

LTL_f Synthesis on Probabilistic Systems

Andrew M. Wells

Rice University, Houston, Texas
andrew.wells@rice.edu

Morteza Lahijanian

University of Colorado, Boulder, Colorado
morteza.lahijanian@colorado.edu

Lydia E. Kavradi

Rice University, Houston, Texas
kavraki@rice.edu

Moshe Y. Vardi

Rice University, Houston, Texas
vardi@rice.edu

Many systems are naturally modeled as Markov Decision Processes (MDPs), combining probabilities and strategic actions. Given a model of a system as an MDP and some logical specification of system behavior, the goal of synthesis is to find a policy that maximizes the probability of achieving this behavior. A popular choice for defining behaviors is Linear Temporal Logic (LTL). Policy synthesis on MDPs for properties specified in LTL has been well studied. LTL, however, is defined over infinite traces, while many properties of interest are inherently finite. Linear Temporal Logic over finite traces (LTL_f) has been used to express such properties, but no tools exist to solve policy synthesis for MDP behaviors given finite-trace properties. We present two algorithms for solving this synthesis problem: the first via reduction of LTL_f to LTL and the second using native tools for LTL_f. We compare the scalability of these two approaches for synthesis and show that the native approach offers better scalability compared to existing automaton generation tools for LTL.

1 Introduction

Many real-world systems are stochastic in nature. They evolve in the world according to action (control) decisions and the uncertainty embedded in the execution of those actions, resulting in stochastic behavior. *Formal synthesis* studies how the system should choose actions so that it can increase the chances of achieving a desirable behavior. To allow such reasoning, *Markov Decision Processes* (MDPs) are typically used to model these systems since MDPs effectively capture sequential decision-making and probabilistic evolutions [5]. The desired behavior is expressed in a formal language, which yields expressive and unambiguous specifications. Temporal logics are a common choice for this language since they allow a combination of temporal and boolean reasoning over the behavior of the system. Most temporal logics are interpreted over behaviors with infinite time durations, but some behaviors are inherently finite and can be expressed more intuitively and more practically using a temporal logic with finite semantics [15, 14]. This work investigates policy synthesis on MDPs for specifications expressed in a formal language called *Linear Temporal Logic over finite traces* (LTL_f) [10], which reasons over system behaviors with finite horizons.

A popular specification language in formal verification is Linear Temporal Logic (LTL) [21]. LTL provides a natural description of temporal properties over an infinite trace. LTL can express properties such as order, e.g., “first *a* and next eventually *b*,” and lack of starvation, e.g., “globally eventually resources are granted.” While infinite-trace properties are essential for reasoning about many systems, other systems require finite-trace properties [15, 14], for example robot planning for finite behaviors. If a robot is to build an object, the interest is in the finite rather than the infinite traces that accomplish this task. In such cases LTL_f [10] is a suitable alternative to LTL, cf. [14].

A number of recent studies has focused on developing efficient frameworks for solving LTL_f-reasoning problems, e.g., [10, 11, 27, 15, 19, 14]. De Giacomo and Vardi [10] presented a translation of LTL_f to LTL, implying that existing tools for LTL satisfiability and synthesis can be used for LTL_f with suitable transformations. In [11], the problem of reactive synthesis from LTL_f specifications that contain system (controllable) and environment (uncontrollable) variables is considered. The work shows that this problem reduces to strategy synthesis in a two-player automaton game, which is a 2EXPTIME-complete problem. To enable practical solutions to this problem, [27] introduces a symbolic approach to the LTL_f synthesis problem. The applications of those results in the field of robot planning are studied in [15, 13, 14]. In [19], the fundamental problem of LTL_f satisfiability checking is studied through a SAT-based approach. The results of those studies show that more scalable tools can be built by treating LTL_f *natively*, rather than reducing to LTL, even though LTL reasoning is a more mature and extensively studied domain. The underlying assumption in all those works is that the system is either purely deterministic or purely nondeterministic. Hence, it is natural to ask whether similar results can be obtained for probabilistic systems. That is, do native LTL_f techniques outperform reasoning by reduction to LTL?

Formal synthesis on probabilistic systems has been extensively studied in the formal-verification literature. In those studies, MDPs are a de facto modeling tool for the probabilistic system, and LTL and PCTL (probabilistic computation temporal logic) [5] are the specification languages of choice. For LTL synthesis, the approach is based on first translating the specification to a *deterministic Rabin automaton* (DRA) and then solving a stochastic shortest path problem [24, 9] on the product of the DRA with the MDP. Tools such as PRISM [18] can solve this LTL synthesis problem efficiently with their symbolic engine. LTL_f synthesis for probabilistic systems, however, has not been studied, and it is not clear whether the same tools can be extended for LTL_f reasoning with a similar efficiency.

Assigning rewards based on temporal goals has also been studied in the context of planning. In [23], a logic similar to LTL_f, \$FLTL, is considered; however, this logic cannot express properties such as “Eventually.” In [3, 4], PastLTL is used to describe finite properties associated with rewards. PastLTL and LTL_f naturally express different properties (see [23]). In [7, 6], rewards can be attached to LTL_f or *linear dynamic logic on finite traces* (LDL_f) [10] formulas. In all of these MDP planning works, the goal is to maximize rewards rather than to study the behavior of the MDP. These approaches use approximations that give lower bounds and converge only in the limit. Thus, they cannot necessarily give a negative answer to a decision query about a synthesis problem (e.g., does a policy with at least 95% probability into success exist?).

In this work, we present the problem of LTL_f synthesis for MDPs. In approaching this problem, we specifically seek to answer the empirical question of which approach is more efficient; an approach based on the mature and well-studied LTL synthesis or an approach based on the unique properties of LTL_f itself. The answer to this question can have a broader impact on the employment of formal methods for probabilistic systems. Hence, we introduce two solutions to this problem. The first one is based on a reduction to an LTL synthesis problem, and the second one is a native approach. For the first approach, we show a translation of the LTL_f specification to LTL and a corresponding augmentation of the MDP that allows us to use standard tools for LTL synthesis. For the second approach, we use specialized tools to obtain an automaton that we input to standard tools in order to solve the problem.

Even though both approaches have the same theoretical complexity bound [5, 11], we demonstrate that the native approach scales better than the translation to LTL through a series of benchmarking case studies. Our code and examples are available on GitHub [25].

2 Problem Formulation

Our goal is to synthesize a policy for an MDP such that the probability of satisfying a given LTL_f property is maximized. First we introduce the formalisms needed to define this problem. (For a detailed treatment, see [5]). In Section 2.1, we introduce Markov processes and define the labeling of a path on these processes and the probability measure associated with a set of paths. Next, we introduce LTL and the finite-trace version LTL_f in Section 2.2. We then define when a path on an MDP satisfies an LTL or LTL_f formula as well as the probability of satisfaction in Section 2.3. Finally, we give our formal problem definition in Section 2.4.

2.1 Markov Processes

Markov processes are frequently used to model systems that exhibit stochastic behavior. While this paper deals with MDPs, it is also important to define *Discrete-Time Markov Chains* (DTMCs) as they are needed to define probability measures.

Definition 1 (DTMC). *A labeled Discrete-Time Markov Chain (DTMC) is a tuple $\mathcal{D} = (S, P, s_{init}, AP, L)$, where S is a countable set of states, $s_{init} \in S$ is the initial state, AP is a finite set of atomic propositions, $L : S \rightarrow 2^{AP}$ is a labeling function, and $P : S \times S \rightarrow [0, 1]$ is a transition probability function where $\sum_{s' \in S} P(s, s') = 1$ for all $s \in S$.*

An execution of a DTMC is given by a path as defined below.

Definition 2 (Path). *A path w through a DTMC is an infinite sequence of states:*

$$w = s_0 s_1 \dots s_n \dots, \quad \text{such that} \quad P(s_i, s_{i+1}) > 0 \quad \forall i \geq 0.$$

$Paths(s)$ is the set of paths starting in s . A finite path is a path with a last state s_n . $Paths_{fin}(s)$ is the set of all finite paths starting in s . For every path $w \in Paths(s)$, $pre(w)$ denotes the set of prefixes of w (similarly for $w \in Paths_{fin}(s)$).

A trace or labeling of a path is the sequence of labels of the states in the path.

Definition 3 (Labeling of a Path). *A labeling (also referred to as valuation or trace) of an infinite path $w = s_0 s_1 \dots s_n \dots$ is the sequence $L(s_0)L(s_1) \dots L(s_n) \dots$. The labeling of a finite path is defined analogously. We use $L(w)$ to denote the labeling of w .*

To reason over the paths of a DTMC probabilistically, we need to define a probability space with a probability measure over infinite paths. To this end, we first define cylinder sets that extend a finite path to a set of infinite paths.

Definition 4 (Cylinder Set for DTMC). *The cylinder set for some finite path $w \in Paths_{fin}(s_{init})$, denoted by $Cyl(w)$, is the set of all infinite paths that share w as a prefix:*

$$Cyl(w) = \{w' \in Paths(s_{init}) \mid w \in pre(w')\}. \quad (1)$$

Definition 5 (Probability Measure over Paths of a DTMC). *For the probability space $(\Omega, \mathcal{E}, Pr)$, where sample space $\Omega = Paths(s_{init})$, event space $\mathcal{E}$ is the smallest σ -algebra on Ω containing the cylinder sets of all finite paths (i.e., $Cyl(w) \in \mathcal{E}$ for all $w \in Paths_{fin}(s_{init})$), the probability measure Pr is defined as:*

$$Pr(Cyl(s_0 \dots s_n)) = \prod_{0 \leq i < n} P(s_i, s_{i+1}), \quad (2)$$

where $s_0 = s_{init}$. This probability measure is unique per Carathéodory's extension theorem. We define $Pr(Cyl(s_{init})) = 1$.

Some systems exhibit not only probabilistic behavior but also non-deterministic behavior. These systems are typically modeled as Markov Decision Processes (MDPs). MDPs extend the definition of DTMCs by allowing choices of actions at each state.

Definition 6 (MDP). A labeled Markov Decision Process (MDP) is a tuple: $\mathcal{M} = (S, A, P, s_{init}, AP, L)$, where s_{init} , AP and L are as in Def. 1 and:

- S is a finite set of states;
- A is a finite set of actions, and $A(s) \subseteq A$ denotes the set of actions enabled at state $s \in S$;
- $P : S \times A \times S \rightarrow [0, 1]$ is the transition probability function where $\sum_{s' \in S} P(s, a, s') = 1$ for all $s \in S$ and $a \in A(s)$.

Example 1. An example of an MDP is shown in Figure 1. Actions ($A = \{a_0, a_1\}$) and probabilities are shown as edge labels. State names are within each state and state labels are above each state.

The notion of path can be straightforwardly extended from DTMCs to MDPs.

Definition 7 (Path through MDP). A path w through an MDP is a state followed by a sequence of action-state pairs:

$$w = s_0 \langle a_0, s_1 \rangle \langle a_1, s_2 \rangle \dots \langle a_{n-1}, s_n \rangle \dots$$

such that $s_0 = s_{init}$, $a_i \in A(s_i)$, and $P(s_i, a_i, s_{i+1}) > 0$ for all $i \geq 0$. $Paths(s)$, $Paths_{fin}(s)$ and $pre(w)$ are defined analogously to the DTMC case.

A policy (also known as an *adversary* or *strategy*) resolves the choice of actions. This induces a (possibly infinite) DTMC.

Definition 8 (Policy). A policy $\pi : Paths_{fin} \rightarrow A$ maps every finite path $w = s_0 \langle a_0, s_1 \rangle \langle a_1, s_2 \rangle \dots \langle a_{n-1}, s_n \rangle$ to an element of $A(s_n)$, i.e., $\pi(w) \in A(s_n)$.¹

The set of all policies is Π . A policy is *stationary* if $\pi(w)$ only depends on the most recent state s_n of w denoted by $last(w)$. Under policy π the action choices on $\mathcal{M}$ are determined. This gives us a DTMC whose states correspond to the finite paths of the MDP. We call this the DTMC induced on $\mathcal{M}$ by policy π and denote this by $\mathcal{D}^\pi$. $Paths^\pi(s)$ is shorthand for $Paths(s)$ of $\mathcal{D}^\pi$. Therefore, the probability of the paths of the MDP under π are defined according to $\mathcal{D}^\pi$, whose probability space includes a probability measure Pr^π over infinite paths via cylinder sets that extend a finite path to a set of infinite paths. See [5] for details.

2.2 Linear Temporal Logic

Linear temporal logic (LTL) is a popular formalism used to specify temporal properties. Here, we are interested in LTL interpreted over finite traces, but we must still define LTL as we use a reduction to LTL as one of our two approaches to synthesis.

Definition 9 (LTL Syntax). An LTL formula is built from a set of propositional symbols $Prop$ and is closed under the boolean connectives as well as the “next” operator X and the “until” operator U :

$$\varphi ::= \top \mid p \mid (\neg \varphi) \mid (\varphi_1 \wedge \varphi_2) \mid (X \varphi) \mid (\varphi_1 U \varphi_2),$$

where $p \in Prop$.

¹Here, we focus on deterministic policies as they are sufficient for optimality of LTL (and LTL_f) properties on MDPs [5].

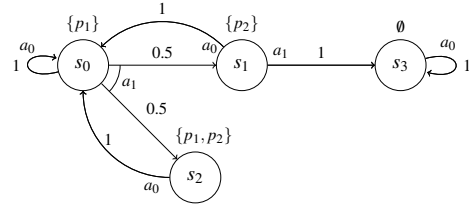


Figure 1: Example MDP.

The common temporal operators “eventually” (F) and “globally” (G) are defined as: $F\varphi = \top U \varphi$ and $G\varphi = \neg F \neg \varphi$. The semantics of LTL are defined over infinite traces (for full definition, see [21]). An LTL formula φ defines an ω -regular language $\mathcal{L}(\varphi)$ over alphabet 2^{Prop} .

Next we define LTL_f . To distinguish between the formulas of the two logics, we use φ for LTL and ϕ for LTL_f formulas.

Definition 10 (LTL_f Syntax & Semantics). *An LTL_f formula has identical syntax to LTL, but the semantics is defined over finite traces. Let $|\rho|$ and ρ_i denote the length of trace ρ and the symbol in the i^{th} position in ρ , respectively, and $\rho, i \models \phi$ is read as: “the i^{th} step of trace ρ is a model of ϕ .” Then,*

- $\rho, i \models \top$;
- $\rho, i \models p$ iff $p \in \rho_i$;
- $\rho, i \models \neg \phi$ iff $\rho, i \not\models \phi$;
- $\rho, i \models \phi_1 \wedge \phi_2$ iff $\rho, i \models \phi_1$ and $\rho, i \models \phi_2$;
- $\rho, i \models X\phi$ iff $|\rho| > i + 1$ and $\rho, i + 1 \models \phi$;
- $\rho, i \models \phi_1 U \phi_2$ iff $\exists j$ s.t. $i \leq j < |\rho|$ and $\rho, j \models \phi_2$ and $\forall k, i \leq k < j, \rho, k \models \phi_1$.

We say finite trace ρ satisfies formula ϕ , denoted by $\rho \models \phi$, iff $\rho, 0 \models \phi$. An LTL_f formula ϕ defines a language $\mathcal{L}(\phi)$ over the alphabet 2^{Prop} . $\mathcal{L}(\phi)$ is a star-free regular language [10].

2.3 Satisfaction of Temporal Logic Specification

Here, we define what it means for an MDP to satisfy an LTL or LTL_f formula.

Definition 11 (Path satisfying LTL). *For a pair $(\mathcal{M}, \varphi)$ of an MDP and an LTL formula where the atomic propositions of $\mathcal{M}$ match the propositions of φ (i.e., $Prop = AP$), we say that an infinite path w on $\mathcal{M}$ satisfies specification φ if the labeling of w is in the language of φ , i.e., $L(w) \in \mathcal{L}(\varphi)$.*

Following [27], we define finite satisfaction (of an LTL_f formula) as follows.

Definition 12 (Path satisfying LTL_f). *For a pair $(\mathcal{M}, \phi)$ of an MDP and an LTL_f formula where the atomic propositions of $\mathcal{M}$ match the propositions of ϕ (i.e., $Prop = AP$), we say that a (possibly finite) path w of $\mathcal{M}$ satisfies specification ϕ if at least one prefix of w is in the language of ϕ , i.e.,*

$$w \models \phi \iff \exists w' \in \text{pre}(w) \text{ s.t. } L(w') \in \mathcal{L}(\phi). \quad (3)$$

Intuitively, this corresponds to a system that can declare its execution complete after satisfying its goals. LTL_f is suitable for specifications that are to be completed in finite time.

Definition 13 (Probability of LTL_f satisfaction). *The probability of satisfying an LTL_f property ϕ in $\mathcal{M}$ under policy π is $Pr(\mathcal{M}^\pi \models \phi) = Pr^\pi(w \in \text{Paths}^\pi(s_{init}) \mid w \models \phi)$.*

2.4 Problem Statement

We formalize the problem of LTL_f synthesis on MDPs as:

Problem 1 (LTL_f synthesis on MDPs). *Given MDP $\mathcal{M}$ and an LTL_f formula ϕ , compute a policy π^* that maximizes the probability of satisfying ϕ , i.e.,*

$$\pi^* = \arg \max_{\pi \in \Pi} Pr(\mathcal{M}^\pi \models \phi),$$

as well as this probability, i.e., $Pr(\mathcal{M}^{\pi^*} \models \phi)$.

3 Synthesis Algorithms

We introduce two approaches to LTL_f policy synthesis on MDPs. The first approach is based on reduction to classical LTL policy synthesis on MDPs. The second approach is through a translation of LTL_f formulas to first order logic formulas, which can be translated to a symbolic deterministic automaton. In this section, we detail these two algorithms, and in Section 4, we show that the second approach scales better than the classical LTL approach.

3.1 Reduction to LTL Synthesis

We can reduce the problem of LTL_f policy synthesis on MDPs to a classical LTL policy synthesis. The algorithm consists of two main steps: (1) construction of an MDP $\mathcal{M}'$ from $\mathcal{M}$ by augmenting it with an additional state and atomic proposition, and (2) translation of LTL_f formula ϕ on the labels of $\mathcal{M}$ to its equivalent LTL formula ϕ on the labels of $\mathcal{M}'$.

3.1.1 Augmented MDP

Recall that the semantics of LTL_f formulas is over finite traces whereas the interpretation of LTL formulas is over infinite traces. In order to reduce the LTL_f synthesis problem to an LTL one, we need to be able to capture the finite paths (traces) of $\mathcal{M}$ that satisfy ϕ and extend them to infinite paths (traces). Specifically, we need those satisfying finite paths that contain no ϕ -satisfying proper prefixes, i.e., satisfy ϕ for the first time. To this end, we allow the system (policy) to decide when to “terminate.” We view the system to be “alive” until termination, at which point it is no longer alive. Then, we define an LTL formula that requires the system to be alive while it has not satisfied ϕ and to terminate after satisfying ϕ .

To this end, we augment MDP $\mathcal{M}$ with a terminal action and state and an atomic proposition *alive*. Formally, we construct MDP $\mathcal{M}' = (S', A', P', s'_{\text{init}}, AP', L')$, where $S' = S \cup \{s_{\text{term}}\}$, $A' = A \cup \{a_{\text{term}}\}$, $s'_{\text{init}} = s_{\text{init}}$, $AP' = AP \cup \{\text{alive}\}$,

$$A'(s) = \begin{cases} A(s) \cup \{a_{\text{term}}\} & \text{if } s \neq s_{\text{term}} \\ \{a_{\text{term}}\} & \text{if } s = s_{\text{term}} \end{cases}, \quad L'(s) = \begin{cases} L(s) \cup \{\text{alive}\} & \text{if } s \neq s_{\text{term}} \\ \emptyset & \text{if } s = s_{\text{term}} \end{cases}, \quad (4)$$

$$P'(s, a, s') = \begin{cases} P(s, a, s') & \text{if } s \in S, a \in A, s' \in S \\ 0 & \text{if } s \in S, a \in A, s' = s_{\text{term}} \\ 1 & \text{if } s \in S', a = a_{\text{term}}, s' = s_{\text{term}} \end{cases}.$$

In this MDP, the system can decide to terminate by taking action a_{term} , in which case it transitions to state s_{term} with probability one and remains there forever. The labeling of the corresponding path includes the atomic proposition *alive* at every time step until s_{term} is visited and is empty thereafter.

Example 2. Figure 2 illustrates the augmented MDP $\mathcal{M}'$ constructed from the example MDP $\mathcal{M}$ in Figure 1. State s_{term} along with the dashed edges and atomic proposition *alive* are added to $\mathcal{M}$. The dashed edges are enabled by action a_{term} and have transition probability of one. The label of s_{term} is the empty set, and the labels of the rest of the states contain *alive*.

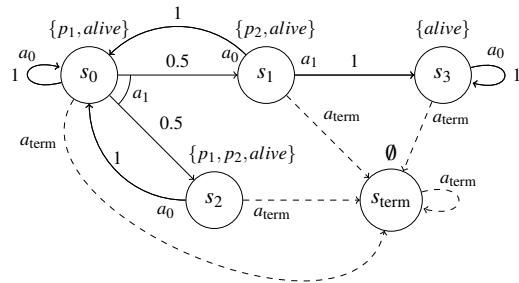


Figure 2: Augmented MDP constructed from $\mathcal{M}$ in Fig. 1.

With this augmentation, the system is able to terminate once a satisfying finite path is generated. We show that, for finite properties, the maximum probability of satisfaction in $\mathcal{M}$ equals the maximum probability of satisfaction in $\mathcal{M}'$ under a policy that enforces a visit to s_{term} .

Let $\text{Paths}_{\text{fin},\phi} \subseteq \text{Paths}_{\text{fin}}(s_{\text{init}})$ be a set of finite paths of interest in arbitrary MDP $\mathcal{M}$. Denote the probability of satisfaction of infinite paths with prefixes from $\text{Paths}_{\text{fin},\phi}$ under policy $\pi \in \Pi$ by:

$$\Pr(\mathcal{M}^\pi \models \text{Paths}_{\text{fin},\phi}) = \Pr^\pi(w \in \text{Paths}^\pi(s_{\text{init}}) \mid \text{pre}(w) \cap \text{Paths}_{\text{fin},\phi} \neq \emptyset). \quad (5)$$

The following lemma states the equivalence of the maximum probability of satisfying this set of paths in $\mathcal{M}$ and $\mathcal{M}'$.

Lemma 1. *Let $\text{Paths}_{\text{fin},\phi} \subseteq \text{Paths}_{\text{fin}}(s_{\text{init}})$ be a set of finite paths of interest of MDP $\mathcal{M}$. Further, let $\mathcal{M}'$ be the augmented version of $\mathcal{M}$ and Π' be the set of policies of $\mathcal{M}'$ that enforce a visit to state s_{term} (take action a_{term}). Then, it holds that*

$$\max_{\pi \in \Pi} \Pr(\mathcal{M}^\pi \models \text{Paths}_{\text{fin},\phi}) = \max_{\pi' \in \Pi'} \Pr(\mathcal{M}'^{\pi'} \models \text{Paths}_{\text{fin},\phi}). \quad (6)$$

Proof. We prove the results by showing that $\text{LHS} \leq \text{RHS}$ and then $\text{LHS} \geq \text{RHS}$.

Case 1. We want to prove that:

$$\max_{\pi \in \Pi} \Pr(\mathcal{M}^\pi \models \text{Paths}_{\text{fin},\phi}) \leq \max_{\pi' \in \Pi'} \Pr(\mathcal{M}'^{\pi'} \models \text{Paths}_{\text{fin},\phi}).$$

Consider policy $\pi \in \Pi$ on $\mathcal{M}$. We can construct policy π' on $\mathcal{M}'$ from π as follows:

$$\pi'(w_{\text{fin}}) = \begin{cases} \pi(w_{\text{fin}}) & \text{if } w_{\text{fin}} \notin \text{Paths}_{\text{fin},\phi} \wedge \text{last}(w_{\text{fin}}) \neq s_{\text{term}} \\ a_{\text{term}} & \text{otherwise.} \end{cases}$$

Recall the probability measure given by the cylinder set $\Pr(\text{Cyl}(s_0 \dots s_n)) = \prod_{0 \leq i < n} P(s_i, \pi^*(s_i), s_{i+1})$ and that the transition probability under a_{term} is always 1. Thus from optimal policy π^* on $\mathcal{M}$, we can construct π' on $\mathcal{M}'$ with equal probability.

Case 2. Next we prove that:

$$\max_{\pi \in \Pi} \Pr(\mathcal{M}^\pi \models \text{Paths}_{\text{fin},\phi}) \geq \max_{\pi' \in \Pi'} \Pr(\mathcal{M}'^{\pi'} \models \text{Paths}_{\text{fin},\phi}),$$

Consider an optimal policy π'^* on $\mathcal{M}'$. By assumption, we know that π'^* chooses a_{term} at some finite point along any path. Consider an arbitrary path and let a_{term} be chosen as the $n+1^{\text{th}}$ action: $w'^* = s_0 \dots s_n s_{\text{term}} \dots$. Let $k \leq n$ be the smallest natural number such that $w' = s_0 \dots s_k \in \text{Paths}_{\text{fin},\phi}$. Then consider policy π' which takes action a_{term} at state s_k and with probability 1 transitioned to state s_{term} (π' and π'^* may be identical). Recall the probability measure given by the cylinder set $\Pr(\text{Cyl}(s_0 \dots s_n)) = \prod_{0 \leq i < n} P(s_i, \pi'^*(s_i), s_{i+1})$. Thus, because a_{term} has probability 1, $\Pr(\mathcal{M}'^{\pi'} \models \text{Paths}_{\text{fin},\phi}) \geq \Pr(\mathcal{M}'^{\pi'^*} \models \text{Paths}_{\text{fin},\phi})$. We can map π' to a policy π for $\mathcal{M}$ with the same probability of satisfaction by changing every choice of a_{term} to an arbitrary available action. Thus, from (optimal policy) π' we can construct a policy π on $\mathcal{M}$ with an equal probability of satisfaction. □

3.1.2 LTL_f to LTL

To translate an LTL_f formula on $\mathcal{M}$ to its equivalent LTL formula on $\mathcal{M}'$, we follow [10]. Let Φ_f be the set of LTL_f formulas ϕ defined over atomic propositions in AP and Φ be the set of LTL formulas φ defined over AP' . Then $g : \Phi_f \rightarrow \Phi$ is defined as:

$$g(\phi) = t(\phi) \wedge (\text{alive} U (G \neg \text{alive})), \quad (7)$$

where $t : \Phi_f \rightarrow \Phi$ is inductively defined as:

- $t(p) = (p \wedge \text{alive})$, where $p \in AP$;
- $t(\neg \phi) = \neg t(\phi)$;
- $t(\phi_1 \wedge \phi_2) = t(\phi_1) \wedge t(\phi_2)$;
- $t(X \phi) = X (\text{alive} \wedge t(\phi))$;
- $t(\phi_1 U \phi_2) = t(\phi_1) U (\text{alive} \wedge t(\phi_2))$.

In this construction, mapping t ensures that *alive* is present in the last letter of every finite trace that satisfies ϕ . Then, g translates ϕ to φ by requiring *alive* to be true until ϕ is satisfied and false thereafter. In other words, the translated LTL formula $\varphi = g(\phi)$ requires the system to terminate after it satisfies ϕ .

Theorem 1. *Given an MDP $\mathcal{M}$ and its augmentation $\mathcal{M}'$ as well as an LTL_f formula ϕ and the corresponding LTL formula $g(\phi)$, then*

$$\max_{\pi \in \Pi} Pr(\mathcal{M} \models \phi) = \max_{\pi' \in \Pi'} Pr(\mathcal{M}' \models g(\phi)),$$

where Π' is the set of all policies in $\mathcal{M}'$.

Proof. This follows through an application of Lemma 1 and a proof that the finite paths on $\mathcal{M}$ satisfying ϕ correspond to the infinite paths on $\mathcal{M}'$ satisfying $g(\phi)$. First we note that the termination assumption of Lemma 1 (i.e., that policies visit s_{term}) is met because $g(\phi)$ contains the clause $(\text{alive} U (G \neg \text{alive}))$. For all $s \in \mathcal{M}$ and $s' \in \mathcal{M}' \setminus \{s_{\text{term}}\} : L(s) = L(s') \setminus \{\text{alive}\}$. Recall that s_{term} is a sink state with label $\emptyset$. Then from [10] it follows that finite paths of $\mathcal{M}$ that satisfy ϕ correspond to paths in $\mathcal{M}'$ that satisfy $g(\phi)$. $\square$

A direct result of the above theorem is the reduction of LTL_f policy synthesis to classical LTL policy synthesis on MDPs.

Corollary 1. *The policy synthesis to maximize the probability of satisfying LTL_f formula ϕ on $\mathcal{M}$ can be reduced to the LTL maximal policy synthesis problem.*

Therefore, to solve Problem 1, we can use LTL synthesis on augmented MDP $\mathcal{M}'$ and property $\varphi = g(\phi)$. The general LTL synthesis algorithm is well-established [5] and follows the following steps: (1) translation of the LTL formula to a DRA, (2) composition of the DRA with the MDP, which results in another MDP called the product MDP, (3) identification of the maximal end-components on the product MDP that satisfy the accepting condition of DRA, and finally (4) solving the maximal reachability probability problem (stochastic shortest path problem) [9] on the product MDP with the accepting end-components as the target states. There exist many tools such as PRISM [18] that can solve the LTL synthesis problem. Specifically, PRISM has a symbolic implementation of this algorithm, enabling fast computations. In Section 4, however, we show that the native approach introduced below outperforms the LTL-reduction approach even by using PRISM's symbolic engine.

3.2 Native Approach

In the native approach, we first convert the LTL_f formula into a *deterministic finite automaton* (DFA) using specialized tools [27]. Then we take the product of this automaton with the MDP. Finally, we synthesize a strategy by solving the maximal reachability probability problem on this product MDP.

3.2.1 Translation to DFA

A Deterministic Finite Automaton (DFA) is a tuple: $\mathcal{A} = (Q, \Sigma, \delta, q_0, F)$, where Q is the set of states, Σ the alphabet, $\delta : Q \times \Sigma \mapsto Q$ the transition function, q_0 the initial state and F the set of accept states. A finite *run* of a DFA on a trace $\rho = \rho_0\rho_1 \dots \rho_n$, where $\rho_i \in \Sigma$, is the sequence of states $q_0q_1 \dots q_{n+1}$ such that $q_{i+1} = \delta(q_i, \rho_i)$ for all $0 \leq i \leq n$. This run is accepting if $q_{n+1} \in F$.

Following [27], we translate LTL_f to a DFA using MONA [16]. LTL_f is expressively equivalent to First-order Logic on finite words, which in turn is a fragment of Weak Second-order Theory of One Successor. We use the translation given in [10] to convert LTL_f formula ϕ to a First-order Logic formula. MONA offers translations from Weak Second-order Theory of One or Two Successors (WS1S/WS2S) to a DFA that accepts precisely the language of our LTL_f formula ϕ . We denote this DFA by $\mathcal{A}_\phi$.

3.2.2 Product of DFA with MDP

Given DFA $\mathcal{A}_\phi$, we can take the product with the MDP $\mathcal{M}$ to achieve a new MDP $\mathcal{M} \times \mathcal{A}_\phi$ as follows. The product of MDP $\mathcal{M} = (S, A, P, s_{\text{init}}, AP, L)$ and DFA $\mathcal{A}_\phi = (Q, \Sigma, q_0, \delta, F)$ is an MDP

$$\mathcal{M} \times \mathcal{A}_\phi = (S \times Q, A, P^{\mathcal{M} \times \mathcal{A}_\phi}, (s_{\text{init}}, q_{\text{init}})),$$

where $q_{\text{init}} = \delta(q_0, L(s_{\text{init}}))$, and

$$P^{\mathcal{M} \times \mathcal{A}_\phi}((s, q), a, (s', q')) = \begin{cases} P(s, a, s') & \text{if } q' = \delta(q, L(s)) \\ 0 & \text{otherwise} \end{cases}.$$

The paths of this product MDP have one-to-one correspondence to the paths of MDP $\mathcal{M}$ as well as the runs of $\mathcal{A}_\phi$. Therefore, the projection of the paths of $\mathcal{M} \times \mathcal{A}_\phi$ that reach state (s, q) , where $q \in F$, on $\mathcal{A}_\phi$ are accepting runs and on $\mathcal{M}$ are ϕ -satisfy paths. Thus, we solve Problem 1 by synthesizing an optimal policy on this product MDP, using standard tools for the maximal reachability probability problem as discussed in Sec. 3.1. The resulting policy is stationary on the product MDP but history-dependent on $\mathcal{M}$.

In theory, both the LTL approach and the native approach exhibit runtimes doubly-exponential in the size of the formula, due to the need to construct deterministic automata [8, 11]. In practice, the native approach's translation to a DFA using specialized tools offers better runtime and memory usage compared to the LTL-based approach. Additionally, it produces a minimal DFA, while the LTL pipeline produces an ω -automaton which cannot be minimized effectively by existing tools. We show the benefits of the native approach experimentally in Sec. 4.

4 Evaluation

We evaluate the proposed synthesis approaches through a series of benchmarking case studies. Below, we provide details on our implementation, experimental scenarios, and obtained results. A version of our tool along with examples is available on GitHub [25].

4.1 Experimental Framework

We run our experiments using the PRISM framework [18]. PRISM uses a symbolic encoding of MDPs as well as an encoding of automata as a list of edges, where the labels of edges are encoded symbolically using BDDs. PRISM supports several tools for the LTL-to-automata translation. We tested PRISM’s built-in translator as well as Rabinizer, LTL3DRA and SPOT [17, 2, 12]. Of these, PRISM’s built in conversion and SPOT performed significantly better than the others, and were used for evaluation.

We note that we also considered a comparison study against MoChiBa [22], a tool based on PRISM that uses Limit-Deterministic Büchi Automata. Unfortunately, it is not possible to conduct this study since the current implementation of MoChiBa does not support loading automata from external tools and only runs with a modified implementation of PRISM’s explicit engine. Our tests use PRISM’s Hybrid engine (enabled by default). We leave a comparison of our DFA-based approach for LTL_f synthesis to an approach based on Limit-Deterministic Büchi Automata to future work.

In the implementation of the LTL-reduction approach, we augment the MDP and convert the LTL_f formula into an equivalent LTL formula as described in Section 3.1. Then, we input both the LTL formula and the modified MDP into PRISM for synthesis. In the implementation of the native approach, we invoke PRISM on the original LTL_f formula and MDP and use an external tool to convert the LTL_f formula into a DFA using MONA [16] then convert from MONA’s format to the HOA format [1]. Note that external tools (SPOT and our native approach) read from hard disk, whereas using PRISM’s built-in conversion avoids this. Nevertheless, even including this time, the native approach (and sometimes SPOT) typically gives better performance as shown below. All experiments are run on a computer with an Intel i7-8550U and 16GB of RAM. PRISM is run using the default settings.

4.2 Experimental Scenarios

4.2.1 Test MDPs

We consider four types of MDPs: *Gridworld*, *Dining Philosophers*, *Nim*, and *Double Counter*. In *Gridworld*, an agent is given some goal as an LTL_f formula and must maximize the probability of satisfaction. The agent has four actions: North, South, East, and West. Under each action, the probability of moving to the cell in the intended direction is 0.69, then 0.01 for its opposite cell, and finally 0.1 for each of the other directions and for remaining in the same cell. If the resulting movement would place the agent in collision with the boundary, the agent remains in its current cell. We model the motion of this agent as an MDP, where the states correspond to the cells of the grid and the set of actions is $\{a_{\text{north}}, a_{\text{south}}, a_{\text{east}}, a_{\text{west}}\}$.

We base the Dining-Philosophers domain on the tutorial from the PRISM website. We consider a ring of five philosophers with two different specifications. Note that typical Dining-Philosophers specifications of interest are infinite-trace; our tests are merely meant to show that our approach works on MDPs other than *Gridworld*. Both *Nim* and *Double Counter* are probabilistic versions of games presented in [20].

4.2.2 Test formulas

Unfortunately, there is no standard set of LTL_f formulas that we can use for benchmarking. In [27], random LTL formulas are used as the basis of LTL_f benchmarks; however, because we also consider an MDP, random formulas are frequently tautologies or non-realizable with respect to the MDP. For instance, consider a randomly generated formula $\phi = p_1 U p_2$, where p_1 and p_2 are randomly assigned to the labels of the states in the *Gridworld* MDP. There is likely no path of the MDP that can satisfy this formula. As a matter of fact, in a test of more than 70 randomly generated LTL_f formulas and *Gridworld* MDPs,

only one was “interesting” (yielding probability between zero and one) on the corresponding MDP. For benchmarking, the interest is in sets of formulas that make sense in a probabilistic setting.

Therefore, as the first set of test formulas for the *Gridworld* MDP, we considered a natural finite-trace specification where the agent has n goals to accomplish in any order, as well as a “safety” property where it must globally avoid some states. This is typical of e.g., a robotics domain [14]. We use **F n** to denote this formula for a given n . For the second set of test formulas, we keep regions to visit and avoid but introduce some ordering and repeat visits. **OS** is a short formula introducing order and **OL** is a longer formula that also contains nested temporal operators. All formulas are given in the appendix of the online version [26].

For the *Dining Philosophers* domains, typical examples focus on infinite-run properties. We test several finite-run properties on these domains. These properties are meant to illustrate our tool on an MDP other than Gridworld, but are not representative of typical finite-trace properties. One version (**D5**) asserts that all philosophers must eat at least once. The other (**D5C**) is a more complex property involving orders of eating. Both formulas are available in the appendix of the online version [26].

For *Nim*, the game of Nim is played against a stochastic opponent. To increase difficulty, the specifications require that randomly chosen stack heights are to be reached or avoided by the system player. For *Double Counter*, two four-bit binary counters are used. One counter is controlled by a stochastic environment, and the other is controlled by the system. The aim of the system is to make its counter match that of the environment.

4.3 Experimental Results

We provide experiments varying the complexity of the formulas and of the MDPs. Total runtime is shown in seconds. We refer to the translation to LTL as the “LTL pipeline” and the translation to a DFA via MONA as the “Native pipeline.” All plots are best viewed in color.

4.3.1 Automata Construction

First we measure how the length of the formula affects the synthesis computation time. For all of these experiments, we use a 10×10 Gridworld as the original MDP. For small formulas, the time needed to read the HOA file from hard disk outweighs the shorter construction time and smaller automata of the native approach as shown in Figure 3a. For formulas longer than **F3**, the native approach offers better computation time (e.g., see Figure 3b and Figure 3c).

PRISM successfully builds automata for formulas up to **F9**. In Figure 3d, we highlight the superiority of the native approach in constructing the automaton for **F10**. For this formula, PRISM runs out of memory when constructing the automaton according to the LTL pipeline. SPOT completes for **F10** but runs out of memory on **F11**. MONA works for formulas up to **F17**, which takes 8.586 seconds to compute, though writing the resulting file to disk is prohibitively expensive (the file is larger than 10GB). MONA runs out of memory constructing the automaton for **F18**.

We also considered a Gridworld with random obstacles (Figure 3e) and with hallways (Figure 3f), both with the formula **F9**. Finally, we consider two other formulas (**OS** and **OL**) in Figure 3g and Figure 3h to demonstrate that our improvement is not tied to the specific form of the specification.

For the Dining-Philosophers example, we consider five philosophers. Again, our results show that the native pipeline is faster and more memory efficient than the other approaches for automata construction (Figure 3i). However, the Dining-Philosophers example illustrates an issue where PRISM’s built-in automata translation sometimes constructs automata such that computing the maximal accepting end-component using BDDs within PRISM is significantly faster (Figure 3k). The automata construction is

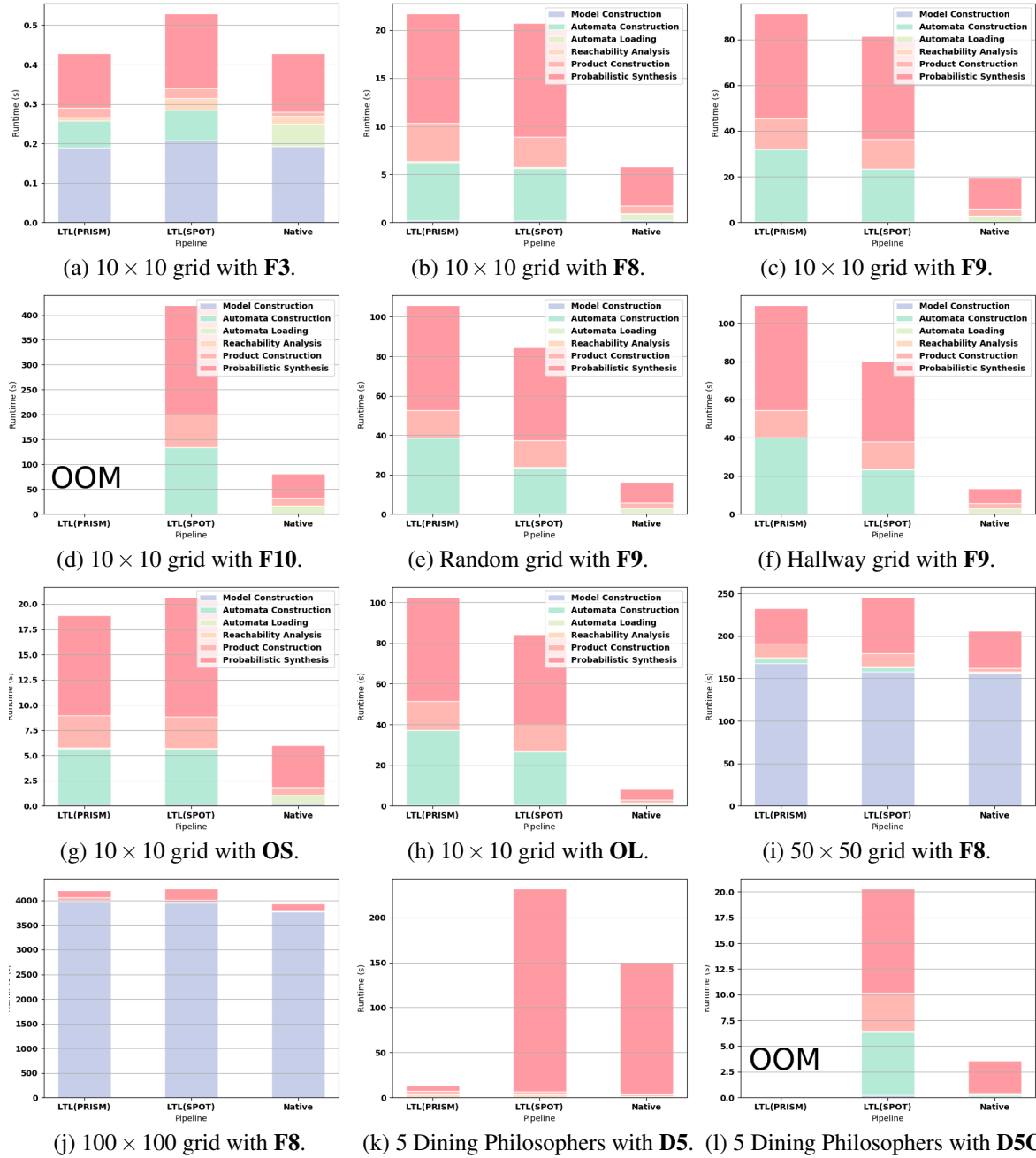


Figure 3: Runtime results for the Gridworld and Dining Philosophers MDPs with the LTL (PRISM and SPOT) and native pipelines.

still slower than both SPOT and our native approach, and the automata returned have more states, but the BDD representation is more efficient. This issue affects some automata, not only for our tool, but also for other external tools we tested. This means not only automata construction and size, but also the BDD representation are important for overall runtime. However, PRISM's built-in translator's memory usage scales more poorly than SPOT or our native approach. On a more complex formula Figure 31, PRISM's built-in construction runs out of memory, even though SPOT and the native approach can

finish construction in less than ten seconds. Thus, for sufficiently complex formulas, PRISM’s built-in translation is not viable. This suggests a need for automata construction methods that are not only more efficient but also produce automata that work well within PRISM.

In the Nim game, the native approach far outperforms both PRISM and PRISM with SPOT (Figure 4a). In the Double-Counter game, PRISM runs out of memory and PRISM with SPOT times out after more than 3 hours, whereas the native approach completes synthesis in less than 5 seconds.

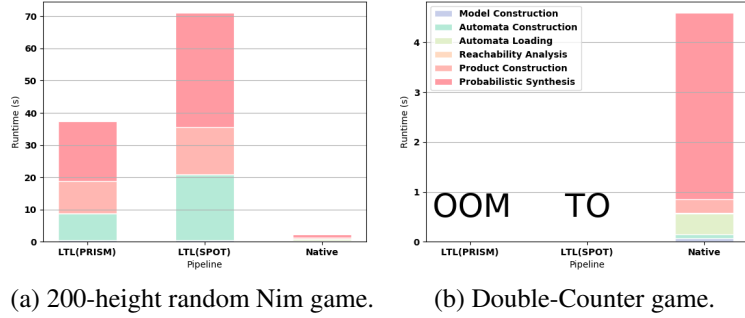


Figure 4: Runtime results for the Nim and Double-Counter games.

Overall we observe that for large formulas the native pipeline offers significantly better scalability than the LTL pipeline (tested with PRISM’s built-in translator, SPOT, Rabinizer, and LTL3DRA automata translators). With an implementation that does not access the hard disk, we expect even better performance of the native pipeline.

4.3.2 Automata Size

The DFA generated by the native pipeline is minimal and typically much smaller than the LTL pipeline’s DRA. SPOT and PRISM typically produce similarly sized DRAs. Because we take the product of the input MDP with these automata, we expect the size of the resulting product to be smaller in the native pipeline. Table 1 shows the automaton sizes for the various formulas. Table 2 shows the sizes of the product MDPs for the LTL and native pipelines respectively.

Interestingly, while the size of the DFA obtained from MONA is roughly half the size of the DRA constructed by PRISM, the final sizes of the products of the MDP and the automata are comparable for both approaches. The number of states, transitions and nondeterministic choices are all measured after reachability analysis is performed. We see time savings in the reachability analysis and product construction phases (both about two times faster), but the final products are similar in size. However, the product from the native pipeline has fewer nondeterministic transitions than the product from the LTL pipeline. The improvements from this are small relative to the time it takes to construct a large gridworld MDP (Figure 3j).

Examples of this are shown in Figure 3i and Figure 3j. Note that the majority of the computation time is spent constructing the MDPs. For the LTL pipeline, this construction takes slightly longer time because of the augmented MDP. It is important to note that the native approach allows us to use larger formulas with these models whereas the LTL pipeline is limited to **F10** or smaller.

In summary, we observe that the native pipeline is the most efficient in terms of both runtime and memory. There are two possible drawbacks to the native pipeline. First, it requires reading from disk, which for small formulas negates the advantage in automata construction time. Second, PRISM’s built-in automata translation sometimes constructs automata whose BDD representations work better within PRISM than automata from any external tools we tested (Figure 3k).

Table 1: Sizes of the automata from LTL (DRA) and native (DFA) pipelines.

	F3	F8	F9	F10	OS	OL	D5	D5C	Nim	Counter
DRA states	19	512	1,027	NA	515	1,028	67	NA	701	NA
DFA states	10	258	514	1,026	258	259	33	29	67	130

Table 2: Sizes of the product MDP in the LTL pipeline and native pipelines.

	10x10 F3	10x10 F8	10x10 F9	50x50 F8	100x100 F8	10x10 F10	10x10 OS
States (LTL)	890	24,678	48,998	641,478	2,568,978	NA	24,678
States (Native)	881	24,421	48,485	641,221	2,568,721	96357	24,421
Transitions (LTL)	17,130	475,030	942,742	13,265,398	53,537,298	NA	475,030
Transitions (Native)	16,256	450,468	893,760	12,623,936	50,968,336	1775424	450,368
Choices (LTL)	4,410	122,358	242,934	3,206,358	12,843,858	NA	122,358
Choices (Native)	3,524	97,684	193,940	2,564,884	10,274,884	385428	97,684
	OL	10x10 rand	10x10 halls	5 Phil D5	5 Phil D5C	Nim	Counter
States (LTL)	24,935	44,253	19,476	4,548,220	NA	2115	NA
States (Native)	24,519	43,740	19,187	1,476,976	93,068	404	449
Transitions (LTL)	475,287	804,097	310,996	46,948,710	NA	8343	NA
Transitions (Native)	452,172	759,856	291,536	7,891,750	494,420	1205	908
Choices (LTL)	122,615	219,209	96,220	44,059,330	NA	6297	NA
Choices (Native)	98,076	174,960	76,748	6,895,580	437,050	808	460

5 Conclusion

We introduced the problem of LTL_f synthesis for probabilistic systems and presented two approaches. The first one is a reduction of LTL_f to LTL with a corresponding augmentation of the MDP. The second approach uses native tools to construct an automaton and then takes the product of this automaton with the MDP to construct a product MDP that can be used for synthesis through standard techniques. We showed that this native approach offers better scalability than the reduction to LTL. Our work opens the door to the use of LTL_f synthesis on practical domains such as robotics, cf. [14]. Our tool is on GitHub [25].

For future work we would like to expand our results to include probability minimization (c.f. Theorem 1). Our native approach extends easily, but there are subtleties that prevent applying the LTL pipeline to minimization. We are also interested in a fully symbolic methodology, where the automaton is represented symbolically and the product is also taken symbolically. The experiment in Figure 3k and our experiences with MoChiBa both indicate that implementing LTL_f methods optimized for PRISM may be more efficient than external tools. We are also interested in applying the work to robotics domains.

Acknowledgments Work on this project by AMW has been supported by NASA 80NSSC17K0162; by ML has been supported by the University of Colorado Boulder Autonomous Systems Interdisciplinary Research Theme and the NSF Center for Unmanned Aircraft Systems; by LEK has been supported in part by NSF grant IIS-1830549; and by MYV has been supported in part by NSF grants IIS-1527668, CCF-1704883, IIS-1830549, and an award from the Maryland Procurement Office. The authors would like to thank Gabriel Santos and Joachim Klein from the University of Oxford for their insightful comments and suggestions for efficient implementation of the framework in PRISM. We would like to thank Salomon Sikert from the Technical University of Munich for his aid in running MoChiBa.

References

- [1] Tomáš Babiak, František Blahoudek, Alexandre Duret-Lutz, Joachim Klein, Jan Křetínský, David Müller, David Parker & Jan Strejček (2015): *The Hanoi Omega-Automata Format*. In: *Computer Aided Verification*, Springer Intl. Publishing, Cham, pp. 479–486, doi:10.1007/978-3-319-21690-4_31.
- [2] Tomáš Babiak, František Blahoudek, Mojmir Křetínský & Jan Strejček (2013): *Effective translation of LTL to deterministic Rabin automata: Beyond the (F, G)-fragment*. In: *Automated Technology for Verification and Analysis*, Springer, pp. 24–39, doi:10.1007/10722167_21.
- [3] Fahiem Bacchus, Craig Boutilier & Adam Grove (1996): *Rewarding Behaviors*. In: *Proc. of the Thirteenth Natl. Conf. on Artificial Intelligence - Volume 2*, AAAI96, AAAI Press, p. 11601167.
- [4] Fahiem Bacchus, Craig Boutilier & Adam Grove (1997): *Structured Solution Methods for Non-Markovian Decision Processes*. In: *Proc. of the (AAAI-97) and (IAAI-97)*, AAAI97/IAAI97, AAAI Press, p. 112117.
- [5] Christel Baier & Joost-Pieter Katoen (2008): *Principles of Model Checking*. The MIT Press.
- [6] Ronen I. Brafman, Giuseppe De Giacomo & Fabio Patrizi (2018): *LTLf/LDLf Non-Markovian Rewards*. In: *Proc. of (AAAI-18), (IAAI-18), and (EAAI-18)*, New Orleans, Louisiana, USA, February 2-7, 2018, AAAI Press, pp. 1771–1778.
- [7] Alberto Camacho, Oscar Chen, Scott Sanner & Sheila A. McIlraith (2017): *Non-Markovian Rewards Expressed in LTL: Guiding Search Via Reward Shaping*. In: *SOCS*.
- [8] Costas Courcoubetis & Mihalis Yannakakis (1995): *The Complexity of Probabilistic Verification*. *J. ACM* 42(4), pp. 857–907, doi:10.1145/210332.210339.
- [9] Luca De Alfaro (1998): *Formal Verification of Probabilistic Systems*. Ph.D. thesis, Stanford, CA, USA.
- [10] Giuseppe De Giacomo & Moshe Y Vardi (2013): *Linear Temporal Logic and Linear Dynamic Logic on Finite Traces*. In: *Intl. Joint Conf. on Artificial Intelligence (IJCAI)*, 13, pp. 854–860.
- [11] Giuseppe De Giacomo & Moshe Y Vardi (2015): *Synthesis for LTL and LDL on Finite Traces*. In: *Intl. Joint Conf. on Artificial Intelligence (IJCAI)*, 15, pp. 1558–1564.
- [12] Alexandre Duret-Lutz, Alexandre Lewkowicz, Amaury Fauchille, Thibaud Michaud, Étienne Renault & Laurent Xu (2016): *Spot 2.0 — a framework for LTL and ω -automata manipulation*. In: *Automated Technology for Verification and Analysis (ATVA'16)*, Lecture Notes in Computer Science 9938, Springer, pp. 122–129, doi:10.1007/3-540-60915-6_6.
- [13] Keliang He, Morteza Lahijanian, Lydia E Kavraki & Moshe Y Vardi (2018): *Automated Abstraction of Manipulation Domains for Cost-Based Reactive Synthesis*. *IEEE Robotics and Automation Letters* 4(2), pp. 285–292, doi:10.1109/LRA.2018.2889191.
- [14] Keliang He, Andrew M Wells, Lydia E Kavraki & Moshe Y Vardi (2019): *Efficient Symbolic Reactive Synthesis for Finite-Horizon Tasks*. In: *2019 Intl. Conf. on Robotics and Automation (ICRA)*, IEEE, pp. 8993–8999, doi:10.1109/ICRA.2019.8794170.
- [15] Keliang He, Morteza Lahijanian, Lydia E. Kavraki & Moshe Y. Vardi (2017): *Reactive Synthesis for Finite Tasks Under Resource Constraints*. In: *Int. Conf. on Intelligent Robots and Systems (IROS)*, IEEE, Vancouver, BC, Canada, pp. 5326–5332, doi:10.1109/IROS.2017.8206426.
- [16] Jesper G Henriksen, Jakob Jensen, Michael Jørgensen, Nils Klarlund, Robert Paige, Theis Rauhe & Anders Sandholm (1995): *Mona: Monadic second-order logic in practice*. In: *Intl. Workshop on Tools and Algorithms for the Construction and Analysis of Systems*, Springer, pp. 89–110, doi:10.1007/3-540-60630-0_5.
- [17] Jan Křetínský, Tobias Meggendorfer, Salomon Sickert & Christopher Ziegler (2018): *Rabinizer 4: From LTL to Your Favourite Deterministic Automaton*. In: *Computer Aided Verification*, Springer International Publishing, Cham, pp. 567–577, doi:10.1007/978-3-319-46520-3_9.
- [18] M. Kwiatkowska, G. Norman & D. Parker (2011): *PRISM 4.0: Verification of Probabilistic Real-time Systems*. In: *Proc. 23rd Intl. Conf. on Computer Aided Verification (CAV'11)*, LNCS 6806, Springer, pp. 585–591, doi:10.1007/3-540-45657-0_17.

- [19] Jianwen Li, Kristin Y Rozier, Geguang Pu, Yueling Zhang & Moshe Y Vardi (2019): *SAT-Based Explicit LTL_f Satisfiability Checking*. In: *Proc. of the AAAI Conf. on Artificial Intelligence*, 33, pp. 2946–2953, doi:10.1007/978-3-030-25543-5_1.
- [20] Lucas Martinelli Tabajara & Moshe Y. Vardi (2019): *Partitioning Techniques in LTL_f Synthesis*. In: *Proceedings of the Twenty-Eighth International Joint Conference on Artificial Intelligence, IJCAI-19*, International Joint Conferences on Artificial Intelligence Organization, pp. 5599–5606, doi:10.24963/ijcai.2019/777.
- [21] Amir Pnueli (1977): *The temporal logic of programs*. In: *Foundations of Computer Science, 1977.*, 18th Annual Symposium on, IEEE, pp. 46–57.
- [22] Salomon Sickert & Jan Kretínský (2016): *MoChiBA: Probabilistic LTL Model Checking Using Limit-Deterministic Büchi Automata*. In: *Automated Technology for Verification and Analysis (ATVA’16), Lecture Notes in Computer Science* 9938, pp. 130–137, doi:10.1007/978-3-319-45856-4.
- [23] Sylvie Thiébaux, Charles Gretton, John Slaney, David Price & Froduald Kabanza (2006): *Decision-Theoretic Planning with Non-Markovian Rewards*. *J. Artif. Int. Res.* 25(1), p. 1774, doi:10.1613/jair.1676.
- [24] Moshe Y. Vardi (1985): *Automatic Verification of Probabilistic Concurrent Finite State Programs*. In: *Proceedings of the 26th Annual Symposium on Foundations of Computer Science, SFCS ’85*, IEEE Computer Society, USA, p. 327338, doi:10.1109/SFCS.1985.12.
- [25] Andrew M. Wells: *LTL_f for PRISM*. Available at https://github.com/andrewmw94/ltlf_prism.
- [26] Andrew M Wells, Morteza Lahijanian, Lydia E Kavraki & Moshe Y Vardi: *LTL_f Synthesis on Probabilistic Systems (Online version)*. Available at <http://www.andrewmwells.com/ltlf-synthesis-on-probabilistic-systems/>.
- [27] Shufang Zhu, Lucas M Tabajara, Jianwen Li, Geguang Pu & Moshe Y Vardi (2017): *Symbolic LTL_f synthesis*. In: *Proc. of the 26th Intl. Joint Conf. on Artificial Intelligence*, AAAI Press, pp. 1362–1369.