
Environment-Independent Task Specifications via GLTL

Michael L. Littman

Ufuk Topcu

Jie Fu

Charles Isbell

Min Wen

James MacGlashan

Abstract

We propose a new task-specification language for Markov decision processes that is designed to be an improvement over reward functions by being environment independent. The language is a variant of Linear Temporal Logic (LTL) that is extended to probabilistic specifications in a way that permits approximations to be learned in finite time. We provide several small environments that demonstrate the advantages of our geometric LTL (GLTL) language and illustrate how it can be used to specify standard reinforcement-learning tasks straightforwardly.

1 Introduction

The thesis of this work is that (1) rewards are an excellent way of controlling the behavior of agents, but (2) rewards are difficult to use for specifying behaviors in an environment-independent way, therefore (3) we need intermediate representations between behavior specifications and reward functions.

The intermediate representation we propose is a novel variant of linear temporal logic that is modified to be probabilistic so as to better support reinforcement-learning tasks. Linear temporal logic has been used in the past to specify reward functions that depend on temporal sequences (Bacchus et al., 1996); here, we expand the role to provide a robust and consistent semantics that allows desired behaviors to be specified in an environment-independent way. Briefly, our approach involves the specification of tasks via temporal operators that have a constant probability of expiring on each step. As such, it bears a close relationship to the notion of discounting in standard Markov decision process (MDP) reward functions (Puterman, 1994).

At a philosophical level, we are viewing behavior specification as a kind of programming problem. That is, if we think of a Markov decision process (MDP) as an input, a

reward function as a program, and a policy as an output, then reinforcement learning can be viewed as a process of program interpretation. We would like the same program to work across all possible inputs.

1.1 Specifying behavior via reward functions

An MDP consists of a finite state space, action space, transition function, and reward function. Given an environment, an agent should behave in a way that maximizes cumulative discounted expected reward. The problems of learning and planning in such environments have been vigorously studied in the AI community for over 25 years (Watkins, 1989; Boutilier et al., 1999; Strehl et al., 2009). A reinforcement-learning (RL) agent needs to learn to maximize cumulative discounted expected reward starting with an incomplete model of the MDP itself.

For “programming” reinforcement-learning agents, the state of the art is to define a reward function and then for the learning agent to interact with the environment to discover ways to maximize its reward. Reward-based specifications have proven to be extremely valuable for optimal planning in complex, uncertain environments (Russell & Norvig, 1994). However, we can show that reward functions, as they are currently structured, are very difficult to work with as a way of reliably specifying tasks. The best use case for reward functions is when the utilities of all actions and outcomes can be expressed in a consistent unit, for example, time or money or energy. In reality, however, putting a meaningful dollar figure on scuffing a wall or dropping a clean fork is challenging. When informally adding negative rewards to undesirable outcomes, it is difficult to ensure a consistent semantics over which planning and reasoning can be carried out. Further, reward values often need to be changed if the environment itself changes—they are not environment independent. Therefore, to get a system to exhibit a desired behavior, it can be necessary to try different reward structures and carry out learning multiple times in the target environment, greatly undermining the purpose of autonomous learning in the first place.

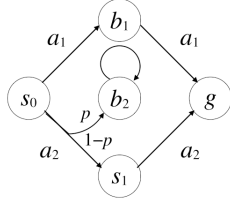


Figure 1: Action a_2 has probability $1 - p$ of transitioning to a non- b state and a probability of p of entering a self-loop in a b state. Action a_1 passes through a b state and then over to the goal.

Consider the simple example MDP in Figure 1. The agent is choosing between a_1 and a_2 in the initial state s_0 . Choosing a_1 causes the agent to pass through *bad* state b_1 for one step, then to continue on to the goal g . Action a_2 , however, results in a probabilistic transition to s_1 (with probability $1 - p$) or bad state b_2 (with slip probability p). From s_1 , the agent can continue on to the goal. If it reaches b_2 , it gets stuck there forever.

Let’s say our desired behavior is “maximize the probability of reaching g without hitting a bad state”. (A bad state could be something like colliding with a wall or bumping up against a table.) The probability of success of a_1 is zero and a_2 is $1 - p$. Thus, for any $0 \leq p < 1$, it is better to take action a_2 .

What reward function encourages this behavior? For concreteness, let’s assume a discount of $\gamma = 0.8$ and a reward of $+1$ for reaching the goal. We can assign bad states a value of $-r$. In the case where $p = 0.1$, setting $r > 0.16$ encourages the desired behavior.

Consider, though, what happens if the slip probability is $p = 0.3$. Now, there is no value of r for which a_2 is preferred to a_1 ¹. That is, it has become impossible to find a reward function that creates the correct incentives for the desired behavior to be optimal.

This example is perhaps a bit contrived, but we have observed the same phenomenon in large and natural state spaces as well. The reason for this result is that reward functions force us to express utility in terms of the discounted expected visit frequency of states. In this case, we are stuck trying to make a tradeoff between the certainty of encountering a bad state once and the possibility of encountering a bad state repeatedly. Since we are trying to maximize the probability of *zero* encounters with a bad state, the expected number of encounters is only useful for distinguishing zero from more than zero—the objective cannot be translated into a reward function when bad states

¹Actually, $r < -24/50$ works for this example, but that is tantamount to rewarding the agent for bumping into things—something bound to result in other problems.

are unavoidable.

1.2 Specifying behavior via LTL

An alternative to specifying tasks via reward functions is to use a formal specification like linear temporal logic or LTL (Manna & Pnueli, 1992; Baier & Katoen, 2008).

Linear temporal logic formulas are built up from a set of atomic propositions; the logic connectives: negation ($\neg$), disjunction ($\vee$), conjunction ($\wedge$) and material implication ($\rightarrow$); and the temporal modal operators: next ($\bigcirc$), always ($\Box$), eventually ($\Diamond$) and until ($\mathcal{U}$). A wide class of properties including safety ($\Box \neg b$), goal guarantee ($\Diamond g$), progress ($\Box \Diamond g$), response ($\Box(b \rightarrow \Diamond g)$), and stability ($\Diamond \Box g$), where b and g are atomic propositions, can be expressed as LTL formulas. More complicated specifications can be obtained from the composition of such simple formulas. For example, the specification of “repeatedly visit certain locations of interest in a given order while avoiding certain other unsafe or undesirable locations” can be obtained through proper composition of simpler safety and progress formulas (Manna & Pnueli, 1992; Baier & Katoen, 2008).

Returning to the example in Figure 1, the task is to avoid b states until g is reached: $\neg b \mathcal{U} g$. Given an LTL specification and an environment, an agent, for example, should adopt a behavior that maximizes the probability that the specification is satisfied. One advantage of this approach is its ability to specify tasks that cannot be expressed using simple reward functions (like the example MDP in Section 1.1). Indeed, in the context of reinforcement-learning problems, we have found it very natural to express standard MDP task specifications using LTL.

Standard MDP tasks can be expressed well using these temporal operators. For example:

- Goal-based tasks like mountain car (Moore, 1991): If p represents the attribute of being at the goal (the top of the hill, say), $\Diamond p$ corresponds to eventually reaching the goal.
- Avoidance-type tasks like cart pole (Barto et al., 1983): If q represents the attribute of being in the failure state (dropping the pole, say), $\Box \neg q$ corresponds to always avoiding the failure state.
- Sequence tasks like taxi (Dietterich, 2000): If p represents some task being completed (getting the passenger, say) and q represents another task being completed (delivering the passenger, say), $\Diamond(p \wedge \Diamond q)$ corresponds to eventually completing the first task, then, from there, eventually completing the second task.
- Stabilizing tasks like pendulum swing up (Atkeson, 1994): If p represents the property that needs to be stabilized (the pendulum being above the vertical, say),

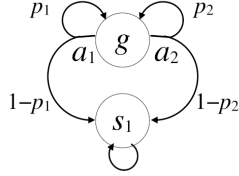


Figure 2: Action a_1 has probability p_1 of a self-loop and $1 - p_1$ of transitioning to a non- g state. Action a_2 has probability p_2 of a self-loop and $1 - p_2$ of transitioning to a non- g state. The policy that maximizes the probability of satisfaction of $\Box g$ is highly dependent on p_1 and p_2 if they are near one.

$\Diamond\Box p$ corresponds to eventually achieving and continually maintaining the desired property.

- Approach-avoid tasks like the 4×3 grid (Russell & Norvig, 1994): If p represents the attribute of being at the goal (the upper right corner the grid, say), and q represents the attribute of being at a bad state (the state below it, say), $\neg q \mathcal{U} p$ corresponds to avoiding the bad state en route to the goal.

On the other hand, there are barriers to straightforwardly adopting temporal logic-based languages in a reinforcement-learning setup. The most significant is that we can show that it is simply impossible to learn to satisfy classical LTL specifications in some cases. A key property for being able to learn near-optimal policies efficiently in the context of reward-based MDPs is what is known as the Simulation Lemma (Kearns & Singh, 1998). Informally, it says that, for any MDP and any $\epsilon > 0$, there exists an $\epsilon' > 0$ such that finding optimal policies in an ϵ' -close model of the real environment results in behavior that is ϵ -close to optimal in the real environment.

Unfortunately, tasks specified via LTL do not have this property. In particular, there is an MDP and an $\epsilon > 0$ such that no ϵ' -close approximation for $\epsilon' > 0$ is sufficient to produce a policy with ϵ -close satisfaction probability.

Consider the MDP in Figure 2. If we want to find a behavior that nearly maximizes the probability of satisfying the specification $\Box g$ (stay in the good state forever), we need accurate estimates of p_1 and p_2 . If $p_1 = p_2 = 1$ or $p_1 < 1$ and $p_2 < 1$, either policy is equally good. If $p_1 = 1$ and $p_2 < 1$, only action a_1 is near optimal. If $p_2 = 1$ and $p_1 < 1$, only action a_2 is near optimal. As there is no finite bound on the number of learning trials needed to distinguish $p_1 = 1$ from $p_1 < 1$, a near optimal behavior cannot be found in worst-case finite time. LTL expressions are simply too unforgiving to be used with any confidence in a learning setting.

In this work, we develop a hybrid approach for specifying behavior in reinforcement learning that combines the

strengths of both reward functions and temporal logic specifications.

2 Learning To Satisfy LTL

While provable guarantees of efficiency and optimality have been at the core of the literature on learning (Fiechter, 1994; Kearns & Singh, 2002; Brafman & Tennenholtz, 2002; Li et al., 2011), correctness with respect to complicated, high-level task specifications—during the learning itself or in the behavior resulting from the learning phase—has attracted limited attention (Abbeel & Ng, 2005).

2.1 Geometric linear temporal logic

We present a variant of LTL we call geometric linear temporal logic (GLTL) that builds on the logical and temporal operators in LTL while ensuring learnability. The idea of GLTL is roughly to restrict the period of validity of the temporal operators to bounded windows—similar to the bounded semantics of LTL (Manna & Pnueli, 1992). To this end, GLTL introduces operators of the form of $\Diamond_\mu b$ with the atomic proposition b , which is interpreted as “ b eventually holds within k time steps where k is a random variable following a geometric distribution with parameter μ .” Similar semantics stochastically restricting the window of validity for other temporal operators are also introduced.

This kind of geometric decay fits very nicely with MDPs for a few reasons. It can be viewed as a generalization of reward discounting, which is already present in many MDP models. It also avoids unnecessarily expanding the specification state space by only requiring extra states to represent events and not simply the passage of time.

Using $G1(\mu)$ to represent the geometric distribution with parameter μ , the temporal operators are:

- $\Diamond_\mu p$: p is achieved in the next k steps, $k \sim G1(\mu)$.
- $\Box_\mu q$: q holds for the next k steps, $k \sim G1(\mu)$.
- $q \mathcal{U}_\mu p$: q must hold at least until p becomes true, which itself must be achieved in the next k steps, $k \sim G1(\mu)$.

Returning to our earlier example from Figure 2, evaluating the probability of satisfaction for $\Box g$ requires infinite precision in the learned transition probabilities in the environment. Consider instead evaluating $\Box_\mu g$ in this environment. An encoding of the specification for this example is shown in Figure 3 (Third). We call it a *specification MDP*, as it specifies the task using states (derived from the formula), actions (representing conditions), and probabilities (capturing the stochasticity of operator expiration). This example says that, from the initial state q_0 , encountering any state where g is not true results in immediately failing the specification. In contrast, encountering any state

where g is true results in either continued evaluation (with probability μ) or success (with probability $1 - \mu$). Success represents the idea that the temporal window in which g must hold true has expired without g being violated.

Composing these two MDPs leads to the composite MDP in Figure 3 (Fourth). The true satisfaction probability for action a_i is $\frac{1-\mu}{(1-\mu+p_i)}$. Thus, if $\mu = .9$, the dependence of this value on ϵ is $\frac{.1}{.1+\epsilon}$, which is well behaved for all values of ϵ . The sensitivity of the computed satisfaction probability has a maximum $1/(1-\mu)^2$ dependence on the accuracy of the estimate of ϵ . Thus, GLTL is considerably more friendly to learning than is LTL.

Returning to the MDP example in Figure 1, we find that GLTL is also more expressive than rewards. The GLTL formula $\neg q \mathcal{U}_\mu p$ can be translated to a specification MDP. Essentially, the idea is that encountering a bad state (q) even once or running out of time results in specification failure. Maximizing the satisfaction of this GLTL formula results in taking action a_1 regardless of the value of p . That is, it is an environment-independent specification of the task.

The reason the GLTL formulation is able to succeed where standard rewards fail is that the GLTL formula results in an augmentation of the state space so that the reward function can depend on whether a bad state has yet been encountered. On the first encounter, a penalty can be issued. After the first encounter, no additional penalty is added. By composing the environment MDP with this bit of internal memory, the task can be expressed provably correctly and in an environment-independent way.

3 Related Work

Discounting has been used in previous temporal models. In quantitative temporal logic, it gives more weight to the satisfaction of a logic property in the near future than the far future. De Alfaro et al. (2003, 2004) augment computation tree logic (CTL) with discounting and develop fixpoint-based algorithms for checking such properties for probabilistic systems and games. Almagor et al. (2014) explicitly refine the “eventually” operator of LTL to a *discounting operator* such that the longer it takes to fulfill the task the smaller the value of satisfaction. Further, they show that discounted LTL is more expressive than discounted CTL. They use both discounted until and undiscounted until for expressing traditional eventually as well as its discounted version. However, algorithms for model checking and synthesis discounted LTL for probabilistic systems and games are yet to be developed.

LTL has been used extensively in robotics domains. Work on the trustworthiness of autonomous robots, automated verification and synthesis with provable correctness with respect to temporal logic-based specifications in motion, task, and mission planning have attracted considerable at-

Table 1: Operator precedence in specification MDP construction.

Precedence	Operator	# of Operands
1	not, $\neg$	1
2	μ -always, $\Box_\mu$	1
	μ -eventually, $\Diamond_\mu$	1
	μ -until, $\mathcal{U}_\mu$	2
3	and, $\wedge$	2
4	or, $\vee$	2

tention recently. The results include open-loop and reactive control of deterministic, stochastic or non-deterministic finite-state models as well as continuous state models through appropriate finite-state abstractions (Wongpiromsarn et al., 2012; Kress-Gazit et al., 2009; Liu et al., 2013; Wolff et al., 2012; Ding et al., 2011; Lahijanian et al., 2011; Kress-Gazit et al., 2011). While temporal logic had initially focused on reasoning about temporal and logical relations, its dialects with probabilistic modalities have been used increasingly for robotics applications (Baier & Katoen, 2008; De Alfaro, 1998; Kwiatkowska et al., 2002).

4 Generating Specification MDPs

Similar to LTL, GLTL formulas are built from a set of atomic propositions AP , Boolean operators $\wedge$ (conjunction), $\neg$ (negation) and temporal operator $\mathcal{U}_\mu$ (μ -until). Useful operators such as $\vee$ (disjunction), $\Diamond_\mu$ (μ -eventually) and $\Box_\mu$ (μ -always) can be derived from these basic operators.

GLTL formulas can be converted to the corresponding specification MDPs recursively, with the operator precedence listed in descending order in Table 1. Operators of the same precedence are read from right to left. For example, $\Box_{\mu_1} \Diamond_{\mu_2} \varphi = (\Box_{\mu_1} (\Diamond_{\mu_2} \varphi))$, $\varphi_1 \mathcal{U}_{\mu_1} \varphi_2 \mathcal{U}_{\mu_2} \varphi_3 = (\varphi_1 \mathcal{U}_{\mu_1} (\varphi_2 \mathcal{U}_{\mu_2} \varphi_3))$.

Assume $\varphi, \varphi_1, \varphi_2$ are GLTL formulas in the following discussion.

- b , where $b \in AP$ is an atomic proposition: A specification MDP $M_b = (\{s^{ini}, acc, rej\}, \{a\}, T, R)$ for b can be constructed such that, if p holds at s^{ini} , the transition (s^{ini}, a, acc) is taken with probability 1; otherwise, the transition (s^{ini}, a, rej) is taken with probability 1.
- $\neg\varphi$: A specification MDP $M_{\neg\varphi}$ can be constructed from a specification MDP M_φ by swapping the terminal states acc and rej .
- $\varphi_1 \wedge \varphi_2$: A specification MDP $M_{\varphi_1 \wedge \varphi_2} = (S, A, T, R)$ can be constructed from specification MDPs $M_{\varphi_1} = (S_1, A_1, T_1, R_1)$

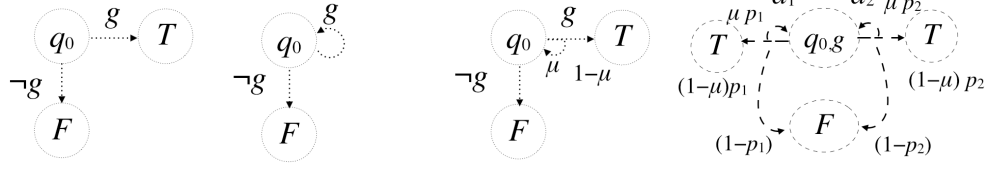


Figure 3: *First:* The specification MDP representation of the LTL formula b . *Second:* The specification MDP representation of the LTL formula $\Box b$. *Third:* The specification MDP representation of the GLTL formula $\Box_\mu b$. *Fourth:* The composition of the specification MDP representation of the GLTL formula $\Box_\mu b$ with the MDP from Figure 2.

and $M_{\varphi_2} = (S_2, A_2, T_2, R_2)$ such that (1) $S = (S_1 \setminus \{rej_1\}) \times (S_2 \setminus \{rej_2\}) \cup \{acc\}$, and the accepting state is $acc = (acc_1, acc_2)$; (2) $A = A_1 \times A_2$; (3) for all transitions (s_1, a_1, s'_1) of M_{φ_1} and (s_2, a_2, s'_2) of M_{φ_2} , if either $s'_1 = rej_1$ or $s'_2 = rej_2$, let $T((s_1, s_2), (a_1, a_2), rej) = T_1(s_1, a_1, s'_1)T_2(s_2, a_2, s'_2)$; otherwise, $T((s_1, s_2), (a_1, a_2), (s'_1, s'_2)) = T_1(s_1, a_1, s'_1)T_2(s_2, a_2, s'_2)$.

- $\varphi_1 \vee \varphi_2 = \neg(\neg\varphi_1 \wedge \neg\varphi_2)$.
- $\varphi_1 \mathcal{U}_\mu \varphi_2$: The operator μ -until has two operands, φ_1 and φ_2 , which generate specification MDPs $M_{\varphi_1} = (S_1, A_1, T_1, R_1)$ and $M_{\varphi_2} = (S_2, A_2, T_2, R_2)$. The new specification MDP $M_{\varphi_1 \mathcal{U}_\mu \varphi_2} = (S, A, T, R)$ is constructed from M_{φ_1} and M_{φ_2} : $S = (S_1 \setminus \{acc_1, rej_1\}) \times (S_2 \setminus \{acc_2, rej_2\}) \cup \{acc, rej\}$, where acc and rej are the accepting and rejecting state, respectively, and $s^{ini} = (s_1^{ini}, s_2^{ini}) \in S$ is the initial state; $A = A_1 \times A_2$; for all $s = (s_1, s_2) \in S \setminus \{acc, rej\}$, $a = (a_1, a_2) \in A$, $s'_1 \in S_1$ and $s'_2 \in S_2$, if $T_1(s_1, a_1, s'_1) > 0$ and $T_2(s_2, a_2, s'_2) > 0$, a transition (s, a, s') is added to $M_{\varphi_1 \mathcal{U}_\mu \varphi_2}$ with probability $T(s, a, s')$ as specified in Table 2.

Here are some intuitions behind the construction of T . The formula $\varphi_1 \mathcal{U}_\mu \varphi_2$ means that, within some stochastically decided time period k , we would like to successfully implement task φ_2 in at most k steps without ever failing in task φ_1 . If we observe a success in M_{φ_2} (that is, the specification reaches acc_2) before φ_1 fails (that is the specification reaches rej_1), $M_{\varphi_1 \mathcal{U}_\mu \varphi_2}$ goes to state acc for sure; if we observe a failure in M_{φ_1} (that, the specification reaches rej_1) before succeeding in M_{φ_2} (that is, the specification reaches acc_2), $M_{\varphi_1 \mathcal{U}_\mu \varphi_2}$ goes to state rej for sure. In all other cases, $M_{\varphi_1 \mathcal{U}_\mu \varphi_2}$ primarily keeps track of the transitions in M_{φ_1} and M_{φ_2} , with a tiny probability of failing immediately, which corresponds to the operator expiring.

- $\Diamond_\mu \varphi_2$: As in the semantics of LTL, μ -eventually $\Diamond_\mu \varphi_2 = \text{True} \mathcal{U}_\mu \varphi_2$. Hence, given a specification

Table 2: Transition (s, a, s') in $M_{\varphi_1 \mathcal{U}_\mu \varphi_2}$ constructed from a transition (s_1, a_1, s'_1) in M_{φ_1} and a transition (s_2, a_2, s'_2) in M_{φ_2} . Here, $p(s'|s'_1, s'_2) = \frac{T(s, a, s')}{T_1(s_1, a_1, s'_1)T_2(s_2, a_2, s'_2)}$. That is, to get the transition probability, multiple the p column by the corresponding T_1 and T_2 transition probabilities.

s'_1	s'_2	s'	$p(s' s'_1, s'_2)$
acc_1	acc_2	acc	1
acc_1	rej_2	(s'_1, s'_2)	$1 - \mu$
		rej	μ
acc_1	$S_2 \setminus \{acc_2, rej_2\}$	(s_1^{ini}, s'_2)	$1 - \mu$
		rej	μ
rej_1	acc_2	acc	1
rej_1	$S_2 \setminus \{acc_2\}$	rej	1
$S_1 \setminus \{acc_1, rej_1\}$	acc_2	acc	1
$S_1 \setminus \{acc_1, rej_1\}$	rej_2	(s'_1, s_2^{ini})	$1 - \mu$
		rej	μ
$S_1 \setminus \{acc_1, rej_1\}$	$S_2 \setminus \{acc_2, rej_2\}$	(s'_1, s'_2)	$1 - \mu$
		rej	μ

MDP $M_{\varphi_2} = (S_2, A_2, T_2, R_2)$ for φ_2 , we can construct a specification MDP $M_{\Diamond_\mu \varphi_2} = (S, A, T, R)$ for $\Diamond_\mu \varphi_2$: $S = S_2$, $s^{ini} = s_2^{ini}$, $acc = acc_2$, $rej = rej_2$; $A = A_2$; transitions of $M_{\Diamond_\mu \varphi_2}$ are modified from those of M_{φ_2} as in Table 3. Informally, $\Diamond_\mu \varphi_2$ is satisfied if we succeed in task φ_2 within the stochastic observation time period.

- $\Box_\mu \varphi_2$: μ -always φ_2 is equivalent to $\neg\Diamond_\mu \neg\varphi_2 = \neg(\Diamond_\mu (\neg\varphi_2))$. In other words, $\Box_\mu \varphi_2$ is satisfied if we did not witness a failure of φ_2 within the stochastic observation time period. The transitions of a specification MDP $M_{\Box_\mu \varphi_2}$ can be constructed from Table 3, or directly from Table 4.

Using the transitions as described, a given GLTL formula can be converted into a specification MDP. To satisfy the specification in a given environment, a joint MDP is created as follows:

1. Take the cross product of the MDP representing the environment and the specification MDP.

Table 3: Transition (s, a, s') in $M_{\diamond_\mu \varphi_2}$ constructed from a transition (s_2, a_2, s'_2) in M_{φ_2} . As above, $p(s'|s'_2) = \frac{T(s, a, s')}{T_2(s_2, a_2, s'_2)}$.

s'_2	s'	$p(s' s'_2)$
acc_2	acc_2	1
rej_2	s_2^{ini}	$1 - \mu$
	rej_2	μ
$S_2 \setminus \{acc_2, rej_2\}$	s'_2	$1 - \mu$
	rej_2	μ

Table 4: Transition (s, a, s') in $M_{\square_\mu \varphi_2}$ constructed from a transition (s_2, a_2, s'_2) in M_{φ_2} . As above, $p(s'|s'_2) = \frac{T(s, a, s')}{T_2(s_2, a_2, s'_2)}$.

s'_2	s'	$p(s' s'_2)$
acc_2	s_2^{ini}	$1 - \mu$
	acc_2	μ
rej_2	rej_2	1
$S_2 \setminus \{acc_2, rej_2\}$	s'_2	$1 - \mu$
	acc_2	μ

- Any state that corresponds to an accepting or rejecting state of the specification MDP becomes a sink state. However, the accepting states also include a reward of +1.
- The resulting MDP is solved to create a policy.

The resulting policy is one that maximizes the probability of satisfying the given formula where the random events are both the transitions in the environment and the stochastic transitions in the specification MDP. Such policies tend to prefer satisfying formulas quickly, as that increases the chance of successful completion before operators expire.

5 Example Domain

Consider the following formula:

$$(\neg \text{blue} \mathcal{U}_\mu \text{red}) \wedge (\diamond_\mu (\text{red} \wedge \diamond_\mu \text{green})).$$

It specifies a task of reaching a red state without encountering a blue state and, once a red state is reached, going to a green state.

Figure 5 illustrates a grid world environment in which this task can be carried out. It consists of different colored grid cells. The agent can move to one of the four adjacent cells to its current position with a north, south, east, or west action. However, selecting an action for one direction has a 0.02 probability of moving in the one of the three other directions. This stochastic movement causes the agent to keep

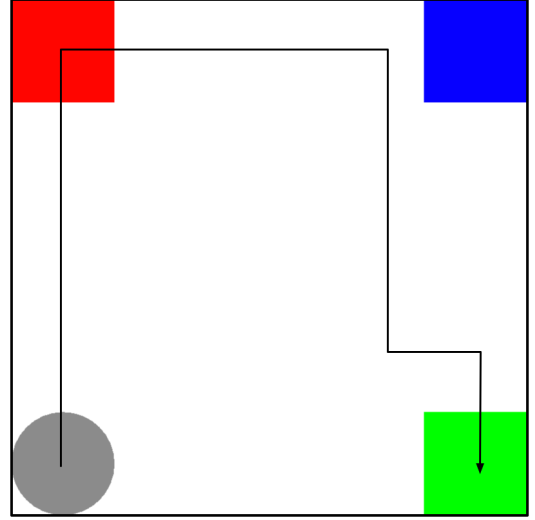


Figure 4: The optimal path in a grid world.

its distance from dangerous grid cells that could result in task failure, whenever possible. The solid line in the figure traces the path of the optimal policy of following this specification in the grid. As can be seen, the agent moves to red and then green. Note that this behavior can be very difficult to encode in a standard reward function as both green and red need to be given positive reward and therefore either would be a sensible place for the agent to stop.

Figure 5 illustrates a grid world environment in which the blue cells create a partial barrier between the red and green cells. As a result of the “until” in the specification, the agent goes around the blue wall to get to the red cell. However, since the prohibition against blue cells is lifted once the red cell is reached, it goes directly through the barrier to reach green.

These 25-state environments become 98-state MDPs when combined with the specification MDP.

6 Conclusion

In contrast to standard MDP reward functions, we have provided an environment-independent specification for tasks. We have shown that this specification language can capture standard tasks used in the MDP community and that it can be automatically incorporated into an environment MDP to create a fixed MDP to solve. Maximizing reward in this resulting MDP maximizes the probability of satisfying the task specification.

Future work includes inverse reinforcement learning of task specifications and techniques for accelerating planning.

- Liu, Jun, Ozay, Necmiye, Topcu, Ufuk, and Murray, Richard M. Synthesis of reactive switching protocols from temporal logic specifications. *IEEE Transactions on Automatic Control*, 58(7):1771–1785, 2013.
- Manna, Zohar and Pnueli, Amir. *The Temporal Logic of Reactive & Concurrent Sys.* . Springer, 1992.
- Moore, Andrew W. Variable resolution dynamic programming: Efficiently learning action maps in multivariate real-valued spaces. In *Proc. Eighth International Machine Learning Workshop*, 1991.
- Puterman, Martin L. *Markov Decision Processes—Discrete Stochastic Dynamic Programming*. John Wiley & Sons, Inc., New York, NY, 1994.
- Russell, Stuart J. and Norvig, Peter. *Artificial Intelligence: A Modern Approach*. Prentice-Hall, Englewood Cliffs, NJ, 1994. ISBN 0-13-103805-2.
- Strehl, Alexander L., Li, Lihong, and Littman, Michael L. Reinforcement learning in finite MDPs: PAC analysis. *Journal of Machine Learning Research*, 10:2413–2444, 2009.
- Watkins, Christopher J. C. H. *Learning from Delayed Rewards*. PhD thesis, King’s College, Cambridge, UK, 1989.
- Wolff, Eric M., Topcu, Ufuk, and Murray, Richard M. Robust control of uncertain markov decision processes with temporal logic specifications. In *Proc. of the IEEE Conference on Decision and Control*, 2012.
- Wongpiromsarn, T., Topcu, U., and Murray, R.M. Receding horizon temporal logic planning. *IEEE T. on Automatic Control*, 57:2817–2830, 2012.