

PAC Learning in Turn-Based Stochastic Games with Reachability Objectives: A Decentralized Private Approach via Expected Conditional Distance

Ali Asadi 

Institute of Science and Technology Austria (ISTA), Klosterneuburg, Austria

Krishnendu Chatterjee 

Institute of Science and Technology Austria (ISTA), Klosterneuburg, Austria

Pavol Kebis 

Institute of Science and Technology Austria (ISTA), Klosterneuburg, Austria

Abstract

Reachability is the most fundamental logical objective, yet it is notoriously difficult to learn in reinforcement learning settings: even for Markov decision processes, PAC learning of reachability is impossible without additional assumptions. This difficulty also holds in turn-based stochastic games (TBSGs), where two adversarial players interact on a finite state space. In this work, we consider turn-based stochastic games with reachability objectives. For such settings, adversarial learning, in which players are adversarial even in the learning phase, is impossible. Therefore, the goal is to consider learning, in which both players learn the unknown model together. In this spirit, previous literature on PAC learning in TBSGs considers (a) public information shared by both players; and (b) centralized learning, which means that players share the same learning algorithm. In this work, our contribution is two-fold. First, we relax these strong assumptions and ensure learning: (i) with private information not shared with the other player; and (ii) decentralized learning where the players do not share the same learning algorithm. To the best of our knowledge, this work is the first positive result for decentralized and private information learning of TBSGs with reachability objectives. Second, we introduce a game-theoretic generalization of the Expected Conditional Distance (ECD) parameter, which measures the expected length of reaching the target set. We establish a polynomial-sample complexity bound with respect to the number of states, actions, ECD parameter, and inverses of error tolerance and failure probability.

2012 ACM Subject Classification Theory of computation → Logic and verification

Keywords and phrases formal methods, games and logic, logical aspects of AI, model checking

Digital Object Identifier 10.4230/LIPIcs.CONCUR.2026.12

Funding The research was partially supported by Austrian Science Fund (FWF) 10.55776/COE12, ERC CoG 863818 (ForM-SMArt), FWF-2022-SFB F8502 (SPyCoDe), and ERC-2020-AdG 101020093 (VAMOS) grants.

1 Introduction

Turn-Based Stochastic Games. Turn-based stochastic games (TBSGs) [7] are zero-sum turn-based games played over a finite state space by two adversarial players, Max and Min, along with randomness in the transition function. The state space is partitioned into two disjoint sets for Max and Min. At each time step, the player owning the current state chooses an action. The subsequent state is then determined by a probabilistic transition function. This model generalizes several classical formalisms such as Markov decision processes (MDP) [19], which have only one player and stochastic uncertainty, and graph games [6, 11], where the transition function collapses to Dirac distributions.



© Ali Asadi, Krishnendu Chatterjee, and Pavol Kebis;
licensed under Creative Commons License CC-BY 4.0

37th International Conference on Concurrency Theory (CONCUR 2026).

Editors: Ana Sokolova and Patrick Totzke; Article No. 12; pp. 12:1–12:23

Leibniz International Proceedings in Informatics



LIPICs Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl Publishing, Germany

Objectives. In TBSGs, the interaction of players is guided by an *objective function*, which formally captures the desired behaviour of the model. Objectives are typically categorized into: (a) logical objectives, e.g., reachability, safety, and parity; and (b) quantitative objectives, e.g., finite-horizon, discounted sum, and mean payoff. This work focuses on reachability objectives, which are the most fundamental logical objectives, i.e., given a set of target states, the objective requires that some target state is eventually visited. It is important to distinguish reachability from discounted sum or finite-horizon objectives. Discounted sum objectives introduce a discount factor $\lambda < 1$, which effectively imposes a “soft” horizon. Finite-horizon objectives strictly bound the interaction to L steps. In contrast, reachability is an unbounded property; a target might be reached after an arbitrarily large number of steps.

Strategies and Values. Strategies are recipes that define the choice of actions of the players. They are functions that, given a game history, return a distribution over actions. Given a TBSG and an objective, the value of player Max at a state is the maximal expectation that the player can guarantee for the objective against all strategies of player Min. A strategy is ϵ -optimal if it guarantees the value up to additive error ϵ .

PAC Learning. While classical model checking assumes a known model, in the reinforcement learning setting the model is unknown. The players must learn near-optimal strategies solely through interaction with a simulator. In this setting, the gold standard is Probably Approximately Correct (PAC) guarantees for learning near-optimal strategies [21]. The PAC-RL problem is defined as follows.

Can we design learning algorithms for both players such that for any error tolerance $\epsilon > 0$ and failure probability $p \in (0, 1)$, the algorithms output strategies that are ϵ -optimal with probability at least $1 - p$?

Crucially, for the problem to be considered tractable, the number of samples required by the algorithm (called *sample complexity*) must be polynomial in number of states and actions, inverse error tolerance $1/\epsilon$ and inverse failure probability $1/p$.

Expected Conditional Distance. Even in MDPs, the PAC-RL problem for reachability objectives is impossible in general [1, 23]. Thus, to circumvent this impossibility the literature considers further assumptions including prior knowledge on (a) the topology of the underlying graph [10]; (b) the minimum non-zero probability [2]; and (c) a parameter called the Expected Conditional Distance (ECD), which was introduced in [20] for MDPs. The ECD parameter provides a measure of the expected number of steps to reach the target. We generalize ECD to the TBSG setting. This generalization is quite subtle, as several natural generalizations of ECD to games fail to achieve PAC guarantees. Intuitively, if a game has a small ECD, it implies that if the target is reachable, it is reachable relatively quickly on average. This assumption excludes pathological games where the only optimal strategies involve waiting for exponentially many steps. Bounding the ECD allows us to truncate the infinite-horizon, converting the intractable reachability problem into a tractable finite-horizon approximation.

Tractable Private and Decentralized Learning. Adversarial learning, in which players are adversarial even in the learning phase, is impossible for TBSGs with reachability objectives: consider a game with an initial player-Min state and an additional action that immediately leads to the target. In the learning phase, Min chooses the trivial target-reaching action,

which is never part of the optimal strategy, rendering learning useless. Since adversarial learning with PAC guarantees is impossible, the goal is to consider learning where both players learn the unknown model together. In this setting, [2] established an anytime algorithm with the prior knowledge on (a) the minimum non-zero transition probability; or (b) the topology of the underlying graph. However, this prior work [2] has two important limitations: First, it assumes (i) public information shared by both players; and (ii) centralized learning where players share the same learning algorithm. Second, while the algorithm is anytime, it does not provide sample-complexity bounds for the PAC-RL problem.

Motivation. The motivation of this work is two-fold. The main motivation is to relax the above two strong assumption and ensure learning: (i) with private information not shared with the other player; and (ii) decentralized learning where the players do not share the same learning algorithm. Second, even in previous setting of centralized learning with public information, sample complexity bound was not established. The goal is to establish a polynomial-sample complexity bound with respect to the number of states, actions, ECD parameter, and inverses of error tolerance and failure probability.

Our Contributions. We address the above gaps by considering the decentralized private information setting of TBSGs with reachability objectives. We present a pair of algorithms for both players which are PAC-RL learnable with prior knowledge on the ECD parameter. The sample complexity of these algorithms is polynomial in the game parameters, the inverses of the error tolerance and failure probability, and the ECD parameter. To the best of our knowledge, this pair of algorithms is the first positive result for decentralized private information PAC-RL learning in TBSGs with reachability objectives.

Technical Contributions. Our technical contributions are as follows.

- We generalize the Expected Conditional Distance (ECD) parameter to the TBSG setting.
- We provide a reduction showing that, if the ECD of a game is small, we can approximate the reachability value using a finite-horizon reachability objective.
- We define a finite-horizon expanded game over state-step pairs that unfolds the horizon into the state space, enabling backward induction and local learning at each state-step.
- We present a learning procedure where both players use backward induction to learn local ϵ -optimal actions one step at a time. The algorithm uses a best-arm identification routine at each state-step to identify ϵ -optimal actions with high confidence.
- The algorithm iteratively constructs a set of strategies in stages. Each newly constructed strategy is added to the set used in subsequent stages to ensure that previously discovered state-steps of the game remain reachable while the players explore new state-steps.
- To explore new state-steps, we maintain a set of unexplored ones. These are treated as auxiliary target sets, incentivising the players to visit more of the state-step space.

Proofs omitted due to space restrictions are provided in the Appendix.

Technical Novelty. The technical novelty of this work is two-fold. The first novelty is the appropriate definition of ECD for games. Second, the previous works rely on estimating the underlying probabilistic transitions, which is infeasible in private decentralized learning. Our approach carefully combines different techniques: best arm identification; tracking strategies for exploration; and backward induction to directly compute near-optimal strategies.

Related Works. The intersection of formal verification and reinforcement learning has recently received significant attention. We summarize some related works as follows.

- *MDPs.* PAC guarantees for complex logical objectives, such as Linear Temporal Logic (LTL) [18], has seen significant development but remains constrained by specific environmental assumptions, due to the inherent impossibility of learnability in general settings without additional assumptions [1, 23]. Early PAC learning algorithm presented in [10] required complete knowledge of the environment's topology. [2] improved upon this by requiring only a lower bound on the minimum non-zero transition probability. More recently, [17] has established PAC results using the mixing time of the environment. [20] has introduced the ECD parameter and established PAC results relying on this parameter.
- *TBSGs.* Early works on learning in TBSGs focused mainly on quantitative objectives and did not provide PAC guarantees [14, 15, 5]. For logical objectives, [22] has established PAC learning algorithms for TBSGs with LTL objectives by combining the special case of almost-sure satisfaction of a specification with optimizing quantitative objectives. [2] obtained PAC guarantees for reachability objectives by computing under- and over-approximation of values, originally introduced in [13]. It is noteworthy that all these works consider the centralized public information setting.

2 Preliminaries

In this section, we define the notations of turn-based stochastic games and PAC learning.

Notation. For a positive integer n , the set $\{0, 1, 2, \dots, n\}$ is denoted by $[n]$. Open and closed intervals of reals are denoted $(x, y) = \{a \in \mathbb{R} \mid x < a < y\}$ and $[x, y] = \{a \in \mathbb{R} \mid x \leq a \leq y\}$, respectively. Sets are denoted by calligraphic letters, e.g., $\mathcal{S}, \mathcal{A}$. Elements of sets are denoted by lowercase letters, e.g., s, a . The set of probability distributions over a set $\mathcal{S}$ is denoted by $\Delta(\mathcal{S})$. The set of natural numbers is $\mathbb{N} = \{0, 1, 2, \dots\}$.

2.1 Turn-Based Stochastic Games

► **Definition 1** (Turn-Based Stochastic Games). *A turn-based stochastic game (TBSG for short) is a tuple $\mathbb{G} = (\mathcal{S}, \mathcal{A}, \delta, \mu)$ where*

- $\mathcal{S} = \mathcal{S}_{\text{Max}} \uplus \mathcal{S}_{\text{Min}}$ *is a finite set of states, partitioned into the set of player-Max states $\mathcal{S}_{\text{Max}}$ and the set of player-Min states $\mathcal{S}_{\text{Min}}$;*
- $\mathcal{A}$ *is a finite set of actions;*
- $\delta: \mathcal{S} \times \mathcal{A} \rightarrow \Delta(\mathcal{S})$ *is a probabilistic transition function which, given a state and an action, assigns a probability distribution over the successor state; and*
- $\mu \in \Delta(\mathcal{S})$ *is a probability distribution over the initial state.*

Dynamic. At the beginning, an initial state $s_0 \sim \mu$ is drawn, and the game proceeds as follows. In each step $\ell \in \mathbb{N}$, the owner of the state s_ℓ selects an action $a_\ell \in \mathcal{A}$, possibly at random, and the successor state $s_{\ell+1} \sim \delta(s_\ell, a_\ell)$ is drawn.

Histories and Plays. A history is a finite sequence $h = (s_0, a_0, s_1, a_1, \dots, s_L)$ of states and actions such that for all $\ell \in [L - 1]$, we have $\delta(s_\ell, a_\ell)(s_{\ell+1}) > 0$. A play is an infinite sequence of states and actions $\omega = (s_0, a_0, s_1, a_1, \dots)$ such that, for all $\ell \in \mathbb{N}$, we have $\delta(s_\ell, a_\ell)(s_{\ell+1}) > 0$. The set of all plays is denoted by Ω .

Strategies. A strategy determines how a player chooses an action based on the history up to a given step. Formally, a strategy for a player $i \in \{\text{Max}, \text{Min}\}$ is a function $\pi_i: (\mathcal{S} \times \mathcal{A})^* \times \mathcal{S}_i \rightarrow \Delta(\mathcal{A})$. The set of all strategies is denoted by Π_i . A strategy is Markovian if it depends on the current state and current step of the play, i.e., $\pi_i: \mathcal{S}_i \times \mathbb{N} \rightarrow \Delta(\mathcal{A})$. A strategy is pure if it prescribes deterministic actions, i.e., it corresponds to a function $\pi_i: (\mathcal{S} \times \mathcal{A})^* \times \mathcal{S}_i \rightarrow \mathcal{A}$. A strategy is memoryless if it decides only based on the current state, i.e., $\pi_i: \mathcal{S}_i \rightarrow \Delta(\mathcal{A})$. A strategy is positional if it is pure and memoryless. Note that in TBSGs with reachability objectives positional strategies are as powerful as general strategies [7]. Given strategies for both players π_{Max} and π_{Min} , we denote the strategy profile by $(\pi_{\text{Max}}, \pi_{\text{Min}})$, and if the context is clear, we simply use π .

Probability Measures. For a history h , its cone is the set of plays where h is their prefix. Given a strategy profile π and an initial belief μ , the unique probability measure over Borel sets of infinite plays is denoted by $\mathbb{P}_\mu^\pi(\cdot)$, which is defined by Carathéodory's extension theorem by extending the natural definition over cones of plays [4].

Reachability Objectives. An *objective* in a TBSG is a Borel set of plays $\Phi \subseteq \Omega$ in the Cantor topology on Ω [12]. In this work, we consider reachability objectives which lie in the first level of the Borel hierarchy. Given a set of target states $\mathcal{T}$, the reachability objective requires that a target state is eventually visited, i.e., $\text{Reach}(\mathcal{T}) := \{\omega \in \Omega: \exists \ell \in \mathbb{N} \quad s_\ell \in \mathcal{T}\}$. The goal of player Max is to maximize the probability of satisfying the objective, while the goal of player Min is to minimize it.

We now recall a fundamental determinacy for TBSGs with reachability objectives.

► **Theorem 2 (Determinacy [7]).** *For all TBSGs with a target set $\mathcal{T} \subseteq \mathcal{S}$, we have*

$$\sup_{\pi_{\text{Max}} \in \Pi_{\text{Max}}} \inf_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \mathbb{P}_\mu^\pi(\omega \in \text{Reach}(\mathcal{T})) = \inf_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \sup_{\pi_{\text{Max}} \in \Pi_{\text{Max}}} \mathbb{P}_\mu^\pi(\omega \in \text{Reach}(\mathcal{T})).$$

Values. Theorem 2 implies that switching the quantifiers does not make a difference and leads to a unique notion of value. Formally, given a target set $\mathcal{T}$, the value is a function of initial distribution

$$\text{Val}_{R(\mathcal{T})}^\mathbb{G}(\mu) := \sup_{\pi_{\text{Max}} \in \Pi_{\text{Max}}} \inf_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \mathbb{P}_\mu^\pi(\omega \in \text{Reach}(\mathcal{T})).$$

We omit writing $\mathbb{G}$ when clear from the context.

Approximately Optimal Strategies. Given $\varepsilon \geq 0$, a strategy π_{Max} for player Max is ε -optimal if it guarantees the value up to an additive error ε , i.e., if $\inf_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \mathbb{P}_\mu^\pi(\omega \in \text{Reach}(\mathcal{T})) \geq \text{Val}_{R(\mathcal{T})}(\mu) - \varepsilon$. We denote the set of ε -optimal strategies by $\Pi_{\text{Max}}^\varepsilon$. In particular, we call a 0-optimal strategy simply optimal. The definition of ε -optimal strategies for player Min is analogous.

Best-responses. For a player-Min strategy π_{Min} , we define the set of best-responses for player Max as

$$\text{BR}(\pi_{\text{Min}}) := \left\{ \pi_{\text{Max}} \in \Pi_{\text{Max}} : \mathbb{P}_\mu^{(\pi_{\text{Max}}, \pi_{\text{Min}})}(\omega \in \text{Reach}(\mathcal{T})) = \sup_{\pi'_{\text{Max}} \in \Pi_{\text{Max}}} \mathbb{P}_\mu^{(\pi'_{\text{Max}}, \pi_{\text{Min}})}(\omega \in \text{Reach}(\mathcal{T})) \right\}.$$

The set of best-responses for player Min is defined analogously.

2.2 Reinforcement Learning for TBSGs

In the reinforcement learning setting for TBSGs with reachability objectives, we consider a scenario where the players have no information about the transition probabilities δ or the initial distribution μ ; only the parameters $\mathcal{S}$ and $\mathcal{A}$ are known to both players, and the players access the TBSG only through a simulator $\mathbb{M}$. The goal of both players is to use *learning algorithms* to find a near-optimal strategy profile. In this work, the learning algorithms are *decoupled*, i.e., each player has its own learning algorithm that does not communicate with the other player's algorithm. The assumption of *private states* is another difference that distinguishes our setting from previously considered settings on TBSGs. We assume that the current state of a play is announced only to its owner and not to the other player. In the rest, we formalize the notion of simulators, learning algorithms, and PAC-RL in this setting.

► **Definition 3 (Simulators).** *Given a TBSG $\mathbb{G} = (\mathcal{S}, \mathcal{A}, \delta, \mu)$ with a target set $\mathcal{T} \subseteq \mathcal{S}$, a simulator $\mathbb{M}$ stores the current state of the game, receives inputs from both players, performs actions, and outputs to players the new state to which the play is proceeded. Precisely, it works as follows:*

1. Any player i can propose to terminate the simulator with a strategy π_i by calling the procedure $\mathbb{M}.\text{propose}(\pi_i)$. The simulator terminates with a strategy profile π only if both players propose.
2. $\mathbb{M}$ informs both players that a new play has started;
3. $\mathbb{M}$ samples the initial state $s \sim \mu$;
4. $\mathbb{M}$ repeats the following:
 - a. the active player i is Max if $s \in \mathcal{S}_{\text{Max}}$ or Min if $s \in \mathcal{S}_{\text{Min}}$;
 - b. if $s \in \mathcal{T}$, both players are informed that the target was reached and the simulator returns to step 1;
 - c. both players are informed who the active player is but **only the active player has access to the current state s** ;
 - d. active player i either (I) chooses an action $a \in \mathcal{A}$ by calling the procedure $\mathbb{M}.\text{step}(a)$; or (II) resets the game by calling the procedure $\mathbb{M}.\text{reset}()$. Note that **a is announced only to the simulator and not the other player**. The play proceeds with transitioning to a new state $s' \sim \delta(s, a)$ and returning to step 4.

► **Definition 4 (Learning Algorithms).** *A learning algorithm $\mathfrak{A}_i$ for a player $i \in \{\text{Max}, \text{Min}\}$ is an algorithm that interacts with the simulator $\mathbb{M}$ by calling the procedures $\mathbb{M}.\text{step}(a)$, $\mathbb{M}.\text{reset}()$, and $\mathbb{M}.\text{propose}(\pi_i)$ where $a \in \mathcal{A}$, and π_i is a player- i strategy. A learning algorithm is decoupled if it does not communicate with the other player's learning algorithm.*

► **Definition 5 (PAC-RL).** *A pair of learning algorithms $(\mathfrak{A}_{\text{Max}}, \mathfrak{A}_{\text{Min}})$ is PAC-RL for reachability objectives if there exists a function f such that for all $\varepsilon, p \in (0, 1)$ and all TBSGs $\mathbb{G} = (\mathcal{S}, \mathcal{A}, \delta, \mu)$ with a target set $\mathcal{T}$, taking $N = f(|\mathcal{S}|, |\mathcal{A}|, \frac{1}{p}, \frac{1}{\varepsilon})$, with probability at least $1 - p$, the simulator terminates with a strategy profile $(\pi_{\text{Max}}, \pi_{\text{Min}})$ after at most N procedure calls where both strategies are ε -optimal.*

Sample Complexity. The function f in the definition of PAC-RL is called the sample complexity of the learning algorithms. If f is a polynomial function, then we say that the learning algorithms have polynomial sample complexity.

General Hardness. It is known that, even for MDPs with reachability objectives, there is no algorithm that is PAC-RL in general [1, 23], meaning that there is no function f that satisfies the condition of PAC-RL. In order to circumvent this hardness, we consider a parameter called *Expected Conditional Distance* (ECD).

3 Expected Conditional Distance

In this section, we introduce a parameter for TBSGs called the *Expected Conditional Distance* (ECD). The ECD parameter was previously studied for MDPs with reachability objectives [20]. The main goal of this parameter is to reduce the PAC-RL for reachability to PAC-RL for finite-horizon reachability. The generalization of this parameter to TBSGs is quite subtle, since we show below that several natural generalizations do not yield a suitable bound on the horizon. We then provide an appropriate generalization to TBSGs and give a reduction from PAC-RL with the ECD parameter to PAC-RL for finite-horizon reachability objectives. Finally, we discuss several key aspects of our parameter which justify the usefulness: (a) its important properties that make it useful for PAC learning; (b) how it can be bounded using other classical parameters from the literature; and (c) how it compares with the well-studied stochastic shortest path parameter.

► **Definition 6** (Alternative Generalizations). *Given a TBSG $G = (\mathcal{S}, \mathcal{A}, \delta, \mu)$ with a target set $\mathcal{T} \subseteq \mathcal{S}$, consider the following alternative definitions of ECD:*

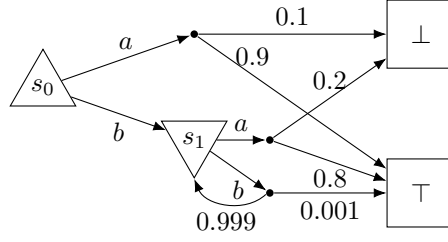
$$\begin{aligned} ECD_G^1 &:= \sup_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \inf_{\pi_{\text{Max}} \in \Pi_{\text{Max}}} ETR_G(\pi_{\text{Max}}, \pi_{\text{Min}}), \\ ECD_G^2 &:= \sup_{\pi_{\text{Min}} \in \Pi_{\text{Min}}^0} \inf_{\pi_{\text{Max}} \in \text{BR}(\pi_{\text{Min}})} ETR_G(\pi_{\text{Max}}, \pi_{\text{Min}}), \\ ECD_G^3 &:= \inf_{\pi_{\text{Max}} \in \Pi_{\text{Max}}^0} \sup_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} ETR_G(\pi_{\text{Max}}, \pi_{\text{Min}}), \\ ECD_G^4 &:= \sup_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \inf_{\pi_{\text{Max}} \in \Pi_{\text{Max}}^0} ETR_G(\pi_{\text{Max}}, \pi_{\text{Min}}), \\ ECD_G^5 &:= \inf_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \sup_{\pi_{\text{Max}} \in \Pi_{\text{Max}}} ETR_G(\pi_{\text{Max}}, \pi_{\text{Min}}), \end{aligned}$$

where $ETR(\pi)$ is the expected time to reach the target set using the strategy profile $(\pi_{\text{Max}}, \pi_{\text{Min}})$:

$$ETR_G(\pi_{\text{Max}}, \pi_{\text{Min}}) := \mathbb{E}_\mu^\pi \left(\arg \inf_{n \in \mathbb{N}} \mathbb{1}(s_n \in \mathcal{T}) \mid (s_n)_{n \in \mathbb{N}} \cap \mathcal{T} \neq \emptyset \right).$$

These definitions fail in the following example for a reduction of PAC-RL for reachability to PAC-RL for finite-horizon reachability.

► **Example 7.** Consider a game, shown in Figure 1, with four states $\mathcal{S} = \{s_0, s_1, \top, \perp\}$ where $\top$ and $\perp$ are absorbing states. State s_0 belongs to Max and s_1 belongs to Min. The action set is $\mathcal{A} = \{a, b\}$. The target set is $\mathcal{T} = \{\top\}$. In state s_0 , playing action a leads to $\top$ with probability 0.9 and leads to $\perp$ with probability 0.1, and playing action b leads to s_1 with probability 1. In state s_1 , playing action a leads to $\top$ with probability 0.8 and leads to $\perp$ with probability 0.2, and playing action b leads to $\top$ with probability 0.001 and self loops with probability 0.999. The initial state is s_0 . Therefore, in the case of the infinite-horizon version of the game, the optimal strategies for both players are to play the action a . However, for any finite-horizon game with horizon $L \leq 100$, the optimal strategy for Min is to play action b . We need the ECD parameter to bound a horizon length for which a near-optimal



■ **Figure 1** A game where alternative generalizations of ECD fail.

strategy in the finite-horizon game is also near-optimal in the infinite-horizon game. However, all of the definitions above fail as their values are at most 2. The values $ECD_{\mathbb{G}}^1, ECD_{\mathbb{G}}^2, ECD_{\mathbb{G}}^3$, and $ECD_{\mathbb{G}}^4$ are equal to 1 for this game since the infimum over player-Max actions selects action a which ends the game immediately. The value of $ECD_{\mathbb{G}}^5$ is 2 since the infimum over player-Min actions selects the action a . It is noteworthy that changing the probabilities of action b in the state s_1 makes the gap between the needed horizon and ECD values larger. In contrast, our definition of ECD provides a suitable bound on the horizon since $ECD_{\mathbb{G}} = 1001$.

► **Definition 8** (Expected Conditional Distance). *Given a TBSG $\mathbb{G} = (\mathcal{S}, \mathcal{A}, \delta, \mu)$ with a target set $\mathcal{T} \subseteq \mathcal{S}$, the expected conditional distance is defined as follows.*

$$ECD_{\mathbb{G}} := \sup_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \inf_{\pi_{\text{Max}} \in \text{BR}(\pi_{\text{Min}})} ETR_{\mathbb{G}}(\pi_{\text{Max}}, \pi_{\text{Min}}),$$

where $ETR_{\mathbb{G}}(\pi_{\text{Max}}, \pi_{\text{Min}})$ is defined as in Definition 6.

Description of ECD. If the ECD parameter is bounded by L , then for all strategies for player Min, there exists a best-response strategy of player Max that can reach the target set $\mathcal{T}$ in expected time at most L . More formally, the ECD parameter is defined as follows. For any player-Min strategy, take the set of player-Max strategies that are best-responses for the reachability objective. Among these reachability-optimal strategies, ECD takes the one that minimizes the expected number of steps to reach the target, conditioned on the target being reached. Finally, ECD takes the maximum of this quantity over all player-Min strategies.

We now define the PAC learning framework with respect to the ECD parameter.

► **Definition 9** (PAC-RL with ECD). *A pair of learning algorithms $(\mathfrak{A}_{\text{Max}}, \mathfrak{A}_{\text{Min}})$ is PAC-RL with ECD if there exists a function f such that for all $\varepsilon, p \in (0, 1)$, all $L \in \mathbb{N}$, and all TBSGs $\mathbb{G} = (\mathcal{S}, \mathcal{A}, \delta, \mu)$ with a target set $\mathcal{T} \subseteq \mathcal{S}$ such that $ECD_{\mathbb{G}} \leq L$, taking $N = f(|\mathcal{S}|, |\mathcal{A}|, \frac{1}{p}, \frac{1}{\varepsilon}, L)$, with probability at least $1 - p$, the simulator terminates with a strategy profile $(\pi_{\text{Max}}, \pi_{\text{Min}})$ after at most N procedure calls where both strategies are ε -optimal.*

► **Remark 10.** The difference between this definition and the standard PAC-RL definition is the inclusion of the ECD parameter L in the function f .

We now define an objective called finite-horizon reachability and show that the problem of PAC-RL with ECD for reachability objectives can be reduced to the problem of PAC-RL for finite-horizon reachability objectives.

► **Definition 11** (Finite-horizon Reachability Objectives). *Given a set of target states $\mathcal{T}$ and a time horizon $L \in \mathbb{N}$, the finite-horizon reachability objective requires that a target state is visited within the first L steps, i.e., $\text{Reach}_L(\mathcal{T}) := \{\omega \in \Omega : \exists \ell \in [L] \quad s_{\ell} \in \mathcal{T}\}$. We also admit*

$L \in \mathbb{R}$ in which case the target has to be visited within the first $\lfloor L \rfloor$ where $\lfloor x \rfloor$ is the biggest number $n \in \mathbb{N}$ such that $n \leq x$. We denote $\text{Val}_{R_L(\mathcal{T})}^{\mathbb{G}}(\mu) := \sup_{\pi_{\text{Max}} \in \Pi_{\text{Max}}} \inf_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \mathbb{P}_{\mu}^{\pi}(\omega \in \text{Reach}_L(\mathcal{T}))$. We omit writing $\mathbb{G}$ when clear from the context.

We similarly define the PAC-RL for finite-horizon reachability objectives.

► **Definition 12** (PAC-RL for Finite-horizon Reachability Objectives). *A pair of learning algorithms $(\mathfrak{A}_{\text{Max}}, \mathfrak{A}_{\text{Min}})$ is PAC-RL for finite-horizon reachability objectives if there exists a function f such that for all $\varepsilon, p \in (0, 1)$, all TBSGs $\mathbb{G} = (\mathcal{S}, \mathcal{A}, \delta, \mu)$ with a target set $\mathcal{T}$, and all time-horizons $L \in \mathbb{N}$, taking $N = f(|\mathcal{S}|, |\mathcal{A}|, \frac{1}{p}, \frac{1}{\varepsilon}, L)$, with probability at least $1 - p$, the simulator terminates with a strategy profile $(\pi_{\text{Max}}, \pi_{\text{Min}})$ after at most N procedure calls where both strategies are ε -optimal.*

► **Proposition 13.** *Let π be a strategy profile such that $\text{ETR}_{\mathbb{G}}(\pi) \leq L$. Then, for all $\varepsilon > 0$ we have*

$$\left| \mathbb{P}_{\mu}^{\pi}(\omega \in \text{Reach}_{\frac{L}{\varepsilon}}(\mathcal{T})) - \mathbb{P}_{\mu}^{\pi}(\omega \in \text{Reach}(\mathcal{T})) \right| \leq \varepsilon.$$

Proof. Recall that $\text{ETR}_{\mathbb{G}}(\pi) = \mathbb{E}_{\mu}^{\pi}(\arg \inf_{n \in \mathbb{N}} \mathbb{1}(s_n \in \mathcal{T}) \mid (s_n)_{n \in \mathbb{N}} \cap \mathcal{T} \neq \emptyset)$. By Markov's inequality and $\text{ETR}_{\mathbb{G}}(\pi) \leq L$, we have $\mathbb{P}_{\mu}^{\pi}(\arg \inf_{n \in \mathbb{N}} \mathbb{1}(s_n \in \mathcal{T}) > \frac{L}{\varepsilon} \mid (s_n)_{n \in \mathbb{N}} \cap \mathcal{T} \neq \emptyset) \leq \varepsilon$, which yields the result. ◀

► **Theorem 14.** *If a pair of learning algorithms $(\mathfrak{A}_{\text{Max}}, \mathfrak{A}_{\text{Min}})$ is PAC-RL for finite-horizon reachability objectives, then it is PAC-RL with ECD for reachability objectives.*

Proof Sketch. Let $H = 2(L + 1)/\varepsilon$. Run the finite-horizon PAC-RL algorithm with horizon H and error tolerance $\varepsilon/2$. With probability at least $1 - p$, it returns a profile π^* that is $\varepsilon/2$ -optimal for the finite-horizon reachability game. Since $\text{ECD}_{\mathbb{G}} \leq L$, for every player-Min strategy π_{Min} there exists a best-response π_{Max} of Max such that $\text{ETR}_{\mathbb{G}}(\pi_{\text{Max}}, \pi_{\text{Min}}) \leq L + 1$. By Proposition 13, truncating reachability to horizon $H = 2(L + 1)/\varepsilon$ changes the reachability probability of such a best response by at most $\varepsilon/2$. Therefore, $|\text{Val}_{R(\mathcal{T})}(\mu) - \text{Val}_{R_{\frac{2(L+1)}{\varepsilon}}(\mathcal{T})}(\mu)| \leq \frac{\varepsilon}{2}$. Combining this with the $\varepsilon/2$ -optimality of π^* in the finite-horizon game gives that π^* is ε -optimal for the original reachability objective. Hence PAC-RL for finite-horizon reachability implies PAC-RL with ECD for reachability. ◀

We now discuss several key aspects of the ECD parameter which justify why we use this parameter in this work.

Properties of ECD. our ECD definition has two important properties: First, for every game it is finite. Second, it has a meaningful intuition to bound the horizon of the game, i.e., it captures that for every strategy of player Min, there exists a counter-strategy of player Max such that (i) the counter-strategy is optimal for the reachability objectives with respect to the strategy of player Min; and (ii) the expected time to reach is small.

Estimation of ECD. A related parameter is the minimum non-zero transition probability p_{min} of a TBSG $\mathbb{G}$. This parameter has been used in the context of PAC learning for TBSGs with reachability objectives [2]. The minimum non-zero transition probability p_{min} provides a bound on the expected time to reach the target set $\mathcal{T}$. Indeed, if $p_{\text{min}} > 0$, then for any strategy profile π , we have $\text{ETR}_{\mathbb{G}}(\pi) \leq (1/p_{\text{min}})^{|\mathcal{S}|}$. However, this worst-case bound is exponential, while the ECD parameter can be much smaller. Better bounds require more information about the game. Since we provide the theoretical foundation in this work, model-dependent estimation of this parameter is subject for future work.

Stochastic Shortest Path. A closely-related parameter to ECD is the *stochastic shortest path* (SSP) parameter [3]. The difference between SSP and ECD is that, in SSP, player Max requires to reach the target set $\mathcal{T}$ as soon as possible, while player Min wants to delay the reachability of the target set $\mathcal{T}$ as much as possible. The SSP parameter measures non-reaching plays as having infinite cost. Therefore, this parameter can be infinite. In contrast, our definition guarantees that the parameter is always finite. Finiteness of the parameter is necessary for the reduction to finite-horizon games.

4 PAC-RL for Finite-horizon Reachability

In this section, we present a pair of algorithms ($\text{LeTuReGa}_{\text{Max}}$, $\text{LeTuReGa}_{\text{Min}}$) for PAC-RL of TBSGs with finite-horizon reachability. LeTuReGa stands for Learning Turn-based Reachability Games.

► **Theorem 15.** *The pair of algorithms ($\text{LeTuReGa}_{\text{Max}}$, $\text{LeTuReGa}_{\text{Min}}$) is PAC-RL for finite-horizon reachability objectives with sample complexity $O\left(\frac{|S|^3 L^7 |\mathcal{A}| \log(|S|^2 L^2 / p)}{\varepsilon^3}\right)$.*

Theorems 14 and 15 imply a result for reachability objectives with ECD assumption.

► **Corollary 16.** *The pair of algorithms ($\text{LeTuReGa}_{\text{Max}}$, $\text{LeTuReGa}_{\text{Min}}$) is PAC-RL with ECD for reachability objectives with sample complexity $O\left(\frac{|S|^3 L^7 |\mathcal{A}| \log(|S|^2 L^2 / (p\varepsilon^2))}{\varepsilon^{10}}\right)$.*

Proof. To obtain ε -optimal strategies for the infinite-horizon reachability, we use finite-horizon algorithms with the length of the game $2(L+1)/\varepsilon$ (see the proof of Theorem 14). ◀

Significance. Corollary 16 establishes that reachability becomes PAC learnable in turn-based stochastic games in a decentralized and private information setting under bounded ECD. To the best of our knowledge, this is the first result that (i) handles decentralized and private learning; or (ii) provides explicit polynomial sample complexity bounds.

This section is organized as follows. We first recall some algorithms from bandit learning literature. Then, we describe the LeTuReGa algorithms and finally, we prove Theorem 15.

4.1 Best-Arm Identification

In this subsection, we recall a problem in the bandit learning literature and an optimal solution for it. The *best arm identification bandit learning* problem asks to find an ε -optimal arm with probability $1 - p$. Let $\mathcal{A}$ be a set of arms where each arm is associated with an unknown value $r : \mathcal{A} \rightarrow [0, 1]$. A player can sample an arm $a \in \mathcal{A}$ to obtain a random reward $R \in \{0, 1\}$ such that $\mathbb{E}(R) = r(a)$. We say an algorithm can identify an ε -optimal arm with sample complexity f and confidence $1 - p$, if for any set of arms $\mathcal{A}$ and any $p, \varepsilon \in [0, 1]$, after $f(|\mathcal{A}|, 1/p, 1/\varepsilon)$ samples, with probability at least $1 - p$ it outputs a candidate arm a' such that $|\max_a r(a) - r(a')| \leq \varepsilon$. A basic approach is to sample each arm $\frac{\log(|\mathcal{A}|/p)}{\varepsilon^2}$ times, estimate its unknown value, and then select the best arm. The guarantees follow directly from Hoeffding's bound. A better sample complexity is achieved by the Median Elimination algorithm [9][Theorem 10] See Lemma 17 for the formal statement. This sample complexity matches the lower bound for this problem [16].

► **Lemma 17** ([9][Theorem 10]). *For a set of arms $\mathcal{A}$, a value function $r : \mathcal{A} \rightarrow [0, 1]$, error and confidence $\varepsilon, p \in (0, 1)$, there exists an algorithm which identifies an ε -optimal arm with confidence $1 - p$ and sample complexity $O\left(\frac{|\mathcal{A}| \log(1/p)}{\varepsilon^2}\right)$.*

4.2 Algorithm

This section presents a pair of algorithms for PAC-RL of TBSGs with finite-horizon reachability objectives. First, we give an overview of the main techniques used in the algorithm. We then provide a more detailed description. The pseudocode of the algorithm is provided in Algorithm 1. The correctness of the algorithm is proven in the next subsection.

Algorithm Overview. Firstly, the algorithm *extends the set of states to state-steps*, i.e., it learns which action is good enough for every state-step pair, where step is bounded by the horizon. The algorithm keeps track of *unexplored state-steps*. Initially, all state-steps are considered unexplored except for the target set. The set of unexplored state-steps shrinks over time, and it is *treated as a target* to enhance exploration. In the learning process, the algorithm learns how to visit more state-steps and identifies good-enough actions for each of them. The procedure follows in stages. In each stage, the algorithm constructs a new strategy by *backward induction*. For each state-step, it uses a *best arm identification routine*, which proposes a local ε -optimal action with high confidence. The algorithm always learns *one step at a time*, fixing the strategy in the rest of the game. At the end of the induction, a candidate strategy is constructed. This strategy is then used to discover new state-steps. If no new state-steps are discovered from the perspective of the player, this player proposes to the simulator to terminate the algorithm with the recently constructed strategy. The procedure terminates only when both players propose to the simulator to terminate.

We now define some notions used in our algorithm.

► **Definition 18** (Expanded Game). *For a given TBSG $\mathbb{G} = (\mathcal{S}, \mathcal{A}, \delta, \mu)$ with a target set $\mathcal{T} \subseteq \mathcal{S}$ and a time horizon $L \in \mathbb{N}$, we define the expanded game $\mathbb{G}' := (\mathcal{S}', \mathcal{A}, \delta', \mu')$ where*

- $\mathcal{S}' := \mathcal{S}'_{\text{Max}} \uplus \mathcal{S}'_{\text{Min}}$ where $\mathcal{S}'_i := \{(s, \ell) : s \in \mathcal{S}_i, \ell \in [L]\}$ for $i \in \{\text{Max}, \text{Min}\}$;
- For all states $s, s' \in \mathcal{S}$, actions $a \in \mathcal{A}$ and steps $\ell, \ell' \in [L]$, the transition function δ' is defined as

$$\delta'((s, \ell), a)(s', \ell') := \begin{cases} \delta(s, a)(s') & \text{if } \ell \in [L-1] \wedge \ell' = \ell + 1 \\ 1 & \text{if } \ell = L \wedge \ell' = L \wedge s' = s \\ 0 & \text{otherwise;} \end{cases}$$

- For all states $s \in \mathcal{S}$ and steps $\ell \in [L]$, the initial distribution μ' is defined as

$$\mu'(s, \ell) := \begin{cases} \mu(s) & \text{if } \ell = 0 \\ 0 & \text{otherwise.} \end{cases}$$

The target set is defined as $\mathcal{T}_{\text{Max}} := \mathcal{T} \times [L]$. We also define a target set for the player Min as $\mathcal{T}_{\text{Min}} := \{(s, L) \mid s \notin \mathcal{T}\}$.

The expanded game is the original game accompanied by a counter. A play starts with the counter value of 0, and in every step the counter is incremented. The counter is bounded by L , which means the state stays invariant after L steps. Consequently, the expanded game is equivalent to the original game for the finite-horizon L . Moreover, the fact that the state is not changed after L steps implies that the finite-horizon variant has the same value as the infinite-horizon for the expanded game. Thus, we obtain the following result.

► **Proposition 19.** *For a given TBSG $\mathbb{G} = (\mathcal{S}, \mathcal{A}, \delta, \mu)$ with a target set $\mathcal{T} \subseteq \mathcal{S}$ and a time horizon $L \in \mathbb{N}$, we have $\text{Val}_{R_L(\mathcal{T})}^{\mathbb{G}} = \text{Val}_{R_L(\mathcal{T}_{\text{Max}})}^{\mathbb{G}'} = \text{Val}_{R(\mathcal{T}_{\text{Max}})}^{\mathbb{G}'}$.*

► **Remark 20.** Recall that positional strategies are as powerful as general strategies for TBSGs with reachability objectives. Thus, by Proposition 19, we only consider positional strategies in the expanded game, and we need to consider Markovian strategies in the original game.

► **Definition 21.** Given an expanded TBSG game $\mathbb{G} = (\mathcal{S}', \mathcal{A}, \delta', \mu')$ and two disjoint sets $U, V \subseteq \mathcal{S}'$, we define $\text{NotUntil}(U, V) := \{\omega \in (\mathcal{S}' \times \mathcal{A})^* \times \mathcal{S}' : \exists \ell \geq 0, s_\ell \in V, \forall j \in [\ell - 1] : s_j \notin U\}$. Intuitively, $\text{NotUntil}(U, V)$ is the set of finite plays that avoid reaching any state from U until a state from V is reached.

In the algorithm, we use constants which we define below.

► **Definition 22 (Constants).** Let $\mathbb{G} = (\mathcal{S}, \mathcal{A}, \delta, \mu)$ be a TBSG with finite-horizon L and $p, \varepsilon \in [0, 1]$ be the confidence and error of the PAC guarantees. We define the constants used in the algorithm as follows.

- $\varepsilon_{\text{emp}} := \frac{\varepsilon}{8|\mathcal{S}|L}$;
- $\varepsilon_{\text{bai}} := \frac{\varepsilon}{2L}$;
- C is the constant given by the best arm identification algorithm constant that is implicitly present in the O -notation [9][Theorem 10]; and
- $K := \frac{C|\mathcal{A}| \log(|\mathcal{S}|^2 L^2 / p)}{\varepsilon_{\text{emp}} \varepsilon_{\text{bai}}^2}$.

We are now able to explain the algorithms in detail.

Algorithm Details. Pseudocode of the algorithms is given in Algorithm 1. We describe the algorithm for player $i \in \{\text{Max}, \text{Min}\}$. The algorithm starts by initialising the set of unexplored state-steps U_i^0 to be any state owned by player i which is not in the last step (Line 1) and the strategies $\pi_i^0, \dots, \pi_i^{|\mathcal{S}'|}$ (Line 2). A strategy π_i^q is constructed in the stage q using backward induction (Line 27) and it is used in all of the following stages $q + 1, q + 2, \dots$ for the purpose of exploration (Line 8). After the initialisation, the algorithm runs at most $|\mathcal{S}'|$ stages (Line 4) and in each stage it performs backward induction on the length of the game L (Line 6). In a stage q , after it performs an induction, it checks whether the set of unexplored state-steps has shrunk or not (Line 31). If not, it means the strategy learnt in the stage $q - 1$ did not explore anything new, which makes it a good candidate for an ε -optimal strategy. However, this holds only if the set of unexplored states is unchanged for both players in the same stage, which results in the termination of the procedure. The backward induction is split into a sampling phase (Line 7 to Line 18) and an analysis phase (Line 22 to Line 27). Let the stage be q and the level of induction be ℓ . In the sampling phase, the algorithm learns the best action for a state-step (s, ℓ) where $s \in \mathcal{S}_i$. To achieve this, the algorithm uses formerly constructed strategies $\pi_i^0, \dots, \pi_i^{q-1}$ in the first $\ell - 1$ steps of the game. In the step ℓ it plays according to a best-arm identification routine that tries various actions to determine the best one with high confidence. In the steps $\ell + 1$ onwards, it plays according to the currently learnt strategy π_i^q which is being inductively constructed (Line 11). The algorithm samples qK plays, that is, K plays for every formerly constructed strategy π_i^r for $r \in [q - 1]$ (Line 12). The algorithm tracks whether player i is successful in a particular play (Line 16). For player Max, this happens when the play reaches a target state or an unexplored state-step. For player Min, it happens either when the target states are completely avoided, or when an unexplored state-step is reached before a target state is reached. For every state s , the algorithm keeps track of how many times the best arm identification routine was called for that state (Line 18). In the analysis phase of the induction, the algorithm changes the strategy π_i^q to play according to the result of the best-arm identification routine for the step ℓ in a state s (Line 27). This happens only if the state was visited a sufficient number of times (Line 26). Furthermore, if a state s was visited even higher number of times (Line 23), the state-step (s, ℓ) is removed from the set of unexplored state-steps (Line 24).

Comparison with existing work. In the algorithm design, we drew inspiration from [8]. However, our work differs significantly from theirs. First, we consider turn-based games with reachability objectives where PAC-RL guarantees are impossible in general, while they consider concurrent discounted-sum games which are easy in PAC learning. Second, they use an adversarial bandit learning routine, while we use a best-arm identification routine. Third, our approach removes the need to estimate the visitation distribution, it is simpler in general, and mainly, the complexity of our algorithm is more efficient than theirs.

4.3 Proof of Theorem 15

Overview of the proof. We proceed as follows: (a) we prove that the learning simulation terminates (Lemma 23); (b) we define some notations that we use in the next steps (Definition 25); (c) we define an event (Definition 28) that occurs with high probability (Lemma 29); and (d) we show that under this event, the strategy profile proposed by the algorithms is ε -optimal for the finite-horizon reachability game (Lemmas 30–32).

► **Lemma 23.** *The learning simulation of algorithms (LeTuReGa_{Max}, LeTuReGa_{Min}) terminates after at most $|S'|$ stages.*

Proof. At the end of every stage q , if $U_i^q = U_i^{q-1}$, then the algorithm LeTuReGa_i proposes π_i^{q-1} as the candidate strategy. The learning simulation terminates (Line 32) if both algorithms propose a candidate strategy. Otherwise, at least one state-step has to be removed either from U_{Max}^q or U_{Min}^q . However, this cannot happen more than $|S'|$ times as $|U_{\text{Max}}^0 \cup U_{\text{Min}}^0| \leq |S'|$ which means the simulation terminates after at most $|S'|$ stages. ◀

► **Definition 24.** *We define $f \in \{1, \dots, |S'|\}$ as the stage $q - 1$, i.e., the second-to-last stage before termination.*

Notions of V and Q values. We define the value $V_{B,W}^\pi(s, \ell)$ which is the probability of reaching a state-step from $W \subseteq S'$ while avoiding $B \subseteq S'$, starting in a state-step (s, ℓ) and playing according to the expanded game $\mathbb{G}'$ and a strategy profile π . The value $Q_{B,W}^\pi(s, \ell, a)$ is defined similarly to $V_{B,W}^\pi(s, \ell)$ but the first action taken is a .

► **Definition 25.** *Let $\mathbb{G}' = (S', \mathcal{A}, \delta', \mu')$ be an expanded game for an original game $\mathbb{G} = (S, \mathcal{A}, \delta, \mu)$. Let $B, W \subseteq S'$ be two disjoint subsets and π a positional strategy profile for the expanded game. For all states $s \in S$ and steps $\ell \in [L]$, we inductively define*

$$V_{B,W}^\pi(s, \ell) := \begin{cases} 1 & \text{if } (s, \ell) \in W \\ 0 & \text{if } (s, \ell) \in B \text{ or } (s, \ell) = (s, L) \notin W \\ Q_{B,W}^\pi(s, \ell, \pi(s, \ell)) & \text{otherwise} \end{cases}$$

and for all states $s \in S$, steps $\ell \in [L - 1]$, and actions $a \in \mathcal{A}$, we have

$$Q_{B,W}^\pi(s, \ell, a) := \sum_{s' \in S} \delta(s, a)(s') V_{B,W}^\pi(s', \ell + 1).$$

We also denote $V_{B,W}^\pi := \sum_{s \in S} \mu(s) V_{B,W}^\pi(s, 0)$.

Proposition 26 connects these notions with the values computed by the algorithms.

► **Proposition 26.** *For all positional strategy profiles π for the expanded game, all states $s \in S$, and all sets $B, W \subseteq S'$, the following statements hold.*

■ **Algorithm 1** Algorithm LeTuReGa_i for player $i \in \{\text{Max}, \text{Min}\}$.

Data: $\mathcal{S}_i, |\mathcal{S}|, \mathcal{A}, p, \varepsilon, L$

```

1  $U_i^0 \leftarrow \mathcal{S}_i \times [L-1]$  ; // set of unexplored states
2  $\pi_i^q \leftarrow \pi^{\text{uniform}}$  for all  $q \in 0, 1, \dots, |\mathcal{S}'|$  ; // memoryless uniform strategy
3  $C_i^q(s, \ell) \leftarrow 0$  for all  $q \in [|\mathcal{S}'|], (s, \ell) \in \mathcal{S}'$  ; // visit counter of  $(s, \ell)$  in stage  $q$ 
4 for  $q \in 1, \dots, |\mathcal{S}'|$  do
5    $U_i^q \leftarrow U_i^{q-1}$ ;
6   for  $\ell \in L-1, \dots, 1$  do
7     Initialise a best-arm identification (BAI) routine for all  $s \in \mathcal{S}_i$  with
        $(\varepsilon_{bai}, \frac{p}{|\mathcal{S}^2|L^2})$  parameters ;
8     for  $\pi_i \in \pi_i^0, \dots, \pi_i^{q-1}$  do
9       for  $k \in 1, \dots, K$  do
10        Let  $\pi_i^{BAI}(s)$  be the action that the BAI routine wants to sample at the
        state  $s$ , for all  $s \in \mathcal{S}_i$ ;
11        Define a positional strategy for any  $s \in \mathcal{S}_i, j \in [L]$ 
        
$$\pi'_i(s, j) \leftarrow \begin{cases} \pi_i(s, j) & \text{if the step } j < \ell \\ \pi_i^{BAI}(s) & \text{if the step } j = \ell; \\ \pi_i^q(s, j) & \text{if the step } j > \ell \end{cases}$$

12        Interact with the simulator  $\mathbb{M}$  to sample a play
         $(s_1, a_1, s_2, a_2, \dots, s_{J-1}, a_{J-1}, s_J)$  using the strategy  $\pi'_i$ . Notice that
         $J < L$  only if  $s_J \in \mathcal{T}$ ;
13        if  $J > \ell$  and  $s_\ell \in \mathcal{S}_i$  then
14           $\text{Play} \leftarrow ((s_{\ell+1}, \ell+1), a_{\ell+1}, \dots, (s_J, J))$ ;
15           $\text{Winning} \leftarrow \begin{cases} \text{NotUntil}(\emptyset, U_{\text{Max}}^q \cup \mathcal{T}_{\text{Max}}) & \text{if } i = \text{Max} ; \\ \text{NotUntil}(\mathcal{T}_{\text{Max}}, U_{\text{Min}}^q \cup \mathcal{T}_{\text{Min}}) & \text{if } i = \text{Min} ; \end{cases}$ 
16           $\text{Result} \leftarrow \mathbb{1}(\text{Play} \in \text{Winning})$ ;
17          Update the BAI routine at the state  $s_\ell$  with  $(a_\ell, \text{Result})$  ;
          // Note that  $a_\ell$  was requested by  $\pi_i^{BAI}(s_\ell)$ 
18          Increment  $C_i^q(s_\ell, \ell)$ ;
19        end
20      end
21    end
22    for  $s \in \mathcal{S}_i$  do
23      if  $C_i^q(s, \ell) \geq 3K\varepsilon_{emp}$  then
24         $U_i^q \leftarrow U_i^q / \{(s, \ell)\}$ ;
25      end
26      if  $C_i^q(s, \ell) \geq K\varepsilon_{emp}$  then
27        Make  $\pi_i^q$  to play the action suggested by BAI at the state  $(s, \ell)$ ;
28      end
29    end
30  end
31  if  $U_i^q = U_i^{q-1}$  then
32    Call  $\mathbb{M}.propose(\pi_i^{q-1})$ ;
33  end
34 end

```

■ For all steps $\ell \in [L]$, we have

$$V_{B,W}^\pi(s, \ell) = \mathbb{P}_{(s,\ell)}^\pi((s_\ell, a_\ell, \dots, s_L) \in \text{NotUntil}(B, W)) ;$$

■ for all steps $\ell \in [L-1]$ and actions $a \in \mathcal{A}$, we have

$$Q_{B,W}^\pi(s, \ell, a) = \mathbb{P}_{\delta'((s,\ell),a)}^\pi((s_{\ell+1}, a_{\ell+1}, \dots, s_L) \in \text{NotUntil}(B, W)) .$$

Proof Sketch. The proof is by backward induction on ℓ . At the last step, the claim is immediate from the definition. For $\ell < L$, by the induction hypothesis and the law of total probability, the quantity $Q_{B,W}^\pi(s, \ell, a) = \sum_{s' \in \mathcal{S}} \delta(s, a)(s') V_{B,W}^\pi(s', \ell+1)$ is the probability of satisfying $\text{NotUntil}(B, W)$ after taking action a . Then $V_{B,W}^\pi(s, \ell)$ follows directly because it is 0 on B , 1 on W , and otherwise equals $Q_{B,W}^\pi(s, \ell, \pi(s, \ell))$. Thus both recursive definitions coincide with the reach-avoid probabilities. ◀

Event E . We define an event E which is a collection of conditions. To do so, we first define some useful notations.

► **Definition 27.** We define $C^q(s, \ell) := \sum_{i \in \{\text{Max}, \text{Min}\}} C_i^q(s, \ell)$ and $U^q := U_{\text{Max}}^q \cup U_{\text{Min}}^q$ for all $q \in [f], s \in \mathcal{S}, \ell \in [L]$. Intuitively, $C^q(s, \ell)$ is the number of times the pair (s, ℓ) is visited in the stage q . By U^q , we denote all unexplored state-step pairs.

► **Definition 28.** After both algorithms terminate, we say an event E happens if for all stages $q \in [f]$ and every state $(s, \ell) \in \mathcal{S}'$, all of the following conditions hold:

(a)

$$\left| C^q(s, \ell) - K \sum_{r \in [q-1]} V_{\emptyset, \{(s, \ell)\}}^{\pi^r} \right| \leq K \varepsilon_{\text{emp}} ; \quad (1)$$

(b) if $C^q(s, \ell) \geq 3K \varepsilon_{\text{emp}}$ then $C^{q'}(s, \ell) \geq K \varepsilon_{\text{emp}}$ for all stages $q' \in \{q, \dots, f\}$; and

(c) for all $a \in \mathcal{A}$, if $s \in \mathcal{S}_{\text{Max}}$ then

$$Q_{\emptyset, U_{\text{Max}}^q \cup \mathcal{T}_{\text{Max}}}^{\pi^q}(s, \ell, a) - Q_{\emptyset, U_{\text{Max}}^q \cup \mathcal{T}_{\text{Max}}}^{\pi^q}(s, \ell, \pi_{\text{Max}}^q(s, \ell)) < \varepsilon_{\text{bai}}$$

and if $s \in \mathcal{S}_{\text{Min}}$ then

$$Q_{\mathcal{T}_{\text{Max}}, U_{\text{Min}}^q \cup \mathcal{T}_{\text{Min}}}^{\pi^q}(s, \ell, a) - Q_{\mathcal{T}_{\text{Max}}, U_{\text{Min}}^q \cup \mathcal{T}_{\text{Min}}}^{\pi^q}(s, \ell, \pi_{\text{Min}}^q(s, \ell)) < \varepsilon_{\text{bai}} .$$

► **Lemma 29.** The probability of the event E is at least $1 - p$.

Proof Sketch. The event E consists of three parts. It is enough to show that (a) and (c) hold with probability at least $1 - p/2$, while (b) follows deterministically from (a). A union bound then gives the result. For part (a), fix a stage q and a state-step (s, ℓ) . The count $C_q(s, \ell)$ is exactly the total number of visits to (s, ℓ) in the K samples taken for each earlier profile π^r . Hence it can be written as a sum of independent indicators whose expectation is $K \sum_{r \in [q-1]} V_{\emptyset, \{(s, \ell)\}}^{\pi^r}$. Applying Hoeffding's inequality shows that $C_q(s, \ell)$ concentrates around this expectation within $K \varepsilon_{\text{emp}}$, and taking a union bound over all stages, steps, and states gives probability at least $1 - p/2$. For part (c), once $(s, \ell) \notin U_i^q$, parts (a) and (b) ensure that the best-arm identification routine at (s, ℓ) has been called enough for its $(\varepsilon_{\text{bai}}, p/(|S|^2 L^2))$ guarantee. Therefore, by Proposition 26, the selected action is ε_{bai} -optimal with the stated confidence. Again, a union bound gives probability at least $1 - p/2$. ◀

Near-optimality of π^f . In the following, we condition on the event E . Lemma 30 bounds the probability of reaching the set of state-steps U^f under the strategy profile π^f . Lemma 31 states that π^f is $L\varepsilon_{emp}$ -optimal when treating U^f as a target. Lemma 32 combines the two lemmas to show that the strategy profile π^f is ε -optimal in the expanded game with finite-horizon reachability objectives.

► **Lemma 30.** *Under the event E , we have $V_{\emptyset, U^f}^{\pi^f} \leq 4|\mathcal{S}'|\varepsilon_{emp}$.*

Proof. First, the termination condition implies that $U^f = U^{f+1}$. For any $(s, \ell) \in \mathcal{S}'$, we have that $(s, \ell) \in U^f$ iff $(s, \ell) \in U^{f+1}$ if and only if (s, ℓ) has not been removed from U_{Max}^{f+1} nor U_{Min}^{f+1} . Therefore, for all $(s, \ell) \in U^f$, we have $C^q(s, \ell) < 3K\varepsilon_{emp}$ for all stages $q \in \{1, \dots, f+1\}$. Using the property (a) of the event E for the stage $f+1$, we get that $\sum_{r \in [f]} V_{\emptyset, \{(s, \ell)\}}^{\pi^r} \leq 4\varepsilon_{emp}$ which implies in particular that $V_{\emptyset, \{(s, \ell)\}}^{\pi^f} \leq 4\varepsilon_{emp}$. Using the fact that $V_{\emptyset, U^f}^{\pi^f} \leq \sum_{(s, \ell) \in U^f} V_{\emptyset, \{(s, \ell)\}}^{\pi^f}$ we obtain the desired inequality. ◀

► **Lemma 31.** *Under the event E , the following statements hold.*

■ *For all strategies π_{Max} for the player Max, we have*

$$V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f} - V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi^f} \leq L\varepsilon_{bai};$$

■ *for all strategies π_{Min} of the player Min, we have*

$$V_{\mathcal{T}_{\text{Max}}, U_{\text{Min}}^f \cup \mathcal{T}_{\text{Min}}}^{\pi_{\text{Max}}^f, \pi_{\text{Min}}} - V_{\mathcal{T}_{\text{Max}}, U_{\text{Min}}^f \cup \mathcal{T}_{\text{Min}}}^{\pi^f} \leq L\varepsilon_{bai}.$$

Proof Sketch. Use backward induction on ℓ and prove the stronger bound

$V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f}(s, \ell) - V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi^f}(s, \ell) \leq (L - \ell)\varepsilon_{bai}$. The base case is immediate because both values are fixed. For the induction step, the continuation error is bounded by $(L - \ell - 1)\varepsilon_{bai}$ by the induction hypothesis. If $s \in S_{\text{Max}}$, there is at most one additional local error from the choice of $\pi_{\text{Max}}^f(s, \ell)$, and event $E(c)$ gives that this action is ε_{bai} -optimal. Hence the total loss is at most $(L - \ell)\varepsilon_{bai}$. If $s \in S_{\text{Min}}$, no extra local loss is incurred, so the same bound follows. Evaluating this at the initial distribution gives the first inequality. The second inequality follows from symmetric arguments. ◀

► **Lemma 32.** *Under the event E , the following statements hold.*

$$\sup_{\pi_{\text{Max}} \in \Pi_{\text{Max}}} \mathbb{P}_{\mu}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f}(\omega \in \text{Reach}_L(\mathcal{T})) - \text{Val}_{R_L}(\mu) \leq \varepsilon;$$

$$\text{Val}_{R_L}(\mu) - \inf_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \mathbb{P}_{\mu}^{\pi_{\text{Max}}^f, \pi_{\text{Min}}}(\omega \in \text{Reach}_L(\mathcal{T})) \leq \varepsilon.$$

Proof Sketch. By Proposition 26, finite-horizon reachability can be written using the V -values: $\mathbb{P}_{\mu}^{\pi}(\omega \in \text{Reach}_L(\mathcal{T})) = V_{\emptyset, \mathcal{T}_{\text{Max}}}^{\pi} = 1 - V_{\mathcal{T}_{\text{Max}}, \mathcal{T}_{\text{Min}}}^{\pi}$. For Max, enlarge the target set from $\mathcal{T}_{\text{max}}$ to $U_{\text{max}}^f \cup \mathcal{T}_{\text{max}}$. Lemma 31 says that π^f is $L\varepsilon_{bai}$ -optimal for this auxiliary target, and Lemma 30 says that the reachability probability to the unexplored set U^f under π^f is at most $4|\mathcal{S}'|\varepsilon_{emp}$. Hence, for any player-Max strategy π_{Max} , we have $V_{\emptyset, \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f} - V_{\emptyset, \mathcal{T}_{\text{Max}}}^{\pi^f} \leq L\varepsilon_{bai} + 4|\mathcal{S}'|\varepsilon_{emp} \leq \varepsilon$. The inequality for Min is proven analogously. Use the auxiliary target $U_{\text{min}}^f \cup \mathcal{T}_{\text{min}}$, apply Lemma 31, and then remove the unexplored-set error using Lemma 30. Thus the strategies of both players are ε -optimal in the finite-horizon game. ◀

Proof of Theorem 15. We prove the correctness and sample complexity of the algorithms. *Correctness.* By Lemma 23, the algorithms terminate and by Lemma 32, the algorithms output ε -optimal strategies. Thanks to Proposition 19, the value of the original game and the expanded game coincide.

Sample Complexity. We bound the number of procedure calls to the simulator $\mathcal{M}$. Every sampled play at Line 12 corresponds to at most $L + 1$ procedure calls (either $\mathcal{M}.step(\cdot)$ or $\mathcal{M}.reset()$). Furthermore, there are at most $|\mathcal{S}'| + 1$ number of $\mathcal{M}.propose(\cdot)$ calls. We bound the number of sampled plays at Line 12. The loop at Line 4 is iterated at most $|\mathcal{S}|L$ number of times, the loop at Line 6 at most L times, the loop at Line 8 at most $|\mathcal{S}|L$ times, the loop at Line 9 at most K times, which means the number of plays sampled is bounded by

$$\begin{aligned} |\mathcal{S}|L \cdot L \cdot |\mathcal{S}|L \cdot K &= |\mathcal{S}|^2 L^3 \frac{C|\mathcal{A}| \log(|\mathcal{S}|^2 L^2 / p)}{\varepsilon_{emp} \varepsilon_{bai}^2} \\ &= 32|\mathcal{S}|^3 L^6 \frac{C|\mathcal{A}| \log(|\mathcal{S}|^2 L^2 / p)}{\varepsilon^3} \in O\left(\frac{|\mathcal{S}|^3 L^6 |\mathcal{A}| \log(|\mathcal{S}|^2 L^2 / p)}{\varepsilon^3}\right) \end{aligned}$$

Multiplying this by $L + 1$ and adding $|\mathcal{S}'| + 1$ yields the result. $\blacktriangleleft$

Concluding Remarks. In this work, we consider the PAC learning of turn-based stochastic games with reachability objectives. We provide algorithms that ensure learning: (a) with private information; and (b) decentralized setting. Moreover, we generalize the ECD parameter from MDPs to games and establish a polynomial-sample complexity bound with respect to the number of states, actions, ECD parameter, and inverses of error tolerance and failure probability. This framework suggests several interesting open problems: (i) extending to concurrent stochastic games; and (ii) the setting where samplings are drawn from an arbitrary state rather than relying on simulator which restarts from the initial distribution.

References

- 1 Rajeev Alur, Suguman Bansal, Osbert Bastani, and Kishor Jothimurugan. A framework for transforming specifications in reinforcement learning. In *Principles of Systems Design*, volume 13660 of *Lecture Notes in Computer Science*, pages 604–624. Springer, 2022. doi:10.1007/978-3-031-22337-2_29.
- 2 Pranav Ashok, Jan Kretínský, and Maximilian Weininger. PAC statistical model checking for markov decision processes and stochastic games. In Isil Dillig and Serdar Tasiran, editors, *CAV 2019, New York City, NY, USA, July 15-18, 2019*, volume 11561 of *Lecture Notes in Computer Science*, pages 497–519. Springer, 2019. doi:10.1007/978-3-030-25540-4_29.
- 3 Dimitri P. Bertsekas and John N. Tsitsiklis. An analysis of stochastic shortest path problems. *Math. Oper. Res.*, 16(3):580–595, 1991. doi:10.1287/MOOR.16.3.580.
- 4 Patrick Billingsley. *Probability and Measure*. Wiley, Hoboken, NJ, USA, 2012.
- 5 Ronen I. Brafman and Moshe Tennenholtz. A near-optimal poly-time algorithm for learning a class of stochastic games. In *IJCAI*, pages 734–739. Morgan Kaufmann, 1999. URL: <http://ijcai.org/Proceedings/99-2/Papers/012.pdf>.
- 6 Ashok K. Chandra, Dexter Kozen, and Larry J. Stockmeyer. Alternation. *J. ACM*, 28(1):114–133, 1981. doi:10.1145/322234.322243.
- 7 Anne Condon. The complexity of stochastic games. *Inf. Comput.*, 96(2):203–224, 1992. doi:10.1016/0890-5401(92)90048-K.
- 8 Constantinos Daskalakis, Noah Golowich, and Kaiqing Zhang. The complexity of markov equilibrium in stochastic games. In Gergely Neu and Lorenzo Rosasco, editors, *COLT 2023, 12-15 July 2023, Bangalore, India*, volume 195 of *Proceedings of Machine Learning Research*, pages 4180–4234. PMLR, 2023. URL: <https://proceedings.mlr.press/v195/daskalakis23a.html>.

- 9 Eyal Even-Dar, Shie Mannor, and Yishay Mansour. Action elimination and stopping conditions for the multi-armed bandit and reinforcement learning problems. *J. Mach. Learn. Res.*, 7:1079–1105, 2006. URL: <https://jmlr.org/papers/v7/evendar06a.html>.
- 10 Jie Fu and Ufuk Topcu. Probably approximately correct MDP learning and control with temporal logic constraints. In *Robotics: Science and Systems*, 2014.
- 11 Yuri Gurevich and Leo Harrington. Trees, automata, and games. In *STOC*, pages 60–65. ACM, 1982. doi:10.1145/800070.802177.
- 12 Alexander S. Kechris. *Classical Descriptive Set Theory*. Springer, New York, NY, USA, 1995. doi:10.1007/978-1-4612-4190-4.
- 13 Edon Kelmendi, Julia Krämer, Jan Kretínský, and Maximilian Weininger. Value iteration for simple stochastic games: Stopping criterion and learning algorithm. In *CAV (1)*, volume 10981 of *Lecture Notes in Computer Science*, pages 623–642. Springer, 2018. doi:10.1007/978-3-319-96145-3_36.
- 14 S. Lakshmivarahan and Kumpati S. Narendra. Learning algorithms for two-person zero-sum stochastic games with incomplete information. *Math. Oper. Res.*, 6(3):379–386, 1981. doi:10.1287/MOOR.6.3.379.
- 15 Michael L. Littman. Markov games as a framework for multi-agent reinforcement learning. In *ICML*, pages 157–163. Morgan Kaufmann, 1994. doi:10.1016/B978-1-55860-335-6.50027-1.
- 16 Shie Mannor and John N. Tsitsiklis. The sample complexity of exploration in the multi-armed bandit problem. *J. Mach. Learn. Res.*, 5:623–648, 2004. URL: <https://jmlr.org/papers/volume5/mannor04b/mannor04b.pdf>.
- 17 Mateo Perez, Fabio Somenzi, and Ashutosh Trivedi. A PAC learning algorithm for LTL and omega-regular objectives in mdps. In *AAAI*, pages 21510–21517. AAAI Press, 2024. doi:10.1609/AAAI.V38I19.30148.
- 18 Amir Pnueli. The temporal logic of programs. In *FOCS*, pages 46–57. IEEE Computer Society, 1977. doi:10.1109/SFCS.1977.32.
- 19 Martin L. Puterman. *Markov Decision Processes: Discrete Stochastic Dynamic Programming*. Wiley Series in Probability and Statistics. Wiley, 1994. doi:10.1002/9780470316887.
- 20 Jakub Svoboda, Suguman Bansal, and Krishnendu Chatterjee. Reinforcement learning from reachability specifications: PAC guarantees with expected conditional distance. In *ICML 2024, Vienna, Austria, July 21-27, 2024*. OpenReview.net, 2024. URL: <https://proceedings.mlr.press/v235/svoboda24a.html>.
- 21 Leslie G. Valiant. A theory of the learnable. *Commun. ACM*, 27(11):1134–1142, 1984. doi:10.1145/1968.1972.
- 22 Min Wen and Ufuk Topcu. Probably approximately correct learning in stochastic games with temporal logic specifications. In *IJCAI*, pages 3630–3636. IJCAI/AAAI Press, 2016. URL: <http://www.ijcai.org/Abstract/16/511>.
- 23 Cambridge Yang, Michael L. Littman, and Michael Carbin. Reinforcement learning for general LTL objectives is intractable. *CoRR*, abs/2111.12679, 2021. arXiv:2111.12679.

A Proofs of Section 3

► **Theorem 14.** *If a pair of learning algorithms $(\mathfrak{A}_{\text{Max}}, \mathfrak{A}_{\text{Min}})$ is PAC-RL for finite-horizon reachability objectives, then it is PAC-RL with ECD for reachability objectives.*

Proof. Consider a TBSG $\mathbb{G} = (\mathcal{S}, \mathcal{A}, \delta, \mu)$, error tolerance $\varepsilon \in (0, 1)$, failure probability $p \in (0, 1)$, a parameter $L \in \mathbb{R}$, and a target set $\mathcal{T} \subseteq \mathcal{S}$ such that $\text{ECD}_{\mathbb{G}} \leq L$. Let π^* be a strategy profile outputted by the pair of algorithms $(\mathfrak{A}_{\text{Max}}, \mathfrak{A}_{\text{Min}})$ such that, with probability at least $1 - p$, both strategies are $\frac{\varepsilon}{2}$ -optimal for the game $\mathbb{G}$ with finite-horizon reachability objective to the set $\mathcal{T}$ of length $\frac{2(L+1)}{\varepsilon}$, i.e., the following two inequalities

$$\text{Val}_{R_{\frac{2(L+1)}{\varepsilon}}}(\mathcal{T})(\mu) - \inf_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \mathbb{P}_{\mu}^{(\pi_{\text{Max}}^*, \pi_{\text{Min}})}(\omega \in \text{Reach}_{\frac{2(L+1)}{\varepsilon}}(\mathcal{T})) \leq \frac{\varepsilon}{2} \quad (2)$$

$$\sup_{\pi_{\text{Max}} \in \Pi_{\text{Max}}} \mathbb{P}_{\mu}^{(\pi_{\text{Max}}, \pi_{\text{Min}}^*)}(\omega \in \text{Reach}_{\frac{2(L+1)}{\varepsilon}}(\mathcal{T})) - \text{Val}_{R_{\frac{2(L+1)}{\varepsilon}}}(\mathcal{T})(\mu) \leq \frac{\varepsilon}{2} \quad (3)$$

hold with probability at least $1 - p$. From this point on, we condition on this event.

First, we prove that

$$\left| \text{Val}_{R(\mathcal{T})}(\mu) - \text{Val}_{R_{\frac{2(L+1)}{\varepsilon}}}(\mathcal{T})(\mu) \right| \leq \frac{\varepsilon}{2}.$$

Indeed, since $\text{Reach}_{\frac{2(L+1)}{\varepsilon}}(\mathcal{T}) \subseteq \text{Reach}(\mathcal{T})$, we have

$$\text{Val}_{R_{\frac{2(L+1)}{\varepsilon}}}(\mathcal{T})(\mu) \leq \text{Val}_{R(\mathcal{T})}(\mu). \quad (4)$$

Therefore, we only need to show that

$$\text{Val}_{R_{\frac{2(L+1)}{\varepsilon}}}(\mathcal{T})(\mu) \geq \text{Val}_{R(\mathcal{T})}(\mu) - \frac{\varepsilon}{2}. \quad (5)$$

From Definition 8, we have that for all $\pi_{\text{Min}} \in \Pi_{\text{Min}}$, there exists $\pi_{\text{Max}} \in \text{BR}(\pi_{\text{Min}})$ such that $ETR_{\mathbb{G}}(\pi) \leq ECD_{\mathbb{G}} + 1$. We define a function $\varphi : \Pi_{\text{Min}} \rightarrow \Pi_{\text{Max}}$ such that $ETR_{\mathbb{G}}(\varphi(\pi_{\text{Min}}), \pi_{\text{Min}}) \leq ECD_{\mathbb{G}} + 1$ and $\varphi(\pi_{\text{Min}}) \in \text{BR}(\pi_{\text{Min}})$ for all π_{Min} . For all $\pi_{\text{Min}} \in \Pi_{\text{Min}}$, we have

$$\begin{aligned} \sup_{\pi_{\text{Max}} \in \Pi_{\text{Max}}} \mathbb{P}_{\mu}^{\pi}(\omega \in \text{Reach}_{\frac{2(L+1)}{\varepsilon}}(\mathcal{T})) &\stackrel{(1)}{\geq} \sup_{\pi_{\text{Max}} \in \text{BR}(\pi_{\text{Min}})} \mathbb{P}_{\mu}^{\pi}(\omega \in \text{Reach}_{\frac{2(L+1)}{\varepsilon}}(\mathcal{T})) \\ &\stackrel{(2)}{\geq} \mathbb{P}_{\mu}^{\varphi(\pi_{\text{Min}}), \pi_{\text{Min}}}(\omega \in \text{Reach}_{\frac{2(L+1)}{\varepsilon}}(\mathcal{T})) \\ &\stackrel{(3)}{\geq} \mathbb{P}_{\mu}^{\varphi(\pi_{\text{Min}}), \pi_{\text{Min}}}(\omega \in \text{Reach}(\mathcal{T})) - \frac{\varepsilon}{2} \\ &\stackrel{(4)}{=} \sup_{\pi_{\text{Max}} \in \Pi_{\text{Max}}} \mathbb{P}_{\mu}^{\pi}(\omega \in \text{Reach}(\mathcal{T})) - \frac{\varepsilon}{2} \end{aligned} \quad (6)$$

where (1) follows from $\text{BR}(\pi_{\text{Min}}) \subseteq \Pi_{\text{Max}}$; (2) follows from the definition of sup and since $\varphi(\pi_{\text{Min}}) \in \text{BR}(\pi_{\text{Min}})$; (3) follows from $ETR_{\mathbb{G}}(\varphi(\pi_{\text{Min}}), \pi_{\text{Min}}) \leq ECD_{\mathbb{G}} + 1 \leq L + 1$ and by Proposition 13; (4) follows from the definition of best response and since $\varphi(\pi_{\text{Min}}) \in \text{BR}(\pi_{\text{Min}})$.

Now, we proceed to prove Inequality 5. We have

$$\begin{aligned} \text{Val}_{R_{\frac{2(L+1)}{\varepsilon}}}(\mathcal{T})(\mu) &\stackrel{(1)}{=} \inf_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \sup_{\pi_{\text{Max}} \in \Pi_{\text{Max}}} \mathbb{P}_{\mu}^{\pi}(\omega \in \text{Reach}_{\frac{2(L+1)}{\varepsilon}}(\mathcal{T})) \\ &\stackrel{(2)}{\geq} \inf_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \sup_{\pi_{\text{Max}} \in \Pi_{\text{Max}}} \mathbb{P}_{\mu}^{\pi}(\omega \in \text{Reach}(\mathcal{T})) - \frac{\varepsilon}{2} \stackrel{(3)}{=} \text{Val}_{R(\mathcal{T})}(\mu) - \frac{\varepsilon}{2}, \end{aligned}$$

where (1) follows from the definition of finite-horizon reachability value; (2) follows from Inequality 6; (3) follows from the definition of reachability value.

We now show that π^* is ε -optimal for reachability objectives. For the strategy π_{Max}^* :

$$\begin{aligned} \text{Val}_{R(\mathcal{T})}(\mu) - \inf_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \mathbb{P}_{\mu}^{(\pi_{\text{Max}}^*, \pi_{\text{Min}})}(\omega \in \text{Reach}(\mathcal{T})) \\ &\stackrel{(1)}{\leq} \text{Val}_{R(\mathcal{T})}(\mu) - \inf_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \mathbb{P}_{\mu}^{(\pi_{\text{Max}}^*, \pi_{\text{Min}})}(\omega \in \text{Reach}_{\frac{2(L+1)}{\varepsilon}}(\mathcal{T})) \\ &\stackrel{(2)}{\leq} \frac{\varepsilon}{2} + \text{Val}_{R_{\frac{2(L+1)}{\varepsilon}}}(\mathcal{T})(\mu) - \inf_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \mathbb{P}_{\mu}^{(\pi_{\text{Max}}^*, \pi_{\text{Min}})}(\omega \in \text{Reach}_{\frac{2(L+1)}{\varepsilon}}(\mathcal{T})) \stackrel{(3)}{\leq} \varepsilon, \end{aligned}$$

where (1) follows from $\text{Reach}_{\frac{2(L+1)}{\varepsilon}}(\mathcal{T}) \subseteq \text{Reach}(\mathcal{T})$; (2) follows from Inequality 5; and (3) follows from Inequality 2.

For the strategy π_{Min}^* , we have

$$\begin{aligned} & \sup_{\pi_{\text{Max}} \in \Pi_{\text{Max}}} \mathbb{P}_{\mu}^{(\pi_{\text{Max}}, \pi_{\text{Min}}^*)}(\omega \in \text{Reach}(\mathcal{T})) - \text{Val}_{R(\mathcal{T})}(\mu) \\ & \stackrel{(1)}{\leq} \sup_{\pi_{\text{Max}} \in \Pi_{\text{Max}}} \mathbb{P}_{\mu}^{(\pi_{\text{Max}}, \pi_{\text{Min}}^*)}(\omega \in \text{Reach}(\mathcal{T})) - \text{Val}_{R_{\frac{2(L+1)}{\varepsilon}}(\mathcal{T})}(\mu) \\ & \stackrel{(2)}{\leq} \sup_{\pi_{\text{Max}} \in \Pi_{\text{Max}}} \mathbb{P}_{\mu}^{(\pi_{\text{Max}}, \pi_{\text{Min}}^*)}(\omega \in \text{Reach}_{\frac{2(L+1)}{\varepsilon}}(\mathcal{T})) + \frac{\varepsilon}{2} - \text{Val}_{R_{\frac{2(L+1)}{\varepsilon}}(\mathcal{T})}(\mu) \stackrel{(3)}{\leq} \varepsilon, \end{aligned}$$

where (1) follows from Inequality 4; (2) follows from Inequality 6, (3) follows from Inequality 3. $\blacktriangleleft$

B Proofs of Section 4

► **Proposition 26.** *For all positional strategy profiles π for the expanded game, all states $s \in \mathcal{S}$, and all sets $B, W \subseteq \mathcal{S}'$, the following statements hold.*

■ *For all steps $\ell \in [L]$, we have*

$$V_{B,W}^{\pi}(s, \ell) = \mathbb{P}_{(s, \ell)}^{\pi}((s_{\ell}, a_{\ell}, \dots, s_L) \in \text{NotUntil}(B, W)) ;$$

■ *for all steps $\ell \in [L-1]$ and actions $a \in \mathcal{A}$, we have*

$$Q_{B,W}^{\pi}(s, \ell, a) = \mathbb{P}_{\delta'((s, \ell), a)}^{\pi}((s_{\ell+1}, a_{\ell+1}, \dots, s_L) \in \text{NotUntil}(B, W)) .$$

Proof. We prove both items by an induction on the step ℓ .

Induction Base ($\ell = L$). We have

$$V_{B,W}^{\pi}(s, L) \stackrel{(1)}{=} \mathbb{1}((s, L) \in W) \stackrel{(2)}{=} \mathbb{P}_{(s, L)}^{\pi}((s_L) \in \text{NotUntil}(B, W)) ,$$

where (1) and (2) are by the definitions.

Induction Step ($\ell < L$). We first show the second item of the result and then we prove the first item. We assume the claims hold for $\ell + 1$. We have

$$Q_{B,W}^{\pi}(s, \ell, a) \stackrel{(1)}{=} \sum_{s' \in \mathcal{S}} \delta(s, a)(s') V_{B,W}^{\pi}(s', \ell + 1) ,$$

where (1) follows from the definition of $Q_{B,W}^{\pi}(s, \ell, a)$.

By the induction hypothesis, we have $V_{B,W}^{\pi}(s', \ell + 1) = \mathbb{P}_{(s', \ell+1)}^{\pi}((s_{\ell+1}, a_{\ell+1}, \dots, s_L) \in \text{NotUntil}(B, W))$. Therefore, we have $Q_{B,W}^{\pi}(s, \ell, a) = \sum_{s' \in \mathcal{S}} \delta(s, a)(s') \cdot \mathbb{P}_{(s', \ell+1)}^{\pi}((s_{\ell+1}, a_{\ell+1}, \dots, s_L) \in \text{NotUntil}(B, W))$. Finally, by the law of total probability, the right-hand side is exactly the probability of the same event when the initial state-step is drawn from $\delta'((s, \ell), a)$ (i.e., $(s', \ell + 1)$ is chosen with probability $\delta(s, a)(s')$):

$$\sum_{s' \in \mathcal{S}} \delta(s, a)(s') \cdot \mathbb{P}_{(s', \ell+1)}^{\pi}(\cdot) = \mathbb{P}_{\delta'((s, \ell), a)}^{\pi}(\cdot) ,$$

which proves the second item. We now prove the first item. There are three cases.

Case 1: $(s, \ell) \in B$. Then $V_{B,W}^\pi(s, \ell) = 0$ by Definition 25. Also, by starting the suffix from (s, ℓ) , the condition of the event $\text{NotUntil}(B, W)$ is violated immediately. Hence

$$\mathbb{P}_{(s,\ell)}^\pi((s_\ell, a_\ell, \dots, s_L) \in \text{NotUntil}(B, W)) = 0.$$

Case 2: $(s, \ell) \in W$. Then $V_{B,W}^\pi(s, \ell) = 1$ by Definition 25. By starting the suffix from (s, ℓ) , the event $\text{NotUntil}(B, W)$ holds immediately. Thus

$$\mathbb{P}_{(s,\ell)}^\pi((s_\ell, a_\ell, \dots, s_L) \in \text{NotUntil}(B, W)) = 1.$$

Case 3: $(s, \ell) \notin B \cup W$. Then Definition 25 gives

$$V_{B,W}^\pi(s, \ell) = Q_{B,W}^\pi(s, \ell, \pi(s, \ell)).$$

By the induction hypothesis for the second item applied at step ℓ with state s and action $\pi(s, \ell)$, we have

$$Q_{B,W}^\pi(s, \ell, \pi(s, \ell)) = \mathbb{P}_{\delta'((s,\ell), \pi(s,\ell))}^\pi((s_{\ell+1}, a_{\ell+1}, \dots, s_L) \in \text{NotUntil}(B, W)).$$

Since $(s, \ell) \notin B \cup W$, the right-hand side is the probability of satisfying $\text{NotUntil}(B, W)$ starting from (s, ℓ) under π , after taking the first deterministic action $\pi(s, \ell)$; hence this equals

$$\mathbb{P}_{(s,\ell)}^\pi((s_\ell, a_\ell, \dots, s_L) \in \text{NotUntil}(B, W))$$

which closes the third case.

Combining the three cases proves the first item and completes the proof. $\blacktriangleleft$

► **Lemma 29.** *The probability of the event E is at least $1 - p$.*

Proof. We prove the probability of sub-events (a) and (c) is at least $1 - p/2$ and the sub-event (b) follows directly from the sub-event (a). Taking a union bound yields the result.

Let $q \in [|S'|]$ be a stage and $(s, \ell) \in S'$ be a state.

(a) For a stage $r \in [q - 1]$ and $k \in [K]$, let $X_{r,k}$ be the event of reaching a state-step (s, ℓ) in a play induced by the strategy profile π^r . We define $X := \sum_{r \in [q-1], k \in [K]} X_{r,k}$. By Proposition 26, we have $\mathbb{E}(X) = K \sum_{r \in [q-1]} V_{\emptyset, \{(s,\ell)\}}^{\pi^r}$.

Lines 11 to 18 show that $C^q(s, \ell) = X$ since the value of $C^q(s, \ell)$ only depends on the first $\ell - 1$ steps of the sampled plays which are driven by the strategy profiles $\pi^0, \dots, \pi^{q-1}$.

Thus, we obtain

$$\begin{aligned} \mathbb{P} \left(\left| C^q(s, \ell) - K \sum_{r \in [q-1]} V_{\emptyset, \{(s,\ell)\}}^{\pi^r} \right| \leq K\varepsilon_{\text{emp}} \right) &\stackrel{(1)}{=} \mathbb{P}(|X - \mathbb{E}(X)| \leq K\varepsilon_{\text{emp}}) \\ &\stackrel{(2)}{\geq} 1 - 2\exp(-(K\varepsilon_{\text{emp}})^2) \stackrel{(3)}{\geq} 1 - 2\exp(-K\varepsilon_{\text{emp}}) \stackrel{(4)}{\geq} 1 - \frac{p}{2|S|^2 L^2}, \end{aligned}$$

where (1) follows from $X = C^q(s, \ell)$ and $\mathbb{E}(X) = K \sum_{r \in [q-1]} V_{\emptyset, \{(s,\ell)\}}^{\pi^r}$, (2) follows from Hoeffding's inequality, (3) follows from $K\varepsilon_{\text{emp}} \geq 1$, and (4) follows from $K\varepsilon_{\text{emp}} \geq \log\left(\frac{|S|^2 L^2}{4p}\right)$. Taking the union bound for all stages, steps and states, we get that the probability of the sub-event is $1 - p/2$.

(b) This is a direct consequence of the sub-event (a). Assuming $3K\varepsilon_{emp} \leq C^q(s, \ell)$ we get

$$3K\varepsilon_{emp} \leq C^q(s, \ell) \stackrel{(1)}{\leq} K \sum_{r \in [q-1]} V_{\{(s, \ell)\}}^{\pi^r} + K\varepsilon_{emp} \stackrel{(2)}{\leq} K \sum_{r \in [q'-1]} V_{\{(s, \ell)\}}^{\pi^r} + K\varepsilon_{emp} \stackrel{(3)}{\leq} C^{q'}(s, \ell) + 2K\varepsilon_{emp},$$

where (1) is due to Equation (1), (2) is due to $q' \geq q$, and (3) is due to Equation (1).

(c) We prove only the first claim as the second one is proven analogously.

The statement holds trivially for s such that $(s, \ell) \in U_{\text{Max}}^q \cup \mathcal{T}_{\text{Max}}$. Hence, we assume $(s, \ell) \notin U_{\text{Max}}^q$ which means there exists $q' < q$ such that $C^{q'}(s, \ell) \geq 3K\varepsilon_{emp}$ (by the condition at Line 23). Assuming sub-events (a) and (b), this implies $C^q(s, \ell) \geq K\varepsilon_{emp}$. Hence, the best arm identification routine (Line 17) was called at least $K\varepsilon_{emp} \geq \frac{C|\mathcal{A}|\log(|\mathcal{S}|^2 L^2/p)}{\varepsilon_{bai}^2}$ number of times which is sufficient to obtain $(\varepsilon_{bai}, \frac{p}{|\mathcal{S}|^2 L^2})$ -PAC guarantees on the selected arm [9][Theorem 10]. First, we need to show that for each action $a \in \mathcal{A}$, the samples were independent random variables from the same Bernoulli distribution. This is clear by the definition of the strategy in Line 11 since the part of the play from the step $\ell + 1$ onwards is driven by the strategy profile π^q which does not change during the sampling. Therefore, sampling an action a in a state-step (s, ℓ) corresponds to sampling a Bernoulli random variable with the value defined by Line 14 to Line 16, i.e., the random play that starts in (s, ℓ) satisfies the condition defined in Line 15. By Proposition 26, this is equal to $Q_{\emptyset, U_{\text{Max}}^q \cup \mathcal{T}_{\text{Max}}}^{\pi^q}(s, \ell, a)$. Since the output of the best-arm identification routine is the action $\pi_{\text{Max}}^q(s, \ell)$, we obtain the desired inequality for all $a \in \mathcal{A}$ with probability $1 - \frac{p}{|\mathcal{S}|^2 L^2}$. Taking a union over all stages, steps, and states, we obtain that the probability of the sub-event is $1 - p/2$. $\blacktriangleleft$

► **Lemma 31.** *Under the event E , the following statements hold.*

■ *For all strategies π_{Max} for the player Max, we have*

$$V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f} - V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi^f} \leq L\varepsilon_{bai};$$

■ *for all strategies π_{Min} of the player Min, we have*

$$V_{\mathcal{T}_{\text{Max}}, U_{\text{Min}}^f \cup \mathcal{T}_{\text{Min}}}^{\pi_{\text{Max}}^f, \pi_{\text{Min}}} - V_{\mathcal{T}_{\text{Max}}, U_{\text{Min}}^f \cup \mathcal{T}_{\text{Min}}}^{\pi^f} \leq L\varepsilon_{bai}.$$

Proof. We prove the first claim. The proof of the second claim is similar. Since there always exists an optimal positional strategy for the expanded game (see Remark 20), we assume π_{Max} and π_{Min} to be positional. We prove the claim by an induction on ℓ , i.e., we prove that for all states $s \in \mathcal{S}$, we have

$$V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f}(s, \ell) - V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi^f}(s, \ell) \leq (L - \ell)\varepsilon_{bai}. \quad (7)$$

Induction Base ($\ell = L$). The base case $\ell = L$ is trivial as both values are either 0 or 1 depending on whether $s \in U^f \cup \mathcal{T}_{\text{Max}}$ and the strategies play no role.

Induction Step ($\ell < L$). We assume Equation (7) for $\ell + 1$. There are two cases.

Case $(s, \ell) \notin U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}$. Recall that $\delta(s, a)(s')$ denotes the probability of reaching a state s' from a state s by playing an action a . First, we assume that $s \in \mathcal{S}_{\text{Max}}$. Therefore,

$$\begin{aligned}
& V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f}(s, \ell) - V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi^f}(s, \ell) \\
& \stackrel{(1)}{=} Q_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f}(s, \ell, \pi_{\text{Max}}(s, \ell)) - Q_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi^f}(s, \ell, \pi_{\text{Max}}^f(s, \ell)) \\
& \stackrel{(2)}{=} \sum_{s' \in \mathcal{S}} \delta(s, \pi_{\text{Max}}(s, \ell))(s') V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f}(s', \ell + 1) - Q_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi^f}(s, \ell, \pi_{\text{Max}}^f(s, \ell)) \\
& \stackrel{(3)}{\leq} \sum_{s' \in \mathcal{S}} \delta(s, \pi_{\text{Max}}(s, \ell))(s') V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi^f}(s', \ell + 1) - Q_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi^f}(s, \ell, \pi_{\text{Max}}^f(s, \ell)) \\
& \quad + (L - \ell - 1)\varepsilon_{bai} \\
& \stackrel{(4)}{=} Q_{U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi^f}(s, \ell, \pi_{\text{Max}}(s, \ell)) - Q_{U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi^f}(s, \ell, \pi_{\text{Max}}^f(s, \ell)) + (L - \ell - 1)\varepsilon_{bai} \\
& \stackrel{(5)}{\leq} (L - \ell)\varepsilon_{bai}
\end{aligned}$$

where (1), (2) and (4) is due to Definition 25, (3) is due to inductive assumption, and (5) is due to part (c) of event E . If $s \in \mathcal{S}_{\text{Min}}$, the same argument yields a tighter bound $(L - \ell - 1)\varepsilon_{bai}$.

Case $(s, \ell) \in U^f \cup \mathcal{T}_{\text{Max}}$. This is trivial due to Definition 25.

Altogether, by Definition 25, we have $V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f} = \sum_{s \in \mathcal{S}} \mu(s) V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f}(s, 1)$. $\blacktriangleleft$

► **Lemma 32.** *Under the event E , the following statements hold.*

$$\begin{aligned}
& \sup_{\pi_{\text{Max}} \in \Pi_{\text{Max}}} \mathbb{P}_{\mu}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f}(\omega \in \text{Reach}_L(\mathcal{T})) - \text{Val}_{R_L}(\mu) \leq \varepsilon; \\
& \text{Val}_{R_L}(\mu) - \inf_{\pi_{\text{Min}} \in \Pi_{\text{Min}}} \mathbb{P}_{\mu}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f}(\omega \in \text{Reach}_L(\mathcal{T})) \leq \varepsilon.
\end{aligned}$$

Proof. Observe that for all strategy profiles π , we have

$$\begin{aligned}
\mathbb{P}_{\mu}^{\pi}(\omega \in \text{Reach}_L(\mathcal{T})) &= \mathbb{P}_{\mu'}^{\pi}((s_0, a_0, \dots, s_L) \in \text{NotUntil}(\emptyset, \mathcal{T}_{\text{Max}})) \\
&= 1 - \mathbb{P}_{\mu'}^{\pi}((s_0, a_0, \dots, s_L) \in \text{NotUntil}(\mathcal{T}_{\text{Max}}, \mathcal{T}_{\text{Min}})).
\end{aligned}$$

Therefore, by Proposition 26 Item (1), we only need to prove the following statements:

(i) for all strategies π_{Max} of the player Max, we have $V_{\emptyset, \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f} - V_{\emptyset, \mathcal{T}_{\text{Max}}}^{\pi^f} \leq \varepsilon$; and (ii) for all strategies π_{Min} of the player Min, we have $V_{\mathcal{T}_{\text{Max}}, \mathcal{T}_{\text{Min}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f} - V_{\mathcal{T}_{\text{Max}}, \mathcal{T}_{\text{Min}}}^{\pi^f} \leq \varepsilon$. Indeed, for the first claim, we have

$$\begin{aligned}
& V_{\emptyset, \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f} - V_{\emptyset, \mathcal{T}_{\text{Max}}}^{\pi^f} \\
& \stackrel{(1)}{\leq} V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f} - V_{\emptyset, \mathcal{T}_{\text{Max}}}^{\pi^f} \stackrel{(2)}{=} V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f} - V_{\emptyset, \mathcal{T}_{\text{Max}}}^{\pi^f} - V_{\emptyset, U_{\text{Max}}^f}^{\pi^f} + V_{\emptyset, U_{\text{Max}}^f}^{\pi^f} \\
& \stackrel{(3)}{\leq} V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f} - V_{\emptyset, \mathcal{T}_{\text{Max}}}^{\pi^f} - V_{\emptyset, U_{\text{Max}}^f}^{\pi^f} + 4|S'|\varepsilon_{emp} \stackrel{(4)}{\leq} V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f} - V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi^f} + 4|S'|\varepsilon_{emp} \\
& \stackrel{(5)}{\leq} L\varepsilon_{bai} + 4|S'|\varepsilon_{emp} \stackrel{(6)}{\leq} \varepsilon,
\end{aligned}$$

where (1) is due to $\mathcal{T}_{\text{Max}} \subseteq U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}$, (2) is due to algebraic manipulation, (3) is due to Lemma 30, (4) is due to $V_{\emptyset, U_{\text{Max}}^f \cup \mathcal{T}_{\text{Max}}}^{\pi_{\text{Max}}, \pi_{\text{Min}}^f} \leq V_{\emptyset, \mathcal{T}_{\text{Max}}}^{\pi^f} + V_{\emptyset, U_{\text{Max}}^f}^{\pi^f}$, (5) is due to Lemma 31, and (6) is due to Definition 22. The second claim proof is similar. $\blacktriangleleft$