



Scale-free adaptive planning for deterministic dynamics & discounted rewards

Peter Bartlett, Victor Gabillon, Jennifer Healey, Michal Valko

► To cite this version:

Peter Bartlett, Victor Gabillon, Jennifer Healey, Michal Valko. Scale-free adaptive planning for deterministic dynamics & discounted rewards. International Conference on Machine Learning, 2019, Long Beach, United States. hal-02387484

HAL Id: hal-02387484

<https://inria.hal.science/hal-02387484>

Submitted on 29 Nov 2019

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Scale-free adaptive planning for deterministic dynamics & discounted rewards

Peter L. Bartlett¹ Victor Gabillon² Jennifer Healey³ Michal Valko⁴

Abstract

We address the problem of planning in an environment with deterministic dynamics and stochastic discounted rewards under a limited numerical budget where the ranges of both rewards and noise are unknown. We introduce **PlaTγPOOS**, an adaptive, robust, and efficient alternative to the OLOP (open-loop optimistic planning) algorithm. Whereas OLOP requires a priori knowledge of the ranges of both rewards and noise, **PlaTγPOOS** *dynamically adapts* its behavior to both. This allows **PlaTγPOOS** to be immune to two vulnerabilities of OLOP: failure when given underestimated ranges of noise and rewards and inefficiency when these are overestimated. **PlaTγPOOS** additionally adapts to the global smoothness of the value function. **PlaTγPOOS** acts in a provably more efficient manner vs. OLOP when OLOP is given an overestimated reward and show that in the case of no noise, **PlaTγPOOS** learns exponentially faster.

1. Introduction

We consider the problem of planning in a general *stochastic environment* with *deterministic dynamics* and *discounted rewards*. Our goal is to recommend the best first action for an agent to take from a given state. We envision that the discount factor γ is known and that our learner has a limited allocation of n interactions to spend querying a generative model of the environment. The objective is to maximize the sum of discounted rewards of the best sequence of actions following from the recommended first action. This is equivalent to minimizing the *simple regret*. We introduce the algorithm **PlaTγPOOS**, Planning wiTh γ Plus an Online Optimization Strategy, as a robust and efficient scale-free alternative to the OLOP algorithm (open-loop optimistic planning, Bubeck & Munos, 2010; Leurent & Maillard, 2019) for this

setting. Our algorithm implements a scale-free function optimization strategy similar to Sequ00L (Bartlett et al., 2019) rather than an upper-confidence-bound approach which allows us to efficiently adapt to the problem space *without prior knowledge of the ranges of the noise or the rewards*.

Planning in a stochastic environment is an important setting often modeled by Markov decision processes (MDPs, Puterman, 1994; Bertsekas & Tsitsiklis, 1996). One approach to solving these settings is to find the optimal policy that maximizes the expected sum of rewards and then generate an action recommendation according to that optimal policy. Unfortunately, in most practical settings where we are limited by computational resources, finding this optimal policy is often not possible, especially when the state space becomes large. Therefore, instead of trying to estimate the optimal policy of the MDP, we focus only on finding the *best first action* given our budget. We evaluate the performance of the recommendation in terms of the simple regret, the difference in reward between choosing the optimal first action vs. choosing our recommended first action and then in both cases choosing an optimal sequence of actions following the first action. This metric is often used to evaluate planning strategies that optimize numeric budgets (Bubeck & Munos, 2010; Buşoniu & Munos, 2012; Grill et al., 2016), in contrast to the cumulative regret where we are penalized during the search for querying sub-optimal actions. Once the agent takes the first action and moves to the next state, our evaluation can be repeated with a new budget allocation and the following best first action can be recommended. This allows us to approximately follow an optimal policy, action by action, in an online way. Previously, there have been several strategies proposed on how to efficiently allocate a numeric budget to search for an optimal value in a stochastic space. Many of these have been successfully implemented using methods based on upper confidence bounds (UCBs) such as UCT (Upper Confidence Trees, Kocsis & Szepesvári, 2006). This approach has been proven to be very efficient in practice (Coulom, 2007; Gelly et al., 2006; Silver et al., 2016), however, UCT can badly misbehave on some problems (Coquelin & Munos, 2007) and more theoretically sound approaches have been proposed (Hren & Munos, 2008; Bubeck & Munos, 2010; Buşoniu & Munos, 2012; Feldman & Domshlak, 2014; Szörényi et al., 2014; Kaufmann & Koolen, 2017; Shah et al., 2019). Some of

¹University of California, Berkeley, USA ²Noah's Ark Lab, Huawei Technologies, London, UK ³Adobe Research, San Jose, USA ⁴SequeL team, INRIA Lille - Nord Europe, France. Correspondence to: Victor Gabillon <victor.gabillon@huawei.com>.

these methods are connected to the ones from function optimization (Bubeck et al., 2011; Munos, 2011; Valko et al., 2013) as shown by Munos (2014), however, one key difference is that in planning, as opposed to function optimization, the structure of the reward is a discounted reward, specifically a *sum* of rewards *discounted by factor* γ . This reward structure influences the behavior of the optimizers (Bubeck & Munos, 2010), in particular, the discount factor brings smoothness to the value function which in turn makes it easier to optimize. **PlatYP00S** exploits the effect of the discount factor to efficiently manage an adaptive planning strategy in the face of unknown ranges of noise and rewards.

This adaptive strategy of **PlatYP00S** makes it more robust and efficient in practice than other planning strategies. For example, even though they are theoretically sound, the empirical performance of UCB-based approaches depends on the careful tuning of the upper confidence bound. If the upper confidence bound is too large then the UCB-based learner plays very conservatively by overestimating sub-optimal options for many rounds. Moreover, these UCBs might depend on instance parameters that are simply not known such as the *range of the rewards* and the *range of the noise*. We build on the function optimization approach of Bartlett et al. (2019) that does not use UCBs and obtains improved results over the state-of-the-art by adapting to the problem difficulty with a scale-free approach. **PlatYP00S** adapts this scale-free optimization to planning. This *scale-free* property becomes a desired feature as machine learning gets closer to applications, whether it is online (Ross et al., 2013; Orabona & Pál, 2018) or deep learning (Orabona & Tommasi, 2017), since many parameters are never known.

In terms of planning strategy, the **PlatYP00S** algorithm is an adaptive, robust, and efficient alternative to OLOP. Whereas OLOP *requires the knowledge of where the ranges of both rewards and noise*, **PlatYP00S** dynamically adapts its behavior to the both ranges, as well as some potential additional global smoothness of the value function. Our algorithm’s ability to adapt allows it to avoid failure in cases where the ranges of noise and rewards are underestimated and to act more efficiently in cases where they are over-estimated. **PlatYP00S** recovers the results of OLOP while allowing improvements in various classes of problems.

Our contributions We show that **PlatYP00S**

- adapts its behavior to an unknown range of rewards,
- requires no apriori assumptions or knowledge on noise,
- empirically learns much faster than UCB approaches,
- gets the fast rate of deterministic planning in low noise for all regime; in particular, it learns exponentially faster than OLOP when there happens to be no noise,
- adapts also the global smoothness ρ and ν beyond the base smoothness provided by γ .

We additionally address a realistic constraint where the

agent can only *reset* to the original state and not to any state it wishes. Our results hold for MDPs with deterministic dynamics and can equally be applied to open loop planning problems (as discussed by Munos 2014) where we search for the best sequence of actions, ignoring the actual states that are reached after each action (Bubeck & Munos, 2010).

Related algorithms, where the objective is to find the value of the state rather than to identify the best action, include TrailBlazer (Grill et al., 2016) and StOP (Szörényi et al., 2014). A key difference is that these algorithms are *fixed confidence* and output a value using a small number of samples given an accuracy/probability, whereas our algorithm does exploration under a fixed budget of samples and guarantees *how good* the found action is. Even for simple multi-arm bandits, these two problems have different complexity (Carpentier & Locatelli, 2016) and can only be equivalent under unrealistic side knowledge (Gabillon et al., 2012). These related algorithms are also impractical for our setting. TrailBlazer uses confidence bounds that are humongous and StOP takes exponential time. Similar to OLOP, both also need to know noise and reward ranges.

2. Background

We model our problem with an MDP with state space X , action space A and dynamics such that taking the chosen action a_t at time t deterministically transitions the system from $x_t \in X$ to state $x_{t+1} \triangleq f(x_t, a_t)$ generating a reward $r_t \triangleq r(x_t, a_t) + \varepsilon_t$, with ε_t being the noise. We consider:

deterministic rewards The evaluations are noiseless, that is for all t , $\varepsilon_t \triangleq 0$ and $r_t \triangleq r(x_t, a_t)$.

stochastic rewards The evaluations are perturbed by a noise of range $b \in \mathbb{R}_+$: At any round, ε_t is a random variable, independent from noise at previous rounds,

$$\mathbb{E}[r_t | x_t] \triangleq r(x_t, a_t) \text{ and } |r_t - r(x_t, a_t)| \leq b. \quad (1)$$

We assume that all rewards lie in the interval $[0, R_{\max}]$ and while the state space may be large and possibly infinite, that the action space is finite, with K available actions. We treat an infinite time-horizon problem with discounted rewards where the discount factor ($0 \leq \gamma < 1$) is *known*. For any possible policy $\pi : X \rightarrow A$, we define the value function $V^\pi : X \rightarrow \mathbb{R}$ associated to π as $V^\pi(x) \triangleq \sum \gamma^t r(x_t, \pi(x_t))$, where x_t is the state of the system at time t when starting from x (i.e., $x_0 \triangleq x$) and following policy π . In the next definition, we also define the Q -value function $Q^\pi : X \times A \rightarrow \mathbb{R}$ associated to policy π , for each state-action pair (x, a) , as the value of playing action a in state x and the following π thereafter.

Definition 1. The Q -value function Q^π of policy π is

$$Q^\pi(x, a) \triangleq r(x, a) + \gamma V^\pi(f(x, a)).$$

Notice that $V^\pi(x) = Q^\pi(x, \pi(x))$. We define the *optimal* value function, and *Q-value* function respectively, as $V^*(x) \triangleq \sup_\pi V^\pi(x)$ and $Q^*(x, a) \triangleq \sup_\pi Q^\pi(x, a)$, which corresponds to playing a first and optimally after. From the dynamic programming, we have the Bellman equations (Bertsekas & Tsitsiklis, 1996; Puterman, 1994),

$$\begin{aligned} V^*(x) &= \max_{a \in A} r(x, a) + \gamma V^*(f(x, a)), \\ Q^*(x, a) &= r(x, a) + \gamma \max_{b \in A} Q^*(f(x, a), b). \end{aligned}$$

Let $[a : c] \triangleq \{a, a+1, \dots, c\}$ with $a, c \in \mathbb{N}$, $a \leq c$, and $[a] \triangleq [1 : a]$. Let $\log_d$ be the logarithm in base d , $d \in \mathbb{R}$ and $\log$ without a subscript be the natural logarithm.

2.1. Optimistic planning under finite numerical budget

We assume that we have a generative model of f and r that generates simulated transitions and rewards. We want to make the best possible use of this model in order to recommend a best next action $a(n)$ such that the sum of the rewards resulting from playing $a(n)$ and then optimally afterwards is as close as possible to playing optimally from the beginning. For that purpose, we define the performance loss that we aim to minimize r_n as

$$r_n \triangleq \max_{a \in A} Q^*(x, a) - Q^*(x, a(n)).$$

2.2. The planning tree

For a given initial state x , consider the (infinite) planning tree defined by all possible sequences of actions (thus all possible reachable states starting from x). Let A^∞ be the set of infinite sequences $(a_0, a_1, a_2, \dots)$ where $a_t \in A$. The branching factor of this tree is the number of actions $K \triangleq |A|$. Since the dynamics are deterministic, to each finite sequence $a \in A^d$ of length d we assign a state that is reachable starting from x by following this sequence of d actions. Using standard notation for alphabets, we write $A^0 \triangleq \{\emptyset\}$ and $A^\bullet$ for the set of finite sequences. For $a \in A^\bullet$, we let $h(a)$ be the length of a , and $aA^h \triangleq \{aa', a' \in A^h\}$, where aa' denotes the sequence a followed by a' . We identify the set of finite sequences $a \in A^\bullet$ with the set of nodes of the tree. With $h' \leq h$ and $a \in A^h$, we denote $a_{[h']}$ the sequence of action composed of the h' first actions from a , i.e., $\{a_0, \dots, a_{h'-1}\}$. We fix $a_{[0]} \triangleq \emptyset$. The value $v(a)$ of an infinite sequence $a \in A^\infty$ is the discounted sum of rewards along the trajectory starting from the initial state x and defined by the choice of this sequence of actions,

$$v(a) \triangleq \sum_{t \geq 0} \gamma^t r(x_t, a_t), \text{ where } x_0 = x; x_{t+1} \triangleq f(x_t, a_t).$$

Now, for any finite sequence $a \in A^\bullet$, or node, we define the value $v(a) \triangleq \sup_{a' \in A^\infty} v(aa')$. We write $v^* \triangleq v(\emptyset) = \sup_{a \in A^\infty} v(a)$ for the optimal value at the initial state which

is the root of the tree, $v^* = V^*(x)$. We denote the set of optimal infinite sequence of action as A^* which contains any $a \in A^\infty$ such that $v(a) = v^*$. We note the set of optimal finite sequence of actions of depth h as $A^{*,h}$ which contains any $a \in A^h$ such that $v(a) = v^*$. We also define the u - and b -values for the lower- and upper- bounds on $v(a)$ as

$$u(a) \triangleq \sum_{t=0}^{h(a)-1} \gamma^t r(x_t, a_t), \text{ and } b(a) \triangleq u(a) + \frac{\gamma^{h(a)} R_{\max}}{1 - \gamma}.$$

Indeed, since all rewards are in $[0, R_{\max}]$, we trivially have that $u(a) \leq v(a) \leq b(a)$. At any finite time t an algorithm has opened a set of nodes, which defines the expanded tree $\mathcal{T}_t$. We say the learner *opens* (or *expands*) a node a with m evaluations if uses the generative model f and r to generate m transitions and rewards for the K children nodes aA . In the deterministic reward feedback, $m = 1$. The bounds reported in this paper are in terms of the total number of openings n , instead of evaluations. The number of function evaluations is upper bounded by Kn . $T_{x,a}$ denotes the total number of evaluations allocated to action $a \in A$ in state x . We define, especially for the noisy case, the estimated value of the reward $\hat{r}(x, a)$ of action $a \in A$ in state x . Given the $T_{x,a}$ evaluations $r_1, \dots, r_{T_{x,a}}$, we let $\hat{r}(x, a) \triangleq \frac{1}{T_{x,a}} \sum_{s=1}^{T_{x,a}} r_s$ be the empirical average of rewards obtained at when performing action $a \in A$ in state x . To ease notation, for $a \in A^m$ and $h \leq m$, we write $T_a \triangleq \mathbb{E}[T_{x_{h(a)-1}, a_{h(a)-1}} | x_{t+1} \sim P(\cdot | x_t, a_t), x_0 = x]$ for the number of pulls to the last action in a . Similarly, $\hat{r}_h(a) \triangleq \mathbb{E}[\hat{r}(x_h, a_h) | x_{t+1} \sim P(\cdot | x_t, a_t), x_0 = x]$ and $r_h(a) \triangleq \mathbb{E}[r(x_h, a_h) | x_{t+1} \sim P(\cdot | x_t, a_t), x_0 = x]$. In the case of deterministic dynamics, x_h is such that $x_{t+1} \sim P(\cdot | x_t, a_t)$ and $x_0 \triangleq x$ is a fixed state from which we can sample from if we have a full access to the generative model. Hence, for a finite sequence $a \in A^\bullet$ or a node,

$$\hat{u}(a) \triangleq \sum_{t=0}^{h(a)-1} \gamma^t \hat{r}(x_t, a_t) = \sum_{t=0}^{h(a)-1} \gamma^t \hat{r}_t(a).$$

We assume the existence of at least one $a^* \in A^\infty$ for which $V^*(x) = \sup_{a \in A^\infty} v(a)$ and define a smoothness for v .

Proposition 1. *There exists $\nu \in (0, R_{\max}/(1 - \gamma)]$ and $\rho \in (0, \gamma]$ such that $\forall h \geq 0, \forall a \in A^h, u(a) \geq v(a) - \nu \rho^h$.*

Note that this holds automatically for $\nu = R_{\max}/(1 - \gamma)$ and $\rho = \gamma$. For problems with an extra regularity this may also hold for some $\nu < R_{\max}/(1 - \gamma)$ and $\rho < \gamma$. Note that our results automatically adapt to ρ without knowing its value. Moreover, note that while having a smoothness ρ means having rewards diminishing geometrically with depth with a ratio of ρ , the constant ν is linked to the scale of variation of the V which can often be realistically smaller than $\nu < R_{\max}/(1 - \gamma)$. We now define a measure of the quantity of near-optimal sequences for the smoothness ν, ρ .

Definition 2. For any $\nu > 0$ and $\rho \in (0, 1)$, the **branching factor** $\kappa^v(\nu, \rho)$ with associated constant C , is defined as

$$\kappa^v(\nu, \rho) \triangleq \inf \{ \kappa \geq 1 : \exists C > 1, \forall h \geq 0, \mathcal{N}_h^v(3\nu\rho^h) \leq C\kappa^h \},$$

where $\mathcal{N}_h^v(\varepsilon)$ is the number of nodes $a \in A^h$ of depth h such that $v(a) \geq v^* - \varepsilon$.

We also define a related quantity but different from $\kappa^v(\nu, \rho)$. In particular, we define $\kappa^u(\nu, \rho)$ that uses $\mathcal{N}_h^u(\varepsilon)$ instead of $\mathcal{N}_h^v(\varepsilon)$ where $\mathcal{N}_h^u(\varepsilon)$ is the number of nodes $a \in A^h$ of depth h such that $u(a) \geq v^* - \varepsilon$. Our results use the new quantity $\kappa^u(\nu, \rho)$, recover the previous results using $\kappa^v(\nu, \gamma)$ in the method of Bubeck & Munos (2010) and go beyond them. Indeed, theirs are only formulated for $\rho = \gamma$ while we prove the following claim in Appendix A.

Proposition 2. $\kappa^u(\nu/2, \gamma) \leq \kappa^v(\nu, \gamma) \leq \kappa^u(2\nu, \gamma)$.

3. Optimization vs. planning

Our **PlaTγPOOS** approach is a very close sibling to the flat optimization algorithm, **Stroqu00L** (Bartlett et al., 2019), however, we explain how the structure of the planning setting and the discount factor γ shaped our approach. The optimal application of an optimization algorithm to the planning setting is not straightforward as discussed by Bubeck & Munos (2010). In their Section 2.2, Bubeck & Munos (2010) show that optimization can be applied to the planning problem either in a naïve way or a *good* way. The authors take as an example the uniform planning problem. The naïve and good strategies are evaluated by comparing the uncertainty $|u(a) - \hat{u}(a)|$ of their estimates $\hat{u}(a)$. Both strategies collect rewards identically, evaluating $u(a)$ for all the K^H nodes $a \in A^H$ at a fixed depth $H \triangleq h(a)$ by allocating one episode (of length H) for each a and receiving $r_{h,a}$, $1 \leq h \leq H$. In the naïve version, for all sequences a , the estimation of $u(a)$ uses only the H samples collected in the one episode related to a . Here $\hat{u}(a) \triangleq \sum_{h=0}^{h(a)-1} \gamma^h r_{h,a}$ and $T_{a[h]} = 1$. In contrast, the good planning strategy *reuses* estimates. For two distinct sequences a and a' , the good strategy reuses any sample of $r_{h,a}$ it gets for the estimation of both $u(a)$ and $u(a')$ if $a[h] = a'[h]$. In this case, $\hat{u}(a) \triangleq \sum_{h=0}^{h(a)-1} \gamma^h \frac{1}{K^{H-h}} \sum_{a': a[h]=a'[h]} r_{h,a'}$ and $T_{a[h]} = K^{H-h}$. This comparison is used to demonstrate the advantage of the use of the cross-sequence information to concentrate the estimate of the mean reward associated with each action more efficiently. It is by using this type of cross-sequence information that OLOP is able to obtain a reduced regret over a naïve application of H00 for optimization (Bubeck et al., 2011) to the planning problem.

The previous discussion on cross-sequence information is tied to the case of uniform exploration strategies. Good uniform strategies guarantee $T_{a[h]} = K^{H-h}$ by exploring

Input: n, A
Initialization: open $t_{0,1}; h_{\max} \leftarrow \lfloor n/\log(n) \rfloor$
For $h = 1$ to $h_{\max}$
 open $\lfloor h_{\max}/h \rfloor$ nodes $a_{h,i}$ of depth h
 with largest values $u(a_{h,i})$
Output $x(n) \leftarrow \arg \max_{a_{h,i} \in \mathcal{T}} u(a_{h,i})$

Figure 1. Algorithm for free planning with no reset condition

until a reasonably shallow depth H_u but sampling all K^{H_u} nodes and sharing cross sequence information. In the case of OLOP, H00, or, particularly **Stroqu00L**, only a subset S of size $|S| \ll K^{H_u}$ of the most promising nodes are explored but at a deeper depth $H_s \gg H_u$. Therefore, obtaining a lower bound on the number of sequence of actions at depth H_s that contains a for a given sub-sequence of actions a at depth $h < H_s$ is complex in general. Actually one can design a problem with two actions that would drastically limit the amount of cross-sequence information for the optimal sequence of actions a^* in **Stroqu00L**. As a result, applying **Stroqu00L** with information sharing may not be enough and we chose to algorithmically ensure that a node at depth $h + 1$ will be pulled γ^2 times less than a node at depth h . Indeed, using Chernoff-Hoeffding inequalities, we have $|u(a) - \hat{u}(a)| \leq \sum_{h=0}^{h(a)-1} \gamma^h / \sqrt{T_{a[h]}}$. However, **PlaTγPOOS**, through some simple reparametrization, remains very close to **Stroqu00L** and we leave as an open question whether equivalent theoretical guarantees could be proved directly for **Stroqu00L** applied to planning.

4. Deterministic dynamics and rewards

In this section, we consider a simpler case of deterministic rewards in order to introduce our new ideas. The evaluations are noiseless, that is $\forall t, \varepsilon_t \triangleq 0$ and $r_t \triangleq r(x_t, a_t)$.

In Figure 1, we provide the **Sequ00L** (Bartlett et al., 2019) algorithm applied to planning. In this case, it is straightforward to follow the analysis of **Sequ00L** in order to obtain the same rates of simple regret as the state of the art algorithm OPD for the doubly deterministic case (Hren & Munos, 2008; Munos, 2014), up to logarithmic factors; and get the result¹ of Theorem 1. This direct usage was already discussed by Munos (2014, Section 5.1). Using **Sequ00L** for planning already permits to have an algorithm that does not use the parameter $R_{\max}$ and that adapts to extra smoothness in the value function ν, ρ as discussed Section 2. Note that to obtain similar adaptations to $R_{\max}, \nu$, and ρ , we could have already used S00 (Munos, 2014). However, S00 *does not* come with optimal simple regret (Bartlett et al., 2019).

¹ $\log(n)$ is the n -th harmonic number

Input: n, A
Initialization: open $t_{0,1}$; $h_{\max} \leftarrow \left\lfloor \frac{n}{\log n} \right\rfloor$
For $h=1$ **to** $h_{\max}$
 open $\lfloor h_{\max}/h^2 \rfloor$ nodes $a_{h,i}$ of depth h
 with largest values $u(a_{h,i})$
Output $x(n) \leftarrow \arg \max_{a_{h,i} \in \mathcal{T}} u(a_{h,i})$

Figure 2. Algorithm for constraint planning (restart)

Theorem 1. For any planning problem with associated (ν, ρ) , and branching factor $\kappa \triangleq \kappa^u(\nu, \rho)$, the simple regret of Sequ00L is after n rounds bounded as follows.

- If $\kappa = 1$, then $r_n \leq \nu \rho^{\frac{1}{C} \left\lfloor \frac{n}{\log n} \right\rfloor}$.
- If $\kappa > 1$, then $r_n = \mathcal{O} \left(\nu \left(\frac{\log \kappa}{C} \left\lfloor \frac{n}{\log n} \right\rfloor \right)^{-\frac{\log 1/\rho}{\log \kappa}} \right)$.

On the reset condition In practice, physical constraints can force the exploration to interact with the unknown environment under a reset condition. This means that there exists a starting state x and a trajectory needs to start from this state and following a given policy. Therefore, if we wish to collect a sample from an arbitrary state x' after a reset, we must first reach that state from the starting state x .

Sequ00L is a strategy that explores the MDP deeper and deeper from an initial starting state x . The original Sequ00L strategy did not consider any reset condition, so to make a comparable analysis we will say that to ‘open’ a node a you first need to reach a at a budget cost of $h(a)$ which is equal to the depth of a . This additional cost has the consequence that Sequ00L with reset will not be able to explore as deeply as without. A naïve extension is shown in Figure 2. Under a total limited budget of n , the number of nodes now open at depth h is $\mathcal{O}(n/h^2)$ instead of $\mathcal{O}(n/h)$ and the maximal depth is now of order of $\sqrt{n}$ instead of n . However, this does not influence the simple regret when $\kappa > 1$ by more than numerical constants as shown in Theorem 2. When $\kappa > 1$, the exponentially diminishing simple regret remains but is changed from ρ^n to $\rho^{\sqrt{n}}$.

Theorem 2. For a planning problem with associated (ν, ρ) , and branching factor $\kappa \triangleq \kappa^u(\nu, \rho)$, after n rounds, the simple regret of Sequ00L with reset condition verifies:

- If $\kappa = 1$, then $r_n \leq \nu \rho^{\frac{1}{C} \left\lfloor \frac{n}{\log n} \right\rfloor}$.
- If $\kappa > 1$, then $r_n \leq \mathcal{O} \left(\nu \left(\frac{\log \kappa}{C} \left\lfloor \frac{n}{\log n} \right\rfloor \right)^{-\frac{\log 1/\rho}{\log \kappa}} \right)$.

5. Deterministic dynamics, stochastic rewards

We now describe the **PlatYP00S** algorithm detailed in Figure 3. In the presence of noise, it is natural to evaluate the

Input: n, A
Initialization: open the root node $\emptyset$, $h_{\max}$ times
 $h_{\max} \leftarrow \left\lfloor \frac{n}{2(\log_2 n + 1)^2} \right\rfloor$, $p_{\max} \leftarrow \lfloor \log_2(h_{\max}) \rfloor$
For $h = 1$ **to** $h_{\max}$ ◀ exploration ▶
 For $p = \lfloor \log_2(h_{\max} / \lceil h^2 \gamma^{2h} \rceil) \rfloor$ **down to** 0
 open $\lceil h^{2p} \gamma^{2h} \rceil$ times the at most $\left\lfloor \frac{h_{\max}}{h \lceil h^{2p} \gamma^{2h} \rceil} \right\rfloor$
 non-opened nodes $a^{h,i} \in A^h$ with highest values $\hat{u}(a^{h,i})$ and given $T_{a^{h,i}} \geq \lceil (h-1)2^p \gamma^{2(h-1)} \rceil$
 For $p \in [0 : p_{\max}]$ ◀ cross-validation ▶
 evaluate $(t+1)\gamma^{2t} h_{\max} (1-\gamma^2)^2$ times
 the actions at round t , a_t^p , of the candidates:
 $a^p \leftarrow \arg \max_{a \in A^p : \forall t \in [2:h(a)], T_{a|t} \geq \lceil (t-1)2^p \gamma^{2(t-1)} \rceil} \hat{u}(a)$
 Output $a^n \leftarrow \arg \max_{\{a^p, p \in [0:p_{\max}]\}} \hat{u}(a^p)$

 Figure 3. The **PlatYP00S** algorithm

cells multiple times, not just one time as in the deterministic case. The amount of times a cell should be evaluated to differentiate its value from the optimal value of the function depends on the gap between these two values as well as the range of noise. As we do not want to make *any* assumptions on knowing these quantities, our algorithm tries to be robust to any potential values by not making a fixed choice on the number of evaluations. Intuitively, we do this following a path similar to Stroqu00L (Bartlett et al., 2019) by using a modified version of Sequ00L, denoted Sequ00L(m), that allows us to evaluate cells m times, whereas for Sequ00L, $m = 1$. Evaluating cells more times (m large) leads to a better quality of the mean estimates in each cell, however, as a trade-off, it uses more evaluations per depth. This would normally limit us from exploring deep depths of the partition, however, **PlatYP00S** takes advantage of the knowledge of γ which gives less weight to reward collected deeper in the tree. In order to obtain the concentration results for a node a on $\hat{u}(a) - u(a)$ in Lemma 2, **PlatYP00S** uses a Chernoff-Hoeffding result that gives with high probability, $\hat{u}(a) - u(a) \leq \sqrt{\sum_{h=0}^{h(a)-1} \gamma^{2h} / T_{a|h}}$ and balances the range of confidence intervals at different depths. Therefore, **PlatYP00S** tends to pull less with deeper depth as the number of pulls for a fixed m is $\lceil hm \gamma^{2h} \rceil$ where the additional h factor ensures that the sum of confidence interval until depth h is bounded for any h . **PlatYP00S** then *implicitly* performs $\log n$ instances of Sequ00L(m) each with a number of evaluations of $m = 2^p$, where $p \in [0 : \log n]$.

In Figure 3, remember that ‘opening’ a node means ‘evaluating’ its children actions. The algorithm opens nodes by sequentially diving them deeper and deeper from the root

node $h = 0$ to a maximal depth of $h_{\max}$. At depth h , we allocate, in an almost decreasing fashion, different number of evaluations $\lceil h2^p\gamma^{2h} \rceil$ to the nodes with highest value of that depth, with p starting at $\lfloor \log_2(h_{\max}/h) \rfloor$ down to 0. The best node that has been evaluated at least $\mathcal{O}(h_{\max}/h)$ times is opened with $\mathcal{O}(h_{\max}/h)$ evaluations, the two next best cells that have been evaluated at least $\mathcal{O}(h_{\max}/(2h))$ times are opened with $\mathcal{O}(h_{\max}/(2h))$ evaluations, the four next best cells that have been evaluated at least $\mathcal{O}(h_{\max}/(4h))$ times are opened with $\mathcal{O}(h_{\max}/(4h))$ evaluations and so on, until some $\mathcal{O}(h_{\max}/h)$ next best cells that have been evaluated at least once are opened with one evaluation. More precisely, given, p and h , we open, with $\lceil h2^p\gamma^{2h} \rceil$ evaluations, the $\lfloor h_{\max}/(h \lceil h2^p\gamma^{2h} \rceil) \rfloor$ non-previously-opened nodes $a^{h,i} \in A^h$ with highest values $\hat{u}(a^{h,i})$ and given that $T_{a^{h,i}} \geq \lceil (h-1)2^p\gamma^{2(h-1)} \rceil$. The maximum number of evaluations of any node is $2^{p_{\max}}$, with $2^{p_{\max}} = \mathcal{O}(h_{\max})$ as $p_{\max} \triangleq \lfloor \log_2(h_{\max}) \rfloor$. For each $p \in [0 : p_{\max}]$, the candidate output a^p is the node a with the highest estimated value such that all actions leading to that node have been evaluated in the following way $\forall t \in [2 : h(a)], T_{a[t]} \geq \lceil (t-1)2^p\gamma^{2(t-1)} \rceil$. We set $h_{\max} \triangleq \lfloor n/(2(\log_2 n + 1)^2) \rfloor$.

5.1. Analysis of PLaTγPOOS

$\perp_{h,p}$ is the depth of the deepest opened node, a with at least $\lceil h2^p\gamma^h \rceil$ evaluations such that there is a $a^* \in A^*$ with $a^* \triangleq ab$, with $b \in A^\infty$, at the end of the opening of depth h .

Lemma 1. *For any planning problem with associated (ν, ρ) as in Property 1), on event ξ defined in Appendix C, for any depth $h \in [h_{\max}]$, for any $p \in [0 : \lfloor \log_2(h_{\max}/(h^2\gamma^{2h})) \rfloor]$, we have $\perp_{h,p} = h$ if (1) and (2) simultaneously hold:*

- (1) $b\sqrt{\log(4n/\delta)/2^{p+1}} \leq \nu\rho^h$
- (2) We distinguish cases and express the condition in each:

Case 1) $h2^p\gamma^{2h} \leq 1$:

$$\frac{h_{\max}}{h} = \frac{h_{\max}}{h \lceil h2^p\gamma^{2h} \rceil} \geq C\kappa(\nu, \rho)^h$$

and for all $h' \in [h]$, $\frac{h_{\max}}{h'^{2p+1}\gamma^{2h'}} \geq C\kappa(\nu, \rho)^{h'}$

Case 2) $h2^p\gamma^{2h} \geq 1$:

Case 2.1) $\gamma^2\kappa^u \geq 1$: $\frac{h_{\max}}{h^2\gamma^{2p+1}\gamma^{2h}} \geq C\kappa(\nu, \rho)^h$

Case 2.2) $\gamma^2\kappa^u \leq 1$: $\frac{h_{\max}}{h^2\gamma^{2p+1}} \geq C$

Lemma 1 gives two conditions so that the cell containing a $a^* \in A^*$ is opened at depth h . This holds if (1) PLaTγPOOS opens, with $\lceil h2^p\gamma^{2h} \rceil$ evaluations, more cells at depth h than the number of near-optimal cells at depth h ($h_{\max}/(h^2\gamma^{2p+1}\gamma^{2h}) \geq C\kappa(\nu, \rho)^h$ if $\gamma^2\kappa^u \geq 1$ and $h_{\max}/h2^p \geq C$ if $\gamma^2\kappa^u \leq 1$) and (2) the $\lceil h2^p\gamma^{2h} \rceil$ evaluations

are sufficient to discriminate the empirical average of near-optimal cells from the empirical average of sub-optimal cells ($b\sqrt{\log(4n/\delta)/2^{p+1}} \leq \nu\rho^h$). To state the next theorems, we introduce $\tilde{h}$, $\tilde{h}_1$, and $\tilde{h}_2$ three positive real numbers satisfying respectively the equations:

$$h_{\max}\nu^2\rho^{2\tilde{h}_1} / \left(\tilde{h}_1 b^2 g_{n,b}^{\delta, R_{\max}} \gamma^{2\tilde{h}_1} \right) = C\kappa^{\tilde{h}_1} \text{ and}$$

$$h_{\max}\nu^2\rho^{2\tilde{h}_2} / \left(\tilde{h}_2 b^2 g_{n,b}^{\delta, R_{\max}} \right) = C, \text{ where}$$

$$\tilde{h}_1 = \frac{1}{\log(\gamma^2\kappa/\rho^2)} \log\left(\frac{\bar{n}_1}{\log \bar{n}_1}\right) + o(1) \text{ and}$$

$$\tilde{h}_2 = \frac{1}{\log(1/\rho^2)} \log\left(\frac{\bar{n}_2}{\log \bar{n}_2}\right) + o(1) \text{ with}$$

$$\bar{n}_1 \triangleq \frac{\nu^2 h_{\max} \log(\gamma^2\kappa/\rho^2)}{C b^2 g_{n,b}^{\delta, R_{\max}}}, \bar{n}_2 \triangleq \frac{\nu^2 h_{\max} \log(1/\rho^2)}{C b^2 g_{n,b}^{\delta, R_{\max}}},$$

where $g_{n,b}^{\delta, R_{\max}} \triangleq p_{\max} \log(R_{\max} n^{3/2}/b(1-\delta))$. $\tilde{h}$ is defined similarly in Equation 5 in Appendix D. The quantities $\tilde{h}_1$ and $\tilde{h}_2$ give the respective depths of deepest cell opened by PLaTγPOOS that contains a a^* with high probability in the cases $\gamma^2\kappa \geq 1$ and $\gamma^2\kappa \leq 1$. Additionally, $\tilde{h}_1$ and $\tilde{h}_2$ also let us characterize for which regime of the noise range b we recover results similar to the loss of the deterministic case. Discriminating on the noise regimes, we now state two of our results, Theorem 3 for a high noise and Theorem 4 for a low one. A more exhaustive list of results is in the Appendix D or in the Table 1.

Theorem 3. High-noise regime *If the noise b is high enough to verify both high-noise conditions as defined in the caption of Table 1, then after n rounds, for any problem with associated (ν, ρ) , and branching factor $\kappa \triangleq \kappa^u(\nu, \rho)$, the simple regret of PLaTγPOOS obeys*

$$\mathbb{E}r_n = \begin{cases} \tilde{\mathcal{O}}\left(\left(\frac{n}{b^2}\right)^{-\frac{1}{2}}\right) & \text{if } \gamma^2\kappa \leq 1, \\ \tilde{\mathcal{O}}\left(\left(\frac{n}{b^2}\right)^{-\frac{\log(1/\rho)}{\log(\gamma^2\kappa/\rho^2)}}\right) & \text{if } \gamma^2\kappa > 1. \end{cases}$$

The proofs are in appendix D. They are quite technical but they are simply based on checking the conditions of Lemma 1 under different b, ρ, γ, κ regimes.

Theorem 4. Low-noise regime *If the noise b is low enough to verify both high-noise conditions as defined in the caption of Table 1, then after n rounds, for any problem with associated (ν, ρ) , and branching factor $\kappa \triangleq \kappa^u(\nu, \rho)$, the simple regret of PLaTγPOOS obeys*

$$\mathbb{E}r_n = \begin{cases} \tilde{\mathcal{O}}(\nu\rho^n) & \text{if } \kappa = 1, \\ \tilde{\mathcal{O}}\left(\nu\left(\frac{n}{b^2}\right)^{-\frac{\log(1/\rho)}{\log(\kappa)}}\right) & \text{if } \kappa > 1. \end{cases}$$

	$\gamma^2 \kappa \leq 1$		$\gamma^2 \kappa \geq 1$	
	<i>High noise</i> (ii)	<i>Low noise</i> (ii)	<i>High noise</i> (iii)	<i>Low noise</i> (iii)
<i>High noise</i> (i)	$m_{\nu,\gamma} \left(\frac{n}{b^2}\right)^{-\frac{1}{2}}$	$\nu \rho \sqrt{n}$	$m_{\nu,\gamma} \left(\frac{n}{b^2}\right)^{-\frac{\log(1/\rho)}{\log(\gamma^2 \kappa / \rho^2)}}$	$\nu \left(\frac{n}{b^2}\right)^{-\frac{\log(1/\rho)}{\log(\kappa)}}$
<i>Low noise</i> (i)	$m_{\nu,\gamma} \left(\frac{n}{b^2}\right)^{-\frac{\log(1/\rho)}{\log(\kappa)}}$	$\kappa = 1 : \nu \rho^n$	$m_{\nu,\gamma} \left(\frac{n}{b^2}\right)^{-\frac{\log(1/\rho)}{\log(\kappa)}}$	$\nu \left(\frac{n}{b^2}\right)^{-\frac{\log(1/\rho)}{\log(\kappa)}}$
		$\kappa > 1 : \nu \left(\frac{n}{b^2}\right)^{-\frac{\log(1/\rho)}{\log(\kappa)}}$		

Table 1. Rates of the our upper bounds on the simple regret of **PlaTγPOOS** for various classes. The condition on noise (i) is whether the noise b verifies $\nu^2 \rho^{2\tilde{h}} / (\gamma^{2\tilde{h}} \tilde{h} b^2 g_{n,b}^{\delta, R_{\max}}) \leq 1$. The condition on noise (ii) is whether $\nu^2 \rho^{2\tilde{h}_2} / (b^2 g_{n,b}^{\delta, R_{\max}}) \leq 1$. The condition on noise (iii) is whether $\nu^2 \rho^{2\tilde{h}_1} / (b^2 g_{n,b}^{\delta, R_{\max}}) \leq 1$. Moreover, $m_{\nu,\gamma} \triangleq \max(1/1 - \gamma^2, \nu)$.

Worst-case comparison with OLOP when b is large and known In Table 1, we give our results for various classes of problems depending on the whether $\gamma^2 \kappa \geq 1$, whether $\kappa = 1$ or $\kappa > 1$, and several conditions for the range of the noise b . The results for OLOP were distinguishing the results based on $\gamma^2 \kappa$ being greater or smaller than 1. For these two cases, we recover the same rate in term of n for the simple regret, as displayed in Table 1 with a grey background color, for instance taking $b = 1$ like in OLOP and having $\rho = \gamma$. However, we provide more specific treatment for sub-cases with associated improvements that we list and detail next.

Adaptation to the range of the noise b without a prior knowledge Our analysis shows that **PlaTγPOOS** adapts favorably to the unknown range of noise. Already, in the standard cases discussed above, where the noise is large, our bound already adapts to the amount of noise as it scales with n/b^2 . OLOP requires an estimate $\tilde{b}$ of b and has a regret scaling with $n/\tilde{b}^2$ which is problematic in case of a wrong estimate $\tilde{b} \gg b$. Moreover, we give technical conditions on the range of noise that shows when **PlaTγPOOS** gets improved rates. When $\gamma^2 \kappa \geq 1$, OLOP was already obtaining rates that were the same as the rates of deterministic reward case. Therefore, beyond the n/b^2 improvement and the adaptation to extra smoothness that will be discussed latter, no more rate improvement should be expected. In the case $\gamma^2 \kappa \leq 1$, the improvement are even more striking. When the noise is very low, then contrary to the OLOP rate of $\mathcal{O}(1/\sqrt{n})$, we obtain the *deterministic* rate of OPD (Hren & Munos, 2008; Munos, 2014) which is either $n^{-\log(1/\rho)/\log \kappa}$ or ρ^n . These improved rates *could not* be obtained by OLOP. Indeed, OLOP relies on upper confidence bound (UCB) that uses a range of the noise $\tilde{b}$ as input,

$$\bar{u}(a) = \mathcal{O} \left(\hat{u}(a) + \sum_{h=0}^{h(a)-1} \gamma^h \tilde{b} \sqrt{\frac{1}{T_{a[a]}}} + \frac{R_{\max} \gamma^{h(a)}}{1 - \gamma} \right).$$

This works if $\tilde{b} = b$. However, the true b is unknown. If $b = 0$, using any $\tilde{b} > 0$ will not result in an improved rate.

Adaptation to additional smoothness ν and ρ beyond γ As defined in Section 2, we aim to adapt to the true smoothness ν, ρ of V which can go beyond γ . We show that **PlaTγPOOS** is able to take advantage of ν, ρ in a large portion of cases. In most cases, the rate γ in OLOP is replaced by ρ in **PlaTγPOOS**. In the case where $\gamma^2 \kappa \geq 1$ we have

$$r^{\text{OLOP}} = \mathcal{O}(n^{-\frac{\log(1/\gamma)}{\log(\kappa)}}) \leq \mathcal{O}(n^{-\frac{\log(1/\rho)}{\log(\gamma^2 \kappa / \rho^2)}}) = r^{\text{PlaTγPOOS}}.$$

Adaptation to the deterministic case and $\kappa = 1$ **PlaTγPOOS** adapts to the branching factor κ of the problem that under low noise conditions, leads to an exponentially decreasing simple regret $r_n^{\text{PlaTγPOOS}} = \mathcal{O}(\nu \rho \sqrt{n})$. This is a light-years improvement over OLOP for these conditions, as OLOP's regret is at best $r_n^{\text{PlaTγPOOS}} = \mathcal{O}(n^{-1/2})$. This result is possible because **PlaTγPOOS** explores much deeper than OLOP, as its maximal depth is of order n . Actually, in most scenarios, the actual larger depth explored will be of order $\sqrt{n}$ due to sampling limitations. On the other hand, OLOP can only go $\log n$ deep.

Moreover, $\kappa = 1$ is a common case in planning. Indeed, as discussed by Bubeck & Munos (2010), $\kappa = 1$ is equivalent to having near-optimal dimension $d = 0$ in an optimization task (Munos, 2014) which is a common value as shown by Valko et al. (2013). Therefore, we expect the case when $\gamma \kappa^2 \leq 1$, that is, where we get the most significant improvement other OLOP, to be the most common in practice.

The reset condition As discussed in Section 4, in the stochastic reward case, the effect of the reset condition affects, **PlaTγPOOS** as follows. First, all polynomial rates stay the same. Next, only the exponential rates change. A ρ^n rate without the condition becomes a $\rho \sqrt{n}$ with it. Next, a $\rho \sqrt{n}$ rate becomes a $\rho^{n^{1/3}}$ one.

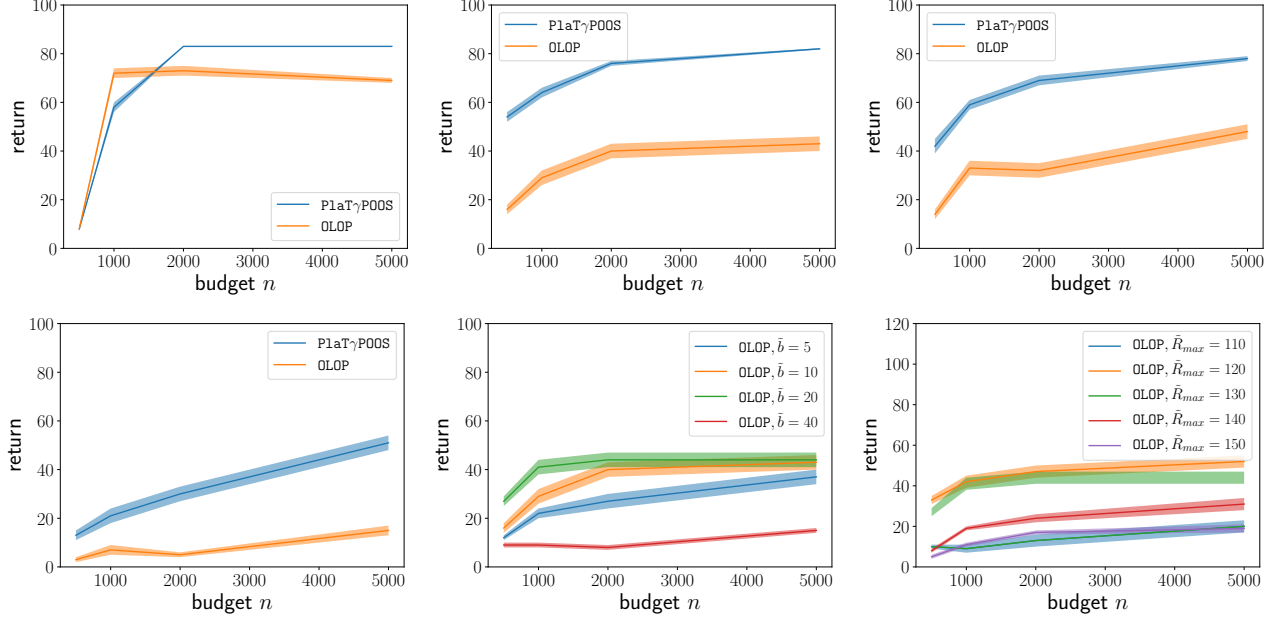


Figure 4. *Top and bottom left*: Average cumulative discounted return collected by OLOP and **PlaTγPOOS** with different range of noise, $b = 1$ (top left), $b = 10$ (top center), $b = 20$, (top right), and $b = 50$ (bottom left). *Bottom center*: the sensitivity of OLOP to different $\tilde{R}_{max}$ parameters. *Bottom right*: the sensitivity of OLOP to different range of the input noise $\tilde{b}$ as parameters while the true b is set to 10.

6. Numerical experiments

In this section, we empirically illustrate the benefits of **PlaTγPOOS**. We chose a simple MDP, shown in Figure 5. In this MDP, a state $x \triangleq (\text{bin}, d)$ is a pair of a binary variable bin and a non-negative integer d . The MDP has two actions that are also binary. If $\text{bin} \neq a$, the base reward is 2, in which case, the next state is $(a, 0)$. Otherwise, if $\text{bin} = a$, then $r = d$ and the next state is $(a, d + 1)$. The reward is then shifted by adding 100 to it so that the noises with different ranges can be added on top without making the reward negative.

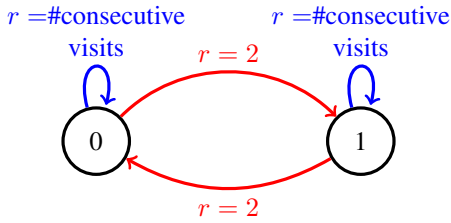


Figure 5. MDP used for our experiments

The initial state is $(0, 0)$. Therefore, the agent has a choice. It can, for instance, remain in the same binary state bin, starting with a null reward but sees its instant reward growing with time if it keeps taking the same action in the future. Alternatively, it could greedily switch to the other binary state bin and obtain a reward of 2 but delaying the hope of obtaining growing reward as in the first scenario. We set $\gamma = 0.95$. Therefore, $R_{max} \approx 130$.

Figure 4 reports the results. All the figures show the cumulative discounted return collected by OLOP and **PlaTγPOOS** after having interacted for 20 steps with the MDP, having chosen each time an action following their planning strategy and then being transferred to the state resulting of applying the recommended action in the current state; therefore also collecting a reward that is composing the final return.

Note that the return reported are shifted in order to not take into account the fixed 100 part of each the reward. The figures in the top row, as well as the figure at the bottom left, reports the comparison between the two returns of OLOP and **PlaTγPOOS** for different ranges of noise b . **PlaTγPOOS** is systematically outperforming OLOP while in this case OLOP is given the correct $\tilde{R}_{max}$ as input, $\tilde{R}_{max} = R_{max}$ and the correct range of the noise $\tilde{b}$, that is $\tilde{b} = b$.

In Figure 4, bottom center and right, we illustrate the sensitivity of OLOP to misleading input parameters. Notice that the performance of OLOP is very vulnerable to these mis-specifications while **PlaTγPOOS** is not using such inputs.

Acknowledgments The research presented was supported by European CHIST-ERA project DELTA, French Ministry of Higher Education and Research, Nord-Pas-de-Calais Regional Council, Inria and Otto-von-Guericke-Universität Magdeburg associated-team north-European project Allocate, and French National Research Agency project BoB (grant n.ANR-16-CE23-0003), FMJH Program PGMO with the support of this program from Criteo.

References

- Bartlett, P. L., Gabillon, V., and Valko, M. [A simple parameter-free and adaptive approach to optimization under a minimal local smoothness assumption](#). In *Algorithmic Learning Theory*, 2019.
- Bertsekas, D. and Tsitsiklis, J. *Neuro-dynamic programming*. Athena Scientific, Belmont, MA, 1996.
- Bubeck, S. and Munos, R. [Open-loop optimistic planning](#). In *Conference on Learning Theory*, 2010.
- Bubeck, S., Munos, R., Stoltz, G., and Szepesvári, C. [\$\mathcal{X}\$ -armed bandits](#). *Journal of Machine Learning Research*, 12:1587–1627, 2011.
- Buşoniu, L. and Munos, R. [Optimistic planning for Markov decision processes](#). In *International Conference on Artificial Intelligence and Statistics*, 2012.
- Carpentier, A. and Locatelli, A. [Tight \(lower\) bounds for the fixed budget best-arm identification bandit problem](#). In *Conference on Learning Theory*, 2016.
- Coquelin, P.-A. and Munos, R. [Bandit algorithms for tree search](#). In *Uncertainty in Artificial Intelligence*, 2007.
- Coulom, R. [Efficient selectivity and backup operators in Monte-Carlo tree search](#). *Computers and games*, 4630: 7283, 2007.
- Feldman, Z. and Domshlak, C. [Simple regret optimization in online planning for Markov decision processes](#). *Journal of Artificial Intelligence Research*, 2014.
- Gabillon, V., Ghavamzadeh, M., and Lazaric, A. [Best-arm identification: A unified approach to fixed budget and fixed confidence](#). In *Neural Information Processing Systems*, 2012.
- Gelly, S., Yizao, W., Munos, R., and Teytaud, O. [Modification of UCT with patterns in Monte-Carlo Go](#). Technical report, Inria, 2006.
- Grill, J.-B., Valko, M., and Munos, R. [Blazing the trails before beating the path: Sample-efficient Monte-Carlo planning](#). In *Neural Information Processing Systems*, 2016.
- Hoorfar, A. and Hassani, M. [Inequalities on the lambert w function and hyperpower function](#). *Journal of Inequalities in Pure and Applied Mathematics*, 9(2):5–9, 2008.
- Hren, J.-F. and Munos, R. [Optimistic planning of deterministic systems](#). In *European Workshop on Reinforcement Learning*, 2008.
- Kaufmann, E. and Koolen, W. M. [Monte-carlo tree search by best-arm identification](#). In *Neural Information Processing Systems*, 2017.
- Kocsis, L. and Szepesvári, C. [Bandit-based Monte-Carlo planning](#). In *European Conference on Machine Learning*, 2006.
- Leurent, E. and Maillard, O.-A. [Practical Open-Loop Optimistic Planning](#). *arXiv preprint arXiv:1904.04700*, 2019.
- Munos, R. [Optimistic optimization of deterministic functions without the knowledge of its smoothness](#). In *Neural Information Processing Systems*, 2011.
- Munos, R. [From bandits to Monte-Carlo tree search: The optimistic principle applied to optimization and planning](#). *Foundations and Trends in Machine Learning*, 7(1):1–130, 2014.
- Orabona, F. and Pál, D. [Scale-free online learning](#). *Theoretical Computer Science*, 2018.
- Orabona, F. and Tommasi, T. [Training deep networks without learning rates through coin betting](#). In *Neural Information Processing Systems*, 2017.
- Puterman, M. L. *Markov Decision Processes: Discrete Stochastic Dynamic Programming*. John Wiley & Sons, New York, NY, 1994.
- Ross, S., Mineiro, P., and Langford, J. [Normalized online learning](#). In *Uncertainty in Artificial Intelligence*, 2013.
- Shah, D., Xie, Q., and Xu, Z. [On reinforcement learning using Monte-Carlo tree search with supervised learning: Non-asymptotic analysis](#). *arXiv preprint arXiv:1902.05213*, 2019.
- Silver, D., Huang, A., Maddison, C. J., Guez, A., Sifre, L., van den Driessche, G., Schrittwieser, J., Antonoglou, I., Panneershelvam, V., Lanctot, M., Dieleman, S., Grewe, D., Nham, J., Kalchbrenner, N., Sutskever, I., Lillicrap, T., Leach, M., Kavukcuoglu, K., Graepel, T., and Hassabis, D. [Mastering the game of Go with deep neural networks and tree search](#). *Nature*, 529(7587):484–489, 2016.
- Szörényi, B., Kedenburg, G., and Munos, R. [Optimistic planning in Markov decision processes using a generative model](#). In *Neural Information Processing Systems*, 2014.
- Valko, M., Carpentier, A., and Munos, R. [Stochastic simultaneous optimistic optimization](#). In *International Conference on Machine Learning*, 2013.

A. On the branching factor

Proposition 2. $\kappa^u(\nu/2, \gamma) \leq \kappa^v(\nu, \gamma) \leq \kappa^u(2\nu, \gamma)$.

Proof. For any global optimum and for $h \geq 0$, we prove that

$$\mathcal{N}_h^v(\varepsilon) \leq \mathcal{N}_h^u\left(\varepsilon + \frac{\gamma^h}{1-\gamma}\right).$$

Let a node $a \in A^h$ be such that $v(a) \geq v^* - \varepsilon$. Then, we have

$$u(a) \geq v(a) - \frac{\gamma^h}{1-\gamma} \geq v^* - \varepsilon - \frac{\gamma^h}{1-\gamma}.$$

Similarly, for $h \geq 0$, we have that for any global optimum,

$$\mathcal{N}_h^u(\varepsilon) \leq \mathcal{N}_h^v\left(\varepsilon + \frac{\gamma^h}{1-\gamma}\right).$$

Using Definition 2, we get the claimed result. $\square$

B. PlaT γ P00S is not using a budget larger than $n + 1$

Notice that for any given depth $h \in [1 : h_{\max}]$, PlaT γ P00S never uses more evaluations than $(p_{\max} + 1) \frac{h_{\max}}{h}$ because

$$\sum_{p=0}^{\lfloor \log_2(h_{\max}/\lceil h\gamma^{2h} \rceil) \rfloor} \left\lfloor \frac{h_{\max}}{h \lceil h2^p\gamma^{2h} \rceil} \right\rfloor \lceil h2^p\gamma^{2h} \rceil \leq (\lfloor \log_2(h_{\max}/\lceil h\gamma^{2h} \rceil) \rfloor + 1) \frac{h_{\max}}{h}.$$

Summing over the depths, PlaT γ P00S never uses more evaluations than the budget $n + 1$ during its depth exploration as

$$\begin{aligned} 1 + (p_{\max} + 1) \sum_{h=1}^{h_{\max}} \left\lfloor \frac{h_{\max}}{h} \right\rfloor &\leq 1 + (p_{\max} + 1) h_{\max} \sum_{h=1}^{h_{\max}} \frac{1}{h} \\ &= 1 + h_{\max} \overline{\log}(h_{\max})(p_{\max} + 1) \leq 1 + h_{\max}(p_{\max} + 1)^2 \\ &\leq \frac{n}{2} + 1. \end{aligned}$$

We need to add the additional evaluation for the cross-validation at the end,

$$\sum_{p=0}^{p_{\max}} \sum_{t=0}^{h_{\max}} \frac{(t+1)\gamma^{2t}h_{\max}}{(1-\gamma^2)^2} \leq \sum_{p=0}^{p_{\max}} \left\lfloor \frac{n}{2(\log_2 n + 1)^2} \right\rfloor \leq \frac{n}{2}.$$

Therefore, the total budget is never more than $n/2 + n/2 + 1 = n + 1$. Again, notice we use the budget of $n + 1$ only for the notational convenience, we could also use $n/4$ for the evaluation in the end to fit under n . Nonetheless, it's important that the amount of openings is *linear* in n .

C. Proofs of the lemmas

We first define favorable event ξ and prove that it holds with high probability.

Lemma 2. *Let $\mathcal{C}$ be the set of sequence of actions evaluated by PlaT γ P00S during one of its runs. $\mathcal{C}$ is a random quantity. Let ξ be the event under which all average estimates for the reward of the state-action pairs receiving at least one evaluation from PlaT γ P00S are within their confidence interval, then $P(\xi) \geq 1 - \delta$, where*

$$\xi \triangleq \left\{ \forall a \in \mathcal{C}, \forall h \in [2 : h(a)], p \in [0 : p_{\max}] : \text{if } T_{a[h]} \geq \left\lceil (h-1)2^p\gamma^{2(h-1)} \right\rceil, \text{ then } |\hat{u}(a) - u(a)| \leq b\sqrt{\frac{p_{\max} \log(4n/\delta)}{2^{p+1}}} \right\}.$$

Proof. The idea of the proof follows the line of proof of the statement given for StoS00 (Valko et al., 2013). The crucial point is that while we have potentially exponentially many combinations of cells that can be evaluated, given any particular execution we need to consider only a polynomial number of estimators, m , for which we can use a Azuma-Hoeffding concentration inequality.

We denote $\forall h \in [0 : h_{\max}], p \in [0 : p_{\max}] : a^{i,h,p} \in \mathcal{C}$, the i -th evaluated node of depth h such that $\forall t \in [2 : h], T_{a_{[t]}^{i,h,p}} \geq \lceil (t-1)2^p \gamma^{2(t-1)} \rceil$. Note that in **PlatYP00S** we have $T_{a_{[1]}^{i,h,p}} = h_{\max}$.

Though $a^{i,h,p}$ is random, we study the quantity $|\hat{u}(a^{i,h,p}) - u(a^{i,h,p})|$. We recall that

$$\hat{u}(a^{i,h,p}) - u(a^{i,h,p}) = \sum_{t=0}^{h-1} \gamma^t (\hat{r}_t(a^{i,h,p}) - r_t(a^{i,h,p})) \quad (2)$$

$$= \sum_{t=0}^{h-1} \gamma^t \sum_{s=0}^{T_{a_{[t+1]}^{i,h,p}}} \frac{\hat{r}_{t,s}^{i,h,p} - r_t(a^{i,h,p})}{T_{a_{[t+1]}^{i,h,p}}} \quad (3)$$

This quantity is composed of the elements $\hat{r}_{t,s}^{i,h,p} - r_t(a^{i,h,p})$ that form a martingale.

Therefore using a Azuma-Hoeffding concentration inequality with a union bound already on the values of T we have

$$\mathbb{P} \left(\hat{u}(a^{i,h,p}) - u(a^{i,h,p}) \leq b \sqrt{\sum_{t=0}^{h-1} \frac{\gamma^{2t} \log(p_{\max}/\delta)}{2T_{a_{[t+1]}^{i,h,p}}}} \right) \geq 1 - \delta/p_{\max}$$

Moreover we have for all $h \geq t > 1$,

$$\frac{\gamma^{2t}}{T_{a_{[t+1]}^{i,h,p}}} \leq \frac{\gamma^{2t}}{\lceil t2^p \gamma^{2t} \rceil} \leq \frac{\gamma^{2t}}{t2^p \gamma^{2t}} = \frac{1}{t2^p} \quad (4)$$

For $t = 0$, $\frac{\gamma^{2t}}{T_{a_{[t+1]}^{i,h,p}}} = \frac{1}{h_{\max}} \leq \frac{1}{2^p}$ for all $p \leq p_{\max}$.

Therefore we have

$$\mathbb{P} \left(\hat{u}(a^{i,h,p}) - u(a^{i,h,p}) \leq b \sqrt{\frac{\log h_{\max} \log(?/\delta)}{2^{p+1}}} \right) \geq 1 - \delta/?$$

Then we had an extra union bound over all cells that is bounded by n

□

Lemma 3. For any planning problem with associated (ν, ρ) (see Property 1), on event ξ , for any depths $h \in [h_{\max}]$, for any $p \in [0 : \lfloor \log_2(h_{\max}/(h^2 \gamma^{2h})) \rfloor]$, we have $\perp_{h,p} = h$ if conditions (1) and (2) simultaneously hold true.

(1) $b\sqrt{\log(4n/\delta)/2^{p+1}} \leq \nu \rho^h$

(2) For all $h' \in [h, h_{\max}/(h' \lceil h' 2^p \gamma^{2h'} \rceil)] \geq C\kappa(\nu, \rho)^{h'}$.

Finally we have $\perp_{0,p} = 0$.

Proof. We place ourselves on event ξ defined in Lemma 2 and for which we proved that $P(\xi) \geq 1 - \delta$. We fix p .

We prove the statement of the lemma, given that event ξ holds, by induction in the following sense. For a given h and p , we assume the hypotheses of the lemma for that h and p are true and we prove by induction that $\perp_{h',p} = h'$ for $h' \in [h]$.

1° For $h = 0$, we trivially have that $\perp_{h,p} \geq 0$.

2° Now consider $h' > 0$, and assume $\perp_{h'-1,p} = h' - 1$ with the objective to prove that $\perp_{h',p} = h'$.

Therefore, at the end of the processing of depth $h' - 1$, during which we were opening the nodes of depth $h' - 1$ we managed to open an optimal node that we denote $a^{*,h'-1} \in A^{*,h'-1}$. Moreover if we consider all the sequence of actions b that one can build by appending any action in A to $a^{*,h'-1} \in A^{*,h'-1}$, we have for all such b that $T_{b[t]} \geq \lceil (t-1)2^p\gamma^{t-1} \rceil$ for $t \in [h']$.

Note that by definition there exist an optimal infinite sequence of actions $a^* \in A^*$ such that $a^{*,h'-1} = a_{[h'-1]}^*$

During phase h' the $\left\lfloor h_{\max}/\left(h' \lceil h'2^p\gamma^{2h'} \rceil\right) \right\rfloor$ evaluated nodes from $A^{h'-1}$ with highest values $\{\widehat{u}(a^{h'-1,i})\}_{h'-1,i}$ are opened.

For the purpose of contradiction, let us assume that $a_{[h']}^*$ is not one of them. This would mean that there exist at least $\left\lfloor h_{\max}/\left(h' \lceil h'2^p\gamma^{2h'} \rceil\right) \right\rfloor$ nodes from $A^{h'}$, distinct from $a_{[h']}^*$, satisfying $\widehat{u}(a^{h',i}) \geq \widehat{u}(a_{[h']}^*)$ and each verifying $T_{a_{[t]}^{h',i}} \geq \lceil t2^p\gamma^t \rceil$ for $t \in [h']$. This means that, for these nodes we have: $u(a_{h',i}) + \nu\rho^{h'} \geq u(a_{h',i}) + \nu\rho^h \stackrel{(a)}{\geq} u(a_{h',i}) + b\sqrt{\log(4n/\delta)/2^{p+1}} \stackrel{(b)}{\geq} \widehat{u}(a_{h',i}) \geq \widehat{u}(a_{h',i_h^*}) \geq u(a_{h',i_h^*}) - b\sqrt{\log(4n/\delta)/2^{p+1}} \stackrel{(a)}{\geq} u(a_{h',i_h^*}) - \nu\rho^h \geq u(a_{h',i_h^*}) - \nu\rho^{h'}$, where (a) is by assumption of the lemma, (b) is because ξ holds. As $u(a_{h',i_h^*}) \geq v^* - \nu\rho^{h'}$ by Proposition 1, this means we have $\mathcal{N}_{h'}^u(3\nu\rho^{h'}) \geq \left\lfloor h_{\max}/\left(h' \lceil h'2^p\gamma^{2h'} \rceil\right) \right\rfloor + 1$ (the +1 is for a_{h,i_h^*}). However, by assumption of the lemma $h_{\max}/\left(h' \lceil h'2^p\gamma^{2h'} \rceil\right) \geq C\kappa(\nu, \rho)^{h'}$. It follows that in general $\mathcal{N}_{h'}^u(3\nu\rho^{h'}) > \left\lfloor C\kappa(\nu, \rho)^{h'} \right\rfloor$. This leads to having a contradiction with the $\kappa^u(\nu, \rho)$ with associated constant C as defined in Definition 2. Indeed, the condition $\mathcal{N}_{h'}^u(3\nu\rho^{h'}) \leq C\kappa(\nu, \rho)^{h'}$ in Definition 2 is equivalent to the condition $\mathcal{N}_{h'}^u(3\nu\rho^{h'}) \leq \left\lfloor C\kappa(\nu, \rho)^{h'} \right\rfloor$ as $\mathcal{N}_{h'}^u(3\nu\rho^{h'})$ is an integer.

□

Lemma 1. *For any planning problem with associated (ν, ρ) as in Property 1), on event ξ defined in Appendix C, for any depth $h \in [h_{\max}]$, for any $p \in [0 : \lfloor \log_2(h_{\max}/(h^2\gamma^{2h})) \rfloor]$, we have $\perp_{h,p} = h$ if (1) and (2) simultaneously hold:*

(1) $b\sqrt{\log(4n/\delta)/2^{p+1}} \leq \nu\rho^h$

(2) We distinguish cases and express the condition in each:

Case 1) $h2^p\gamma^{2h} \leq 1$:

$$\frac{h_{\max}}{h} = \frac{h_{\max}}{h \lceil h2^p\gamma^{2h} \rceil} \geq C\kappa(\nu, \rho)^h$$

$$\text{and for all } h' \in [h], \frac{h_{\max}}{h'2^{p+1}\gamma^{2h'}} \geq C\kappa(\nu, \rho)^{h'}$$

Case 2) $h2^p\gamma^{2h} \geq 1$:

$$\textbf{Case 2.1)} \gamma^2\kappa^u \geq 1 : \frac{h_{\max}}{h^22^{p+1}\gamma^{2h}} \geq C\kappa(\nu, \rho)^h$$

$$\textbf{Case 2.2)} \gamma^2\kappa^u \leq 1 : \frac{h_{\max}}{h^22^{p+1}} \geq C$$

Proof. To prove this statement we just need to show that we verify the hypotheses of Lemma 3. This means we need to prove that for all $h' \in [h]$, $h_{\max}/\left(h' \lceil h'2^p\gamma^{2h'} \rceil\right) \geq C\kappa(\nu, \rho)^{h'}$.

We first consider the case 2) where $h2^p\gamma^h \geq 1$. If $h = 1$ we already know $\perp_{0,p} \geq 0$. Let us now look at the case $h > 1$. First notice that $h2^p\gamma^{2h} \geq 1$ gives $(h-1)2^p\gamma^{2(h-1)} \geq 1$. If $\gamma^2\kappa^u \geq 1$ we have that for all $h' \in [h-1]$,

$$\frac{h_{\max}}{h'2^{p+1}} \geq \frac{h_{\max}}{h^22^{p+1}} \geq C(\gamma^2\kappa(\nu, \rho))^h \geq C(\gamma^2\kappa(\nu, \rho))^{h'}$$

If $h > 1$, and if $\gamma^2\kappa^u \leq 1$ we have that for all $h' \in [h-1]$,

$$\frac{h_{\max}}{h'2^{p+1}} \geq \frac{h_{\max}}{h^22^{p+1}} \geq C \geq C(\gamma^2\kappa(\nu, \rho))^{h'}.$$

For both $\gamma^2 \kappa^u \leq 1$ and $\gamma^2 \kappa^u \geq 1$, we then have,

$$\frac{h_{\max}}{h' \lceil h' 2^p \gamma^{2h'} \rceil} \geq \frac{h_{\max}}{h' h' 2^{p+1} \gamma^{2h'}}$$

as $h 2^p \gamma^{2h} \geq 1$.

For both $\gamma^2 \kappa^u \leq 1$ and $\gamma^2 \kappa^u \geq 1$, the previous equations mean that for $h' \in [h-1]$, h' verifies $h_{\max}/h'^2 2^{p+1} \gamma^{2h'} \geq C\kappa(\nu, \rho)^{h'} \geq 1$. Therefore $p \leq \lfloor \log_2(h_{\max}/(h'^2 \gamma^{2h'})) \rfloor$.

We now consider case 1) where $h 2^p \gamma^{2h} \leq 1$. We prove by induction that for all $h' \in [h]$, $h_{\max}/(h' \lceil h' 2^p \gamma^{2h'} \rceil) \geq C\kappa(\nu, \rho)^{h'}$.

1° By assumption of the lemma we say: $h_{\max}/(h \lceil h 2^p \gamma^{2h} \rceil) \geq C\kappa(\nu, \rho)^h$

2° We further assume $h_{\max}/(h' \lceil h' 2^p \gamma^{2h'} \rceil) \geq C\kappa(\nu, \rho)^{h'}$ is true for some $h' \leq h$ with $h' 2^p \gamma^{2h'} \leq 1$

We want to prove that either:

both $(h' - 1) 2^p \gamma^{2(h'-1)} \leq 1$

and $\frac{h_{\max}}{(h'-1) \lceil (h'-1) 2^p \gamma^{2(h'-1)} \rceil} \geq C\kappa(\nu, \rho)^{h'-1}$

or $(h' - 1) 2^p \gamma^{2(h'-1)} \geq 1$

then $h_{\max}/((h'') \lceil (h'') 2^p \gamma^{2(h'')} \rceil) \geq C\kappa(\nu, \rho)^{h''}$ is already true for all $h'' \in [h']$. If $\lceil (h' - 1) 2^p \gamma^{2(h'-1)} \rceil = 1$ then we have

$$\frac{h_{\max}}{(h' - 1)} \geq \frac{h_{\max}}{h'} \geq C\kappa(\nu, \rho)^{h'} \geq C\kappa(\nu, \rho)^{h'-1}$$

If $\lceil (h' - 1) 2^p \gamma^{2(h'-1)} \rceil > 1$, then we have that

$$\begin{aligned} h_{\max}/(h' \lceil (h' - 1) 2^p \gamma^{2(h'-1)} \rceil) &\geq \\ h_{\max}/((h' - 1) 2^{p+1} \gamma^{2(h'-1)}) &\geq C\kappa(\nu, \rho)^{h'-1} \end{aligned}$$

Using this inequality we can now use Case 2) to have that: $h_{\max}/((h'') \lceil (h'') 2^p \gamma^{2(h'')} \rceil) \geq C\kappa(\nu, \rho)^{h''}$ is already true for all $h'' \in [h']$.

The previous equations mean that for $h' \in [h-1]$, h' verifies $h_{\max}/(h'^2 2^{p+1} \gamma^{2h'}) \geq C\kappa(\nu, \rho)^{h'} \geq 1$. Therefore $p \leq \lfloor \log_2(h_{\max}/(h'^2 \gamma^{2h'})) \rfloor$. $\square$

D. Proof of Theorem 3 and Theorem 4

Theorem 3. High-noise regime *If the noise b is high enough to verify both high-noise conditions as defined in the caption of Table 1, then after n rounds, for any problem with associated (ν, ρ) , and branching factor $\kappa \triangleq \kappa^u(\nu, \rho)$, the simple regret of **PlaTγPOOS** obeys*

$$\mathbb{E}r_n = \begin{cases} \tilde{\mathcal{O}}\left(\left(\frac{n}{b^2}\right)^{-\frac{1}{2}}\right) & \text{if } \gamma^2 \kappa \leq 1, \\ \tilde{\mathcal{O}}\left(\left(\frac{n}{b^2}\right)^{-\frac{\log(1/\rho)}{\log(\gamma^2 \kappa / \rho^2)}}\right) & \text{if } \gamma^2 \kappa > 1. \end{cases}$$

Theorem 4. Low-noise regime *If the noise b is low enough to verify both high-noise conditions as defined in the caption of Table 1, then after n rounds, for any problem with associated (ν, ρ) , and branching factor $\kappa \triangleq \kappa^u(\nu, \rho)$, the simple regret of **PlatYP00S** obeys*

$$\mathbb{E}r_n = \begin{cases} \tilde{\mathcal{O}}(\nu \rho^n) & \text{if } \kappa = 1, \\ \tilde{\mathcal{O}}\left(\nu \left(\frac{n}{b^2}\right)^{-\frac{\log(1/\rho)}{\log(\kappa)}}\right) & \text{if } \kappa > 1. \end{cases}$$

Proof of Theorem 3 and Theorem 4. We first place ourselves on the event ξ defined in Lemma 2 and where it is proven that $P(\xi) \geq 1 - \delta$. We bound the simple regret of **PlatYP00S** on ξ .

Step 1) General definition of the regret

We chose $\delta = \frac{4b(1-\delta)}{R_{\max}\sqrt{n}}$ for the bound. We consider a problem with associated (ν, ρ) . For simplicity we write $\kappa = \kappa^u(\nu, \rho)$. We have for all $p \in [0 : p_{\max}]$

$$\begin{aligned} v(a^n) + \frac{b}{1-\gamma^2} \sqrt{\frac{p_{\max} \log(R_{\max} n^{3/2}/b)}{2h_{\max}}} \\ &\geq u(a^n) + \frac{b}{1-\gamma^2} \sqrt{\frac{p_{\max} \log(R_{\max} n^{3/2}/b)}{2h_{\max}}} \stackrel{\text{(a)}}{\geq} \hat{u}(a^n) \\ &\stackrel{\text{(c)}}{\geq} \hat{u}(a^p) \stackrel{\text{(b)}}{\geq} \hat{u}(a_{[\perp_{h_{\max}, p+1}]}^p) \\ &\stackrel{\text{(a)}}{\geq} u(a_{[\perp_{h_{\max}, p+1}]}^p) - \frac{b}{1-\gamma^2} \sqrt{\frac{p_{\max} \log(R_{\max} n^{3/2}/b)}{2h_{\max}}} \\ &\stackrel{\text{(d)}}{\geq} v^* - \nu \rho^{\perp_{h_{\max}, p+1}} - \frac{b}{1-\gamma^2} \sqrt{\frac{p_{\max} \log(R_{\max} n^{3/2}/b)}{2h_{\max}}} \end{aligned}$$

where **(a)** is because the actions at time t , $a_t(n, p)$, of the candidate $a(n, p)$ have been evaluated $\frac{(t+1)\gamma^{2t}h_{\max}}{(1-\gamma)^2}$ times and because ξ holds, **(b)** is because $a_{[\perp_{h_{\max}, p+1}]}^p \in \{a \in A^\bullet : \forall t \in [2 : h(a)], T_{a[t]} \geq \lceil (t-1)2^p\gamma^{2(t-1)} \rceil\}$ and $a^p = \arg \max_{a \in A^\bullet : \forall t \in [2 : h(a)], T_{a[t]} \geq \lceil (t-1)2^p\gamma^{2(t-1)} \rceil} \hat{u}(a)$, **(c)** is because $a^n = \arg \max_{\{a^p, p \in [0 : p_{\max}]\}} \hat{u}(a^p)$, and **(d)** is by Assumption 1.

From the previous inequality we have $r_n = v^* - Q^*(x, a^n) \leq \nu \rho^{\perp_{h_{\max}, p+1}} + 2 \frac{b}{1-\gamma^2} \sqrt{\frac{p_{\max} \log(R_{\max} n^{3/2}/b)}{2h_{\max}}}$, for $p \in [0 : p_{\max}]$.

Step 2) Defining some important depths For the rest of proof we want to lower bound $\max_{p \in [0 : p_{\max}]} \perp_{h_{\max}, p}$. Lemma 3 and 1 provide some sufficient conditions on p and h to get lower bounds. These conditions are inequalities in which as p gets smaller (fewer samples) or h gets larger (more depth) these conditions are more and more likely not to hold. For our bound on the regret of **PlatYP00S** to be small, we want quantities p and h where the inequalities hold but using as few samples as possible (small p) and having h as large as possible. Therefore we are interested in determining when the inequalities flip signs which is when they turn to equalities. This is what we solve next.

We set the notation $g_{n,b}^{\delta, R_{\max}} = p_{\max} \log(R_{\max} n^{3/2}/b(1-\delta))$.

In its most general form we are interested in the real numbers $\tilde{h}$ and $\tilde{p}$ are such that $\tilde{h}$ is the larger real number such that for all $h \leq \tilde{h}'$

$$\frac{h_{\max}}{h^2 2^{\tilde{p}+1} \gamma^{2h}} \geq C \kappa(\nu, \rho)^h \text{ while } b \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{2^{\tilde{p}}}} = \nu \rho^{\tilde{h}} \quad (5)$$

In the case $\gamma^2\kappa \geq 1$ we can simply solve the following equations. We denote $\tilde{h}_1, \tilde{p}_1$ the real numbers satisfying

$$\frac{h_{\max}\nu^2\rho^{2\tilde{h}_1}}{2\tilde{h}_1^2b^2g_{n,b}^{\delta,R_{\max}}\gamma^{2\tilde{h}_1-2}} = C\kappa^{\tilde{h}_1} \quad \text{and} \quad b\sqrt{\frac{g_{n,b}^{\delta,R_{\max}}}{2\tilde{p}_1}} = \nu\rho^{\tilde{h}_1}\gamma^2. \quad (6)$$

In the case $\gamma^2\kappa \leq 1$ the previous equation can possess two solutions where the largest of these two solutions will not verify Equation 5. Additionally the smallest solution might be hard to express in a close form when $\gamma^2\kappa \leq 1$. Therefore for simplicity we define for the case $\gamma^2\kappa \leq 1$, $\tilde{h}_2, \tilde{p}_2$ the real numbers satisfying

$$\frac{h_{\max}\nu^2\rho^{2\tilde{h}_2}}{2\tilde{h}_2^2b^2g_{n,b}^{\delta,R_{\max}}} = C \quad \text{and} \quad b\sqrt{\frac{g_{n,b}^{\delta,R_{\max}}}{2\tilde{p}_2}} = \nu\rho^{\tilde{h}_2}. \quad (7)$$

$\tilde{h}_1$ and $\tilde{p}_1$ are defined for the case $\gamma^2\kappa \geq 1$ while $\tilde{h}_2$ and $\tilde{p}_2$ are defined for the case $\gamma^2\kappa \leq 1$. Our approach is to solve Equation 6 and 7 and then verify that it gives a valid indication of the behavior of our algorithm in term of its optimal p and h . We have

$$\begin{aligned} \tilde{h}_1 &= \frac{2}{\log(\gamma^2\kappa/\rho^2)} W \left(\log(\gamma^2\kappa/\rho^2)/2 \sqrt{\frac{\gamma^2\nu^2h_{\max}}{2Cb^2g_{n,b}^{\delta,R_{\max}}}} \right) \\ \tilde{h}_2 &= \frac{2}{\log(1/\rho^2)} W \left(\log(1/\rho^2)/2 \sqrt{\frac{\nu^2h_{\max}}{2Cb^2g_{n,b}^{\delta,R_{\max}}}} \right) \end{aligned}$$

where standard W is the Lambert W function.

However after a close look at the Equation 7, we notice that it is possible to get values of $\tilde{p}$ and $\tilde{h}$ which would lead to a number of evaluations $\tilde{h}2^{\tilde{p}}\gamma^{\tilde{h}} < 1$. This actually corresponds to an interesting case when the noise has a small range and where we can expect to obtain an improved result, that is: obtain a regret rate close to the deterministic case. This low range of noise case then has to be considered separately.

Therefore, we distinguish two cases which corresponds to different noise regimes depending on the value of b . Looking at the equation on the right of (7), we have that $\tilde{h}2^{\tilde{p}}\gamma^{\tilde{h}} < 1$ if $\frac{\nu^2\rho^{2\tilde{h}}}{\gamma^{2\tilde{h}}\tilde{h}b^2g_{n,b}^{\delta,R_{\max}}} > 1$. Based on this condition we now consider the two cases. However for both of them we define some generic $\tilde{h}$ and $\tilde{p}$.

Case 1) $\gamma^2\kappa \geq 1$: Note that in this case then $\kappa > 1$. We subdivide this case into multiple subcases:

Case 1.1) Noise regime $\frac{\nu^2\rho^{2\tilde{h}_1}}{\gamma^{2\tilde{h}_1}\tilde{h}_1b^2g_{n,b}^{\delta,R_{\max}}} \leq 1$

Case 1.1.1) High-noise regime $\frac{\nu^2\rho^{2\tilde{h}_1}}{b^2g_{n,b}^{\delta,R_{\max}}} \leq 1$

In this case, we denote $\tilde{h}_1 = \tilde{h}$ and $\tilde{p}_1 = \tilde{p}$. As $\frac{\nu^2\rho^{2\tilde{h}_1}}{b^2g_{n,b}^{\delta,R_{\max}}} \leq 1$ by construction, we have $\tilde{p}_1 \geq 0$. Using standard properties of the $\lfloor \cdot \rfloor$ function, we have

$$b\sqrt{\frac{g_{n,b}^{\delta,R_{\max}}}{2\lfloor \tilde{p}_1 \rfloor + 1}} \leq b\sqrt{\frac{g_{n,b}^{\delta,R_{\max}}}{2\tilde{p}_1}} \leq \nu\rho^{\tilde{h}_1} \leq \nu\rho^{\lfloor \tilde{h}_1 \rfloor} \quad (8)$$

$$\begin{aligned} \text{and, } & \frac{h_{\max}}{\lfloor \tilde{h}_1 \rfloor \lfloor \tilde{h}_1 \rfloor 2^{\lfloor \tilde{p}_1 \rfloor + 1} \gamma^{2\lfloor \tilde{h}_1 \rfloor}} \geq \frac{h_{\max}}{\lfloor \tilde{h}_1 \rfloor \lfloor \tilde{h}_1 \rfloor 2^{\tilde{p}_1 + 1} \gamma^{2\lfloor \tilde{h}_1 \rfloor}} \\ & \geq \frac{h_{\max}}{\lfloor \tilde{h}_1 \rfloor \tilde{h}_1 2^{\tilde{p}_1 + 1} \gamma^{2\tilde{h}_1 - 2}} = \frac{h_{\max}\nu^2\rho^{2\tilde{h}_1}}{2\tilde{h}_1^2b^2g_{n,b}^{\delta,R_{\max}}\gamma^{2\tilde{h}_1}} \\ & = C\kappa^{\tilde{h}_1} \geq C\kappa^{\lfloor \tilde{h}_1 \rfloor}. \end{aligned}$$

We will verify that $\lfloor \ddot{h} \rfloor$ is a reachable depth by **PlaTγPOOS** in the sense that $\ddot{h} \leq h_{\max}$ and $\lfloor \ddot{p} \rfloor \leq \lfloor \log_2(h_{\max}/(h^2\gamma^{2h})) \rfloor$ and . As $\kappa < 1$, and $\ddot{h} \geq 0$ we have $\kappa^{\ddot{h}} \geq 1$. This gives $C\kappa^{\ddot{h}} \geq 1$. Finally as $\frac{h_{\max}}{h^2\gamma^{2h}} \geq C\kappa^{\ddot{h}}$, we have $\ddot{h}^2\gamma^{2\ddot{h}} \leq h_{\max}/2^{\ddot{p}}$.

Case 1.1.2) Low-noise regime 1 $\frac{\nu^2 \rho^{2\ddot{h}_1}}{b^2 g_{n,b}^{\delta, R_{\max}}} \geq 1$

We denote $\ddot{h} = \bar{h}_1$ and $\ddot{p} = \bar{p}_1$ where $\bar{h}$ and $\bar{p}$ verify,

$$\frac{h_{\max}}{2\bar{h}_1^2\gamma^{2\bar{h}_1}} = C\kappa^{\bar{h}_1} \quad \text{and} \quad \bar{p}_1 = 0. \quad (9)$$

Again, $\frac{h_{\max}}{2\bar{h}_1^2\gamma^{2\bar{h}_1}} \geq 1$.

$$\bar{h}_1 = \frac{2}{\log(\gamma^2\kappa)} W \left(\log(\gamma^2\kappa)/2\sqrt{\frac{h_{\max}}{2C}} \right)$$

Using standard properties of the $\lfloor \cdot \rfloor$ function, we have

$$b\sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{2^{\lfloor \bar{p}_1 \rfloor + 1}}} \leq b\sqrt{g_{n,b}^{\delta, R_{\max}}} < \nu\rho^{\bar{h}_1} \stackrel{\text{(a)}}{\leq} \nu\rho^{\bar{h}_1} \leq \nu\rho^{\lfloor \bar{h}_1 \rfloor} \quad (10)$$

where (a) is because of the following reasoning. As we have $\frac{h_{\max}\nu^2\rho^{2\bar{h}_1}}{2\bar{h}_1^2b^2g_{n,b}^{\delta, R_{\max}}\gamma^{2\bar{h}_1}} = C\kappa^{\bar{h}_1}$ and $\frac{\nu^2\rho^{2\bar{h}_1}}{b^2g_{n,b}^{\delta, R_{\max}}} \geq 1$, then, $\frac{h_{\max}}{2\bar{h}_1^2\gamma^{2\bar{h}_1}} \leq C\kappa^{\bar{h}_1}$. From the inequality $\frac{h_{\max}}{2\bar{h}_1^2} \leq C\kappa^{\bar{h}_1}\gamma^{2\bar{h}_1}$ and the fact that $\bar{h}_1$ corresponds to the case of equality $\frac{h_{\max}}{2\bar{h}_1^2} = C\kappa^{\bar{h}_1}\gamma^{2\bar{h}_1}$, we deduce that $\bar{h}_1 \leq \tilde{h}_1$, since the left term of the inequality decreases with h while the right term increases (as $\gamma^2\kappa \geq 1$). Having $\bar{h}_1 \leq \tilde{h}_1$ gives $\rho^{\bar{h}_1} \geq \rho^{\tilde{h}_1}$.

Moreover, the term $\log(\gamma^2\kappa)$ of $\bar{h}_1$ could lead to think that we could potentially obtain a better rate than in the deterministic case where the term is $\log(\gamma^2\kappa)$. However this is not true because as $\bar{h}_1$ is the solution of $\bar{h}_1 = \frac{h_{\max}}{2\bar{h}_1^2\gamma^{2\bar{h}_1}} = C\kappa^{\bar{h}_1}$ and we have by assumption in this case $\bar{h}_1\gamma^{2\bar{h}_1} \geq 1$ then $\bar{h}_1 \leq h_3$ where h_3 is defined as the solution of $h_3 = \frac{h_{\max}}{2h_3\gamma^{2h_3}} = C\kappa^{h_3}$. We have $h_3 = \frac{1}{\log(\kappa)} W \left(\log(\kappa) \frac{h_{\max}}{2C} \right)$. Therefore one can see that this rate is not better than the deterministic rates.

Case 1.2) Low noise regime 2 $\frac{\nu^2 \rho^{2\hat{h}_1}}{\gamma^{2\hat{h}_1} \hat{h}_1 b^2 g_{n,b}^{\delta, R_{\max}}} \geq 1$

We denote $\ddot{h} = \hat{h}_1$ and $\ddot{p} = \hat{p}_1$ where $\hat{h}$ and $\hat{p}$ verify,

$$\frac{h_{\max}}{2\hat{h}_1} = C\kappa^{\hat{h}_1} \quad \text{and} \quad \hat{p}_1 = \max(0, \tilde{p}_1). \quad (11)$$

$$\hat{h}_1 = \frac{1}{\log(\kappa)} W \left(\frac{h_{\max} \log(\kappa)}{2C} \right)$$

Using standard properties of the $\lfloor \cdot \rfloor$ function, we have

$$b\sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{2^{\lfloor \hat{p}_1 \rfloor + 1}}} \leq b\sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{2^{\tilde{p}_1}}} < \nu\rho^{\hat{h}_1} \stackrel{\text{(a)}}{\leq} \nu\rho^{\hat{h}_1} \leq \nu\rho^{\lfloor \hat{h}_1 \rfloor} \quad (12)$$

where (a) is because of the following reasoning. As we have $\frac{h_{\max}\nu^2\rho^{2\hat{h}_1}}{2\hat{h}_1^2b^2g_{n,b}^{\delta, R_{\max}}\gamma^{2\hat{h}_1}} = C\kappa^{\hat{h}_1}$ and $\frac{\nu^2\rho^{2\hat{h}_1}}{\gamma^{2\hat{h}_1}\hat{h}_1b^2g_{n,b}^{\delta, R_{\max}}} \geq 1$, then, $\frac{h_{\max}}{2\hat{h}_1} \leq C\kappa^{\hat{h}_1}$. From the inequality $\frac{h_{\max}}{2\hat{h}_1} \leq C\kappa^{\hat{h}_1}$ and the fact that $\hat{h}_1$ corresponds to the case of equality $\frac{h_{\max}}{2\hat{h}_1} = C\kappa^{\hat{h}_1}$, we deduce that $\hat{h}_1 \leq \tilde{h}_1$, since the left term of the inequality decreases with h while the right term increases. Having $\hat{h}_1 \leq \tilde{h}_1$ gives $\rho^{\hat{h}_1} \geq \rho^{\tilde{h}_1}$.

Case 2) $\gamma^2 \kappa \leq 1$

Case 2.1) Noise regime $\frac{\nu^2 \rho^{2\tilde{h}}}{\gamma^{2\tilde{h}} \tilde{h} b^2 g_{n,b}^{\delta, R_{\max}}} \leq 1$

Case 2.1.1) High-noise regime $\frac{\nu^2 \rho^{2\tilde{h}_2}}{b^2 g_{n,b}^{\delta, R_{\max}}} \leq 1$

In this case, we denote $\ddot{h} = \tilde{h}_2$ and $\ddot{p} = \tilde{p}_2$. As $\frac{\nu^2 \rho^{2\tilde{h}_2}}{b^2 g_{n,b}^{\delta, R_{\max}}} \leq 1$ by construction, we have $\tilde{p}_2 \geq 0$. Using standard properties of the $\lfloor \cdot \rfloor$ function, we have

$$b \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{2^{\lfloor \tilde{p}_2 \rfloor + 1}}} \leq b \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{2^{\tilde{p}_2}}} = \nu \rho^{\tilde{h}_2} \leq \nu \rho^{\lfloor \tilde{h}_2 \rfloor} \quad (13)$$

$$\begin{aligned} \text{and, } \frac{h_{\max}}{\lfloor \tilde{h}_2 \rfloor \lfloor \tilde{h}_2 \rfloor 2^{\lfloor \tilde{p}_2 \rfloor + 1}} &\geq \frac{h_{\max}}{\lfloor \tilde{h}_2 \rfloor \lfloor \tilde{h}_2 \rfloor 2^{\tilde{p}_2 + 1}} \\ &\geq \frac{h_{\max}}{\tilde{h}_2^2 2^{\tilde{p}_2 + 1}} = \frac{h_{\max} \nu^2 \rho^{2\tilde{h}_2}}{2 \tilde{h}_2^2 b^2 g_{n,b}^{\delta, R_{\max}}} \\ &= C. \end{aligned}$$

We will verify that $\lfloor \tilde{h} \rfloor$ is a reachable depth by **PlaT γ P00S** in the sense that $\ddot{h} \leq h_{\max}$ and $\lfloor \tilde{p} \rfloor \leq \lfloor \log_2(h_{\max}/(h^2 \gamma^{2h})) \rfloor$. As $\kappa < 1$, and $\ddot{h} \geq 0$ we have $\kappa^{\ddot{h}} \geq 1$. This gives $C \kappa^{\ddot{h}} \geq 1$. Finally as $\frac{h_{\max}}{\tilde{h}^2 2^{\tilde{p}} \gamma^{2\tilde{h}}} \geq C \kappa^{\ddot{h}}$, we have $\ddot{h}^2 \gamma^{2\ddot{h}} \leq h_{\max}/2^{\ddot{p}}$.

Case 2.1.2) Low-noise regime 1 $\frac{\nu^2 \rho^{2\tilde{h}_2}}{b^2 g_{n,b}^{\delta, R_{\max}}} \geq 1$

We denote $\ddot{h} = \bar{h}_2$ and $\ddot{p} = \bar{p}_2$ where $\bar{h}$ and $\bar{p}$ verify,

$$\frac{h_{\max}}{2\bar{h}_2^2} = C \quad \text{and} \quad \bar{p}_2 = 0. \quad (14)$$

Again, $\frac{h_{\max}}{2\bar{h}_2^2 2^{\bar{p}_2} \gamma^{2\bar{h}_2}} \geq 1$.

$$\bar{h}_2 = \sqrt{\frac{h_{\max}}{2C}}$$

Using standard properties of the $\lfloor \cdot \rfloor$ function, we have

$$b \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{2^{\lfloor \tilde{p}_2 \rfloor + 1}}} \leq b \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{2^{\tilde{p}_2}}} < \nu \rho^{\tilde{h}_1} \stackrel{(a)}{\leq} \nu \rho^{\bar{h}_2} \leq \nu \rho^{\lfloor \bar{h}_2 \rfloor} \quad (15)$$

where (a) is because of the following reasoning. As we have $\frac{h_{\max} \nu^2 \rho^{2\tilde{h}_2}}{2 \tilde{h}_2^2 b^2 g_{n,b}^{\delta, R_{\max}}} = C$ and $\frac{\nu^2 \rho^{2\tilde{h}_2}}{b^2 g_{n,b}^{\delta, R_{\max}}} \geq 1$, then, $\frac{h_{\max}}{2\tilde{h}_2^2} \leq C$. From the inequality $\frac{h_{\max}}{2\tilde{h}_2^2} \leq C$ and the fact that $\bar{h}_2$ corresponds to the case of equality $\frac{h_{\max}}{2\bar{h}_2^2} = C$, we deduce that $\bar{h}_2 \leq \tilde{h}_2$, since the left term of the inequality decreases with h while the right term stays constant. Having $\bar{h}_2 \leq \tilde{h}_2$ gives $\rho^{\bar{h}_2} \geq \rho^{\tilde{h}_2}$.

Case 2.2) Low noise regime 2 $\frac{\nu^2 \rho^{2\tilde{h}}}{\gamma^{2\tilde{h}} \tilde{h} b^2 g_{n,b}^{\delta, R_{\max}}} \geq 1$

We denote $\ddot{h} = \hat{h}_2$ and $\ddot{p} = \hat{p}_2$ where $\hat{h}$ and $\hat{p}$ verify,

$$\frac{h_{\max}}{\hat{h}_2} = C \kappa^{\hat{h}_2} \quad (16)$$

$$\hat{h}_2 = \frac{1}{\log(\kappa)} W\left(\frac{h_{\max} \log(\kappa)}{C}\right)$$

By construction, we have $\tilde{h}_2 \leq \tilde{h}$. We set

$$\hat{p}_2 = \max(0, \tilde{p}). \quad (17)$$

Using standard properties of the $\lfloor \cdot \rfloor$ function, we have

$$b \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{2^{\lfloor \hat{p}_2 \rfloor + 1}}} \leq b \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{2^{\tilde{p}}}} = \nu \rho^{\tilde{h}} \stackrel{(a)}{\leq} \nu \rho^{\hat{h}_2} \leq \nu \rho^{\lfloor \hat{h}_2 \rfloor} \quad (18)$$

where (a) is because of the following reasoning. As we have $\frac{h_{\max} \nu^2 \rho^{2\tilde{h}}}{\tilde{h}^2 b^2 g_{n,b}^{\delta, R_{\max}} \gamma^{2\tilde{h}}} = C \kappa^{\tilde{h}}$ and $\frac{\nu^2 \rho^{2\tilde{h}}}{\gamma^{2\tilde{h}} \tilde{h} b^2 g_{n,b}^{\delta, R_{\max}}} \geq 1$, then, $\frac{h_{\max}}{\tilde{h}} \leq \kappa^{\tilde{h}}$. From the inequality $\frac{h_{\max}}{\tilde{h}} \leq C \kappa^{\tilde{h}}$ and the fact that $\hat{h}_2$ corresponds to the case of equality $\frac{h_{\max}}{2\hat{h}_2} = C \kappa^{\hat{h}_2}$, we deduce that $\hat{h}_2 \leq \tilde{h}$, since the left term of the inequality decreases with h while the right term increases. Having $\hat{h}_2 \leq \tilde{h}$ gives $\rho^{\hat{h}_2} \geq \rho^{\tilde{h}}$.

Step 3 Given these particular definitions of $\tilde{h}$ and $\tilde{p}$ in two distinct cases we now bound the regret.

We always have $\perp_{h_{\max}, \lfloor \tilde{p} \rfloor} \geq 0$. If $\tilde{h} \geq 1$, as discussed above $\lfloor \tilde{h} \rfloor \in [h_{\max}]$, therefore $\perp_{h_{\max}, \lfloor \tilde{p} \rfloor} \geq \perp_{\lfloor \tilde{h} \rfloor, \lfloor \tilde{p} \rfloor}$, as $\perp_{\cdot, \lfloor p \rfloor}$ is increasing for all $p \in [0, p_{\max}]$. Moreover on event ξ , and for the cases 1.1.1, 1.1.2, 2.1.1 and 2.1.2 described above, $\perp_{\lfloor \tilde{h} \rfloor, \lfloor \tilde{p} \rfloor} = \lfloor \tilde{h} \rfloor$ because of Lemma 1 (Case 2)) which assumptions on $\lfloor \tilde{h} \rfloor$ and $\lfloor \tilde{p} \rfloor$ are verified in each cases as detailed above and, in general, $\lfloor \tilde{h} \rfloor \in [h_{\max}/2^{\tilde{p}}]$ and $\lfloor \tilde{p} \rfloor \in [0 : p_{\max}]$. So, for the aforementioned cases, we have $\perp_{\lfloor h_{\max}/2^{\tilde{p}} \rfloor, \lfloor \tilde{p} \rfloor} \geq \lfloor \tilde{h} \rfloor$. Very similarly cases 1.2 and 2.2. lead to $\perp_{\lfloor h_{\max}/2^{\tilde{p}} \rfloor, \lfloor \tilde{p} \rfloor} \geq \lfloor \tilde{h} \rfloor$ by using Lemma 1 (Case 1)).

We bound the regret now discriminating on whether or not the event ξ holds. We have

$$\begin{aligned} r_n &\leq (1 - \delta) \left(\nu \rho^{\perp_{h_{\max}, \tilde{p}} + 1} + 2 \frac{b}{1 - \gamma^2} \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{2 h_{\max}}} \right) + \delta \times \frac{R_{\max}}{1 - \gamma} \\ &\leq \nu \rho^{\perp_{h_{\max}, \tilde{p}} + 1} + 2 \frac{b}{1 - \gamma^2} \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{2 h_{\max}}} + \frac{4b}{\sqrt{n}} \\ &\leq \nu \rho^{\perp_{h_{\max}, \tilde{p}} + 1} + 6 \frac{b}{1 - \gamma^2} \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}}. \end{aligned}$$

We can now bound the regret in the two regimes.

Case 1) $\gamma^2 \kappa \geq 1$: Note that in this case then $\kappa > 1$. We subdivide this case into multiple subcases:

Case 1.1) Noise regime $\frac{\nu^2 \rho^{2\tilde{h}_1}}{\gamma^{2\tilde{h}_1} \tilde{h}_1 b^2 g_{n,b}^{\delta, R_{\max}}} \leq 1$

Case 1.1.1) High-noise regime In general, we have

$$r_n \leq \nu \rho^{\frac{2}{\log(\gamma^2 \kappa / \rho^2)} W \left(\log(\gamma^2 \kappa / \rho^2) / 2 \sqrt{\frac{\gamma^2 \nu^2 h_{\max}}{2 C b^2 g_{n,b}^{\delta, R_{\max}}}} \right)} + 6 \frac{b}{1 - \gamma} \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}}.$$

Moreover, as proved by [Hoorfar & Hassani \(2008\)](#), the Lambert $W(x)$ function verifies for $x \geq e$, $W(x) \geq \log \left(\frac{x}{\log x} \right)$.

Therefore, if $\log(\gamma^2 \kappa / \rho^2) / 2 \sqrt{\frac{\gamma^2 \nu^2 h_{\max}}{2Cb^2 g_{n,b}^{\delta, R_{\max}}}} > e$ we have,

$$\begin{aligned}
 r_n & - 6 \frac{b}{1-\gamma} \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}} \\
 & \frac{2}{\log(\gamma^2 \kappa / \rho^2)} \left(\log \left(\frac{\log(\gamma^2 \kappa / \rho^2) / 2 \sqrt{\frac{\gamma^2 \nu^2 h_{\max}}{2Cb^2 g_{n,b}^{\delta, R_{\max}}}}}{\log \left(\log(\gamma^2 \kappa / \rho^2) / 2 \sqrt{\frac{\gamma^2 \nu^2 h_{\max}}{2Cb^2 g_{n,b}^{\delta, R_{\max}}}} \right)} \right) \right) \\
 & \leq \nu \rho \\
 & \frac{1}{\log(\gamma^2 \kappa / \rho^2)} \left(\log \left(\frac{\log^2(\gamma^2 \kappa / \rho^2) / 2 \frac{\gamma^2 \nu^2 h_{\max}}{2Cb^2 g_{n,b}^{\delta, R_{\max}}}}{\log^2 \left(\log(\gamma^2 \kappa / \rho^2) / 2 \sqrt{\frac{\gamma^2 \nu^2 h_{\max}}{2Cb^2 g_{n,b}^{\delta, R_{\max}}}} \right)} \right) \right) \log(\rho) \\
 & = \nu e \\
 & = \nu \left(\frac{\log^2(\gamma^2 \kappa / \rho^2) / 2 \frac{\gamma^2 \nu^2 h_{\max}}{2Cb^2 g_{n,b}^{\delta, R_{\max}}}}{\log^2 \left(\log(\gamma^2 \kappa / \rho^2) / 2 \sqrt{\frac{\gamma^2 \nu^2 h_{\max}}{2Cb^2 g_{n,b}^{\delta, R_{\max}}}} \right)} \right)^{\frac{\log(\rho)}{\log(\gamma^2 \kappa / \rho^2)}}.
 \end{aligned}$$

Case 1.1.2) Low-noise regime 1 $\frac{\nu^2 \rho^{2\tilde{h}_1}}{b^2 g_{n,b}^{\delta, R_{\max}}} \geq 1$

$$r_n \leq \nu \rho^{\frac{2}{\log(\gamma^2 \kappa)}} W\left(\log(\gamma^2 \kappa) / 2 \sqrt{\frac{h_{\max}}{2C}}\right) + 6 \frac{b}{1-\gamma} \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}}.$$

Moreover, as proved by [Hoorfar & Hassani \(2008\)](#), the Lambert $W(x)$ function verifies for $x \geq e$, $W(x) \geq \log\left(\frac{x}{\log x}\right)$.

Therefore, if $\log(\gamma^2 \kappa) / 2 \sqrt{\frac{h_{\max}}{2C}} > e$ we have,

$$\begin{aligned}
 r_n & - 6 \frac{b}{1-\gamma} \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}} \\
 & \frac{2}{\log(\gamma^2 \kappa)} \left(\log \left(\frac{\log(\gamma^2 \kappa) / 2 \sqrt{\frac{h_{\max}}{2C}}}{\log \left(\log(\gamma^2 \kappa) / 2 \sqrt{\frac{h_{\max}}{2C}} \right)} \right) \right) \\
 & \leq \nu \rho \\
 & \frac{1}{\log(\gamma^2 \kappa)} \left(\log \left(\frac{\log^2(\gamma^2 \kappa) / 2 \frac{h_{\max}}{2C}}{\log^2 \left(\log(\gamma^2 \kappa) / 2 \sqrt{\frac{h_{\max}}{2C}} \right)} \right) \right) \log(\rho) \\
 & = \nu e \\
 & = \nu \left(\frac{\log^2(\gamma^2 \kappa) / 2 \frac{h_{\max}}{2C}}{\log^2 \left(\log(\gamma^2 \kappa) / 2 \sqrt{\frac{h_{\max}}{2C}} \right)} \right)^{\frac{\log(\rho)}{\log(\gamma^2 \kappa)}}.
 \end{aligned}$$

We have $6b \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}} \leq 6 \frac{\nu \rho^{\tilde{h}_1}}{\sqrt{g_{n,b}^{\delta, R_{\max}}}} \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}} \leq 6\nu \rho^{\tilde{h}_1} \leq 6\nu \rho^{\bar{h}_1}$.

Therefore $r_n \leq \nu \rho^{\perp_{h_{\max}, \bar{p}}+1} + 6b \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}} \leq 7\nu \rho^{\bar{h}_1}$.

Case 1.2) Low noise regime 2 $\frac{\nu^2 \rho^{2\tilde{h}_1}}{\gamma^{2\tilde{h}_1} h_1 b^2 g_{n,b}^{\delta, R_{\max}}} \geq 1$

$$r_n - 6 \frac{b}{1-\gamma} \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}} \leq \nu \left(\frac{\frac{h_{\max} \log(\kappa)}{2C}}{\log \left(\frac{h_{\max} \log(\kappa)}{2C} \right)} \right)^{\frac{\log(\rho)}{\log(\kappa)}}.$$

Moreover if $\frac{\nu^2 \rho^{2\tilde{h}_1}}{b^2 g_{n,b}^{\delta, R_{\max}}} \geq 1$, we have $6b \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}} \leq 6 \frac{\nu \rho^{\tilde{h}_1}}{\sqrt{g_{n,b}^{\delta, R_{\max}}}} \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}} \leq 6\nu \rho^{\tilde{h}_1} \leq 6\nu \rho^{\hat{h}_1}$.

Therefore $r_n \leq \nu \rho^{\perp_{h_{\max}, \bar{p}}+1} + 6b \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}} \leq 7\nu \rho^{\hat{h}_1}$.

Case 2) $\gamma^2 \kappa \leq 1$

Case 2.1) Noise regime $\frac{\nu^2 \rho^{2\tilde{h}}}{\gamma^{2\tilde{h}} h b^2 g_{n,b}^{\delta, R_{\max}}} \leq 1$

Case 2.1.1) High-noise regime $\frac{\nu^2 \rho^{2\tilde{h}_2}}{b^2 g_{n,b}^{\delta, R_{\max}}} \leq 1$

$$r_n - 6 \frac{b}{1-\gamma} \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}} \leq \nu \left(\frac{\log^2(1/\rho^2)/2 \frac{\nu^2 h_{\max}}{2C b^2 g_{n,b}^{\delta, R_{\max}}}}{\log^2 \left(\log(1/\rho^2)/2 \sqrt{\frac{\nu^2 h_{\max}}{2C b^2 g_{n,b}^{\delta, R_{\max}}}} \right)} \right)^{-\frac{1}{2}}.$$

Case 2.1.2) Low-noise regime 1 $\frac{\nu^2 \rho^{2\tilde{h}_2}}{b^2 g_{n,b}^{\delta, R_{\max}}} \geq 1$

Here with a similar reasoning as in the case 1.1.2) we have $r_n \leq 7\nu \rho^{\bar{h}_1} \leq 7\nu \rho^{\sqrt{\frac{h_{\max}}{2C}}}$.

Case 2.2) Low noise regime 2 $\frac{\nu^2 \rho^{2\tilde{h}}}{\gamma^{2\tilde{h}} h b^2 g_{n,b}^{\delta, R_{\max}}} \geq 1$

$$r_n - 6 \frac{b}{1-\gamma} \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}} \leq \nu \left(\frac{\frac{h_{\max} \log(\kappa)}{C}}{\log \left(\frac{h_{\max} \log(\kappa)}{C} \right)} \right)^{\frac{\log(\rho)}{\log(\kappa)}}.$$

Moreover if $\frac{\nu^2 \rho^{2\tilde{h}}}{b^2 g_{n,b}^{\delta, R_{\max}}} \geq 1$, we have $6b \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}} \leq 6 \frac{\nu \rho^{\tilde{h}}}{\sqrt{g_{n,b}^{\delta, R_{\max}}}} \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}} \leq 6\nu \rho^{\tilde{h}} \leq 6\nu \rho^{\hat{h}_2}$.

Therefore $r_n \leq \nu \rho^{\perp_{h_{\max}, \bar{p}}+1} + 6b \sqrt{\frac{g_{n,b}^{\delta, R_{\max}}}{h_{\max}}} \leq 7\nu \rho^{\hat{h}_1}$.

Moreover if $\kappa = 1$ then $r_n \leq 7\nu \rho^{\frac{h_{\max}}{C}}$ □

E. Use of the budget

Remark 1. The algorithm can be made anytime and agnostic to n using the standard doubling trick.

Remark 2 (More efficient use of the budget). Because of the use of the floor functions $\lfloor \cdot \rfloor$, the budget used in practice can be significantly smaller than n . While this only affects numerical constants in the bounds, in practice, it can noticeably influence the performance. Therefore one should consider, for instance, having $h_{\max}$ replaced by $c \times h_{\max}$ with c been the largest number such that the budget is still smaller than n . Additionally, the use of the budget n could be slightly optimized by taking into account that the necessary number of pulls at depth h cannot be larger than K^h .