

Uniform Sampling for Timed Automata with Application to Language Inclusion Measurement^{*}

Benoît Barbot^{1**}, Nicolas Basset¹, Marc Beunardeau^{2***}, and Marta Kwiatkowska¹

¹ Department of Computer Science, University of Oxford, United Kingdom

² Ingenico Labs, Paris, France, and École Normale Supérieure, Paris, France

Abstract. Monte Carlo model checking introduced by Smolka and Grosu is an approach to analyse non-probabilistic models using sampling and draw conclusions with a given confidence interval by applying statistical inference. Though not exhaustive, the method enables verification of complex models, even in cases where the underlying problem is undecidable. In this paper we develop Monte Carlo model checking techniques to evaluate quantitative properties of timed languages. Our approach is based on uniform random sampling of behaviours, as opposed to isotropic sampling that chooses the next step uniformly at random. The uniformity is defined with respect to volume measure of timed languages previously studied by Asarin, Basset and Degorre. We improve over their work by employing a zone graph abstraction instead of the region graph abstraction and incorporating uniform sampling within a zone-based Monte Carlo model checking framework. We implement our algorithms using tools PRISM, SageMath and COSMOS, and demonstrate their usefulness on statistical language inclusion measurement in terms of volume.

1 Introduction

Since the seminal work of Alur and Dill [1], timed automata (TAs) have been widely studied in the context of real-time systems verification. Several algorithms from the classical automata-theoretic verification were successfully lifted to the timed case. In spite of this, many problems become undecidable, the most important being the inclusion of timed languages. One way to circumvent undecidability is to employ statistical methods, where results are given with some confidence level. However, timed automata are non-stochastic models and it is not clear a priori with what probability to sample runs when performing statistical experiments. A natural answer is given by the maximal entropy principle: “without knowledge a priori on the distribution of probability to be taken, the

^{*} This work is supported by ERC AdG VERIWARE.

^{**} Now in LACL, Université Paris Est Créteil, France

^{***} Contributed to the work during an internship funded by ERC AdG VERIWARE

one with maximal entropy should be preferred” [15]. A maximal entropy stochastic process for timed automata was recently proposed in [7]. Essentially, this is the stochastic process that yields the most uniform sampling when the length of the timed words tends to infinity. By uniform sampling we mean that all timed words of a given length have the same density of probability to be chosen.

In this paper we propose several algorithms to achieve uniform sampling of timed words in timed languages. The methods are based on the theory of volumetry of timed languages recently developed by Asarin, Degorre and Basset [3], which provides means for quantitative measurement of languages in terms of volume. Here, we employ this theory to achieve statistical estimation of volume and demonstrate its usefulness for language inclusion measurement. The accuracy of statistical estimation depends on the ability to uniformly sample the executions. The method provided in [7], where the transitions of a TA were annotated with probability functions so that the resulting stochastic process enables random simulation in the most uniform way possible, is based on spectral attributes of a functional operator Ψ (an analogue in the TA context of the adjacency matrix of a graph) [3]. Unfortunately, it is not practical, as it relies on the region graph abstraction and the computation of eigenfunctions. In this paper, we overcome this problem by adopting a zone-based approach and approximating the probability functions of [7] with quotients of the volume functions.

Contributions. (i) We provide a zone-based computation of volume functions for TAs, which enables the first practical implementation of volumetry of timed languages. (ii) We develop three methods (Method 1-3) to sample in a (quasi) uniform manner timed words in a language recognised by a deterministic timed automaton (DTA). In particular, we propose a receding horizon framework that allows us to approximate the maximal entropy stochastic process discussed above. (iii) We apply uniform sampling for DTAs to uniform sampling and volume measurement for arbitrary timed languages, provided the membership problem for the language is decidable. (iv) We have implemented the algorithms presented here in PRISM [16] (for the splitting of the DTA into zones), SageMath [20] (for the computation of volume functions) and COSMOS [4] (for the random generation of timed words and property checking) and illustrate them on several examples, with encouraging results. Omitted proofs and further details can be found in [5].

Related work. The theory of volumetry of timed languages has been studied and applied to robustness analysis [3], timed channel coding [2] and combinatorics of permutations [6], but has not yet been applied in practice.

The *recursive method* for uniform sampling is a well-known method in discrete combinatorics [12] whose generalisation to the timed case (Method 1 here) was already done for very specific timed languages in [6].

Monte Carlo model checking was proposed in [13] for discrete models to randomly explore their behaviour by means of simulating execution paths. Similarly, statistical model checking [21] uses simulation to verify temporal logic properties with statistical guarantees, and has been applied to stochastic timed/hybrid

systems [10]. This avoids state-space explosion, thus ensuring the feasibility of verification of complex models, and has also been used to check undecidable properties [10]. Here we implement Monte Carlo techniques for TAs.

Monte Carlo or statistical model checking usually employs an *isotropic* random walk to explore the executions (as explained in [19,11] for discrete models). This involves choosing uniformly at random, at each step of the simulation, the next transition from those available. It has been argued that the isotropic methods are not able to efficiently perform uniform sampling of the behaviours (see e.g. the pathological examples in [19] for sampling of lassos and [11] for sampling paths in a finite-state automaton). Here we implement uniform sampling based on the tool COSMOS, but the techniques are more generally applicable and can be implemented in other tools, for example UPPAAL-SMC [10], which supports user-defined distributions.

Statistical model checkers such as UPPAAL-SMC consider timed automata augmented with probability distribution on transitions that are either user-defined or given “by default”. Thus, the model to verify is already probabilistic and specifications are written in temporal logic with probabilistic operators. Our work addresses a different and novel question: how can one use statistical experiments on a non-probabilistic timed language and draw conclusions about that language, without being given probability distributions on it?

2 Preliminaries

2.1 Timed languages and volumetry

A *timed word* $\alpha = (t_1, a_1) \dots (t_n, a_n)$ is a word over the alphabet $\mathbb{R}_{\geq 0} \times \Sigma$, where $\mathbb{R}_{\geq 0}$ denotes the set of non-negative reals and Σ is a finite alphabet of *events*. Times t_i represent *delays* between events a_{i-1} and a_i . Throughout this paper, delays will be bounded³ by an integer constant M . A *timed language* L is a set of timed words. Given $n \geq 0$, we denote by L_n the timed language L restricted to timed words of length n . For every timed language L and every word $w = a_1 \dots a_n \in \Sigma^n$, we define $P_w^L = \{(t_1, \dots, t_n) \mid (t_1, a_1) \dots (t_n, a_n) \in L\}$, and denote by $\text{Vol}(P_w^L)$ its (hyper-)volume.

Example 1 (Running example). Examples of such hyper-volumes are given in Fig. 1. Anticipating what follows, these sets correspond to the timed language restricted to timed words of length 2 of the TA depicted in Fig. 2 (Left).

For a fixed n , we define the *n-volume* of L as follows:

$$\text{Vol}(L_n) = \sum_{w \in \Sigma^n} \text{Vol}(P_w^L) = \sum_{a_1 \in \Sigma} \int_0^M \dots \sum_{a_n \in \Sigma} \int_0^M \mathbf{1}_{P_w^L}(t) dt_1 \dots dt_n.$$

Continuing the example; the hyper-volume for dimension 2 is calculated as

$$\text{Vol}(L_2) = \text{Vol}(P_{ab}^L) + \text{Vol}(P_{aa}^L) + \text{Vol}(P_{ba}^L) + \text{Vol}(P_{bb}^L) = 3.5 + 2 + 4 + 2 = 11.5.$$

³ Our approach to timed languages is based on volume and does not apply, in its present form, to unbounded delays that result in infinite volume.

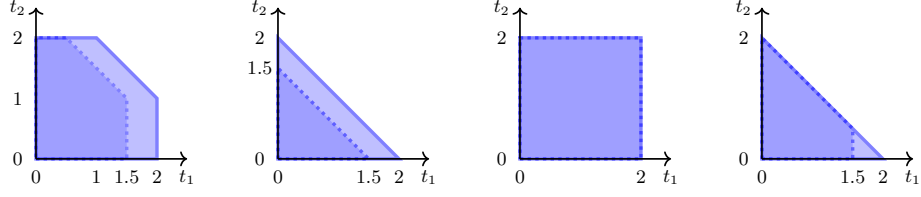


Fig. 1. From left to right, languages P_{ab}^L , P_{aa}^L , P_{ba}^L and P_{bb}^L for the running example (Example 1). The darker areas corresponds to initial clock vector $(x, y) = (0.5, 0)$.

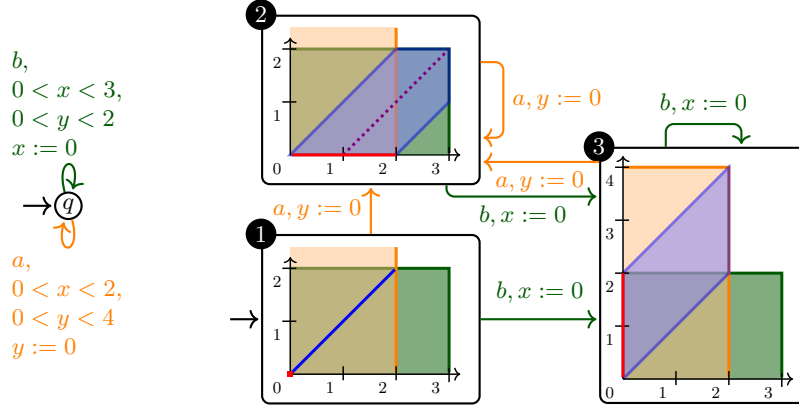


Fig. 2. Left: a DTA. Right: the same DTA obtained after applying the forward reachability algorithm. Entry zones are represented in red. Guards for a and b are the same in the two TAs. The blue part represents clock vectors reachable through entry zones by time elapsing. In location 2, the guard of transition b should be split along the dotted line to obtain the split DTA of Fig 3.

We define the *uniform probability distribution* on a timed language L by assigning weight $1/\text{Vol}(L_n)$ to every timed word of length n . The main purpose of this article is to show how to sample according to that distribution when the language is recognised by a timed automaton. For instance, the probability of a uniformly sampled timed word to fall in the set $E = \{(t_1, b)(t_2, a) \mid t_1 \in (0, 1), t_2 \in (0, 2)\}$ is $\text{Vol}(E)/\text{Vol}(L_2) = 2/11.5 \approx 0.17$.

Given two timed languages L, L' over the same alphabet of events Σ , we say that L' is a *tight under-approximation* of L if, for all $w \in \Sigma^*$, $P_w^{L'} \subseteq P_w^L$ and $\text{Vol}(P_w^L \setminus P_w^{L'}) = 0$; hence $\text{Vol}(P_w^L) = \text{Vol}(P_w^{L'})$. In particular, timed words uniformly sampled in L' are uniformly sampled in L .

2.2 Timed automata

Let X be a finite set of non-negative real-valued variables called *clocks*. Here we assume that clocks remain bounded by a constant $M \in \mathbb{N}$. A *clock constraint* has the form $x \sim c$ or $x - y \sim c$ where $\sim \in \{\leq, <, =, >, \geq\}$, $x, y \in X$, $c \in \mathbb{N}$. A *guard* is a finite conjunction of clock constraints; it is called open if its constraints

involve only strict inequalities. A *zone* is a set of clock vectors $\mathbf{x} \in [0, M]^X$ satisfying a guard. For a clock vector $\mathbf{x} \in [0, M]^X$ and a non-negative real t , we denote by $\mathbf{x} + t$ (resp. $\mathbf{x} - t$) the vector $\mathbf{x} + (t, \dots, t)$ (resp. $\mathbf{x} - (t, \dots, t)$).

A *timed automaton* (TA) $\mathcal{A}$ is a tuple $(\Sigma, X, Q, i_0, F, \Delta)$ where Σ is a finite set of events; X is a finite set of clocks; Q is the finite set of *locations*; i_0 is the initial location; $F \subseteq Q$ is the set of final locations; and Δ is the finite set of *transitions*. Any transition $\delta \in \Delta$ has an *origin* $\delta^- \in Q$, a *destination* $\delta^+ \in Q$, a label $a_\delta \in \Sigma$, a *guard* $\mathbf{g}_\delta$ and a *reset* function $\mathbf{r}_\delta$ determined by a subset of clocks $B \subseteq X$: it resets to 0 all the clocks in B and does not modify the value of the other clocks.

A *timed transition* is an element (t, δ) of $\mathbb{A} \stackrel{\text{def}}{=} [0, M] \times \Delta$. The *delay* t represents the time before firing the transition δ . A *state* $s = (q, \mathbf{x}) \in Q \times [0, M]^X$ is a pair of a location and a clock vector. Given a state $s = (q, \mathbf{x})$ and a timed transition $\alpha = (t, \delta) \in \mathbb{A}$, the *successor* of s by α is denoted by s_α and defined as follows. If $\delta^- = q$ and $\mathbf{x} + t$ satisfies the guard $\mathbf{g}_\delta$ then $s_\alpha = (\delta^+, \mathbf{r}_\delta(\mathbf{x} + t))$ else $s_\alpha = \perp$. Here and in the rest of the paper $\perp$ represents undefined states. A sequence of timed transitions is called a *timed path*. We extend the successor action to timed paths by induction: $s_\varepsilon = s$ and $s_{(\alpha\alpha')} = (s_\alpha)_{\alpha'}$ for all states s , timed transitions $\alpha \in \mathbb{A}$ and timed paths $\alpha' \in \mathbb{A}^*$. The initial state of the timed automaton is $s = (i_0, \mathbf{0})$. The *labelling* of a timed path $(t_1, \delta_1) \dots (t_n, \delta_n)$ is the timed word $(t_1, a_{\delta_1}) \dots (t_n, a_{\delta_n}) \in ([0, M] \times \Sigma)^*$. The *timed language* $L(\mathcal{A})$ of a timed automaton $\mathcal{A}$ is the set of timed words that are labellings of timed paths α such that $s_\alpha \in F \times [0, M]^X$. We also write $L_n(\mathcal{A})$ instead of $(L(\mathcal{A}))_n$.

For a guard g , we denote by $\text{TE}^{-1}(g)$ the set of clock vectors from which g can be reached when time elapses; formally, $\text{TE}^{-1}(g) = \{\mathbf{x} \mid \exists t \geq 0, \mathbf{x} + t \in g\}$. Given a state $s = (q, \mathbf{x})$ we denote by $\Delta(s)$ the set of transition *available* from s , that is such that $\delta^- = q$ and $\mathbf{x} \in \text{TE}^{-1}(\mathbf{g}_\delta)$. Given a state $s = (q, \mathbf{x})$ and a transition $\delta \in \Delta(s)$, we define $\text{lb}_\delta(s) \stackrel{\text{def}}{=} \inf \{t \mid \mathbf{x} + t \in \mathbf{g}_\delta\}$ and $\text{ub}_\delta(s) \stackrel{\text{def}}{=} \sup \{t \mid \mathbf{x} + t \in \mathbf{g}_\delta\}$ so that the condition $\mathbf{x} + t \in \mathbf{g}_\delta$ is equivalent to $t \in (\text{lb}_\delta(s), \text{ub}_\delta(s))$.

A *deterministic timed automaton* (DTA) is a TA such that no clock vector can satisfy guards of pairwise distinct transitions with the same label and origin. This implies that timed words and timed paths of a DTA are in one-to-one correspondence. We are interested in the prefixes of infinite timed words of a DTA. To be sure that $L_n(\mathcal{A})$ contains exactly the prefixes of size n , we consider only DTAs that satisfy the two following conditions: (i) every location is final, (ii) from every reachable state, there is a timed transition that can be taken.

2.3 Equations on timed languages and volumes

Given a DTA $\mathcal{A}$, we denote by $L_n(s)$ the n -th *timed language* recognised from a state s and defined inductively as follows: $L_0(s) = \{\varepsilon\}$, and

$$L_{n+1}(s) = \bigcup_{\delta \in \Delta(s)} \bigcup_{t \in I(s, \delta)} (t, a_\delta) L_n(s_{(t, \delta)}). \quad (1)$$

For the running example and initial state $[q, (0.5, 0)]$ we have:

$$L_2([q, (0.5, 0)]) = \bigcup_{t \in (0, 1.5)} (t, a) L_1([q, (0.5 + t, 0)]) \cup \bigcup_{t \in (0, 2)} (t, b) L_1([q, (0, t)]). \quad (2)$$

The language $L_2([q, (0.5, 0)])$ is depicted in Fig 1.

We also parametrise the volume by the initial state and define the n -th volume function as $v_n(s) = \mathbf{Vol}(L_n(s))$. These functions can be defined recursively by replacing union over intervals by integrals and union over transitions by finite sums in (1). We obtain $v_0(s) = 1$ and

$$v_{n+1}(s) = \sum_{\delta \in \Delta(s)} \int_{\text{lb}_\delta(s)}^{\text{ub}_\delta(s)} v_n(s_{(t, \delta)}) dt. \quad (3)$$

For the running example, passing to volumes in (2) yields

$$v_2([q, (0.5, 0)]) = \int_0^{1.5} v_1([q, (0.5 + t, 0)]) dt + \int_0^2 v_1([q, (0, t)]) dt. \quad (4)$$

A key idea used in [3,7] is to rewrite (3) as

$$v_{n+1}(s) = \Psi(v_n)(s) \quad (5)$$

where Ψ is an integral operator defined by

$$\Psi(f)(s) = \sum_{\delta \in \Delta(s)} \Psi_\delta(f)(s) \quad \text{with} \quad (6)$$

$$\Psi_\delta(f)(s) = \int_{\text{lb}_\delta(s)}^{\text{ub}_\delta(s)} f(s_{(t, \delta)}) dt. \quad (7)$$

Thus, volume functions are defined via iteration of the operator Ψ on the constant function $\mathbf{1}$: $v_n = \Psi^n(\mathbf{1})$. In [3,7], the state space was decomposed into regions, which guaranteed algebraic properties such as polynomial volume functions at the price of an explosion of the number of locations of the TA. A TA before such a decomposition into regions has volume functions that are complicated (piecewise defined), and hence difficult to handle in practice. Here we want to keep volume functions simple (polynomial) while keeping the set of locations small. For this we adopt a zone-based approach.

The idea of the zone-based decomposition described in the next section is to split the state space into several pieces in which the functions $\text{lb}_\delta(s)$ and $\text{ub}_\delta(s)$ have simple form, ensuring that every volume function $v_n = \Psi^n(\mathbf{1})$ restricted to any location is polynomial (see Table 1).

3 Volume function computation for DTAs

In this section we explain how to transform a DTA $\mathcal{A}$ into a DTA $\mathcal{A}'$ called *split DTA* that facilitates efficient volume computation.

Table 1. First volume functions $v_n[l_i, (x, y)]$ associated to the TA of Fig. 3.

	$[l_0, (0, 0)]$	$[l_1, (x, 0)]$	$[l_2, (0, y)]$	$[l_3, (x, 0)]$
v_0	1	1	1	1
v_1	4	$-x + 4$	$-y + 4$	$-2x + 5$
v_2	15	$-4x + 15$	$\frac{1}{2}y^2 - 4y + 15$	$-\frac{1}{2}x^2 - 6x + \frac{35}{2}$
v_3	$\frac{335}{6}$	$-15x + \frac{335}{6}$	$-\frac{1}{6}y^3 + 2y^2 - 15y + \frac{335}{6}$	$-\frac{1}{6}x^3 - \frac{1}{2}x^2 - 25x + \frac{133}{2}$

Decomposition into zones. We first apply a forward reachability algorithm, implemented for instance in PRISM [16], which returns the so-called *forward-reachability graph*, that is, a finite graph with annotations, which we view as a DTA (the annotations are essentially, for each edge δ , the guard $\mathbf{g}_\delta$ and label a_δ and, for each location l , the zone Z_l which is entered). Formally, we say that a TA is *decomposed into zones* if, for every $l \in Q$, there is a zone Z_l called the *entry zone* of l , such that the entry zone of the initial state is $\{\mathbf{0}\}$ and, for every transition δ , the successors of states in $\{\delta^-\} \times Z_{\delta^-}$ through δ with some delay are in $\{\delta^+\} \times Z_{\delta^+}$, that is, $\{\mathbf{r}_\delta(\mathbf{x} + t) \mid \mathbf{x} \in Z_{\delta^-}, \mathbf{x} + t \in \mathbf{g}_\delta\} \subseteq Z_{\delta^+}$. We denote by $\mathbb{S} = \cup_{l \in Q} \{l\} \times Z_l$ the set of states corresponding to entry zones. The forward-reachability graph for the running example is given in Fig. 2 (Right).

Guard split. Let δ be the transition from location 2 to location 3 in the automaton of Fig. 2 (Right), then $g_\delta \stackrel{\text{def}}{=} (0 < x < 3) \wedge (0 < y < 2)$. Then one can see that $\text{ub}_\delta(2, (x, 0)) = 2$ if $x \in (0, 1)$ (due to guard $y < 2$) and $\text{ub}_\delta(2, (x, 0)) = 3 - x$ if $x \in (1, 2)$ (due to guard $x < 3$). The guard g_δ thus needs to be split into two (along the dotted line in the figure) to achieve a simpler form for ub_δ . It is well known how to get the tightest constraints of a guard and get rid of redundant constraints using the Floyd-Warshall algorithm (see e.g. [8]). A guard is said to be *upper-split* (*lower-split*) if there is at most one useful constraint (that is, not implied by other constraints) of the form $x_j < a$ ($x_j > a$). The guard g_δ discussed above is not upper-split as the two constraints $x < 3$ and $y < 2$ are both useful. Analogous definitions hold for lower-bounds and a guard is said to be *split* if it is both lower-split and upper-split.

Pre-stability. A second phenomenon we want to avoid is when the set of available transitions $\Delta(q, \mathbf{x})$ is not constant on the entry zone of q . A TA decomposed into zones is called *pre-stable* if, for every location q and clock vector $\mathbf{x} \in Z_q$, the set of transitions $\Delta(q, \mathbf{x})$ is exactly the set of transitions δ whose origin is q . Equivalently, a TA is pre-stable if $Z_{\delta^-} \subseteq \text{TE}^{-1}(\mathbf{g}_\delta)$ for every δ . In case we detect a transition such that $Z_{\delta^-} \not\subseteq \text{TE}^{-1}(\mathbf{g}_\delta)$ we will split the zone Z_{δ^-} to isolate $\text{TE}^{-1}(\mathbf{g}_\delta) \cap Z_{\delta^-}$ from its complement. Continuing the example above, after splitting g_δ the functions associated to each new guard are null for $x \in (0, 1)$ or $x \in (1, 2)$. Location 2 is split into two locations of the final TA of Fig. 3: l_1 for $(0, 1)$ and l_3 for $(1, 2)$. Every incoming transition to location 2 is split accordingly into two transitions (one orange to l_1 and one purple to l_3).

Trimming. Last but not least, we say that a TA is *trimmed* if the set of outgoing transitions of each location is non-empty. A TA is called *split* if it is pre-stable, trimmed and all the guards of its transitions are split and open. It implies, in particular, that, for every entry state $s \in \mathbb{S}$, $\Delta(s)$ is not empty and for all transition $\delta \in \Delta(s)$ it holds that $\text{ub}_\delta(s) - \text{lb}_\delta(s) > 0$. Note that *opening guards*, that is, transforming non-strict inequalities into strict ones is made wlog., as it only removes part of the language that has a null volume measure.

Splitting algorithm. We propose an algorithm to transform a DTA into a split DTA such that the language of the latter is a tight under-approximation of the language of the former (see Theorem 1). First, we apply a forward reachability algorithm to obtain a DTA decomposed into zones and open its guards. Then we successively split zones that falsify pre-stability and guard split conditions, until the conditions are satisfied in the DTA. The splitting algorithm maintains a stack of transitions that need to be checked, which initially contains all the transitions. As the algorithm proceeds, transitions are popped from the stack and are checked against pre-stability and guard split conditions. If one test fails, the zone (or guard) is split accordingly into several zones (or open guards) and the transitions that are affected are added to the stack (incoming transitions to, and outgoing transitions from the split zone). When no more transition need to be checked (i.e. the stack is empty), the TA is split and the algorithm terminates. This occurs in a finite number of steps since transitions are added to the stack only when a zone is split into strictly smaller sub-zones, and there are finitely many zones (as the clocks are bounded by a constant M).

Theorem 1. *Given a DTA $\mathcal{A}$, one can construct (using the algorithm sketched above) a split DTA $\mathcal{A}'$ that recognises a tight under-approximation of $L(\mathcal{A})$.*

The splitting algorithm and the proof can be found in the technical report [5].

Volume function of a split DTA. We have the following result.

Proposition 1. *Given a split DTA $\mathcal{A}$ and $n \in \mathbb{N}$, denote by c the maximal affine dimension of an entry zone of $\mathcal{A}$. One can compute the volume function v_k for $k \leq n$ in time and space complexity $O(n^{c+2}|Q_{\mathcal{A}}|)$ using dynamic programming based on the recursive equation (3). Each volume function v_k restricted to a location q is a polynomial of degree at most k that is positive on Z_q .*

Example 2. We have implemented the splitting algorithm sketched in Sect. 3 and applied it to the DTA of Fig. 2 (Right) to obtain the DTA of Fig. 3. Our program also returns for each transition δ of the output DTA the interval $(\text{lb}_\delta, \text{ub}_\delta)$, allowing us to compute with SageMath the operator Ψ as well as the volume functions. On the example, for $f : \mathbb{S} \rightarrow \mathbb{R}$, $(x, y) \in Z_l$ with $l \in \{l_0, \dots, l_3\}$,

$$\begin{aligned} \Psi(f)[l_0, (0, 0)] &= \int_0^1 f(l_1, (t, 0))dt + \int_0^2 f(l_2, (0, t))dt + \int_1^2 f(l_3, (t, 0))dt; \\ \Psi(f)[l_1, (x, 0)] &= \int_0^{1-x} f(l_1, (x+t, 0))dt + \int_0^2 f(l_2, (0, t))dt + \int_{1-x}^{2-x} f(l_3, (x+t, 0))dt; \\ \Psi(f)[l_2, (0, y)] &= \int_0^1 f(l_1, (t, 0))dt + \int_0^{2-y} f(l_2, (0, y+t))dt + \int_1^2 f(l_3, (t, 0))dt; \\ \Psi(f)[l_3, (x, 0)] &= \int_0^{3-x} f(l_2, (0, t))dt + \int_0^{2-x} f(l_3, (x+t, 0))dt. \end{aligned}$$

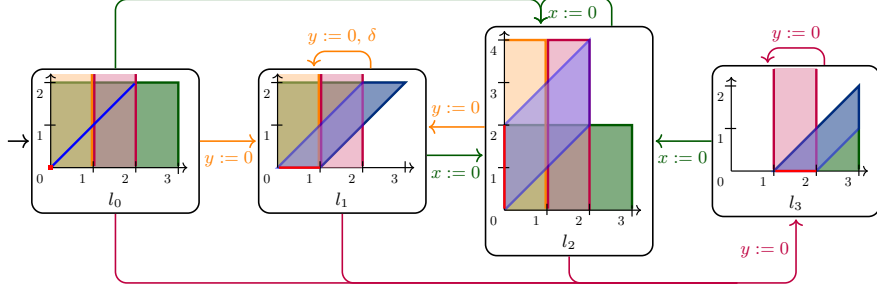


Fig. 3. The split form of the running example (Example 1).

First volume functions computed using Equation (5) are given in Table 1.

4 Sampling methods for timed languages of DTAs

In this section we consider random sampling of timed words. We first give a method that achieves exact uniform sampling when the length of timed words to be generated is finite; we speak of finite horizon. When the length is infinite or too long to be treated by the previous method, we consider a receding horizon method, where, at the k -th step of the generation, the next timed letter is chosen according to the volume of the timed words for the next m steps; these possible futures constitute a finite receding horizon. At the limit, where the receding horizon becomes infinite ($m \rightarrow \infty$), this can be interpreted as a stochastic process over runs of maximal entropy [7].

Parametric probability distributions. A *discrete probability distribution* (DPD) on a finite set A is a function $\mathbf{dpd} : A \rightarrow [0, 1]$ such that $\sum_{a \in A} \mathbf{dpd}(a) = 1$. A *probability density function* (PDF) on an interval (a, b) is a Lebesgue measurable function $\mathbf{pdf} : (a, b) \rightarrow \mathbb{R}_{\geq 0}$ such that $\int_a^b \mathbf{pdf}(t) dt = 1$. Values of DPD and PDF are referred to as weights. The DPD $\mathbf{isoDPD}(A)$ on a set A (resp. the PDF $\mathbf{isoPDF}(a, b)$ on an interval (a, b)) that attributes the same weight to every $a \in A$ (resp. $t \in (a, b)$) is called *isotropic*. In other words, $\mathbf{isoDPD}(A)(a) = 1/|A|$ for every $a \in A$ (resp. $\mathbf{isoPDF}(a, b)(t) = 1/(b - a)$ for every $t \in (a, b)$). PDFs considered in the following are just polynomials on the delay variable t . Their coefficients depend on the current state (location and clock values) and on the transition to fire. Choosing a delay t according to a PDF can be done using the *inverse method*: a random number r is drawn uniformly in $(0, 1)$, and the output $t \in (a, b)$ is the unique solution of $\int_a^t \mathbf{pdf}(t') dt' - r = 0$. In the case of the isotropic PDF on (a, b) , the output t is just $a + r(b - a)$.

Random generation of timed words in $L_n(s)$ for a given state $s \in \mathbb{S}$ is done as follows: for $k = 1..n$, pick randomly the next transition δ according to a DPD $\mathbf{dpd}_s^k$ parametrised by the current state s , then chose the delay t in $(\text{lb}_\delta(s), \text{ub}_\delta(s))$ according to a PDF $\mathbf{pdf}_{s,\delta}^k$ parametrised by the current state s and the transition just chosen; take the successor of s by (t, δ) as the new current state s ; output (t, a_δ) and repeat the loop.

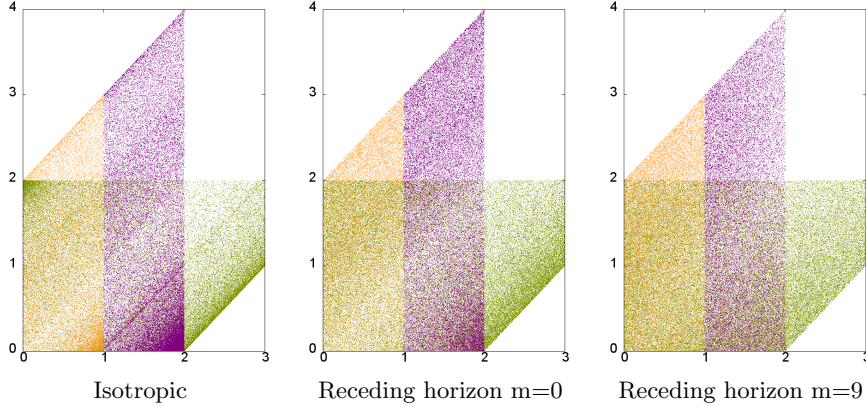


Fig. 4. Trajectories of the running example (Fig. 3) sampled using isotropic sampling (Left) and Method 2 with receding horizon $m = 0$ (Middle) and $m = 9$ (Right). Each point of a given colour corresponds to a clock vector where a transition of that colour occurs. Each plot visualises a single trajectory with 200,000 transitions. The receding horizon $m = 9$ visibly yields the most uniform sampling. The receding horizon sampling with $m = 0$ is already more uniform than the isotropic sampling as the former assigns weights to transitions proportional to lengths of intervals $\left(\text{dpd}_s = \frac{\text{ub}_\delta(s) - \text{lb}_\delta(s)}{v_1(s)}\right)$.

This random generation outputs timed words of $L_n(s)$ with weights given by

$$\text{Weight}[(t_1, a_1) \cdots (t_n, a_n)] \stackrel{\text{def}}{=} \prod_{k=1}^n \text{dpd}_{s_{k-1}}^k(\delta_k) \text{pdf}_{s_{k-1}, \delta_k}^k(t_k). \quad (8)$$

where, for every $k = 1..n$, s_{k-1} is the state before the k th sampling loop, (t_k, δ_k) is the k th timed transition randomly picked during the k th sampling loop and a_k is the label of δ_k .

Isotropic and uniform sampling. *Isotropic sampling*⁴ relies on using in each step the isotropic DPD $\text{isoDPD}(\Delta(s))$ and the isotropic PDF $\text{isoPDF}(I(s, \delta))$. These distributions are particularly simple to sample, but when the length of samples grows the probability concentrates on small sections of the runs, see Fig. 4 (Left). By contrast, *uniform sampling* for $L_n(s)$ assigns the same weight $1/v_n(s)$ to every timed word. In other words, for any measurable set $B \subseteq L_n(s)$ the probability $\text{Vol}(B)/v_n(s)$ to fall in this set is proportional to its measure.

The recursive method for uniform sampling. The idea of the recursive method for uniform sampling of n -length timed words from a state s is to choose the first delay t and transition δ according to well chosen DPD and PDF that depend on the volume functions v_n and v_{n-1} , and then recursively apply uniform sampling to generate an $(n - 1)$ -length timed word from the updated state $s_{(t, \delta)}$.

⁴ Note that some works, consider instead sampling the delay first and then the transitions available in the state updated by the delay (see [9]).

Define, for every function $f : \mathbb{S} \rightarrow \mathbb{R}_{>0}$ and state s , the DPD $\omega(f, s) : \delta \mapsto \frac{\Psi_\delta(f)(s)}{\Psi(f)(s)}$. If moreover δ is given, define the PDF $\varphi(f, s, \delta) : t \mapsto \frac{f(s_{(t,\delta)})}{\Psi_\delta(f)(s)}$ from $(\text{lb}_\delta(s), \text{ub}_\delta(s))$ to $\mathbb{R}_{>0}$.

Method 1 (Exact uniform sampling) *Given a split DTA and $n \in \mathbb{N}$, precompute the volume functions $v_0 = \mathbf{1}, \dots, v_n = \Psi^n(\mathbf{1})$ (see Proposition 1), then the uniform sampling of n -length timed words can be achieved in linear time using the following sequences of DPDs and PDFs: $(\omega(v_{n-k}, s), \varphi(v_{n-k}, s, \delta))_{k=1..n}$*

Proof. Using the same notation as in (8), it holds that

$$\begin{aligned} \text{Weight}[(t_1, a_1) \cdots (t_n, a_n)] &= \prod_{k=1}^n \omega(v_{n-k}, s_{k-1})(\delta_k) \varphi(v_{n-k}, s_{k-1}, \delta_k)(t_k) \\ &= \prod_{k=1}^n \frac{\Psi_\delta(v_{n-k})(s_{k-1})}{v_{n-k+1}(s_{k-1})} \frac{v_{n-k}(s_k)}{\Psi_\delta(v_{n-k})(s_{k-1})} = \frac{v_0(s_{n-1})}{v_n(s_0)} = \frac{1}{v_n(s_0)}. \end{aligned}$$

Example 3. We illustrate the DPDs and PDFs used in the last but one step of the uniform random sampling for the running example, obtained from volume functions of Table 1. Consider the state $s = (l_1, (x, 0))$ with $x \in (0, 1)$ and δ the self-loop on l_1 (see Fig. 3). Then $(\text{lb}_\delta(s), \text{ub}_\delta(s)) = (0, 1 - x)$ and $s_{(t,\delta)} = (l_1, (x + t, 0))$. The DPD used to choose δ is

$$\text{dpd}_s^{n-1}(\delta) = \frac{1}{v_2(s)} \int_{\text{lb}_\delta(s)}^{\text{ub}_\delta(s)} v_1(s_{(t,\delta)}) dt = \int_0^{1-x} \frac{4 - x - t}{15 - 4x} dt = \frac{7 - 8x + x^2}{30 - 8x}$$

The PDF used to choose t is

$$\text{pdf}_{s,\delta}^{n-1}(t) = \frac{1_{t \in (\text{lb}_\delta(s), \text{ub}_\delta(s))} v_1(s_{(t,\delta)})}{\text{dpd}_s^{n-1}(\delta) v_2(s)} = 1_{t \in (0, x)} \frac{8 - 2x - 2t}{7 - 6x + x^2}$$

Random sampling with finite receding horizon. With the previous method, the k -th timed transition of a run of length n is sampled according to DPD and PDF that depend on k and n . This dependency on k and n is not suitable for large n as it requires storage of as many polynomials as the length of the run to generate n . Also, one might wish to randomly generate arbitrarily long runs without a prescribed bound on the length. To take the k th timed transition in the recursive method for uniform sampling, we use DPD and PDF that depend on v_{n-k} , that is, on the volume measure of the possible $(n - k)$ step future. The idea of the following method is to replace $(n - k)$ by a fixed $m \ll n$ at every step of the sampling. The constant m can be seen as a receding horizon used in control theory [17]. At each step we consider only the possible m step future to generate the next timed transition.

Method 2 (Random sampling with finite receding horizon m) *Given a split DTA, $n \in \mathbb{N}$ and $m \in \mathbb{N}$, precompute the volume functions $v_0 = \mathbf{1}, \dots, v_m = \Psi^m(\mathbf{1})$ (see Proposition 1), then sample n -length timed words in linear time using the same DPD $\omega(v_m, s)$ and PDF $\varphi(v_m, s, \delta)$ for every $k = 1..n$.*

The precomputation is polynomial in m . Hence this methods is more efficient than Method 1 when $m \ll n$, but it does not yield exact uniform sampling.

Table 2. Table for Example 4.

m	$(C^+/C^-) - 1$	$n_{0.01}$	m	$(C^+/C^-) - 1$	$n_{0.01}$	m	$(C^+/C^-) - 1$	$n_{0.01}$
0	3	1	4	3.272×10^{-4}	35	8	2.364×10^{-7}	42 098
1	0.3229	2	5	8.431×10^{-5}	124	9	2.520×10^{-8}	394 801
2	1.659×10^{-2}	3	6	9.308×10^{-6}	1 076	10	5.304×10^{-9}	1.8760×10^6
3	4.444×10^{-3}	6	7	1.409×10^{-6}	7 069	11	4.487×10^{-10}	2.2178×10^7

Quasi-uniform random sampling. We now present a trade-off between exact uniform sampling (Method 1) and the finite receding horizon sampling (Method 2). We give bounds on the distance to uniformity for this method, which we conjecture to be small in practice for small horizon m . This conjecture is supported by theoretical results of previous works [3,7,5] and by practical experiments (notably in Example 4 below).

Method 3 (Switching method for quasi-uniform sampling) *Given a split DTA, $n \in \mathbb{N}$ and $m \in \mathbb{N}$, precompute the volume functions $v_0 = \mathbf{1}, \dots, v_m = \Psi^m(\mathbf{1})$ (see Proposition 1), then generate the $n - m$ first letters as in Method 2 and use Method 1 from the current state for the last m steps.*

This method ensures quasi-uniform sampling in the following sense.

Theorem 2. *If in Method 3 there exist constants $C^-, C^+ \in \mathbb{R}_{>0}$ such that $C^- v_{m+1} \leq v_m \leq C^+ v_{m+1}$, then the weight of every timed word lies in the interval $[(1 - \varepsilon_{m,n})/v_n(s_0), (1 + \varepsilon_{m,n})/v_n(s_0)]$, with $\varepsilon_{m,n} = (C^+/C^-)^{(n-m-1)} - 1$.*

Example 4. For the running example (Example 1) we determine the tightest constraints $C^- \stackrel{\text{def}}{=} \inf_{s \in \mathbb{S}} v_m(s)/v_{m+1}(s)$ and $C^+ \stackrel{\text{def}}{=} \sup_{s \in \mathbb{S}} v_m(s)/v_{m+1}(s)$ for $m = 0..11$. We observe empirically that C^+/C^- tends to 1 exponentially fast when m grows (see Table 2). Given a maximal tolerated error of ε , one can determine for every m the maximal n , called n_ε , such that $\varepsilon_{m,n} \leq \varepsilon$ for every $n \leq n_\varepsilon$; formally, $n_\varepsilon \stackrel{\text{def}}{=} m + 1 + \lfloor \log_{(C^+/C^-)}(1 + \varepsilon) \rfloor$. First values of $n_{0.01}$ as a function of m are given in Table 2; for instance, using receding horizon for $m = 11$ one can generate timed words of length 20,000,000 with a divergence to uniformity less than 1%.

Our sampling method requires the computation of a complete zone graph, as opposed to on-the-fly techniques used in state-of-the-art statistical model checkers; this is the price we pay for statistical evaluation of quantities of timed words in complex sub-languages as described in the next section.

5 Applications and experiments

5.1 Tackling general timed languages

It is well known that language inclusion for languages recognised by non-deterministic TAs (NTAs) is undecidable, even when a robust semantics is considered [14].

The situation is even worse for stopwatch automata, hybrid automata, etc., for which the reachability problem is undecidable. However, we can handle a statistical variant of the inclusion problem when, first, an overapproximation of the language described by a DTA is known and, second, the languages admit decision procedures for the membership problem defined as: given a language $\mathcal{L}$ and a word w , is $w \in \mathcal{L}$? Our method is based on statistical volume estimation that relies on the quasi-uniform random sampling developed in the previous section. The complexity results given below are expressed in terms of the number of membership queries one has to solve.

Application 1 (Statistical volume estimation) *Given a timed language $\mathcal{L}$, $n \in \mathbb{N}$, a confidence level θ , an error bound ε , and an over-approximation of the language recognised by a DTA $\mathcal{C}$, that is, $\mathcal{L}_n \subset L_n(\mathcal{C})$, define $N \geq (1/\varepsilon^2) \log(\theta/2)$ (Chernoff-Hoeffding bound); draw N samples uniformly at random in $L_n(\mathcal{C})$ and answer N queries for membership in $\mathcal{L}$ to return a value p such that $\text{Vol}(\mathcal{L}_n)/\text{Vol}(L_n(\mathcal{C}))$ lies in $[p - \varepsilon, p + \varepsilon]$ with confidence $1 - \theta$.*

Application 2 (Inclusion measurement) *Given two timed languages $\mathcal{L}'$, $\mathcal{L}''$ and an over-approximation of the two languages recognised by a DTA $\mathcal{C}$ one can use the previous application with $\mathcal{L} = \mathcal{L}' \setminus \mathcal{L}''$ to evaluate the volume $\text{Vol}(\mathcal{L}' \setminus \mathcal{L}'')$. If a positive value is returned, a timed word in $\mathcal{L}'_n \setminus \mathcal{L}''_n$ has been detected and one can surely claim $\mathcal{L}'_n \not\subseteq \mathcal{L}''_n$. Otherwise, a null value allows one to claim with confidence $1 - \theta$ that either the inclusion holds or the difference of the two languages is smaller than $\varepsilon \text{Vol}(L_n(\mathcal{C}))$.*

Application 3 (Uniform sampling) *Given a timed language $\mathcal{L}$ and $n \in \mathbb{N}$, and an over-approximation of the language recognised by a DTA $\mathcal{C}$, that is, $\mathcal{L}_n \subseteq L_n(\mathcal{C})$, draw samples uniformly at random in $L_n(\mathcal{C})$ until one falls in $\mathcal{L}_n$.*

The sampling is uniform: every timed word of $\mathcal{L}_n$ has the same density of probability to be output. The expected number of samplings in $L_n(\mathcal{C})$ to sample one timed word in $\mathcal{L}_n$ is $\text{Vol}(\mathcal{L}_n)/\text{Vol}(L_n(\mathcal{C}))$. The choice of $\mathcal{C}$ is crucial, since if $L_n(\mathcal{C})$ is a too coarse approximation of $\mathcal{L}_n$ the probability of a sample from $L_n(\mathcal{C})$ to be in $\mathcal{L}_n$ is small and the methods become inefficient. We leave as future work the design of heuristics that, given a general timed language $\mathcal{L}$, automatically generate a DTA that recognises a good over-approximation of $\mathcal{L}$.

5.2 Implementation and experiments

We implemented the techniques using three tools: PRISM [16], SageMath [20] and COSMOS [4]. The workflow is depicted in Fig. 5. We modify the tools to meet our needs. We adapted PRISM's forward reachability algorithm to implement the splitting algorithm of Sect. 3. We also export the split zone graph in a file format easy to read for SageMath. We use SageMath to compute distributions and weights of transitions as rational functions of clock valuations, which are exported and read by COSMOS in the form of a Stochastic Petri Net with general

distributions. COSMOS then samples trajectories of this model, checks the membership of the language of a given NTA, and returns the probability. We have modified COSMOS to handle distributions given by arbitrary rational functions and to compute the membership of a timed word in an NTA. Our implementation can be found at <http://www.prismmodelchecker.org/files/qest16>.

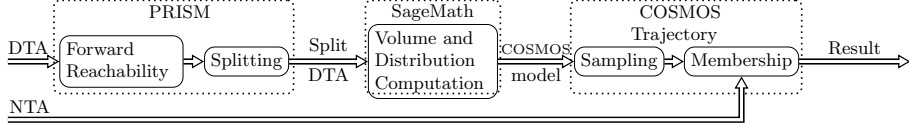


Fig. 5. Tool workflow. For the running example (Example 1), the DTA is the automaton in Fig. 2 (Left), the zone graph is the automaton in Fig. 3, the COSMOS model is the zone graph annotated with probability distributions as described in Example 3, and examples of trajectories are depicted in Fig. 4.

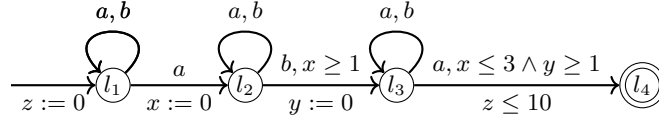


Fig. 6. The NTA $\mathcal{B}$ for Example 5. Every transition has a guard $z \leq 10$ omitted.

Example 5. Let $\mathcal{A}$ be the DTA of the running example (Example 1). The NTA $\mathcal{B}$ of Fig. 6 recognises the timed words that contain aba as a subword within the first 10 time units, where the latter a occurs at most 3 time units after the former and there is at least 1 time unit between b and both as . We have estimated $\text{Vol}(L_{10}(\mathcal{A}) \cap L_{10}(\mathcal{B})) / \text{Vol}(L_{10}(\mathcal{A}))$ by implementing Application 1. Sampling was performed using Method 2 with $m = 5$. The result is in the interval $[0.679, 0.688]$ with confidence level 0.99; 58,000 simulations were used in 5s.

A case study. We additionally consider a larger case study of a failure and repair system modelled as an NTA (see [5] for more details). We consider a model with K machines that need to be fully repaired for the overall system to work properly. Each machine contains N levels of failure and can fail at most n_b times between two full repairs. The model is implemented by an NTA $\mathcal{A}$ with Nn_b locations and $K + 1$ clocks. The property we are interested in is encoded in another NTA $\mathcal{B}$ with 4 locations and 2 clocks. We apply our method by over-approximating the NTA $\mathcal{A}$ with a DTA $\mathcal{C}$ with $R \stackrel{\text{def}}{=} KN$ locations and 2 clocks. The results are reported in Table 3. We use our approach to sample timed words of length 50 of the DTA $\mathcal{C}$ and check their membership in $L_{50}(\mathcal{A})$ and $L_{50}(\mathcal{B})$. We compare receding horizon sampling to isotropic sampling. We observe that for isotropic sampling the probability for a timed word in $L_{50}(\mathcal{A})$ to be in $L_{50}(\mathcal{B})$ (denoted by $P_{50}(\mathcal{B}|\mathcal{A})$) tends to 1 quickly when R increases, which, for large values of R , might be interpreted as an inclusion of the languages. On the other

Table 3. Result of receding horizon sampling compared to isotropic sampling for the case study with two machines ($K = 2$). “Pre Time” is the pre-computation time, “Sim Time” is the simulation time. The meaning of R , $P_{50}(\mathcal{A}|\mathcal{C})$ and $P_{50}(\mathcal{B}|\mathcal{A})$ is described in the text. The receding horizon is $8 + R$. The number of samples is 100,000.

R	Receding horizon					Isotropic		
	Pre Time	#Zones	Sim Time	$P_{50}(\mathcal{B} \mathcal{A})$	$P_{50}(\mathcal{A} \mathcal{C})$	Sim Time	$P_{50}(\mathcal{B} \mathcal{A})$	$P_{50}(\mathcal{A} \mathcal{C})$
4	45s	380	133s	0.999977	0.86539	36s	0.990439	0.03347
6	99s	581	369s	0.997717	0.58701	39s	0.975795	0.05123
8	219s	783	5005s	0.930944	0.06111	56s	0.995179	0.07052
10	417s	985	5773s	0.509091	0.00275	55s	0.999893	0.09325
12	745s	1187	7954s	0.0344828	0.00029	64s	1	0.1019

hand, with the receding horizon sampling the same probability ($P_{50}(\mathcal{B}|\mathcal{A})$) tends to zero, which shows that the model does not satisfy the property. This result demonstrates the necessity of (quasi)-uniform sampling to explore the behaviour of the model, since the results of isotropic simulation significantly diverge from those of (quasi)-uniform simulation, and thus do not yield reliable information about the system.

We also observe that the probability for timed words in the over-approximation $L_{50}(\mathcal{C})$ to fall in $L_{50}(\mathcal{A})$ (denoted by $P_{50}(\mathcal{A}|\mathcal{C})$) tends to zero, meaning that it becomes too crude for large values of R . Thus, tight over-approximations are important to obtain efficient simulation of an NTA through a DTA.

The time required for receding horizon simulation is high compared to isotropic, since it requires sampling of complex distributions involving many polynomials.

6 Conclusion and further work

We have developed the foundations for the practical application of volumetry of timed languages to quantitative and statistical verification of complex properties for TAs. We implemented our work in a tool chain and provide first experiments.

On the theoretical side, we want to show that constants in Method 3 and Theorem 2 can be chosen to guarantee arbitrarily small divergence from exact uniform sampling and consider extending the theory to probabilistic TAs. We would also like to implement membership checking in COSMOS for general timed languages (e.g. recognised by stopwatch automata, LHA, etc.). We also plan to use our random sampling algorithms to detect forgetful cycles described in [3], which are needed to synthesise controllers robust to timing imprecision [18].

References

1. Rajeev Alur and David L. Dill. A theory of timed automata. *Theoretical Computer Science*, 126:183–235, 1994.
2. Eugene Asarin, Nicolas Basset, Marie-Pierre Béal, Aldric Degorre, and Dominique Perrin. Toward a timed theory of channel coding. In *FORMATS’12*, LNCS 7595, 2012.

3. Eugene Asarin, Nicolas Basset, and Aldric Degorre. Entropy of regular timed languages. *Information and Computation*, 241:142–176, 2015.
4. Paolo Ballarini, Benoît Barbot, Marie Duflot, Serge Haddad, and Nihal Peker-gin. Hasl: A new approach for performance evaluation and model checking from concepts to experimentation. *Performance Evaluation*, 90(0):53 – 77, 2015.
5. Benoît Barbot, Nicolas Basset, Marc Beunardeau, and Marta Kwiatkowska. Uni-form sampling for timed automata with application to language inclusion measure-ment. Technical Report CS-RR-16-04, University of Oxford, 2016.
6. Nicolas Basset. Counting and generating permutations using timed languages. In *LATIN*, LNCS 8392, pages 502–513. Springer, 2014.
7. Nicolas Basset. A maximal entropy stochastic process for a timed automaton. *Information and Computation*, 243:50–74, 2015.
8. Johan Bengtsson and Wang Yi. Timed automata: Semantics, algorithms and tools. In *Lectures on Concurrency and Petri Nets, Advances in Petri Nets*, 2003.
9. Dimitri Bohlender, Harold Brientjes, Sebastian Junges, Jens Katelaan, VietYen Nguyen, and Thomas Noll. A review of statistical model checking pitfalls on real-time stochastic models. In *Leveraging Applications of Formal Methods, Verification and Validation.*, volume LNCS 8803. Springer, 2014.
10. Alexandre David, Kim G. Larsen, Axel Legay, Marius Mikucionis, and Danny Bøgsted Poulsen. Uppaal SMC tutorial. *STTT*, 17(4):397–415, 2015.
11. Alain Denise, Marie-Claude Gaudel, Sandrine-Dominique Gouraud, Richard Las-saigne, Johan Oudinet, and Sylvain Peyronnet. Coverage-biased random explo-ration of large models and application to testing. *STTT*, 14(1):73–93, 2012.
12. Philippe Flajolet, Paul Zimmerman, and Bernard Van Cutsem. A calculus for the random generation of labelled combinatorial structures. *Theoretical Computer Science*, 132(1):1–35, 1994.
13. Radu Grosu and Scott A. Smolka. Monte carlo model checking. In *TACAS’05*, 2005.
14. Thomas A. Henzinger and Jean-François Raskin. Robust undecidability of timed and hybrid systems. In *HSCC 2000*, 2000.
15. E. T. Jaynes. Information Theory and Statistical Mechanics. II. *Physical Review Online Archive (Prola)*, 108(2):171–190, October 1957.
16. M. Kwiatkowska, G. Norman, and D. Parker. PRISM 4.0: Verification of proba-bilistic real-time systems. In *Proc. CAV’11*, 2011.
17. Richard M Murray, John Hauser, Ali Jadbabaie, Mark B Milam, Nicolas Pe-tit, William B Dunbar, and Ryan Franz. Online control customization via optimization-based control. *Software-Enabled Control: Information Technology for Dynamical Systems*, page 149, 2003.
18. Youssef Oualhadj, Pierre-Alain Reynier, and Ocan Sankur. Probabilistic robust timed games. In *CONCUR’14*, 2014.
19. Johan Oudinet, Alain Denise, Marie-Claude Gaudel, Richard Lassaig-ne, and Syl-vain Peyronnet. Uniform monte-carlo model checking. In *FASE*, volume 6603, pages 127–140, 2011.
20. W.A. Stein et al. *Sage Mathematics Software (Version 6.9)*. The Sage Develop-ment Team, 2015. <http://www.sagemath.org>.
21. H.L.S. Younes and R.G. Simmons. Statistical probabilistic model checking with a focus on time-bounded properties. *Inf. and Comput.*, 204(9):1368–1409, 2006.