\mathcal{E}(i)$ for $\mathcal{E}_P(P, i)$, when P is clear from the context. We write $g \sqsubseteq \mathcal{E}(i)$ if $\mathcal{E}(i) = g_1 \vee \dots \vee g \vee \dots \vee g_n$, with $n \geq 1$.

Example 2. Consider the infon $i = \text{can_read}(a, \text{file})$ along with the type-1 theory of example 1. Then,

$$\mathcal{E}(i) = \{ \text{user}(a), \text{public}(\text{file}) \} \vee \{ \text{admin}(a), \text{public}(\text{file}) \} \vee \{ \text{admin}(a), \text{classified}(\text{file}) \} \vee \{ \text{can_read}(a, \text{file}) \}$$

Write $\mathcal{E}(i) = g_1 \vee g_2 \vee g_3 \vee g_4$, with $g_1 = \{ \text{user}(a), \text{public}(\text{file}) \}$, etc. The guard $\mathcal{E}(i)$ is interpreted as: $\text{can_read}(a, \text{file})$ holds, i.e. a can read file in example 1, iff at least one of the following conditions holds: $[g_1]$ $\text{user}(a)$ and $\text{public}(\text{file})$ are known, or $[g_2]$ $\text{admin}(a)$ and $\text{public}(\text{file})$ are known, or $[g_3]$ $\text{admin}(a)$ and $\text{classified}(\text{file})$ are known, or $[g_4]$ $\text{can_read}(a, \text{file})$ is known via an inference outside the RBAC system of example 1. Similarly, we get $\mathcal{E}(\text{can_write}(a, \text{file})) = \{ \text{admin}(a) \} \vee \{ \text{can_write}(a, \text{file}) \}$. •

We refine the function ζ (introduced above) by incorporating the expansion function $\mathcal{E}$ into ζ . This intuitively ensures that $\zeta(i)$, for infon i , is obtainable from $\zeta(\sigma(a, i))$ if there exist at least one $g \sqsubseteq \mathcal{E}(\theta(a, i))$ such that $\zeta(g)$ can be obtained first. Hence, we define:

$$\zeta(i) = \begin{cases} \llbracket \zeta(X), \text{sig}(\bar{A}, \zeta(X)) \rrbracket_{\zeta(\mathcal{E}_P(P, \theta(A, X)))} & \text{if } i = \sigma(A, X) \\ \theta(\bar{A}, \zeta(X)) & \text{if } i = \theta(A, X) \\ i & \text{otherwise} \end{cases}$$

Here, $\llbracket x \rrbracket_{k_1 \vee \dots \vee k_\ell}$ stands for the tuple $\llbracket x \rrbracket_{k_1}, \dots, \llbracket x \rrbracket_{k_\ell}$, function ζ distributes over $\vee$, and $P = P^1 \cup \{TD\}$ with P^1 being the type-1 theory at hand. For a finite set of infons g , $\zeta(g)$ is defined as the concatenation of $\zeta(i)$, for all $i \in g$. We remark that elements of $\mathcal{E}_P(P, \theta(A, X))$ in the definition of ζ are singletons. This is because in any $\mathbf{A}_1$ architecture, $\mathcal{E}_P(P, \theta(a, i)) = \mathcal{E}_{\{TD\}}(\{TD\}, \theta(a, i))$, as θ does not appear in P^1 .

Correctness of the encoding. We remark that the purpose of the proposed encoding is to replace the logic programs of services with the derivation rules of the Dolev-Yao model.

The following theorem ensures that if a policy fact is derivable in the logic program of a service, then its corresponding encoded term can be derived using the DY inference rules, and vice versa. We consider the standard DY capabilities for the term algebra $\mathcal{T}_{\Sigma(\mathcal{V})}$, which comprises both infon and message constructors. That is, the infon constructors are seen as uninterpreted functions, while message constructors (e.g. $\llbracket \cdot \rrbracket$) have their standard meaning in the DY model.

In the rest of the paper, $T \vdash u$ denotes that u is derivable from a set of terms T with the standard DY inference rules as formalized, e.g., in appendix A and [7].

Theorem 1. Let P be the intensional knowledge of a service in an $\mathbf{A}_1$ architecture, with $P = \{TA\} \cup Q$, $Q = \{TD\} \cup P^1$

and P^1 being a type-1 theory. For any (ground) infon f and finite set of (ground) infons G ,

$$\mathcal{K}(f) \in [\mathcal{K}(G)]^P \iff \exists g \sqsubseteq \mathcal{E}_Q(Q, f). \quad \zeta(G) \vdash \zeta(g),$$

where $\mathcal{K}(G)$ stands for the set $\{\mathcal{K}(f_i) \mid f_i \in G\}$.

Proof: Fix the policy set P . We write $G \Vdash f$ for $\mathcal{K}(f) \in [\mathcal{K}(G)]^P$, suppress $\mathcal{K}$ when confusion is unlikely, and write $\mathcal{E}(f)$ for $\mathcal{E}_Q(Q, f)$. Below, we talk about *proof trees* for f , given G . The correspondence between finding proof trees and computing closures is immediate. The proof is split into two directions.

$\Rightarrow$ We use structural induction on proof trees for f , given G . If $f \in G$, then the implication is trivial. Otherwise, consider the last rule applied in the proof tree:

- (TA) Then $G \Vdash \sigma(a, f), \theta(a, f)$, for some $a \in \text{Agents}$. By induction hypotheses,

$$\exists s \sqsubseteq \mathcal{E}(\sigma(a, f)), t \sqsubseteq \mathcal{E}(\theta(a, f)). \quad \zeta(G) \vdash \zeta(s), \zeta(t).$$

Observe that $s = \sigma(a, f)$. The term $\zeta(\sigma(a, f))$ is the tuple $\llbracket \zeta(f), \text{sig}(\bar{a}, \zeta(f)) \rrbracket_{\mathcal{E}(\theta(a, f))}$. Since $t \sqsubseteq \mathcal{E}(\theta(a, f))$, through unpairing, we obtain the ciphertext $\llbracket \zeta(f), \text{sig}(\bar{a}, \zeta(f)) \rrbracket_{\zeta(t)}$ from $\zeta(\sigma(a, f))$. From $\zeta(G) \vdash \zeta(t)$, by applying the *Sdec* rule and unpairing we get $\zeta(G) \vdash \zeta(f)$. Clearly $f \sqsubseteq \mathcal{E}(f)$.

- (TD) Then $f = \theta(a, \theta(b, i))$ for some $a, b \in \text{Agents}$ and $i \in \text{Infons}$, and $G \Vdash \theta(a, i)$. By induction hypotheses, $\exists t \sqsubseteq \mathcal{E}(\theta(a, i))$. $\zeta(G) \vdash \zeta(t)$. Now, the claim follows since for any infon t , $t \sqsubseteq \mathcal{E}(\theta(a, i))$ implies $t \sqsubseteq \mathcal{E}(\theta(a, \theta(b, i)))$.
- (Type-1) Let $R = r \leftarrow r_1, \dots, r_\ell \in P^1$ be the last rule applied. Then $f = r\rho$ and $G \Vdash r_1\rho, \dots, r_\ell\rho$, for some grounding substitution ρ (cf. condition (b) in definition 2). By induction hypotheses, $\exists r'_1 \sqsubseteq \mathcal{E}(r_1\rho), \dots, r'_\ell \sqsubseteq \mathcal{E}(r_\ell\rho)$. $\zeta(G) \vdash \zeta(r'_1), \dots, \zeta(r'_\ell)$. By definition of $\mathcal{E}$, $\{r'_1, \dots, r'_\ell\} \sqsubseteq \mathcal{E}(f)$, hence follows the claim.

$\Leftarrow$ First, we claim that $\zeta(G) \vdash \zeta(g)$ implies $G \Vdash g$. Notice that the $\zeta(g)$ is either of the form $\llbracket \zeta(x), \text{sig}(\bar{a}, \zeta(x)) \rrbracket_{\mathcal{E}(\theta(a, x))}$, or of the form $i(x)$, with i being an infon constructor. The claim follows by case analysis on the DY attacker's message (de)composition abilities. In particular, note that (1) to fabricate $\llbracket \zeta(x), \text{sig}(\bar{a}, \zeta(x)) \rrbracket_{\mathcal{E}(\theta(a, x))}$, the attacker needs to construct $\text{sig}(\bar{a}, \zeta(x))$, which is impossible as the attacker does not own the private key for any $\bar{a} \in \overline{\text{Agents}}$, and (2) infon constructors are uninterpreted functions in the DY model, i.e. they can neither be applied by the attacker, nor their application can be deconstructed. The other cases are straightforward; we thus omit them here. Finally,

notice that if $G \Vdash g$ and $g \in_v \mathcal{E}(f)$, then $G \Vdash f$; hence follows the claim.

This completes our proof. $\blacksquare$

B. A constraint solving algorithm for deciding reachability

We begin with a brief description of Millen and Shmatikov's constraint solving algorithm for deciding reachability in cryptographic protocols [7]. Recall that participants are specified as sequences of communication (i.e. send and receive) events in [7], and the reachability problem, given a message s , asks whether there exists a reachable configuration z of the protocol where $T(z) \vdash s$, with $T(z)$ being the attacker's knowledge in z .

The algorithm of [7] searches the finite set of interleavings of (symbolic) actions performed by the participants and a fictitious *test* process, which receives s and then sends stop to the search algorithm. For each interleaving, a sequence C of *attacker constraints* is constructed. An attacker constraint is a pair $\langle m : T \rangle$, where m is a term that the attacker should derive from the set of terms T , using her inference capabilities. The constraint sequence is built for each interleaving as: when a participant sends a message term, the term is added to the attacker term set, and when a receive action occurs, a constraint $\langle m : T \rangle$ is enqueued to C , with m being the term that is to be received and T is the current attacker term set.

A solution for constraint sequence $C = \langle m_1 : T_1 \rangle \cdots \langle m_i : T_i \rangle \cdots \langle m_n : T_n \rangle$, with $m_i = s$, is a (grounding) substitution $\sigma : \text{var}(\langle m_1 : T_1 \rangle \cdots \langle m_i : T_i \rangle) \rightarrow \mathcal{T}_{\Sigma(\emptyset)}$ such that $T_j \sigma \vdash m_j \sigma$, for $1 \leq j \leq i$; here, $\text{var}(c_1 \cdots c_i)$ is the set of variables appearing in the constraints $c_1, \dots, c_i$. In our presentation, therefore, we account for *partial* executions as well, cf. [18]. Millen and Shmatikov's algorithm applies a number of reduction rules which reduce C to a sequence of immediately (un)satisfiable constraints. In the following, we refer to their reduction procedure as MSReduce. If MSReduce does not succeed for C (i.e. C is unsatisfiable), next interleaving is considered, until all the interleavings are exhausted. If one of the constraint sequences is satisfiable, then the (supposedly) secret message s is revealed to the attacker. That is, an attack is found. Otherwise the protocol is correct, i.e. s is not revealed to the attacker, for the instantiation at hand.

The following two observations enable us to use Millen and Shmatikov's procedure (with minor extensions) for deciding reachability for $\mathbf{A}_1$ architectures:

- 1) Checking guards in, and making updates to, policy engines of the services can be emulated by communication actions. Namely, sending an infon to the knowledge set of a service reflects updating the policy engine of that service, while receiving an infon derived from the knowledge set reflects querying the policy engine of the service for evaluating a guard.
- 2) The encoding presented in section VI-A entails that the same inference rules which are used for the attacker (namely the standard DY message derivation capabilities) can model the computations of the participants at their policy level. Therefore, the send and receive actions

which would emulate checking guards and updating knowledge sets of services (mentioned above) can be treated with the constraint solving procedure that one would use for the attacker knowledge (here, MSReduce).

Algorithm 1 Constraint solving for deciding reachability in $\mathbf{A}_1$ fragment

```

REQUIRES: REACH(arch, a, f), arch = ((Σ, V, P), Π, A)
expand(Π)
I := interleavings(Π, a, f)
for all  $\iota \in I$  do
  flatten( $\iota$ )
  C := constraint_sequence( $\iota$ )
  trace := MSReduce⊞(C)
  if trace ≠ ∅ then
    return (reach : trace)
return (unreach)

```

Algorithm 1 details our constraint solving procedure for deciding reachability in $\mathbf{A}_1$ architectures. The input to the algorithm is a reachability problem $\text{REACH}(\text{arch}, a, f)$, with $\text{arch} = ((\Sigma, \mathcal{V}, \mathcal{P}), \Pi, \mathbb{A})$ being an $\mathbf{A}_1$ architecture. The algorithm either returns a *trace* witnessing $\text{REACH}(\text{arch}, a, f) = \text{T}$, or returns *unreach* if $\text{REACH}(\text{arch}, a, f) = \text{F}$.

In algorithm 1, *expand*(Π) expands the guards in each process, for all the services in Π . A guard $g_1 \vee \cdots \vee g_n$ in a service with policy P , is expanded to $\bigvee_{1 \leq j \leq n} \mathcal{E}_Q(Q, g_j)$, where $Q = P \setminus \{TA\}$ and $\mathcal{E}_Q(Q, g)$ for a finite set of infons g is defined as the union of $\mathcal{E}_Q(Q, i)$, for all $i \in g$. Recall that $\cup$ distributes over $\vee$ (cf. section VI-A). The *expand* procedure thus rewrites guards into guards. For example, the guard $\{a, b\}$ in a service with intensional knowledge $\{TA, TD, a \leftarrow c, b \leftarrow d\}$ is expanded to $\{a, b\} \vee \{c, b\} \vee \{a, d\} \vee \{c, d\}$.

The procedure *interleavings*(Π, a, f) computes the finite set of interleavings of events of the participants in Π . Furthermore, this procedure (1) adds a *testing* process whose sole purpose is to indicate that the search has reached a configuration in which $f \in [\Omega(a)]^{\mathbb{A}_a}$, given $\text{REACH}(\text{arch}, a, f)$; that is, if $a = \mathbb{A}$ and $f = \mathcal{K}(m)$, then the testing process simply receives m and then sends stop to the search algorithm. If $a \neq \mathbb{A}$, then the testing process $\{f\} \blacktriangleright \text{s(stop)}$ is added to service a . (2) The *interleavings* procedure treats the disjunction operator $\vee$ inside guards as branching points, e.g. the interleaving $e_1 \cdot (g \vee g' \blacktriangleright \text{s}(m)) \cdot e_2$, with $e_1, e_2 \in E^*$, gives rise to two interleavings $e_1 \cdot (g \blacktriangleright \text{s}(m)) \cdot e_2$ and $e_1 \cdot (g' \blacktriangleright \text{s}(m)) \cdot e_2$. Consequently, no $\vee$ appears in guards for any interleaving $\iota \in I$ in algorithm 1.

For each interleaving, the procedure *flatten* intuitively “flattens” the two levels of specification into one. That is, a sequence of events (i.e. guarded sends, and receives coupled with updates) is translated into a sequence of *annotated* send and receive actions. The annotations indicate the knowledge set with which the communication is carried out: $\mathbb{A}$ for network communications through the attacker, and $\pi \in \Pi$ for

each service π .

$(\{a_1, \dots, a_\ell\} \blacktriangleright s(m))$ maps to $r^\pi(\zeta(a_1), \dots, \zeta(a_\ell)) \cdot s^\mathbb{A}(m)$
 $(r(m) \triangleright \{a_1, \dots, a_\ell\})$ maps to $r^\mathbb{A}(m) \cdot s^\pi(\zeta(a_1), \dots, \zeta(a_\ell))$

The annotated send and receive actions are used in the reduction procedure $\text{MSReduce}^\boxtimes$, which is a barely syntactic modification of MSReduce . The modification consists in allocating one term set for each service in Π , and one set for the attacker. This is in contrast to MSReduce where only one term set, denoting the attacker knowledge, is considered. In $\text{MSReduce}^\boxtimes$, annotated sends and receives communicate with the term set that is determined by their annotation.

After *flattening*, the procedure *constraint_sequence* generates a constraint sequence for the interleaving (as it is done in [7]). The resulting constraint sequence is fed to the procedure $\text{MSReduce}^\boxtimes$.

Example 3. Consider a file server service fs , with intensional knowledge $P = \{TA, TD\} \cup P^1$, where P^1 is the type-1 theory of example 1. A typical event of the main process of the fs service is $\{can_read(A, F)\} \blacktriangleright s(\{F\}_{pk(A)})$, where A and F denote, respectively, a client of the service and a file stored on fs . For this event the *expand* procedure returns $\{user(A), public(F)\} \vee \{admin(A), public(F)\} \vee \{admin(A), classified(F)\} \vee \{can_read(A, F)\} \blacktriangleright s(\{F\}_{pk(A)})$; see example 2.

The *interleavings* procedure then creates a branch for each $g \in_v$ the resulting guard, while interleaving this event with other events of the architecture. For each of the branches the *flatten* procedure maps the events into communication actions. For example, the guarded send $\{user(A), public(F)\} \blacktriangleright s(\{F\}_{pk(A)})$ maps to $r^{fs}(user(A), public(F)) \cdot s^\mathbb{A}(\{F\}_{pk(A)})$. •

The following theorem states that algorithm 1 is terminating on decision problems for $\mathbf{A}_1$ architectures; moreover the algorithm is correct (i.e. sound and complete) w.r.t. the semantics of $\mathbf{A}_1$ architectures.

Theorem 2. *Given decision problem $\text{REACH}(\text{arch}, a, f)$, with arch being an $\mathbf{A}_1$ architecture, algorithm 1 terminates, and returns a trace iff $\text{REACH}(\text{arch}, a, f) = \text{T}$.*

The proof of theorem 2 is relegated to appendix A. Below, we sketch the main idea of the proof. Let $\mathbf{KS} = (S, S^0, T)$ be the Kripke structure induced by arch. A *trace* in $\mathbf{KS}$ is a sequence $z_0 e_1 z_1 \dots e_n z_n$, where $z_0 = S^0$, $(z_{j-1}, z_j) \in T$, for $1 \leq j \leq n$, and the system evolves from z_{j-1} into z_j when event e_j occurs, as defined in section III-B. A *symbolic trace* st is a trace which may contain non-ground terms. We say st is *realizable* in $\mathbf{KS}$ iff there exists a grounding substitution σ , such that $\text{st}\sigma$ is a trace in $\mathbf{KS}$. Given a symbolic trace st and a sequence of constraints C we define their *correspondence* inductively: if $\text{st} = z_0$, then the empty sequence, i.e. $C = \text{nil}$, corresponds to st . Now, let $\text{st} = z_0 e_1 \dots z_n e z$, with $n \geq 0$, and C' be a sequence of constraints that corresponds to $z_0 \dots z_n$. If $e = g \blacktriangleright s(m)$, then any $C = C' \cdot \langle \zeta(g_i) : \zeta(\Omega_\pi^{\otimes z_n}) \rangle$ corresponds to st ,

where π is the service that performs e , $g_i \in_v \mathcal{E}(g)$ is calculated w.r.t. the intensional knowledge of π , and $\Omega_\pi^{\otimes z_n}$ refers to the knowledge of π at symbolic configuration z_n . The correspondence between event e and constraint $\langle \zeta(g_i) : \zeta(\Omega_\pi^{\otimes z_n}) \rangle$ hinges upon theorem 1, which tells us that the encoding function ζ is such that $\zeta(G) \vdash \zeta(g_i)$ for at least one $g_i \in_v \mathcal{E}(g)$ if $\mathcal{K}(g) \in [\mathcal{K}(G)]^P$, given any $G \subseteq \mathcal{A}_{\Sigma(\emptyset)}$. If $e = r(m) \triangleright u$, then $C = C' \cdot \langle m : \Omega_\mathbb{A}^{\otimes z_n} \rangle$ corresponds to st , where $\Omega_\mathbb{A}^{\otimes z_n}$ refers to the attacker knowledge at symbolic configuration z_n . We remark that for the attacker knowledge the (suppressed) predicate $\mathcal{K}$ ranges directly over messages.

From the definition above, clearly there are finitely many constraint sequences (created due to the $\vee$ operator in guards) corresponding to any symbolic trace st . The proof of theorem 2 intuitively goes by showing that for each symbolic trace st , created by an interleaving of events of services in arch, st is realizable in $\mathbf{KS}$ iff a constraint sequence corresponding to st is satisfiable in algorithm 1. In the following we give an informal explanation of why the constraint reduction system of Millen and Shmatikov can be applied to the constraint sequences generated from architectures in $\mathbf{A}_1$. See the appendix for a formal proof.

Any constraint sequence C generated from architectures in $\mathbf{A}_1$ can be partitioned into two subsets C_A and C_P . These subsets respectively refer to *attacker constraints* and *policy constraints*, and these are syntactically distinguishable. Intuitively, for each symbolic trace, constraints in C_A correspond to messages that must be generated by the attacker so that the symbolic trace is realizable, and constraints in C_P correspond to guards that need to be evaluated to true in honest services so that the symbolic trace is realizable.

The Millen-Shmatikov reduction system is readily applicable to the elements of the set C_A , intuitively because these are attacker constraints. For the set C_P , however, we notice that for any constraint of the form $\langle q : K \rangle$ in C_P , with V being a variable appearing in K , there exists a constraint $\langle m : T \rangle$ in C_A , where V is a subterm of m . This is because all variables in our specifications are originally instantiated at a receive event in an honest process. The key idea in using Millen-Shmatikov's reduction procedure for architectures in $\mathbf{A}_1$ is that all the variables appearing in policy constraints are wrapped with infon constructors (which are seen as uninterpreted function constructors by the reduction procedure). Hence, the only applicable reduction rule on a policy constraint $\langle q : K \rangle$ is the unification rule (see the appendix) as far as q is an infon. Therefore, the policy constraints are ultimately either removed from C as they can be solved using unification, or they are unsatisfiable as the infon constructors are not available to the attacker. Recall that due to our expansion procedure (which is a “backward proof search”) the set of premises from which infons can be inferred have already been closed under the policy rules, before the constraint reduction procedure starts.

We proceed with the CRP case study.

C. Verification results

We have implemented algorithm 1 in Prolog. The implementation extends the constraint solver developed by Millen and Shmatikov [7]. The tool and the formal specifications of two case studies (including CRP) are available on-line at <http://www.infsec.ethz.ch/people/fraus>. Below, we briefly describe our verification results.

1) *CRP*: We have verified the following properties of the CRP case study (specified in section V):

- *Executability*. We replaced the last event in the specification of the citizen with

$$r([h(\text{mike}, \text{doc}), \text{success_token}]_{\text{piet}}) \triangleright \{\text{stored}(\text{doc})\}$$

and checked if $\text{REACH}(\text{crp}, \text{mike}, \text{stored}(\text{doc})) = \text{T}$; that is, whether *mike* knows that his document has been successfully stored in the *cr*. The tool returns a trace showing that this property indeed holds in the CRP architecture.

- *Secrecy*. We have checked if $\text{REACH}(\text{crp}, \text{eve}, \text{doc}) = \text{T}$, i.e. whether the attacker *eve* can discover *doc*. No (attack) trace is found.
- *Safety*. We have checked whether the attacker *eve* can obtain the right to store in the *cr*. That is, if $\text{REACH}(\text{crp}, \text{cr}, \text{can_store}(\text{eve})) = \text{T}$. No (attack) trace is found.

It is worth mentioning that in $\mathbf{A}_1$ architectures, e.g. the formalization of CRP in section V, only a finite number of services are allowed. Our decidability result would in fact fall apart if an unbounded number of services were considered in $\mathbf{A}_1$ architectures. This immediately follows from the undecidability of reachability in security protocols with an unbounded number of sessions.

2) *A flawed variant of CRP*: Minor changes in the CRP scenario (explained in section V) may lead to misinterpretation of the received messages, and ultimately to wrong policy decisions. For instance, if the document *D* is omitted from the signed message in the main process of *cr* (and correspondingly omitted from the main process of *Piet*) as:

$$\begin{aligned} r(\{C, D\}_{pk(\text{cr})}, [H, h(C)]_E) &\triangleright \{\theta(H, \text{empl}(E))\} \\ \{head(H)\} &\blacktriangleright s(E, \text{empl_status}) \\ r([E, \text{is_empl}]_F) &\triangleright \{\sigma(F, \text{empl}(E))\} \\ \{can_store(E)\} &\blacktriangleright s([h(C, D)]_{cr}) \end{aligned}$$

Then the attacker *eve* can pass a fake document to be stored at *cr*, although she does not have (nor gain) the right to store documents at the repository. This attack has indeed been found using our tool.

3) *An electronic health record system*: We have formalized a simple scenario of the electronic health records system described in [10]. The formalization is available on-line (see above). In this health records system, there are three “services” involved: the doctor *d* who is treating a patient *p*, the specialist doctor *s* who is meant to further check the health status of *p*, and the hospital’s data center *dc*. The attacker *eve* is also present in the system. The health records of the patients are

stored in the hospital’s data center, and they are accessible to their treating doctors. Any other access to health records is prohibited, except when a treating doctor delegates the right to access the records to another doctor or specialist. We consider the following scenario: The doctor *d*, who is treating patient *p*, asks the specialist *s* whether *s* would examine the health status of *p*. We assume *s* is willing to do so. Then, *d* delegates to *s* the right to access *p*’s health record. Meanwhile *s* requests *p*’s health record from *dc*. The data center runs two threads: one for handling delegation requests, and one for managing access to the health records. The data center grants access to the health records of *p* iff at least one of the following conditions hold: the requester is the doctor who is treating *p*, or the doctor who is treating *p* has indeed delegated the right to read the health record associated to *p* to the requester. These rules are encoded in the policy engine of *cr* as:

$$\begin{aligned} can_read_ehr(X, P) &\leftarrow delegated(X, P) \\ can_read_ehr(X, P) &\leftarrow doctor(X, P) \end{aligned}$$

Here $can_read_ehr(X, P)$ states that *X* can read the health record of patient *P*, $doctor(X, P)$ means that *X* is the treating doctor of *P*, while $delegated(X, P)$ stands for the fact that the right to read the record of *P* has been delegated to *X*. We remark that the right to access patient’s records cannot be further delegated.

We have verified three properties for this scenario: (1) *Executability*. We have checked if the specialist can indeed access the health record of the patient. That is, the delegation mechanism operates as expected. The tool returns a trace showing that this property does hold. (2) *Non-transitive delegation*. We have checked that the right to read the patient’s health record cannot be further delegated. That is, a specialist cannot delegate this right to another entity. No attack traces are found for this property. (3) *Secrecy*. We have checked if the attacker can learn the health record of the patient. No attack traces are found in this case.

VII. CONCLUSION

We have presented a language for formal specification of service-oriented architectures. The language allows us to specify communication level events, policy level decisions, and the interface between the two. We have shown that the reachability problem is decidable for a fragment of architectures specified in the language. The decidable fragment is of immediate practical relevance. Deciding reachability in service-oriented architectures is an NP-hard problem, because it subsumes the reachability problem for security protocols (which is an NP-complete problem [8]). The decision procedure for the fragment of architectures given in this paper uses an encoding which in worst-case is exponential in the size of the input. The efficiency of the decision algorithm has so far not been a major concern in our study. For the case studies we report in this paper, the tool can find attacks in less than a minute on a typical machine, while for showing that no attacks exists (hence exploring the state space exhaustively) the tool can take up to a day, depending on the specification. We intend to

investigate more efficient (in asymptotic terms) decision algorithms for this fragment. In particular, partial order reduction techniques to reduce the number of interleavings would be effective in this respect.

In this decidable fragment, the policies of the (honest) services are limited to trust application and trust delegation rules, besides a finite set of Horn clauses of *type-1*. Type-1 Horn theories are characterized here by placing certain syntactic conditions on the form of the clauses. As we show in the paper, these conditions are indeed sufficient; they are however in general not necessary for deciding reachability. We intend to investigate how, without undermining our decidability result, the policies of honest services can be extended to theories beyond type-1. A more ambitious goal is to formally characterize the set of policies for which the reachability problem is decidable.

Revocation of rights is not expressible in a natural way in authorization logics which are based on Horn theories, cf. [19] for an overview. We are currently working towards accommodating rights revocation in our language.

We (similar to, e.g., SECPAL and DKAL) see “know” as a predicate, as opposed to a modality as it is common in epistemic logics. This causes a discrepancy between the syntactic form of “knowledge”, represented by predicates, and the semantic notion of knowledge. It must therefore be interesting to investigate how our proposed formalism can be connected to the standard epistemic models.

Acknowledgements: We are grateful to D. Basin, S. Burri, S. Ranise and E. Zalescu for their comments on this paper. The work has been supported by the EU FP7 projects AVANTSSAR (no. 216471) and SPACIOS (no. 257876).

REFERENCES

- [1] R. Parikh and R. Ramanujam, “A knowledge based semantics of messages,” *Journal of Logic, Language and Information*, vol. 12, no. 4, 2003.
- [2] H. Ditmarsch, W. Hoek, and B. Kooi, *Dynamic epistemic logic*, 1st ed. Springer, 2007.
- [3] J. Guttman, F. Thayer, J. Carlson, J. Herzog, J. Ramsdell, and B. Sniffen, “Trust management in strand spaces,” in *ESOP '04*, ser. LNCS, vol. 2986, 2004, pp. 325–339.
- [4] J. DeTreville, “Binder, a logic-based security language,” in *IEEE Symposium on Security and Privacy '02*, 2002, p. 105.
- [5] M. Becker, C. Fournet, and A. Gordon, “Design and semantics of a decentralized authorization language,” in *CSF '07*. IEEE Computer Society, 2007, pp. 3–15.
- [6] Y. Gurevich and I. Neeman, “DKAL: Distributed-knowledge authorization language,” in *CSF '08*. IEEE Computer Society, 2008, pp. 149–162.
- [7] J. Millen and V. Shmatikov, “Constraint solving for bounded-process cryptographic protocol analysis,” in *CCS '01*. ACM Press, 2001, pp. 166–175.
- [8] M. Rusinowitch and M. Turuani, “Protocol insecurity with finite number of sessions is NP-complete,” in *CSFW '01*. IEEE CS, 2001, p. 174.
- [9] D. Dolev and A. Yao, “On the security of public key protocols,” *IEEE Trans. on Information Theory*, vol. IT-29, no. 2, pp. 198–208, 1983.
- [10] AVANTSSAR, “Deliverable D5.1: Problem cases and their trust and security requirements,” 2008, available at <http://www.avantssar.eu>.
- [11] M. Harrison, W. Ruzzo, and J. Ullman, “Protection in operating systems,” *Commun. ACM*, vol. 19, no. 8, pp. 461–471, 1976.
- [12] A. Armando and S. Ponta, “Model checking of security-sensitive business processes,” in *FAST '09*, ser. LNCS, vol. 5983. Springer, 2010, pp. 66–80.
- [13] A. Schaad, V. Lotz, and K. Sohr, “A model-checking approach to analysing organisational controls in a loan origination process,” in *SACMAT '06*, 2006, pp. 139–149.
- [14] M. Barletta, S. Ranise, and L. Viganò, “Verifying the interplay of authorization policies and workflow in service-oriented architectures,” in *CSE (3)*. IEEE Computer Society, 2009, pp. 289–296.
- [15] M. van Emden and R. Kowalski, “The semantics of predicate logic as a programming language,” *J. ACM*, vol. 23, no. 4, pp. 733–742, 1976.
- [16] Y. Gurevich and I. Neeman, “The logic of infons,” *Bulletin of the EATCS*, vol. 98, pp. 150–178, 2009.
- [17] B. Blanchet, “An efficient cryptographic protocol verifier based on prolog rules,” in *CSFW '01*. IEEE Computer Society, 2001, pp. 82–96.
- [18] R. Corin and S. Etalle, “An improved constraint-based system for the verification of security protocols,” in *SAS '02*, ser. LNCS, vol. 2477. Springer, 2002, pp. 326–341.
- [19] J. Halpern and V. Weissman, “Using first-order logic to reason about policies,” *ACM Trans. Inf. Syst. Secur.*, vol. 11, no. 4, 2008.

APPENDIX

Below, we list the capabilities of the Dolev-Yao attacker model. The following rules describe the intensional knowledge (i.e. the deduction capabilities) of the attacker, cf. [17]. The attacker is called *eve* here.

<i>Analysis</i>		
x	$\leftarrow (x, y)$	ϕ_{proj^1}
y	$\leftarrow (x, y)$	ϕ_{proj^2}
x	$\leftarrow \{x\}_{pk(eve)}$	ϕ_{pdec}
x	$\leftarrow \{x\}_y, y$	ϕ_{sdec}
<i>Synthesis</i>		
(x, y)	$\leftarrow x, y$	ϕ_{pair}
$\{x\}_y$	$\leftarrow x, y$	ϕ_{penc}
$\{x\}_y$	$\leftarrow x, y$	ϕ_{senc}
$h(x)$	$\leftarrow x$	ϕ_{hash}
$sig(pk(eve), x)$	$\leftarrow x$	ϕ_{sig}

Below, we present the proof of theorem 2. The proof relies upon a lemma on correctness and termination of $MSReduce^{\boxtimes}$.

Proof of theorem 2: Termination of algorithm 1 is immediate, as the procedure *interleavings* produces a finite number of (symbolic) interleavings, and all the functions applied to interleavings, in particular $MSReduce^{\boxtimes}$ (see lemma 2, below) and *expand* (due to lemma 1), terminate in finitely many steps.

Notice that algorithm 1 exhaustively considers all possible symbolic traces of *arch*, and for each symbolic trace $\mathfrak{st}$, the algorithm considers all constraint sequences corresponding to $\mathfrak{st}$. We write $\hat{\mathfrak{st}}$ for the set of all constraint sequences which corresponds to $\mathfrak{st}$.

The proof is split into two directions.

$\Rightarrow$ Assume $REACH\langle arch, a, f \rangle = T$. Then, there exists a symbolic trace $\mathfrak{st} = z_0 e_1 \cdots e_n z_n$, and a nonempty set of substitutions $\mathfrak{S} = \{\sigma_1, \sigma_2, \dots\}$ where $\mathfrak{st}\sigma_j$ is a trace in **KS**, for any $\sigma_j \in \mathfrak{S}$. Then, there exists at least one constraint sequence $C \in \hat{\mathfrak{st}}$, which is satisfiable by any substitution $\sigma_j \in \mathfrak{S}$, and that is considered by algorithm 1. It now remains to show that $MSReduce^{\boxtimes}$ returns a witness trace for C . This holds due to lemma 2, below.

$\Leftarrow$ Let C be a constraint sequence generated in algorithm 1. Then, $C \in \hat{\mathfrak{st}}$, for some symbolic trace $\mathfrak{st}$.

Suppose $\text{MSReduce}^\boxtimes$ returns a witness trace, showing that C is satisfiable under (non-ground) substitution ρ (see lemma 2 on applicability of $\text{MSReduce}^\boxtimes$, below). By theorem 1 and the fact that solving constraint $\langle m:T \rangle$ implies $T \vdash m$, it follows that $\text{st}\sigma$ is a trace in **KS**, for any ground substitution σ that refines ρ . That is, $\text{REACH}(\text{arch}, a, f) = T$.

This completes the proof. $\blacksquare$

Lemma 2, below, concerns the applicability of $\text{MSReduce}^\boxtimes$. Given a reduction system for constraint sequences, we write $C \rightarrow C'$ iff C is reduced, using a rule in the reduction system, to C' . A reduction system is

- *terminating* iff it admits no infinite reductions;
- *sound* iff $C \rightarrow C'$ implies that for any σ that is a solution of C' , σ is also a solution of C ;
- *complete* iff for any σ that is a solution of C , there exists a reduction step $C \rightarrow C'$ such that σ is a solution of C' .

Lemma 2. $\text{MSReduce}^\boxtimes$ is terminating, sound and complete for the constraint sequences generated in algorithm 1.

We prove lemma 2 by revising the proofs of termination and correctness (i.e. soundness and completeness) of MSReduce , originally presented in [7]. We start by giving an overview of MSReduce in appendix A. There, we also single out the differences between MSReduce and $\text{MSReduce}^\boxtimes$. Then, the proofs of termination and correctness of $\text{MSReduce}^\boxtimes$ are given, respectively, in appendices B and C.

A. Overview of MSReduce

Extended attacker model. The attacker model considered in [7] is based on the standard Dolev-Yao attacker model, as formalized in appendix A. In addition, Millen and Shmatikov consider an *encryption hiding* operator, here denoted $\| \cdot \|$, that serves as a technical means to avoid non-termination of the constraint solving algorithm. Consequently, the following attacker capabilities are also considered:

$$\begin{array}{lll} \text{Encryption hiding} & & \\ \|x\|_y & \leftarrow & \|x\|_y \quad \phi_{\text{open}} \\ \|x\|_y & \leftarrow & \{x\}_y \quad \phi_{\text{hide}} \end{array}$$

Given a set of messages T we denote by $\mathcal{F}(T)$ the set of messages that the attacker can derive from T using her capabilities. In terms of the formalization introduced in the paper, $\mathcal{F}(T) = [T]^{\mathbf{IA}}$.

Reduction procedure. In the following we give a short description of the MSReduce algorithm (see algorithm 2). The algorithm takes an initial constraint sequence IC as input, and builds the tree of all possible reductions for IC . Each node is labelled by a pair (C, σ) whose first element is a constraint sequence, and the second element is a substitution for the variables in C . Starting from the root $(IC, \emptyset)$, the tree is explored with a depth first search. Notice that the algorithm 2 uses a stack structure, rather than a tree structure, as it is usual in the depth first search. The exploration terminates successfully (i.e. node (C, σ) corresponds to an attack) when C is *simple*; that is, every constraint in C is of the form $\langle V:T \rangle$

where V is a variable. Simple constraint sequences (under the monotonicity and origination assumptions, see below) are immediately solvable.

Algorithm 2 MSReduce

REQUIRES: initial constraint sequence IC

```

stack :=  $\emptyset$ 
stack.push( $(IC, \emptyset)$ )
repeat
   $(C, \sigma) := \text{stack.pop}$ 
  let  $c = \langle m:t \rangle$  be the first constraint in  $C$ 
  s.t.  $m$  is not a variable
  if  $c$  not found then
    return (satisfiable :  $(C, \sigma)$ )
  apply rule (elim) to  $c$  until no longer applicable
  for all  $r \in R$  do
    if  $r$  is applicable to  $C$  then
      stack.push( $r(C, \sigma)$ )
until stack.empty
return (unsatisfiable)

```

Initially, the root of the tree is set to the pair $(IC, \emptyset)$ (i.e. $(IC, \emptyset)$ is the only element of the stack). Then, the tree is explored as follows. A node (C, σ) is popped from the stack, and the *active* constraint (that is, the first constraint $c = \langle m:T \rangle \in C$ with a non-variable m) is picked. If no such constraint exists in C , then C is simple hence the node is returned as a solution for the constraint sequence. If c is found, then all occurrences of stand-alone variables are removed from the term set T . Subsequently, for any applicable reduction rule r (see paragraph Reduction rules, below), the new node resulting from the application of r to the current node is pushed on top of the stack. When no nodes are left in the stack, all possible reductions have been considered, so the algorithm returns “unsatisfiable” and terminates.

Reduction rules. In the following we present the reduction rules used by MSReduce . The rules read from top to bottom, that is, a rule r

$$\frac{(C_{<} \cdot c \cdot C_{>}, \sigma)}{(C'_{<} \cdot c' \cdot C'_{>}, \sigma')} r$$

says that constraint sequence $C_{<} \cdot c \cdot C_{>}$, where c is the active constraint and $C_{<}$ and $C_{>}$ are respectively the constraints preceding and following c , is reduced to $C'_{<} \cdot c' \cdot C'_{>}$, and substitution σ is refined by σ' .

$$\frac{(C_{<} \cdot \langle m:T \sqcup \{V\} \rangle \cdot C_{>}, \sigma)}{(C_{<} \cdot \langle m:T \rangle \cdot C_{>}, \sigma)} \text{ elim}$$

where V is a variable

$$\frac{(C_{<} \cdot \langle m:T \rangle \cdot C_{>}, \sigma)}{(C_{<}\tau \cdot C_{>}\tau, \sigma \cup \tau)} \text{ un}$$

where $\tau = \text{mgu}(m, t)$, for some $t \in T$

$$\frac{(C_{<} \cdot \langle (m_1, m_2):T \rangle \cdot C_{>}, \sigma)}{(C_{<} \cdot \langle m_1:T \rangle \cdot \langle m_2:T \rangle \cdot C_{>}, \sigma)} \text{ pair}$$

$$\begin{array}{c}
\frac{(C_{<} \cdot \langle h(m):T \rangle \cdot C_{>}, \sigma)}{(C_{<} \cdot \langle m:T \rangle \cdot C_{>}, \sigma)} \text{hash} \\
\frac{(C_{<} \cdot \langle \{m\}_k:T \rangle \cdot C_{>}, \sigma)}{(C_{<} \cdot \langle k:T \rangle \cdot \langle m:T \rangle \cdot C_{>}, \sigma)} \text{penc} \\
\frac{(C_{<} \cdot \langle \{m\}_k:T \rangle \cdot C_{>}, \sigma)}{(C_{<} \cdot \langle k:T \rangle \cdot \langle m:T \rangle \cdot C_{>}, \sigma)} \text{senc} \\
\frac{(C_{<} \cdot \langle \text{sig}(\text{pk}(\text{eve})):T \rangle \cdot C_{>}, \sigma)}{(C_{<} \cdot \langle m:T \rangle \cdot C_{>}, \sigma)} \text{sig} \\
\frac{(C_{<} \cdot \langle m:T \cup \{(t_1, t_2)\} \rangle \cdot C_{>}, \sigma)}{(C_{<} \cdot \langle m:T \cup \{t_1, t_2\} \rangle \cdot C_{>}, \sigma)} \text{split} \\
\frac{(C_{<} \cdot \langle m:T \cup \{\{t\}_{\text{pk}(\text{eve})}\} \rangle \cdot C_{>}, \sigma)}{(C_{<} \cdot \langle m:T \cup \{t\} \rangle \cdot C_{>}, \sigma)} \text{pdec} \\
\frac{(C_{<} \cdot \langle m:T \cup \{\{t\}_k\} \rangle \cdot C_{>}, \sigma)}{(C_{<} \tau \cdot \langle m\tau:T\tau \cup \{\{t\}_k\} \rangle \cdot C_{>}\tau, \sigma\tau)} \text{ksub} \\
\text{where } \tau = \text{mgu}(k, \text{pk}(\text{eve})), k \neq \text{pk}(\text{eve}) \\
\frac{(C_{<} \cdot \langle m:T \cup \{\{t\}_k\} \rangle \cdot C_{>}, \sigma)}{(C_{<} \cdot \langle k:T \rangle \cdot \langle m:T \cup \{k, t\} \rangle \cdot C_{>}, \sigma)} \text{pdec}
\end{array}$$

Properties of MSReduce. For applying MSReduce to a sequence of constraints C , it is required that [7]:

- (Monotonicity) If constraint $\langle m:T \rangle$ precedes $\langle m':T' \rangle$ in C , then $\mathcal{F}(T) \subseteq \mathcal{F}(T')$.
- (Origination) For any constraint of the form $\langle m:T \rangle$ in C , with V being a variable appearing in T , there exists a preceding constraint $\langle m':T' \rangle$ in C , where V is a subterm of m' and $\mathcal{F}(T') \subset \mathcal{F}(T)$.

MSReduce preserves the aforementioned properties at each reduction step. Notice that the monotonicity and the origination property are necessary to grant solvability of simple constraint sequences.

Overview of MSReduce[⊞]. MSReduce and MSReduce[⊞] differ only in the properties of the constraint sequences that are passed to them as input.

Let C be a constraint sequence generated by algorithm 1. C can always be partitioned into two disjoint sequences of constraints C_A and C_P , that we call respectively “attacker constraints” and “policy constraints”.

Formally, an attacker constraint $\langle m:T \rangle$ models a receive event at the communication level, where m is a term that an attacker should construct from a finite set of terms T . Here, both m and the terms in T are message terms, i.e. $m \in \mathcal{T}_{\Sigma(V)}$ and $T \subseteq \mathcal{T}_{\Sigma(V)}$.

A policy constraint models the evaluation of a query at the policy level. Differently from an attacker constraint, for a policy constraint $\langle q:K \rangle$ the term q and the terms in K are the result of the application of function ζ to *Infons*. In particular, all terms in policy constraints contain constructors belonging to Σ_{Infons} (cf. sections IV and VI-A), while terms in attacker constraints never contain infon constructors. It is hence possible to distinguish between them with a simple syntactical check.

We note that C_A satisfies the monotonicity and origination properties. This is immediate from the syntax and semantics of our language; see section III. The origination and monotonicity properties do not hold for C_P , in general. However, for any constraint of the form $\langle q:K \rangle$ in C_P , with V being a variable appearing in K , there exists a preceding constraint $\langle m':T' \rangle$ in C_A , where V is a subterm of m' . This is because all variables in specifications of service-oriented architectures are originally instantiated at a receive event in an honest process.

We remark that the monotonicity and the origination properties hold for simple constraint sequences obtained by applying MSReduce[⊞] to constraint sequences generated by algorithm 1. The resulting simple constraint sequences are thus immediately solvable. This is because:

- policy constraints, for which the monotonicity property does not hold, are eliminated by the reduction procedure. This is due to the fact that the variables in policy constraints are always subterms of applications of infon constructors. Infon constructors are uninterpreted functions for the attacker’s inference rules; hence they are never constructed or deconstructed by the reduction system. If such constraints are satisfiable, then they will be eliminated using the unification rule *un*.
- the reduction procedure preserves the monotonicity of the attacker constraints, as shown in [7].

B. Proof of termination

The proof of termination of MSReduce[⊞] follows closely the termination proof of MSReduce. The proof is based on a termination measure (N_v, N_s) of a constraint sequence C . Here, N_v and N_s are naturals. In particular, N_v is the number of distinct variables occurring in C and N_s is a special *expansion* measure. Tuples are ordered lexicographically.

The expansion measure hinges upon another measure, the size $|m|$ of a term m , that is the number of operator applications plus the number of constants and variable in m . Then the expansion measure N_s of constraint sequence C is the sum of the expansion measures of its constraints; in turn, the expansion measure of a constraint $\langle m:T \rangle$ is $|m| \cdot \chi(T)$, where χ is defined as follows:

$$\begin{aligned}
\chi(t) &= 2 \quad \text{if } t \text{ is a variable or constant} \\
\chi(\{t_1, \dots, t_n\}) &= \chi(t_1) \cdots \chi(t_n) \\
\chi((t_1, t_2)) &= \chi(t_1)\chi(t_2) + 1 \\
\chi(\{t\}_k) &= \chi(t) \\
\chi(\|t\|_k) &= 1 \\
\chi(\text{sig}(k, t)) &= \chi(t) + 1 \\
\chi(h(t)) &= \chi(t) + 1 \\
\chi(\{t\}_k) &= \chi(t)\chi(k) + |k| + 1
\end{aligned}$$

We show now that the termination measure decreases strictly at each application of a reduction rule. Rule *elim* removes a stand-alone variable, hence reduces N_s ; rule *un* either substitutes a variable, hence decreasing N_v , or decreases N_s by removing the constraint; rules *sig*, *pair*, *hash*, *penc* and *senc* reduce N_s by splitting the constraint in constraints whose sum of expansion measures is smaller; rules *split* and *pdec*

decrease N_s by replacing a term with terms whose product of expansion measures is smaller; rule *ksub* substitutes a variable, hence decreases N_v . Finally rule *sdec* replaces a constraint $c = \langle m : T \cup \{t\}_k \rangle$ with constraints $c' = \langle k : T \cup \{t\}_k \rangle$ and $c'' = \langle m : T \cup \{t, k\} \rangle$. The expansion measure of c is $|m|\chi(T)(\chi(t)\chi(k) + |k| + 1)$, while the product of the expansion measures of c' and c'' is $|k|\chi(T) + |m|\chi(T)\chi(t)\chi(k)$. Since $|k| < |m|(|k| + 1)$, the expansion is strictly decreased in each *sdec* reduction step. $\text{MSReduce}^\boxtimes$ thus terminates.

C. Proof of correctness

In the following, for ease of presentation, we ignore the encryption hiding operator and related attacker capabilities. The encryption hiding operator is used in [7] merely to avoid non-termination, as mentioned above.

1) *Proof of soundness*: We show here that if a rule of $\text{MSReduce}^\boxtimes$ reduces constraint sequence C to constraint sequence C' , then any solution σ of C' is also a solution of C . In other words, $\text{MSReduce}^\boxtimes$ does not introduce new solutions. We condition on the rule applied:

- Rule *elim* removes a stand-alone variable V from $T \cup \{V\}$, for the active constraint being $c = \langle m : T \cup \{V\} \rangle$. We need to distinguish two cases:
 - c is an attacker constraint. We show that $\mathcal{F}(T \cup V) = \mathcal{F}(T)$. It is obvious that $\mathcal{F}(T) \subseteq \mathcal{F}(T \cup \{V\})$, we show then that $\mathcal{F}(T \cup \{V\}) \subseteq \mathcal{F}(T)$. By origination property, there exists an earlier constraint $c' = \langle m' : T' \rangle$, such that V does not appear in T' and V appears in m' . In particular, $m' = V$ (all constraints earlier than c are simple). Due to idempotency of closure, it suffices to show that $T \cup \{V\} \subseteq \mathcal{F}(T)$, and we only need to show that $V \subseteq \mathcal{F}(T)$, since $T \subseteq \mathcal{F}(T)$. By the monotonicity property $T' \subseteq T \cup \{V\}$, and since $V \notin T'$ then $T' \subseteq T$. But $V \in \mathcal{F}(T')$, therefore also $V \in \mathcal{F}(T)$.
 - c is a policy constraint. Then rule *elim* is never applied, because all variables in T appear under the application of an infon constructor. Furthermore, infon constructors are uninterpreted functions, hence can not be deconstructed to yield the variables they contain.
- Rules *split* and *pdec* are sound since $\mathcal{F}$ is closed under ϕ_{pair} and ϕ_{penc} .
- Rule *un* removes $c = \langle m : T \rangle$ when m is unifiable with some term $t \in T$. Let $\tau = \text{mgu}(m, t)$. If σ is a solution for C' , then $\sigma\tau$ is a solution for C provided that $T\sigma\tau \vdash m\sigma\tau$. This is obvious since $T\tau \vdash m\tau$.
- Rules *pair*, *hash*, *penc*, *senc* and *sig* are sound since $\mathcal{F}$ is closed under the corresponding ϕ rules of the attacker.
- Rule *sdec* replaces the active constraint $c = \langle m : T \rangle$, with $\{t\}_k \in T$, with the constraints $\langle k : T \rangle$ and $\langle m : T \cup \{k, t\} \rangle$. This rule is sound because if $k \in \mathcal{F}(T)$, then $\mathcal{F}(T) = \mathcal{F}(T \cup \{k\})$; moreover, since $\mathcal{F}$ is closed under ϕ_{sdec} , we have $\mathcal{F}(T) = \mathcal{F}(T \cup \{k, t\})$ given $k \in \mathcal{F}(T)$.

2) *Proof of completeness*: We show here that for every constraint sequence C and solution σ of C , there exists a rule r that reduces C to C' , and σ is a solution of C' . In other words, all solutions for C are preserved in at least one reduction path.

Let $\langle m : T \rangle$ be the active constraint in C . The proof relies on the existence of a *normal* proof of $T\sigma \vdash m\sigma$, that is, a proof tree such that no label appears more than once in any path from the root to a leaf [7]:

Proposition 1. *Let t be a ground term and T a set of ground terms. If $t \notin T$ and $t \in \mathcal{F}(T)$ then there exists a normal sequence $\phi_1, \dots, \phi_n$ such that $t \in \phi_n(\dots \phi_1(T))$ and one of the following conditions holds:*

- ϕ_n is a synthesis rule
- ϕ_1 is an analysis rule
- ϕ_i , for some $1 \leq i \leq n$, is ϕ_{sdec} and $\phi_1, \dots, \phi_{i-1}$ are synthesis rules

The proposition intuitively states that any normal sequence $\phi_1, \dots, \phi_n$ such that $t \in \phi_n(\dots \phi_1(T))$ can be reordered so that analysis rules always appear earlier than synthesis rules, except in the case where ϕ_{sdec} is used (i.e. synthesis rules appear before ϕ_{sdec} to construct a non-atomic key).

Finding an applicable rule. Let $c = \langle m : T \rangle$ be the active constraint in C and σ a solution of C . Then $m\sigma \in \mathcal{F}(T\sigma)$. Intuitively, $\text{MSReduce}^\boxtimes$ tries to apply a sequence of reduction rules that reflects the order of a normal proof of $T\sigma \vdash m\sigma$ as indicated by proposition 1.

Recall that by definition c does not contain stand-alone variables neither on the left hand side (it is chosen as the first constraint whose left hand side is not a variable) nor on the right hand side (as all stand-alone variables are removed by application of the *elim* rule). It follows that every term in $\{m\} \cup T$ has a well-defined *top level structure*, i.e. outermost function application, against which applicability of a rule can be checked.

If $m\sigma \in T\sigma$, then rule *un* is applicable. If $m\sigma \notin T\sigma$ then we can assume a normal sequence of operators $\phi_1, \dots, \phi_n$ as described in proposition 1. If ϕ_1 is an analysis rule then there is a term $t \in T$ with corresponding top level structure, hence the corresponding analysis rule can be applied. Similarly for ϕ_n being a synthesis rule, since m has a well-defined top structure the corresponding synthesis rule can be applied. Finally, if ϕ_i is ϕ_{sdec} then there is a term $\{x\}_y \in T$ that enables rule *sdec*. *Preserving the solution.* Proving that each applicable rule preserves the solution proceeds by cases, on ϕ_1 if the rule is an analysis rule, or on ϕ_n if the rule is a synthesis rule. For brevity, we omit the details and explain only the proof for the case of *sdec*, i.e. in which ϕ_i is ϕ_{sdec} and $\phi_1 \dots \phi_{i-1}$ are all synthesis rules. Rule *sdec* replaces the active constraint $c = \langle m : T \rangle$ with constraints $c' = \langle k : T \rangle$ and $c'' = \langle m : T \cup \{k, t\} \rangle$, for term $\{t\}_k \in T$. Observe that $k\sigma \in \phi_{i-1}(\dots \phi_1(T\sigma))$, otherwise rule *sdec* would not be applicable, and consequently σ is also a solution for c' . Also, since $k\sigma \in \mathcal{F}(T\sigma)$ and $\mathcal{F}$ is idempotent, then $\mathcal{F}(T\sigma) = \mathcal{F}(T\sigma \cup \{k\sigma, t\sigma\})$, which shows that σ is a solution for c'' .