

Synthesis from Satisficing and Temporal Goals

Suguman Bansal,¹ Lydia Kavraki,² Moshe Y. Vardi,² Andrew Wells^{2,3}

¹ University of Pennsylvania

² Rice University

³ Tesla

Abstract

Reactive synthesis from high-level specifications that combine *hard* constraints expressed in Linear Temporal Logic (LTL) with *soft* constraints expressed by discounted-sum (DS) rewards has applications in planning and reinforcement learning. An existing approach combines techniques from LTL synthesis with optimization for the DS rewards but has failed to yield a sound algorithm. An alternative approach combining LTL synthesis with satisficing DS rewards (rewards that achieve a threshold) is sound and complete for integer discount factors, but, in practice, a fractional discount factor is desired. This work extends the existing satisficing approach, presenting the first sound algorithm for synthesis from LTL and DS rewards with fractional discount factors. The utility of our algorithm is demonstrated on robotic planning domains.

1 Introduction

Reactive synthesis is the automated construction, from a high-level description of its desired behavior, of a reactive system that continuously interacts with an uncontrollable external environment (Church 1957).

Recent applications of reactive synthesis have emerged in AI for planning and robotics tasks (Camacho, Bienvenu, and McIlraith 2019; He et al. 2019; Kress-Gazit, Lahijanian, and Raman 2018). These applications can be formulated as a deterministic turn-based interaction between a controllable system player and an uncontrollable environment player. Given a specification, the synthesis task is to generate a system strategy such that all resulting interactions with the environment satisfy the specification. A large focus in this line of work has been on synthesis from Linear Temporal Logic (LTL) specifications (Pnueli 1977; Pnueli and Rosner 1989).

Yet, several desired specifications either cannot be expressed using LTL or doing is cumbersome. Examples include specifications about the quantitative properties of systems, such as rewards, costs, degrees of satisfaction, and so on. In fact, the combination of LTL with quantitative properties is used to express more nuanced and complex specifications (see Figure 1). Subsequently, synthesis algorithms from combination specifications have followed (Ding et al. 2014; He et al. 2017; Lahijanian et al. 2015).

Copyright © 2022, Association for the Advancement of Artificial Intelligence (www.aaai.org). All rights reserved.

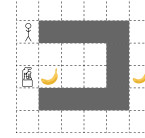


Figure 1: Example scenario: The (controlled) robot must retrieve objects and avoid the (uncontrolled) human in a grocery store. The robot’s hard constraint is to retrieve objects from its grocery list without colliding with the walls (in grey) or human. Its soft constraint is to socially (Manhattan) distance itself from the human. A fractional discount factor makes the robot less “greedy.”

This work investigates the problem of reactive synthesis from specifications that combine *hard qualitative constraints* expressed by LTL with *soft quantitative constraints* expressed by *discounted-sum rewards*. Discounted-sum rewards are well-suited for infinite-horizon executions because the discounted-sum is guaranteed to converge on infinite-sequence of costs whereas other aggregation functions such as limit-average may not. Discounted-sum encodes diminishing returns. As a result, the combination of LTL with discounted-sum rewards frequently appears in the automated construction of systems using planning and reinforcement learning (Bozkurt et al. 2020; Camacho et al. 2017, 2019; Hasanbeig et al. 2019; Kalagarla, Jain, and Nuzzo 2021; Kwiatkowska, Parker, and Wiltsche 2017). Note, however, these works only deal with a single player case (controllable system agent) while reactive synthesis also assumes the presence of an uncontrollable environment.

Broadly speaking, there are two approaches to reactive synthesis from LTL and discounted-sum rewards. The first approach is based on *optimization* of the discounted-sum reward. A strategy that optimizes the discounted-sum reward alone is guaranteed to exist in deterministic, turn-based settings (Shapley 1953). This existence guarantee, however, is lost upon combination with LTL constraints. For example, consider a two-state game where state s_0 gives negative reward and state s_1 positive. Each state can transition to all other states. Our LTL objective is (*Globally Eventually s_0*). Clearly, there exists no strategy that simultaneously maximizes the discounted-sum reward and satisfies the LTL objective (Chatterjee et al. 2017). To

this end, an alternate synthesis task is to compute an optimal strategy from those that satisfy the LTL constraint (Wen, Ehlers, and Topcu 2015). Unfortunately, even here existing synthesis algorithms may generate a sub-optimal strategy. Overall, synthesis algorithms from LTL and discounted-sum rewards that optimize the discounted-sum reward, in one way or another, have hitherto failed to provide guarantees of correctness or completeness.

The second approach to synthesis from LTL and discounted-sum rewards is based on *satisficing* the discounted-sum reward. A strategy is satisficing with respect to a given threshold value $v \in \mathbb{Q}$ if it guarantees the discounted-sum reward of all executions will exceed v . The synthesis task, therefore, is to compute a strategy that satisfies the LTL specification and is satisficing w.r.t. the threshold value. The advantage of this approach is that when the discount factor is an integer, an existing synthesis algorithm is both sound and complete (Bansal, Chatterjee, and Vardi 2021). The method builds on novel automata-based technique for quantitative reasoning called *comparator automata* (Bansal, Chaudhuri, and Vardi 2018a,b). The central result of comparator automata is that for integer discount factors, examining whether the discounted-sum of an execution exceeds a given threshold reduces to determining the membership of the execution in an (Büchi) automaton. Thus, satisficing goals are precisely captured by a comparator. This insight allows for elegant combination of satisficing and temporal goals since both are automata-based. The disadvantage of this method is that it cannot be applied with non-integer discount factors since the comparator for non-integer discount factors are not represented by automata. This is a severe limitation because in practice the discount factor is taken to be a fractional value between 1 and 2 in order to reason over a long-horizon (Sutton and Barto 2018). Consider Fig. 1. If an integer discount factor greater than 1 is used, the robot will be “greedy” and obtain the immediate reward at the cost of becoming “trapped” by the human. The fractional discount factor is necessary so the robot recognizes the longer-term benefits to avoid becoming “trapped.”

The central contribution of this work is a theoretically sound algorithm for reactive synthesis from LTL and satisficing discounted-sum goals for the case when the discount factor ranges between 1 and 2 (specifically of the form $1 + 2^{-k}$ for positive integer values of k). To the best of our knowledge, this is the first synthesis algorithm from LTL and discount-sum rewards that offers theoretical guarantees of correctness and is practically applicable.

Our solution is also based on comparator automata. We bypass the issue with fractional discount factors by introducing approximations into the comparator framework. We show that comparators for approximations of discounted-sum with fractional discount factors can be represented by Büchi automata. In brief, we show that for fractional discount factors, examining whether the discounted-sum of an execution *approximately exceeds* a threshold value can be determined by membership of the execution in a Büchi automaton. This combined with synthesis techniques for LTL gives rise to a purely automata-based algorithm for LTL and discounted-sum rewards, and thus preserves soundness.

Due to the use of approximation, our algorithm is no longer complete. To this end, we evaluate the practical utility of our algorithm on case studies from robotics planning. Our evaluation demonstrates that our sound but incomplete procedure succeeds in efficiently constructing high-quality strategies in complex domains from nuanced constraints.

2 Preliminaries

2.1 Automata and Formal Specifications

Büchi Automata and Co-Safety Automata. A *Büchi automaton* is a tuple $\mathcal{A} = (S, \Sigma, \delta, s_{\mathcal{I}}, \mathcal{F})$, where S is a finite set of *states*, Σ is a finite *input alphabet*, $\delta \subseteq (S \times \Sigma \times S)$ is the *transition relation*, state $s_{\mathcal{I}} \in S$ is the *initial state*, and $\mathcal{F} \subseteq S$ is the set of *accepting states*. A Büchi automaton is *deterministic* if for all states s and inputs a , $|\{s' \mid (s, a, s') \in \delta \text{ for some } s'\}| \leq 1$. For a word $w = w_0 w_1 \dots \in \Sigma^\omega$, a *run* ρ of w is a sequence of states $s_0 s_1 \dots$ s.t. $s_0 = s_{\mathcal{I}}$, and $\tau_i = (s_i, w_i, s_{i+1}) \in \delta$ for all i . Let $\text{inf}(\rho)$ denote the set of states that occur infinitely often in run ρ . A run ρ is an *accepting run* if $\text{inf}(\rho) \cap \mathcal{F} \neq \emptyset$. A word w is an *accepting word* if it has an accepting run. Büchi automata are closed under set-theoretic union, intersection, and complementation (Thomas, Wilke et al. 2002).

A *co-safety automata* is a deterministic Büchi automata with a single accepting state. Additionally, the accepting state is a sink state (Kupferman and Vardi 1999).

Comparator Automata. Given an aggregate function $f : \mathbb{Z}^\omega \rightarrow \mathbb{R}$, equality or inequality relation $R \in \{<, >, \leq, \geq, =, \neq\}$, and a threshold value $v \in \mathbb{Q}$, the *comparator automaton for f with upper bound μ , relation R , and threshold $v \in \mathbb{Q}$* is an automaton that accepts an infinite word A over the alphabet $\Sigma = \{-\mu, -\mu + 1, \dots, \mu\}$ iff $f(A) R v$ holds (Bansal, Chaudhuri, and Vardi 2018b,c).

The discounted-sum of an infinite-length weight-sequence $W = w_0 w_1 \dots$ with discount factor $d > 1$ is given by $DS(W, d) = \sum_{i=0}^{\infty} \frac{w_i}{d^i}$. The comparator automata for the discounted-sum has been shown to be a safety or co-safety automata when the discount factor $d > 1$ is an integer, for all values of R , μ and v . It is further known to not form a Büchi automata for non-integer discount factors $d > 1$, for all values of R , μ and v (Bansal and Vardi 2019; Bansal, Chatterjee, and Vardi 2021).

Linear Temporal Logic. Linear Temporal Logic (LTL) extends propositional logic with infinite-horizon temporal operators. The syntax of LTL is defined as $\varphi := a \in \mathcal{AP} \mid \neg\varphi \mid \varphi \wedge \varphi \mid \varphi \vee \varphi \mid X\varphi \mid \varphi U \varphi \mid F\varphi \mid G\varphi$. Here X (Next), U (Until), F (Eventually), G (Always) are temporal operators. The semantics of LTL can be found in (Pnueli 1977).

2.2 Two-Player Graph Games

Reachability Games. A reachability game $G = (V, v_{\text{init}}, E, \mathcal{F})$ consists of a directed graph (V, E) , initial state v_{init} , and non-empty set of accepting states $\mathcal{F} \subseteq V$. The set V is partitioned into V_0 and V_1 . For convenience, we assume every state has at least one successor. A game is played between two players P_0 and P_1 .

A *play* in the game is created by the players moving a token along the edges as follows: at the beginning, the token is at the initial state. If the token's current position v belongs to V_i , then P_i chooses the next position from the successors of v . Formally, a play $\rho = v_0 v_1 v_2 \dots$ is an infinite sequence of states such that $v_0 = v_{\text{init}}$ and $(v_k, v_{k+1}) \in E$ for all $k \geq 0$. A play is *winning for player P_1* in the game if it visits an accepting state, and *winning for player P_0* otherwise.

A *strategy* for a player is a recipe that guides the player on which state to go next to based on the history of a play. A *strategy is winning for a player P_i* if for all strategies of the opponent player P_{1-i} , all resulting plays are winning for P_i . To solve a graph game is to determine whether there exists a winning strategy for player P_1 . Reachability games are solved in $\mathcal{O}(|V| + |E|)$ (Thomas, Wilke et al. 2002).

Quantitative Graph Games. A quantitative graph game (quantitative game, in short) is given by $G = (V = V_0 \uplus V_1, v_{\text{init}}, E, \gamma, \mathcal{L}, d)$ where $V, V_0, V_1, v_{\text{init}}, E$, plays, and strategies are defined as earlier. Each edge is associated with a *cost* determined by the cost function $\gamma : E \rightarrow \mathbb{Z}$, and $d > 1$ is the *discount factor*. The cost-sequence of a play ρ is the sequence $w_0 w_1 w_2 \dots$ where $w_k = \gamma((v_k, v_{k+1}))$ for all $i \geq 0$. The cost of play ρ is the discounted-sum of its cost sequence with discount factor $d > 1$. A labelling function $\mathcal{L} : V \rightarrow 2^{\mathcal{AP}}$ maps states to propositions from the set $\mathcal{AP}$. The *label sequence* of a play ρ is given by $\mathcal{L}(v_0)\mathcal{L}(v_1) \dots$.

3 Problem Formulation and Overview

The two players, the controllable system and uncontrollable environment, interact in a domain described by a quantitative game G . The specification of the system player is a combination of hard and soft constraints.

The hard constraint is given as by an LTL formula φ . A play in G satisfies formula φ if its labelled sequence satisfies the formula. We say, a strategy for the system player satisfies a formula φ if it guarantees that all resulting plays will satisfy the formula. We call such a strategy φ -*satisfying*.

The soft constraints are given by satisficing goals. W.l.o.g, the system and environment players maximize and minimize the cost of plays, respectively. Given a threshold value $v \in \mathbb{Q}$, a play is v -satisficing for the system (maximizing) player if its cost is greater than or equal to v . Conversely, a play is v -satisficing for the environment (minimizing) player if its cost is less than v . A strategy is v -satisficing for a player if it guarantees all resulting plays are v -satisficing for the player.

We are interested in solving the following problem:

Problem (Reactive Synthesis from Satisficing and Temporal Goals). *Given a quantitative game G , a threshold value $v \in \mathbb{Q}$, and an LTL formula φ , the problem of reactive synthesis from satisficing and temporal goals is to compute a strategy for the system player that is φ -satisfying and v -satisficing for the player, if such a strategy exists.*

The problem is solved for integer discount factors (Bansal, Chatterjee, and Vardi 2021).

Algorithm Overview. In this paper, we extend to fractional discount factors $1 < d < 2$, yielding practical applications of the synthesis problem. In particular, we solve the

problem for $d = 1 + 2^{-k}$ where $k > 0$ is an integer. Since the comparator for discounted-sum with fractional discount factors are not representable by Büchi automata, we construct a comparator automata for lower approximations of discounted-sum. This comparator soundly captures the criteria for v -satisficing for system player. In particular, if the comparator accepts the weight sequence of a play, then the play must be v -satisficing for the player. Therefore, just like LTL goals, the satisficing goal is also soundly captured by an automaton. Thus, we can reduce the synthesis problem to parity games via appropriate synchronized product constructions of both the automata-based goals.

The comparator construction for approximation of discounted-sum is presented in Section 4 and the reduction to games on graphs is presented in Section 5.

4 Comparator Construction

This section develops the key machinery required to design our theoretically sound algorithm for synthesis from temporal and satisficing goals with fractional discount factors. We construct comparator automata for a lower approximation of discounted-sum for fractional discount factors of the form $d = 1 + 2^{-k}$ where $k > 0$ is an integer. We show these comparators are represented by co-safety automata.

This section is divided in two parts. Section 4.1 defines an aggregate function that approximates the discounted-sum. Section 4.2 constructs a comparator for this function.

Unless stated otherwise, we assume the approximation factor is of the form $\varepsilon = 2^{-p}$ where $p > 0$ is an integer. Please refer to the Appendix for missing proofs and details.

4.1 Approximation of Discounted-Sum

Given parameters $k, p > 0$ of the discount factor and the approximation factor, respectively, let $\text{roundLow}(x, k, p)$ be the largest integer multiple of $2^{-(p+k)}$ that is less than or equal to x , where $x \in \mathbb{R}$. Let $W[\dots n]$ denote the n -length prefix of a weight-sequence W .

Then, the *lower approximation of discounted-sum* of an infinite-length weight-sequence W with discount factor $d > 1$ and approximation factor $\varepsilon > 0$ is defined as

$$\text{DSLow}(W, k, p) = \lim_{n \rightarrow \infty} \frac{\text{gapLow}(W[\dots n], k, p)}{d^{n-1}}$$

where the *lower gap value* of a finite-length weight-sequence U is defined as

$$\text{gapLow}(U, k, p) = \begin{cases} 0, & \text{if } |U| = 0 \\ \text{roundLow}(d \cdot \text{gapLow}(V, k, p) + v, & k, p), & \text{if } U = V \cdot v \end{cases}$$

Finally, the definition of DSLow is completed by proving DSLow approximates the discounted-sum of sequences within an additive factor of $d \cdot \varepsilon$:

Theorem 1. *Let $d = 1 + 2^{-k}$ be the discount factor and $\varepsilon = 2^{-p}$ be the approximation factor, for rational parameters $p, k > 0$. Let W be an infinite-length weight sequence. Then, $0 \leq \text{DS}(W, d) - \text{DSLow}(W, k, p) < d \cdot \varepsilon$.*

Proof Sketch. While the definition of the lower approximation of discounted-sum may look notationally dense, it is inspired by an alternate definition of discounted-sum:

$$DS(W, d) = \lim_{n \rightarrow \infty} \frac{\text{gap}(W[\dots n], d)}{d^{n-1}}$$

where $\text{gap}(U, d) = 0$ if $|U| = 0$ and $\text{gap}(U, d) = d \cdot \text{gap}(V, d) + v$ if $U = V \cdot v$.

Intuitively, the lower gap value approximates gap. Subsequently, DS_{Low} approximates the discounted-sum. $\square$

4.2 Comparator for Approximation of DS

This section presents the construction of a comparator for lower approximation of discounted-sum defined above.

Definition 1 (Comparator automata for lower approximation of DS). *Let $\mu > 0$ be an integer bound, and $k, p > 0$ be integers. The comparator automata for lower approximation of discounted sum with discount factor $d = 1 + 2^{-k}$, approximation factor $\varepsilon = 2^{-p}$, upper bound μ , threshold value $v \in \mathbb{Q}$, and inequality relation $R \in \{\leq, \geq\}$ is an automaton over infinite weight sequences W over the alphabet $\Sigma = \{-\mu, \dots, \mu\}$ that accepts W iff $\text{DS}_{\text{Low}}(W, k, p) R v$.*

Construction Sketch. We sketch the construction of the comparator for lower approximation of discounted-sum. For sake of exposition, we begin the construction for threshold value $v = 0$. W.l.o.g., we present for the relation $\geq$. Notations $\mu, d = 1 + 2^{-k}, \varepsilon = 2^{-p}$, and W are from Definition 1.

The lower gap value of prefixes of a weight-sequence can be used as a proxy for acceptance of a weight-sequence in the comparator for the following two observations:

1. $\text{DS}_{\text{Low}}(W, k, p) \geq 0$ for an infinite-length weight sequence W iff there exists a finite prefix A of W such that $\text{gap}_{\text{Low}}(A, k, p) \geq \mu \cdot 2^k + 2^{-p}$. Let us denote $\mu \cdot 2^k + 2^{-p}$ by upperLimit.
2. $\text{DS}_{\text{Low}}(W, k, p)$ cannot be greater than or equal to 0 iff there exists a finite prefix A of W such that $\text{gap}_{\text{Low}}(A, k, p) \leq -\mu \cdot 2^k$. Let us denote $-\mu \cdot 2^k$ by lowerLimit.

So, the core idea behind our construction is two fold: (a) use states of the comparator to record the lower gap value of finite-length prefixes, and (b) assign transitions so that the final state of finite-prefix corresponds to its lower gap value.

To this end, we set the initial state to 0 as the lower gap value of the 0-length prefix is 0. The transition relation mimics the inductive definition of lower gap value, i.e. there is a transition from a state s on the alphabet (weight) a to state t if $t = \text{round}_{\text{Low}}(d \cdot s + a, k, p)$. These ensure that the lower gap value of a finite-state word (finite-length weight-sequence) is detected from the final state in its run. Clearly, the transition relation is deterministic.

The final piece of the construction is to restrict the automata to finitely many states and to determine its accepting states. Note that due to the enumerated observations it is sufficient to track the lower gap value for only as long as it lies between lowerLimit and upperLimit. Observe that there are only finitely many such values of interest since lower gap value is always an integer multiple of $2^{-(p+k)}$. Thus, we

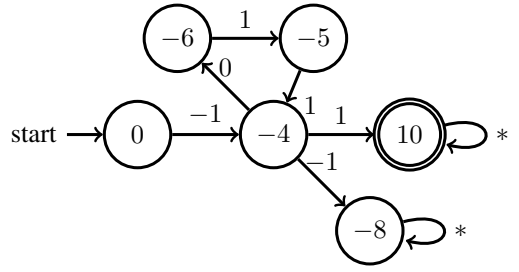


Figure 2: Snippet of comparator for $d = 1.5, \varepsilon = 0.5, \mu = 1, v = 0$, and $\geq$. Labels on states have been simplified. A state labelled by s refers to a lower gap value of $s \cdot 2^{-(p+k)}$

have obtained a finite number of states. By the first observation, state upperLimit is made an accepting sink since every weight-sequence that visits upperLimit must be accepted by the comparator. Similarly, by the second observation, the state lowerLimit is made a non-accepting sink. This completes the construction for threshold value $v = 0$.

To extend the construction to a non-zero threshold value $v \in \mathbb{Q}$, let V be a lasso weight-sequence s.t. $\text{DS}(V, d) = v$, the comparator incorporates V into its construction. Specifically, we construct a comparator that accepts W iff $\text{DS}_{\text{Low}}(W - V, k, p) \geq 0$. So, when W is accepted then $\text{DS}(W, d) \geq v$, otherwise $\text{DS}(W, d) \leq v + d \cdot \varepsilon$.

As an example, Figure 2 illustrates a snippet of the comparator with discount factor $d = 1 + 2^{-1}$, approximation factor $\varepsilon = 2^{-1}$, upper bound $\mu = 1$, threshold value $v = 0$, and relation $\geq$. As one can see, weight sequence $A = -1, 0, 1^\omega$ with discounted-sum $\frac{1}{3}$ is accepting and weight sequence $B = -1, -1, 1^\omega$ with discounted-sum $\frac{-1}{3}$ is non-accepting.

Theorem 2. *The comparator automata for lower approximation of discounted sum with discount factor $d = 1 + 2^{-k}$, approximation factor $\varepsilon = 2^{-p}$, upper bound μ , threshold 0, and inequality relation $R \in \{\leq, \geq\}$ is a co-safety automata with $\mathcal{O}(\frac{\mu}{(d-1)^2 \cdot \varepsilon})$ states, where $k, p > 0$ are integers.*

Therefore, $\text{DS}(W, d) \geq 0$ if a weight-sequence W is accepted by the comparator constructed above, and $\text{DS}(W, d) < d \cdot \varepsilon$ otherwise (Theorem 1-2).

5 Reactive Synthesis from Satisficing and Temporal Goals

This section presents the central contribution of this work. We present a theoretically sound algorithm for reactive synthesis from LTL and satisficing discounted-sum goals for the case when the discount factor ranges between 1 and 2, referred to as fractional discount factors hereon.

Our algorithm utilizes the comparator automata for the lower approximation of discounted-sum for fractional discount factors constructed in Section 4. For ease of exposition, we present our solution in two parts. First, we present an algorithm for reactive synthesis from satisficing goals only in Section 5.1. Next, we extend this algorithm to solve our original problem in Section 5.2.

5.1 Satisficing Goals

We describe an automata-based solution for reactive synthesis from satisficing goals with fractional discount factor. Our solution reduces to reachability games using the comparator for lower approximation of discounted sum.

The key idea behind our solution is that the said comparator can be treated as a sufficient criteria to compute a satisficing strategy for the system player. We explain this further. Take a comparator for the lower approximation for discounted-sum with discount factor d , approximation factor ε , threshold $v \in \mathbb{Q}$, and relation $\geq$. Then, a play in the quantitative game is v -satisficing for the system player if the comparator accepts the cost sequence of the play. This can be derived directly from Theorem 1-2. So, a strategy is v -satisficing for the system player if it is winning with respect to the comparator. To this end, we construct a *synchronized product* of the quantitative game with the comparator. The resulting product game is a reachability game since the comparator is represented by a co-safety automata. Formally,

Theorem 3. *Let G be a quantitative game with discount factor $d = 1 + 2^{-k}$, for integer $k > 0$. Let $v \in \mathbb{Q}$ be the threshold value and $\varepsilon = 2^{-p}$ be the approximation factor. There exists a reachability game GA such that*

- *If the system has a winning strategy in GA , then the system has a v -satisficing strategy in G .*
- *If the environment has a winning strategy in GA , then the environment has a $v + d \cdot \varepsilon$ -satisficing strategy in G .*

Proof. The product game synchronizes costs along edges in the quantitative game with the alphabet of the co-safety comparator. Let $G = (V = V_0 \uplus V_1, v_{\text{init}}, E, \gamma)$ be a quantitative game. Let $\mu > 0$ be the maximum absolute value of costs along transitions in G . Then, let $\mathcal{A} = (S, s_I, \Sigma, \delta, \mathcal{F})$ be the co-safety comparator for lower approximation of discounted-sum with upper bound μ , discount factor $d = 1 + 2^{-k}$, approximation factor $\varepsilon = 2^{-p}$, threshold value v , and relation $\geq$. Then, the reachability game is $GA = (W = W_0 \uplus W_1, s_0 \times \text{init}, \delta_W, \mathcal{F}_W)$. Here, $W = V \times S$, $W_0 = V_0 \times S$, and $W_1 = V_1 \times S$. Clearly, W_0 and W_1 partition W . The edge relation $\delta_W \subseteq W \times W$ is defined such that edge $((v, s), (v', s')) \in \delta_W$ synchronizes between transitions $(v, v') \in E$ and $(s, a, s') \in \delta$ if $a = \gamma((v, v'))$ is the cost of transition (v, v') in G . State $s_0 \times \text{init}$ is the initial state and $\mathcal{F}_W = V \times \mathcal{F}$.

It suffices to prove that a play is winning for the system in GA iff its cost sequence A in G satisfies $\text{DSLow}(A, k, p) \geq 0$. This is ensured by the standard synchronized product construction and Theorem 2. The reachability game GA is linear in size of the quantitative graph and the comparator. $\square$

Theorem 3 describes a sound algorithm for reactive synthesis from satisficing goals when the discount factor is fractional. The algorithm is not complete since it is possible that there is a $v + d \cdot \varepsilon$ -satisficing strategy for the environment even when the system has a v -satisficing strategy.

5.2 Satisficing and Temporal Goals

Finally, we present our theoretically sound algorithm for synthesis from LTL and discounted-sum satisficing goals for fractional discount factors.

The algorithm is essentially a sum of two parts. The algorithm combines the automata-based solution for satisficing goals (presented in Section 5.1) with the classical automata-based solutions for LTL goals (Pnueli and Rosner 1989). Solving satisficing goals forms a reachability game while solving LTL goals forms a parity game. Thus, the final game which combines both of the goals will be a parity game. Lastly, the algorithm will inherit the soundness guarantees from both of its parts.

Theorem 4. *Let G be a quantitative game with discount factor $d = 1 + 2^{-k}$, for integer $k > 0$. Let φ be an LTL formula and $v \in \mathbb{Q}$ be a threshold value. Let $\varepsilon = 2^{-p}$ be the approximation factor. There exists a parity game GA such that*

- *If the system has a winning strategy in GA , then the system has a v -satisficing and φ -satisfying strategy in G .*
- *If the environment has a winning strategy in GA , then then either it has a $v + d \cdot \varepsilon$ -satisficing strategy or it has a winning strategy w.r.t. LTL formula in G .*

Proof Sketch. The reduction consists of two steps of synchronized products: first with the comparator to fulfil the v -satisficing goal and then with the automaton corresponding to the LTL goal. The first step conducts the reduction from Theorem 3 while lifting the labelling function from the quantitative game to the reachability game: If a state s is labeled by l in the quantitative game, the all states of the form (s, q) will be labelled by l in the reachability game. The second product synchronizes between the atomic propositions in the reachability game (with a labelling function) and the deterministic parity automaton (DPA) corresponding to the LTL specification, thus combining their winning conditions.

Observe that the product construction is commutative, i.e., one can first construct the product of G with the DPA of the LTL goal and then with the comparator automata.

In either case, we generate a parity game of size linear in $|G|$, DPA of the LTL specification, and the comparator. A winning strategy for the system player in this game is also v -satisficing and φ -satisfying for the same player in G . $\square$

A salient feature of our algorithm is that the complexity to solve the final product game is primarily governed by the temporal goal and not the satisficing goal. What we mean is that if the temporal goal is given by a fragment of LTL, such as co-safe LTL (Lahijanian et al. 2015), then the final product game would be reachability game. This is because co-safe LTL formulas are represented by co-safety automata and thus their combination with comparators would also be a co-safety automata. More generally, if the temporal goal is a conjunction of safety and reachability goals, the resulting game would be a *weak-Büchi game*, which are also solved in linear time in size of the game (Chatterjee 2008). This demonstrates that even though the comparator contributes to growing the state-space of the game linearly,

whether the game is solved using efficient linear-time algorithms or higher complexity algorithms for parity games is determined by the temporal goal.

This feature has implications on the practicality of our algorithm. In practice, it has been observed that wide-ranging temporal goals in robotics domains can be expressed in simpler fragments and variants of LTL, such as co-safe LTL (Lahijanian et al. 2015) and LTLf (He et al. 2017). These fragments can be expressed as conjunctions of safety and reachability goals. For this fragment synthesis from temporal and satisficing goals can be solved in linear-time.

6 Case Studies

The objective of our case studies is to demonstrate the utility of reactive synthesis from LTL and satisficing goals in realistic domains inspired from robotics and planning. Since ours is the first algorithm to offer theoretical guarantees with fractional discount factors, there are not any baselines to compare to. So, we focus on our scalability trends and identify future scalability opportunities.

6.1 Design and Set-up

We examine our algorithm on two challenging domains inspired from robotic navigation and manipulation problems. Source code and benchmarks are open source¹.

Grid World. The robot-human interaction is based on a classic $n \times n$ grid world domain (see Fig 1). The grid simulates a grocery store with static obstacles, e.g., placements of aisles. Each agent controls its own location and is only permitted to move in the cardinal directions

The robot’s LTL constraint is to reach the locations of all items on its grocery list without colliding with the walls (in grey) or the dynamic human, thus combining safety and reachability goals. The robot’s soft constraints are modelled to achieve two behaviors. The first one is to distance itself from the human. The second is to encode promptness in fulfilling its reachability goal `reach_banana`. We model distancing with quantitative rewards using the Manhattan distance between the two agents. Suppose, the locations of the players are (x_0, y_0) and (x_1, y_1) , then the reward received by the robot is given by $\left\lfloor \frac{\text{negative_reward}}{|x_0 - x_1| + |y_0 - y_1|} \right\rfloor$, where `negative_reward` < 0 is an integer parameter. We model promptness with an integer `positive_reward` > 0 which the robot receives only when it reaches a location of each item on its grocery list for the first time. The rewards are additive, i.e., the robot receives the sum of both rewards in every grid configuration. Then, it is reasonable to say that a play accomplishes these two behaviors if the discounted-sum reward of the robot is greater than or equal to 0, i.e., 0-satisficing plays/strategies are good for the robot.

Conveyor Belt. Our second case study is inspired by cutting-edge applications in manipulation tasks (Wells et al. 2021). A robot must operate along a $r \times c$ conveyor belt with r rows and c columns across from a human. When out

Dimensions	Number of States
Grid World	
$n = 4$	397
$n = 6$	2407
$n = 8$	8093
$n = 10$	20572
Conveyor Belt	
$r \times c = 4 \times 3$, 2 blocks	9966
$r \times c = 5 \times 3$, 2 blocks	31547
$r \times c = 5 \times 3$, 3 blocks	60540

Table 1: Complexity of Benchmarks: Number of states in product of the labelled quantitative game with the automata of its LTL specification.

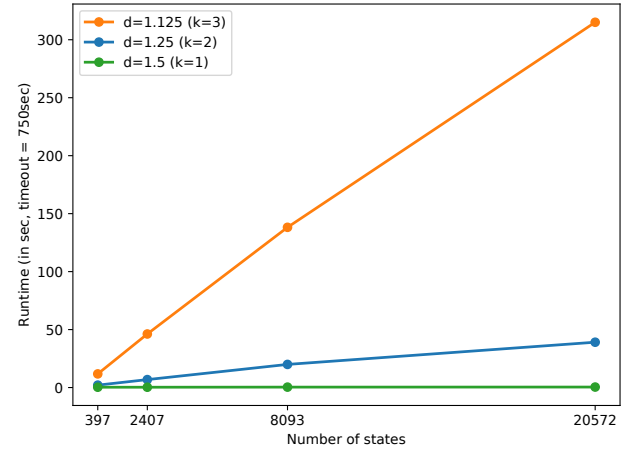


Figure 3: Scalability (Number of states). Plots runtime on Grid World with `positive_reward` = 10 and `negative_reward` = -2. x -axis are $n = 4, 6, 8, 10$.

of reach of the human, the robot can move quickly. Otherwise, it must proceed more slowly. The blocks move down the conveyor belt at a constant speed. Each agent controls the location of its arm. The human also controls the placement of new objects. New blocks are placed whenever a block is removed to maintain a constant number of blocks on the belt at all times.

The robot’s LTL goal is to avoid interfering with the human. As soft constraints, the robot gains a `positive_reward` for every object it grasps and a `negative_reward` for every object that falls off the belt. The rewards are additive in every belt configuration. The robot’s goal is to ensure its total discounted-sum reward exceeds 0.

On grid world, we take $n = 4, 6, 8, 10$. On conveyor belt, we take $r \times c = 4 \times 3, 5 \times 3$ with 2 or 3 blocks. The hardness of our benchmarks is illustrated Table 1. The benchmarks have so many states since both scenarios have a large number of unique configurations.

Combined with values for `positive_reward` and `negative_reward`, we create 20 grid world and 7 conveyor belt benchmarks. Every benchmark is run with

¹<https://github.com/suguman/NonIntegerGames>

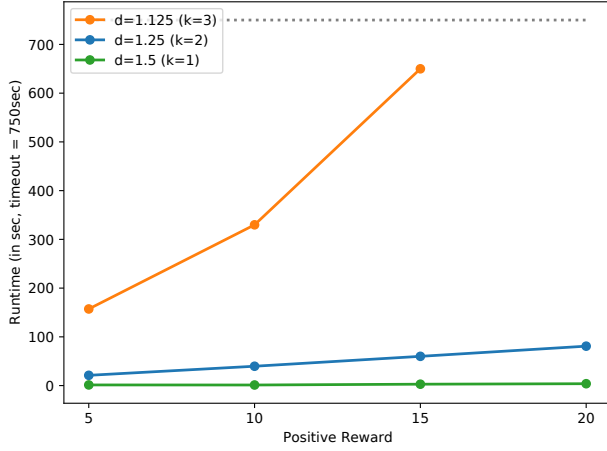


Figure 4: Scalability plot in positive_reward (affects the size of the comparator). Plotting runtime on Grid world with $n = 10$ (20572 states) and negative_reward = -2 .

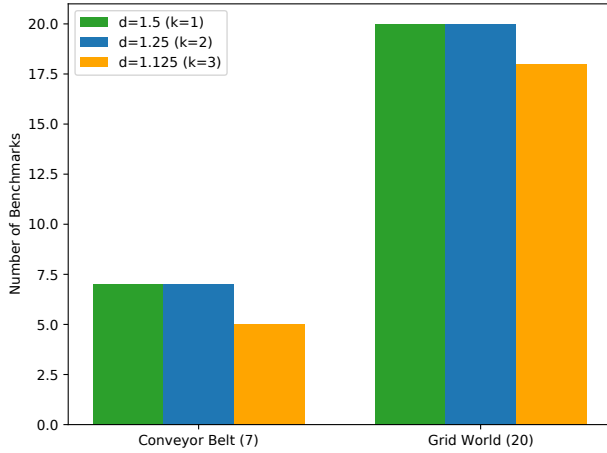


Figure 5: Number of benchmarks solved

$d = 1.5, 1.25, 1.125$ ($k = 1, 2, 3$), approx. factor $\varepsilon = 0.5$ ($p = 1$) and threshold $v = 0$. Our prototype is implemented in C++ on Ubuntu 18.04LTS. Experiments are run on an i7-4770 with 32GBs of RAM with a timeout of 750 sec.

6.2 Observations

Our evaluation demonstrates that our solution successfully scales to very large benchmarks. Despite their difficulty, we solve almost all of our benchmarks (Figure 5). Runtime examination indicates that our algorithm is linear in size of the game and the comparator, in practice. The scalability trends in size of the game for varying discount factors are shown in Figure 3. Determining the dependence on the comparator automata is more involved since its size depends on several parameters, namely positive_reward, the discount factor, and the approximation factor. Figure 4 suggests the algorithm is linear in positive_reward. The margin between the three discount factor curves on Fig 3-4 suggests a significant blow-up

as the discount factor nears 1. Additional experiments (see Appendix) that vary the approximation factor also display a significant blow-up as the approximation factor decreases. These are not alarming since the size of the comparator is in the order of $\mathcal{O}(\text{positive_reward})$, $\mathcal{O}((d-1)^{-2})$ and $\mathcal{O}(\varepsilon^{-1})$. These reflect that our current implementation is faithful to the theoretical analysis of the algorithm.

These are encouraging results as our implementation uses explicit state representation. The overhead of this state representations can be very high. In some cases, we observed that for the large benchmarks about 70% of the total compute time may be spent in constructing the product explicitly. Despite these issues with explicit-state representation, our algorithm efficiently scales to large and challenging benchmarks. This indicates potential for further improvements.

In terms of quality of solutions, the resulting strategies are of better quality. For example, in Figure 1 we observed that as the discount factor becomes smaller the robot is able to reason for a longer horizon and not get "trapped". Another benefit are the soundness guarantees. They are especially valuable in environments such as the Conveyor belt which are so complex that they preclude a manual analysis.

To conclude, our case studies demonstrates the promise of our approach in terms of its ability to scale and utility in practical applications, and encourage future investigations.

7 Conclusion

Combining hard constraints (qualitative) with soft constraints (quantitative) is a challenging problem with many applications to automated planning. This paper presents the first sound algorithm for reactive synthesis from LTL constraints with soft discounted-sum rewards when the discount factor is fractional. Our approach uses an automata-based method to solve the soft constraints, which is then elegantly combined with existing automata-based methods for LTL constraints to obtain the sound solution. Case studies on classical and modern domains of robotics planning demonstrate use cases, and also, shed light on recommendations for future work to improve scalability to open up exciting applications in robotics e.g. warehouse robotics, autonomous driving, logistics (supply-chain automation).

Acknowledgements

We thank anonymous reviewers. This work is supported in part by NSF grant 2030859 to the CRA for the CIFellows Project, NSF grants IIS-1527668, CCF-1704883, IIS-1830549, CNS-2016656, DoD MURI grant N00014-20-1-2787, and an award from the Maryland Procurement Office.

References

- Bansal, S.; Chatterjee, K.; and Vardi, M. Y. 2021. On Satisficing of Quantitative Games. In *Proc. of TACAS*.
- Bansal, S.; Chaudhuri, S.; and Vardi, M. Y. 2018a. Automata vs Linear-Programming Discounted-Sum Inclusion. In *Proc. of CAV*.
- Bansal, S.; Chaudhuri, S.; and Vardi, M. Y. 2018b. Comparator automata in quantitative verification. In *Proc. of FOSSACS*.

- Bansal, S.; Chaudhuri, S.; and Vardi, M. Y. 2018c. Comparator automata in quantitative verification (full version). *CoRR*, abs/1812.06569.
- Bansal, S.; and Vardi, M. Y. 2019. Safety and Co-safety Comparator Automata for Discounted-Sum Inclusion. In *Proc. of CAV*.
- Bozkurt, A. K.; Wang, Y.; Zavlanos, M. M.; and Pajic, M. 2020. Control synthesis from linear temporal logic specifications using model-free reinforcement learning. In *2020 IEEE International Conference on Robotics and Automation (ICRA)*, 10349–10355. IEEE.
- Camacho, A.; Bienvenu, M.; and McIlraith, S. A. 2019. Towards a unified view of AI planning and reactive synthesis. In *Proc. of ICAPS*.
- Camacho, A.; Chen, O.; Sanner, S.; and McIlraith, S. A. 2017. Non-markovian rewards expressed in LTL: guiding search via reward shaping. In *In Proc. of SOCS*.
- Camacho, A.; Icarte, R. T.; Klassen, T. Q.; Valenzano, R. A.; and McIlraith, S. A. 2019. LTL and Beyond: Formal Languages for Reward Function Specification in Reinforcement Learning. In *Proc. of IJCAI*.
- Chatterjee, K. 2008. Linear time algorithm for weak parity games. *arXiv preprint arXiv:0805.1391*.
- Chatterjee, K.; Henzinger, T. A.; Otop, J.; and Velner, Y. 2017. Quantitative fair simulation games. *Information and Computation*.
- Church, A. 1957. Applications of recursive arithmetic to the problem of circuit synthesis. *Institute for Symbolic Logic, Cornell University*.
- Ding, X.; Smith, S. L.; Belta, C.; and Rus, D. 2014. Optimal control of Markov decision processes with linear temporal logic constraints. *TACON*.
- Hasanbeig, M.; Kantaros, Y.; Abate, A.; Kroening, D.; Pappas, G. J.; and Lee, I. 2019. Reinforcement learning for temporal logic control synthesis with probabilistic satisfaction guarantees. In *2019 IEEE 58th Conference on Decision and Control (CDC)*, 5338–5343. IEEE.
- He, K.; Lahijanian, M.; Kavraki, L.; and Vardi, M. 2017. Reactive synthesis for finite tasks under resource constraints. In *Proc. of IROS*.
- He, K.; Wells, A. M.; Kavraki, L. E.; and Vardi, M. Y. 2019. Efficient symbolic reactive synthesis for finite-horizon tasks. In *Proc. of ICRA*.
- Kalagarla, K. C.; Jain, R.; and Nuzzo, P. 2021. Optimal Control of Discounted-Reward Markov Decision Processes Under Linear Temporal Logic Specifications. In *2021 American Control Conference (ACC)*, 1268–1274. IEEE.
- Kress-Gazit, H.; Lahijanian, M.; and Raman, V. 2018. Synthesis for robots: Guarantees and feedback for robot behavior. *Annual Review of Control, Robotics, and Autonomous Systems*.
- Kupferman, O.; and Vardi, M. Y. 1999. Model checking of safety properties. In *Proc. of CAV*.
- Kwiatkowska, M.; Parker, D.; and Wiltsche, C. 2017. PRISM-games: Verification and Strategy Synthesis for Stochastic Multi-player Games with Multiple Objectives. *STTT*.
- Lahijanian, M.; Almagor, S.; Fried, D.; Kavraki, L.; and Vardi, M. 2015. This Time the Robot Settles for a Cost: A Quantitative Approach to Temporal Logic Planning with Partial Satisfaction. In *Proc. of AAAI*.
- Pnueli, A. 1977. The temporal logic of programs. In *Proc. of FOCS*.
- Pnueli, A.; and Rosner, R. 1989. On the synthesis of a reactive module. In *Proc. of POPL*.
- Shapley, L. S. 1953. Stochastic games. *Proceedings of the National Academy of Sciences of the United States of America*, 39(10): 1095.
- Sutton, R.; and Barto, A. 2018. *An Introduction to Reinforcement Learning, Second Edition*. MIT press Cambridge.
- Thomas, W.; Wilke, T.; et al. 2002. *Automata, logics, and infinite games: A guide to current research*.
- Wells, A. M.; Kingston, Z.; Lahijanian, M.; Kavraki, L. E.; and Vardi, M. Y. 2021. Finite-Horizon Synthesis for Probabilistic Manipulation Domains. In *IEEE Int. Conf. Robot. Autom.*
- Wen, M.; Ehlers, R.; and Topcu, U. 2015. Correct-by-synthesis reinforcement learning with temporal logic constraints. In *Proc. of IROS*.