

PANDA: An Attack Synthesis Tool for Distributed Protocols

Abstract—Distributed protocols underpin the modern internet, making their correctness and security critical. Formal methods provide rigorous tools for analyzing protocol correctness and cryptographic security, yet existing tools fall short for denial of service (DoS) analysis. We introduce PANDA, a tool that synthesizes attacks on distributed protocols by targeting communication channels to violate linear temporal logic (LTL) specifications. PANDA provides sound, complete analysis, synthesizing attacks or proving their absence through exhaustive state-space search. With support for pre-defined and custom attacker models, PANDA enables targeted DoS analysis and broader LTL-based verification, demonstrated through various case studies.

Index Terms—Protocols, Attack Synthesis, Denial of Service, Model Checking

I. INTRODUCTION

Distributed protocols are the foundation for the modern internet, and therefore ensuring their correctness and security is paramount. To this end, formal methods, the use of mathematically rigorous techniques for reasoning about software, has been increasingly employed to analyze and study distributed protocols. Historically, formal methods has been employed for reasoning about concurrency and distributed algorithms [1]–[3], and in recent years formal methods have been employed at scale to reason about the security of cryptographic protocols and primitives [?], [4]–[7]. This myriad of formal methods tooling applicable to secure protocols has enabled reasoning about security-relevant properties involving secrecy, authentication, indistinguishability in addition to concurrency, safety, and liveness. However, no previous formal methods tooling offered an effective solution for rigorously studying an attacker that controls communication channels. That is, how do you reason about an attacker that can arbitrarily drop, reorder, replay, or insert messages onto a communication channel?

To fill this gap, we introduce PANDA¹, a tool for synthesizing attacks on distributed protocols that implements and extends the theoretical framework proposed in [8]. In particular, PANDA targets the communication channels between the protocol endpoints, and synthesizes attacks to violate arbitrary linear temporal logic (LTL) specifications. PANDA either synthesizes attack, or proves the absence of such via an exhaustive state-space search. PANDA is sound and complete, meaning if there exists an attack PANDA will find it, and PANDA will never have false positives. PANDA supports pre-defined attacker models, including attackers that can replay, reorder, or drop messages on channels, as well as custom user-defined attacker models. Although PANDA best lends itself

for reasoning about denial of service attacks, it can target any specification expressible in LTL.

In this work we take an approach rooted in *formal methods* and *automated reasoning* to construct PANDA. In particular, we employ *model checking*, a sub-discipline of formal methods, to decidably and automatically find attacks in protocols or prove the absence of such.

We summarize our contributions:

- We present PANDA, a tool for synthesizing attacks against communication protocols. PANDA supports four general attacker model gadgets: an attacker that can drop, replay, reorder, or insert messages on a channel.
- We provide an overview of PANDA and show how it can be used through an example of the ABP protocol.
- We present two case studies for two well-known protocols TCP and Raft, illustrating the usefulness of PANDA.

We release our code and our models as open source at <https://anonymous.4open.science/r/attacksynth-artifact-1B5D>.

II. PANDA ARCHITECTURE

In this section we discuss the details behind the design, formal guarantees, implementation, and usage of PANDA.

A. Mathematical Preliminaries

Linear Temporal Logic (LTL) is a model logic for reasoning about program executions. In LTL, we say a program P *models* a property ϕ (notationally, $P \models \phi$). That is, ϕ holds over every execution of P . If ϕ does not hold over every execution of P , we say $P \not\models \phi$. The LTL language is given by predicates over a first-order logic with additional temporal operators: *next*, *always*, *eventually*, and *until*.

An LTL model is a tool that, given P and ϕ , can automatically check whether or not $P \models \phi$; in general, LTL is a *decidable* logic, and LTL model checkers will always be able to decide whether $P \models \phi$ given enough time and resources.

We use $\parallel$ to denote rendezvous composition. That is, if $S = P \parallel Q$, processes P and Q are composed together into a singular state machine by matching their equivalent transitions.

LTL program synthesis is the problem of, given an LTL specification ϕ , automatically deriving a program P that satisfies ϕ (that is, $P \models \phi$). *LTL attack synthesis* is logically dual to LTL program synthesis. In attack synthesis, the problem is flipped: given a program P and a property ϕ such that $P \models \phi$, we ask whether there exists some “attack” A such that $(P \parallel A) \not\models \phi$. Fundamentally, PANDA is a synthesizer for such an A .

¹PANDA is a fictitious name for our system, for double-blind submission.

B. High-level design

As aforementioned, PANDA is based on *LTL attack synthesis*; in particular, PANDA synthesizes attacks with respect to *imperfect* channels. That is, PANDA is designed to synthesize attacks that involve replaying, dropping, reordering, or inserting messages on a communication channel.

PANDA is designed to attack user-specified communication channels in state machine-based formal models of distributed protocols. To use PANDA, the user inputs a formal model of a distributed protocol in the PROMELA language, the communication channel(s) in the formal model they wish to attack, the desired attacker model, and a formalized correctness property for the formal model. The formal model should satisfy the correctness property in absence of PANDA.

Once PANDA is invoked, it will modify the user-inputted PROMELA model such that it integrates the desired attacker model. Then, PANDA passes the updated PROMELA model to the model checker, which performs the exhaustive search or provides an explicit counterexample.

A high-level visual overview of the PANDA pipeline is given in the Figure 1.

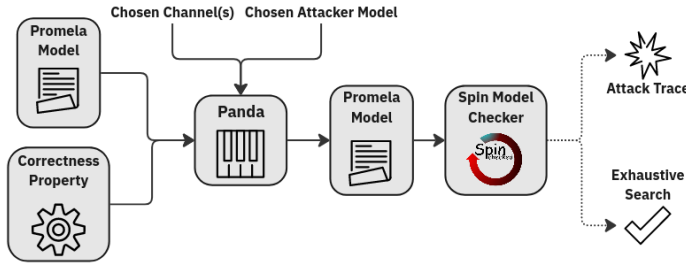


Fig. 1: A high-level overview of the PANDA workflow

C. Supported Attacker Models

PANDA supports four general attacker models: an attacker that can drop, replay, reorder, or insert messages on a channel. In this section we discuss the various details that went into the implementation of the gadgets that encapsulate the behavior of the respective attacker models.

Drop Attacker Model Gadget The most simple attacker model PANDA supports is an attacker that can *drop* messages from a channel. The user specifies a "drop limit" value that limits the number of packets the attacker can drop from the channel. Note, a higher drop limit will increase the search space of possible attacks, thereby increasing execution time. The dropper attacker model gadget PANDA synthesizes works as follows. The gadget will nondeterministically choose to observe a message on a channel. Then, if the drop limit variable is not zero, it will consume the message. An example is shown in Figure 1.

Replay Attacker Model Gadget The next attacker model PANDA supports is an attacker that can observe and *replay* messages back onto a channel. Similarly to the drop limit for the dropping attacker model, the user can specify a "replay limit" that caps the number of observed messages the attacker

```
chan cn = [8] of { int, int, int };

active proctype attacker_drop() {
  int b_0, b_1, b_2;
  byte lim = 3; // drop limit
MAIN:
  do
    :: cn ? [b_0, b_1, b_2] -> atomic {
      if
        :: lim == 0 -> goto BREAK;
        :: else ->
          cn ? b_0, b_1, b_2; // consume message
          on the channel
          lim = lim - 1;
          goto MAIN;
      fi
    }
  od
BREAK:
}
```

Listing 1: Example dropping attacker model gadget with drop limit of 3, targetting channel "cn"

can replay back onto the specified channel. The replay attacker model gadget PANDA employs works as follows. The gadget has two states, CONSUME and REPLAY. The gadget starts in the CONSUME state and nondeterministically reads (but not consumes) messages on the target channel, sending them into a local storage buffer. Once the gadget read the number of messages on the channel equivalent to the defined replay limit, its state changes to REPLAY. In the REPLAY state, the gadget nondeterministically selects messages from its storage buffer to replay onto the channel until out of messages. An example is shown in Figure 2.

Reorder Attacker Model Gadget PANDA supports synthesizing attackers that can *reorder* messages on a channel. Like the drop and replay attacker model gadgets, the user can specify a "reordering limit" that caps the number of messages that can be reordered by the attacker on the specified channel. The reordering attacker model gadget PANDA synthesizes works as follows. The gadget has three states, INIT, CONSUME, and REPLAY. The gadget begins in the INIT state, where it arbitrarily chooses a message to start consuming by transitioning to the CONSUME state. When in the CONSUME state, the gadget consumes all messages that appear on the channel, filling up a local buffer, until hitting the defined reordering limit. Once this limit is hit, the gadget transitions into the REPLAY state. In the REPLAY state, the gadget nondeterministically selects messages from its storage buffer to replay onto the channel until out of messages. An example is shown in Figure 3.

Insert Attacker Models PANDA supports the synthesis of attackers that can simply insert messages onto a channel. While the drop, replay, and reordering attacker model gadgets as previously described have complex gadgets that PANDA synthesizes with respect to a user-specified channel, the insert

```

chan cn = [8] of { int, int, int };

// local memory for the gadget
chan gadget_mem = [3] of { int, int, int };

active proctype attacker_replay() {
  int b_0, b_1, b_2; int i = 3;
  CONSUME:
  do
    // read messages until the limit is passed
    :: cn ? [b_0, b_1, b_2] -> atomic {
      cn ? <b_0, b_1, b_2> -> gadget_mem ! b_0,
        b_1, b_2;
      i--;
      if
        :: i == 0 -> goto REPLAY;
        :: i != 0 -> goto CONSUME;
      fi }
  od
  REPLAY:
  do
    :: atomic {
      // nondeterministically select a random
      value from the storage buffer
      int am;
      select(am : 0 .. len(gadget_mem)-1);
      do
        :: am != 0 ->
          am = am-1;
          gadget_mem ? b_0, b_1, b_2 ->
            gadget_mem ! b_0, b_1, b_2;
        :: am == 0 ->
          gadget_mem ? b_0, b_1, b_2 -> cn ! b_0,
            b_1, b_2;
          break;
        od }
    // doesn't need to use all messages on the
    channel
    :: atomic { gadget_mem ? b_0, b_1, b_2; }
    // once mem has no more messages, we're done
    :: empty(gadget_mem) -> goto BREAK;
  od
  BREAK:
}

```

Listing 2: Example replay attacker model gadget with the selected replay limit as 3, targeting channel "cn"

attacker model gadget is synthesized with respect to a user-defined *IO-file*. This file denotes the specific outputs and channels the attacker is capable of sending, and PANDA generates a gadget capable of synthesizing attacks using the given inputs. An example I/O file is given in Figure 4, and the generated gadget is given in Figure 5.

These attacker models can be mixed and matched as desired by the PANDA user. For example, a user can specify a drop attacker and replay attacker to target channel 1, a reordering attacker to target channel 2, and an insert attacker to target channel 3. If multiple attacker models are declared, PANDA will synthesize attacks where the attackers on different channel *coordinate* to construct a unifying attack.

```

chan cn = [8] of { int, int, int };

chan gadget_mem = [3] of { int, int, int };
active proctype attacker_reordering()
  priority 255 {
  byte b_0, b_1, b_2, blocker; int i = 3;
  INIT:
  do
    :: { // arbitrarily choose a message to
      start consuming on
      blocker = len(cn);
      do :: b != len(c) -> goto INIT; od
    }
    :: goto CONSUME;
  od
  CONSUME:
  do
    // consume messages with high priority
    :: c ? [b_0] -> atomic {
      c ? b_0 -> gadget_mem ! b_0; i--;
      if
        :: i == 0 -> goto REPLAY;
        :: i != 0 -> goto CONSUME;
      fi }
  od
  REPLAY:
  do
    // replay messages back onto the channel,
    also with priority
    :: atomic {
      int am;
      select(am : 0 .. len(gadget_mem)-1);
      do
        :: am != 0 ->
          am = am-1;
          gadget_mem ? b_0 -> attacker_mem_0 !
            b_0;
        :: am == 0 ->
          gadget_mem ? b_0 -> c ! b_0;
          break;
        od }
    :: atomic { empty(gadget_mem) -> goto
      BREAK; }
  od
  BREAK:
}

```

Listing 3: Example reordering attacker model gadget with the selected replay limit as 3, targeting channel "cn"

D. PANDA Implementation

We implemented PANDA on top of the SPIN, a popular and robust model checker for reasoning about distributed and concurrent systems. SPIN has existed for over 40 years, and has been applied to dozens of real systems including the Mars Rover [9], Path-Star Access server [10], and an avionics operating system [11]. Additionally, SPIN has spawned a dedicated formal methods symposium, currently in its 32nd year², and earned the 2002 ACM Software System award.

²<https://spin-web.github.io/SPIN2025/>

```

cn:
  I:
    0:1-1-1, 1-2-3, 3-4-5

```

Listing 4: Example I/O file targetting channel "cn"

```

chan cn = [8] of { int, int, int };

active proctype daisy() {
INIT:
  do
    :: cn ! 1,1,1;
    :: cn ! 1,2,3;
    :: cn ! 3,4,5;
    :: goto RECOVERY;
  od
RECOVERY:
}

```

Listing 5: Example gadget synthesized from an I/O file targetting the channel "cn"

Intuitively, models written in PROMELA, the modeling language of SPIN, are communicating state machines whose messages are passed over defined *channels*. Channels in PROMELA can either be unbuffered *synchronous* channels, or buffered *asynchronous* channels. PANDA generates attacks *with respect* to these defined channels.

```

// channel of buffer size 0
chan msg_channel = [0] of { int }

active proctype Peer1() {
  msg_channel ! 1;
}

active proctype Peer2() {
  int received_msg;
  msg_channel ? received_msg;
}

```

Listing 6: Example PROMELA model of peers communicating over a channel. ! indicates sending a message onto a channel, ? indicates receiving a message from a channel.

PANDA is designed to parse user-chosen channels and generate gadgets for sending, receiving, and manipulating messages on them. PANDA has built-in gadgets that are designed to emulate various real-world attacker models. Once one or multiple gadgets are generated, PANDA invokes SPIN to check if a given property of interest remains satisfied in the presence of the attacker gadgets.

Preventing PANDA Livelocks In general, there are two types of LTL properties: safety, and liveness. Informally, safety properties state "a bad thing never happens," and liveness properties state "a good thing always happens." Therefore, safety properties can be violated by finite traces, while liveness properties require infinite traces to be violated. When evalu-

ating a PANDA attacker model gadget against a PROMELA model and a liveness property, it is crucial to ensure the gadget has no cyclic behavior. If a PANDA gadget has cyclic behavior in any way, it will trivially violate the liveness property and produce a garbage attack trace. To prevent this, we make the following considerations.

First, we design our PANDA gadgets such that they never arbitrarily send and consume messages to a single channel. Second, we allow PANDA gadgets, which are always processing messages on channels, to arbitrarily "skip" messages on a channel if need be. To demonstrate the latter, consider the extension of the drop attacker model gadget in Figure 7. We implement message skipping by arbitrarily stopping and waiting after observing a message on a channel; once the channel is observed changing lengths, the message is considered skipped and future messages can be consumed.

```

chan cn = [8] of { int, int, int };

active proctype attacker_drop() {
  int b_0, b_1, b_2, blocker;
  byte lim = 3; // drop limit
MAIN:
  do
    :: cn ? [b_0, b_1, b_2] -> atomic {
      if
        :: lim == 0 -> goto BREAK;
        :: else ->
          cn ? b_0, b_1, b_2; // consume message
                           on the channel
          lim = lim - 1;
          goto MAIN;
      fi
    }
    // pass over a message on a channel as
    // needed
    :: cn ? [b_0, b_1, b_2] -> atomic {
      // wait for the channel to change lengths
      // then, once it does, go to MAIN
      blocker = len(cn);
      do
        :: blocker != len(cn) -> goto MAIN;
      od
    }
    :: goto BREAK;
  od
BREAK:
}

```

Listing 7: Example dropping attacker model gadget with message skipping

E. Usage

To demonstrate the usage of PANDA, we provide a step-by-step example of proving the alternate bit protocol (ABP) is secure with respect to attackers that can replay messages. ABP is a simple communication protocol that provides reliable communication between two peers over an unreliable communication by continually agreeing on a bit value.

To use PANDA, the user first authors a PROMELA model and a correctness property in LTL. For example, take the PROMELA model as shown in Listing 8. The sender repeatedly sends its stored bit, `A_curr`, to the receiver. The receiver changes its internal bit, `B_curr`, and sends an acknowledgment to the sender. When the sender receives the acknowledgment, it will bitflip `A_curr` and repeatedly send the updated bit. A natural specification for this protocol, formalized into the LTL property `eventually_agrees`, states that if the sender and receiver do not currently agree on a bit, they eventually will be able to reach an agreement.

```
chan StoR = [2] of { bit };
chan RtoS = [2] of { bit };

bit A_curr = 0, B_curr = 1, rcv_a, rcv_b;

active proctype Sender() {
  do
    :: StoR ! A_curr;
    :: RtoS ? rcv_a ->
      if :: rcv_a == A_curr ->
        A_curr = (A_curr + 1) % 2;
      fi
    od
}

active proctype Receiver() {
  do
    :: RtoS ! B_curr;
    :: StoR ? rcv_b ->
      :: rcv_b != B_curr ->
        B_curr = rcv_b;
      fi
    od
}

ltl eventually_agrees {
  (A_curr != B_curr) implies eventually
  (A_curr == B_curr)
}
```

Listing 8: Example (simplified) PROMELA model of the alternating bit protocol.

Next, the user selects a *channel* to generate an attacker on, and an attacker model of choice. For example, we select `StoR` and `RtoS` as our channels of choice, `replay` as our attacker model of choice, and assume the ABP model is in the file `abp.pml`. Then, we run PANDA via command line.

```
$ ./panda --model=abp.pml --attacker=replay
--channel=StoR,RtoS --eval
```

PANDA will then modify the `abp.pml` file to include the `replay` attacker gadgets attacking channels `StoR` and `RtoS`, and model-check it with SPIN. PANDA outputs the following text, cut down for readability, indicating an exhaustive search for attacks:

```
Full statespace search for:
  never claim + (eventually_agrees)
```

```
ltl eventually_agree ((A_curr!=B_curr))
  implies (eventually ((A_curr==B_curr)))

PANDA's exhaustive search is complete, no
attacks found!
```

If desired, `--output` can also be specified so the PANDA-modified `abp.pml` can be more closely examined and modified. A full shell-script replicating this example is available in the artifact.

III. CASE STUDIES

In this section we describe two case studies: the Transmission Control Protocol (TCP), a data transfer protocol, and Raft, a state machine replication protocol.

A. TCP

Transmission Control Protocol (TCP) is a transport-layer protocol designed to establish reliable, ordered communications between two peers. TCP is ubiquitous in today's internet, and therefore has seen ample formal verification efforts [12]–[14], including using PROMELA and SPIN [14]. We construct a TCP PROMELA model referencing the set of TCP RFCs. For our analysis, we borrow the four LTL properties used in [14], as detailed below:

- ϕ_1 = No half-open connections.
- ϕ_2 = Passive/active establishment eventually succeeds.
- ϕ_3 = Peers don't get stuck.
- ϕ_4 = SYN_RECEIVED is eventually followed by ESTABLISHED, FIN_WAIT_1, or CLOSED.

We evaluated the TCP PROMELA model against PANDA's drop, replay, and reordering attacker models on a single unidirectional communication channel. The resulting breakdown of attacks discovered is shown in Figure 2.

	Drop Attacker	Replay Attacker	Reorder Attacker
ϕ_1			
ϕ_2	x	x	
ϕ_3			
ϕ_4			

Fig. 2: Automatically discovered attacks against our TCP model for ϕ_1 through ϕ_4 . "x" indicates an attack was discovered, and no "x" indicates PANDA proved the absence of an attack via an exhaustive search. These experiments were ran on a laptop with an eighth generation i7 and 16gb of memory. Full attack traces are available in the artifact.

B. Raft

Raft is a consensus algorithm designed to replicate a state machine across distributed peers, and sees broad usage in distributed databases, key-value stores, distributed file systems, distributed load-balancers, and container orchestration. Historically, verification efforts of Raft using both constructive, mechanized proving techniques [15]–[17] and automated verification [17] have reasoned about the protocol under certain

assumptions about the stability of the communication channels. Previously, Raft has been proven to maintain properties of interest with respect volatile, attacker-controlled channels constructively using Rocq³ [16]. However, no previous approach to Raft verification has reasoned explicitly about a coordinated, arbitrary on-channel attacker *external* to the protocol itself. Uniquely, PANDA enables us to study Raft in this context.

Referencing the original Raft thesis [17] and other raft models [15], we constructed a PROMELA model of the Raft protocol. Additionally, we derived and formalized the following properties, which our PROMELA model satisfies:

- ϕ_1 = No two servers can be leaders in the same term.
- ϕ_2 = Entries committed to the log at the same index must be equivalent.
- ϕ_3 = Only leaders may append entries to the log.
- ϕ_4 = If a leader commits at an index, any server that becomes leader afterwards must follow that commit.
- ϕ_5 = If any two servers commit the same log entry, the log entry at the previous index must be equivalent

We construct our Raft model such that we can model-check an arbitrary number of peers. We also designed our model such that each peer maintains separate channels for receiving AppendEntry requests, AppendEntry responses, RequestVote requests, and RequestVote responses. This gives PANDA ample handle to reason about Raft. In particular, we study Raft in the presence of drop and replay attackers on all four aforementioned channel types, attacking both a minority and majority of peers.

To test PANDA, we altered our original Raft model to introduce a subtle bug in the raft consensus mechanism by not ensuring votes come from unique peers. We'll refer to our original, correct Raft model as `raft.pml`, and our buggy Raft model as `raft-bug.pml`. Both `raft.pml` and `raft-bug.pml` passed on ϕ_1 - ϕ_5 (that is, assuming the channels are perfect). We assess `raft-bug.pml` with PANDA, and a breakdown of our findings is shown in Figure 3.

Scenario	Attack found?
Dropping AppendEntries messages	no
Dropping RequestVote messages	no
Replaying RequestVote messages	yes (ϕ_1, ϕ_4 violated)
Replaying AppendEntry messages	no
Dropping RequestVoteResponse messages	no
Dropping AppendEntryResponse messages	no

Fig. 3: Breakdown of the attacker scenarios assessed with PANDA against our buggy Raft PROMELA model, `raft-bug.pml`. In all experiments, the Raft model was set to five peers and the drop/replay limits of the gadgets PANDA synthesized were set to two. We conducted our experiments on a research computing cluster, allocating 250GB of memory to each verification run. The full models and attacker traces are included in the artifact.

³Previously known as Coq

In our experiments, we found just one attack on our `raft-bug.pml` PROMELA model, violating election safety in particular. In this scenario, peer A and peer B are candidates for election. Peer A receives three votes, one from itself and two from other peers, and Peer B receives two votes, one from itself and one from another peer. The replay attacker simply replays the vote sent to peer B. Then, both Peer A and Peer B are convinced they won the election and change their state to leader. Following this, leader completeness is also naturally violated. In this scenario, PANDA demonstrates its ability to discover subtle bugs in protocol logic, exploiting the buggy Raft implementation.

IV. PROOFS OF SOUNDNESS AND COMPLETENESS

PANDA is an implementation of the theoretical attack synthesis framework proposed by [8]. This framework enjoys soundness and completeness guarantees for attacks discovered; that is, if there exists an attack, it is discovered, and if an attack is discovered, it is valid. However, the attack synthesis framework proposed by [8] reasons about an abstracted, theoretical process construct. Therefore, in order to correctly claim PANDA is also sound and complete, it is necessary to demonstrate discovering an attack within the theoretical framework reduces to the semantics of SPIN, the model checker PANDA is built on top of.

There exists a semantic gap between the theoretical attack synthesis framework proposed by [8], and the semantics of PANDA. Therefore, in order to correctly claim PANDA maintains the soundness and completeness of the theoretical framework it implements, it suffices to demonstrate finding an attack within the theoretical attack synthesis framework precisely reduces to the semantics of SPIN.

Definition 1 (Büchi Automata). A Büchi Automata is a tuple $B = (Q, \Sigma, \delta, Q_0, F)$ where:

- Q is a finite set of states,
- Σ is a finite alphabet,
- $\delta \subseteq Q \times \Sigma \times Q$ is a transition relation,
- $Q_0 \subseteq Q$ is a set of initial states,
- $F \subseteq Q$ is a set of accepting states.

A run of a Büchi Automata is an infinite sequence of states $q_0, q_1, q_2, \dots$ such that $q_0 \in Q_0$ and $(q_i, a, q_{i+1}) \in \delta$ for some $a \in \Sigma$ at each step i . The run is considered accepting if it visits states in F infinitely often.

Definition 2 (Process). A Process is a tuple $P = \langle AP, I, O, S, s_0, T, L \rangle$, where:

- AP is a finite set of atomic propositions,
- I is a set of inputs,
- O is a set of output, such that $I \cap O = \emptyset$,
- S is a finite set of states,
- $s_0 \in S$ is the initial state,
- $T \subseteq S \times (I \cup O) \times S$ is the transition relation,
- $L : S \rightarrow 2^{AP}$ is a labeling function mapping each state to a subset of atomic propositions.

A transition $(s, x, s') \in T$ is called an input transition if $x \in I$ and an output transition if $x \in O$.

Theorem 1. A process, as defined in [8], always directly corresponds to a Büchi Automata.

Proof. Given a Büchi Automata $B = (Q, \Sigma, \delta, Q_0, F)$, we construct a corresponding Process $P = \langle AP, I, O, S, s_0, T, L \rangle$ as follows:

- Atomic Propositions: $AP = \{\text{accept}\}$, a singleton set containing a special proposition indicating acceptance.
- Inputs and Outputs: $I = \Sigma$ and $O = \emptyset$.
- States: $S = Q$ and $s_0 \in Q_0$.
- Transition Relation: $T = \delta$.
- Labeling Function: $L : S \rightarrow 2^{AP}$ defined by

$$L(s) = \begin{cases} \{\text{accept}\} & \text{if } s \in F, \\ \emptyset & \text{otherwise.} \end{cases}$$

In this mapping, the states and transitions of the BA are preserved in the Process, and the accepting states F are identified via the labeling function L .

Conversely, given a Process $P = \langle AP, I, O, S, s_0, T, L \rangle$ with an acceptance condition defined by a distinguished proposition $p \in AP$, we define a Büchi Automata $B = (Q, \Sigma, \delta, Q_0, F)$ as follows:

- States: $Q = S$ and $Q_0 = \{s_0\}$.
- Alphabet: $\Sigma = I \cup O$.
- Transition Relation: $\delta = T$.
- Accepting States: $F = \{s \in S \mid p \in L(s)\}$.

Here, the accepting states in the BA correspond to those states in the Process that are labeled with the distinguished proposition p .

In both structures, a run is an infinite sequence of states connected by transitions:

- In the Büchi Automata: $q_0, q_1, q_2, \dots$ with $q_0 \in Q_0$ and $(q_i, a_i, q_{i+1}) \in \delta$ for some $a_i \in \Sigma$.
- In the Process: $s_0, s_1, s_2, \dots$ with $s_0 = s_0$ and $(s_i, x_i, s_{i+1}) \in T$ for some $x_i \in I \cup O$.

An accepting run in the Büchi Automata visits states in F infinitely often. Similarly, an accepting run in the Process visits states labeled with p infinitely often. Since $F = \{s \in S \mid p \in L(s)\}$, the acceptance conditions are preserved under the mappings. $\square$

Definition 3 (Threat Model). A threat model is a tuple $(P, (Q_i)_{i=0}^m, \phi)$ where:

- $P, Q_0, \dots, Q_m$ are processes.
- Each process Q_i has no atomic propositions (i.e., its set of atomic propositions is empty).
- ϕ is an LTL formula such that $P \parallel Q_0 \parallel \dots \parallel Q_m \models \phi$.
- The system $P \parallel Q_0 \parallel \dots \parallel Q_m$ satisfies the formula ϕ in a non-trivial manner, meaning that $P \parallel Q_0 \parallel \dots \parallel Q_m$ has at least one infinite run.

Theorem 2. Checking whether there exists an attacker under a given threat model, the $R\text{-}\exists\text{ASP}$ problem as proposed in [8], is equivalent to Büchi Automata language inclusion (which is in turn solved by the SPIN model checker).

Proof. For a given threat model $(P, (Q_i)_{i=0}^m, \phi)$, checking $\exists\text{ASP}$ is equivalent to checking

$$R = MC(P \parallel \text{Daisy}(Q_0) \parallel \dots \parallel \text{Daisy}(Q_m), \phi)$$

Where MC is a model checker, and $\text{Daisy}(Q_i)$ is for intents of this proof, equivalent to a process. Therefore, via the previous theorem we can construct Büchi Automata $BA_P, BA_{\text{Daisy}(Q_0)}, \dots, BA_{\text{Daisy}(Q_m)}$ from the processes $P, \text{Daisy}(Q_0), \dots, \text{Daisy}(Q_m)$. Then, we check

$$\text{SPIN}(BA_P \parallel BA_{\text{Daisy}(Q_0)} \parallel \dots \parallel BA_{\text{Daisy}(Q_m)}, \phi)$$

Or equivalently, translating ϕ to the equivalent Büchi Automata BA_ϕ via [2], we equivalently check

$$(BA_P \parallel BA_{\text{Daisy}(Q_0)} \parallel \dots \parallel BA_{\text{Daisy}(Q_m)}) \subseteq BA_\phi$$

Where rendezvous composition for I/O Büchi Automata is precise the same as for I/O Kripke Automata; that is, input and output transitions are matched. It's easy to see these composition operations are equivalent. $\square$

Theorem 3. Checking whether there exists an attacker for a given threat model, the $R\text{-}\exists\text{ASP}$ problem as proposed in [8], is in PSPACE.

Proof. By the previous argument the $\exists\text{ASP}$ problem corresponds to Büchi Automata language inclusion, which is well-known to be PSPACE-complete [18]. $\square$

V. RELATED WORK

Similar Tools. Several formal methods tools reason about attackers on secure protocols, primarily in the cryptographic context: ProVerif, VerifPal, Tamarin, and Scyther are *Symbolic* and abstract away cryptographic primitives as terms [5], [19]–[21], while CryptoVerif and EasyCrypt are *computational* and reason about game-based cryptographic security proofs [6], [22]. For a general overview, see [23], [24]. Before PANDA, model checker-based approaches for reasoning about secure protocols have typically employed SPIN or TLA+ and only reasoned about correctness [3], [25]–[28].

Reasoning About Channels. There is a long history of using formal methods tools ad-hoc to reason about on-channel attackers, particularly in the context of Byzantine protocols [16], [29]. Formal methods tools have also been applied to reason about message tampering [30], delays [31], and congestion control [32].

VI. CONCLUSION

In conclusion, PANDA addresses a critical gap in the formal verification of distributed protocols by enabling the synthesis of communication channel-based attacks against arbitrary linear temporal logic specifications. By leveraging SPIN, PANDA ensures soundness and completeness in attack synthesis. Its modular support for pre-defined attacker models enhances its versatility, enabling thorough protocol analysis across diverse and interesting scenarios. We demonstrate the effectiveness of PANDA by employing it to study TCP and Raft, marking it as an invaluable tool for ensuring the validity and security of distributed protocols.

REFERENCES

- [1] L. Lamport, “The temporal logic of actions,” *ACM Transactions on Programming Languages and Systems*, vol. 16, no. 3, p. 872–923, May 1994.
- [2] G. Holzmann, “The model checker spin,” *IEEE Transactions on Software Engineering*, vol. 23, no. 5, p. 279–295, May 1997.
- [3] E. M. Clarke and Q. Wang, “25 years of model checking.”
- [4] D. Basin, C. Cremers, J. Dreier, and R. Sasse, “Tamarin: Verification of large-scale, real-world, cryptographic protocols,” *IEEE Security & Privacy*, vol. 20, no. 3, p. 24–32, May 2022.
- [5] N. Kobeissi, G. Nicolas, and M. Tiwari, “Verifpal: Cryptographic protocol analysis for the real world.”
- [6] B. Blanchet and C. Jacomme, “Cryptoverif: a computationally-sound security protocol verifier.”
- [7] D. Basin, F. Linker, and R. Sasse, “A formal analysis of the imessage pq3 messaging protocol.”
- [8] Anonym, “Anonymized for blinded submission,” XXX.
- [9] G. J. Holzmann, “Mars code,” *Communications of the ACM*, vol. 57, no. 2, p. 64–73, Feb. 2014.
- [10] G. J. Holzmann and M. H. Smith, “Automating software feature verification,” *Bell Labs Technical Journal*, vol. 5, no. 2, p. 72–87, 2000.
- [11] W. Visser, K. Havelund, G. Brat, and S. Park, “Model checking programs,” in *Proceedings ASE 2000. Fifteenth IEEE International Conference on Automated Software Engineering*. Grenoble, France: IEEE, 2000, p. 3–11. [Online]. Available: <http://ieeexplore.ieee.org/document/873645/>
- [12] G. Cluzel, K. Georgiou, Y. Moy, and C. Zeller, “Layered formal verification of a tcp stack,” in *2021 IEEE Secure Development Conference (SecDev)*. Atlanta, GA, USA: IEEE, Oct. 2021, p. 86–93. [Online]. Available: <https://ieeexplore.ieee.org/document/9652642/>
- [13] M. A. S. Smith, “Formal verification of tcp and t/tcp,” Thesis, Massachusetts Institute of Technology, 1997, accepted: 2008-09-03T18:09:43Z. [Online]. Available: <https://dspace.mit.edu/handle/1721.1/42779>
- [14] M. L. Pacheco, M. V. Hippel, B. Weintraub, D. Goldwasser, and C. Nita-Rotaru, “Automated attack synthesis by extracting finite state machines from protocol specification documents,” in *2022 IEEE Symposium on Security and Privacy (SP)*. San Francisco, CA, USA: IEEE, May 2022, p. 51–68. [Online]. Available: <https://ieeexplore.ieee.org/document/9833673/>
- [15] D. Woos, J. R. Wilcox, S. Anton, Z. Tatlock, M. D. Ernst, and T. Anderson, “Planning for change in a formal verification of the raft consensus protocol,” in *Proceedings of the 5th ACM SIGPLAN Conference on Certified Programs and Proofs*. St. Petersburg FL USA: ACM, Jan. 2016, p. 154–165. [Online]. Available: <https://dl.acm.org/doi/10.1145/2854065.2854081>
- [16] J. R. Wilcox, D. Woos, P. Panchekha, Z. Tatlock, X. Wang, M. D. Ernst, and T. Anderson, “Verdi: A framework for implementing and formally verifying distributed systems.”
- [17] D. Ongaro, “Consensus: Bridging theory and practice.”
- [18] D. Kozen, “Lower bounds for natural proof systems,” in *18th Annual Symposium on Foundations of Computer Science (sfcs 1977)*. Providence, RI, USA: IEEE, Sep. 1977, p. 254–266. [Online]. Available: <http://ieeexplore.ieee.org/document/4567949/>
- [19] B. Blanchet, B. Smyth, V. Cheval, and M. Sylvestre, “Proverif 2.05: Automatic cryptographic protocol verifier, user manual and tutorial.”
- [20] D. Basin, C. Cremers, J. Dreier, and R. Sasse, “Tamarin: Verification of large-scale, real-world, cryptographic protocols,” *IEEE Security & Privacy*, vol. 20, no. 3, p. 24–32, May 2022.
- [21] C. J. F. Cremers, *The Scyther Tool: Verification, Falsification, and Analysis of Security Protocols*, ser. Lecture Notes in Computer Science. Berlin, Heidelberg: Springer Berlin Heidelberg, 2008, vol. 5123, p. 414–418. [Online]. Available: http://link.springer.com/10.1007/978-3-540-70545-1_38
- [22] V. Pereira, “Easycrypt - a (brief) tutorial.”
- [23] M. Barbosa, G. Barthe, K. Bhargavan, B. Blanchet, C. Cremers, K. Liao, and B. Parno, “Sok: Computer-aided cryptography,” in *2021 IEEE Symposium on Security and Privacy (SP)*, May 2021, p. 777–795. [Online]. Available: https://ieeexplore.ieee.org/document/9519449/?ar_number=9519449
- [24] D. Basin, C. Cremers, and C. Meadows, *Model Checking Security Protocols*. Cham: Springer International Publishing, 2018, p. 727–762. [Online]. Available: http://link.springer.com/10.1007/978-3-319-10575-8_22
- [25] A. S. Khan, M. Mukund, and S. P. Suresh, *Generic Verification of Security Protocols*, ser. Lecture Notes in Computer Science. Berlin, Heidelberg: Springer Berlin Heidelberg, 2005, vol. 3639, p. 221–235. [Online]. Available: http://link.springer.com/10.1007/11537328_18
- [26] H. Wayne, “Modeling adversaries with tla+,” <https://www.hillelwayne.com/post/adversaries/>, 2019, accessed: 2024-12-03. [Online]. Available: <https://www.hillelwayne.com/post/adversaries/>
- [27] P. Narayana, R. Chen, Y. Zhao, Y. Chen, Z. Fu, and H. Zhou, “Automatic vulnerability checking of ieee 802.16 wimax protocols through tla+,” in *2006 2nd IEEE Workshop on Secure Network Protocols*, Nov. 2006, p. 44–49. [Online]. Available: <https://ieeexplore.ieee.org/document/4110436/?arnumber=4110436>
- [28] G. Delzanno, M. Tatarek, and R. Traverso, “Model checking paxos in spin,” *Electronic Proceedings in Theoretical Computer Science*, vol. 161, p. 131–146, Aug. 2014.
- [29] M. Castro and B. Liskov, “Practical byzantine fault tolerance and proactive recovery,” *ACM Transactions on Computer Systems*, vol. 20, no. 4, p. 398–461, Nov. 2002.
- [30] N. Ben Henda, “Generic and efficient attacker models in spin,” in *Proceedings of the 2014 International SPIN Symposium on Model Checking of Software*. San Jose CA USA: ACM, Jul. 2014, p. 77–86. [Online]. Available: <https://dl.acm.org/doi/10.1145/2632362.2632378>
- [31] J. Ginesin, M. von Hippel, E. Defloor, C. Nita-Rotaru, and M. Tüxen, “A formal analysis of sctp: Attack synthesis and patch verification,” no. arXiv:2403.05663, Mar. 2024, arXiv:2403.05663 [cs]. [Online]. Available: <http://arxiv.org/abs/2403.05663>
- [32] S. Jero, E. Hoque, D. Choffnes, A. Mislove, and C. Nita-Rotaru, “Automated attack discovery in tcp congestion control using a model-guided approach,” in *Proceedings 2018 Network and Distributed System Security Symposium*. San Diego, CA: Internet Society, 2018. [Online]. Available: https://www.ndss-symposium.org/wp-content/uploads/2018/02/ndss2018_02A-1_Jero_paper.pdf