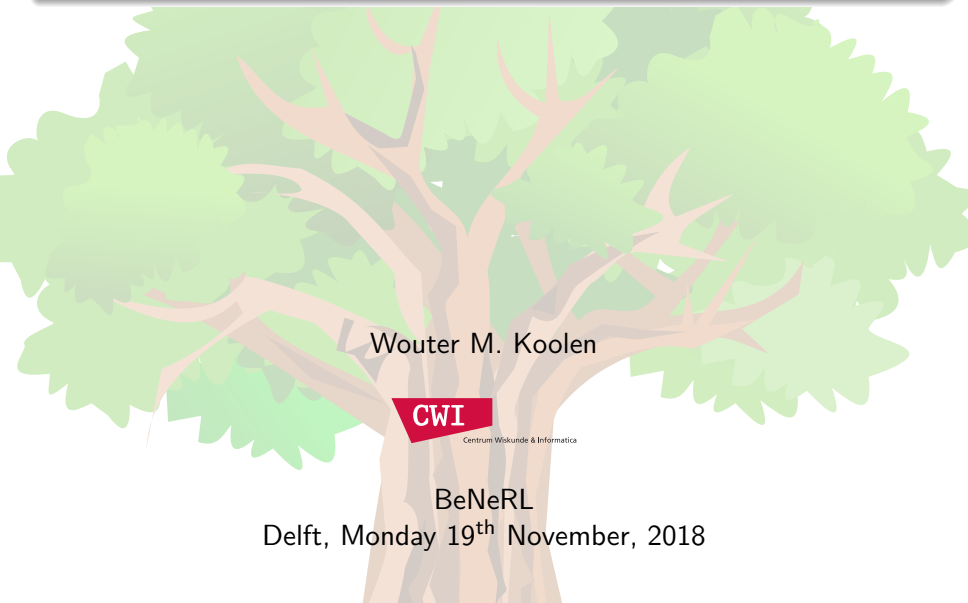


A Pure Exploration perspective on Game Tree Search



Wouter M. Koolen

CWI

Centrum Wiskunde & Informatica

BeNeRL

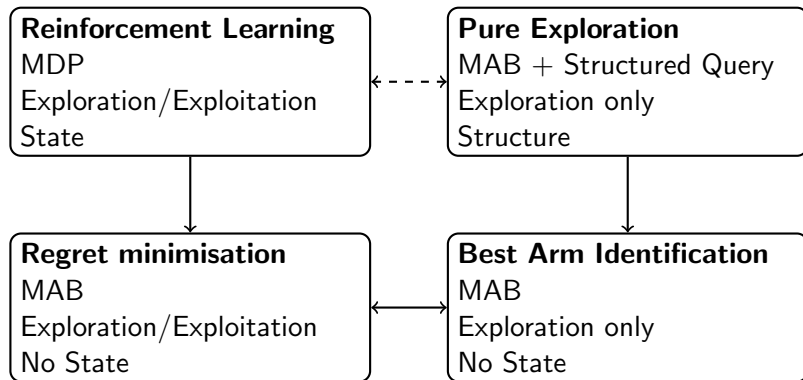
Delft, Monday 19th November, 2018

Menu

- What is Pure Exploration?
- Relation to Reinforcement Learning?
- Why care?
 - ▶ Pure Exploration problems occur as **sub-problems** in RL!
 - ▶ Novel/interesting/powerful PE learning algorithms! “Fresh”
- My focus: Pure Exploration **Renaissance** (2016)
 - ① Track-and-Stop algorithm for Best Arm Identification
 - ② BAI-MCTS approach for Game Tree Search
 - ③ Murphy Sampling for Games Trees of “depth 1.5”

- 1 Introduction
- 2 Relation of RL and PE
- 3 Pure Exploration Intro: Best Arm Identification
 - Model
 - Sample Complexity Lower Bound
 - Algorithms
- 4 Game Tree Search
 - Game Trees of Arbitrary Depth
 - Confidence Intervals on Min/Max
 - Game Trees of Depth 1.5 (Maximum/Minimum)
 - Results
- 5 Conclusion

Relation of RL and PE



Example 1: Phased Q-Learning

[Even-Dar, Mannor, and Mansour, 2002]

Early example(s) of Pure Exploration as sub-module in RL

```

Initialise  $V_0(s) := 0$  for each state  $s$ .
for phase  $i = 1, 2, \dots$  do
  for each state  $s$  do
    Run Best Arm Identification algorithm on
       $a \mapsto r + \gamma V_i(s')$  where  $(r, s') \sim \mathbb{P}(r, s' | s, a)$ 

    Store estimate in  $V_{i+1}(s)$ .
  end for
end for
  
```

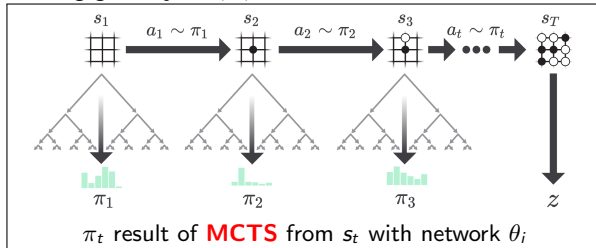
Example 2: AlphaZero

MCTS as Policy/Value Improvement Operator

Randomise network parameters θ_1 .

for iteration $i = 1, 2, \dots$ **do**

for training games $j = 1, 2, \dots$ **do**



Store $(s_1, \pi_1, z), \dots, (s_T, \pi_T, z)$.

end for

Train network θ_{i+1} on stored data.

end for

Formal model



Environment (Multi-armed bandit model)

K distributions parameterised by their means $\boldsymbol{\mu} = (\mu_1, \dots, \mu_K)$.

The **best arm** is

$$i^* = \operatorname{argmax}_{i \in [K]} \mu_i$$

Formal model



Environment (Multi-armed bandit model)

K distributions parameterised by their means $\boldsymbol{\mu} = (\mu_1, \dots, \mu_K)$.

The **best arm** is

$$i^* = \operatorname{argmax}_{i \in [K]} \mu_i$$

Strategy

- **Stopping rule** $\tau \in \mathbb{N}$
- In round $t \leq \tau$ **sampling rule** picks $I_t \in [K]$. See $X_t \sim \mu_{I_t}$.
- **Recommendation rule** $\hat{I} \in [K]$.

Formal model



Environment (Multi-armed bandit model)

K distributions parameterised by their means $\boldsymbol{\mu} = (\mu_1, \dots, \mu_K)$.

The **best arm** is

$$i^* = \operatorname{argmax}_{i \in [K]} \mu_i$$

Strategy

- **Stopping rule** $\tau \in \mathbb{N}$
- In round $t \leq \tau$ **sampling rule** picks $I_t \in [K]$. See $X_t \sim \mu_{I_t}$.
- **Recommendation rule** $\hat{I} \in [K]$.

Realisation of interaction: $(I_1, X_1), \dots, (I_\tau, X_\tau), \hat{I}$.

Formal model



Environment (Multi-armed bandit model)

K distributions parameterised by their means $\mu = (\mu_1, \dots, \mu_K)$.

The **best arm** is

$$i^* = \operatorname{argmax}_{i \in [K]} \mu_i$$

Strategy

- **Stopping rule** $\tau \in \mathbb{N}$
- In round $t \leq \tau$ **sampling rule** picks $I_t \in [K]$. See $X_t \sim \mu_{I_t}$.
- **Recommendation rule** $\hat{I} \in [K]$.

Realisation of interaction: $(I_1, X_1), \dots, (I_\tau, X_\tau), \hat{I}$.

Two objectives: **sample efficiency** τ and **correctness** $\hat{I} = i^*$.



Objective

On bandit μ , strategy $(\tau, (I_t)_t, \hat{I})$ has

- **error probability** $\mathbb{P}_\mu(\hat{I} \neq i^*(\mu))$, and
- **sample complexity** $\mathbb{E}_\mu[\tau]$.

Idea: constrain one, optimise the other.



Objective

On bandit μ , strategy $(\tau, (I_t)_t, \hat{I})$ has

- **error probability** $\mathbb{P}_\mu(\hat{I} \neq i^*(\mu))$, and
- **sample complexity** $\mathbb{E}_\mu[\tau]$.

Idea: constrain one, optimise the other.

Definition

Fix small confidence $\delta \in (0, 1)$. A strategy is δ -**correct** if

$$\mathbb{P}_\mu(\hat{I} \neq i^*(\mu)) \leq \delta \quad \text{for every bandit model } \mu.$$

(Generalisation: output ϵ -best arm)



Objective

On bandit μ , strategy $(\tau, (I_t)_t, \hat{I})$ has

- **error probability** $\mathbb{P}_\mu(\hat{I} \neq i^*(\mu))$, and
- **sample complexity** $\mathbb{E}_\mu[\tau]$.

Idea: constrain one, optimise the other.

Definition

Fix small confidence $\delta \in (0, 1)$. A strategy is δ -**correct** if

$$\mathbb{P}_\mu(\hat{I} \neq i^*(\mu)) \leq \delta \quad \text{for every bandit model } \mu.$$

(Generalisation: output ϵ -best arm)

Goal: minimise $\mathbb{E}_\mu[\tau]$ over **all δ -correct strategies**.



Algorithms

- Sampling rule I_t ?
- Stopping rule τ ?
- Recommendation rule $\hat{I}$?

$$\hat{I} = \operatorname{argmax}_{i \in [K]} \hat{\mu}_i(\tau)$$

where $\hat{\mu}(t)$ is **empirical mean**.

Instance-Dependent Sample Complexity Lower bound

Define the **alternatives** to μ by $\text{Alt}(\mu) = \{\lambda | i^*(\lambda) \neq i^*(\mu)\}$.

Instance-Dependent Sample Complexity Lower bound

Define the **alternatives** to μ by $\text{Alt}(\mu) = \{\lambda \mid i^*(\lambda) \neq i^*(\mu)\}$.

Theorem (Castro 2014, Garivier and Kaufmann 2016)

Fix a δ -correct strategy. Then for every bandit model μ

$$\mathbb{E}_{\mu}[\tau] \geq T^*(\mu) \ln \frac{1}{\delta}$$

*where the **characteristic time** $T^*(\mu)$ is given by*

$$\frac{1}{T^*(\mu)} = \max_{w \in \Delta_K} \min_{\lambda \in \text{Alt}(\mu)} \sum_{i=1}^K w_i \text{KL}(\mu_i \parallel \lambda_i).$$

Instance-Dependent Sample Complexity Lower bound

Define the **alternatives** to μ by $\text{Alt}(\mu) = \{\lambda \mid i^*(\lambda) \neq i^*(\mu)\}$.

Theorem (Castro 2014, Garivier and Kaufmann 2016)

Fix a δ -correct strategy. Then for every bandit model μ

$$\mathbb{E}_{\mu}[\tau] \geq T^*(\mu) \ln \frac{1}{\delta}$$

*where the **characteristic time** $T^*(\mu)$ is given by*

$$\frac{1}{T^*(\mu)} = \max_{w \in \Delta_K} \min_{\lambda \in \text{Alt}(\mu)} \sum_{i=1}^K w_i \text{KL}(\mu_i \parallel \lambda_i).$$

Intuition (going back to Lai and Robbins [1985]): if observations are likely under both μ and λ , yet $i^*(\mu) \neq i^*(\lambda)$, then learner cannot stop and be correct in both.

Example

$K = 5$ arms, Bernoulli $\mu = (0, 0.1, 0.2, 0.3, 0.4)$.

$$T^*(\mu) = 200.4 \quad w^*(\mu) = (0.45, 0.46, 0.06, 0.02, 0.01)$$

At $\delta = 0.05$, the time gets multiplied by $\ln \frac{1}{\delta} = 3.0$.

Sampling Rule

Look at the lower bound again. Any good algorithm **must** sample with optimal (**oracle**) proportions

$$w^*(\mu) = \operatorname{argmax}_{w \in \Delta_K} \min_{\lambda \in \text{Alt}(\mu)} \sum_{i=1}^K w_i \text{KL}(\mu_i \| \lambda_i)$$

Sampling Rule

Look at the lower bound again. Any good algorithm **must** sample with optimal (**oracle**) proportions

$$w^*(\mu) = \operatorname{argmax}_{w \in \Delta_K} \min_{\lambda \in \text{Alt}(\mu)} \sum_{i=1}^K w_i \text{KL}(\mu_i \| \lambda_i)$$

Track-and-Stop

Idea: draw $I_t \sim w^*(\hat{\mu}(t))$.

- Ensure $\hat{\mu}(t) \rightarrow \mu$ hence $N_i(t)/t \rightarrow w_i^*$ by “forced exploration”
- Draw arm with $N_i(t)/t$ below w_i^* (tracking)
- Computation of w^* (reduction to 1d line search)

All in all

Final result: lower and upper bound meet on **every problem instance**.

Theorem (Garivier and Kaufmann 2016)

For **Track-and-Stop** algorithm, for any bandit μ

$$\limsup_{\delta \rightarrow 0} \frac{\mathbb{E}_{\mu}[\tau]}{\ln \frac{1}{\delta}} = T^*(\mu)$$

All in all

Final result: lower and upper bound meet on **every problem instance**.

Theorem (Garivier and Kaufmann 2016)

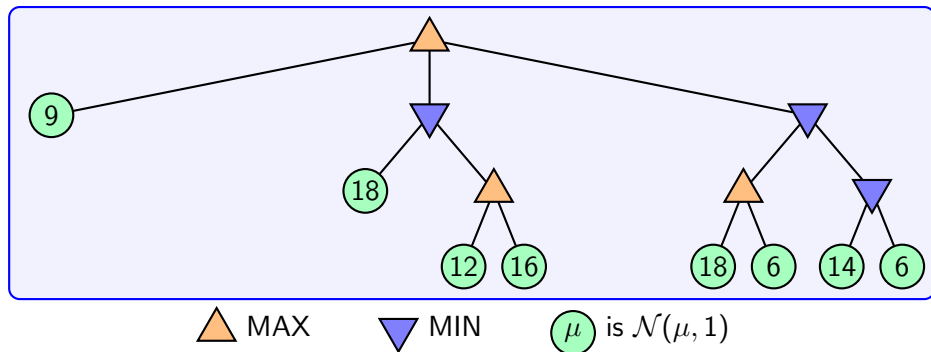
For **Track-and-Stop** algorithm, for any bandit μ

$$\limsup_{\delta \rightarrow 0} \frac{\mathbb{E}_{\mu}[\tau]}{\ln \frac{1}{\delta}} = T^*(\mu)$$

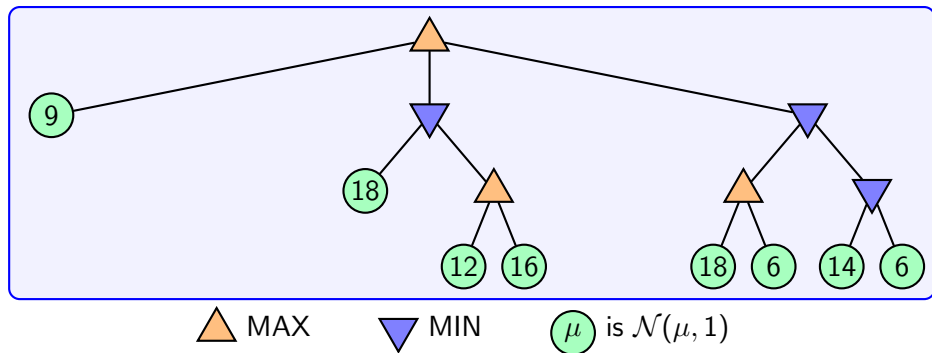
Very similar optimality result for **Top Two Thompson Sampling** by Russo [2016]. Here $N_i(t)/t \rightarrow w_i^*$ result of posterior sampling.

- 1 Introduction
- 2 Relation of RL and PE
- 3 Pure Exploration Intro: Best Arm Identification
 - Model
 - Sample Complexity Lower Bound
 - Algorithms
- 4 Game Tree Search
 - Game Trees of Arbitrary Depth
 - Confidence Intervals on Min/Max
 - Game Trees of Depth 1.5 (Maximum/Minimum)
 - Results
- 5 Conclusion

Challenge Environment: Stochastic Game Tree Search

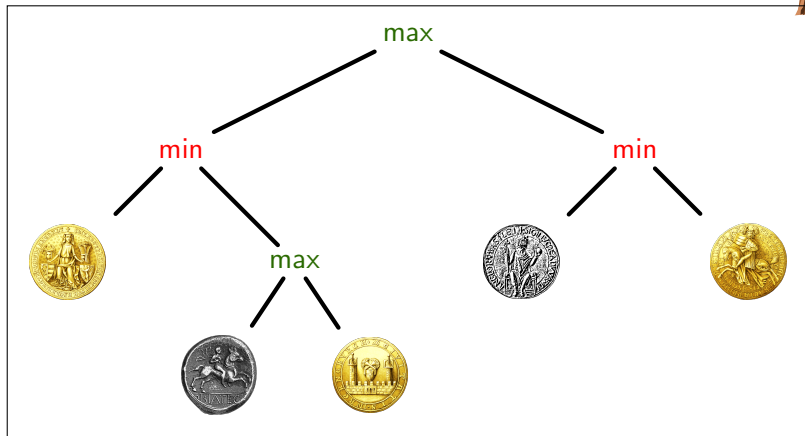


Challenge Environment: Stochastic Game Tree Search



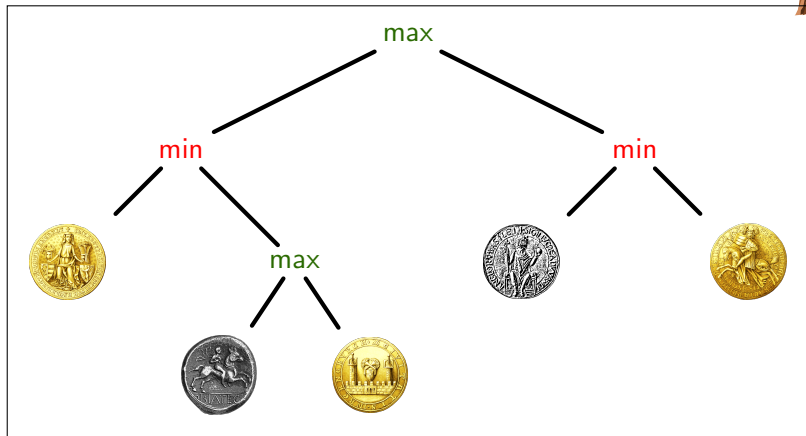
Problem

Determine the **optimal move** at the root
 From **sample access** to the leaf payoffs



Maximin Action Identification Problem

Find **best move at root** from samples of **leaves**.



Maximin Action Identification Problem

Find **best move at root** from samples of **leaves**.

guarantee?

Methods for Best Arm Identification



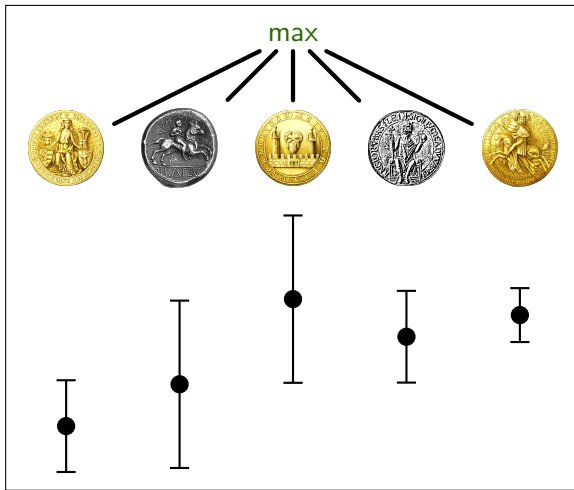
Methods for Best Arm Identification



LUCB,
UGapE,

...

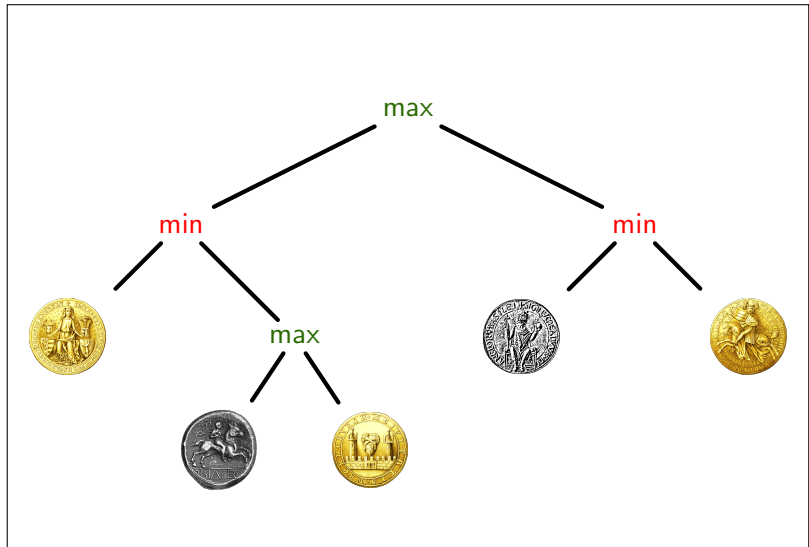
Methods for Best Arm Identification



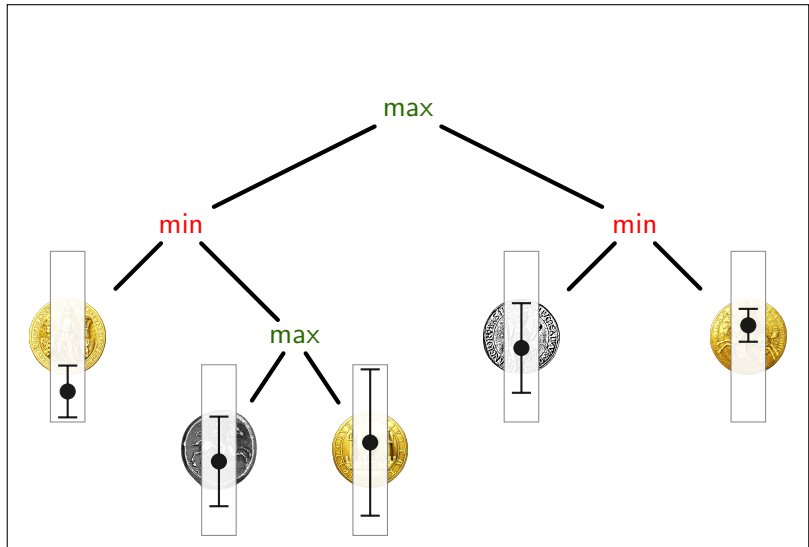
LUCB,
UGapE,

...

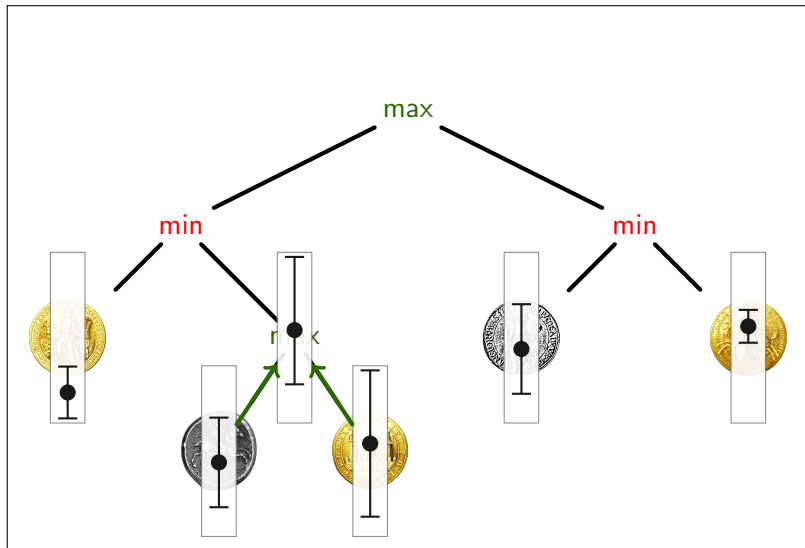
Reduction of MCTS to BAI



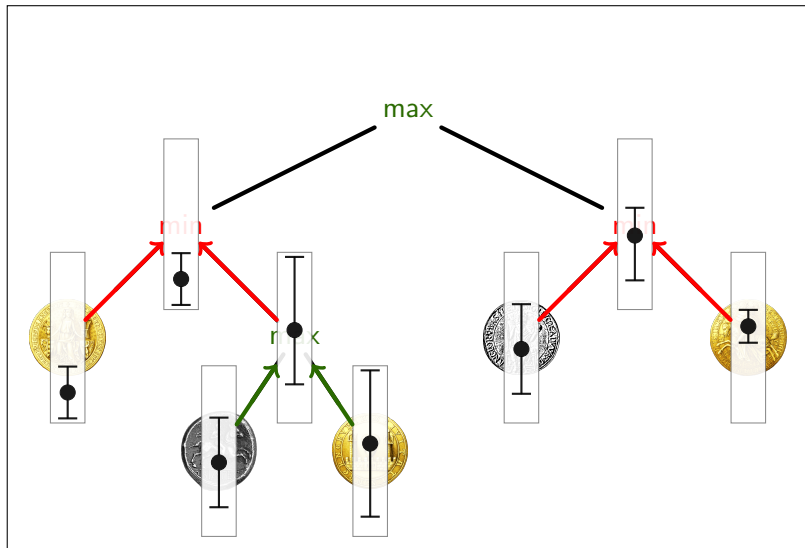
Reduction of MCTS to BAI



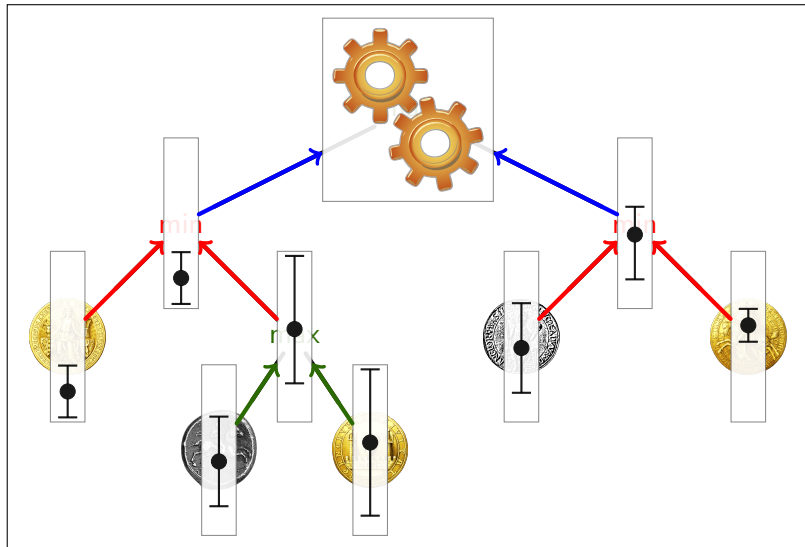
Reduction of MCTS to BAI



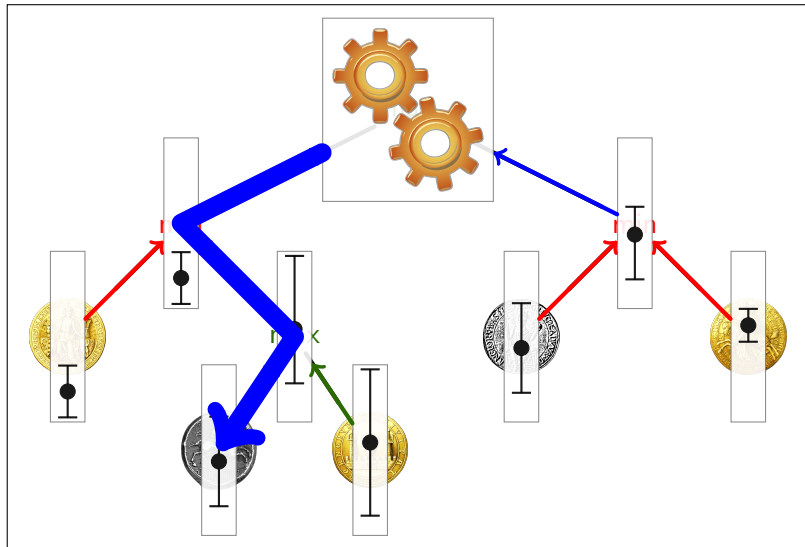
Reduction of MCTS to BAI



Reduction of MCTS to BAI



Reduction of MCTS to BAI





Correctness

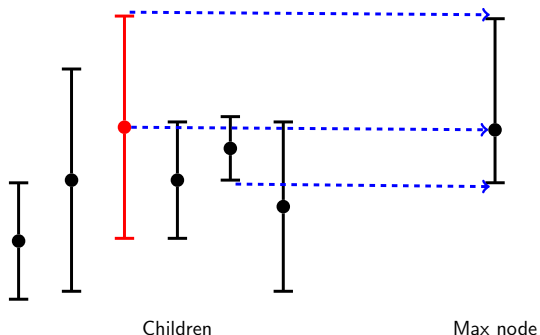
(ϵ, δ) -PAC algorithms

Efficiency

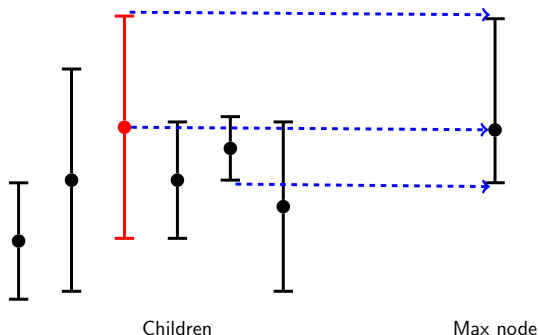
Sample complexity function of **leaf gaps** Δ_ℓ

$$O\left(\sum_{\ell \in \mathcal{L}} \frac{1}{\Delta_\ell^2 \vee \epsilon^2} \log\left(\frac{1}{\delta}\right)\right)$$

How to build confidence intervals on min/max nodes



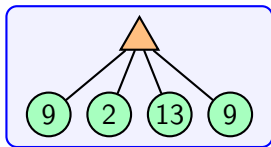
How to build confidence intervals on min/max nodes



More principled approach in [Kaufmann, Koolen, and Garivier, 2018].
 Many equal intervals $\Rightarrow$ higher lower bound.

- 1 Introduction
- 2 Relation of RL and PE
- 3 Pure Exploration Intro: Best Arm Identification
 - Model
 - Sample Complexity Lower Bound
 - Algorithms
- 4 Game Tree Search
 - Game Trees of Arbitrary Depth
 - Confidence Intervals on Min/Max
 - Game Trees of Depth 1.5 (Maximum/Minimum)
 - Results
- 5 Conclusion

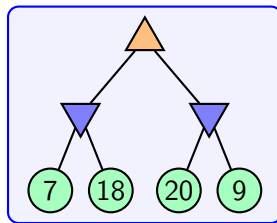
Simplify



Best Arm Identification

[Garivier and Kaufmann, 2016]

Solved

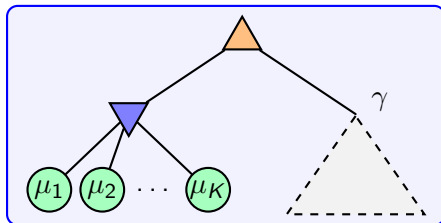


Depth 2

[Garivier, Kaufmann, and Koolen, 2016]

Open

Simple Instance: Minimum Threshold Identification



Fix threshold γ .

$$\mu^* := \min_i \mu_i \leq \gamma?$$

For $t = 1, \dots, \tau$

- Pick leaf A_t
- See $X_t \sim \mu_{A_t}$

Recommend $\hat{m} \in \{<, >\}$

Goal: **fixed confidence** $\mathbb{P}_\mu \{\text{error}\} < \delta$
and small **sample complexity** $\mathbb{E}_\mu[\tau]$



Lower Bound

Generic lower bound [Castro, 2014, Garivier and Kaufmann, 2016] shows **sample complexity** for **any** δ -correct algorithm is at least

$$\mathbb{E}_{\mu}[\tau] \geq T^*(\mu) \ln \frac{1}{\delta}.$$

Lower Bound

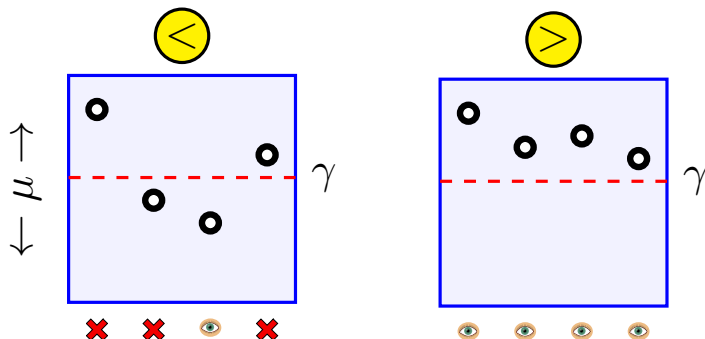
Generic lower bound [Castro, 2014, Garivier and Kaufmann, 2016] shows **sample complexity** for **any** δ -correct algorithm is at least

$$\mathbb{E}_{\mu}[\tau] \geq T^*(\mu) \ln \frac{1}{\delta}.$$

For our problem the **characteristic time** and **oracle weights** are

$$T^*(\mu) = \begin{cases} \frac{1}{\text{KL}(\mu^*, \gamma)} & \mu^* < \gamma, \\ \sum_a \frac{1}{\text{KL}(\mu_a, \gamma)} & \mu^* > \gamma, \end{cases} \quad w_a^*(\mu) = \begin{cases} \mathbf{1}_{a=a^*} & \mu^* < \gamma, \\ \frac{1}{\sum_j \frac{1}{\text{KL}(\mu_j, \gamma)}} & \mu^* > \gamma. \end{cases}$$

Dichotomous Oracle Behaviour! Sampling Rule?



Sampling Rules

- **Lower Confidence Bounds**

Play $A_t = \arg \min_a \text{LCB}_a(t)$

- **Thompson Sampling** (Π_{t-1} is posterior after $t - 1$ rounds)

Sample $\theta \sim \Pi_{t-1}$, then play $A_t = \arg \min_a \theta_a$.

- **Murphy Sampling** **condition on low minimum mean**

Sample $\theta \sim \Pi_{t-1} (\cdot | \min_a \theta_a < \gamma)$, then play $A_t = \arg \min_a \theta_a$.



Intuition for Murphy Sampling

- When $\mu^* < \gamma$ conditioning is immaterial: $\theta \approx \mu$ and $MS \equiv TS$.
- When $\mu^* > \gamma$ conditioning results in $\theta \approx (\mu_1, \dots, \gamma, \dots, \mu_K)$.
Index a lowered to γ with probability $\propto \frac{1}{KL(\mu_a, \gamma)}$ [Russo, 2016].

Murphy Sampling Rule [KKG, NIPS'18]

Theorem

Asymptotic optimality: $N_a(t)/t \rightarrow w_a^*(\mu)$ for all μ

Sampling rule		
Thompson Sampling		
Lower Confidence Bounds		
Murphy Sampling		

Lemma

Any anytime sampling strategy $(A_t)_t$ ensuring $\frac{N_t}{t} \rightarrow w^*(\mu)$ and good stopping rule τ_δ guarantee $\limsup_{\delta \rightarrow 0} \frac{\tau_\delta}{\ln \frac{1}{\delta}} \leq T^*(\mu)$.

Conclusion

- Pure Exploration currently going through a renaissance
- Instance-optimal identification algorithms
 - ▶ Best Arm
 - ▶ Game Tree Search
 - ▶ ...
- Moving toward more complex queries. RL on the horizon ...
- Useful submodules

Conclusion

- Pure Exploration currently going through a renaissance
- Instance-optimal identification algorithms
 - ▶ Best Arm
 - ▶ Game Tree Search
 - ▶ ...
- Moving toward more complex queries. RL on the horizon ...
- Useful submodules

Thank you! And let's talk!