



An interpolation-based method for the verification of security protocols

Marco Rocchetto, Luca Viganò, Marco Volpe

► To cite this version:

Marco Rocchetto, Luca Viganò, Marco Volpe. An interpolation-based method for the verification of security protocols. 2015. hal-01245442

HAL Id: hal-01245442

<https://hal.science/hal-01245442>

Preprint submitted on 4 Jul 2016

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

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

Copyright

An interpolation-based method for the verification of security protocols *

Marco Rocchetto¹, Luca Viganò², and Marco Volpe³

¹Dipartimento di Informatica, Università di Verona, Italy

²Department of Informatics, King's College London, UK

³INRIA and LIX, École Polytechnique, France

July 4, 2016

Abstract

Interpolation has been successfully applied in formal methods for model checking and test-case generation for sequential programs. Security protocols, however, exhibit such idiosyncrasies that make them unsuitable to the direct application of interpolation. We address this problem and present an interpolation-based method for security protocol verification. Our method starts from a protocol specification and combines Craig interpolation, symbolic execution and the standard Dolev-Yao intruder model to search for possible attacks on the protocol. Interpolants are generated as a response to search failure in order to prune possible useless traces and speed up the exploration. We illustrate our method by means of concrete examples and discuss the results obtained by using a prototype implementation.

1 Introduction

A number of tools (e.g., [2, 4, 8, 9, 10, 15, 17, 21, 30] just to name a few) have been developed for the analysis of security protocols at *design time*: starting from a formal specification of a protocol and of a security property it should achieve, these tools typically carry out model checking or automated reasoning to either *falsify* the protocol (i.e., find an attack with respect to that property) or, when possible, *verify* it (i.e., prove that it does indeed guarantee that property, perhaps under some assumptions such as a bounded number of interleaved protocol sessions [34]). While verification is, of course, the optimal result, falsification is also extremely useful as one can often employ the discovered attack trace to directly carry out an attack on the protocol implementation

*Work partially supported by the FP7-ICT-2009-5 Project no. 257876, “SPaCioS: Secure Provision and Consumption in the Internet of Services” and the PRIN 2010-11 project “Security Horizons”. Much of this work was carried out while Luca Viganò and Marco Volpe were at the Università di Verona, Italy. We thank Giacomo Dalle Vedove, Marco Palamà and Fabio Pettenuzzo.

(e.g., [3]) or exploit the trace to devise a suite of test cases so as to be able to analyze the implementation at *run-time* (e.g., [5, 6, 11]).

Such an endeavor has already been undertaken in the programming languages community, where, for instance, *interpolation* has been successfully applied in formal methods for model checking and test-case generation for sequential programs, e.g., [23, 25], with the aim of reducing the dimensions of the search space. Since a state space explosion often occurs in security protocol verification, we expect interpolation to be useful also in this context. Security protocols, however, exhibit such idiosyncrasies that make them unsuitable to the direct application of the standard interpolation-based methods, most notably, the fact that, in the presence of a Dolev-Yao intruder [16], a security protocol is not a sequential program (since the intruder, who is in complete control of the network, can freely interleave his actions with the normal protocol execution).

In this paper, we address this problem and present an interpolation-based method for security protocol verification. Our method starts from the formal specification of a protocol and of a security property and combines Craig interpolation [13], symbolic execution [20] and the standard Dolev-Yao intruder model [16] to search for goals (representing attacks on the protocol). Interpolation is used to prune possible useless traces and speed up the exploration. More specifically, our method proceeds as follows: starting from a specification of the input system, including protocol, property to be checked and a finite number of session instances (possibly generated automatically by using a preprocessor), it first creates a corresponding sequential non-deterministic program, according to a procedure that we have devised, and then defines a set of goals and searches for them by symbolically executing the program. When a goal is reached, an attack trace can be extracted from the constraints that the execution of the path has produced; such constraints represent conditions over parameters that allow one to reconstruct the attack trace found. When the search fails to reach a goal, a backtrack phase starts, during which the nodes of the graph are annotated (according to an adaptation of the algorithm defined in [25] for sequential programs) with formulas obtained by using Craig interpolation. Such formulas express conditions over the program variables, which, when implied from the program state of a given execution, ensure that no goal will be reached by going forward and thus that we can discard the current branch. The output of the method is a proof of (bounded) correctness in the case when no goal location can be reached starting from a finite-state specification; otherwise all the discovered (one or more) attack traces are produced.

In order to show that our method concretely speeds up the validation, we have implemented a Java prototype called *SPiM* (*Security Protocol interpolation Method*). We report here also on some experiments that we have performed: we considered 7 case studies and compared the analysis of SPiM with and without interpolation, thereby showing that interpolation does indeed speed up security protocol verification by reducing the search space and the execution time. We also compare the SPiM tool with the three state-of-the-art model checkers for security protocols that are part of the AVANTSSAR platform [4], namely, CL-AtSe [36], OFMC [8] and SATMC [1]. This comparison shows, as we expected, that SPiM is not yet as efficient as these mature tools but that there is considerable room for improvement, e.g., by enhancing our interpolation-based method with some of the optimization techniques that are inte-

grated in the other tools.

We proceed as follows. In Section 2, we provide some (fairly standard) background on security protocol verification, discussing the algebra of protocol messages, the Dolev-Yao intruder, the two security protocol specification languages ASLan++ and ASLan that we consider in our method (which is however open to the integration with other protocol specification languages), and the running example (the NSL protocol) that we will consider in the rest of the paper. In Section 3, we introduce SiL, the input language of our SPiM tool, which is a simple imperative programming language that we use to define the sequential programs to be analyzed by the verification algorithm. We also give the details of the translation procedure from security protocols into sequential programs, for one and more protocol sessions, and prove the correctness of the translation (i.e., that it does not introduce nor delete attacks with respect to the input ASLan++ specification). In Section 4, we present our interpolation algorithm, which is a slightly simplified version of McMillan’s IntraLA algorithm [25], and show it at work for our running example. In Section 5, we introduce the SPiM tool and discuss the experiments we have performed. In Section 6, we discuss further related work (in addition to the works already considered in the rest of the paper), and we conclude in Section 7 by summarizing our main results and discussing future work. Additional details (examples and a proof of one of the lemmas) are given in appendix.

This paper extends and supersedes [33].

2 Background

We provide some (fairly standard) background on security protocol verification.

2.1 Messages

Security protocols describe how agents exchange messages, built using cryptographic primitives, in order to obtain security guarantees such as confidentiality or authentication. Protocol specifications are parametric and prescribe a general recipe for communication that can be used by different agents playing in the protocol roles (sender, receiver, server, etc.). The *algebra of messages* tells us how messages are constructed. Following standard practice (e.g., [8, 29]), we consider a countable *signature* Σ and a countable set *Var* of *variable symbols* disjoint from Σ , and write Σ^n for the symbols of Σ with arity n ; thus Σ^0 is the set of *constants*, which we assume to have disjoint subsets that we refer to as *agent names* (or just *agents*), *public keys*, *private keys*, *symmetric keys* and *nonces*. The variables are, however, untyped (unless denoted otherwise) and can be instantiated with arbitrary types, yielding an *untyped model*. We will use upper-case letters to denote variables (e.g., $A, B, \dots$ for agents, N for nonces, etc.) and lower-case letters to denote the corresponding constants (concrete agents names, concrete nonces, etc.). All these may be possibly annotated with subscripts and superscripts.

The symbols of Σ that have arity greater than zero are partitioned into the set Σ_p of (*public*) *operations* and the set Σ_m of *mappings*. The public operations represent all those operations that every agent (including the intruder) can perform on messages

they know. In this paper, we consider the following public operations:¹

- $\{M_1\}_{M_2}$ represents the *asymmetric encryption* of M_1 with public key M_2 ;
- $\{M_1\}_{inv(M_2)}$ represents the *asymmetric encryption* of M_1 with private key $inv(M_2)$ (the mapping $inv(\cdot)$ is discussed below);
- $\{[M_1]\}_{M_2}$ represents the symmetric encryption of M_1 with symmetric key M_2 ;
- $[M_1, M_2]$ (or simply M_1, M_2 when there is no risk of confusion) represents the concatenation of M_1 and M_2 .

In contrast to the public operations, the mappings of Σ_m are those functions that do not correspond to operations that agents can perform on messages, but that map between constants. In this paper, we use the following two mappings. First, $inv(M)$ gives the private key that corresponds to public key M . Second, for long-term key infrastructures, we assume that every agent A has a public key $pk(A)$ and corresponding private key $inv(pk(A))$; thus $pk(\cdot)$ is a mapping from agents to public keys. In the same way, one may model further long-term key infrastructures, e.g., using $sk(A, B)$ to denote a shared key of agents A and B .

Since the mappings map from constants to constants, we consider a term like $inv(pk(a))$ as atomic as its construction does not involve any operation performed by an honest agent or the intruder, nor there is a way to “decompose” such a message into smaller parts. Since we will also deal with terms that contain variables, let us call *atomic* all terms that are built from constants in Σ^0 , variables in Var , and the mappings of Σ_m . The set $\mathcal{T}_\Sigma(Var)$ of all *terms* is the closure of the atomic terms under the operations of Σ_p . A *ground term* is a term without variables, and we denote the set of ground terms with $\mathcal{T}_\Sigma$.

As is often done in security protocol verification, we interpret terms in the *free algebra*, i.e., every term is interpreted by itself and thus two terms are equal iff they are syntactically equal.² For instance, two constant symbols n_1 and n_2 immediately represent different values.

2.2 The Dolev-Yao Intruder

For concreteness and brevity, we consider here the standard Dolev and Yao [16] model of an active intruder, denoted by i , who controls the network but cannot break cryptography, but note that our approach is independent of the actual strength of the intruder and weaker (or stronger, e.g., being able to attack the cryptography) intruder models could be considered.

i can intercept messages and analyze them if he possesses the corresponding keys for decryption, and he can generate messages from his knowledge and send them under

¹We could, of course, quite straightforwardly add other operations, e.g., for hash functions, but refrain from doing so for the sake of brevity.

²Numerous algebras have been considered in security protocol verification, e.g. [12, 28], ranging from the free algebra to various formalizations of algebraic properties of the cryptographic operators employed. Here, we focus, for simplicity, on the free algebra, but nothing in our interpolation method would prevent us from considering a more complex algebra (e.g., for protocols that make use of modular exponentiation or xor).

$$\begin{array}{c}
\frac{M \in IK}{M \in DY(IK)} G_{\text{axiom}} \quad \frac{M_1 \in DY(IK) \quad M_2 \in DY(IK)}{[M_1, M_2] \in DY(IK)} G_{\text{pair}} \\
\\
\frac{M_1 \in DY(IK) \quad M_2 \in DY(IK)}{\{M_1\}_{M_2} \in DY(IK)} G_{\text{crypt}} \quad \frac{M_1 \in DY(IK) \quad M_2 \in DY(IK)}{\{[M_1]\}_{M_2} \in DY(IK)} G_{\text{scrypt}} \\
\\
\frac{[M_1, M_2] \in DY(IK)}{M_i \in DY(IK)} A_{\text{pair}_i} \quad \frac{\{[M_1]\}_{M_2} \in DY(IK) \quad M_2 \in DY(IK)}{M_1 \in DY(IK)} A_{\text{scrypt}} \\
\\
\frac{\{M_1\}_{M_2} \in DY(IK) \quad \text{inv}(M_2) \in DY(IK)}{M_1 \in DY(IK)} A_{\text{crypt}} \quad \frac{\{M_1\}_{\text{inv}(M_2)} \in DY(IK) \quad M_2 \in DY(IK)}{M_1 \in DY(IK)} A_{\text{crypt}}^{-1}
\end{array}$$

Figure 1: The system $\mathcal{N}_{DY}$ of rules of the Dolev-Yao intruder.

any agent name. For a set IK of messages, we define $DY(IK)$ (for “Dolev-Yao” and “Intruder Knowledge”) to be the smallest set closed under the standard *generation* (G) and *analysis* (A) rules of the system $\mathcal{N}_{DY}$ given in Fig. 1. The G rules express that the intruder can compose messages from known messages using pairing, asymmetric and symmetric encryption. The A rules describe how the intruder can decompose messages.

2.3 ASLan++ and ASLan

We give here a brief overview of the security protocol specification languages ASLan++ [32] and ASLan [7], focusing on the aspects relevant to our method. We remark that our methodology can be easily adapted to work with other protocol specification languages (which, like ASLan++, typically specify the different protocol roles as interacting processes) by providing a translator to the SiL input language as described in Section 3.2.

ASLan++ is a formal and typed security protocol specification language, whose semantics is defined in terms of the more low-level language ASLan, which we describe below. An ASLan++ specification consists in a hierarchy of *entity declarations*, which are similar to Java classes. The top-level entity is usually called *Environment* (similar to the “main” procedure of a program) and it typically contains the definition of a *Session* entity, which in turn contains a number of sub-entities (and their instantiations, i.e., `new subentity(<parameters>);`) that represent all the parties involved in a protocol. Each entity of an ASLan++ specification is composed by two main sections: *symbols*, in which there is the instantiation of all the variables and constants used in the entity, and *body*, in which the behavior of the entity is described (e.g., message exchange). Inside the body of an entity we use three different types of statements: assignment, message send and message receive. The only type of *assignment* that we use here is of the form `Var:=fresh()`, which assigns to the variable `Var` a new constant of the proper type. A *message send* statement, `Snd -> Rcv: M`, is composed by two variables `Snd` and `Rcv` representing sender and receiver, respectively, and a message `M` exchanged between the two parties. In *message receive*, `Snd` and `Rcv` are swapped and usually, in order to assign a value to the variable `M`, a `?` precedes the message `M`, i.e., `Snd -> Rcv: ?M`. However,

in ASLan++, the **Actor** keyword refers to the entity itself (similar to “this” or “self” in object-oriented languages) and thus we actually write the send and receive statements as **Actor** $\rightarrow$ Rcv: M and Snd $\rightarrow$ **Actor**: ?M, respectively.

Finally, we describe here two kinds of protocol goals in ASLan++. A *channel goal*, $\text{label}(_) : \text{Snd} \langle \text{chn} \rangle \text{Rcv}$, defines a property $\langle \text{chn} \rangle$ that holds on all (the “_” is a wildcard) the exchanged messages labeled with label between the two entities Snd and Rcv. For example, we use *authentication goals* defined as $\text{auth_goal}(_) : \text{Snd} \star \rightarrow \text{Rcv}$; where $\star \rightarrow$ defines sender authenticity. A *secrecy goal* is defined with $\text{label}(_) : \{\text{Snd}, \text{Rcv}\}$, which states that each message labeled with label can only be shared between the two entities Snd and Rcv.

As discussed in [4], an ASLan++ specification can be automatically translated into a more low-level ASLan specification, which ultimately defines a *transition system* $M = \langle S, I, \rightarrow \rangle$, where S is the *set of states*, $I \subseteq S$ is the set of initial states, and $\rightarrow \subseteq S \times S$ is the (*reflexive*) *transition relation*. The structure of an ASLan specification is composed by six different sections: signature of the predicates, types of variables and constants, initial state, Horn clauses, transition rules of $\rightarrow$ and protocol goals. The content of the sections is intuitively described by their names. In particular, an initial state $I \in I$ is composed by the concatenation of all the predicates that hold before running any rule (e.g., the agent names and the intruder’s own keys).

The specifications that we consider in this paper do not use Horn clauses, but a so called *Prelude* file, in which all the actions of the DY intruder are defined as a set H of Horn clauses, is automatically imported during the translation from ASLan++ into ASLan (see [7]).

The transition relation $\rightarrow$ is defined as follows. For all $S \in S$, $S \rightarrow S'$ iff there exist

- a rule such that

$$PP.NP \& PC \& NC \models [V] \Rightarrow R,$$

where PP and NP are sets of positive and negative predicates, PC and NC conjunctions of positive and negative atomic conditions, and

- a substitution $\sigma : \{v_1, \dots, v_n\} \rightarrow T_\Sigma$, where $v_1, \dots, v_n$ are the variables that occur in PP and PC such that:

1. $PP\sigma \subseteq [S]^H$, where $[S]^H$ is the closure of S with respect to the set of clauses H ,
2. $PC\sigma$ holds,
3. $NP\sigma\sigma' \cap [S]^H = \emptyset$ for all substitutions σ' such that $NP\sigma\sigma'$ is ground,
4. $NC\sigma\sigma'$ holds for all substitutions σ' such that $NC\sigma\sigma'$ is ground and
5. $S' = (S \setminus PP\sigma) \cup R\sigma\sigma''$, where σ'' is any substitution such that for all $v \in V$, $v\sigma''$ does not occur in S .

We now define the translation of the ASLan++ constructs we have considered here. Every ASLan++ entity is translated into a new *state predicate* and added to the section signature. This predicate is parametrized with respect to a *step label* (that uniquely identifies every instance) and it mainly keeps track of the local state of an instance

(current values of whose variables) and expresses the control flow of the entity by means of step labels. As an example, if we have the ASLan++ entity

```
1 entity Snd(Actor, Rcv: agent){
2   symbols
3   Var: message;
4 }
```

the predicate `stateSnd` is added to the section signature and, supposing an instantiation of the entity `new Snd(snd, rcv)`, the new predicate

```
1 state_Snd(snd, iid, sl_0, rcv, dummy_message)
```

is used in transition rules to store all the informations of an entity, where the ID `iid` identifies a particular instance, `sl_0` is the step label, the parameters `Actor`, `Rcv` are replaced with constants `snd` and `rcv`, respectively, and the message variable `Var` is initially instantiated with `dummy_message`.

Given that an ASLan++ specification is a hierarchy of entities, when an entity is translated into ASLan, this hierarchy is preserved by a `child(id_1, id_0)` predicate that states `id_0` is the parent entity of `id_1` and both `id_0` and `id_1` are entity IDs.

A variable assignment statement is translated into a transition rule inside the rules section. As an example, if in the body of the entity `Snd` defined above there is an assignment `Var := constant;`, where `constant` is of the same type of `Var`, then we obtain the following transition rule:

```
1 state_Snd(Actor, IID, sl, Rcv, Var)
2 =>
3 state_Snd(Actor, IID, succ(sl), Rcv, constant)
```

In the case of assignments to `fresh()`, the variable `Var` is assigned to a new variable.

In the case of a message exchange (sending or receiving statements) the `iknows(message)` predicate is added to the right-hand side of the corresponding ASLan rule. This states that the message `message` has been sent over the network, where `iknows` stands for *intruder knows* and is used because, as is usual, the Dolev-Yao intruder is identified with the network itself.

The last point we discuss is the translation of goals focusing on authentication and secrecy described above. The label in a send statements (e.g., `Actor -> Rcv: auth:(Na)`) generates a new predicate `witness(Actor, Rcv, label, Payload)` that is inserted into the ASLan transition rule representing the send statement. An equivalent `request(Actor, Snd, label, Payload, IID)` predicate is added for receive statements. These predicates are used in the translation of goals. In fact, an authentication goal is translated into the state (i.e., attack state)

```
1 not(dishonest(Snd)).
2 not(witness(Snd, Rcv, auth, Payload)).
3 request(Rcv, Snd, auth, Payload, IID)
```

where `not(dishonest(Snd))` states the sender `Snd` has not to be the intruder, `not(witness(Snd, Rcv, auth, Payload))` states the payload of the authentication message has not to be sent by the honest agent `Snd` and the last `request` predicate states the receiver `Rcv` has received the authentication message. A secrecy goal is translated into the attack state

```
1 iknows(Payload).
2 not(contains(i, Knowers)).
3 secret(Payload, label, Knowers)
```


$A \rightarrow i : \{N_A, A\}_{pk(i)}$	$Alice_1.Actor \rightarrow Alice_1.B : \{Alice_1.Na, Alice_1.Actor\}_{pk(Alice_1.B)}$	$a \rightarrow i : \{c_1, a\}_{pk(i)}$
$i(A) \rightarrow B : \{N_A, A\}_{pk(B)}$	$? \rightarrow Bob_2.Actor : \{Bob_2.Na, Bob_2.A\}_{pk(Bob_2.Actor)}$	$i(a) \rightarrow b : \{c_1, a\}_{pk(b)}$
$B \rightarrow i(A) : \{N_A, N_B\}_{pk(A)}$	$Bob_2.Actor \rightarrow Bob_2.A : \{Bob_2.Na, Bob_2.Nb\}_{pk(Bob_2.A)}$	$b \rightarrow i(a) : \{c_1, c_2\}_{pk(i(a))}$
$i \rightarrow A : \{N_A, N_B\}_{pk(A)}$	$Alice_1.B \rightarrow Alice_1.Actor : \{Alice_1.Na, Alice_1.Nb\}_{pk(Alice_1.Actor)}$	$i \rightarrow a : \{c_1, c_2\}_{pk(a)}$
$A \rightarrow i : \{N_B\}_{pk(i)}$	$Alice_1.Actor \rightarrow Alice_1.B : \{Alice_1.Nb\}_{pk(Alice_1.B)}$	$a \rightarrow i : \{c_2\}_{pk(i)}$
$i(A) \rightarrow B : \{N_B\}_{pk(B)}$	$Bob_2.A \rightarrow Bob_2.Actor : \{Bob_2.Nb\}_{pk(Bob_2.Actor)}$	$i(a) \rightarrow b : \{c_2\}_{pk(b)}$

Figure 2: Man-in-the-middle attack on the NSPK protocol (left), symbolic attack trace at state 15 of the algorithm execution (middle) and instantiated attack trace obtained with our method (right).

where `iknows(Payload)` states that the `Payload` has to be sent over the network, that the set of knowers (`Snd` and `Rcv` in the example above) does not contain the intruder `i` and the `secret` predicate is used to check the goal only when the rule containing the secrecy goal label is fired. This is because a `secret(Payload, label, Knowers)` predicate is added to all the transition rules that are translations of statements in which the payload of the secrecy goal is used. The declaration of an attack state `AS` amounts to adding a rule `AS => AS.attack` for a nullary predicate `attack`.

2.4 Running example

As a running example, we will use NSL, the Needham-Schroeder Public Key (NSPK) protocol with Lowe's fix [21], which aims at mutual authentication between *A* and *B*:

$$\begin{aligned}
A \rightarrow B & : \{N_A, A\}_{pk(B)} \\
B \rightarrow A & : \{N_A, N_B, B\}_{pk(A)} \\
A \rightarrow B & : \{N_B\}_{pk(B)}
\end{aligned}$$

The presence of *B* in the second message prevents the man-in-the-middle attack that NSPK suffers from, which is shown on the left of Fig. 2, where we write $i(A)$ to denote that the intruder is impersonating the honest agent *A* (that is, $i(x)$ denotes the intruder playing the role of x , for x an agent name.)

We give the overall ASLan++ specifications for the protocol NSL in Appendix A; here we briefly describe only the section modeling the behavior of the two entities involved. Note that, for readability, from now on, we use math fonts instead of mixing math and typewriter fonts (e.g., we write $iknows(Payload)$ instead of `iknows(Payload)`) in the text, while we use typewriter in code listings.

```

1 entity Alice(Actor, B: agent) {
2   symbols
3   Na, Nb: text;
4   body{
5     Na := fresh();
6     Actor -> B: {Na, Actor}_pk(B);
7     B -> Actor: {Na, ?Nb, B}_pk(Actor);
8     Actor -> B: {auth: (Nb)}_pk(B);
9   }
10 }

1 entity Bob(A, Actor: agent) {
2   symbols
3   Na, Nb: text;
4   body{
5     ? -> Actor: {?Na, ?A}_pk(Actor);
6     Nb := fresh();
7     Actor -> A: {Na, Nb, Actor}_pk(A);
8     A -> Actor: {auth: (Nb)}_pk(Actor);
9   }
10 }

```

The two roles are *Alice*, who is the *initiator* of the protocol, and *Bob*, the *responder*. The elements between parentheses in line 1 declare which variables are used to denote the agents playing the different roles along the specification of the role *Alice*: *Actor* refers to the agent playing the role of *Alice* itself, while *B* is the variable referring to the agent who plays the role of *Bob*. Similarly, the section *symbols* declares that *Na* and *Nb* are variables of type *text*. The section *body* specifies the behavior of the role. First, the operation *fresh()* assigns to the nonce *Na* a value that is different from the value assigned to any other nonce. Then *Alice* sends the nonce, together with her name, to the agent *B*, encrypted with *B*'s public key. In line 7, *Alice* receives her nonce back together with a further variable (expected to represent *B*'s nonce in a regular session of the protocol) and the name of *B*, all encrypted with her own public key. As a last step, *Alice* sends to *B* the nonce *Nb* encrypted with *B*'s public key.

The variable declarations and the behavior of *Bob* are specified by the listing on the right. We omit a full description of the code and only remark that the “?” in the beginning of line 5 denotes the fact that the sender of such a message can be any agent, though no assignment is made for ? in that case.

In this example, we want to verify whether the man-in-the-middle attack known for the NSPK protocol can be still applied after Lowe's fix. The scenario we are interested in can be obtained by the following ASLan++ instantiation:

```

1 body { % of Environment
2   any Session(a,i);
3   any Session(a,b);
4 }

```

In session 1, the roles of *Alice* and *Bob* are played by the agents *a* and *i*, respectively, whereas in session 2 they are played by *a* and *b*.

Finally, a set of goals needs to be specified. For simplicity, here we only require to check the authentication property with respect to the nonce of *Bob*, i.e., we will verify that the responder *Bob* authenticates the initiator *Alice*.

```

1 goals { auth:(_) A *-> B; }

```

As an example of the equivalent ASLan specification, we show the ASLan code of the translation of line 7 of the *Alice* entity

```

1 iknows(crypt(pk(E_S_A_Actor), pair(Na, pair(Nb_1, E_S_A_B))))).
2 state_Alice(E_S_A_Actor, E_S_A_IID, 3, E_S_A_B, Na, Nb)
3 =>
4 state_Alice(E_S_A_Actor, E_S_A_IID, 4, E_S_A_B, Na, Nb_1)

```

where, after receiving the message in the *iknows* predicate, the nonce *Nb* is updated in the state fact in the right-hand side of the transition rule.

3 Translating security protocols into sequential programs

3.1 The SPiM Input Language SiL

In Fig. 3, we present the full grammar of the SPiM Input Language SiL, a simple imperative programming language that we will use to define the sequential programs to be analyzed by the verification algorithm.

Definition 1. The *SPiM Input Language SiL* is defined by the grammar in Fig. 3, where X ranges over a set of *variable locations* Loc and c ranges over the set $\Sigma^0 \cup \mathbb{N}$. $\square$

The basic *terms* of the language are in the syntactic category E . A *message* M is a constant, a variable, a concatenated message or some form of encrypted message.

The category L denotes *lists of messages*, whereas S stands for a *set of messages*: here IK is a special identifier referring to the intruder knowledge and $+$ is used to denote the union operation between sets.

B denotes the class of *Booleans*. In addition to the standard Boolean constants and operators, SiL contains two specific predicates: $IK \vdash M$, which intuitively evaluates to true when the message M is derivable from the set of messages in IK , and *witness*, with three arguments (a sender, a receiver, and a message), which is used in order to verify a goal of authentication.

Finally, the *statements* of SiL, in the category C , comprise standard constructs (like assignments, conditionals and concatenation) together with mechanisms used to handle specific aspects of security protocols, like the possibility of setting the values of the set IK and of the predicates *witness* and *attack*, which takes Boolean values and is set to *true* when an attack is found.³

Definition 2. We denote with $V = Loc \cup \{IK, attack\}$ the set of *program variables* and with $D = \Sigma^0 \cup \mathbb{N} \cup \mathcal{P}(\mathcal{T}_\Sigma) \cup \{true, false\} \cup \mathcal{P}(\Sigma^0 \times \Sigma^0 \times \mathcal{T}_\Sigma)$ the set of *possible data values*, i.e., natural numbers, ground messages, sets of ground messages, Boolean values and sets of triples (agent, agent, message) for the *witness* predicate. $\square$

Note that here, in order to simplify the presentation, we do not use an explicitly typed model. However, the implementation described in Section 5 does make use of a typed model in order to improve the efficiency of the tool (at the small expenses of not being able to find type-flaw attacks).

Definition 3. A *SiL data state* (that we will sometimes refer to only as “state”) is a partial function $\varsigma : V \rightarrow D$ and we denote with $\mathbb{D}$ the set of all such functions. $\square$

In order to specify the behavior of SiL constructs, we present a *big-step structural operational semantics* for it. As it is the case for any structural operational semantics, the definition is given by means of a proof system. Rules manipulate judgments of the form $\langle T, \varsigma \rangle \Downarrow v$, where T denotes an element in any of the syntactic categories of

³Two remarks are in order. First, for simplicity, we give the syntax in the case of a single goal to be considered; in case of more goals, a distinct *attack* variable can be added for each goal. Second, by the definition of the translation procedure into a SiL program, an authentication goal is verified immediately after the receipt of the message on which the authentication is based. Thus, we do not need in SiL an equivalent of the ASLan predicate *request*.

$$\begin{aligned}
E &::= X \mid c \\
M &::= E \mid [M, M] \mid \{M\}_M \mid \{M\}_{inv(M)} \mid \{|M|\}_M \\
L &::= M \mid L, M \\
S &::= \{L\} \mid IK \mid S + S \\
B &::= true \mid false \mid IK \vdash M \mid E = E \mid witness(E, E, M) \mid not(B) \mid B \text{ or } B \mid B \text{ and } B \\
C &::= X := E \mid IK := S \mid attack := B \mid witness(E, E, M) := true \mid C; C \mid \text{if } B \text{ then } C \text{ else } C \mid skip \mid end
\end{aligned}$$

Figure 3: The grammar of SiL.

SiL and v is a data value of the corresponding type (in particular, v is a state in the case when T is a statement). In a big-step semantics [19] formulation, $\langle T, \varsigma \rangle \Downarrow v$ means that by the complete evaluation of T in the state ς , we obtain v . (This is in opposition to what happens in the case of the so-called small-step semantics, where each sequent denotes a minimal, atomic step of evaluation.) For instance, $\langle C, \varsigma \rangle \Downarrow \varsigma'$ denotes that by evaluating the statement C in a state ς , we move to a state ς' . Given the simplicity of the language and the kind of analysis that we intend to carry out on it, we chose to give a big-step semantics, which typically has the advantage of needing fewer inference rules and allowing for a more concise presentation.

Definition 4. The *big-step operational semantics of SiL* is defined by the proof system in Fig. 4, where we use the following meta-variables: m ranges over $\mathcal{T}_\Sigma$, l ranges over lists of elements of $\mathcal{T}_\Sigma$, p ranges over $\mathcal{P}(\mathcal{T}_\Sigma)$, and $b \in \{true, false\}$. We denote with $\varsigma[m/X]$ the state obtained from ς by replacing the content of X by m , i.e., $\varsigma[m/X](Y) = m$ if $Y = X$ and $\varsigma[m/X](Y) = \varsigma(Y)$ otherwise. $\square$

Rules for the evaluation of basic terms are quite simple: a constant evaluates to itself and a variable to the value associated to it in a given data state.

Rules for compound messages evaluate the single components and then merge the results in a message of the appropriate form.

Rules of the third class show how lists are evaluated by concatenating single messages and how sets of messages are built by using lists. In particular, the special set variable IK is evaluated in the same way as any other variable.

Evaluation of Booleans is standard: constants evaluate to themselves; predicates (equality and witness) evaluate either to *true* or *false*, according to a side condition referring to the values of the arguments; compound Boolean expressions are evaluated by functionally composing the truth values of the components.

Finally, rules for statements operate by modifying the data state on which they are applied. Assignments modify the state value of the variable considered (be it a generic variable, IK or a variable referring to a predicate). Concatenation and conditional statements are treated as usual. *skip* and *end* do not alter the data state: the first one is just introduced in order to simplify the proof of some results, while the latter allows one to ignore the statements that follow.

3.2 The translation procedure

Definition 5. Given a protocol $\mathcal{P}$ involving a set $\mathcal{R}$ of *roles* (*Alice, Bob, ...*, a.k.a. *entities*), a *session instance* (or *session*, for short) of $\mathcal{P}$ is a function si assigning an agent

BASIC TERMS

$$\langle X, \varsigma \rangle \Downarrow \varsigma(X) \quad \langle c, \varsigma \rangle \Downarrow c$$

MESSAGES

$$\frac{\langle M_1, \varsigma \rangle \Downarrow m_1 \quad \langle M_2, \varsigma \rangle \Downarrow m_2}{\langle [M_1, M_2], \varsigma \rangle \Downarrow [m_1, m_2]} \quad \frac{\langle M_1, \varsigma \rangle \Downarrow m_1 \quad \langle M_2, \varsigma \rangle \Downarrow m_2}{\langle \{M_1\}_{M_2}, \varsigma \rangle \Downarrow \{m_1\}_{m_2}} \quad \frac{\langle M_1, \varsigma \rangle \Downarrow m_1 \quad \langle M_2, \varsigma \rangle \Downarrow m_2}{\langle \{M_1\}_{inv(M_2)}, \varsigma \rangle \Downarrow \{m_1\}_{inv(m_2)}} \quad \frac{\langle M_1, \varsigma \rangle \Downarrow m_1 \quad \langle M_2, \varsigma \rangle \Downarrow m_2}{\langle \{[M_1]\}_{M_2}, \varsigma \rangle \Downarrow \{[m_1]\}_{m_2}}$$

LISTS AND SETS OF MESSAGES

$$\frac{\langle L, \varsigma \rangle \Downarrow l \quad \langle M, \varsigma \rangle \Downarrow m}{\langle L, M, \varsigma \rangle \Downarrow l, m} \quad \frac{\langle L, \varsigma \rangle \Downarrow l}{\langle \{L\}, \varsigma \rangle \Downarrow \{l\}} \quad \langle IK, \varsigma \rangle \Downarrow \varsigma(IK) \quad \frac{\langle S_1, \varsigma \rangle \Downarrow p_1 \quad \langle S_2, \varsigma \rangle \Downarrow p_2}{\langle S_1 + S_2, \varsigma \rangle \Downarrow p_1 \cup p_2}$$

BOOLEAN EXPRESSIONS

$$\begin{aligned} \langle true, \varsigma \rangle \Downarrow true \quad \langle false, \varsigma \rangle \Downarrow false \quad & \frac{\langle IK, \varsigma \rangle \Downarrow \varsigma(IK) \quad \langle M, \varsigma \rangle \Downarrow m}{\langle IK \vdash M, \varsigma \rangle \Downarrow true} \quad m \in DY(\varsigma(IK)) \quad \frac{\langle IK, \varsigma \rangle \Downarrow \varsigma(IK) \quad \langle M, \varsigma \rangle \Downarrow m}{\langle IK \vdash M, \varsigma \rangle \Downarrow false} \quad m \notin DY(\varsigma(IK)) \\ \\ \frac{\langle E_1, \varsigma \rangle \Downarrow c_1 \quad \langle E_2, \varsigma \rangle \Downarrow c_2}{\langle E_1 = E_2, \varsigma \rangle \Downarrow true} \quad c_1 = c_2 \quad & \frac{\langle E_1, \varsigma \rangle \Downarrow c_1 \quad \langle E_2, \varsigma \rangle \Downarrow c_2}{\langle E_1 = E_2, \varsigma \rangle \Downarrow false} \quad c_1 \neq c_2 \\ \\ \frac{\langle E_1, \varsigma \rangle \Downarrow c_1 \quad \langle E_2, \varsigma \rangle \Downarrow c_2 \quad \langle M, \varsigma \rangle \Downarrow m}{\langle witness(E_1, E_2, M), \varsigma \rangle \Downarrow true} \quad & (c_1, c_2, m) \in \varsigma(witness) \\ \\ \frac{\langle E_1, \varsigma \rangle \Downarrow c_1 \quad \langle E_2, \varsigma \rangle \Downarrow c_2 \quad \langle M, \varsigma \rangle \Downarrow m}{\langle witness(E_1, E_2, M), \varsigma \rangle \Downarrow false} \quad & (c_1, c_2, m) \notin \varsigma(witness) \\ \\ \frac{\langle B, \varsigma \rangle \Downarrow b}{\langle not(B), \varsigma \rangle \Downarrow \neg b} \quad \frac{\langle B_1, \varsigma \rangle \Downarrow b_1 \quad \langle B_2, \varsigma \rangle \Downarrow b_2}{\langle B_1 \text{ or } B_2, \varsigma \rangle \Downarrow b_1 \vee b_2} \quad & \frac{\langle B_1, \varsigma \rangle \Downarrow b_1 \quad \langle B_2, \varsigma \rangle \Downarrow b_2}{\langle B_1 \text{ and } B_2, \varsigma \rangle \Downarrow b_1 \wedge b_2} \end{aligned}$$

STATEMENTS

$$\begin{aligned} \frac{\langle E, \varsigma \rangle \Downarrow c}{\langle X := E, \varsigma \rangle \Downarrow \varsigma[c/X]} \quad \frac{\langle S, \varsigma \rangle \Downarrow p}{\langle IK := S, \varsigma \rangle \Downarrow \varsigma[p/IK]} \quad & \frac{\langle B, \varsigma \rangle \Downarrow b}{\langle attack := B, \varsigma \rangle \Downarrow \varsigma[b/attack]} \\ \\ \frac{\langle E_1, \varsigma \rangle \Downarrow c_1 \quad \langle E_2, \varsigma \rangle \Downarrow c_2 \quad \langle M, \varsigma \rangle \Downarrow m}{\langle witness(E_1, E_2, M) := true, \varsigma \rangle \Downarrow \varsigma[witness \cup \{(c_1, c_2, m)\}/witness]} \quad & \frac{\langle C_0, \varsigma \rangle \Downarrow \varsigma'' \quad \langle C_1, \varsigma'' \rangle \Downarrow \varsigma'}{\langle C_0; C_1, \varsigma \rangle \Downarrow \varsigma'} \\ \\ \frac{\langle B, \varsigma \rangle \Downarrow true \quad \langle C_0, \varsigma \rangle \Downarrow \varsigma'}{\langle if B then C_0 else C_1, \varsigma \rangle \Downarrow \varsigma'} \quad \frac{\langle B, \varsigma \rangle \Downarrow false \quad \langle C_1, \varsigma \rangle \Downarrow \varsigma'}{\langle if B then C_0 else C_1, \varsigma \rangle \Downarrow \varsigma'} \quad & \langle skip, \varsigma \rangle \Downarrow \varsigma \quad \langle end, \varsigma \rangle \Downarrow \varsigma \quad \frac{\langle end, \varsigma \rangle \Downarrow \varsigma \quad \langle C_1, \varsigma' \rangle \Downarrow \varsigma''}{\langle end; C_1, \varsigma \rangle \Downarrow \varsigma} \end{aligned}$$

Figure 4: A big-step semantics for SiL.

(honest agent or the intruder i) to each element of $\mathcal{R}$.

A *scenario of a protocol* $\mathcal{P}$ is a finite number of session instances of $\mathcal{P}$. □

The input of our method is then:

1. an ASLan++ specification of a protocol $\mathcal{P}$,
2. a scenario $\mathcal{S}$ of $\mathcal{P}$, and
3. a set of goals (i.e., properties to be verified) in $\mathcal{S}$.

We will first describe how to obtain a program for a single session and then how to decorate it with goal locations used to verify security properties. In Section 3.3, finally, we will explain how to combine more sessions in a single graph.

3.2.1 Translating a single session

First of all, we notice that in our translation, and according to the ASLan++/ASLan instantiation mechanism, a session instance between two honest agents is represented as the composition of two sessions, where each of the honest agents communicates with the intruder. We will refer to the session instances obtained after such a translation as *program instances*.

Example 1. For example, the second session of our running example (Section 2.4), i.e., the one between a and b , is obtained by the composition of two sessions, the first played by a and $i(b)$ and the second by $i(a)$ and b , thus giving rise to the following three program instances

Program	Alice	Bob
1	a	i
2	a	$i(b)$
3	$i(a)$	b

We also remark that an intruder is not obliged to play dishonestly; thus, e.g., a program instance between a and $i(b)$ does also capture the case of an honest session between a and b . □

To simplify notation, for the variables and constants of the resulting program we will use the same names as the ones used in the ASLan++ specification. However, in order to distinguish between variables with the same name occurring in the specification of different roles, program variables have the form $E.V$ where E denotes the role and V the variable name in the specification. In the case when more than one session are considered, we also prefix an index denoting the session to the program variable name, e.g., as in $SI_E.V$.

The behavior of the intruder introduces a form of non-determinism, which we capture by representing the program as a procedure depending on a number of input values, one for each choice of the intruder. Along a single program, input variables are denoted by the symbol Y , possibly subscripted with an index. Finally, symbols of the form c_i , for i an integer, are used to denote constants to be assigned to nonces.

Structure of the program The exchange of messages in a session follows a given flow of execution that can be used to determine an order between the instructions contained in the different roles. Such a sequence of instructions will constitute the skeleton of the program.

After a first section that concerns the initialization of the variables, the program will indeed contain a proper translation, based on the semantics of ASLan++, of the instructions in such a sequence. For each program instance, we will follow the flow of execution of the honest agents, as we can think of the intruder actions as not being driven by any protocol, and model the intruder interaction with the honest agents by means of $IK \vdash M$ statements and updates of IK .

In the next paragraphs, we will describe more specifically: (i) how variables are initialized and (ii) how each statement is translated.

Initialization of the variables A first section of the program consists in the initialization of the variables. Let pi be the program instance of the program we are considering. For each role *Alice* such that $pi(Alice) = a$, for some agent name $a \neq i$, we have an initialization instruction $Alice.Actor := a$. Furthermore, for the same *Alice*, and for each other role *Bob*, with B being the variable referring to the role *Bob* amongst the agent variables of *Alice*: if $si(Bob) = b$, then we have the assignment $Alice.B := b$. Finally, it is necessary to initialize the intruder knowledge. A typical initialization instruction for IK has the form:

$$IK := \{a_1, \dots, a_n, i, pk(a_1), \dots, pk(a_n), pk(i), inv(pk(i))\}.$$

That is, i knows each agent a_j involved in the scenario and his public keys $pk(a_j)$, as well as his own public and private keys $pk(i)$ and $inv(pk(i))$. Specific protocols might require a specific initial intruder knowledge or the initialization of further variables, depending on the context, such as symmetric keys or hash functions, which are possibly defined in the Prelude section of the ASLan++ specification.

Sending of a message The sending of a message $Actor \rightarrow B : M$ defined in a role *Alice* is translated into the instruction $IK := IK + \{M\}$, where the symbol $+$ denotes set union (corresponding to $\cup$) so that the intruder knowledge is increased with the message M .

Receipt of a message Consider the receipt of a message $R \rightarrow Actor : M$ in a role *Alice*. Assume the message is sent from a role *Bob*. Then the instruction is translated into the following code, where $Q_1, \dots, Q_n$ are all the variables occurring preceded by $?$ in M and $Y_1, \dots, Y_n$ are distinct input variables not introduced elsewhere:

```

1  If (IK  $\vdash$  M)
2    then Alice.Q_1 := Y_1;
3     ...
4     Alice.Q_n := Y_n;
5  else end;
```

Generation of fresh values Finally, an instruction of the form $N := \text{fresh}()$ in *Alice*, which assigns a fresh value to a nonce, can be translated into the instruction $\text{Alice}.N := c_I$, where c_I is a constant not introduced elsewhere.

3.2.2 Defining goals for the verification of security properties

Introducing goal locations The next step consists in decorating the program with a goal location for each security property to be verified. As it is common when performing symbolic execution [20], we express such properties as correctness assertions, typically placed at the end of a program. Once we have represented a protocol session as a program (or more programs in the case when a session instance is split into more program instances), and defined the properties we are interested in as correctness assertions in such a program, the problem of verifying security properties over (a session of) the protocol is reduced to verifying the correctness of the program with respect to those assertions.

We consider here three common security properties (authentication, confidentiality and integrity) and show how to represent them inside the program in terms of assertions. They are expressed by means of a statement of the form *if* ($\text{not}(\text{expr})$) *then* $\text{attack} := \text{true}$, where expr is an expression referring to the goal considered, as described below.

Authentication Assume that we want to verify that *Alice* authenticates *Bob* with respect to a message M in the specification of the protocol, in a given program instance by the ASLan++ statement: $B \rightarrow \text{Actor} : \text{auth} : (M)$ where auth is the label of the goal and a corresponding sending statement is included in the specification.

We can restrict our attention to the case when according to the program instance under consideration *Bob* is played by i , since otherwise the authentication property is trivially satisfied. The problem thus reduces to verifying whether the agent i is playing under his real name (in which case authentication is again trivially satisfied) or whether i is pretending to be someone else, i.e., whether the agent playing *Alice* believes she is speaking to someone who is not i . Hence, one of the conditions required in order to reach the goal is $\text{not}(\text{Alice}.B = i)$, where B is the agent variable referring to the role *Bob* inside *Alice*.

A second condition is necessary and concerns the fact that the message M has not been sent by *Alice*. B to *Alice*.*Actor*. This can be verified by using the witness predicate, which is set to true when the message is sent and whose state is checked when a goal is searched for, i.e., immediately after the receipt of the message M .

Example 2. In *NSL*, we are interested in verifying a property of authentication in the session that assigns i to *Alice* and b to *Bob*: namely, we want *Bob* to authenticate *Alice* with respect to the nonce *Bob*.*Nb* contained in the reception in line 8 on the right of the *NSL* example (Section 2.4). Such a receipt corresponds to the sending of line 8 on the left. Thus we can add a witness assignment of the form $\text{witness}(\text{Alice}.\text{Actor}, \text{Alice}.B, [\text{Alice}.Nb, \text{pk}(\text{Alice}.B)]) := \text{true}$ after the sending, and the instruction

```

1 if ( $\text{not}(\text{Bob}.A = i)$  and  $\text{not}(\text{witness}(\text{Bob}.A, \text{Bob}.\text{Actor}, \{\text{Bob}.Nb\}_{\text{pk}(\text{Bob}.\text{Actor})}))$ )
2   then  $\text{attack} := \text{true}$ ;
3   else skip;

```


after the receipt of the message. $\square$

Confidentiality Assume that we want to verify that the message corresponding to a variable M , in the specification of a role $Alice$ of the protocol, is confidential between a given set of roles $\mathcal{R} = \{Alice_1, \dots, Alice_n\}$ in a session si , i.e., we have a sending statement $Actor \rightarrow B : \{secret : (M)\}$, where $secret$ is the goal label, for a confidentiality goal expressed as $secret : (_) \{Alice_1, \dots, Alice_n\}$. This amounts to checking whether the agent i got to know the confidential message M even though i is not included in $\mathcal{R}$. Inside the program, this corresponds to verifying whether the message $Alice.M$ can be derived from the intruder knowledge and whether any honest agent playing a role in $\mathcal{R}$ believes that at least one of the other roles in $\mathcal{R}$ is indeed played by i , which we can read as having indeed $i \in \mathcal{R}$. The following assertion is added at the end of the SiL program:

```

1 if ((IK  $\vdash$  Alice.M) and (not((Alice_1.B^1_1 = i) or
2   ... (Alice_1.B^1_m = i) or ...
3   (Alice_n.B^n_1 = i) or ... (Alice_n.B^n_m = i))))
4   then attack := true;
5   else skip;

```

where $Alice_j$, for $1 \leq j \leq n$, is a role such that $Alice_j \in \mathcal{R}$ and $si(Alice_j) \neq i$, $\{Bob_1, \dots, Bob_m\} \subseteq \mathcal{R}$ is the subset of those roles in $\mathcal{R}$ that are instantiated with i by si and B^j_l , for $1 \leq j \leq n$ and $1 \leq l \leq m$, is the variable referring to the role Bob_l in the specification of the role $Alice_j$.

Example 3. For NSL, assume that we want to verify the confidentiality of the variable Nb (contained in the specification of Bob) between the roles in the set $\{Alice, Bob\}$. We can express this goal by appending the assertion

```

1 if ((IK  $\vdash$  Bob.Nb) and (not (Bob.A = i)))
2   then attack := true;
3   else skip;

```

at the end of the program. $\square$

Integrity In this case, we assume that two variables (possibly of two different roles) are specified in input as the variables containing the value whose integrity needs to be checked. The check will consist in verifying whether the two variables, at a given point of the session execution, also given in input, have the same value. Let M in the role $Alice$ and M' in the role Bob be the two variables; then the corresponding correctness assertion will be

```

1 if (not (Alice.M = Bob.M'))
2   then attack := true;

```

Example 4. The program instances described in Example 1 give rise to the following three SiL programs, for which a single IK initialization instruction holds:

```

IK := {a, b, i, pk(a), pk(b), pk(i), inv(pk(i))}

```

Program 1

```

1  S1_Alice.Actor := a;
2  S1_Alice.B := i;
3  S1_Alice.Na := c_1;
4
5  IK := IK + {{S1_Alice.Na, S1_Alice.Actor}}_pk(S1_Alice.B));
6
7  if (IK |- {{S1_Alice.Na, {S1_Alice.Y_1, S1_Alice.B}}}_pk(S1_Alice.Actor)
8      then S1_Alice.Nb := S1_Alice.Y_1;
9      else end;
10
11 IK := IK + {{S1_Alice.Nb}_pk(S1_Alice.B)};
12 witness(S1_Alice.Actor, S1_Alice.B, {S1_Alice.Nb}_pk(S1_Alice.B)) := true;

```

Program 2

```

1  S2_Alice.Actor := a
2  S2_Alice.B := b
3
4  S2_Alice.Na := c_1
5
6  IK := IK + {{S2_Alice.Na, S2_Alice.Actor}}_pk(S2_Alice.B)}
7
8  if (IK |- {{S2_Alice.Na, {S2_Alice.Y_1, S2_Alice.B}}}_pk(S2_Alice.Actor))
9      then S2_Alice.Nb := S2_Alice.Y_1;
10     else end;
11
12 IK := IK + {{S2_Alice.Nb}_pk(S2_Alice.B)}
13 witness(S2_Alice.Actor, S2_Alice.B, {S2_Alice.Nb}_pk(S2_Alice.B)) := true;

```

Program 3

```

1  S2_Bob.A := a
2  S2_Bob.Actor := b
3
4  if (IK |- {{S2_Bob.Y_1, S2_Bob.Y_2}}_pk(S2_Bob.Actor))
5      then S2_Bob.Na := S2_Bob.Y_1;
6           S2_Bob.A := S2_Bob.Y_2;
7      else end;
8
9  S2_Bob.Nb := c_2
10
11 IK := IK + {{S2_Bob.Na, {S2_Bob.Nb, S2_Bob.Actor}}}_pk(S2_Bob.A)}
12
13 if (IK |- {S2_Bob.Nb}_pk(S2_Bob.Actor))
14     then
15         if (not (witness(S2_Bob.A, S2_Bob.Actor, {S2_Bob.Nb}_pk(S2_Bob.Actor)))
16             and
17             (not (S2_Bob.A = i)));
18         then attack := true;
19     else end;

```

□

3.3 Combining more sessions

Now we need to define a global program that properly “combines” the programs related to all the sessions in the scenario. The idea is that such a program allows for executing, in the proper order, all the instructions of all the sessions in the scenario; the way in which instructions of different sessions are interleaved will be determined by the value of further input variables, denoted by X (possibly subscripted), which can be seen as choices of the intruder with respect to the flow of the execution. Namely, we start to execute each session sequentially and we get blocked when we encounter the receipt of

a message sent by a role that is played by the intruder. When all the sessions are blocked on instructions of that form, the intruder chooses which session has to be reactivated.

For what follows, it is convenient to see a sequential program as a graph (which can be simply obtained by representing its control flow) on which the algorithm of Section 4 for symbolic execution and annotation will be executed. We recall here some notions concerning programs and program runs.

Definition 6. A (SiL) *program graph* is a finite, rooted, labeled graph (Λ, l_0, Δ) where Λ is a *finite set of program locations*, l_0 is the *initial location* and $\Delta \subseteq \Lambda \times \mathcal{A} \times \Lambda$ is a *set of transitions* labeled by actions from a set $\mathcal{A}$, containing the assignments and conditional statements provided by the language SiL.

A (SiL) *program path* of length k is a sequence of the form $l_0, a_0, l_1, a_1, \dots, l_k$, where each *step* $(l_j, a_j, l_{j+1}) \in \Delta$ for $0 \leq j < k-1$.

Let ζ_0 be the initial data state. A (SiL) *program run* of length k is a pair (π, ω) , where π is a program path $l_0, a_0, l_1, a_1, \dots, l_k$ and $\omega = \zeta_0, \dots, \zeta_{k+1}$ is a sequence of data states such that $\langle a_j, \zeta_j \rangle \Downarrow \zeta_{j+1}$ for $0 \leq j \leq k$. $\square$

Let $\mathcal{S}$ be a scenario of a protocol $\mathcal{P}$ with m program instances $pi_1, \dots, pi_m$. We can associate to each program instance pi_j , for $1 \leq j \leq m$, a sequential program by following the procedure described in Section 3.2.

For each $1 \leq j \leq m$, we have a program graph $\mathcal{G}^j = (\Lambda^j, l_0^j, \Delta^j)$ corresponding to the program of pi_j . The program graph corresponding to a given scenario is obtained by composing the graphs of the single program instances. Below we describe an algorithm for concretely obtaining such a graph for $\mathcal{S}$. For simplicity, we will assume that the original specification of $\mathcal{P}$ is such that no receipts of messages are contained inside a loop statement or an if-statement.

Definition 7. Given a program graph, an *intruder location* is a location of the graph corresponding to the receipt of a message sent from a role played by i .

A *block of a program graph* $\mathcal{G}'$ is a subgraph of $\mathcal{G}'$ such that its initial location is either the initial location of $\mathcal{G}'$ or an intruder location.

The *exit locations of a block* $\mathcal{B}$ are the locations of $\mathcal{B}$ with no outgoing edges.

A program graph can simply be seen as a sequence of blocks. Namely, we can associate to the program graph $\mathcal{G}^j$, for each $1 \leq j \leq m$, its *block structure*, i.e., a sequence $\mathcal{B}_1^j, \dots, \mathcal{B}_n^j$ of blocks of $\mathcal{G}^j$, such that: (i) the initial location of $\mathcal{B}_1^j$ is the initial location of $\mathcal{G}^j$; (ii) each intruder location of $\mathcal{G}^j$ is the initial location of $\mathcal{B}_k^j$ for some $1 \leq k \leq n$; (iii) for $1 \leq k < n$, the initial location of $\mathcal{B}_{k+1}^j$ coincides, in $\mathcal{G}^j$, with an exit location of $\mathcal{B}_k^j$; (iv) the graph obtained by composing $\mathcal{B}_1^j, \dots, \mathcal{B}_n^j$, i.e., by letting the initial location of $\mathcal{B}_{k+1}^j$ coincide with the corresponding exit location of $\mathcal{B}_k^j$, is $\mathcal{G}^j$ itself. $\square$

Intuitively, we decompose a session program graph $\mathcal{G}^i$ into sequential blocks starting at each intruder location. The idea is that each such a block will occur as a subgraph in the general scenario graph $\mathcal{G}$ (possibly with more than one occurrence). Namely, the procedure for generating the scenario graph will create a graph that allows one to execute all the blocks of the scenario just once, in any possible sequence that respects

```

create a location  $l$ ;
 $\Lambda := \{l\}$ ;
 $l_0 := l$ ;
 $\Delta := \emptyset$ ;
 $pc(l, j) := 1$  for  $1 \leq j \leq m$ ;
 $ic(l) := 1$ ;
for  $h = 1$  to  $m$  do {
    if (initial location of  $\mathcal{B}_1^h$  is not intruder location)
    then {
        attach  $\mathcal{B}_1^h$  to  $l$ ;
        let  $l'$  be the main exit location of  $\mathcal{B}_1^h$ ;
         $pc(l', j) := pc(l, j)$  for all  $j \neq h$ ;
         $pc(l', h) := pc(l, h) + 1$ ;
         $ic(l') := 1$ ;
         $l := l'$ 
    }
}

 $T := \{l\}$ ;
do {
    pick a location  $l \in T$ ;
    for  $h = 1$  to  $m$  do {
        if ( $\mathcal{B}_{pc(l,h)}^h$  does exist) then {
            create a location  $l^*$ ;
             $\Lambda := \Lambda \cup \{l^*\}$ ;
             $\Delta := \Delta \cup \{(l, \mathbf{i} \mathbf{f} \ \mathbf{x}_k = \mathbf{i}, l^*)\}$ , where  $k = ic(l)$ ;
            attach  $\mathcal{B}_{pc(l,h)}^h$  to  $l^*$ ;
            let  $l'$  be the main exit location of  $\mathcal{B}_{pc(l,h)}^h$ ;
             $pc(l', j) := pc(l, j)$  for all  $h \neq j$ ;
             $pc(l', h) := pc(l, h) + 1$ ;
             $ic(l') := ic(l) + 1$ ;
             $T := T \cup \{l'\}$ ;
        }
    }
     $T := T \setminus \{l\}$ ;
} while ( $T \neq \emptyset$ );

```

Figure 5: An algorithm for building the graph $\mathcal{G} = (\Lambda, l_0, \Delta)$.

the order of the single sessions. For instance, given the block structures $(\mathcal{B}_1^1, \mathcal{B}_2^1)$ and $(\mathcal{B}_1^2)$, the resulting graph will contain a path corresponding to the execution of $\mathcal{B}_1^1, \mathcal{B}_1^1, \mathcal{B}_2^1$ in this order, as well as a path for $\mathcal{B}_1^1, \mathcal{B}_1^2, \mathcal{B}_2^1$, as well as a path for $\mathcal{B}_1^2, \mathcal{B}_1^1, \mathcal{B}_2^1$. Note also that, under the restriction on $\mathcal{P}$ introduced above (i.e., no receipts inside loops and if-statements), each block has at most one “interesting” exit location, in the sense that at most one of its exit locations did not correspond to a location with no outgoing edges even in the original session graph. In the algorithm of Fig. 5, we will refer to such an exit location as the *main exit location*.

In Fig. 5, we give an algorithm that we have devised to incrementally build the graph $\mathcal{G} = (\Lambda, l_0, \Delta)$ starting from the root and adding blocks step by step. We assume the number of program instances m given. In the algorithm we use a procedure *attach*, which given a block $\mathcal{B}$ and a location l , adds the subgraph $\mathcal{B}$ to $\mathcal{G}$ (by letting the initial location of $\mathcal{B}$ coincide with l) and updates the sets Λ and Δ accordingly. During the construction, the set $T \subseteq \Lambda$ contains the locations of the graph to be still expanded. Two functions $pc : \Lambda \times \{1, \dots, m\} \rightarrow \mathbb{N}$ and $ic : \Lambda \rightarrow \mathbb{N}$ are used to keep track of the status of the construction. Their intended meaning is the following: assume that the location l in the graph is still to be expanded; then for each $1 \leq j \leq m$, $\mathcal{B}_{pc(l,j)}^j$ is the next block to be added for what concerns the program instance pi_j (i.e., each path going from the root to l has already executed $\mathcal{B}_h^j$ for $1 \leq h < pc(l, j)$) and the next input variable to be used is $X_{ic(l)}$.

The first **for** loop in the pseudo-code of the algorithm composes, in a sequence, the first blocks of each session program graph. Then the **while** loop expands the graph by

adding a fork at each intruder choice.

The resulting graph $\mathcal{G} = (\Lambda, l_0, \Delta)$ can be finally simplified by making indistinguishable nodes collapse into one, according to standard graph and transition system optimization techniques.

Example 5. Fig. 8 shows a path of the program graph for NSL in the scenario described in the previous examples. The entire graph (Appendix B) is obtained by unifying some equivalent nodes in the graph produced by the algorithm of Fig. 5. $\square$

3.4 Correctness of the translation

Now, we show that the translation into SiL, defined in Sections 3.2 and 3.3, preserves important properties of the original specification. In particular, we show that given an ASLan++ specification, an attack state can be reached by analyzing its ASLan translation if and only if an attack state can be found by executing its SiL translation.

Equivalence of single steps

Definition 8. We say that an ASLan term M' and a SiL term M'' are *equivalent*, $M' \sim M''$, iff one of the following conditions holds:

- $M' \equiv c', M'' \equiv c''$ and $c' = c''$;
- $M' \equiv \text{pair}(M'_1, M'_2), M'' \equiv [M''_1, M''_2]$ and $M'_1 \sim M''_1, M'_2 \sim M''_2$;
- $M' \equiv \text{crypt}(M'_1, M'_2), M'' \equiv \{M''_2\}_{M''_1}$ and $M'_1 \sim M''_1, M'_2 \sim M''_2$;
- $M' \equiv \text{scrypt}(M'_1, M'_2), M'' \equiv \{|M''_2|\}_{M''_1}$ and $M'_1 \sim M''_1, M'_2 \sim M''_2$;
- $M' \equiv \text{inv}(M'_1), M'' \equiv \text{inv}(M''_1)$ and $M'_1 \sim M''_1$. $\square$

In the following, we consider an ASLan++ program and the corresponding ASLan translation. As described in Section 2.3, for each signature in the *SignatureSection* we will have a corresponding state fact.

Definition 9. We define a *variable mapping* as a function $f(E, A)$ that given an entity name E and a variable name A returns the value i corresponding to the index of the position of variable A in the state fact $\text{state_}E$. $\square$

Note that such a function always exists and it is implicitly created at translation time from the translation procedure from ASLan++ into ASLan described in Section 2.3.

In order to handle multiple sessions, let $pi_1, \dots, pi_n$ be the program instances of the considered protocol scenario; we can assume to have a function g such that $g(j) = \text{SID}$ where SID is the identifier of the state fact $\text{state_Session_j}(\dots, \text{SID}, \dots) \subseteq S$ representing the symbolic session corresponding to the program instance pi_j ; note that such a function is implicitly created when a symbolic session is instantiated (Section 2.3) and is bijective. Then we will write $S(E_j, i)$ to indicate the value v_i of the state predicate $\text{state_}E(v_1, \text{SID}, \dots, v_n)$ such that $\text{child}(g(j), \text{SID}) \subseteq S$.

Definition 10. We say that an ASLan state S and a SiL state ς are *equivalent*, $S \sim \varsigma$, iff:

- for each SiL ground term M' and ASLan ground term M'' such that $M' \sim M''$, $M' \in DY(\varsigma(IK)) \Leftrightarrow \text{iknows}(M'') \subseteq \lceil S \rceil^H$;
- $\varsigma(Sj_E.A) = S(E_j, f(E, A))$ for each E representing an entity name involved in the protocol, for each A representing an ASLan++ variable name or parameter name of entity E , for each session instance si_j
- $\varsigma(\text{attack}) = \text{true} \Leftrightarrow \text{attack} \subseteq \lceil S \rceil^H$;
- $(M, M_1, M_2) \in \varsigma(\text{witness}) \Leftrightarrow \text{witness}(M', M'_1, M'_2, \dots) \subseteq \lceil S \rceil^H$, where M, M_1 and M_2 are SiL ground terms and M', M'_1 and M'_2 are ASLan ground terms such that $M \sim M', M_1 \sim M'_1$ and $M_2 \sim M'_2$. $\square$

We notice that while an ASLan transition occurs when there exists a substitution (of values for variables) that makes a rule applicable, in SiL we simulate, and in a sense make more explicit, such a substitution by using the Y input variables. This establishes a correspondence between ASLan substitutions and assignments of values to SiL input variables, which will be important in the following proofs, and that we will handle by means of the following notion of *extension* of a SiL state.

Definition 11. Given a SiL state ς and a set of input variables $Y_1, \dots, Y_n$ such that $\varsigma(Y_i)$ is undefined, we define an *extension* $\bar{\varsigma}$ of ς as a SiL state where $\bar{\varsigma}$ is defined for $Y_1, \dots, Y_n$ and for each other variable A , $\bar{\varsigma}(A) = \varsigma(A)$. $\square$

Since the input variables of the form Y_i are not involved in the definition of equivalence, if an ASLan state S and a SiL state ς are equivalent, that is $S \sim \varsigma$, and $\bar{\varsigma}$ is an extension of ς , then also S and $\bar{\varsigma}$ are equivalent, that is $S \sim \bar{\varsigma}$.

Let r be an ASLan rule; we will write $S \xrightarrow{r} S'$ iff there exists a transition from an ASLan state S to an ASLan state S' obtained by applying the rule r .

Lemma 1. Let I be an ASLan++ statement, r the corresponding ASLan rule and w the corresponding SiL code, as defined in Section 2.3 and 3.2, respectively. Given an ASLan state S and a SiL state ς such that $S \sim \varsigma$ we have:

1. If $S \xrightarrow{r} S'$ then there exists an extension $\bar{\varsigma}$ of ς such that $\langle w, \bar{\varsigma} \rangle \Downarrow \varsigma'$ and $S' \sim \varsigma'$;
2. If there exists an extension $\bar{\varsigma}$ of ς such that $\langle w, \bar{\varsigma} \rangle \Downarrow \varsigma'$, then either there exists an S' such that $S \xrightarrow{r} S'$ and $S' \sim \varsigma'$ or $S = S'$.

Proof. The proof proceeds by considering all the possible ASLan++ statements and is given in Appendix C. $\square$

Equivalence of runs We have showed that, starting from equivalent states, the application of ASLan rules and SiL code fragments that have been generated from the same ASLan++ statements brings to states that are still equivalent. Now we will show

that given an ASLan++ specification, for each run in the SiL translation, there exists a sequence of corresponding ASLan rules in the ASLan translation.

First, we note that, strictly speaking, the translation of an ASLan++ statement into SiL is not always an atomic action, e.g., in the case of a receipt, the corresponding SiL action comprises both a conditional and some assignments. For simplicity, in the remainder of this section, we will use the term actions also to refer to such *compound actions*, i.e., small sequences of atomic actions arising from the translation of a single ASLan++ statement.

Definition 12. Let $\mathcal{P}$ be a protocol and $E_1, \dots, E_n$ the entity names involved in $\mathcal{P}$. We denote with $I_e \equiv I_{e,1}, \dots, I_{e,m_e}$ the sequence of ASLan++ statements corresponding to the entity E_e .

Given a scenario $\mathcal{S}$, for each program instance $pi(j)$, we denote with $r_{e,1}^j, \dots, r_{e,m_e}^j$ the sequence of ASLan rules and with $w_{e,1}^j, \dots, w_{e,m_e}^j$, the sequence of SiL actions corresponding to I_e . We will call sequences of the last form *SiL action paths*⁴.

Finally, we define a *SiL action run* as a pair (π, ω) where $\pi = w_0, \dots, w_k$ is a SiL action path and $\omega = \varsigma_0, \dots, \varsigma_{k+1}$ is a sequence of data states such that $\langle a_j, \varsigma_j \rangle \Downarrow \varsigma_{j+1}$ for $0 \leq j \leq k$. $\square$

It is easy to see that given a program graph, each SiL path corresponds to a SiL action path (obtained by ignoring the locations in the SiL path, removing the X_i -conditionals and possibly grouping some consecutive atomic actions).

Definition 13. An *ASLan path* (for a protocol scenario $\mathcal{S}$) is a sequence $r_0, \dots, r_k$ of ASLan rules such that:

- for each entity E_e , program instance $pi(j)$ and $1 \leq l \leq m_e$, there is one and only one $0 \leq i \leq k$ such that $r_i \equiv r_{e,l}^j$;
- for $0 \leq i \leq k$, $r_i \equiv r_{e,l}^j$ for some e, l and j ;
- for $0 \leq i \leq k$, if $state_E(\dots, sl, \dots)$, where sl is the index referring to the step label, is in the left-hand side of $r_i \equiv r_{e,l}^j$ then either $sl = 1$ or there exists $h < i$ such that $state_E(\dots, sl, \dots)$ is in the right-hand side of r_h and $r_h \equiv r_{e,l-1}^j$. $\square$

The intuition behind this definition is that, given an ASLan transition system, the set of ASLan paths collects all the “potential” sequences of applications of ASLan rules, i.e., those admissible by only taking care of respecting the order given by the step labels inside the rules, no matter how the rest of the state evolves.

Definition 14. An *ASLan run* (for a protocol scenario $\mathcal{S}$) is a pair (τ, ρ) where τ is an ASLan path $r_0, \dots, r_k$ and $\rho = S_0, \dots, S_{k+1}$ is a sequence of ASLan states such that $S_i \xrightarrow{r_i} S_{i+1}$ for $0 \leq i \leq k$. $\square$

⁴Notice however that in a sequence $w_{e,1}^j, \dots, w_{e,m_e}^j$ defining a SiL action path, we ignore the conditionals with respect to X_i variables, i.e., those used in SiL to handle the interleaving between sessions.

Definition 15. We say that an ASLan path $r_0, \dots, r_k$ and a SiL action path $w_0, \dots, w_k$ are *equivalent* iff for each $0 \leq i \leq k$, r_i and w_i can be obtained as the translation of the same ASLan++ statement. $\square$

Lemma 2. Let $\mathcal{S}$ be a protocol scenario and $\mathcal{G}$ the corresponding program graph. Then there exists a SiL action path $w_0, \dots, w_k$ for $\mathcal{G}$ iff there exists an ASLan path $r_0, \dots, r_k$ for $\mathcal{S}$ and the paths are equivalent.

Proof. It is enough to observe that SiL action paths and ASLan paths follow, for a given program instance, the order in which the actions are executed in the protocol: this is obtained by the definition of the graph construction in the case of SiL, and by using step labels inside the rules in the case of ASLan. Furthermore, in both cases, each possible interleaving between sessions is admitted, i.e., whenever in a SiL path an action of the program instance $pi(i)$ is followed by an action of the program instance $pi(j)$, there is a corresponding possible choice for a next rule r to be applied in ASLan such that $r = r_{e,l}^j$ for some e and l ; vice versa, for each ASLan rule in an ASLan path letting one switch from a session i to a session j , there is a corresponding branch where $X_h = j$ giving rise to a corresponding SiL path. $\square$

Theorem 1. There exists a SiL action run (π, ω) of graph $\mathcal{G}$ corresponding to the protocol scenario $\mathcal{S}$, where $\omega = \varsigma_0, \dots, \varsigma_{k+1}$, iff there exists an ASLan run (τ, ρ) for $\mathcal{S}$, where $\rho = S_0, \dots, S_{k+1}$, and $\varsigma_i \sim S_i$ for $0 \leq i \leq k+1$.

Proof. Let ς_0 be the data state obtained after the initialization block of the SiL program graph and S_0 the ASLan initial state, as defined in Section 2. It is easy to check that $\varsigma_0 \sim S_0$. Then, the thesis follows by using Lemma 2 (for each SiL action path, there is an equivalent ASLan path) and Lemma 1 (equivalent steps preserve equivalence of states). $\square$

Finally, we can use the previous theorem to show that an attack state can be found in an ASLan path iff a goal location can be reached in the corresponding SiL path.

Corollary 1. Let $\mathcal{S}$ be a protocol scenario and $\mathcal{G}$ the corresponding program graph. An attack state can be found in an ASLan path for $\mathcal{S}$ iff a goal location can be reached in a SiL path for $\mathcal{G}$.

Proof. Let S be an ASLan attack state, i.e., $attack \subseteq \lceil S \rceil^H$. By Theorem 1, S is in an ASLan run for $\mathcal{S}$ iff there exists $\varsigma \sim S$ in a SiL run for $\mathcal{G}$. By Definition 10, $\varsigma(attack) = true$, i.e., a goal location referring to the given attack has been reached. $\square$

4 An interpolation-based algorithm for verification

In this section, we present the interpolation-based algorithm that we use for verification and describe, in particular, how we can calculate interpolants in our specific setting.

Our algorithm is a slightly simplified version of the IntraLA algorithm of [25], obtained by removing some fields only used there to deal with program procedures. In a nutshell, the idea underlying our algorithm is as follows. The input of our algorithm

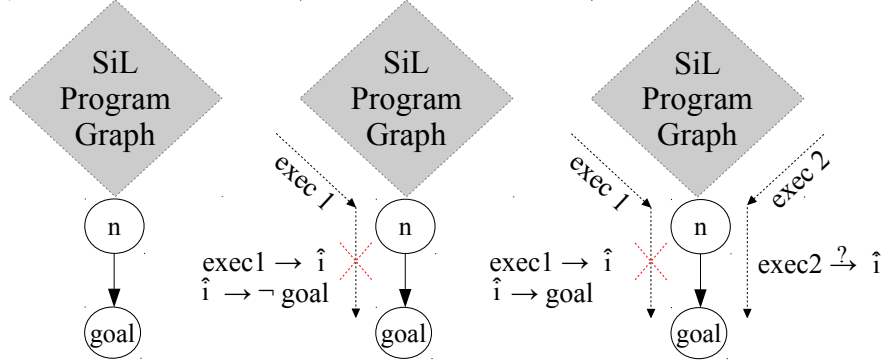


Figure 6: A SiL program graph (left). A first phase of symbolic execution with generation of an annotation (center). A second phase of symbolic execution with annotation check (right).

is a SiL program graph, as defined in Section 3.3, together with a set of attacks (goals) to search for; the output is either the proof that no attack has been found or an abstract attack trace for each attack found. The algorithm executes symbolically the program graph searching for given goal locations, which in our case represent attacks found on the given scenario of the protocol. In Fig. 6-left, we have depicted a simplified version of a generic program graph, highlighting a location n from which a path leading to a goal location starts. In the case when we fail to reach a goal during a search along an edge (Fig. 6-center), an annotation, i.e., a formula expressing a condition under which no goal can be reached, is produced by using Craig interpolation. Informally speaking, the annotation, $\hat{i}$ in the figure, will be a formula implied by (a formula describing the state originated by) the execution $exec1$ and inconsistent with (a formula describing the state reached at) the goal location. Through a backtrack phase, such an annotation is propagated to the preceding nodes of the edge and can be used to block a later phase of symbolic execution along an uninteresting run, i.e., a run for which the information contained in the annotation allows one to foresee that it will not reach a goal (Fig. 6-right).

4.1 Preliminary definitions

4.1.1 The annotation language

In what follows, it seems convenient to use a *two-sorted first-order language with equality*, in which the graph annotations will be expressed. The signature of the first sort is based on the algebra of messages defined in Section 2, over which we also allow a set of unary predicates DY_{IK}^j for $1 \leq j \leq n$ with a fixed $n \in \mathbb{N}$, whose meaning will be clarified below, and a ternary predicate *witness*. The signature of the second sort is based on a signature containing a set of variables (denoted in our examples by X possibly subscripted) and uninterpreted constants (for which we use integers as labels), and allows no functions and no predicates other than equality. We assume fixed the sets

of constants and denote by $\mathcal{L}(\mathcal{V})$ the set of well-formed formulas of such a two-sorted first-order language defined over a (also two-sorted) set $\mathcal{V}$ of variables, which we will instantiate with the concrete program variables of our SiL programs.

4.1.2 Symbolic execution notions

Before presenting the algorithm, we introduce some notions concerning symbolic execution. In the following, we will assume given a program graph (Λ, l_0, Δ) .

Definition 16. Let V be the set of program variables. A *symbolic data state* is a triple (P, C, E) , where P is a (again, two-sorted) *set of parameters*, i.e., variables not in V , $C \in \mathcal{L}(P)$ is a *constraint over the parameters*, and the *environment* E is a map from the program variables V to terms of the corresponding sort defined over P , where, in particular, that IK is mapped to a set of message terms and *witness* to a set of triples of message terms. We write Ξ to denote the set of symbolic data states. $\square$

Given its definition, a symbolic data state ξ can be characterized by the predicate

$$\chi(\xi) = C \wedge (\bigwedge_{v \in V \setminus \{IK\}} (v = E(v))) \wedge (\bigwedge_{m \in E(IK)} DY_{IK}^0(m)) \wedge (\bigwedge_{(m_1, m_2, m_3) \in E(witness)} witness(m_1, m_2, m_3)).$$

Note that the variable IK is treated in a particular way, i.e., we translate the fact that $E(IK) = M$ for some set M of parametric messages into a formula expressing that a predicate DY_{IK}^0 holds for all the messages in M .

A symbolic data state ξ can be associated to the set $\mathcal{E}(\xi)$ of data states produced by the map E for some valuation of the parameters satisfying the constraint C . We assume a defined initial symbolic data state $\mathcal{E}(\xi_0) = \{d_0\}$.

Definition 17. A *symbolic state* is a pair $(l, \xi) \in \Lambda \times \Xi$. A *symbolic interpreter* SI is a total map from the set $\mathcal{A}$ of SiL actions to $\Xi \times \Xi$ such that for each symbolic data state ξ and action a , $\cup \mathcal{E}(SI(a)(\xi)) = Sem(a)(\sigma(\xi))$. $\square$

Intuitively, SI takes a symbolic data state ξ and an action a and returns a non-empty set of symbolic data states, which represent the set of states obtained by executing the action a on ξ .

4.1.3 IntraLa basic notions

Definition 18. An *algorithm state* is a triple (Q, A, G) where Q is the set of *queries* (where a query is a symbolic state), A is a *program annotation* (or simply annotation, for short) and $G \subseteq \Lambda$ is the set of *goal locations* that have not been reached. $\square$

During the execution of the algorithm, the set of queries is used to keep track of which symbolic states still need to be considered, i.e., of those symbolic states whose location has at least one outgoing edge that has not been symbolically executed, and the annotation is a decoration of the graph used to prune the search. Formally:

Definition 19. A program annotation is a set of pairs in $(\Lambda \cup \Delta) \times \mathcal{L}(V)$. We will write these pairs in the form $l : \phi$ or $e : \phi$, where l is a location, e is an edge and ϕ is a formula called the *label*. When we have more than one label on a given location, we can read them as a disjunction of conditions: we define $A(l) = \bigvee \{\phi \mid l : \phi \in A\}$. $\square$

Definition 20. For an edge $e = (l_n, a, l_{n+1})$, the label $e : \phi$ is justified in A whenever starting from the precondition formula ϕ and by executing the action a , the postcondition produced is $A(l_{n+1})$, i.e., when it implies the annotation of l_{n+1} after executing a . In that case, we write $\mathcal{J}(e : \phi, A)$.

Let $Out(l)$ be the set of outgoing edges from a location l . The label $l : \phi$ is justified in A when, for all edges $e \in Out(l)$, there exists $e : \psi \in A$ such that ψ is a logical consequence of ϕ .

An annotation is justified when all its elements are justified. $\square$

A justified annotation is inductive and if it is initially true, then it is an inductive invariant. The algorithm maintains the invariant that A is always justified.

Definition 21. A query $q = (l, \xi)$ is blocked by a formula ϕ when $\xi \models \phi$ and we then write $Bloc(q, A(\phi))$.

With respect to q , the edge e is blocked when $Bloc(q, A(e))$ and the location l is blocked when $Bloc(q, A(l))$. $\square$

4.2 The rules of our algorithm

The rules of our algorithm are given in Fig. 7.

4.2.1 Initialization

The first rule applied is always *Init*, which initializes the algorithm state, i.e., the algorithm starts from the initial location, the initial symbolic data state, an empty annotation and a set G_0 of goals to search for, which is given as input together with the graph. After the application of *Init*, the rules *Decide*, *Learn* and *Conjoin* can be applied whenever their side-conditions are satisfied.

4.2.2 Symbolic execution steps

The *Decide* rule is used to perform symbolic execution. By symbolically executing one program action, it generates a new query from an existing one. It may choose any edge that is not blocked and any symbolic successor state generated by the action a . If the generated query is itself not blocked, it is added to the query set. In the definition of the rule, SI is a symbolic interpreter, l_n and ξ_n denote the currently considered location and symbolic data state, respectively, and l_{n+1} and ξ_{n+1} the location and symbolic data state obtained after executing the action a . The side conditions of the *Decide* rule require that, when we move from ξ_n to ξ_{n+1} , the first needs to be into the query set and the branch between the two nodes must exist and not be blocked (neither the edge nor the location).

4.2.3 Backtracking steps

During a backtracking phase, two rules are used:

- (i) *Learn*, which generates annotations; and

INITIALIZATION

$$\frac{}{\{(l_0, \xi_0)\}, \emptyset, G_0} \textit{Init}$$

SYMBOLIC EXECUTION STEPS

$$\frac{Q, A, G}{Q + (l_{n+1}, \xi_{n+1}), A, G} \textit{Decide} \quad \begin{array}{l} q = (l_n, \xi_n) \in Q \\ e = (l_n, a, l_{n+1}) \in \Delta \\ \neg \textit{Bloc}(q, A(e)) \\ \xi_{n+1} \in SI(a)(\xi_n) \\ \neg \textit{Bloc}((l_{n+1}, \xi_{n+1}), A(l_{n+1})) \end{array}$$

BACKTRACKING STEPS

$$\frac{Q, A, G}{Q, A + e : \phi, G} \textit{Learn} \quad \begin{array}{l} q = (l_n, \xi_n) \in Q \\ e = (l_n, a, l_{n+1}) \in \Delta \\ \textit{Bloc}(q, \phi) \\ \mathcal{J}(e : \phi, A) \end{array}$$

$$\frac{Q, A, G}{Q - q, A + l_n : \phi, G - l_n} \textit{Conjoin} \quad \begin{array}{l} q = (l_n, \xi) \in Q \\ \neg \textit{Bloc}(q, A(l_n)) \\ (\forall e \in \textit{Out}(l_n). \\ e : \phi_e \in A \wedge \textit{Bloc}(q, \phi_e)) \\ \phi = \bigwedge \{\phi_e \mid e \in \textit{Out}(l_n)\} \end{array}$$

Figure 7: Rules of the algorithm IntraLA with corresponding side conditions

(ii) *Conjoin*, which merges annotations coming from distinct branches.

If some outgoing edge $e = (l_n, a, l_{n+1})$ is not blocked, but every possible symbolic step along that edge leads to a blocked state, then the rule *Learn* is applied. Such a rule infers a new label ϕ that blocks the edge, where the formula ϕ can be any formula ϕ that both blocks the current query and is justified. In Section 4.3, we will explain how the formula ϕ can be obtained by exploiting the Craig interpolation lemma [13], which states that given two first-order formulas α and β such that $\alpha \wedge \beta$ is inconsistent, there exists a formula ϕ (their *interpolant*) such that α implies ϕ , ϕ implies $\neg\beta$ and $\phi \in \mathcal{L}(\alpha) \cap \mathcal{L}(\beta)$.

Finally, the rule *Conjoin* is used when all the outgoing edges of the location in a query q are blocked. The rule blocks the query q by labeling its location with the conjunction of the labels that block the outgoing edges. If the location is a goal, then we can remove it from the set of remaining goals. Finally, the query is discarded from the set q .

4.3 The generation of interpolants

We have seen in Section 4.2 that the rule *Learn* (Fig. 7) requires the generation of a formula ϕ that blocks the current query and is justified, to be used as an annotation. Here we will first introduce some notions concerning the description of data states and actions in our annotation language and then describe how to obtain the formula ϕ as an appropriate interpolant.

Let μ be a term, a formula, or a set of terms or of formulas. We write μ' for the result of adding one prime to all the non-logical symbols in μ . Intuitively, the prime is used to refer to the value of a same variable in a later step and it is used in *transition formulas*, i.e., formulas in $\mathcal{L}(V \cup V')$. Since the semantics of a SiL action (see Section 3.1) expresses how we move from a data state to another, we can easily associate to it a transition formula. In the following, we will write $Sem(a)$ to denote the transition formula corresponding to the action a .

In the context of our graphs, the most interesting case is when the action a is represented by a conditional statement, with a condition of the form $IK \vdash M$ for some message M , which intuitively means that the message M can be derived from a set of messages IK by using the rules of $\mathcal{N}_{DY}$ of Fig. 1. In our treatment, we fix a value n as the maximum number of inference steps that the intruder can execute in order to derive M . We observe that this is not a serious limitation of our method since several results (e.g., [34]) show that, when considering a finite number of sessions, as in our case, it is indeed possible to set an upper bound on the number of inference steps needed. Such a value can be established a-priori by observing the set of messages exchanged along the protocol scenario; we assume such an n to be fixed for the whole scenario.⁵

We use formulas of the form $DY_{IK}^j(M)$, for $0 \leq j \leq n$, with the intended meaning that M can be derived in n steps of inference by using the rules of $\mathcal{N}_{DY}$. In particular, the predicate DY_{IK}^0 is used to represent the initial knowledge IK , before any inference step is performed. Under the assumption on the n mentioned above, the statement $IK \vdash M$ can be expressed in our language as the formula $DY_{IK}^n(M)$.

The formula

$$\begin{aligned} \phi_j = \quad & \forall M. (DY_{IK}^{j+1}(M) \leftrightarrow (DY_{IK}^j(M) \\ & \vee (\exists M'. DY_{IK}^j([M, M']) \vee DY_{IK}^j([M', M]))) \\ & \vee (\exists M_1, M_2. M = [M_1, M_2] \wedge DY_{IK}^j(M_1) \wedge DY_{IK}^j(M_2)) \\ & \vee (\exists M_1, M_2. M = \{M_1\}_{M_2} \wedge DY_{IK}^j(M_1) \wedge DY_{IK}^j(M_2)) \\ & \vee (\exists M'. DY_{IK}^j(\{M\}_{M'}) \wedge DY_{IK}^j(inv(M'))) \\ & \vee (\exists M'. DY_{IK}^j(\{M\}_{inv(M')}) \wedge DY_{IK}^j(M')) \\ & \vee (\exists M_1, M_2. M = \{|M_1|\}_{M_2} \wedge DY_{IK}^j(M_1) \wedge DY_{IK}^j(M_2)) \\ & \vee (\exists M'. DY_{IK}^j(\{|M|\}_{M'}) \wedge DY_{IK}^j(M'))), \end{aligned}$$

in which $\leftrightarrow$ denotes the double implication and every quantification has to be intended

⁵The ability of the intruder of generating new messages can be simulated by enriching his initial knowledge with a set of constants not occurring elsewhere in the protocol specification. Since we consider finite scenarios, the size of such a set can also be bounded a-priori.

over the sort of messages, expresses (as a disjunction) all the ways in which a given message can be obtained by the intruder in one inference step, i.e., by a single application of one of the rules in the system $\mathcal{N}_{DY}$, thus moving from a knowledge (denoted by the predicate) DY_{IK}^j to a knowledge (denoted by the predicate) DY_{IK}^{j+1} .

A theory $\mathcal{T}_{Msg}(n)$ over the sort of messages is obtained by enriching classical first-order logic with equality with the axioms φ_j , for $1 \leq j < n$, together with an additional set of axioms that formalize that in the free algebra of messages any two distinct ground terms are not equal, e.g., $\forall M_1.M_2.M_3.M_4. ([M_1, M_2] = [M_3, M_4]) \supset (M_1 = M_3 \wedge M_2 = M_4)$.

Our translation of the program statement $IK \vdash M$ into the formula $DY_{IK}^n(M)$ is justified by the following result. This is proved by induction on the *height of a derivation* Π in the system $DY(IK)$, which is defined as the greatest number of successive applications of rules in Π .

Theorem 2. *Let M be a ground message, $n \in \mathbb{N}$, IK a set of ground messages and $\mathcal{I}$ an interpretation of $\mathcal{T}_{Msg}(n)$ such that $IK = \mathcal{I}(DY_{IK}^0)$. Then $\mathcal{I}$ satisfies the formula $DY_{IK}^n(M)$ iff there exists a derivation of $M \in DY(IK)$ of height at most $n + 1$ in the system $\mathcal{N}_{DY}$.*

Proof. ($\Rightarrow$) Assume that the interpretation $\mathcal{I}$ satisfies the formula $DY_{IK}^n(M)$, denoted $\mathcal{I} \models DY_{IK}^n(M)$. We proceed by induction on n . If $n = 0$, then we have $\mathcal{I} \models DY_{IK}^0(M)$, i.e., $M \in \mathcal{I}(DY_{IK}^0)$ which by hypothesis gives $M \in IK$. But then there exists a derivation in $\mathcal{N}_{DY}$ of $M \in DY(IK)$, obtained by a single application of the rule G_{axiom} . Now assume we have proved the assert for $n = j - 1$ and consider $n = j$. Since $\mathcal{I}$ satisfies the premise of the left-to-right implication in φ_{j-1} , i.e., $DY_{IK}^j(M)$, then it must also satisfy one of the disjuncts in the conclusion. We have a case for each disjunct. We consider two of them; the others are similar. (i) Let $\mathcal{I} \models DY_{IK}^{j-1}(M)$. By induction hypothesis, there exists a derivation of $M \in DY(IK)$ in $\mathcal{N}_{DY}$ of height at most j , which is the derivation we were looking for. (ii) Let $\mathcal{I} \models \exists M'. DY_{IK}^{j-1}([M, M']) \vee DY_{IK}^{j-1}([M', M])$. We can assume there exists a message M' such that $\mathcal{I} \models DY_{IK}^{j-1}([M, M'])$ (the other case is symmetrical). By induction hypothesis, there exists a derivation of $[M, M'] \in DY(IK)$ in $\mathcal{N}_{DY}$ of height at most j . A further application of A_{pair_i} gives a derivation of $M \in DY(IK)$ of height at most $j + 1$.

($\Leftarrow$) Again, we proceed by induction on n . If $n = 0$, the only admissible derivation of $M \in DY(IK)$ is the one given by an application of G_{axiom} . It follows that $M \in IK$. Then $IK = \mathcal{I}(DY_{IK}^0)$ implies $\mathcal{I} \models DY_{IK}^0(M)$. Now let us consider $n = j$ and assume we have a derivation of $M \in DY(IK)$ of length at most $j + 1$. Let r be the last rule applied. We have one case for each rule in $\mathcal{N}_{DY}$. Let r be G_{pair} . It follows that we have two derivations, of length at most j , of $M_1 \in DY(IK)$ and $M_2 \in DY(IK)$, respectively, where $M = [M_1, M_2]$. By induction hypothesis, we have $\mathcal{I} \models DY_{IK}^{j-1}(M_1)$ and $\mathcal{I} \models DY_{IK}^{j-1}(M_2)$, which implies that $\mathcal{I}$ satisfies one of the disjuncts in the premise of the right-to-left implication of φ_{j-1} . It follows that its conclusion must also be satisfied, i.e., $\mathcal{I} \models DY_{IK}^j(M)$. The other cases can be treated similarly. $\square$

Now let $\alpha = \chi(\xi_n)$ and $\beta = Sem(a) \wedge \neg A(l_{n+1})'$. Then we can obtain the formula ϕ we are looking for, during an application of the rule *Learn*, as an interpolant for α and β , possibly by using an interpolating theorem prover. With regard to this, we observe

that, in the presence of our finite scenario assumption, when mechanizing such a search, the problem can be simplified by restricting the domain to a finite set of messages.

4.4 Output and correctness of the algorithm

The algorithm terminates when no rules can be applied, which implies that the query set is empty. In [25], the correctness of the algorithm, with respect to the goal search, is proved: the proof given there applies straightforwardly for the slightly simplified version we have given here.

Theorem 3. *Let G_0 be the set of goal locations provided in input. If the algorithm terminates with the algorithm state (Q, A, G) , then all the locations in $G_0 \setminus G$ are reachable and all the locations in G are unreachable.*

The output of our method can be of two types. If no goal has been reached, then we have a proof of the fact that no attack can be found, with respect to the security property of interest, in the finite scenario that we are considering. Otherwise, for each goal location that has been found, we can generate an abstract attack trace. We also note that, by a trivial modification of the rule *Conjoin*, we can easily obtain an algorithm that keeps searching for a given goal even when this has already been reached through a different path, thus allowing for extracting more attack traces for the same goal on a given scenario.

Such traces can be inferred from the information deducible from the symbolic data state (P, C, E) corresponding to the last step of execution. We proceed as follows. First of all, we can reconstruct the order in which sessions have been interleaved. Such an information is completely contained in the value of the parameters corresponding to the variables X_j , for j an integer. The value of such parameters is specified in C . This allows for obtaining the sequence of messages exchanged, expressed in terms of program variables. Then, by using the maps in E , each such a variable can be associated to a function over the set of parameters P , and possibly further specified by the constraints over the parameters in C . It follows that the final result will be a sequence of messages where all the variables have been replaced by (functions over) parameters. Such a sequence constitutes our attack trace. In the case when the value of some parameter is not fully specified by the conditions in C , we have a parametrical attack trace, which can be instantiated in more than one way. A concrete example of this can be found in Example 6.

Example 6. *We continue our running example by showing the execution of the algorithm on some interesting paths of the graph defined in Section 3.2 for the protocol NSL: Table 1 summarizes the algorithm execution.*

*For readability, we have not reported the evolution of parameters and goals set. We remark that each new parameter is added to the parameters set once used and the goals set is initialized with the goal locations corresponding to the translation of the authentication goal *auth* (see Section 3.2 for details) but, given that no goal is reached, the goals set does not change during the execution of the algorithm. Note that in the table we use statements of the form $IK \vdash M$ in the constraint set as an abbreviation for the set of constraints over the parameters that make the (translation of the) statement*

Table 1: Execution of the algorithm on the program graph for the protocol NSL.

#	R	Query	Edge	Q	A	C	E
0	Init	(l_0, s_0)	-	l_0, ξ_0	$\emptyset$	$\emptyset$	$\emptyset$
1	Decide	(l_0, ξ_0)	(l_0, l_1)	$(l_0, \xi_0), (l_1, \xi_1)$	$\emptyset$	C_0	$E_0 \oplus \{(S2_Alice.B.y_0), (S2_Alice.Na.c_1), (S2_Alice.Actor.a), (S2_Alice.Y_0.y_0), (S2_Bob.Actor.b), (S1_Alice.Na.c_0), (S1_Alice.Actor.a), (S1_Alice.B.i), (IK, \{a, b, i, pk(a), pk(b), pk(i), inv(pk(i))\})\}$
2	Decide	(l_1, ξ_1)	(l_1, l_2)	$Q_1 \cup \{(l_2, \xi_2)\}$	$\emptyset$	C_1	$E_1 \oplus \{(Alice_1.Na.c_1), (IK, IK_1 \cup \{c_1, a\}_{pk(y_0)})\}$
3	Decide	(l_2, ξ_2)	(l_2, l_3)	$Q_2 \cup \{(l_3, \xi_3)\}$	$\emptyset$	$C_2 \cup \{(x_1 = 3)\}$	$E_2 \oplus \{(X_1, x_1)\}$
4	Decide	(l_3, ξ_3)	(l_3, l_4)	$Q_3 \cup \{(l_4, \xi_4)\}$	$\emptyset$	$C_3 \cup \{IK_2 \vdash \{(y_1, y_2)\}_{pk(b)}\}$	$E_3 \oplus \{(S2_Bob.Y_1.y_1), (S2_Bob.Y_2.y_2)\}$
5	Decide	(l_4, ξ_4)	(l_4, l_5)	$Q_4 \cup \{(l_5, \xi_5)\}$	$\emptyset$	C_4	$E_4 \oplus \{(S2_Bob.Na.y_1)\}$
6	Decide	(l_5, ξ_5)	(l_5, l_6)	$Q_5 \cup \{(l_6, \xi_6)\}$	$\emptyset$	C_5	$E_5 \oplus \{(S2_Bob.A.y_2)\}$
7	Decide	(l_6, ξ_6)	(l_6, l_7)	$Q_6 \cup \{(l_7, \xi_7)\}$	$\emptyset$	C_6	$E_6 \oplus \{(S2_Bob.Nb.c_2)\}$
8	Decide	(l_7, ξ_7)	(l_7, l_8)	$Q_7 \cup \{(l_8, \xi_8)\}$	$\emptyset$	C_7	$E_7 \oplus \{(IK, IK_7 \cup \{y_1, c_2, b\}_{pk(y_2)})\}$
9	Decide	(l_8, ξ_8)	(l_8, l_9)	$Q_8 \cup \{(l_9, \xi_9)\}$	$\emptyset$	$C_8 \cup \{(x_{11} = 2)\}$	$E_8 \oplus \{(X_{11}, x_{11})\}$
10	Decide	(l_9, ξ_9)	(l_9, l_{10})	$Q_9 \cup \{(l_{10}, \xi_{10})\}$	$\emptyset$	$C_9 \cup \{IK_8 \vdash \{c_1, y_1, y_0\}_{pk(a)}\}$	E_9
11	Decide	(l_{10}, ξ_{10})	(l_{10}, l_{11})	$Q_{10} \cup \{(l_{11}, \xi_{11})\}$	$\emptyset$	C_{10}	$E_{10} \oplus \{(S2_Alice.Nb.y_3), (S2_Alice.Y_3.y_3)\}$
12	Decide	(l_{11}, ξ_{11})	(l_{11}, l_{12})	$Q_{11} \cup \{(l_{12}, \xi_{12})\}$	$\emptyset$	C_{11}	$E_{11} \oplus \{(IK, IK_{10} \cup \{y_3\}_{pk(y_0)})\}$
13	Decide	(l_{12}, ξ_{12})	(l_{12}, l_{13})	$Q_{12} \cup \{(l_{13}, \xi_{13})\}$	$\emptyset$	$C_{12} \cup \{witness(a, y_0, \{y_3\}_{pk(y_2)})\}$	E_{12}
14	Decide	(l_{13}, ξ_{13})	(l_{13}, l_{14})	$Q_{13} \cup \{(l_{14}, \xi_{14})\}$	$\emptyset$	$C_{13} \cup \{(x_9 = 1)\}$	$E_{13} \oplus \{(X_9, x_9)\}$
15	Decide	(l_{14}, ξ_{14})	(l_{14}, l_{15})	$Q_{14} \cup \{(l_{15}, \xi_{15})\}$	$\emptyset$	$C_{14} \cup \{IK \vdash \{c_0, y_4, i\}_{pk(a)}\}$	$E_{14} \oplus \{(S1_Alice.Y_1.y_4)\}$
16	Decide	(l_{15}, ξ_{15})	(l_{15}, l_{16})	$Q_{15} \cup \{(l_{16}, \xi_{16})\}$	$\emptyset$	C_{15}	$E_{15} \oplus \{(S1_Alice.Nb.y_4)\}$
17	Decide	(l_{16}, ξ_{16})	(l_{16}, l_{17})	$Q_{16} \cup \{(l_{17}, \xi_{17})\}$	$\emptyset$	C_{16}	$E_{16} \oplus \{(IK, IK_{15} \cup \{y_4\}_{pk(i)})\}$
18	Decide	(l_{17}, ξ_{17})	(l_{17}, l_{18})	$Q_{17} \cup \{(l_{18}, \xi_{18})\}$	$\emptyset$	$C_{17} \cup \{witness(a, i, \{y_4\}_{pk(i)})\}$	E_{17}
19	Decide	(l_{18}, ξ_{18})	(l_{18}, l_{19})	$Q_{18} \cup \{(l_{19}, \xi_{19})\}$	$\emptyset$	$C_{18} \cup \{c_2\}_{pk(b)}$	E_{18}
20	Learn	(l_{19}, ξ_{19})	-	Q_{19}	$(l_{19}, l_{20}) : S2_Bob.A = i$	C_{19}	E_{19}
21	Conjoin	(l_{19}, ξ_{19})	(l_{19}, l_{20})	Q_{18}	$A_{20} \cup \{l_{19} : S2_Bob.A = i\}$	C_{20}	E_{20}
22	Learn	(l_{18}, ξ_{18})	-	Q_{18}	$A_{21} \cup \{l_{18}, l_{19} : S2_Bob.A = i\}$	C_{21}	E_{21}
23	Conjoin	(l_{18}, ξ_{18})	(l_{18}, l_{19})	Q_{17}	$A_{22} \cup \{l_{18} : S2_Bob.A = i\}$	C_{22}	E_{22}
...							
33	Conjoin	(l_{14}, ξ_{14})	(l_{14}, l_{15})	Q_{27}	$A_{32} \cup \{l_{14} : S2_Bob.A = i\}$	C_{32}	E_{32}

satisfiable, according to the definition above. Q_i , C_i and E_i denote, respectively, the set of queries, the set of constraints and the environment at step i of the execution. We have also used $\#$ to indicate the step number and R to indicate which rule is applied using the first letter of the rule.

The first path we show (summarized by the Message Sequence Chart (MSC) in Fig. 8) reaches a goal location with an unsatisfiable state and then annotates it with an interpolant, while the other ones reach the previously annotated path and then block their executions (thus saving some execution steps). The algorithm starts, as described in Table 1, by using the *Init* rule to initialize the algorithm state and then it symbolically executes the program graph from query (l_0, ξ_0) to (l_{18}, ξ_{18}) using the *Decide* rule (steps 0–19). In step 20, the algorithm blocks its symbolic execution because the edge (l_{19}, l_{20}) is labeled with the goal action for an authentication goal and any possible symbolic execution step leads to a blocked symbolic data state. The backtrack phase starts and, until step 33, the algorithm creates interpolants to annotate the program graph and then it propagates annotations up to the location l_{14} (where the symbolic execution restarts with the *Decide* rule but we have not reported it in Table 1 for lack of space).

As shown in Fig. 9, there are other two paths that reach location l_{18} . Each path that reaches this location has already executed an action of the form $IK \vdash \{N_A, N_B, B\}_{pk(A)}$. As described in [21], it is impossible for the *DY* intruder to create a message of the form $\{N_A, N_B, B\}_{pk(A)}$ from its knowledge (*IK*) if the intruder is not explicitly playing the role of the sender, i.e., *A*. This means that each symbolic state that reaches location l_{18} implies the interpolant $S2_Bob.A = i$. This is a concrete example of how the annotation method can help (and improve) the search procedure: in *NSL* we can stop following every path that reaches location l_{18} as the annotation method ensures that we will never reach a goal location.

While with *NSL* the algorithm concludes with no attacks found, if we consider the original protocol *NSPK* (i.e., remove Lowe’s addition of “*B*” in the second message of the protocol), then our method reaches the goal location with an execution close to the one we have just provided. In fact, in *NSPK*, when we compute the step after the 19th, the intruder rules lead to the goal with the inequality $S2_Bob.A \neq i$. This is because the intruder *i* can perform a man-in-the-middle attack using the initiator entity of the first session in order to decrypt the messages that the receiver sends to *i* in the second one [21]. To show the attack trace, we first check the path that is used during the algorithm execution to reach the goal location and that is represented by the values of X_j parameters contained in the C_{19} set. In this case, $\{X_{11} = 2, X_9 = 1\} \subseteq C_{19}$, which produces the symbolic attack trace (at state 19 of the algorithm execution) shown in the middle of Fig. 2.

Now, by using the information in Ξ_{19} , we can instantiate this trace using parameter and constant values, and thus obtain the instantiated attack trace shown on the right of Fig. 2. We can note from IK_{19} that Y_2 has no constraints on the fact that it has to be *i*, i.e., the intruder acts as if it were an honest agent (under his real name) in the first session, and then we write the concretization as $i(a)$ to show that the intruder is acting as the honest agent *a* in the second session and this makes possible the man-in-the-middle-attack.

It is also not difficult to extract from this instantiated attack trace a test case, which can then be applied to test the actual protocol implementation. In fact, the constraint

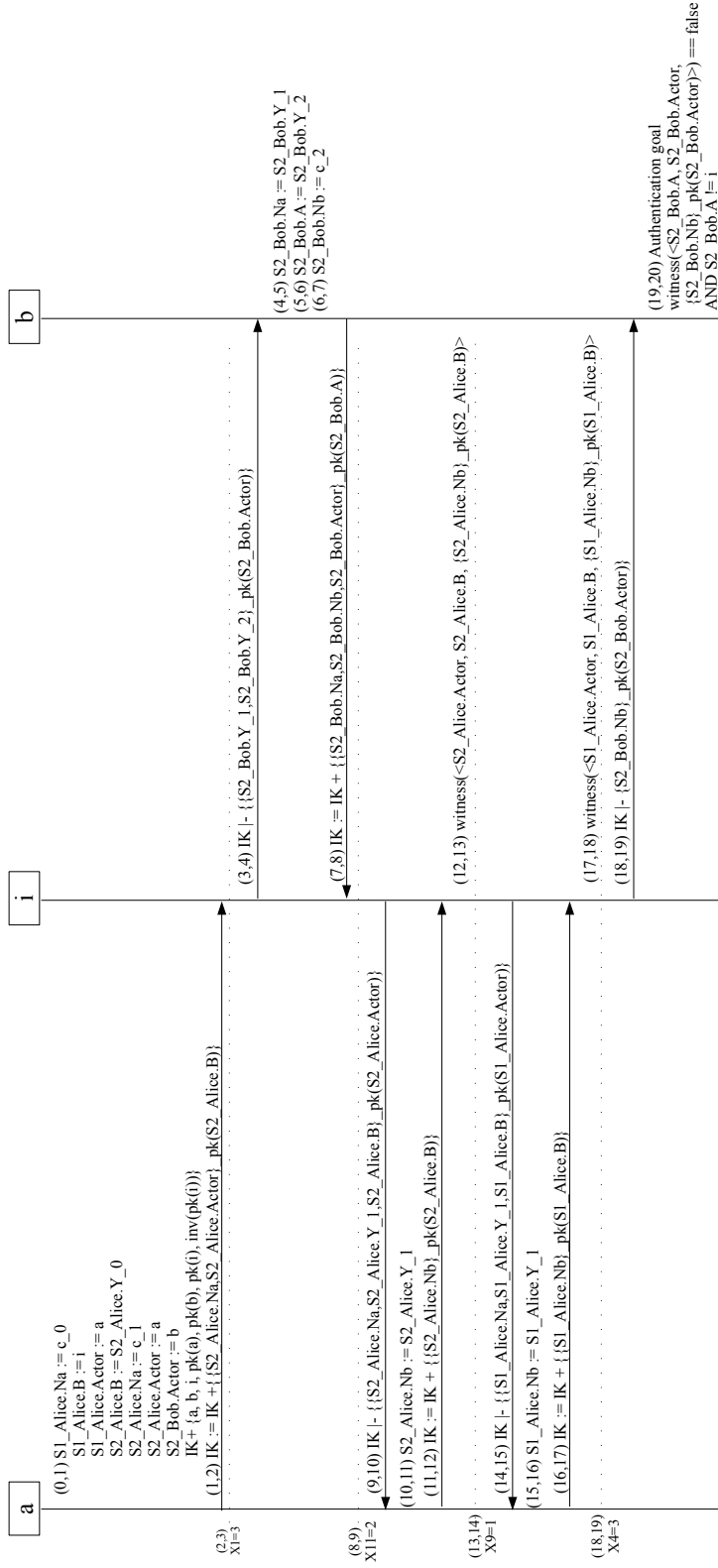


Figure 8: NSL execution path.

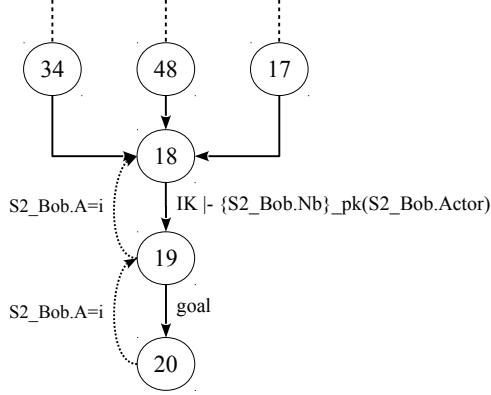


Figure 9: NSL sub-graph.

set contains a sequence of equalities of the form $X_i = n$, which specify the session to be followed at each branch of the executed path. $\square$

5 The SPiM tool

In order to show that our method concretely speeds up the validation, we have implemented a Java prototype called *SPiM* (*Security Protocol interpolation Method*), which is available at <http://regis.di.univr.it/spim.php>. As shown in Fig. 10, SPiM takes an ASLan++ specification as input that is automatically translated into a SiL program graph by the translator *ASLan++2SiL*. The program graph is then given as input to the *Verification Engine* (VE), which verifies the protocol by searching for goal locations that represent attacks on the protocol. The VE is composed by three main components:

- (i) a *quantifier elimination* module,
- (ii) *DY intruder* and *EUf* (*Equalities and Uninterpreted Functions*) theories and
- (iii) the tools Z3 [31] and iZ3 [26], used for SAT solving and interpolant generation, respectively.

Both Z3 and iZ3 are invoked by *SPiA* (*Security Protocol interpolation Algorithm*), which is our implementation of the algorithm in Section 4. Quantifier elimination and the definition of theories are related to the usage of Z3 and iZ3. In fact, as shown in Section 4, our algorithm needs to handle many quantifications and, for performance issues, a module that unfolds each quantifier over the finite set of possible messages has been developed. Moreover, the DY theory has been properly axiomatized (with respect to each formula produced by SPiA) in Z3 and iZ3, which do not support it by default.

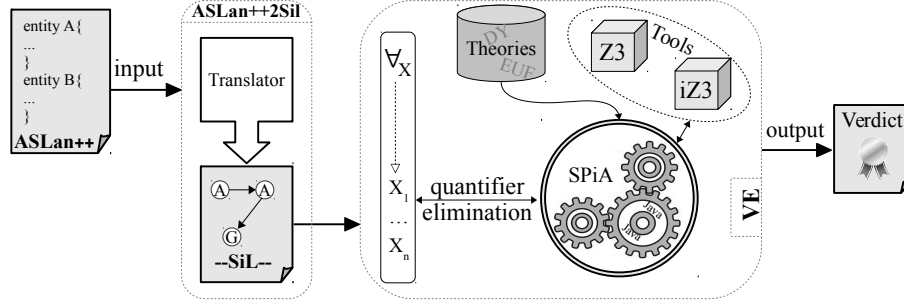


Figure 10: The SPiM tool.

More specifically, the VE symbolically executes a program graph. After the execution of an action branching from a node to the next one, it produces a formula, which represents the symbolic state reached. Z3 is then used for a satisfiability check on the newly produced formula. When the symbolic execution of a given path fails to reach a goal, the VE calls iZ3, which generates an annotation (i.e., a formula expressing a condition under which no goal can be reached) by using Craig’s interpolation. By a backtracking phase, SPiA propagates the annotation through the program graph. Such an annotation is possibly used to block a later phase of symbolic execution along an uninteresting run, as explained in Section 4. SPiM concludes reporting either all the different reachable attack states (from which abstract attack traces can be extracted) or that no attack has been found for the given specification.

5.1 Experiments and results

We considered 7 case studies and compared the results obtained by using interpolation-driven exploration (SPiA) and full exploration (*Full-explore*) of the program graph. Full-explore explores the entire graph checking, at each step, if the state is satisfiable or not. If there is an inconsistency, SPiM blocks the execution of the path resuming from the first unexplored path, until it has explored all paths.⁶

Table 2 shows the results obtained (with a general purpose computer), by making explicit the time required for symbolic execution steps (applications of *Decide*) and for interpolant generation (applications of *Learn*). The usage of SPiA has allowed us to speed up the validation (in the context of security protocols, i.e, using the DY intruder) by (i) reducing the number of states to explore and then (ii) lowering the execution time. The relation between (i) and (ii) is due to the fact that the time needed to perform a *Decide* is comparable to the one required to perform a *Learn*, and the time used to propagate the annotations (*Conjoin* rule) is negligible. For example, the time needed to symbolically execute a (sub-)path twice, using Full-explore, is comparable to the

⁶It would be possible to modify the Full-explore algorithm and check for inconsistencies at the end of the path instead of at any step but this would lead to an unfair comparison. In fact, a similar improvement could have been implemented also for SPiM, but then it would have been difficult to distinguish between the steps pruned by interpolation and those pruned by such an improvement.

Table 2: SPiA vs Full-explore.

Specification (sessions)	SPiA: Decide+Learn (time)	Full-explore: Decide (time)	Speedup %	Result
ISO6 (ab,ab)	311+274 (205m6s)	467* (278m12s)	-26.28 %	no attack found
NSL (ab,ab)	257+234 (57m37s)	631 (173m7s)	-66.71 %	no attack found
NSL (ai,ab)	89+22 (1m30s)	119 (1m49)	-17.43 %	no attack found
NSPK (ab,ab)	257+234 (26m5s)	631 (76m20s)	-65.82 %	no attack found
NSPK (ai,ab)	101+22 (0m56s)	123 (0m51s)	+8.92 %	attack found
Helsinki (ab,ab)	311+274 (112m7s)	660* (261m47s)	-57.17 %	no attack found
Helsinki (ai,ab)	167+88 (13m41)	407 (46m44s)	-70.72 %	attack found

time used to execute and annotate the same (sub-)path. But from that point on, if the annotation blocks the execution, only the Full-explore will execute that (sub-)graph again. We have observed that, in the case studies analyzed, the annotations block the executions of all those (sub-)paths that do not reach a goal location, thus ensuring a clear improvement of the performances. In particular, when applying a *Decide* moving from a node l_1 to a node l_2 , we generate a formula, which describes the state of the execution at node l_2 and the axiomatization of the DY theory; this formula is then given to Z3 that “decides” whether it is satisfiable or not. On the other hand, in order to execute a *Learn* between the same S_1 and S_2 , we translate the state S_1 with the axiomatized DY theory into a formula α and the semantics of the action a together with all previous annotations into a formula β . In order to find an interpolant we use iZ3 that performs a satisfiability check on the formula $\alpha \wedge \beta$ (very similar to what a *Decide* would do) and from the refutation by resolution steps an interpolant can be calculated in linear time [22, 23]. Finally, the *Conjoin* rule propagates these interpolants without performing other satisfiability checks.

Empirically, the more the program graph grows the more the annotations prune the search space. This is due to the fact that the number of states pruned by interpolation is usually related to the size of a program graph; this is confirmed by the results in Table 2 and, in particular, by the case studies for which Full-explore has not concluded the execution (marked with an asterisk).

We have also compared the SPiM tool with the three state-of-the-art model checkers for security protocols that are part of the AVANTSSAR platform [4]: CL-AtSe [36], OFMC [8] and SATMC [1].⁷ Not surprisingly, Table 3 shows that their average computational times of execution are in general better than ours. This is mainly due to several speed-up techniques implemented by these model checkers and to empirical conditions that can stop the execution (both not implemented yet in SPiM). In Table 3, we have also reported the number of transitions and/or nodes reached during the validations with the exception of SATMC, which does not report them as output. However, for each safe specification (in which no attacks are found), SATMC reached the maximum

⁷For this comparison, given that all these tools support ASLan++, we have used the same input files and we have also used the same general purpose computer used to generate the results in Table 2. We have considered OFMC v2012c, which is the last version that supports ASLan++ although it only supports untyped analysis, while for SATMC and CL-AtSe we have considered versions 3.4 and 2.5-21, respectively, which support typed analysis as SPiM does. Note that the times shown in Table 2 also consider the translation from ASLan++ to SiL program graph (usually several seconds), while in Table 3 we do not show the translation time from ASLan++ to ASLan (the input supported by the three tools), which is usually less than one second.

Table 3: SATMC, CL-AtSe and OFMC.

Specification (sessions)	SATMC (v.3.4)	CL-AtSe (v.2.5-21)			OFMC (v.2012c)		Result
	time	transitions	states	time	nodes	time	
ISO6 (ab,ab)	6.318s	452	236	0.034s	8432	3.804s	no attack found
NSL (ab,ab)	14m28s	794	534	0.052s	3236	3.295s	no attack found
NSL (ai,ab)	6m51s	93	69	0.015s	575	0.327s	no attack found
NSPK (ab,ab)	14m10s	794	534	0.053s	8180	3.208s	no attack found
NSPK (ai,ab)	1m56s	14	10	0.014s	96	0.134s	attack found
Helsinki (ab,ab)	7.01s	794	534	0.061s	8180	3.795s	no attack found
Helsinki (ai,ab)	50.8s	14	10	0.017s	96	0.121s	attack found

number of steps (80) permitted as default and the reported timings are comparable to those obtained by SPiM for some specifications; in the case when they are not comparable, it is interesting to observe that SPiM executes a number of rules much higher than 80. For both CL-AtSe and OFMC, on safe specifications, the number of transitions and nodes explored is, in most cases, higher than the number of rules (transitions) of SPiM (Table 2). On unsafe specifications (where an attack is found), these numbers seem to be in disfavor of SPiM but this is because SATMC, OFMC and CL-AtSe stop their executions once a goal is found, while SPiM searches for every possible attack trace in the program graph (i.e., SPiM features a multi-attack-trace support).

We remark that the aim of SPiM is mainly to show that Craig’s interpolation can be used as a speed-up technique also in the context of security protocols and not (yet) to propose an *efficient* implementation of a model checker for security protocol verification. In fact, we do not see our approach as an alternative to such more mature and widespread tools, but we actually expect some interesting and useful interaction. For example, CL-AtSe implements many optimizations, like simplification and rewriting of input specifications, and OFMC implements some optimizations at the intruder level as well as a specific technique, called *constraint differentiation* (CDiff), which considerably prunes the state space (it is more or less the equivalent of partial-order reduction techniques typical of model checking, where the reduction is “pushed” to the constraint solving procedure). Moreover, both CL-AtSe and OFMC implement the *step compression* and *protocol simplifications* techniques, which merge together some of the actions performed in the protocol.

We do not see any incompatibility in using interpolation together with such optimization techniques. For instance, CDiff prunes the state space by not considering the same state twice, whereas interpolation works on reducing the search space by excluding some paths during the analysis (i.e., it prunes the execution of some of the paths). Moreover, based on the idea that the intruder controls the network, when the intruder sends a message ($IK \vdash M$) to an honest agent and the honest agent sends back a reply ($IK := IK + \{M\}$), step compression merges the two into a single step. This would reduce the state space but not prevent SPiM to generate and use interpolants.

The only possible side effect that we foresee in using interpolation together with such optimization techniques is that the number of paths pruned by interpolation could decrease when we use it in combination with other techniques. In general, however, although we do not have experimental evidence yet, we expect that if enhanced with such techniques, SPiM could then reach even higher speed-up rates. We are currently

working in this direction.

6 Related work

To the best of our knowledge, there is no other tool for security protocol analysis that uses a speed-up technique based on Craig’s interpolation. We now discuss some further related work on interpolation, in addition to the works we already considered in detail in the other sections of the paper.

In [25], McMillan presented the IntraLA algorithm that we have used as a basis for this work. However, our application field is network security whereas IntraLA has been developed for software verification, and this has led to a number of substantial differences between the two works. First of all, our case studies are security protocols, and thus parallel programs, whereas IntraLA works on sequential ones. For this reason, we have defined a simple programming language (SiL) with some protocol-oriented features and provided a translation procedure from protocol specifications (expressed in ASLan++) into SiL programs (proving the correctness of the translation with respect to the semantics of ASLan++). In particular, given the object of our study, SiL allows to express statements aimed at handling the actions of the DY intruder. The DY theory has then been used both in the symbolic execution of a program graph (Decide rule, Section 4) and for interpolants generation (Learn rule, Section 4). The nature of the goals that we verify also differ from the ones in [25], as they are directly related to security goals like authentication and integrity. The same differences can be found between SPiM and IMPACT II (the implementation of [25]): IMPACT II takes as input control flow graphs from C programs and has been tested on the source codes of drivers. The algorithm implementations do also have some differences. In particular, in SPiA, we have implemented an optimization according to which an interpolant is calculated, at a given node or edge, only when the graph presents an unexplored path that can be blocked by such an interpolant.

Recently, McMillan has proposed in [27] a variation of IntraLA that mainly adapts IntraLA to large-block encoding (LBE). This technique reduces the abstract reachability tree used by the IntraLA algorithm, for example by simplifying the tree produced from very long sequences of if statements. Moving from original trees to the ones produced with LBE is not a trivial task and requires further investigation. Introducing LBE could speed up our tool too but, as we have already discussed in Section 5.1, we implemented SPiM mainly to show that interpolation can concretely be used as a speed-up technique together with the DY intruder model in the context of security protocols. Other works by McMillan that exploit the use of Craig interpolation in model checking are [22, 24], but interpolants are used there in a different way, i.e., to apply interpolant-based image approximation.

Besides for McMillan’s works on interpolation applied to model checking, there are a number of model checkers that implement different techniques to speed-up the search for goal locations. In particular, for the purpose of the comparison with SPiM and in addition to the tools already considered in Section 5.1, we consider here four security protocol analysis tools that implement the DY intruder theory: Maude-NPA [18], ProVerif [9], Scyther [14] and Tamarin [35].

Besides for DY, Maude-NPA supports a wide range of theories such as the “associative-commutative plus identity” theory. Maude-NPA has been implemented with particular focus on performances and in fact, during the analysis, it takes advantage of various state-space reduction techniques. These range from a modified version of the lazy intruder (called “super lazy intruder”) to a partial-order reduction technique. The ideas behind the speed-up techniques of Maude-NPA are very similar to the ones of SPiM: reduce the number of states to explore and try to not explore a state after having the evidence that from this state the model checker will never reach the goal location (i.e., will never reach the initial state given that Maude-NPA performs a backward reachability search). As for all the back-ends of the AVANTSSAR Platform (discussed in Section 5.1), in principle we do not see any incompatibility in combining the interpolation-based technique we have proposed in this paper with the speed-up techniques implemented in Maude-NPA. However, Maude-NPA performs backward reachability analysis whereas our technique has been defined for forward reachability analysis. This does not prevent possible useful interaction between the two approaches but it might require a non-trivial adaptation of the interpolation-based algorithm.

In ProVerif, security protocols are represented using Prolog rules in order to handle multiple executions. It implements an efficient algorithm that, combined with a unification technique along with rule optimization procedures, handles the problem of state-space explosion. Due to the particular nature of the techniques it implements, it is not clear if ProVerif could further improve its performance by integrating an interpolation-based technique.

Scyther uses a pattern-refinement algorithm that provides concise representations of (infinite) sets of traces. It does not use approximation methods nor abstraction techniques and it could thus benefit from including our technique, in particular, when unbounded verification is performed. However, as for Maude-NPA, due to Scyther’s backward searching algorithm, this integration would require further study.

Tamarin uses a constraint-solving algorithm and a symbolic representation of states like SPiM, but supports analysis for an unbounded number of protocol sessions. Intruder capabilities and protocols are specified jointly as a set of (labeled) multiset rewriting rules. Tamarin is particularly well suited for the analysis of protocols that use the Diffie-Hellman key exchange, which SPiM does not handle. One of the main difficulties one might have in implementing our speed-up technique in Tamarin is thus with the Diffie-Hellman key representation. However, since Tamarin uses a (labeled) operational semantics that is similar to the one used in SPiM, it might still be feasible to adapt the interpolation technique successfully.

7 Concluding remarks

We believe that our interpolation-based method, together with its prototype implementation in the SPiM tool and our experimental evaluation, shows that we can indeed use interpolation to reduce the search space and speed up the execution also in the case of security protocol verification. In particular, as we have shown, we can use a standard security protocol specification language (ASLan++, but, we believe that with little effort, also other languages that specify the different protocol roles as interacting pro-

cesses could be used) and translate automatically into SPiM’s input language SiL with the guarantee that in doing so we will not introduce nor lose any attack. The tool then proceeds automatically and concludes reporting either all the different reachable states (from which one or more abstract attack traces can be extracted) or that no attack has been found for the given specification.

As future work, we plan to increment our experimental results by considering further (and more complex) security protocols, such as those described in [10] and in the standard literature. This will allow us to collect further evidence on to what extent interpolation can indeed increase the performance of SPiM.

More importantly, as we remarked above, we are not aware of any other tool for security protocol verification that uses an interpolation-based speed-up technique, and we believe that actually interpolation might be proficiently used in addition (and not in alternative) to other optimization techniques for security protocol verification. We are thus currently investigating possible useful interactions between interpolation and such optimization techniques, given that there are no theoretical or technical incompatibilities between them. This will allow us to enhance SPiM and promote its performance closer to the level of the more mature tools. Symmetrically, it would be interesting to investigate also whether such mature tools might benefit from the integration of interpolation-based techniques such as ours to provide an additional boost to their performance. This will of course be a much more challenging endeavor to undertake, as it will possibly require some internal changes to already deployed tools, but given our close scientific relations to some of the tool developers, we are hopeful that we will be able to carry out some attempts in this direction.

References

- [1] Armando A, Compagna L (2004) SATMC: a SAT-based model checker for security protocols. In: JELIA, Springer, LNAI 3229, pp 730–733
- [2] Armando A, Basin D, Boichut Y, Chevalier Y, Compagna L, Cuellar J, Hanks Drielsma P, Héam PC, Mantovani J, Mödersheim S, von Oheimb D, Rusinowitch M, Santiago J, Turuani M, Viganò L, Vigneron L (2005) The AVISPA Tool for the Automated Validation of Internet Security Protocols and Applications. In: CAV, Springer, LNCS 3576, pp 281–285, URL <http://www.avispa-project.org>
- [3] Armando A, Carbone R, Compagna L, Cuéllar J, Tobarra Abad L (2008) Formal Analysis of SAML 2.0 Web Browser Single Sign-On: Breaking the SAML-based Single Sign-On for Google Apps. In: FMSE, ACM Press
- [4] Armando A, Arsac W, Avanesov T, Barletta M, Calvi A, Cappai A, Carbone R, Chevalier Y, Compagna L, Cuéllar J, Erzse G, Frau S, Minea M, Mödersheim S, von Oheimb D, Pellegrino G, Ponta SE, Rocchetto M, Rusinowitch M, Torabi Dashti M, Turuani M, Viganò L (2012) The AVANTSSAR Platform for the Automated Validation of Trust and Security of Service-Oriented

- Architectures. In: TACAS, Springer, LNCS 7214, pp 267–282, URL <http://www.avantssar.eu>
- [5] Armando A, Pellegrino G, Carbone R, Merlo A, Balzarotti D (2012) From Model-Checking to Automated Testing of Security Protocols: Bridging the Gap. In: TAP, Springer, LNCS 7305, pp 3–18
 - [6] Armando A, Basin D, Blome A, Büchler M, Compagna L, Costa G, Cuéllar J, Groz R, Guardini D, Hossen K, Li K, Mihancea PF, Minea M, Ochoa M, von Oheimb D, Pellegrino G, Peroli M, Petric Maretic G, Pretschner A, Torabi Dashti M, Viganò L (2014) The SPaCioS Tool for the Secure Provision and Consumption in the Internet of Services, submitted (more information on SPaCioS: <http://www.spacios.eu>)
 - [7] AVANTSSAR (2011) Deliverable 2.3 (update): ASLan++ specification and tutorial. Available at <http://www.avantssar.eu>
 - [8] Basin D, Mödersheim S, Viganò L (2005) OFMC: A symbolic model checker for security protocols. *International Journal of Information Security* 4(3):181–208
 - [9] Blanchet B (2001) An efficient cryptographic protocol verifier based on prolog rules. In: CSFW, IEEE CS, pp 82–96
 - [10] Boyd C, Mathuria A (2010) *Protocols for Authentication and Key Establishment*. Springer
 - [11] Büchler M, Oudinet J, Pretschner A (2011) Security mutants for property-based testing. In: TAP, Springer, LNCS 6706, pp 69–77
 - [12] Cortier V, Delaune S, Lafourcade P (2006) A survey of algebraic properties used in cryptographic protocols. *Journal of Computer Security* 1:1–43
 - [13] Craig W (1957) Three Uses of the Herbrand-Gentzen Theorem in Relating Model Theory and Proof Theory. *The Journal of Symbolic Logic* 22(3):pp. 269–285
 - [14] Cremers C (2008) The Scyther Tool: Verification, falsification, and analysis of security protocols. In: *Proceedings of CAV 2008*, Springer, LNCS 5123, pp 414–418
 - [15] Cremers C, Mauw S (2012) *Operational Semantics and Verification of Security Protocols*. Springer
 - [16] Dolev D, Yao A (1983) On the Security of Public-Key Protocols. *IEEE Transactions on Information Theory* 2(29)
 - [17] Escobar S, Meadows C, Meseguer J (2007) Maude-NPA: Cryptographic Protocol Analysis Modulo Equational Properties. In: *FOSAD*, Springer, pp 1–50
 - [18] Escobar S, Meadows C, Meseguer J, Santiago S (2014) State space reduction in the Maude-NRL protocol analyzer. *Information and Computation* 238:157–186

- [19] Kahn G (1987) Natural semantics. In: STACS, 4th Annual Symposium, pp 22–39, doi: 10.1007/BFb0039592, URL <http://dx.doi.org/10.1007/BFb0039592>
- [20] King JC (1976) Symbolic execution and program testing. CACM 19(7):385–394
- [21] Lowe G (1996) Breaking and Fixing the Needham-Shroeder Public-Key Protocol Using FDR. In: TACAS, Springer, LNCS 1055, pp 147–166
- [22] McMillan KL (2004) An interpolating theorem prover. In: Tools and Algorithms for the Construction and Analysis of Systems, Springer, pp 16–30
- [23] McMillan KL (2005) Applications of Craig Interpolants in Model Checking. In: TACAS, LNCS 3440, Springer, pp 1–12
- [24] McMillan KL (2006) Lazy abstraction with interpolants. In: CAV, 18th International Conference, Springer, pp 123–136
- [25] McMillan KL (2010) Lazy annotation for program testing and verification. In: CAV, LNCS 6174, Springer, pp 104–118
- [26] McMillan KL (2011) Interpolants from Z3 proofs. In: FMCAD, pp 19–27
- [27] McMillan KL (2014) Lazy annotation revisited. In: Computer Aided Verification - 26th International Conference, pp 243–259
- [28] Mödersheim S (2009) Algebraic Properties in Alice and Bob Notation. In: Ares, IEEE CS, pp 433–440
- [29] Mödersheim S, Viganò L (2009) Secure Pseudonymous Channels. In: Esorics, Springer, LNCS 5789, pp 337–354
- [30] Mödersheim S, Viganò L (2009) The Open-Source Fixed-Point Model Checker for Symbolic Analysis of Security Protocols. In: FOSAD 2008/2009, LNCS 5705, Springer, pp 166–194
- [31] de Moura L, Bjorner N (2008) Z3: An Efficient SMT Solver. In: TACAS, LNCS 4963, Springer, pp 337–340
- [32] von Oheimb D, Mödersheim S (2010) ASLan++ — a formal security specification language for distributed systems. In: FMCO, Springer, LNCS 6957, pp 1–22
- [33] Rocchetto M, Viganò L, Volpe M, Dalle Vedove G (2013) Using interpolation for the verification of security protocols. In: STM, Springer, pp 99–114
- [34] Rusinowitch M, Turuani M (2003) Protocol insecurity with a finite number of sessions and composed keys is NP-complete. Theor Comput Sci 299(1-3):451–475
- [35] Schmidt B, Sasse R, Cremers C, Basin DA (2014) Automated verification of group key agreement protocols. In: IEEE Symposium on Security and Privacy, SP, pp 179–194, doi: 10.1109/SP.2014.19

- [36] Turuani M (2006) The CL-Atse Protocol Analyser. In: Term Rewriting and Applications, LNCS 4098, Springer, pp 277–286

A ASLan++ specification of NSL.

```

1 specification NSL
2 channel_model CCM
3
4 entity Environment {
5   symbols
6     a,b:agent;
7
8   entity Session (A, B: agent) {
9
10    entity Alice (Actor, B: agent) {
11
12      symbols
13        Na, Nb: text;
14
15      body {
16        Na := fresh();
17        Actor -> B: {Na.Actor}_pk(B);
18        B -> Actor: {Na.?Nb.B}_pk(Actor);
19        Actor -> B: {auth:(Nb) }_pk(B));
20      }
21    }
22
23    entity Bob (A, Actor: agent) {
24
25      symbols
26        Na, Nb: text;
27
28      body {
29        ? -> Actor: {?Na.?A}_pk(Actor);
30        Nb := fresh();
31        Actor -> A: {Na.Nb.Actor}_pk(A);
32        A -> Actor: {auth:(Nb) }_pk(Actor);
33      }
34    }
35
36    body { % of Session
37      new Alice(A,B);
38      new Bob(A,B);
39    }
40
41    goals
42      auth(_) A *-> B;
43  }
44
45  body { % of Environment
46    any Session(a,i);
47    any Session(a,b);
48  }
49 }

```

B NSL Program Graph

Fig. 11 shows the SiL program graph for the protocol NSL with respect to the instantiation of Example 1.

C Proof of Lemma 1

Proof. We show two representative cases; the other ones can be treated similarly. (i)
Let the statement I considered be the receipt of a message having the form:

```

1 entity Environment {
2 ...
3 entity Session (A, B: agent) {
4 ...
5 entity Alice(Actor, B: agent) {
6 ...
7 body {
8 ...
9   B -> Actor: M(?A_1, ..., ?A_n)
10 ...
11 }

```

The corresponding ASLan rule r has the form:

```

1 step ... (...) :=
2   PF'.
3   iknows(M'(N_1, ..., N_n)).
4   state_Alice(B_1, ..., B_m)
5   =>
6   R'.
7   state_Alice(B'_1, ..., B'_m)

```

where M' is the ASLan translation of M , $n \leq m$ and $\forall j. 1 \leq j \leq m$ if $j = f(\text{Alice}, A_i)$ for some $1 \leq i \leq n$, then $B'_j = N_i$, otherwise $B'_j = B_j$.

For simplicity, we ignore in the variable names the prefixes referring to the session instance. w has the form:

```

1 if (IK | - M''(Y_1, ..., Y_n))
2 then
3   Alice.A_1 = Y_1;
4   ...
5   Alice.A_N = Y_N;
6 else
7   end

```

where M'' is the SiL translation of M where we have replaced $?A_1, \dots, ?A_n$ with $Y_1, \dots, Y_n$.

($\Rightarrow$) Let S' be such that $S \xrightarrow{r} S'$. By the semantics of ASLan, there must exist a substitution σ such that:

$$\text{iknows}(M'(N_1, \dots, N_n)).\text{state_Alice}(B_1, \dots, B_m)\sigma \subseteq \lceil S \rceil^H$$

Furthermore, there exists a substitution σ'' such that:

$$\text{state_Alice}(B'_1, \dots, B'_m)\sigma\sigma'' \subseteq \lceil S' \rceil^H$$

Then we can build an extension $\bar{\zeta}$ of ζ such that:

- $\bar{\zeta}(Y_i) = \sigma(N_i)$ for $1 \leq i \leq n$;
- $\bar{\zeta}(A) = \sigma(A)$ for any other variable A .

It follows that $M'(N_1, \dots, N_n)\sigma \sim M''(Y_1, \dots, Y_n)\bar{\zeta}$ and since $\text{iknows}(M'(N_1, \dots, N_n))\sigma \subseteq \lceil S \rceil^H$ then, by hypothesis, $M''(Y_1, \dots, Y_n) \in DY(\bar{\zeta}(IK))$ which implies $\langle IK \vdash M''(Y_1, \dots, Y_n), \bar{\zeta} \rangle \Downarrow$

true. By using this fact in the following derivation:

$$\begin{array}{c}
\frac{\frac{\frac{\langle Y_1, \bar{\zeta} \rangle \Downarrow \bar{\zeta}(Y)}{\langle \Phi_1, \bar{\zeta} \rangle \Downarrow \zeta_1} \quad \vdots \quad \frac{\langle Alice.A_2 := Y_2, \zeta_1 \rangle \Downarrow \zeta_2}{\langle \Phi_2, \bar{\zeta} \rangle \Downarrow \zeta_2} \quad \dots}{\langle \Phi_{n-1}, \bar{\zeta} \rangle \Downarrow \Sigma_{n-1}} \quad \frac{\langle Y_n, \zeta_{n-1} \rangle \Downarrow \zeta(Y_n)}{\langle \Psi_n, \zeta_{n-1} \rangle \Downarrow \zeta'_n} \\
\frac{\langle IK \vdash M''(Y_1, \dots, Y_n), \bar{\zeta} \rangle \Downarrow true}{\langle if (IK \vdash M''(Y_1, \dots, Y_n)) then \Phi_n else end, \bar{\zeta} \rangle \Downarrow \zeta'}
\end{array}$$

we get that $\langle w, \bar{\zeta} \rangle \Downarrow \zeta' \equiv \zeta_n$, where we have used the abbreviations:

$$\begin{aligned}
\Phi_i &\equiv Alice.A_1 := Y_1; \dots; Alice.A_i := Y_i; \\
\Psi_i &\equiv Alice.A_1 := Y_i; \dots; Alice.A_i := Y_n; \\
\zeta_i &\equiv \bar{\zeta}[Alice.A_1 \leftarrow \bar{\zeta}(Y_1), \dots, Alice.A_i \leftarrow \bar{\zeta}(Y_i)].
\end{aligned}$$

We have that $S'(Alice, f(Alice, A_i)) = \sigma(B'_f(Alice, A_i)) = \sigma(N_i) = \zeta'(Alice.A_i)$, for $1 \leq i \leq n$. Since S' and ζ' coincide with S and ζ , respectively, for what concerns the other variables, we can conclude $S' \sim \zeta'$.

($\Leftarrow$) Assume there exists an extension $\bar{\zeta}$ of ζ such that $\langle w, \bar{\zeta} \rangle \Downarrow \zeta'$. The case when $\langle IK \vdash M''(Y_1, \dots, Y_n), \bar{\zeta} \rangle \Downarrow false$ is trivial, since $\zeta' \equiv \bar{\zeta}$ and we can easily take $S' \equiv S$. Let $\langle IK \vdash M''(Y_1, \dots, Y_n), \bar{\zeta} \rangle \Downarrow true$. Then $M''(Y_1, \dots, Y_n)\bar{\zeta} \in DY(\zeta(IK))$. It follows that we can choose a substitution σ such that $\sigma(N_i) = \bar{\zeta}(Y_i)$, for $1 \leq i \leq n$, and thus $iknows(M'(N_1, \dots, N_n))\sigma \subseteq \lceil S \rceil^H$. By applying the rules of SiL semantics as above and the rule r , we get an S' such that $S \xrightarrow{r} S'$ and $S' \sim \zeta'$.

(ii) Let us assume that Alice wants to authenticate Bob and consider, without loss of generality, a program instance pi where $pi(Alice) = a$ and $pi(Bob) = i$, since if Bob is played by an honest agent, then the authentication property is trivially satisfied. I has the form:

```

1 entity Environment {
2   ...
3   entity Session (A, B: agent) {
4     ...
5     entity Alice(Actor, B: agent) {
6       ...
7       body {
8         ...
9         B -> Actor: auth:(M);
10        ...
11      }
12    }
13  }
14  ...
15  goals
16    auth:(_) B *-> A;
17    ....
18 }

```

and is a particular case of a receipt. As such, it is translated as a common receipt, treated in case (i), plus special constructs/rules aimed at handling the goal conditions, which will be treated here. The corresponding ASLan attack state is described by:

```

1 attack_state auth(M', Actor, B, ...) :=
2   not(dishonest(B)).
3   not(witness(B, Actor, M', auth)).
4   request(Actor, B, M', auth, ...)

```

where M' is the ASLan translation of M (for simplicity, we assume here that the payload on which authentication is based is the whole message). We also add a corresponding ASLan rule r of the form:

```

1 AS => AS.attack

```

which simply adds the 0-ary predicate *attack* to an attack state AS containing the predicates described above.

The corresponding SiL statement w has the form:

```

1 if(not(Alice.B = i) and not(witness(Alice.B, Alice.Actor, M')))
2 then
3   attack := true;
4 else
5   skip;

```

where M'' is the SiL translation of M . First, we notice that while the rule r can be applied at any step in an ASLan run, the corresponding SiL statement w is placed, by the translation procedure, immediately after the receipt instruction. For simplicity, we will restrict to consider those ASLan runs where attack rules concerning authentication goals, like r above, are only applied immediately after the receipt of the corresponding message. This can be done without loss of generality (and is also the reason why we do not need a *request* predicate in SiL).

($\Rightarrow$) In order to apply the rule r , by the semantics of ASLan, there must exist a substitution σ such that:

$$request(Actor, B, M', auth, \dots).state_Alice(\dots, B, \dots)\sigma \subseteq [S]^H$$

where, in particular, $\sigma(B) = S(Alice, f(Alice, B))$. At the same time, we have: $dishonest(B)\sigma \not\subseteq [S]^H$ and $witness(B, Actor, M', auth)\sigma \not\subseteq [S]^H$.

Since, as for every ASLan state, $dishonest(i) \subseteq [S]^H$, we get $\sigma(B) \neq i$. Let $\bar{\zeta}$ be an extension of ζ . By hypothesis, $S \sim \zeta$, from which we infer $\sigma(B) = S(Alice, f(Alice, B)) = \zeta(Alice.B) = \bar{\zeta}(Alice.B) \neq i$. With analogous arguments, we infer $(\bar{\zeta}(Alice.B), \bar{\zeta}(Alice.Actor), \bar{\zeta}(M'')) \notin \bar{\zeta}(witness)$. By using these facts, we obtain the derivation in Figure 12.

We have that S' and ζ' differ from S and $\bar{\zeta}$, respectively, only for the value of the predicate *attack*. By observing that $attack \subseteq [S']^H$ and $\zeta'(attack) = true$, we conclude $S' \sim \zeta'$.

($\Leftarrow$) Let $\bar{\zeta}$ be an extension of ζ such that $\langle w, \zeta \rangle \Downarrow \zeta'$. The case when $\langle \Psi, \bar{\zeta} \rangle \Downarrow false$, where Ψ is defined as in ($\Rightarrow$) above, is trivial. Let us consider $\langle \Psi, \bar{\zeta} \rangle \Downarrow true$. By hypothesis, $S \sim \zeta$ and thus the preconditions of r concerning the predicates *dishonest* and *witness* are enabled in S . As for the condition on the *request*, it is enabled by the fact that the corresponding receipt has just been encountered, by construction of a SiL graph. It follows that r can be applied and we get an ASLan state S' , which differs from S only in the fact that $attack \subseteq [S']^H$. Moreover, by applying the same derivation as in case ($\Rightarrow$) above, we have $\zeta'(attack) = true$, from which we conclude $S' \sim \zeta'$. $\square$

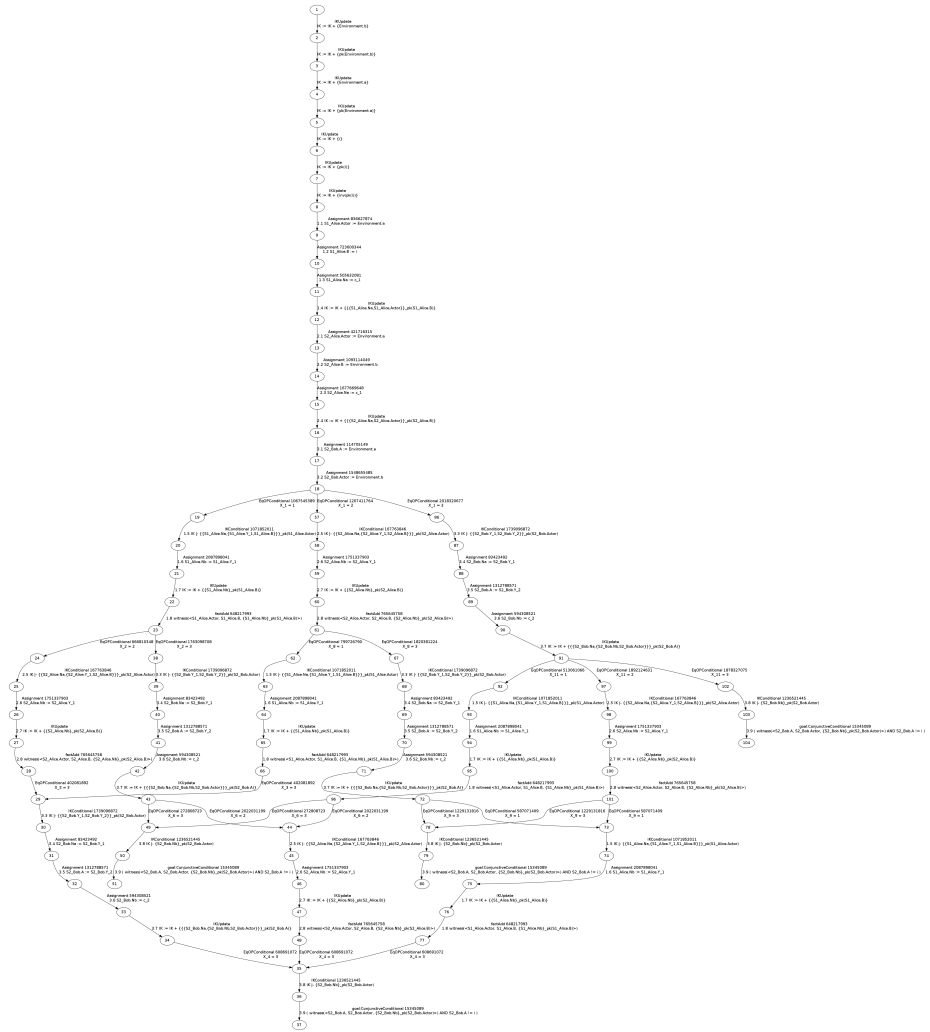


Figure 11: A SiL program graph for NSL.

$$\begin{array}{c}
\frac{\frac{\frac{\langle Alice.B, \xi \rangle \Downarrow \xi(Alice.B) \quad \langle i, \bar{\xi} \rangle \Downarrow i}{\langle Alice.B = i \rangle \Downarrow \bar{\xi} \Downarrow false} \quad \frac{\langle Alice.B, \bar{\xi} \rangle \Downarrow \xi(Alice.B) \quad \langle Alice.Actor, \bar{\xi} \rangle \Downarrow \xi(Alice.Actor) \quad \langle M', \bar{\xi} \rangle \Downarrow \xi(M'')}{\langle \Phi, \bar{\xi} \rangle \Downarrow false}}{\langle Alice.B = i \rangle \Downarrow \bar{\xi} \Downarrow true} \quad \frac{\langle \Psi, \bar{\xi} \rangle \Downarrow true \quad \frac{\langle not(Alice.B = i), \bar{\xi} \rangle \Downarrow true \quad \frac{\langle not(\Phi), \bar{\xi} \rangle \Downarrow true}{\langle if \Psi \text{ then } attack := true \text{ else } skip, \bar{\xi} \rangle \Downarrow \bar{\xi}[true/attack] \equiv \zeta'}}{\langle attack := true, \bar{\xi} \rangle \Downarrow \bar{\xi}[true/attack]} \\
\frac{\langle \Psi, \bar{\xi} \rangle \Downarrow true \quad \frac{\langle \Phi, \bar{\xi} \rangle \Downarrow false \quad \langle not(\Phi), \bar{\xi} \rangle \Downarrow true}{\langle if \Psi \text{ then } attack := true \text{ else } skip, \bar{\xi} \rangle \Downarrow \bar{\xi}[true/attack] \equiv \zeta'}}{\langle attack := true, \bar{\xi} \rangle \Downarrow \bar{\xi}[true/attack]}
\end{array}$$

In the derivation, we used $\Phi \equiv witness(Alice.B, Alice.Actor, M'')$ and $\Psi \equiv not(Alice.B = i)$ and $(not(\Phi))$ as abbreviations.

Figure 12: A derivation for an authentication goal checking by the operational semantics of SiL.