



Model Checking Finite-Horizon Markov Chains with Probabilistic Inference

Steven Holtzen¹✉, Sebastian Junges², Marcell Vazquez-Chanlatte²,
Todd Millstein¹, Sanjit A. Seshia², and Guy Van den Broeck¹

¹ University of California, Los Angeles, CA, USA
sholtzen@cs.ucla.edu

² University of California, Berkeley, CA, USA



Abstract. We revisit the symbolic verification of Markov chains with respect to finite horizon reachability properties. The prevalent approach iteratively computes step-bounded state reachability probabilities. By contrast, recent advances in probabilistic inference suggest symbolically representing all horizon-length paths through the Markov chain. We ask whether this perspective advances the state-of-the-art in probabilistic model checking. First, we formally describe both approaches in order to highlight their key differences. Then, using these insights we develop RUBICON, a tool that transpiles PRISM models to the probabilistic inference tool Dice. Finally, we demonstrate better scalability compared to probabilistic model checkers on selected benchmarks. All together, our results suggest that probabilistic inference is a valuable addition to the probabilistic model checking portfolio, with RUBICON as a first step towards integrating both perspectives.

1 Introduction

Systems with probabilistic uncertainty are ubiquitous, e.g., probabilistic programs, distributed systems, fault trees, and biological models. Markov chains replace nondeterminism in transition systems with probabilistic uncertainty, and *probabilistic model checking* [4, 7] provides model checking algorithms. A key property that probabilistic model checkers answer is: *What is the (precise) probability that a target state is reached (within a finite number of steps h)?* Contrary to classical *qualitative* model checking and approximate variants of probabilistic model checking, precise probabilistic model checking must find the total probability of *all* paths from the initial state to any target state.

S. Holtzen and S. Junges—Contributed equally.

This work is partially supported by NSF grants #IIS-1943641, #IIS-1956441, #CCF-1837129, DARPA grant #N66001-17-2-4032, a Sloan Fellowship, a UCLA Dissertation Year Fellowship, and gifts by Intel and Facebook Research.

This work is partially supported by NSF grants 1545126 (VeHiCaL), 1646208 and 1837132, by the DARPA contracts FA8750-18-C-0101 (AA) and FA8750-20-C-0156 (SDCPS), by Berkeley Deep Drive, and by Toyota under the iCyPhy center.

© The Author(s) 2021

A. Silva and K. R. M. Leino (Eds.): CAV 2021, LNCS 12760, pp. 577–601, 2021.

https://doi.org/10.1007/978-3-030-81688-9_27

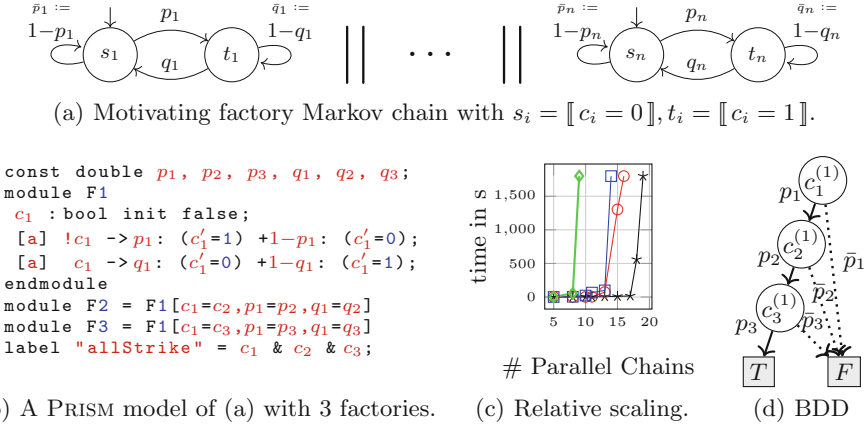


Fig. 1. Motivating example. Figure 1(c) compares the performance of RUBICON ($\text{---}\star\text{---}$), STORM’s explicit engine ($\text{---}\circ\text{---}$), STORM’s symbolic engine ($\text{---}\square\text{---}$) and PRISM ($\text{---}\diamond\text{---}$) when invoked on a (b) with arbitrarily fixed (different) constants for p_i, q_i and horizon $h = 10$. Times are in seconds, with a time-out of 30 min.

Nevertheless, the prevalent ideas in probabilistic model checking are generalizations of qualitative model checking. Whereas qualitative model checking tracks the states that can reach a target state (or dually, that can be reached from an initial state), probabilistic model checking tracks the i -step reachability probability for each state in the chain. The $i+1$ -step reachability can then be computed via multiplication with the *transition matrix*. The scalability concern is that this matrix grows with the state space in the Markov chain. Mature model checking tools such as STORM [36], Modest [34], and PRISM [51] utilize a variety of methods to alleviate the state space explosion. Nevertheless various natural models cannot be analyzed by the available techniques.

In parallel, within the AI community a different approach to representing a distribution has emerged, which on first glance can seem unintuitive. Rather than marginalizing out the paths and tracking reachability probabilities per state, the probabilistic AI community commonly aggregates all *paths* that reach the target state. At its core, inference is then a weighted sum over all these paths [16]. This hinges on the observation that this set of paths can often be stored more compactly, and that the probability of two paths that share the same prefix or suffix can be efficiently computed on this concise representation. This *inference technique* has been used in a variety of domains in the artificial intelligence (AI) and verification communities [9, 14, 27, 39], but is not part of any mature probabilistic model checking tools.

This paper theoretically and experimentally compares and contrasts these two approaches. In particular, we describe and motivate RUBICON, a probabilistic model checker that *leverages the successful probabilistic inference techniques*. We

begin with an example that explains the core ideas of RUBICON followed by the paper structure and key contributions.

Motivating Example. Consider the example illustrated in Fig. 1(a). Suppose there are n factories. Each day, the workers at each factory collectively decide whether or not to strike. To simplify, we model each factory (i) with two states, striking (t_i) and not striking (s_i). Furthermore, since no two factories are identical, we take the probability to begin striking (p_i) and to stop striking (q_i) to be different for each factory. Assuming that each factory transitions synchronously and in parallel with the others, we query: “what is the probability that all the factories are simultaneously striking within h days?”

Despite its simplicity, we observe that state-of-the-art model checkers like STORM and PRISM do not scale beyond 15 factories.¹ For example, Fig. 1(b) provides a PRISM encoding for this simple model (we show the instance with 3 factories), where a Boolean variable c_i is used to encode the state of each factory. The “allStrike” label identifies the target state. Figure 1(c) shows the run time for an increasing number of factories. While all methods eventually time out, RUBICON scales to systems with an order of magnitude more states.

Why is This Problem Hard? To understand the issue with scalability, observe that tools such as STORM and PRISM store the transition matrix, either explicitly or symbolically using algebraic decision diagrams (ADDs). Every distinct entry of this transition matrix needs to be represented; in the case of ADDs using a unique leaf node. Because each factory in our example has a different probability of going on strike, that means each subset of factories will likely have a unique probability of jointly going on strike. Hence, the transition matrix then will have a number of distinct probabilities that is exponential in the number of factories, and its representation as an ADD must blow up in size. Concretely, for 10 factories, the size of the ADD representing the transition matrix has 1.9 million nodes. Moreover, the explicit engine fails due to the dense nature of the underlying transition matrix. We discuss this method in Sect. 3.

How to Overcome This Limitation? This problematic combinatorial explosion is often unnecessary. For the sake of intuition, consider the simple case where the horizon is 1. Still, the standard transition matrix representations blow up exponentially with the number of factories n . Yet, the probability of reaching the “allStrike” state is easy to compute, even when n grows: it is $p_1 \cdot p_2 \cdots p_n$.

RUBICON aims to compute probabilities in this compact *factorized* way by representing the computation as a binary decision diagram (BDD). Figure 1(d) gives an example of such a BDD, for three factories and a horizon of one. A key property of this BDD, elaborated in Sect. 3, is that it can be interpreted as a *parametric Markov chain*, where the weight of each edge corresponds with the probability of a particular factory striking. Then, the probability that the goal state is reached is given by the weighted sum of paths terminating in T : for this instance, there is a single such path with weight $p_1 \cdot p_2 \cdot p_3$. These BDDs are tree-like Markov-chains, so model checking can be performed in time linear in the size

¹ Section 6 describes the experimental apparatus and our choice of comparisons.

of the BDD using dynamic programming. Essentially, the BDD represents the set of paths that reach a target state—an idea common in probabilistic inference.

To construct this BDD, we propose to encode our reachability query symbolically as a *weighted model counting* (WMC) query on a logical formula. By compiling that formula into a BDD, we obtain a diagram where computing the query probability can be done efficiently (in the size of the BDD). Concretely for Fig. 1(d), the BDD represents the formula $c_1^{(1)} \wedge c_2^{(1)} \wedge c_3^{(1)}$, which encodes all paths through the chain that terminate in the goal state (all factories strike on day 1). For this example and this horizon, this is a single path. WMC is a well-known strategy for probabilistic inference and is currently the among the state-of-the-art approaches for discrete graphical models [16], discrete probabilistic programs [39], and probabilistic logic programs [27].

In general, the exponential growth of the number of paths might seem like it dooms this approach: for $n = 3$ factories and horizon $h = 1$, we need to only represent 8 paths, but for $h = 2$, we would need to consider 64 different paths, and so on. However, a key insight is that, for many systems – such as the factory example – the structural compression of BDDs allows a concise representation of exponentially many paths, all *while* being parametric over path probabilities (see Sect. 4). To see why, observe that in the above discussion, the state of each factory is *independent* of the other factories: independence, and its natural generalizations like *conditional* and *contextual* independence, are the driving force behind many successful probabilistic inference algorithms [47]. Succinctly, the key advantage of RUBICON is that it exploits a form of structure that has thus far been under-exploited by model checkers, which is why it scales to more parallel factories than the existing approaches on the hard task. In Sect. 6 we consider an extension to this motivating example that adds dependencies between factories. This dependency (or rather, the accompanying increase in the size of the underlying MC) significantly decreases scalability for the existing approaches but negligibly affects RUBICON.

This leads to the task: *how does one go from a PRISM model to a concise BDD efficiently?* To do this, RUBICON leverages a novel translation from PRISM models into a probabilistic programming language called *Dice* (outlined in Sect. 5).

Contribution and Structure. Inspired by the example, we contribute conceptual and empirical arguments for leveraging BDD-based probabilistic inference in model checking. Concretely:

1. We demonstrate fundamental advantages in using probabilistic inference on a natural class of models (Sect. 1 and 6).
2. We explain these advantages by showing the fundamental differences between existing model checking approaches and probabilistic inference (Sect. 3 and 4). To that end, Sect. 4 presents probabilistic inference based on an operational and a logical perspective and combines these perspectives.
3. We leverage those insights to build RUBICON, a tool that transpiles PRISM to *Dice*, a probabilistic programming language (Sect. 5).

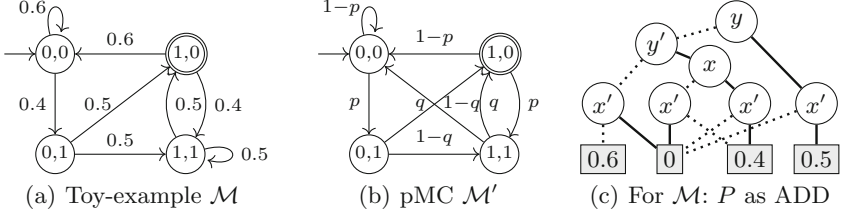


Fig. 2. (a) MC toy example (b) (distinct) pMC toy example (c) ADD transition matrix

4. We demonstrate that RUBICON indeed attains an order-of-magnitude scaling improvement on several natural problems including sampling from parametric Markov chains and verifying network protocol stabilization (Sect. 6).

Ultimately we argue that RUBICON makes a valuable contribution to the portfolio of probabilistic model checking backends, and brings to bear the extensive developments on probabilistic inference to well-known model checking problems.

2 Preliminaries and Problem Statement

We state the problem formally and recap relevant concepts. See [7] for details. We sometimes use $\bar{p}$ to denote $1-p$. A *Markov chain* (MC) is a tuple $\mathcal{M} = \langle S, \iota, P, T \rangle$ with S a (finite) set of *states*, $\iota \in S$ the *initial state*, $P: S \rightarrow \text{Distr}(S)$ the *transition function*, and T a set of *target states* $T \subseteq S$, where $\text{Distr}(S)$ is the set of distributions over a (finite) set S . We write $P(s, s')$ to denote $P(s)(s')$ and call P a *transition matrix*. The successors of s are $\text{Succ}(s) = \{s' \mid P(s, s') > 0\}$. To support MCs with billions of states, we may describe MCs symbolically, e.g., with PRISM [51] or as a probabilistic program [42, 48]. For such a symbolic description $\mathcal{P}$, we denote the corresponding MC with $\llbracket \mathcal{P} \rrbracket$. States then reflect assignments to symbolic variables.

A *path* $\pi = s_0 \dots s_n$ is a sequence of states, $\pi \in S^+$. We use $\pi_{\downarrow}$ to denote the *last state* s_n , and the *length* of π above is n and is denoted $|\pi|$. Let Paths_h denote the paths of length h . The probability of a path is the product of the transition probabilities, and may be defined inductively by $\Pr(s) = 1$, $\Pr(\pi \cdot s) = \Pr(\pi) \cdot P(\pi_{\downarrow}, s)$. For a fixed *horizon* h and set of states T , let the set $\llbracket s \rightarrow \Diamond^{\leq h} T \rrbracket = \{\pi \mid \pi_0 = s \wedge |\pi| \leq h \wedge \pi_{\downarrow} \in T \wedge \forall i < |\pi|. \pi_i \notin T\}$ denote paths from s of length at most h that terminate at a state contained in T . Furthermore, let $\Pr_{\mathcal{M}}(s \models \Diamond^{\leq h} T) = \sum_{\pi \in \llbracket s \rightarrow \Diamond^{\leq h} T \rrbracket} \Pr(\pi)$ describe the probability to reach T within h steps. We simplify notation when $s = \iota$ and write $\llbracket \Diamond^{\leq h} T \rrbracket$ and $\Pr_{\mathcal{M}}(\Diamond^{\leq h} T)$, respectively. We omit $\mathcal{M}$ whenever that is clear from the context.

Formal Problem: Given an MC $\mathcal{M}$ and a horizon h , compute $\Pr_{\mathcal{M}}(\Diamond^{\leq h}T)$.

Example 1. For conciseness, we introduce a toy example MC $\mathcal{M}$ in Fig. 2(a). For horizon $h = 3$, there are three paths that reach state $\langle 1, 0 \rangle$: For example the path $\langle 0, 0 \rangle \langle 0, 1 \rangle \langle 1, 0 \rangle$ with corresponding reachability probability $0.4 \cdot 0.5$. The reachability probability $\Pr_{\mathcal{M}}(\Diamond^{\leq 3}\{\langle 1, 0 \rangle\}) = 0.42$.

It is helpful to separate the topology and the probabilities. We do this by means of a *parametric MC* (pMC) [22]. A pMC over a fixed set of parameters $\mathbf{p}$ generalises MCs by allowing for a transition function that maps to $\mathbb{Q}[\mathbf{p}]$, i.e., to polynomials over these variables [22]. A pMC and a *valuation* of parameters $\mathbf{u}: \mathbf{p} \rightarrow \mathbb{R}$ describe a MC by replacing $\mathbf{p}$ with $\mathbf{u}$ in the transition function P to obtain $P[\mathbf{u}]$. If $P[\mathbf{u}](s)$ is a distribution for every s , then we call $\mathbf{u}$ a *well-defined* valuation. We can then think about a pMC $\mathcal{M}$ as a generator of a set of MCs $\{\mathcal{M}[\mathbf{u}] \mid \mathbf{u} \text{ well-defined}\}$. Figure 2(b) shows a pMC; any valuation $\mathbf{u}$ with $\mathbf{u}(p), \mathbf{u}(q) \in [0, 1]$ is well-defined. We consider the following associated problem:

Parameter Sampling: Given a pMC $\mathcal{M}$, a finite set of well-defined valuations U , and a horizon h , compute $\Pr_{\mathcal{M}[\mathbf{u}]}(\Diamond^{\leq h}T)$ for each $\mathbf{u} \in U$.

We recap binary *decision diagrams* (BDDs) and their generalization into algebraic decision diagrams (ADDs, a.k.a. multi-terminal BDDs). ADDs over a set of variables X are directed acyclic graphs whose vertices V can be partitioned into *terminal nodes* V_t without successors and *inner nodes* V_i with two successors. Each terminal node is labeled with a polynomial over some parameters $\mathbf{p}$ (or just to constants in $\mathbb{Q}$), $\text{val}: V_t \rightarrow \mathbb{Q}[\mathbf{p}]$, and each inner node V_i with a variable, $\text{var}: V_i \rightarrow X$. One node is the root node v_0 . Edges are described by the two successor functions $E_0: V_i \rightarrow V$ and $E_1: V_i \rightarrow V$. A BDD is an ADD with exactly two terminals labeled T and F . Formally, we denote an ADD by the tuple $\langle V, v_0, X, \text{var}, \text{val}, E_0, E_1 \rangle$. ADDs describe functions $f: \mathbb{B}^X \rightarrow \mathbb{Q}[\mathbf{p}]$ (described by a path in the underlying graph and the label of the corresponding terminal node). As finite sets can be encoded with bit vectors, ADDs represent functions from (tuples of) finite sets to polynomials.

Example 2. The transition matrix P of the MC in Fig. 2(a) maps states, encoded by bit vectors, $\langle x, y \rangle, \langle x', y' \rangle$ to the probabilities to move from state $\langle x, y \rangle$ to $\langle x', y' \rangle$. Figure 2(c) shows the corresponding ADD.²

3 A Model Checking Perspective

We briefly analyze the de-facto standard approach to symbolic probabilistic model checking of finite-horizon reachability probabilities. It is an adaptation of qualitative model checking, in which we track the (backward) reachable states. This set can be thought of as a mapping from states to a Boolean indicating whether a target state can be reached. We generalize the mapping to a function that maps every state s to the probability that we reach T within i steps,

² The ADD also depends on the variable order, which we assume fixed for conciseness.

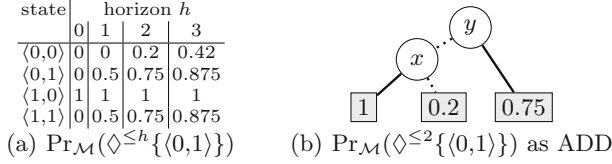


Fig. 3. Bounded reachability and symbolic model checking for the MC $\mathcal{M}$ in Fig. 2(a).

denoted $\Pr_{\mathcal{M}}(s \models \Diamond^{\leq i}T)$. First, it is convenient to construct a transition relation in which the target states have been made absorbing, i.e., we define a matrix with $A(s, s') = P(s, s')$ if $s \notin T$ and $A(s, s') = [s = s']$ ³ otherwise. The following *Bellman equations* characterize that aforementioned mapping:

$$\begin{aligned} \Pr_{\mathcal{M}}(s \models \Diamond^{\leq 0}T) &= [s \in T], \\ \Pr_{\mathcal{M}}(s \models \Diamond^{\leq i}T) &= \sum_{s' \in \text{Succ}(s)} A(s, s') \cdot \Pr_{\mathcal{M}}(s' \models \Diamond^{\leq i-1}T) \quad \text{with } i > 0. \end{aligned}$$

The main aspect model checkers take from these equations is that to compute the h -step reachability from state s , one only needs to combine the $h-1$ -step reachability from any state s' and the transition probabilities $P(s, s')$. We define a vector $\mathbf{T}$ with $\mathbf{T}(s) = [s \in T]$. The algorithm then iteratively computes and stores the i step reachability for $i = 0$ to $i = h$, e.g. by computing $A^3 \cdot \mathbf{T}$ using $A \cdot (A \cdot (A \cdot \mathbf{T}))$. This reasoning is thus *inherently backwards* and *implicitly marginalizing out paths*. In particular, rather than storing the i -step paths that lead to the target, one only stores a vector $\mathbf{x} = A^i \cdot \mathbf{T}$ that stores for every state s the sum over all i -long paths from s .

Explicit representations of matrix A and vector $\mathbf{x}$ require memory at least in the order $|S|$.⁴ To overcome this limitation, *symbolic* probabilistic model checking stores both A and $A^i \cdot \mathbf{T}$ as an ADD by considering the matrix as a function from a tuple $\langle s, s' \rangle$ to $A(s, s')$, and $\mathbf{x}$ as a function from s to $\mathbf{x}(s)$ [2].

Example 3. Reconsider the MC in Fig. 2(a). The h -bounded reachability probability $\Pr_{\mathcal{M}}(\Diamond^{\leq h}\{\langle 1,0 \rangle\})$ can be computed as reflected in Fig. 3(a). The ADD for P is shown in Fig. 2(c). The ADD for $\mathbf{x}$ when $h = 2$ is shown in Fig. 3(b).

The performance of symbolic probabilistic model checking is directly governed by the sizes of these two ADDs. The size of an ADD is bounded from below by the number of leafs. In qualitative model checking, both ADDs are in fact BDDs, with two leafs. However, for the ADD representing A , this lower bound is given by the number of different probabilities in the transition matrix. In the running example, we have seen that a small program $\mathcal{P}$ may have an underlying MC $\llbracket \mathcal{P} \rrbracket$ with an exponential state space S and equally many different transition probabilities. Symbolic probabilistic model checking also scales

³ Where $[x]=1$ if x holds and 0 otherwise.

⁴ Excluding e.g., partial exploration or sampling which typically are not exact.

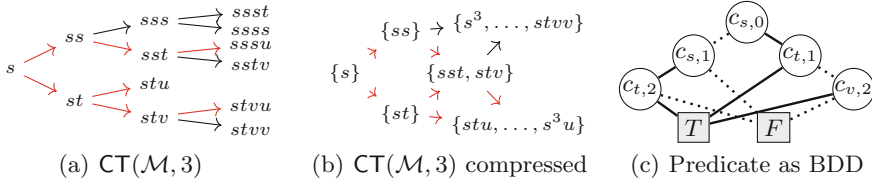


Fig. 4. The computation tree for $\mathcal{M}$ and horizon 3 and its compression. We label states as $s=\langle 0,0\rangle$, $t=\langle 0,1\rangle$, $u=\langle 1,0\rangle$, $v=\langle 1,1\rangle$. Probabilities are omitted for conciseness.

badly on some models where A has a concise encoding but $\mathbf{x}$ has too many different entries.⁵ Therefore, model checkers may store $\mathbf{x}$ partially explicit [49].

The insights above are not new. Symbolic probabilistic model checking has advanced [46] to create small representations of both A and $\mathbf{x}$. In competitions, STORM often applies a bisimulation-to-explicit method that extracts an explicit representation of the bisimulation quotient [26, 36]. Finally, game-based abstraction [32, 44] can be seen as a predicate abstraction technique on the ADD level. However, these methods do not change the computation of the finite horizon reachability probabilities and thus do not overcome the inherent weaknesses of the iterative approach in combination with an ADD-based representation.

4 A Probabilistic Inference Perspective

We present four key insights into probabilistic inference. **(1)** Sect. 4.1 shows how probabilistic inference takes the classical definition as summing over the set of paths, and turns this definition into an algorithm. In particular, these paths may be stored in a computation tree. **(2)** Sect. 4.2 gives the traditional reduction from probabilistic inference to the classical weighted model counting (WMC) problem [16, 57]. **(3)** Sect. 4.3 connects this reduction to point (1) by showing that a BDD that represents this WMC is *bisimilar* to the computation tree assuming that the out-degree of every state in the MC is two. **(4)** Sect. 4.4 describes and compares the computational benefits of the BDD representation. In particular, we clarify that enforcing an out-degree of two is a key ingredient to overcoming one of the weaknesses of symbolic probabilistic model checking: the number of different probabilities in the underlying MC.

4.1 Operational Perspective

The following perspective frames (an aspect of) probabilistic inference as a model transformation. By definition, the set of all paths – each annotated with the transition probabilities – suffices to extract the reachability probability. These sets of paths may be represented in the computation tree (which is itself an MC).

⁵ For an interesting example of this, see the “Queue” example in Sect. 6.

Example 4. We continue from Example 1. We put all paths of length three in a computation tree in Fig. 4(a) (cf. the caption for state identifiers). The three paths that reach the target are highlighted in red. The MC is highly redundant. We may compress to the MC in Fig. 4(b).

Definition 1. For MC $\mathcal{M}$ and horizon h , the computation tree (CT) $\text{CT}(\mathcal{M}, h) = \langle \text{Paths}_h, \iota, P', T' \rangle$ is an MC with states corresponding to paths in $\mathcal{M}$, i.e., $\text{Paths}_h^\mathcal{M}$, initial state ι , target states $T' = \llbracket \Diamond^{\leq h} T \rrbracket$, and transition relation

$$P'(\pi, \pi') = \begin{cases} P(\pi_\downarrow, s) & \text{if } \pi_\downarrow \notin T \wedge \pi' = \pi.s, \\ [\pi_\downarrow \in T \wedge \pi' = \pi] & \text{otherwise.} \end{cases} \quad (1)$$

The CT contains (up to renaming) the same paths to the target as the original MC. Notice that after h transitions, all paths are in a sink state, and thus we can drop the step bound from the property and consider either finite or indefinite horizons. The latter considers all paths that eventually reach the target. We denote the probability mass of these paths with $\Pr_{\mathcal{M}}(s \models \Diamond T)$ and refer to [7] for formal details.⁶ Then, we may compute bounded reachability probabilities in the original MC by analysing unbounded reachability in the CT:

$$\Pr_{\mathcal{M}}(\Diamond^{\leq h} T) = \Pr_{\text{CT}(\mathcal{M}, h)}(\Diamond^{\leq h} T') = \Pr_{\text{CT}(\mathcal{M}, h)}(\Diamond T').$$

The nodes in the CT have a natural topological ordering. The unbounded reachability probability is then computed (efficiently in CT's size) using dynamic programming (i.e., topological value iteration) on the Bellman equation for $s \notin T$:

$$\Pr_{\mathcal{M}}(s \models \Diamond T) = \sum_{s' \in \text{Succ}(s)} P(s, s') \cdot \Pr_{\mathcal{M}}(s' \models \Diamond T).$$

For pMCs, the right-hand side naturally is a factorised form of the *solution function* f that maps parameter values to the induced reachability probability, i.e. $f(\mathbf{u}) = \Pr_{\mathcal{M}[\mathbf{u}]}(\Diamond^{\leq h} T)$ [22, 24, 33]. For bounded reachability (or acyclic pMCs), this function amounts to a sum over all paths with every path reflected by a term of a polynomial, i.e., the sum is a polynomial. In sum-of-terms representation, the polynomial can be exponential in the number of parameters [5].

For computational efficiency, we need a smaller representation of the CT. As we only consider reachability of T , we may simplify [43] the notion of (weak) bisimulation [6] (in the formulation of [40]) to the following definition.

Definition 2. For $\mathcal{M}$ with states S , a relation $\mathcal{R} \subseteq S \times S$ is a (weak) bisimulation (with respect to T) if $s\mathcal{R}s'$ implies $\Pr_{\mathcal{M}}(s \models \Diamond T) = \Pr_{\mathcal{M}}(s' \models \Diamond T)$. Two states s, s' are (weakly) bisimilar (with respect to T) if $\Pr_{\mathcal{M}}(s \models \Diamond T) = \Pr_{\mathcal{M}}(s' \models \Diamond T)$.

Two MCs $\mathcal{M}, \mathcal{M}'$ are bisimilar, denoted $\mathcal{M} \sim \mathcal{M}'$ if the initial states are bisimilar in the disjoint union of the MCs. It holds by definition that if $\mathcal{M} \sim \mathcal{M}'$, then $\Pr_{\mathcal{M}}(\Diamond T) = \Pr_{\mathcal{M}'}(\Diamond T')$. The notion of bisimulation can be lifted to pMCs [33].

⁶ Alternatively, on acyclic models, a large step bound $h > |S|$ suffices.

Idea 1: Given a symbolic description $\mathcal{P}$ of a MC $\llbracket \mathcal{P} \rrbracket$, efficiently construct a concise MC $\mathcal{M}$ that is bisimilar to $\text{CT}(\llbracket \mathcal{P} \rrbracket, h)$.

Indeed, the (compressed) CT in Fig. 4(b) and Fig. 4(a) are bisimilar. We remark that we do not necessarily compute the bisimulation quotient of $\text{CT}(\llbracket \mathcal{P} \rrbracket, h)$.

4.2 Logical Perspective

The previous section defined weakly bisimilar chains and showed computational advantages, but did not present an algorithm. In this section we frame the finite horizon reachability probability as a logical query known as *weighted model counting* (WMC). In the next section we will show how this logical perspective yields an algorithm for constructing bisimilar MCs.

Weighted model counting is well-known as an effective reduction for probabilistic inference [16, 57]. Let φ be a logical sentence over variables C . The *weight function* $W_C: C \rightarrow \mathbb{R}_{\geq 0}$ assigns a weight to each logical variable. A *total variable assignment* $\eta: C \rightarrow \{0, 1\}$ by definition has weight $\text{weight}(\eta) = \prod_{c \in C} W_C(c)\eta(c) + (1 - W_C(c)) \cdot (1 - \eta(c))$. Then the *weighted model count* for φ given W is $\text{WMC}(\varphi, W_C) = \sum_{\eta \models \varphi} \text{weight}(\eta)$. Formally, we desire to compute a reachability query using a WMC query in the following sense:

Idea 2: Given an MC $\mathcal{M}$, efficiently construct a predicate $\varphi_{\mathcal{M},h}^C$ and a weight-function W_C such that $\Pr_{\mathcal{M}}(\Diamond^{\leq h} T) = \text{WMC}(\varphi_{\mathcal{M},h}^C, W_C)$.

Consider initially the simplified case when the MC $\mathcal{M}$ is *binary*: every state has at most two successors. In this case producing $(\varphi_{\mathcal{M},h}^C, W_C)$ is straightforward:

Example 5. Consider the MC in Fig. 2(a), and note that it is binary. We introduce logical variables called *state/step coins* $C = \{c_{s,i} \mid s \in S, i < h\}$ for every state and step. Assignments to these coins denote choices of transitions at particular times: if the chain is in state s at step i , then it takes the transition to the lexicographically first successor of s if $c_{s,i}$ is true and otherwise takes the transition to the lexicographically second successor. To construct the predicate $\varphi_{\mathcal{M},3}^C$, we will need to write a logical sentence on coins whose models encode accepting paths (red paths) in the CT in Fig. 4(a).

We start in state $s = \langle 0, 0 \rangle$ (using state labels from the caption of Fig. 4). We order states as $s = \langle 0, 0 \rangle < t = \langle 0, 1 \rangle < u = \langle 1, 0 \rangle < v = \langle 1, 1 \rangle$. Then, $c_{s,0}$ is true if the chain transitions into state s at time 0 and false if it transitions to state t at time 0. So, one path from s to the target node $\langle 1, 0 \rangle$ is given by the logical sentence $(c_{s,0} \wedge \neg c_{s,1} \wedge c_{t,2})$. The full predicate $\varphi_{\mathcal{M},3}^C$ is therefore:

$$\varphi_{\mathcal{M},3}^C = (c_{s,0} \wedge \neg c_{s,1} \wedge c_{t,2}) \vee (\neg c_{s,0} \wedge c_{t,1}) \vee (\neg c_{s,0} \wedge \neg c_{t,1} \wedge c_{v,2}).$$

Each model of this sentence is a single path to the target. This predicate $\varphi_{\mathcal{M},h}^C$ can clearly be constructed by considering all possible paths through the chain, but later on we will show how to build it more efficiently.

Finally, we fix W_C : The weight for each coin is directly given by the transition probability to the lexicographically first successor: for $0 \leq i < h$, $W_C(c_{s,i}) = 0.6$ and $W_C(c_{t,i}) = W_C(c_{v,i}) = 0.5$. The WMC is indeed 0.42, reflecting Example 1.

When the MC is not binary, it suffices to limit the out-degree of an MC to be at most two by adding auxiliary states, hence binarizing all transitions, cf. [38].

4.3 Connecting the Operational and the Logical Perspective

Now that we have reduced bounded reachability to weighted model counting, we reach a natural question: how do we perform WMC?⁷ Various approaches to performing WMC have been explored; a prominent approach is to compile the logical function into a binary decision diagram (BDD), which supports fast weighted model counting [21]. In this paper, we investigate the use of a BDD-driven approach for two reasons: (i) BDDs admit straightforward support for parametric models. (ii) BDDs provide a direct connection between the logical and operational perspectives. To start, observe that the graph of the BDD, together with the weights, can be interpreted as an MC:

Definition 3. Let φ^X be a propositional formula over variables X and $<_X$ an ordering on X . Let $\text{BDD}(\varphi^X, <_X) = \langle V, v_0, X, \text{var}, \text{val}, E_0, E_1 \rangle$ be the corresponding BDD, and let W be a weight function on X with $0 \leq W(x) \leq 1$. We define the MC $\text{BDD}_{\text{MC}}(\varphi^X, <_X, W) = \langle S, \iota, P, T \rangle$ with $S = V$, $\iota = v_0$, $P(s) = \{E_0(s) \mapsto W(\text{var}(s)), E_1(s) \mapsto 1 - W(\text{var}(s))\}$ and $T = \{v \in V \mid \text{val}(v) = 1\}$.

These BDDs are intimately related to the computation trees discussed before. For a binary MC $\mathcal{M}$, the tree $\text{CT}(\mathcal{M}, h)$ is binary and can be considered as a (not necessarily reduced) BDD. More formally, let us construct $\text{BDD}_{\text{MC}}(\varphi_{\mathcal{M},h}^C, <_C)$. We fix a total order on states. Then we fix *state/step coins* $C = \{c_{s,i} \mid s \in S, i < h\}$ and the weights as in Example 5. Finally, let $<_C$ be an order on C such that $i < j$ implies $c_{s,i} <_C c_{s,j}$. Then:

$$\text{CT}(\mathcal{M}, h) \sim \text{BDD}_{\text{MC}}(\varphi_{\mathcal{M},h}^C, <_C, W). \quad (2)$$

In the spirit of Idea 1, we thus aim to construct $\text{BDD}_{\text{MC}}(\varphi_{\mathcal{M},h}^C, <_C, W)$, a representation as outlined in Idea 2, efficiently. Indeed, the BDD (as MC) in Fig. 4(c) is bisimilar to the MC in Fig. 4(b).

Idea 3: Represent a bisimilar version of the computation tree using a BDD.

⁷ In this paper, we concentrate on reductions to *exact* WMC, leaving approximate approaches for future work [14].

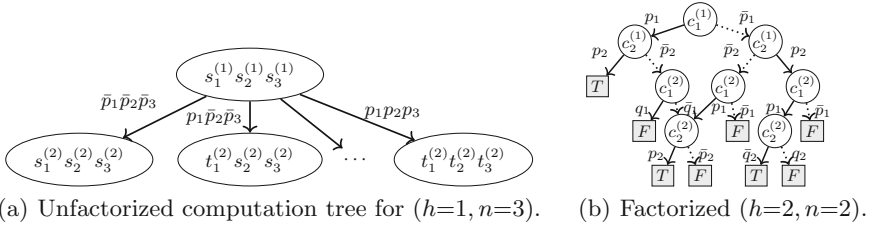


Fig. 5. Two computation trees for the motivating example in Sect. 1.

4.4 The Algorithmic Benefits of BDD Construction

Thus far we have described how to construct a binarized MC bisimilar to the CT. Here, we argue that this construction has algorithmic benefits by filling in two details. First, the binarized representation is an important ingredient for compact BDDs. Second, we show how to choose a variable ordering that ensures that the BDDs grow linearly in the horizon. In sum,

Idea 4: WMC encodings of binarized Markov Chains may increase compression of computation trees.

To see the benefits of binarized transitions, we return to the factory example in Sect. 1. Figure 5(a) gives a bisimilar computation tree for the 3-factory $h = 1$ example. However, in this tree, the states are *unfactorized*: each node in the tree is a joint configuration of factories. This tree has 8 transitions (one for each possible joint state transition) with 8 distinct probabilities. On the other hand, the bisimilar computation tree in Fig. 1(d) has binarized transitions: each node corresponds to a single factory’s state at a particular time-step, and each transition describes an update to only a single factory. This binarization enables the exploitation of new structure: in this case, the independence of the factories leads to smaller BDDs, that is otherwise lost when considering only joint configurations of factories.

Recall that the size of the ADD representation of the transition function is bounded from below by the number of distinct probabilities in the underlying MC: in this case, this is visualized by the number of distinct outgoing edge probabilities from all nodes in the unfactorized computation tree. Thus, a good binarization can have a drastically positive effect on performance. For the running example, rather than 2^n different transition probabilities (with n factories), the system now has only $4 \cdot n$ distinct transition probabilities!

Causal Orderings. Next, we explore some of the *engineering choices* RUBICON makes to exploit the sequential structure in a MC when constructing the BDD for a WMC query. First, note that the transition matrix $P(s, s')$ implicitly encodes a distribution over state transition functions, $S \rightarrow S$. To encode P as a BDD, we must encode each transition as a logical variable, similar to the situation in Sect. 4.2. In the case of binary transitions this is again easy. In the case of non-binary transitions, we again introduce additional logical variables [16, 27, 39, 57].

This logical function has the following form:

$$f_P: \{0, 1\}^C \rightarrow (S \rightarrow S). \quad (3)$$

Whereas the computation tree follows a fixed (temporal) order of states, BDDs can represent the same function (and the same weighted model count) using an arbitrary order. Note that the BDD's size and structure drastically depends both on the construction of the propositional formula *and* the order of the variables in that encoding. We can bound the size of the BDD by enforcing a variable order based on the temporal structure of the original MC. Specifically, given h coin collections $\mathbf{C} = C \times \dots \times C$, one can generate a function f describing the h -length paths via repeated applications of f_P :

$$f: \{0, 1\}^C \rightarrow \text{Paths}_h \quad f(C_1, \dots, C_h) = \left(f_P(C_h) \circ \dots \circ f_P(C_1) \right)(\iota) \quad (4)$$

Let ψ denote an indicator for the reachability property as a function over paths, $\psi: \text{Paths}_h \rightarrow \{0, 1\}$ with $\psi(\pi) = [\pi \in \llbracket \Diamond^{\leq h} T \rrbracket]$. We call predicates formed by composition with f_P , i.e., $\varphi = \psi \circ f_P$, *causal encodings* and orderings on $c_{i,t} \in \mathbf{C}$ that are lexicographically sorted in time, $t_1 < t_2 \implies c_{i,t_1} < c_{j,t_2}$, *causal orderings*. Importantly, causally ordered / encoded BDDs grow linearly in horizon h , [61, Corollary 1]. More precisely, let $\varphi_{\mathcal{M},h}^C$ be causally encoded where $|\mathbf{C}| = h \cdot m$. The causally ordered BDD for $\varphi_{\mathcal{M},h}^C$ has at most $h \cdot |S \times S_\psi| \cdot m \cdot 2^m$ nodes, where $|S_\psi| = 2$ for reachability properties.⁸ However, while the worst-case growth is linear in the horizon, constructing that BDD may induce a super-linear cost in the size, e.g., function composition using BDDs is super-linear!

Figure 5(b) shows the motivating factory example with 2 factories and $h = 2$. The variables are causally ordered: the factories in time step 1 occur before the factories in time step 2. For n factories, a fixed number $f(n)$ of nodes are added to the BDD upon each iteration, guaranteeing growth on the order $\mathcal{O}(f(n) \cdot h)$. Note the factorization that occurs: the BDD has node sharing (node $c_2^{(2)}$ is reused) that yields additional computational benefits.

Summary and Remaining Steps. The operational view highlights that we want to compute a transformation of the original input MC $\mathcal{M}$. The logical view presents an approach to do so efficiently: By computing a BDD that stores a predicate describing all paths that reach the target, and interpreting and evaluating the (graph of the) BDD as an MC. In the following section, we discuss the two steps that we follow to create the BDD: (i) From $\mathcal{P}$ generate $\mathcal{P}'$ such that $\text{CT}(\llbracket \mathcal{P} \rrbracket, h) \sim \llbracket \mathcal{P}' \rrbracket$. (ii) From $\mathcal{P}'$ generate $\mathcal{M}$ such that $\mathcal{M} = \llbracket \mathcal{P}' \rrbracket$.

5 RUBICON

We present RUBICON which follows the two steps outlined above. For exposition, we first describe a translation of *monolithic* PRISM programs to *Dice* programs

⁸ Generally, it is the smallest number of states required for a DFA to recognize ψ .

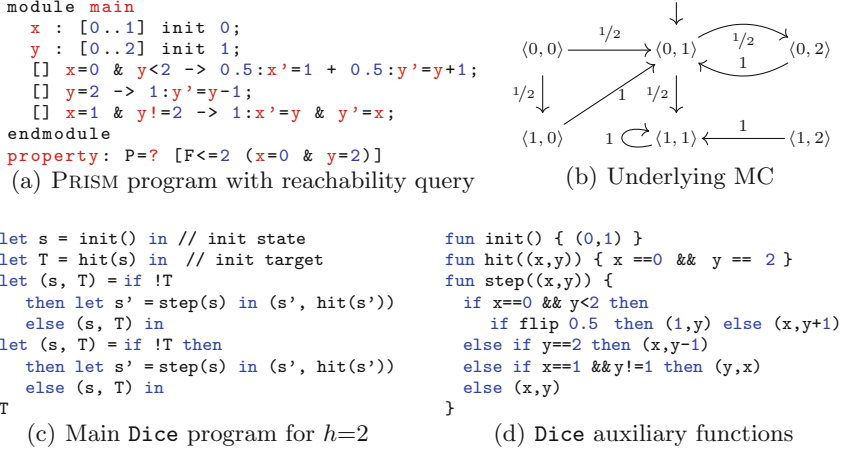


Fig. 6. From PRISM to Dice using RUBICON.

and then extend this translation to admit modular programs. Technical steps and extensions are deferred to [38, Appendix].

Dice Preliminaries. We give a brief description of *Dice*, a probabilistic programming language (PPL) introduced in [39]. A PPL is a programming language augmented with a primitive notion of random choice: for instance, in *Dice*, a Bernoulli random variable is introduced by the syntax `flip 0.5`. The syntax of *Dice* is similar to the programming language *OCaml*: local variables are introduced by the syntax `let x = e1 in e2`, where *e₁* and *e₂* are *expressions*, i.e., sub-programs. *Dice* supports procedures, bounded integers, bounded loops, and standard control flow via *if*-statements.

One goal of a PPL is to perform *probabilistic inference*: compute the probability that the program returns a particular value. Inference on the tiny *Dice* program `let x = flip 0.1 in x` would yield that `true` is returned with probability 0.1. The *Dice* compiler performs probabilistic inference via weighted model counting and BDD compilation. In doing so, it accomplishes the *non-trivial* tasks of: (i) choosing a logical encoding for probabilistic programs (ii) establishing good variable orderings (iii) efficiently manipulating and constructing BDDs (iv) performing WMC. For details, we refer the reader to [39].

RUBICON uses *Dice* to effectively construct a BDD and perform WMC on a *Dice* program that reflects a description of some computation tree. This implementation exploits the structure that was described in Sect. 4.4: in particular, the BDD generated in Fig. 5(b) is exactly the BDD that will be generated by *Dice* from the output of RUBICON. The variable ordering used by *Dice* is given by the order in which program variables are introduced, and RUBICON’s translation was designed with this variable ordering in mind.

Transpiling PRISM to Dice. We present the core translation routine implemented in RUBICON. We note that the ultimate performance of RUBICON is

heavily dependent on the quality of this translation. We evaluate the performance in the next section.

The PRISM specification language consists of one or more reactive *modules* (or partially synchronized state machines) that may interact with each other. Our example in Fig. 1(b) illustrates fully synchronized state machines. While PRISM programs containing multiple modules can be flattened into a single monolithic program, this yields an exponential blow-up: If one flattens the n modules in Fig. 1(b) to a single module, the resulting program has 2^n updates per command. This motivates our direct translation of PRISM programs containing multiple modules.

Monolithic Prism Programs. We explain most ideas on PRISM programs that consist of a single “monolithic” module before we address the modular translation at the end of the subsection. A module has a set of bounded variables, and the valuations of these variables span the state space of the underlying MC. Its transitions are described by guarded *commands* of the form:

$$[\mathbf{act}] \text{ guard} \rightarrow p_1 : \mathbf{update}_1 + \dots + p_n : \mathbf{update}_n$$

The *action* name **act** is only relevant in the modular case and can be ignored for now. The *guard* is a Boolean expression over the module’s variables. If the guard evaluates to **true** for some state (a valuation), then the module evolves into one of the n successor states by updating its variables. An *update* is chosen according to the probability distribution given by the expressions $p_1, \dots, p_n$. In every state enabling the guard, the evaluation of $p_1, \dots, p_n$ must sum up to one. A set of guards *overlap* if they all evaluate to **true** on a given state. The semantics of overlapping guards in the monolithic setting is to first uniformly select an active guard and then apply the corresponding stochastic transition. Finally, a self-loop is implicitly added to states without an enabled guard.

Example 6. We present our translation primarily through example. In Fig. 6(a), we give a PRISM program for a MC. The program contains two variables x and y , where x is either zero or one, and y between zero and two. There are thus 6 different states. We denote states as tuples with the x - and y -value. We depict the MC in Fig. 6(b). From state $\langle 0, 0 \rangle$, (only) the first guard is enabled and thus there are two transitions, each with probability a half: one in which x becomes one and one in which y is increased by one. Finally, there is no guard enabled in state $\langle 1, 1 \rangle$, resulting in an implicit self-loop.

Translation. All Dice programs consist of two parts: a *main* routine, which is run by default when the program starts, and *function declarations* that declare auxiliary functions. We first define the auxiliary functions. For simplicity let us temporarily assume that no guards overlap and that probabilities are constants, i.e., not state-dependent.

The main idea in the translation is to construct a Dice function **step** that, given the current state, outputs the next state. Because a monolithic PRISM

<pre> module main x : [0..2] init 1; y : [0..2] init 1; [] x>1 -> 1:x'=y&y'=x; [] y<2 -> 1:x'=min(x+1,2); endmodule </pre>	<pre> fun step((x,y)) { let aEn =(x>1) let bEn =(y<2) let act = selectFrom(aEn, bEn) if act==1 then (y,x) else if act==2 then (min(x+1,2),y) else (x,y)} ... </pre>
(a)	(b)

Fig. 7. PRISM program with overlapping guards and its translation (conceptually).

<pre> module m1 x : [0..1] init 0; [a] x=1 -> 1:x'=1-y; [b] x=0 -> 1:x'=0; endmodule module m2 y : [0..1] init 0; [b] y=1 -> 0.5:y'=0 +0.5:y'=1; [c] true -> 1:x'=1-x; endmodule </pre>	<pre> fun step((x,y)) { let aEn =(x==1) in let bEn =(x=0 &&y=1) in let cEn =true in let act =selectFrom(aEn, bEn, cEn) in if act==1 then (1-y, y) else if act==2 then (0, flip 0.5) else if act==3 then (1-x, y) else (x, y) } </pre>
(a)	(b)

Fig. 8. Modular PRISM and resulting Dice step function.

program is almost a sequential program, in its most basic version, the **step** function is straightforward to construct using built-in Dice language primitives: we simply build a large if-else block corresponding to each command. This block iteratively considers each command's guard until it finds one that is satisfied. To perform the corresponding update we flip a coin – based on the probabilities corresponding to the updates – to determine which update to perform. If no command is enabled, we return the same state in accordance with the implicit self-loop. Figure 6(d) shows the program blocks for the PRISM program from Fig. 6(a) with target state $\llbracket x = 0, y = 2 \rrbracket$. There are two other important auxiliary functions. The **init** function simply returns the initial state by translating the initialization statements from PRISM, and the **hit** function checks whether the current state is a target state that is obtained from the property.

Now we outline the main routine, given for this example in Fig. 6(c). This function first initializes the state. Then, it calls **step** 2 times, checking on each iteration using **hit** if the target state is reached. Finally, we return whether we have been in a target state. The probability to return true corresponds to the reachability probability on the underlying MC specified by the PRISM program.

Overlapping Guards. PRISM allows multiple commands to be enabled in the same state, with semantics to uniformly at random choose one of the enabled commands to evaluate. Dice has no primitive notion of this construct.⁹ We illustrate the translation in Fig. 7(a) and Fig. 7(b). It determines which guards **aEn**, **bEn**, **cEn** are enabled. Then, we randomly select one of the commands which are enabled, i.e., we uniformly at random select a true bit from a given tuple

⁹ One cannot simply condition on selecting an enabled guard as this redistributes probability mass over all paths and not only over paths with the same prefix.

of bits. We store the index of that bit and use it to execute the corresponding command.

Modular Prism Programs. For modular PRISM programs, the *action names* at the front of PRISM commands are important. In each module, there is a set of action names available. An action is *enabled* if each module that contains this action name has (at least) one command with this action whose guard is satisfied. Commands with an empty action are assumed to have a globally unique action name, so in that case the action is enabled iff the guard is enabled. Intuitively, once an action is selected, we randomly select a command per module in all modules containing this action name. Our approach resembles that for overlapping guards described above. See Fig. 8 for an intuitive example. To automate this, the updates require more care, cf. [38] for details.

Implementation. RUBICON is implemented on top of STORM’s Python API and translates PRISM to Dice fully automatically. RUBICON supports all MCs in the PRISM benchmark suite and a large set of benchmarks from the PRISM website and the QVBS [35], with the note that we require a single initial state and ignore reward declarations. Furthermore, we currently do not support the hide/restrict process-algebraic compositions and some integer operations.

6 Empirical Comparisons

We compare and contrast the performance of STORM against RUBICON to empirically demonstrate the following strengths and weaknesses:¹⁰

Explicit Model Checking (STORM) represents the MC explicitly in a sparse matrix format. The approach suffers from the state space explosion, but has been engineered to scale to models with many states. Besides the state space, the sparseness of the transition matrix is essential for performance.

Symbolic Model Checking (STORM) represents the transition matrix and the reachability probability as an ADD. This method is strongest when the transition matrix and state vector have structure that enables a small ADD representation, like symmetry and sparsity.

RUBICON represents the set of paths through the MC as a (logical) BDD. This method excels when the state space has structure that enables a compact BDD representation, such as conditional independence, and hence scales well on examples with many (asymmetric) parallel processes or queries that admit a compact representation.

The sources, benchmarks and binaries are archived.¹¹

There is no clear-cut model checking technique that is superior to others (see QCOMP [12]). We demonstrate that, while RUBICON is not competitive on some

¹⁰ All experiments were conducted with STORM version 1.6.0 on the same server with 512 GB of RAM, using a single thread of execution. Time was reported using the built-in Unix `time` utility; the total wall-clock time is reported.

¹¹ <http://doi.org/10.5281/zenodo.4726264> and <http://github.com/sjunges/rubicon>.

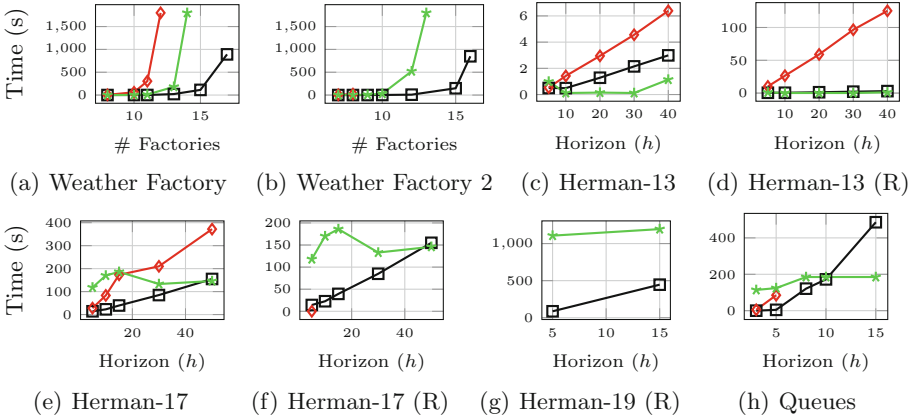


Fig. 9. Scaling plots comparing RUBICON (—■—), STORM’s symbolic engine (—◆—), and STORM’s explicit engine (—★—). An “(R)” in the caption denotes random parameters.

commonly used benchmarks [52], it improves a modern model checking portfolio approach on a significant set of benchmarks. Below we provide several natural models on which RUBICON is superior to one or both competing methods. We also evaluated RUBICON on standard benchmarks, highlighting that RUBICON is applicable to models from the literature. We see that RUBICON is effective on HERMAN (elaborated below), has mixed results on BRP [38, Appendix], and is currently not competitive on some other standard benchmarks (NAND, EGL, LeaderSync). While not exhaustive, our selected benchmarks highlight specific strengths and weaknesses of RUBICON. Finally, a particular benefit of RUBICON is fast sampling of parametric chains, which we demonstrate on HERMAN and our factory example.

Scaling Experiments. In this section, we describe several scaling experiments (Fig. 9), each designed to highlight a specific strength or weakness.

Weather Factories. First, Fig. 9(a) describes a generalization of the motivating example from Sect. 1. In this model, the probability that each factory is on strike is dependent on a common random event: whether or not it is raining. The rain on each day is dependent on the previous day’s weather. We plot runtime for an increasing number of factories for $h=10$. Both STORM engines eventually fail due to the state explosion and the number of distinct probabilities in the MC. RUBICON is orders of magnitude faster in comparison, highlighting that it does not depend on complete independence among the factories. Figure 9(b) shows a more challenging instance where the weather includes *wind* which, each day, affects whether or not the sun will shine, which in turn affects strike probability.

Herman. Herman is based on a distributed protocol [37] that has been well-studied [1, 53] and which is one of the standard benchmarks in probabilistic model checking. Rather than computing the expected steps to ‘stabilization’, we consider the step-bounded probability of stabilization. Usually, all participants in

the protocol flip a coin with the same bias. The model is then highly symmetric, and hence is amenable to symbolic representation with ADDs. Figures 9(c) and 9(e) show how the methods scale on Herman examples with 13 and 17 parallel processes. We observe that the explicit approach scales very efficiently in the number of iterations but has a much higher up-front model-construction cost, and hence can be slower for fewer iterations.

To study what happens when the coin biases vary over the protocol participants, we made a version of the Herman protocol where each participant's bias is randomly chosen, which ruins the symmetry and so causes the ADD-based approaches to scale significantly worse (Figs. 9(d) and 9(f), and 9(g)); we see that symbolic ADD-based approaches completely fail on Herman 17 and Herman 19 (the curve terminating denotes a memory error). RUBICON and the explicit approach are unaffected by varying parameters.

Queues. The Queues model has K queues of capacity Q where every step, tasks arrive with a particular probability. Three queues are of type 1, the others of type 2. We ask the probability that all queues of type 1 and at least one queue of type 2 is full within k steps. Contrary to the previous models, the ADD representation of the transition matrix is small. Figure 9(h) shows the relative scaling on this model with $K = 8$ and $Q = 3$. We observe that ADDs quickly fail due to inability to concisely represent the probability vector $\mathbf{x}$ from Sect. 3. RUBICON outperforms explicit model checking until $h = 10$.

Sampling Parametric Markov Chains. We evaluate performance for the pMC sampling problem outlined in Sect. 2. Table 1 gives for four models the time to construct the BDD and to perform WMC, as well as the time to construct an ADD in STORM and to perform model checking with this ADD. Finally, we show the time for STORM to compute the solution function of the pMC (with the explicit representation). The pMC sampling in STORM – symbolic and explicit – computes the reachability probabilities with concrete probabilities. RUBICON, in contrast, constructs a ‘parametric’ BDD once, amortizing the cost of repeated efficient evaluation. The ‘parametric BDD’ may be thought of as a solution function, as discussed in Sect. 4.1. STORM cannot compute these solution functions as efficiently. We observe in Table 1 that fast parametric sampling is realized in RUBICON: for instance, after a 40s up-front compilation of the factories example with 15 factories, we have a solution function in factorized form and it costs an order of magnitude less time to draw a sample. Hence, sampling and computation of solution functions of pMCs is a major strength of RUBICON.

7 Discussion, Related Work, and Conclusion

We have demonstrated that the probabilistic inference approach to probabilistic model checking can improve scalability on an important class of problems. Another benefit of the approach is for sampling pMCs. These are used to evaluate e.g., robustness of systems [1], or to synthesise POMDP controllers [41]. Many state-of-the-art approaches [17, 19, 24] require the evaluation of various instantiated MCs, and RUBICON is well-suited to this setting. More generally, support

Table 1. Sampling performance comparison and pMC model checking, time in seconds.

Model	RUBICON		STORM (w/ ADD)		STORM (explicit)
	Build	WMC	Build	Solve	pMC solving
Herman R 13 ($h = 10$)	3	<1	32	18	>1800
Herman R 17 ($h = 10$)	45	28	>1800	–	>1800
Factories 12 ($h = 15$)	2	<1	59	286	>1800
Factories 15 ($h = 15$)	40	4	>1800	–	>1800

of inference techniques opens the door to a variety of algorithms for additional queries, e.g., computing *conditional probabilities* [3,8].

An important limitation of probabilistic inference is that only finitely many paths can be stored. For infinite horizon properties in cyclic models, an infinite set of arbitrarily long paths would be required. However, as standard in probabilistic model checking, we may soundly approximate infinite horizons. Additionally, the inference algorithm in *Dice* does not support a notion of non-determinism. It thus can only be used to evaluate MCs, not Markov decision processes. However, [61] illustrates that this is not a conceptual limitation. Finally, we remark that RUBICON achieves its performance with a straightforward translation. We are optimistic that this is a first step towards supporting a larger class of models by improving the transpilation process for specific problems.

Related Work. The tight connection with inference has been recently investigated via the use of model checking for Bayesian networks, the prime model in probabilistic inference [56]. Bayesian networks can be described as probabilistic programs [10] and their operational semantics coincides with MCs [31]. Our work complements these insights by studying how symbolic model checking can be sped up by probabilistic inference.

The path-based perspective is tightly connected to *factored state spaces*. Factored state spaces are often represented as (bipartite) Dynamic Bayesian networks. ADD-based model checking for DBNs has been investigated in [25], with mixed results. Their investigation focuses on using ADDs for factored state space representations. We investigate using BDDs representing paths. Other approaches also investigated a path-based view: The symbolic encoding in [28] annotates propositional sub-formulae with probabilities, an idea closer to ours. The underlying process implicitly constructs an (uncompressed) CT leading to an exponential blow-up. Likewise, an explicit construction of a computation tree without factorization is considered in [62]. Compression by grouping paths has been investigated in two *approximate* approaches: [55] discretises probabilities and encodes into a satisfiability problem with quantifiers and bit-vectors. This idea has been extended [60] to a PAC algorithm by purely propositional encodings and (approximate) model counting [14]. Finally, factorisation exploits symmetries, which can be exploited using symmetry reduction [50]. We highlight that the latter is not applicable to the example in Fig. 1(d).

There are many techniques for exact probabilistic inference in various forms of probabilistic modeling, including probabilistic graphical models [20, 54]. The semantics of graphical models make it difficult to transpile PRISM programs, since commonly used operations are lacking. Recently, *probabilistic programming languages* have been developed which are more amenable to transpilation [13, 23, 29, 30, 59]. We target *Dice* due to the technical development that it enables in Sect. 4, which enabled us to design and explain our experiments. Closest related to *Dice* is ProbLog [27], which is also a PPL that performs inference via WMC; ProbLog has different semantics from *Dice* that make the translation less straightforward. The paper [61] uses an encoding similar to *Dice* for inferring specifications based on observed traces. ADDs and variants have been considered for probabilistic inference [15, 18, 58], which is similar to the process commonly used for probabilistic model checking. The planning community has developed their own disjoint sets of methods [45]. Some ideas from learning have been applied in a model checking context [11].

8 Conclusion

We present RUBICON, bringing probabilistic AI to the probabilistic model checking community. Our results show that RUBICON can outperform probabilistic model checkers on some interesting examples, and that this is not a coincidence but rather the result of a significantly different perspective.

References

1. Aflaki, S., Volk, M., Bonakdarpour, B., Katoen, J.P., Storjohann, A.: Automated fine tuning of probabilistic self-stabilizing algorithms. In: SRDS, pp. 94–103. IEEE (2017)
2. de Alfaro, L., Kwiatkowska, M., Norman, G., Parker, D., Segala, R.: Symbolic model checking of probabilistic processes using MTBDDs and the Kronecker representation. In: Graf, S., Schwartzbach, M. (eds.) TACAS 2000. LNCS, vol. 1785, pp. 395–410. Springer, Heidelberg (2000). https://doi.org/10.1007/3-540-46419-0_27
3. Andrés, M.E., van Rossum, P.: Conditional probabilities over probabilistic and nondeterministic systems. In: Ramakrishnan, C.R., Rehof, J. (eds.) TACAS 2008. LNCS, vol. 4963, pp. 157–172. Springer, Heidelberg (2008). https://doi.org/10.1007/978-3-540-78800-3_12
4. Baier, C., de Alfaro, L., Forejt, V., Kwiatkowska, M.: Model checking probabilistic systems. In: Clarke, E., Henzinger, T., Veith, H., Bloem, R. (eds.) Handbook of Model Checking, pp. 963–999. Springer, Cham (2018). https://doi.org/10.1007/978-3-319-10575-8_28
5. Baier, C., Hensel, C., Hutschenreiter, L., Junges, S., Katoen, J.P., Klein, J.: Parametric Markov chains: PCTL complexity and fraction-free Gaussian elimination. Inf. Comput. **272**, 104504 (2020)
6. Baier, C., Hermanns, H.: Weak bisimulation for fully probabilistic processes. In: Grumberg, O. (ed.) CAV 1997. LNCS, vol. 1254, pp. 119–130. Springer, Heidelberg (1997). https://doi.org/10.1007/3-540-63166-6_14

7. Baier, C., Katoen, J.P.: Principles of Model Checking. MIT Press, Cambridge (2008)
8. Baier, C., Klein, J., Klüppelholz, S., Märcker, S.: Computing conditional probabilities in Markovian models efficiently. In: Ábrahám, E., Havelund, K. (eds.) TACAS 2014. LNCS, vol. 8413, pp. 515–530. Springer, Heidelberg (2014). https://doi.org/10.1007/978-3-642-54862-8_43
9. Baluta, T., Shen, S., Shinde, S., Meel, K.S., Saxena, P.: Quantitative verification of neural networks and its security applications. In: CCS, pp. 1249–1264. ACM (2019)
10. Batz, K., Kaminski, B.L., Katoen, J.-P., Matheja, C.: How long, O Bayesian network, will I sample thee? – a program analysis perspective on expected sampling times. In: Ahmed, A. (ed.) ESOP 2018. LNCS, vol. 10801, pp. 186–213. Springer, Cham (2018). https://doi.org/10.1007/978-3-319-89884-1_7
11. Brázdil, T., et al.: Verification of Markov decision processes using learning algorithms. In: Cassez, F., Raskin, J.-F. (eds.) ATVA 2014. LNCS, vol. 8837, pp. 98–114. Springer, Cham (2014). https://doi.org/10.1007/978-3-319-11936-6_8
12. Budde, C.E., et al.: On correctness, precision, and performance in quantitative verification: QComp 2020 competition report. In: ISOLA. LNCS. Springer, Heidelberg (2020)
13. Carpenter, B., et al.: Stan: a probabilistic programming language. J. Stat. Soft. **VV**(Ii) (2016)
14. Chakraborty, S., Fried, D., Meel, K.S., Vardi, M.Y.: From weighted to unweighted model counting. In: IJCAI, pp. 689–695. AAAI Press (2015)
15. Chavira, M., Darwiche, A.: Compiling Bayesian networks using variable elimination. In: IJCAI, pp. 2443–2449 (2007)
16. Chavira, M., Darwiche, A.: On probabilistic inference by weighted model counting. Artif. Intell. **172**(6–7), 772–799 (2008)
17. Chen, T., Hahn, E.M., Han, T., Kwiatkowska, M.Z., Qu, H., Zhang, L.: Model repair for Markov decision processes. In: TASE, pp. 85–92. IEEE (2013)
18. Claret, G., Rajamani, S.K., Nori, A.V., Gordon, A.D., Borgström, J.: Bayesian inference using data flow analysis. In: FSE, pp. 92–102 (2013)
19. Cubuktepe, M., Jansen, N., Junges, S., Katoen, J.P., Topcu, U.: Scenario-based verification of uncertain MDPs. In: TACAS. LNCS, vol. 12078, pp. 287–305. Springer, Heidelberg (2020)
20. Darwiche, A.: SDD: a new canonical representation of propositional knowledge bases. In: IJCAI, pp. 819–826 (2011)
21. Darwiche, A., Marquis, P.: A knowledge compilation map. JAIR **17**, 229–264 (2002)
22. Daws, C.: Symbolic and parametric model checking of discrete-time Markov chains. In: Liu, Z., Araki, K. (eds.) ICTAC 2004. LNCS, vol. 3407, pp. 280–294. Springer, Heidelberg (2005). https://doi.org/10.1007/978-3-540-31862-0_21
23. De Raedt, L., Kimmig, A., Toivonen, H.: ProbLog: a probabilistic prolog and its application in link discovery. In: IJCAI, vol. 7, pp. 2462–2467 (2007)
24. Dehnert, C., et al.: PROPhESY: a PRObabilistic ParamETER SYNthesis tool. In: Kroening, D., Păsăreanu, C.S. (eds.) CAV 2015. LNCS, vol. 9206, pp. 214–231. Springer, Cham (2015). https://doi.org/10.1007/978-3-319-21690-4_13
25. Deininger, D., Dimitrova, R., Majumdar, R.: Symbolic model checking for factored probabilistic models. In: Artho, C., Legay, A., Peled, D. (eds.) ATVA 2016. LNCS, vol. 9938, pp. 444–460. Springer, Cham (2016). https://doi.org/10.1007/978-3-319-46520-3_28
26. van Dijk, T., van de Pol, J.: Multi-core symbolic bisimulation minimisation. STTT **20**(2), 157–177 (2018)

27. Fierens, D., et al.: Inference and learning in probabilistic logic programs using weighted Boolean formulas. *Theory Pract. Log. Prog.* **15**(3), 358–401 (2015)
28. Fränzle, M., Hermanns, H., Teige, T.: Stochastic satisfiability modulo theory: a novel technique for the analysis of probabilistic hybrid systems. In: Egerstedt, M., Mishra, B. (eds.) *HSCC 2008*. LNCS, vol. 4981, pp. 172–186. Springer, Heidelberg (2008). https://doi.org/10.1007/978-3-540-78929-1_13
29. Gehr, T., Misailovic, S., Vechev, M.: PSI: exact symbolic inference for probabilistic programs. In: Chaudhuri, S., Farzan, A. (eds.) *CAV 2016*. LNCS, vol. 9779, pp. 62–83. Springer, Cham (2016). https://doi.org/10.1007/978-3-319-41528-4_4
30. Gordon, A.D., Henzinger, T.A., Nori, A.V., Rajamani, S.K.: Probabilistic programming. In: *FOSE*, pp. 167–181. ACM (2014)
31. Gretz, F., Katoen, J.P., McIver, A.: Operational versus weakest pre-expectation semantics for the probabilistic guarded command language. *Perform. Eval.* **73**, 110–132 (2014)
32. Hahn, E.M., Hermanns, H., Wachter, B., Zhang, L.: PASS: abstraction refinement for infinite probabilistic models. In: Esparza, J., Majumdar, R. (eds.) *TACAS 2010*. LNCS, vol. 6015, pp. 353–357. Springer, Heidelberg (2010). https://doi.org/10.1007/978-3-642-12002-2_30
33. Hahn, E.M., Hermanns, H., Zhang, L.: Probabilistic reachability for parametric Markov models. *STTT* **13**(1), 3–19 (2011)
34. Hartmanns, A., Hermanns, H.: The modest toolset: an integrated environment for quantitative modelling and verification. In: Ábrahám, E., Havelund, K. (eds.) *TACAS 2014*. LNCS, vol. 8413, pp. 593–598. Springer, Heidelberg (2014). https://doi.org/10.1007/978-3-642-54862-8_51
35. Hartmanns, A., Klauck, M., Parker, D., Quatmann, T., Ruijters, E.: The quantitative verification benchmark set. In: *TACAS*. LNCS, vol. 11427, pp. 344–350. Springer, Heidelberg (2019)
36. Hensel, C., Junges, S., Katoen, J.P., Quatmann, T., Volk, M.: The probabilistic model checker storm. *STTT* (2021, to appear)
37. Herman, T.: Probabilistic self-stabilization. *Inf. Process. Lett.* **35**(2), 63–67 (1990)
38. Holtzen, S., Junges, S., Vazquez-Chanlatte, M., Millstein, T., Seshia, S.A., Van den Broeck, G.: Model checking finite-horizon Markov chains with probabilistic inference. *CoRR* abs/2105.12326 (2021)
39. Holtzen, S., Van den Broeck, G., Millstein, T.: Scaling exact inference for discrete probabilistic programs. *PACMPL OOPSLA*, November 2020
40. Jansen, D.N., Groote, J.F., Timmers, F., Yang, P.: A near-linear-time algorithm for weak bisimilarity on Markov chains. In: *CONCUR. LIPIcs*, vol. 171, pp. 8:1–8:20. Schloss Dagstuhl - LZI (2020)
41. Junges, S., et al.: Finite-state controllers of POMDPs using parameter synthesis. In: *UAI*, pp. 519–529. AUAI Press (2018)
42. Katoen, J.-P., Gretz, F., Jansen, N., Kaminski, B.L., Olmedo, F.: Understanding probabilistic programs. In: Meyer, R., Platzer, A., Wehrheim, H. (eds.) *Correct System Design*. LNCS, vol. 9360, pp. 15–32. Springer, Cham (2015). https://doi.org/10.1007/978-3-319-23506-6_4
43. Katoen, J.-P., Kemna, T., Zapreev, I., Jansen, D.N.: Bisimulation minimisation mostly speeds up probabilistic model checking. In: Grumberg, O., Huth, M. (eds.) *TACAS 2007*. LNCS, vol. 4424, pp. 87–101. Springer, Heidelberg (2007). https://doi.org/10.1007/978-3-540-71209-1_9
44. Kattenbelt, M., Kwiatkowska, M.Z., Norman, G., Parker, D.: A game-based abstraction-refinement framework for Markov decision processes. *FMSD* **36**(3), 246–280 (2010)

45. Klauck, M., Steinmetz, M., Hoffmann, J., Hermanns, H.: Bridging the gap between probabilistic model checking and probabilistic planning: survey, compilations, and empirical comparison. *JAIR* **68**, 247–310 (2020)
46. Klein, J., et al.: Advances in probabilistic model checking with PRISM: variable reordering, quantiles and weak deterministic büchi automata. *STTT* **20**(2), 179–194 (2018)
47. Koller, D., Friedman, N.: *Probabilistic Graphical Models: Principles and Techniques*. MIT Press, Cambridge (2009)
48. Kozen, D.: Semantics of probabilistic programs. *JCSS* **22**(3), 328–350 (1981)
49. Kwiatkowska, M., Norman, G., Parker, D.: Probabilistic symbolic model checking with PRISM: a hybrid approach. In: Katoen, J.-P., Stevens, P. (eds.) *TACAS 2002*. LNCS, vol. 2280, pp. 52–66. Springer, Heidelberg (2002). https://doi.org/10.1007/3-540-46002-0_5
50. Kwiatkowska, M., Norman, G., Parker, D.: Symmetry reduction for probabilistic model checking. In: Ball, T., Jones, R.B. (eds.) *CAV 2006*. LNCS, vol. 4144, pp. 234–248. Springer, Heidelberg (2006). https://doi.org/10.1007/11817963_23
51. Kwiatkowska, M., Norman, G., Parker, D.: PRISM 4.0: verification of probabilistic real-time systems. In: Gopalakrishnan, G., Qadeer, S. (eds.) *CAV 2011*. LNCS, vol. 6806, pp. 585–591. Springer, Heidelberg (2011). https://doi.org/10.1007/978-3-642-22110-1_47
52. Kwiatkowska, M.Z., Norman, G., Parker, D.: The PRISM benchmark suite. In: *QEST*, pp. 203–204. IEEE (2012)
53. Kwiatkowska, M.Z., Norman, G., Parker, D.: Probabilistic verification of Herman’s self-stabilisation algorithm. *Formal Aspects Comput.* **24**(4–6), 661–670 (2012)
54. Pearl, J.: *Probabilistic Reasoning in Intelligent Systems: Networks of Plausible Inference*. Morgan Kaufmann (1988)
55. Rabe, M.N., Wintersteiger, C.M., Kugler, H., Yordanov, B., Hamadi, Y.: Symbolic approximation of the bounded reachability probability in large Markov chains. In: Norman, G., Sanders, W. (eds.) *QEST 2014*. LNCS, vol. 8657, pp. 388–403. Springer, Cham (2014). https://doi.org/10.1007/978-3-319-10696-0_30
56. Salmani, B., Katoen, J.-P.: Bayesian inference by symbolic model checking. In: Gribaudo, M., Jansen, D.N., Remke, A. (eds.) *QEST 2020*. LNCS, vol. 12289, pp. 115–133. Springer, Cham (2020). https://doi.org/10.1007/978-3-030-59854-9_9
57. Sang, T., Beame, P., Kautz, H.A.: Performing Bayesian inference by weighted model counting. In: *AAAI*, vol. 5, pp. 475–481 (2005)
58. Smolka, S., et al.: Scalable verification of probabilistic networks. In: *PLDI*, pp. 190–203. ACM (2019)
59. van de Meent, J.W., Paige, B., Yang, H., Wood, F.: *An Introduction to Probabilistic Programming*. [arXiv:1809.10756](https://arxiv.org/abs/1809.10756) (2018)
60. Vazquez-Chanlatte, M., Rabe, M.N., Seshia, S.A.: A model counter’s guide to probabilistic systems. *CoRR abs/1903.09354* (2019)
61. Vazquez-Chanlatte, M., Seshia, S.A.: Maximum causal entropy specification inference from demonstrations. In: Lahiri, S.K., Wang, C. (eds.) *CAV 2020*. LNCS, vol. 12225, pp. 255–278. Springer, Cham (2020). https://doi.org/10.1007/978-3-030-53291-8_15
62. Wimmer, R., Braitling, B., Becker, B.: Counterexample generation for discrete-time Markov chains using bounded model checking. In: Jones, N.D., Müller-Olm, M. (eds.) *VMCAI 2009*. LNCS, vol. 5403, pp. 366–380. Springer, Heidelberg (2008). https://doi.org/10.1007/978-3-540-93900-9_29

Open Access This chapter is licensed under the terms of the Creative Commons Attribution 4.0 International License (<http://creativecommons.org/licenses/by/4.0/>), which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if changes were made.

The images or other third party material in this chapter are included in the chapter's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the chapter's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.

