

Robot Perception and Learning

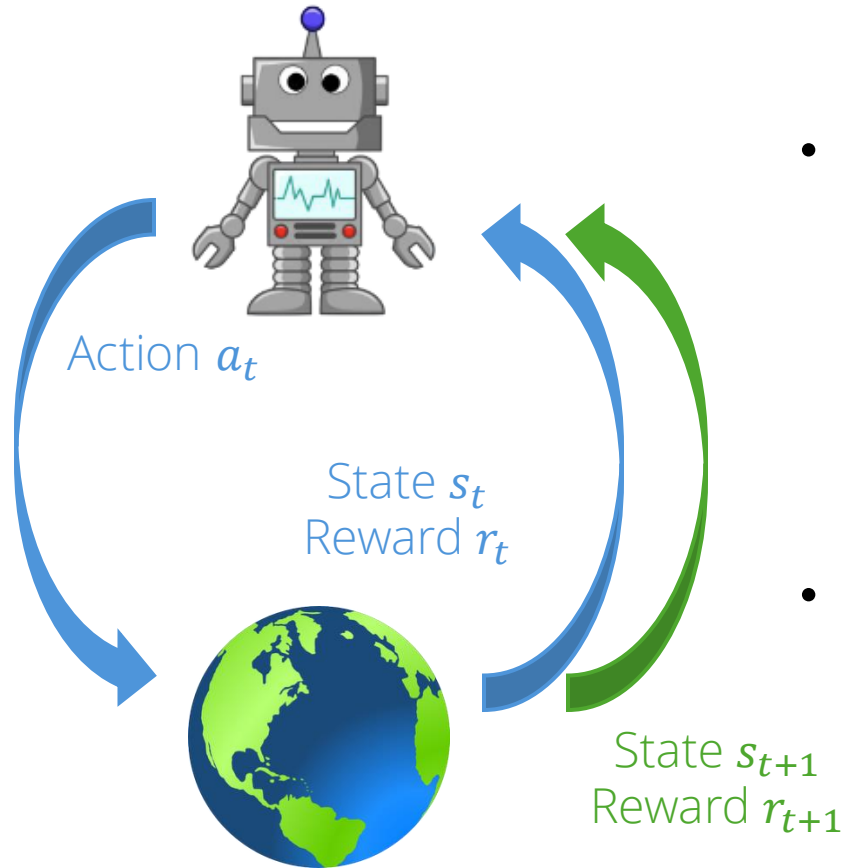
Model Based RL

Tsung-Wei Ke

Fall 2025



Recap: the Objective of Reinforcement Learning



- A trajectory of interaction in the environment

$$\underbrace{p_{\theta}(\mathbf{s}_1, \mathbf{a}_1, \dots, \mathbf{s}_T, \mathbf{a}_T)}_{p_{\theta}(\tau)} = p(\mathbf{s}_1) \prod_{t=1}^T \pi_{\theta}(\mathbf{a}_t | \mathbf{s}_t) \boxed{p(\mathbf{s}_{t+1} | \mathbf{s}_t, \mathbf{a}_t)}$$

The dynamic function of the environment

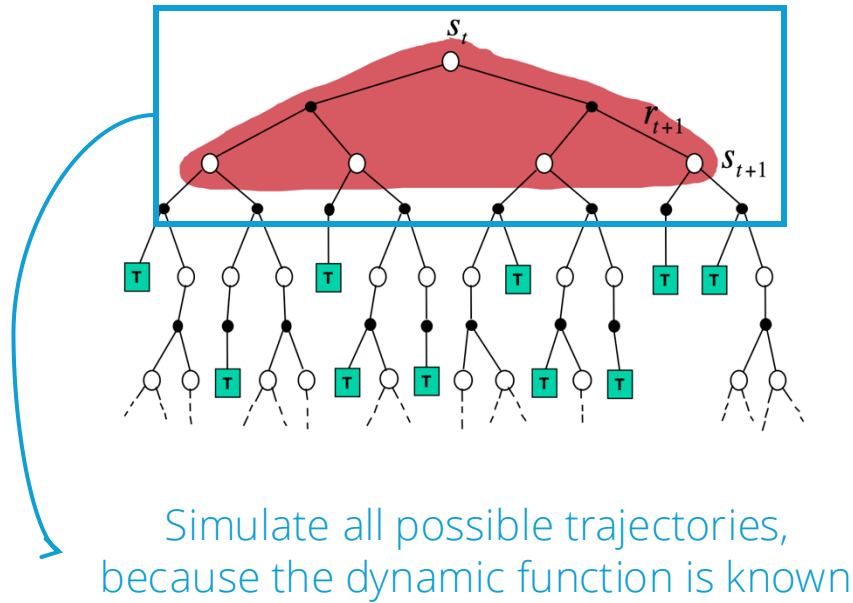
- Maximize the expected value of the cumulative sum of reward

$$\theta^* = \arg \max_{\theta} E_{\tau \sim p_{\theta}(\tau)} \left[\sum_t r(\mathbf{s}_t, \mathbf{a}_t) \right]$$

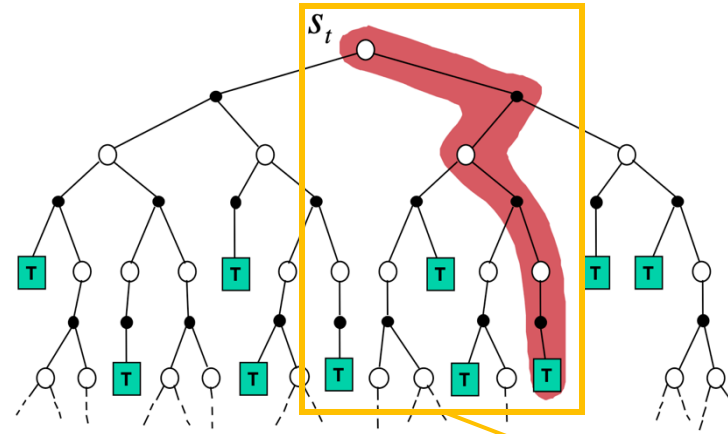
Why do We Use MC / TD, rather than DP?

We Don't Know the Dynamic Function!

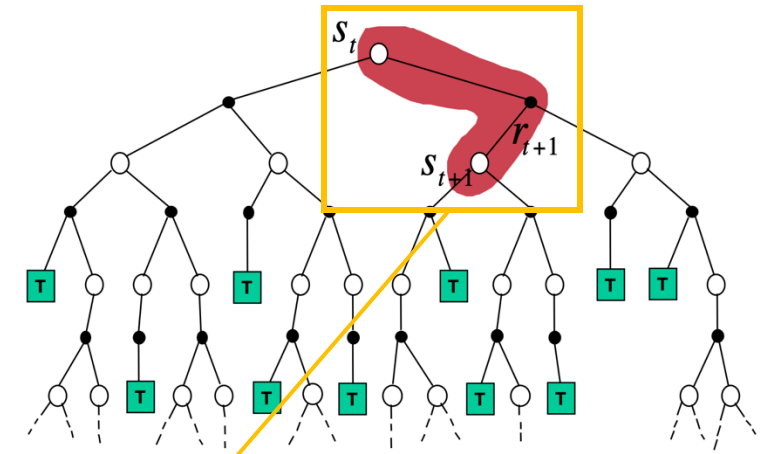
Dynamic Programming Backup



MC Backup



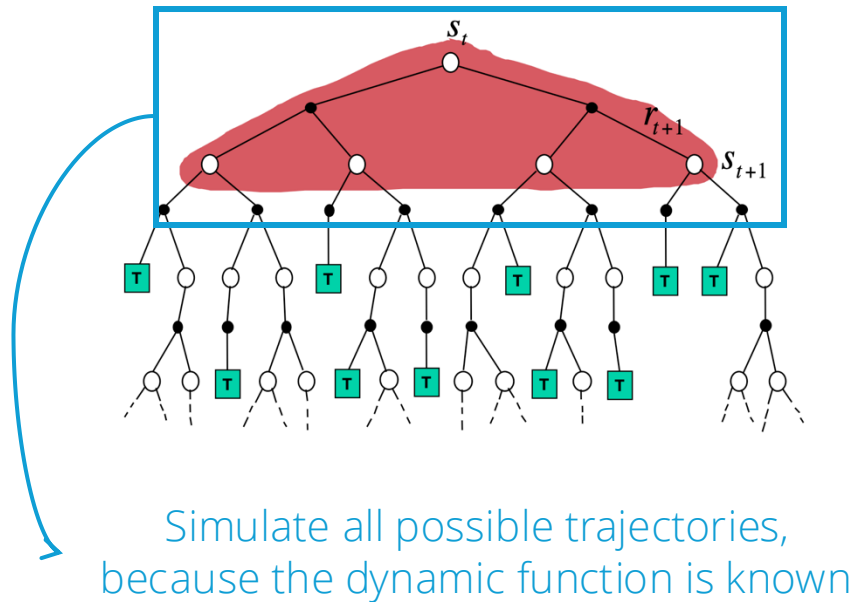
TD Backup



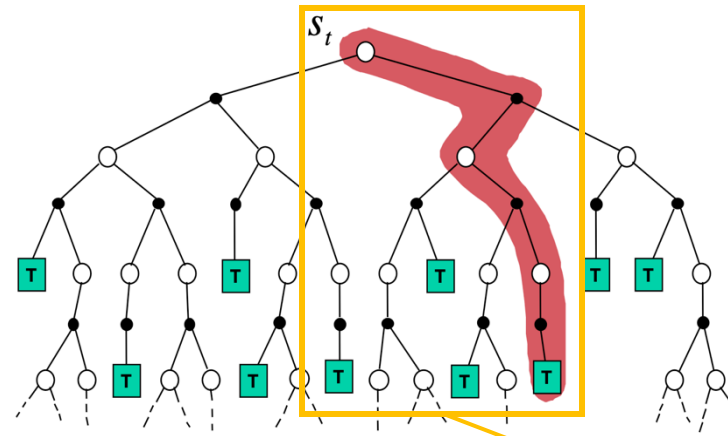
Why do We Use MC / TD, rather than DP?

We Don't Know the Dynamic Function!

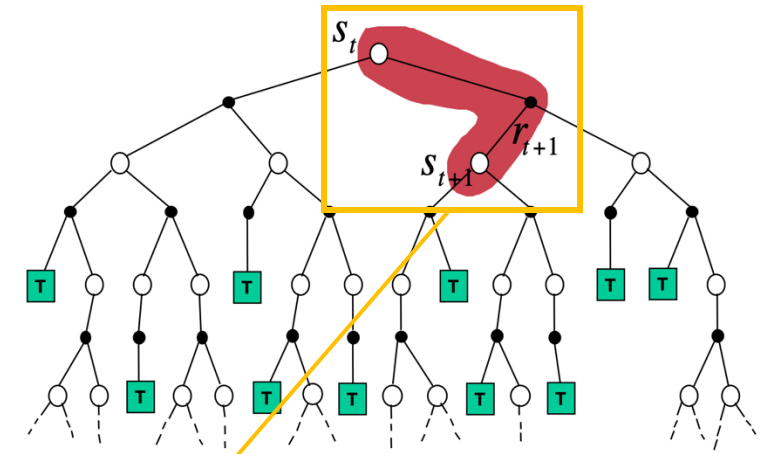
Dynamic Programming Backup



MC Backup



TD Backup



- Policy Gradient: $\nabla J(\theta) = E_{\tau \sim p_{\theta}(\tau)} \left[\left(\sum_{t=0}^T \nabla \log \pi_{\theta}(a_t | s_t) \right) \left(\sum_{t=0}^T r(s_t, a_t) - b \right) \right]$
- Actor Critic: $\nabla J(\theta) = E_{\tau \sim p_{\theta}(\tau)} \left[\left(\sum_{t=0}^T \nabla \log \pi_{\theta}(a_t | s_t) \right) \left(Q^{\pi}(s_t, a_t) - b \right) \right]$

What is a “Dynamics/World Model”?

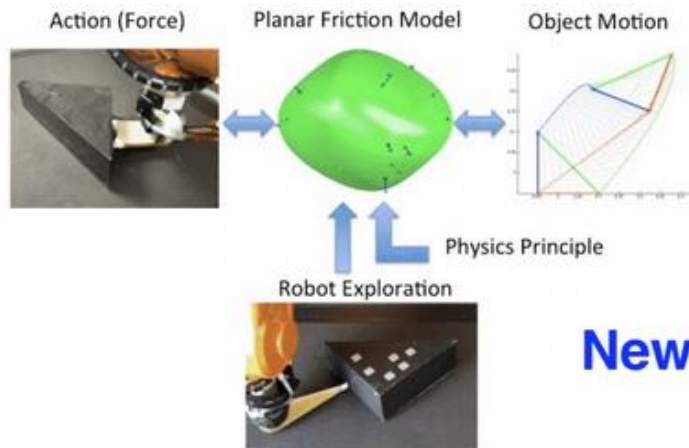


Occupancy map when navigating

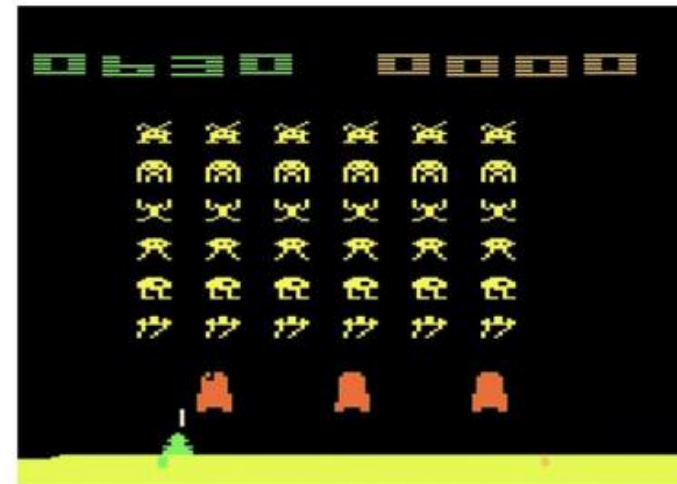


Game Rules when playing

Helps us figure out what to do next



Newton's Physics when Pushing



What is a “Dynamics / World Model”?

The real world: expensive to interact with

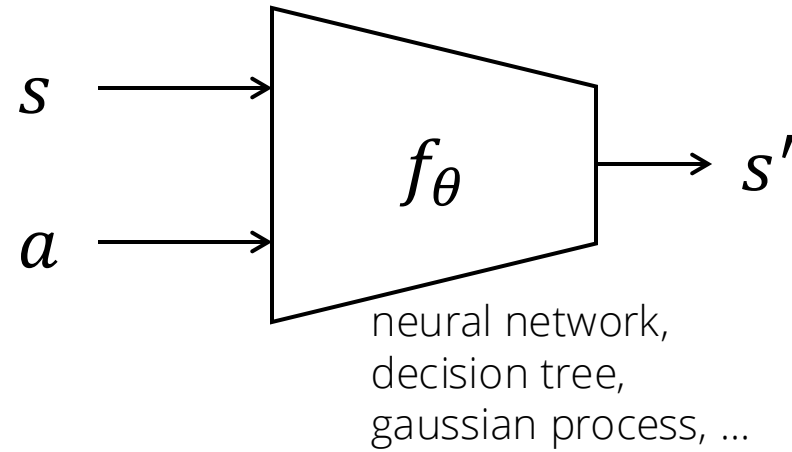


The simulator world: expensive to build and almost impossible to be identical to the real world

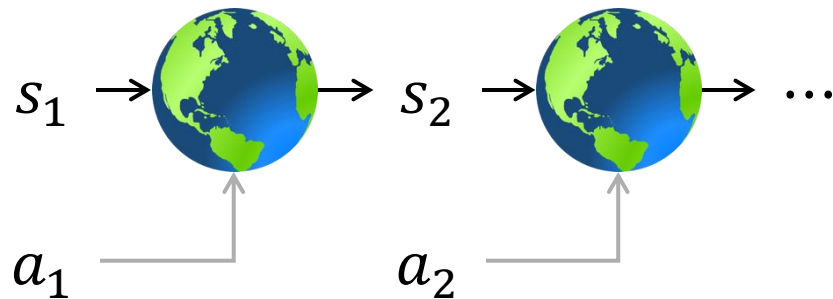


Isaac Sim by Nvidia

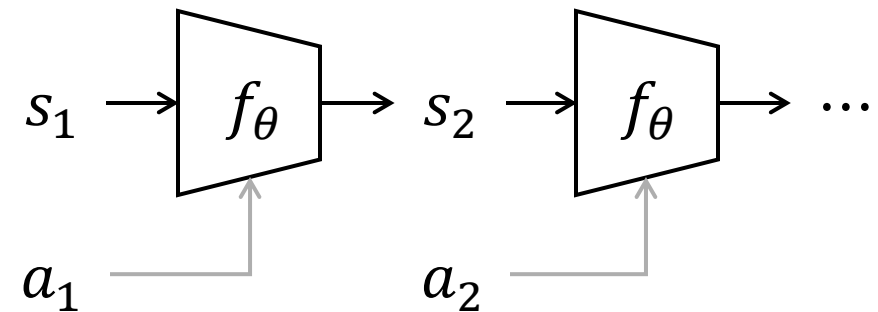
How to Model a “Dynamics / World Model”?



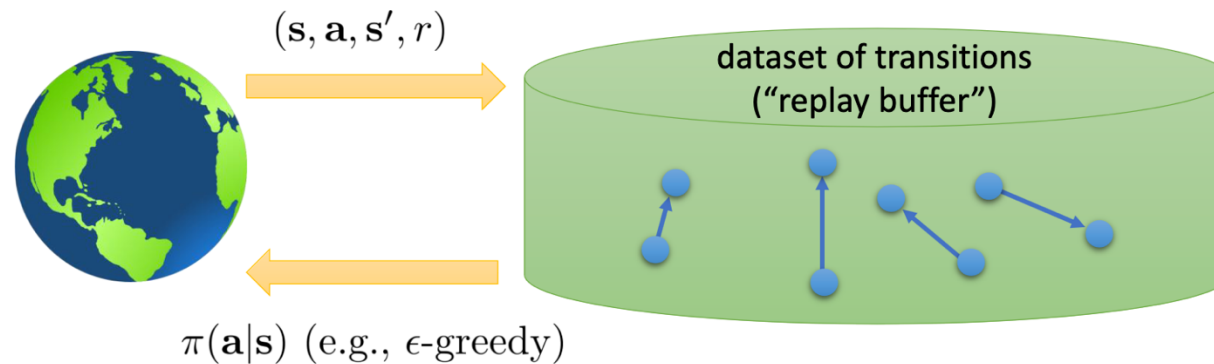
Execute to unroll a sequence of actions in the real world



Imagine to unroll a sequence of actions with a learned world model

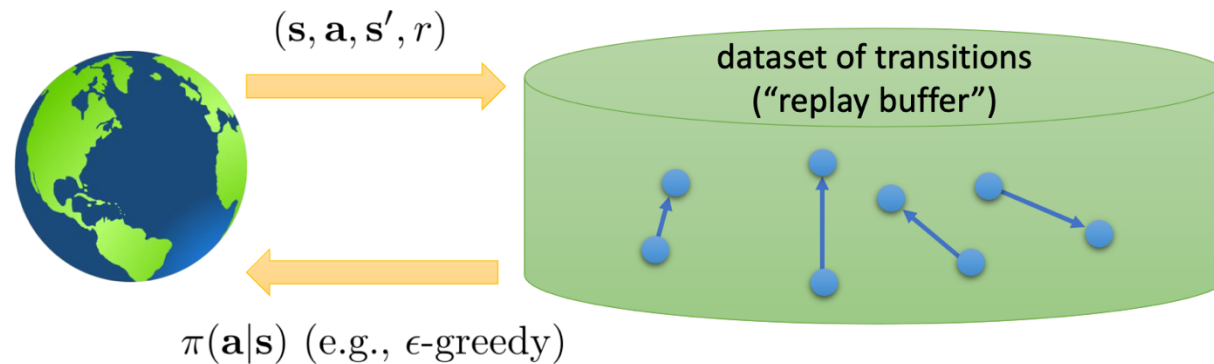


How to Learn a “Dynamics/World Model”?



- For value estimation, we learn value functions from the dataset of transitions:
- We can “re-use” the same dataset of transition for learning a world model.

How to Learn a “Dynamics/World Model”?



Algorithm:

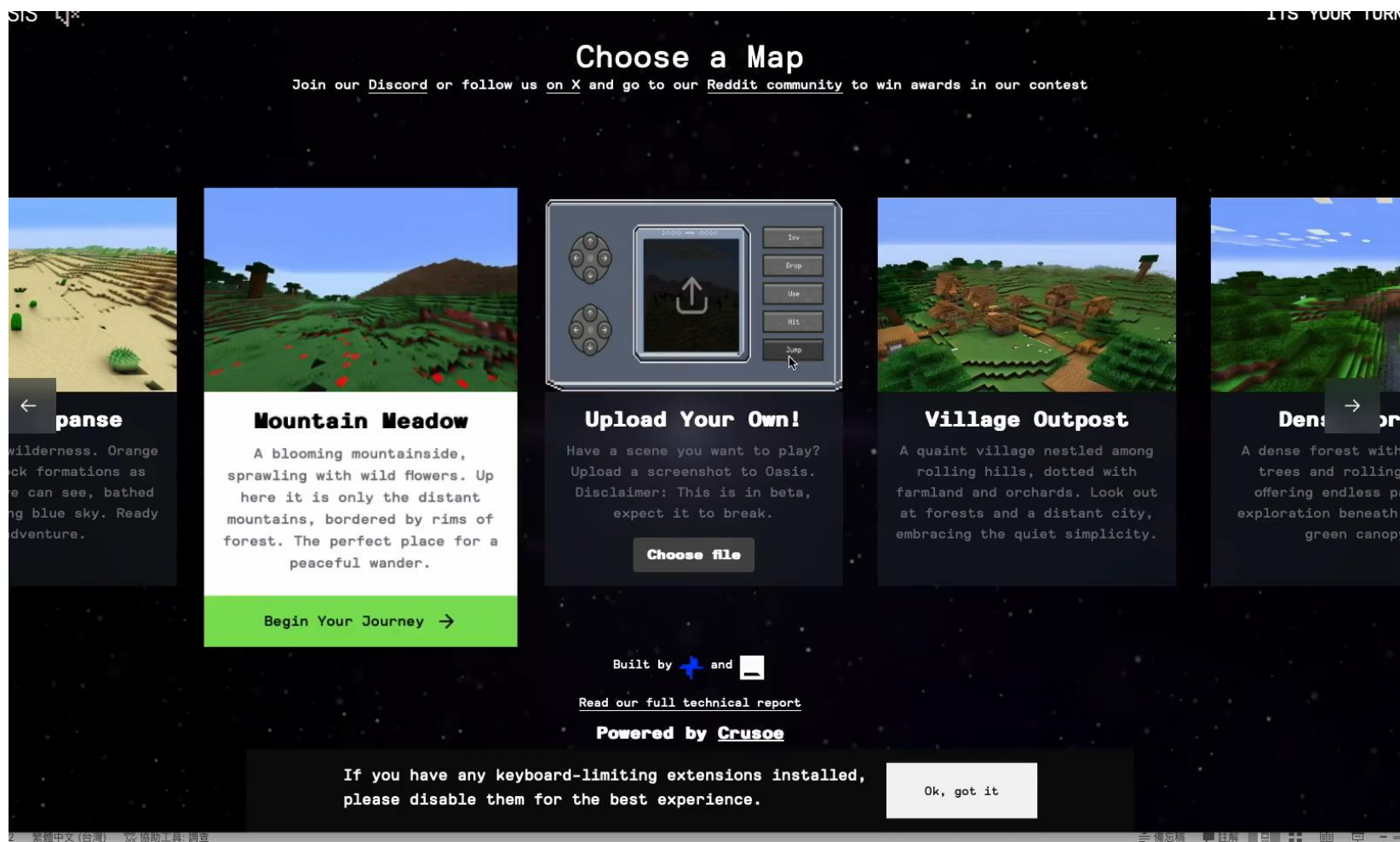
1. Run base policy $\pi_0(a_t|s_t)$ (e.g. random policy) to collect $\mathcal{D} = \{(s_t, a_t, s_{t+1})\}$
2. (Training) Learn dynamics/world model:

$$\min_{\theta} \mathbb{E}_{s_t, a_t, s_{t+1}} [\|s_{t+1} - f_{\theta}(s_t, a_t)\|]$$

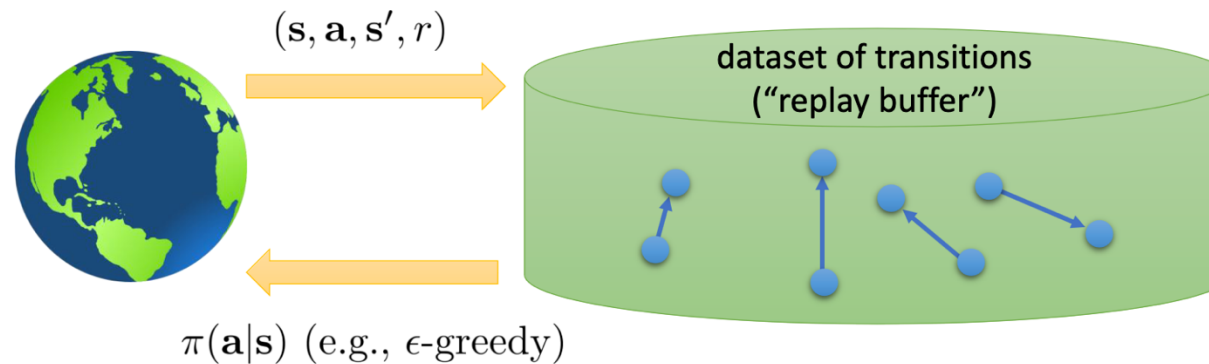
3. (Testing) Plan an optimal sequence of actions $a_{0:T}^*$:

$$a_{0:T}^* = \max_{a_{0:T}} r(s_0, a_0) + r(f(s_0, a_0), a_1) + r(f(f(s_0, a_0), a_1), a_2) + \dots$$

But Learned Dynamics Models Still Have Out-of-Distribution Problems...



Improved the “Dynamics/World Model” with More Experience



Algorithm:

1. Run base policy $\pi_0(a_t|s_t)$ (e.g. random policy) to collect $\mathcal{D} = \{(s_t, a_t, s_{t+1})\}$

2. (Training) Learn dynamics/world model:

$$\min_{\theta} \mathbb{E}_{s_t, a_t, s_{t+1}} [\|s_{t+1} - f_{\theta}(s_t, a_t)\|]$$

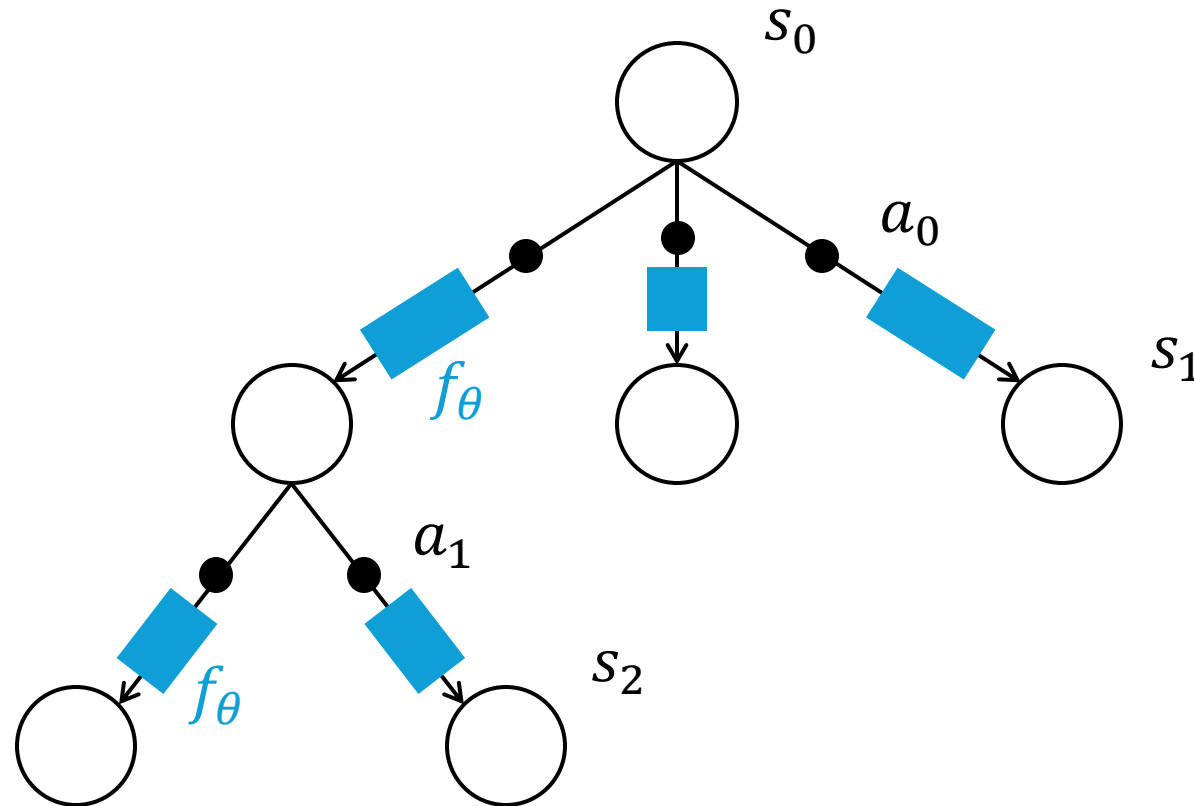
3. (Testing) Plan an optimal sequence of actions $a_{0:T}^*$:

$$a_{0:T}^* = \max_{a_{0:T}} r(s_0, a_0) + r(f(s_0, a_0), a_1) + r(f(f(s_0, a_0), a_1), a_2) + \dots$$

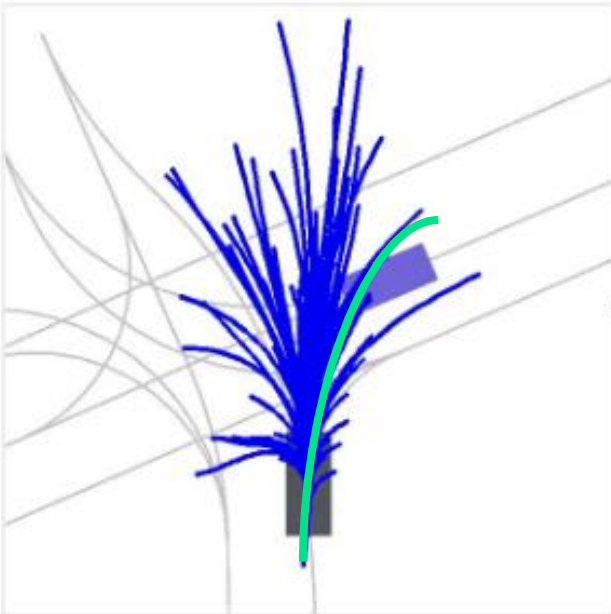
4. Execute the planned actions $a_{0:T}^*$ and add the resulting data $\{(s_t, a_t, s_{t+1})\}$ to $\mathcal{D}$

What Can We Do with a Learned Dynamics Model?

Plan/Search Optimal Action Sequences by Simulation

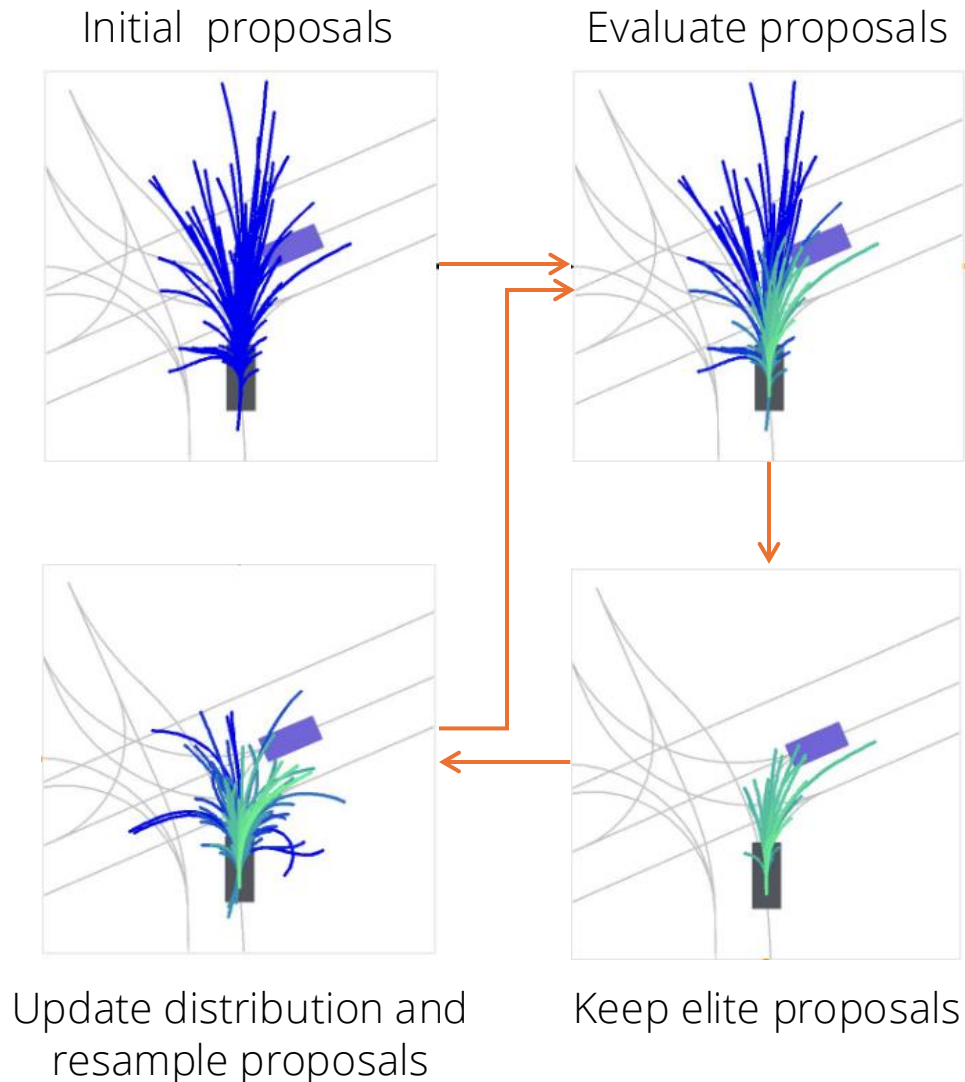


Random Search



1. Sample N sequences of actions $a_{0:T}^1, a_{0:T}^2, \dots, a_{0:T}^N$
 2. Unroll each action sequence with the dynamics model f_θ
 3. Evaluate each unrolled trajectory and pick the best one
- Need a lot of samples to obtain an optimal plan

Can We Search Smartly? Evolutionary Search

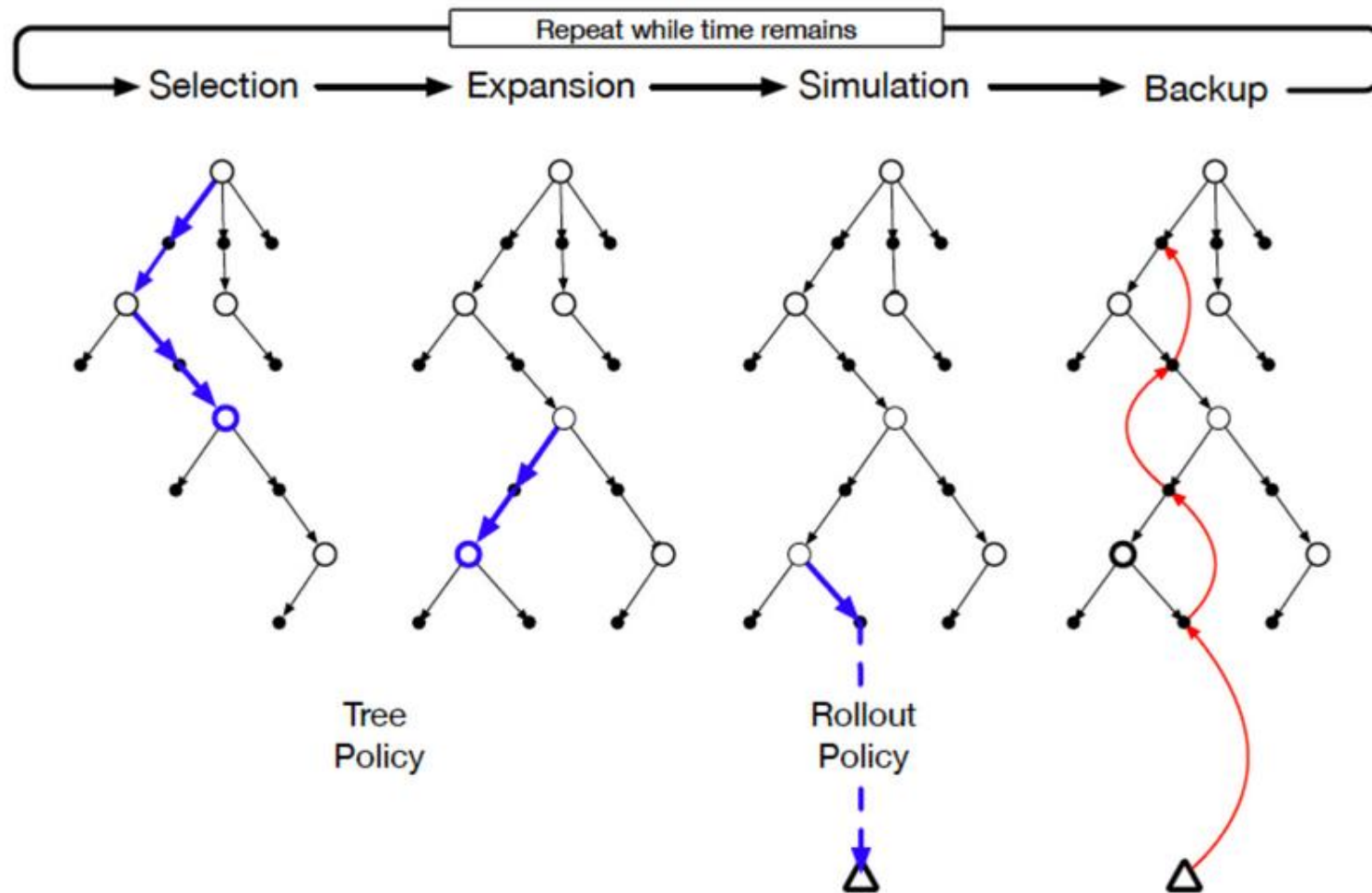


1. Sample N sequences of actions $a_{0:T}^1, a_{0:T}^2, \dots, a_{0:T}^N$ from an initial distribution v_0
2. Unroll each action sequence with the dynamics model f_θ
3. Evaluate each unrolled trajectory and keep the top K sequences. Update the distribution to v_1 with the elite sequences.
4. Resample N sequences of actions $a_{0:T}^1, a_{0:T}^2, \dots, a_{0:T}^N$ from the updated distribution v_1
5. Pick and execute the best action sequence

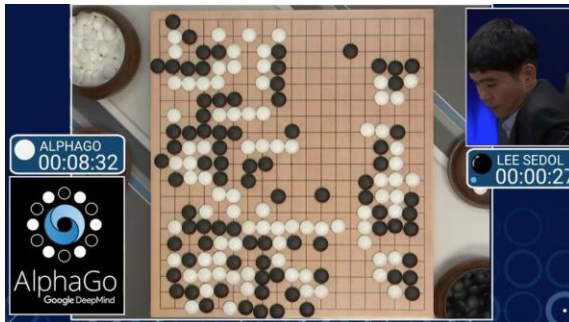
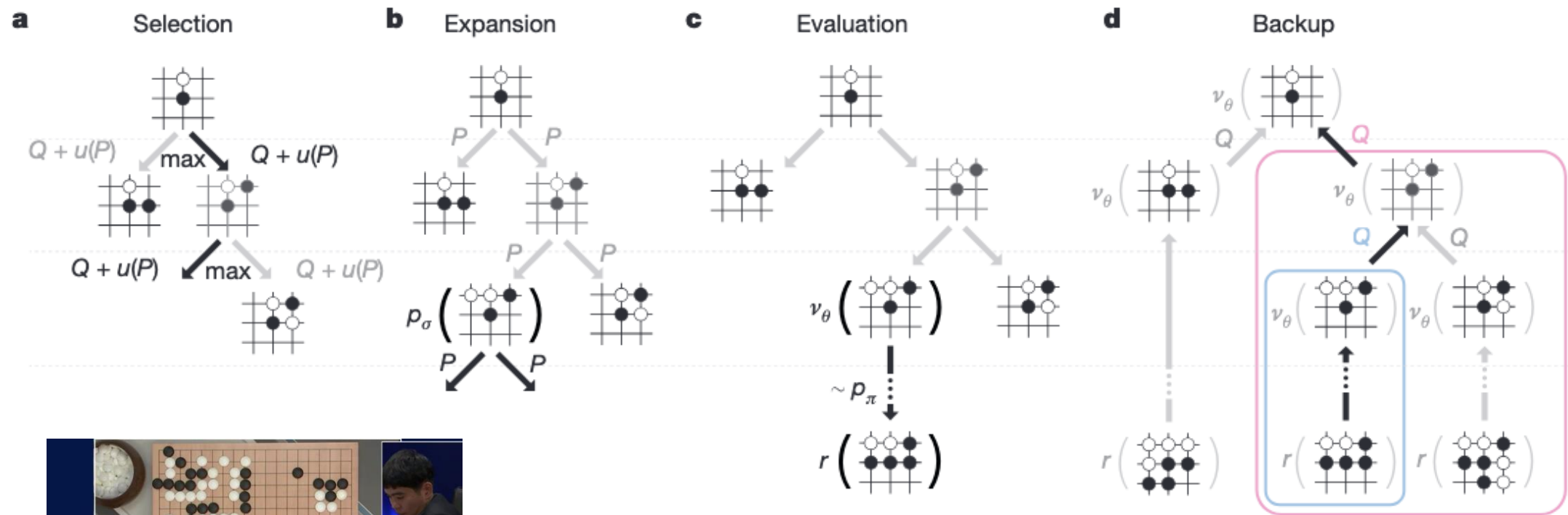
An Example of Cross Entropy Method with Gaussian Distribution

```
// Initialize parameters
μ := -6
σ2 := 100
t := 0
maxits := 100
N := 100
Ne := 10
// While maxits not exceeded and not converged
while t < maxits and σ2 > ε do
    // Obtain N samples from current sampling distribution
    X := SampleGaussian(μ, σ2, N)
    // Evaluate objective function at sampled points
    S := exp(-(X - 2) ^ 2) + 0.8 exp(-(X + 2) ^ 2)
    // Sort X by objective function values in descending order
    X := sort(X, S)
    // Update parameters of sampling distribution via elite samples
    μ := mean(X(1:Ne))
    σ2 := var(X(1:Ne))
    t := t + 1
// Return mean of final sampling distribution as solution
return μ
```

Can We Search Smartly? Monte-Carlo Tree Search



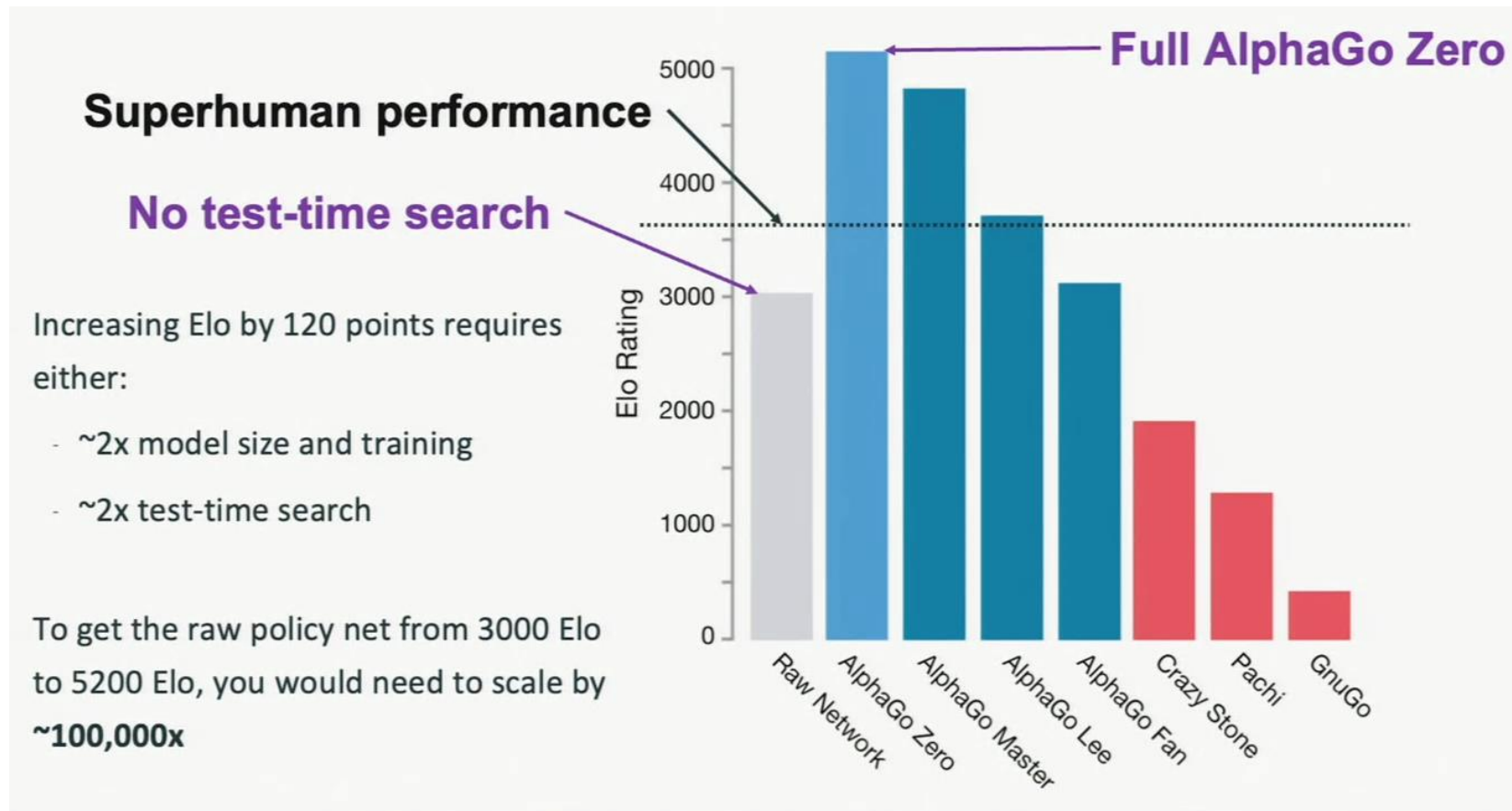
Can We Search Smartly? Monte-Carlo Tree Search



We'll go over Alpha Go today!

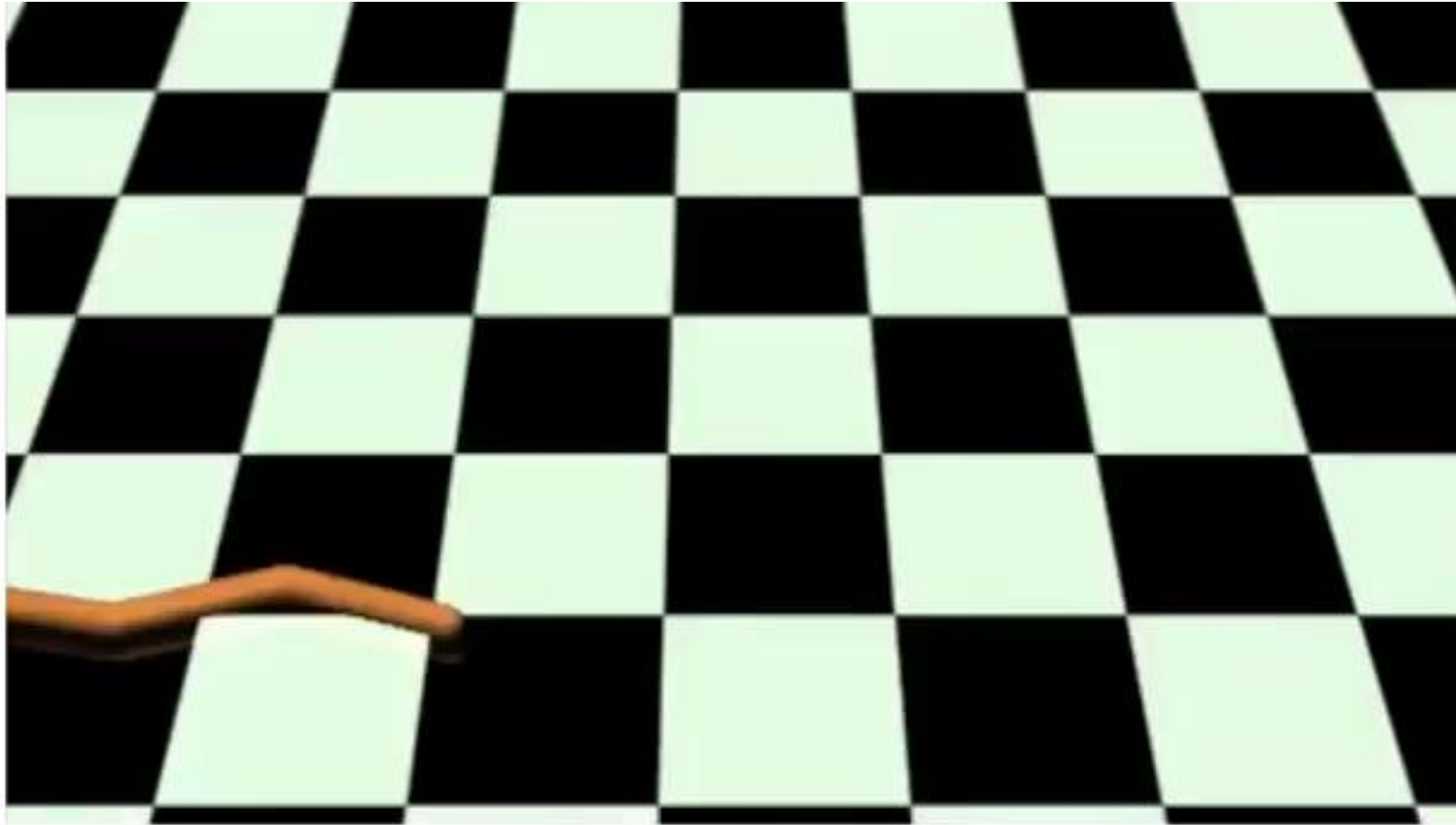
Why Do We Need Planning/Searching??

- Planning enables a sub-optimal policy to outperform a stronger policy w/o planning
- Planning facilitates learning of value functions



Why Do We Need Planning/Searching?

- Planning to achieve multiple (unseen) goals without re-training the policy



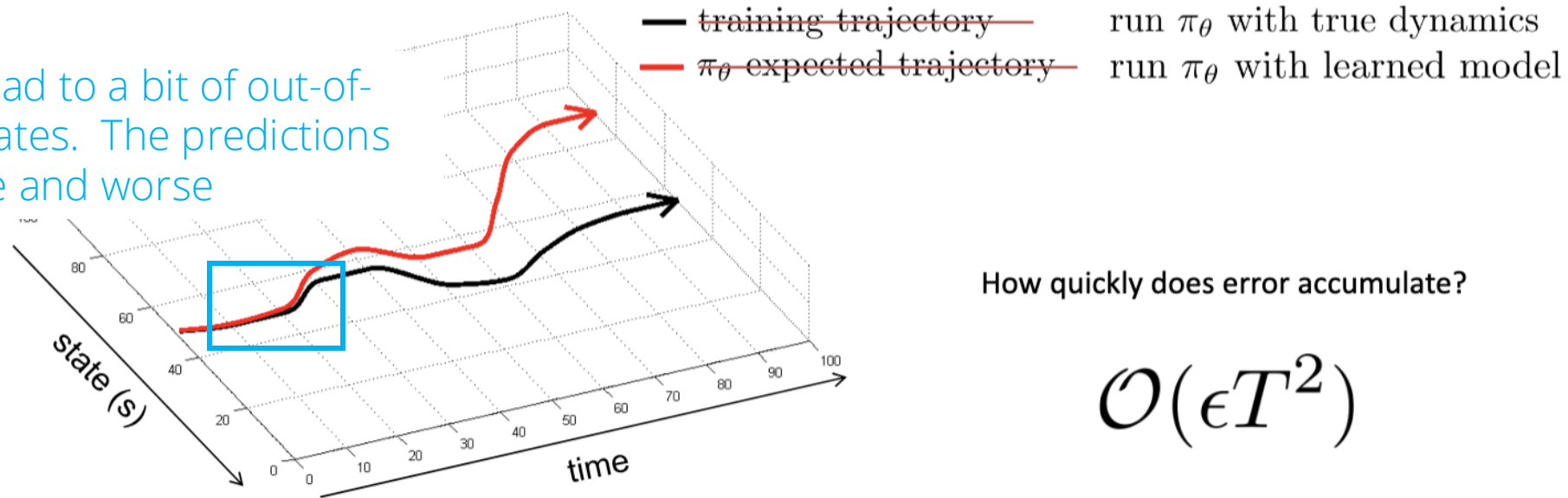
© Authors of ICRA 2018 Paper 1644

Thu PM

Pod Q.8

But Learned Dynamics Models Still Have Error Accumulation Problem (Compounding Errors)...

Small errors lead to a bit of out-of-distribution states. The predictions become worse and worse



How quickly does error accumulate?

$$\mathcal{O}(\epsilon T^2)$$

Solution 1: Augment the Training Data with Rollouts / Noise Perturbation

- Augment training with rollouts

$$\min_{\theta} \mathbb{E}_{s_t, a_t, s_{t+1}} [\|s_{t+1} - f_{\theta}(f_{\theta}(f_{\theta}(\dots), a_{t-2}), a_{t-1}), a_t)\|]$$

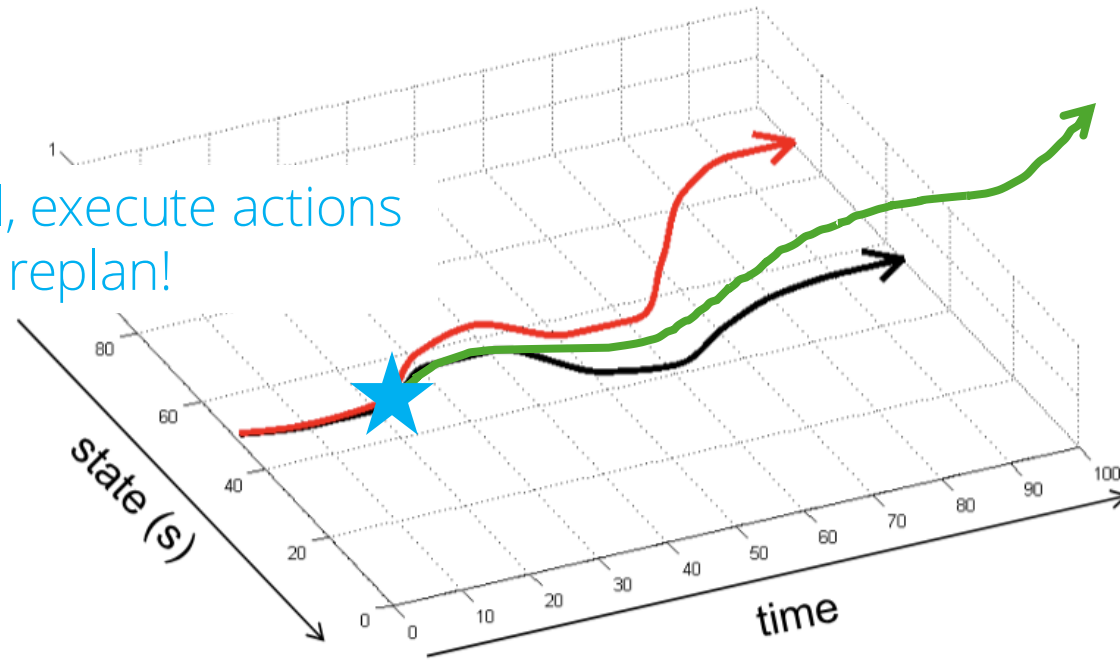
- Augment training with noise perturbation

$$\min_{\theta} \mathbb{E}_{s_t, a_t, s_{t+1}} [\|s_{t+1} - f_{\theta}(s_t + \epsilon, a_t)\|], \epsilon \sim \mathcal{N}(0, 1)$$

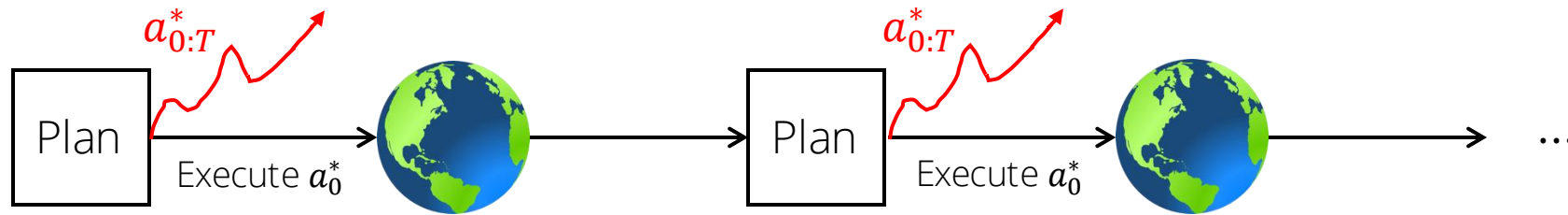
Solution 2: Use the First Few Actions, where Errors are Still Small.

- Idea: Execute the first few actions, where the accumulation errors are still small, and then replan!

Errors are still small, execute actions until here and then replan!



Model Predictive Control



Algorithm:

1. Run base policy $\pi_0(a_t|s_t)$ (e.g. random policy) to collect $\mathcal{D} = \{(s_t, a_t, s_{t+1})\}$

2. (Training) Learn dynamics/world model:

$$\min_{\theta} \mathbb{E}_{s_t, a_t, s_{t+1}} [\|s_{t+1} - f_{\theta}(s_t, a_t)\|]$$

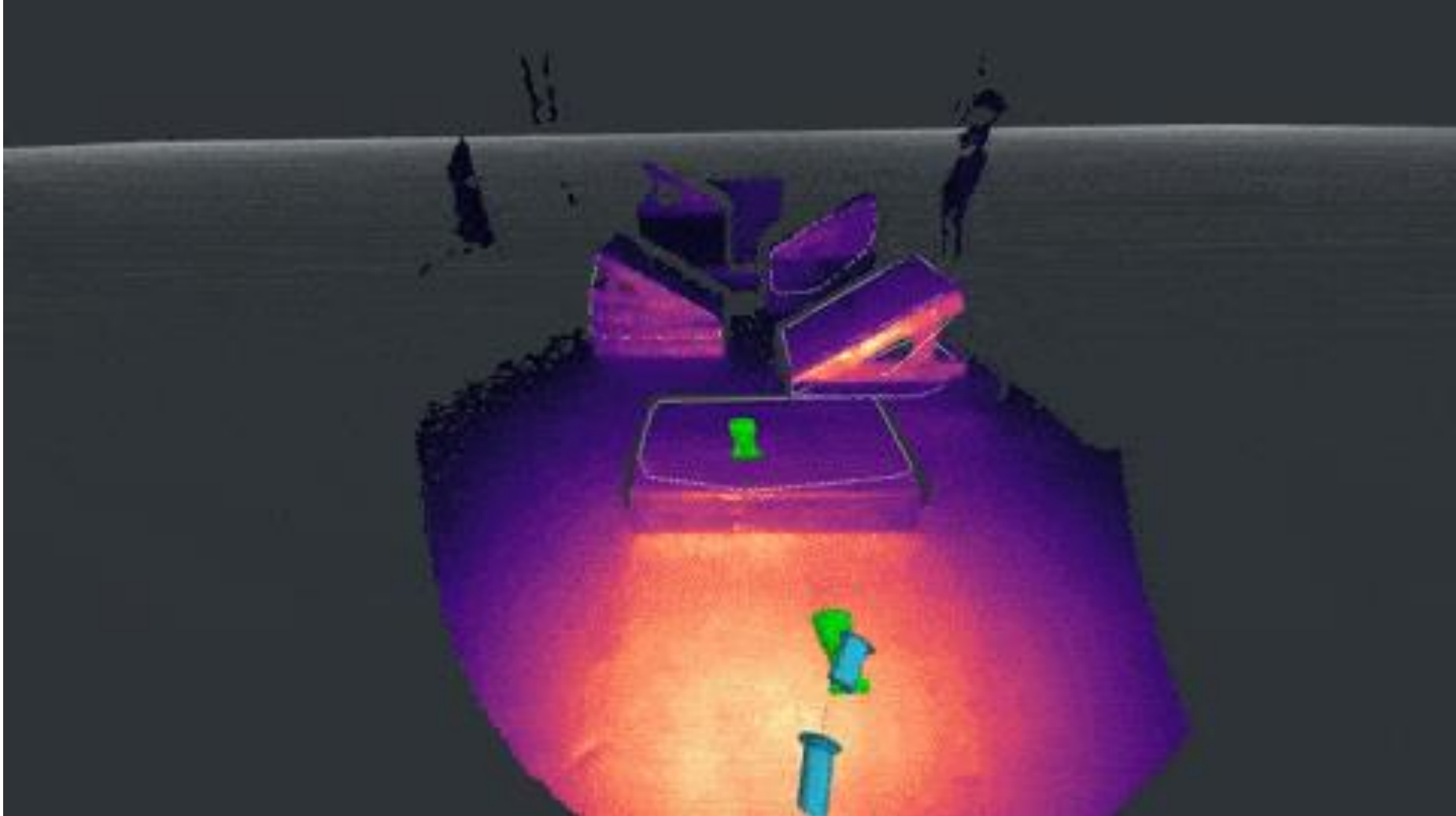
3. (Testing) Plan an optimal sequence of actions $a_{0:T}^*$:

$$a_{0:T}^* = \max_{a_{0:T}} r(s_0, a_0) + r(f(s_0, a_0), a_1) + r(f(f(s_0, a_0), a_1), a_2) + \dots$$

4. Execute the first planned action a_0^* and observe resulting state s_1

5. add the resulting data $\{(s_t, a_t, s_{t+1})\}$ to $\mathcal{D}$

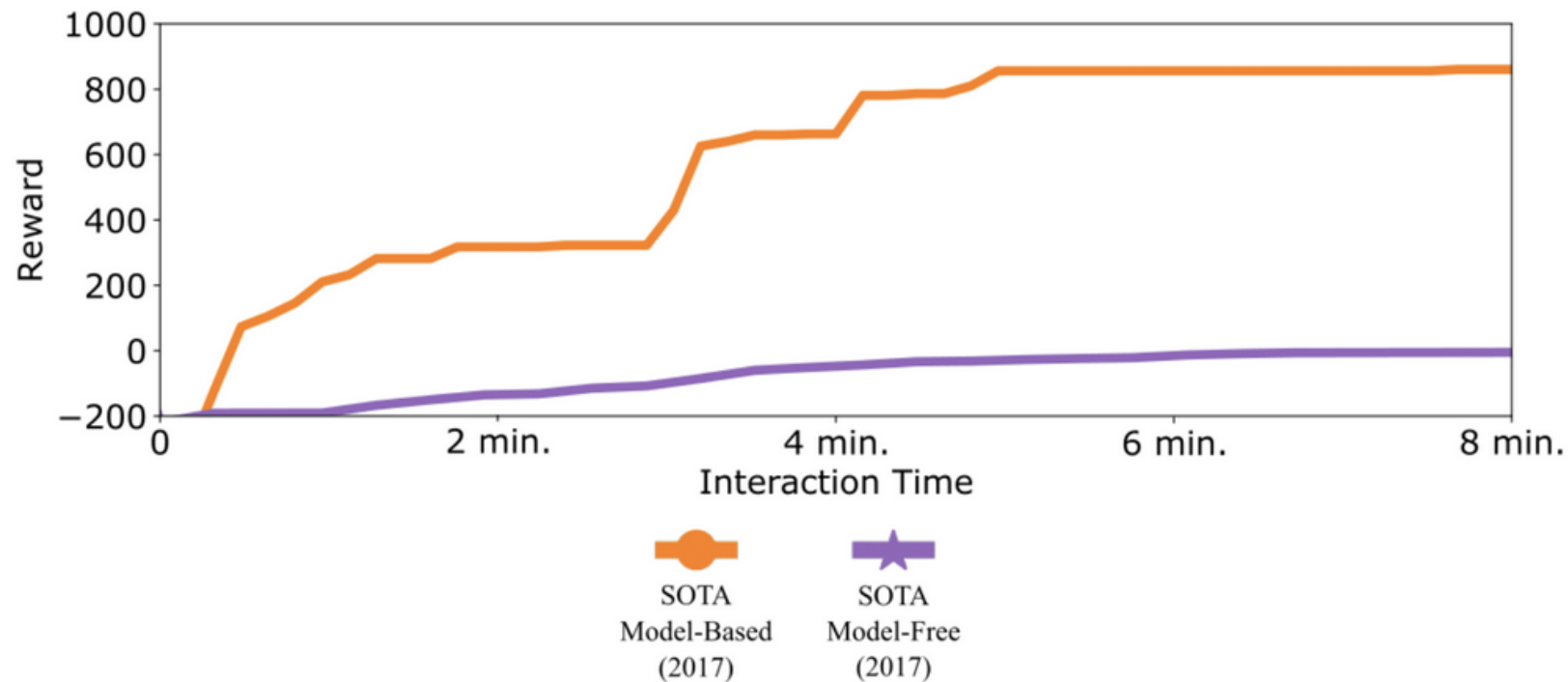
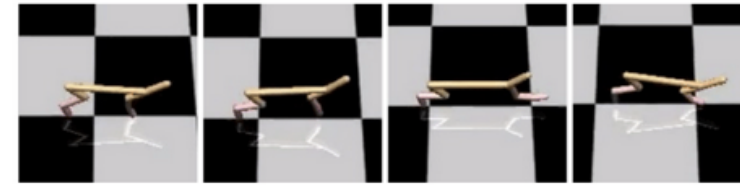
Remember We've Talked About MPC in Optimal Control



What are Model-based RL Good For?

- Model-based RL methods are sample efficient, which require much less interaction than model-free RL methods

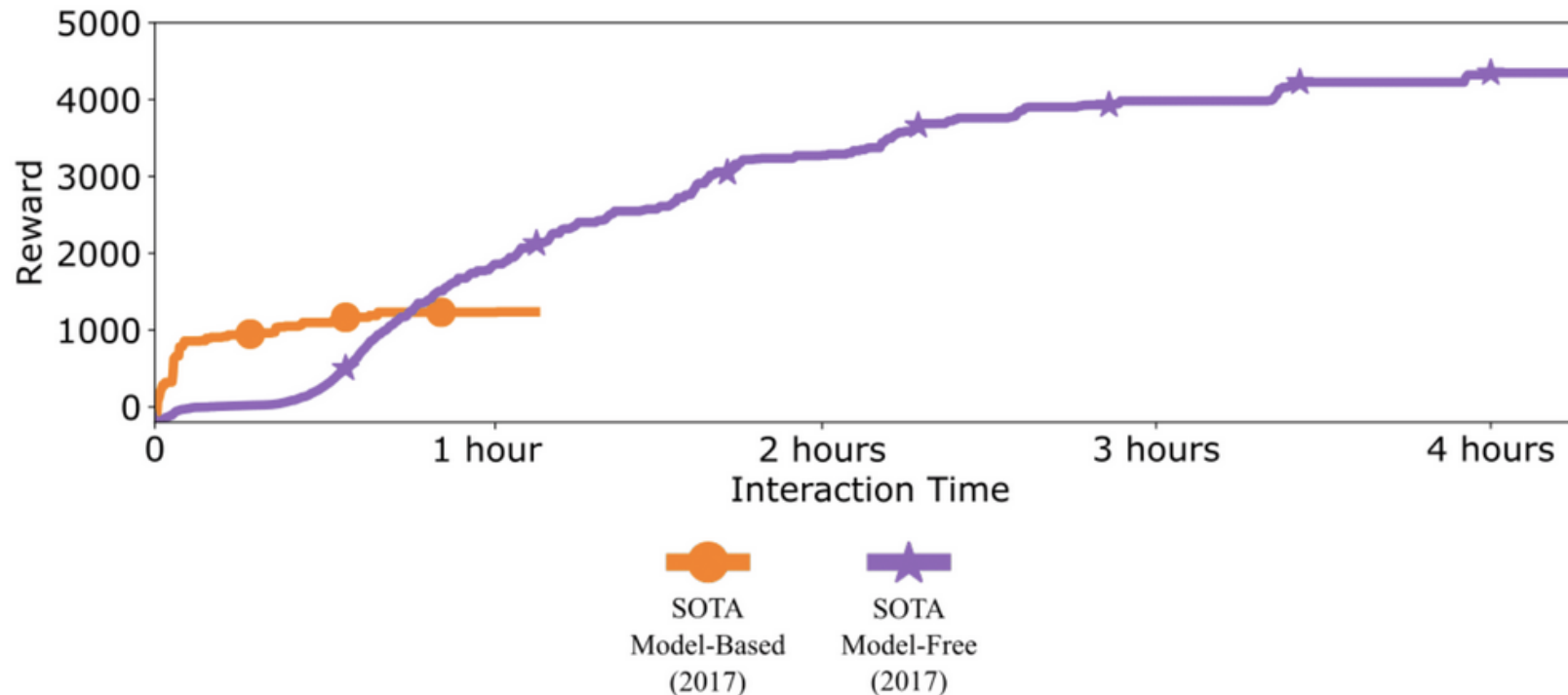
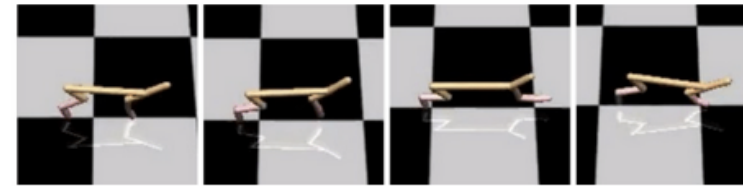
Comparative Performance on HalfCheetah



But, Pure Model-based RL Has Problems...

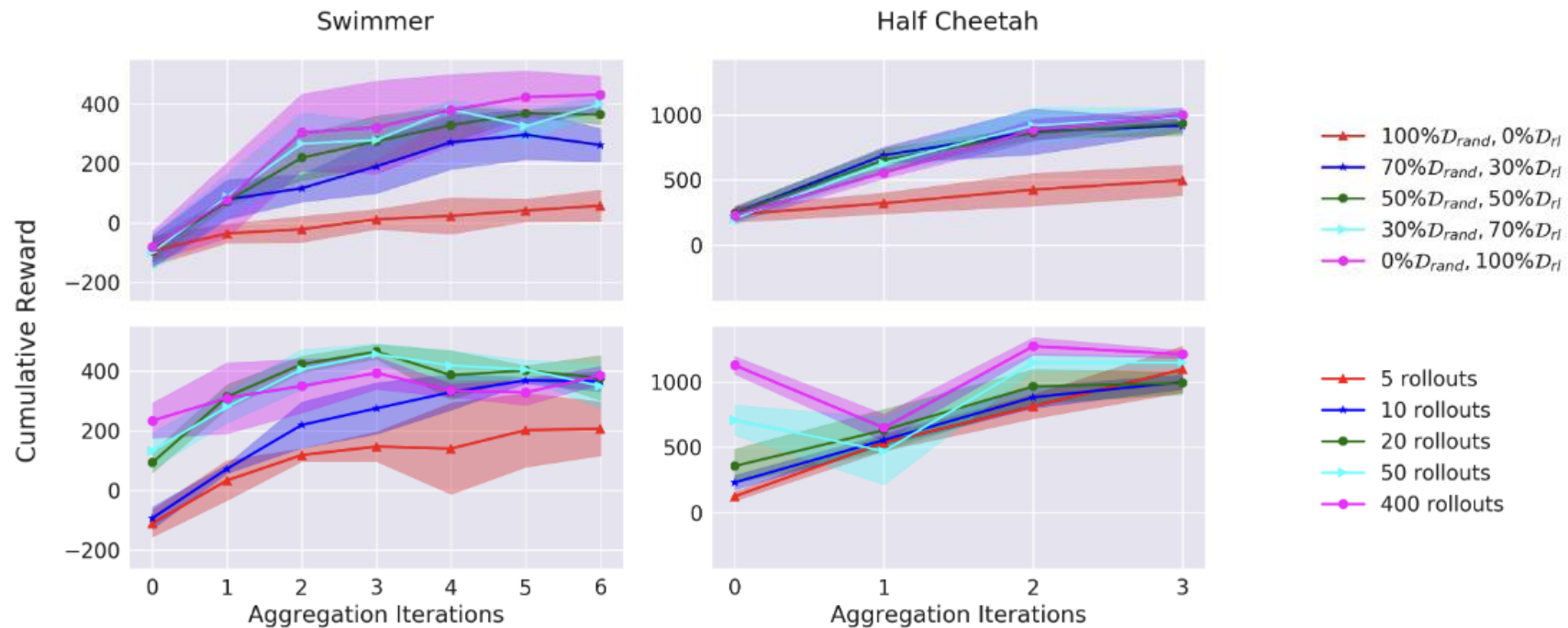
- Pure model-based RL methods are sub-optimal than model-free RL....

Comparative Performance on HalfCheetah



But, Pure Model-based RL Has Problems...

- For MPC to achieve optimality, the learned dynamic models need to know the state transitions of optimal action sequences
- MPC is sub-optimal when the learned dynamic models are trained with random policy rollouts

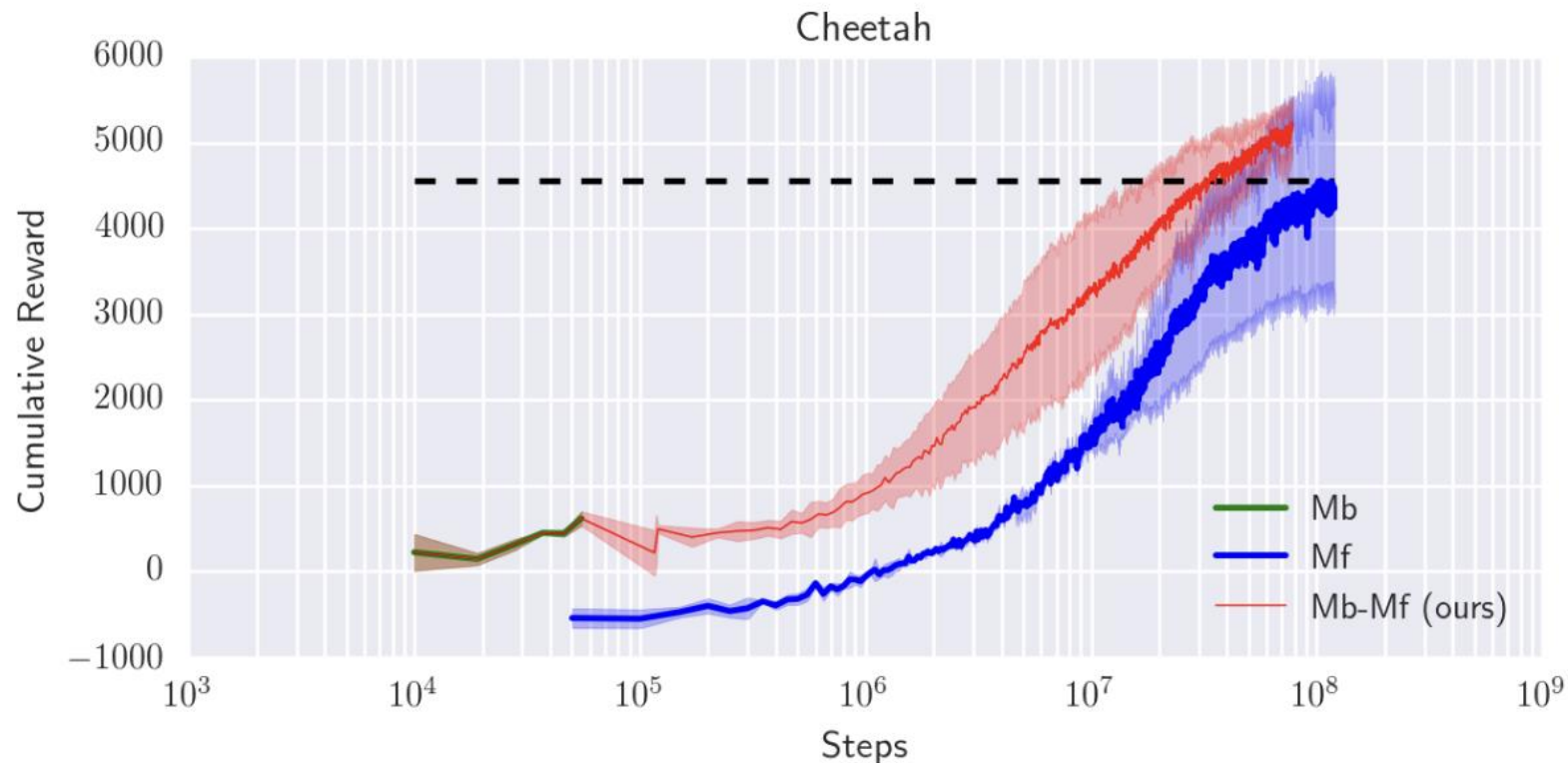


Idea: Combine Model-based with Model-free RL

- Initialize the policy with DAGGER, using the rollouts $\mathcal{D}^*$ of MPC controllers as the training data

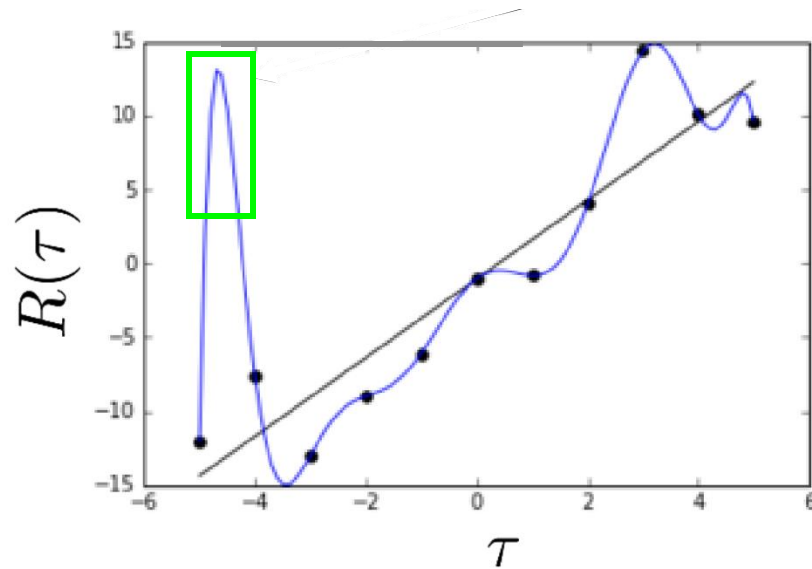
$$\min_{\phi} \frac{1}{2} \sum_{(\mathbf{s}_t, \mathbf{a}_t) \in \mathcal{D}^*} \|\mathbf{a}_t - \mu_{\phi}(\mathbf{s}_t)\|_2^2,$$

- Fine-tune the policy with model-free RL (e.g. TRPO / PPO ...)

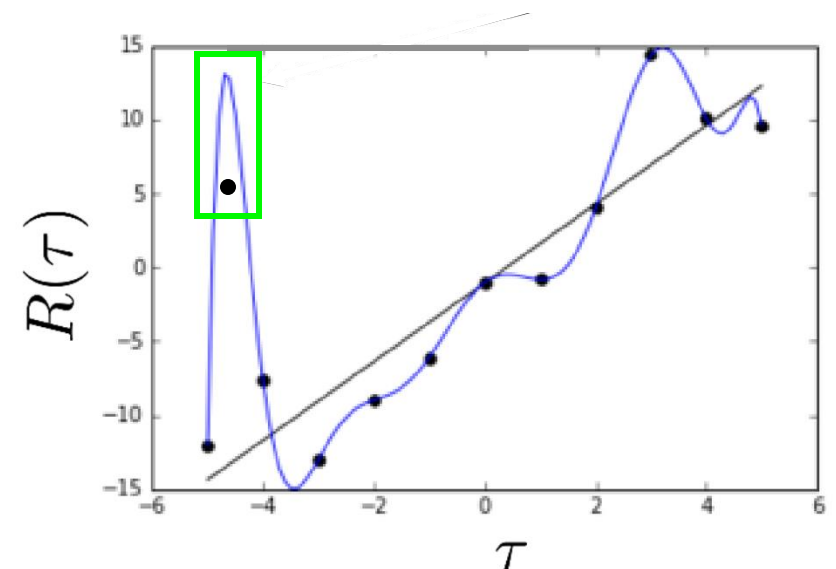


But Why Pure Model-based RL Doesn't Work?

High-variance reward trajectory
induced by overfitting



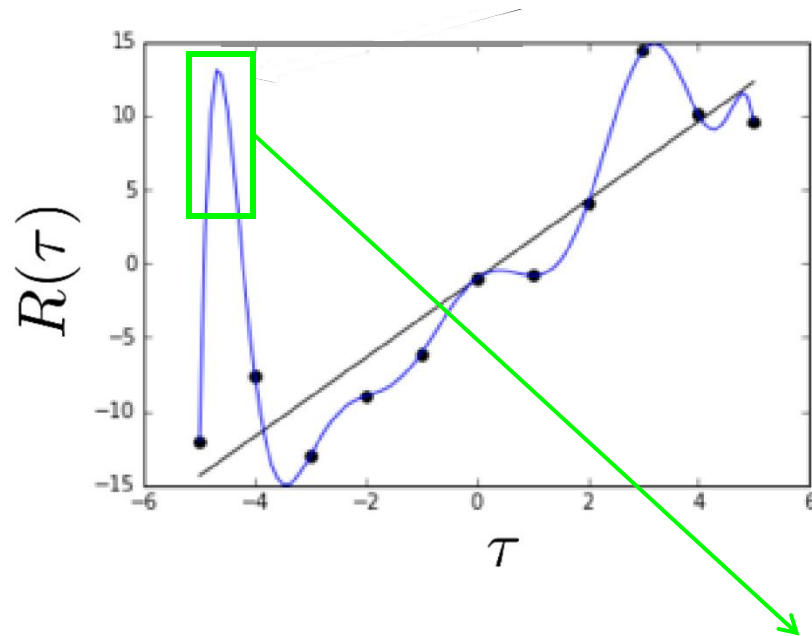
High-variance reward trajectory
induced by underfitting



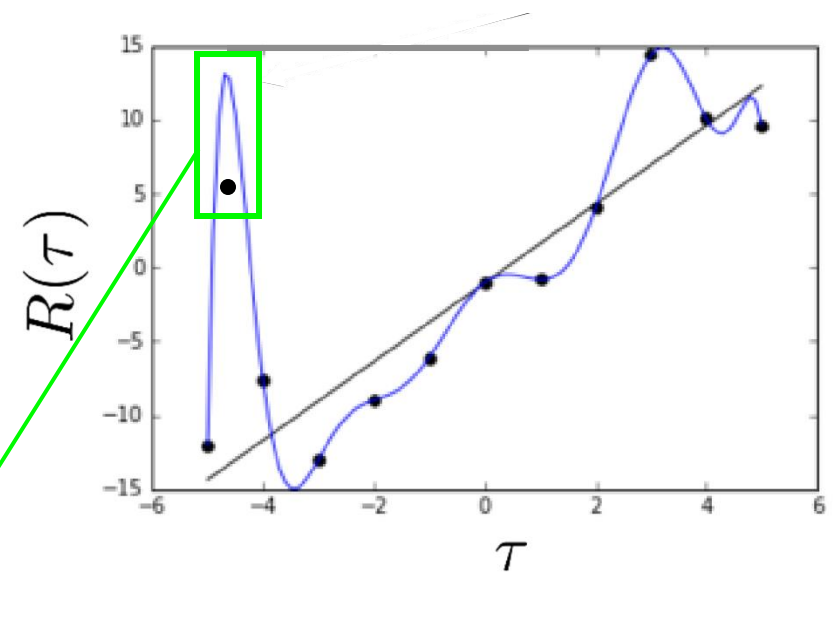
- Incorrect dynamics models lead to inaccurate reward estimates for planned rollouts

But Why Pure Model-based RL Doesn't Work?

High-variance reward trajectory
induced by overfitting



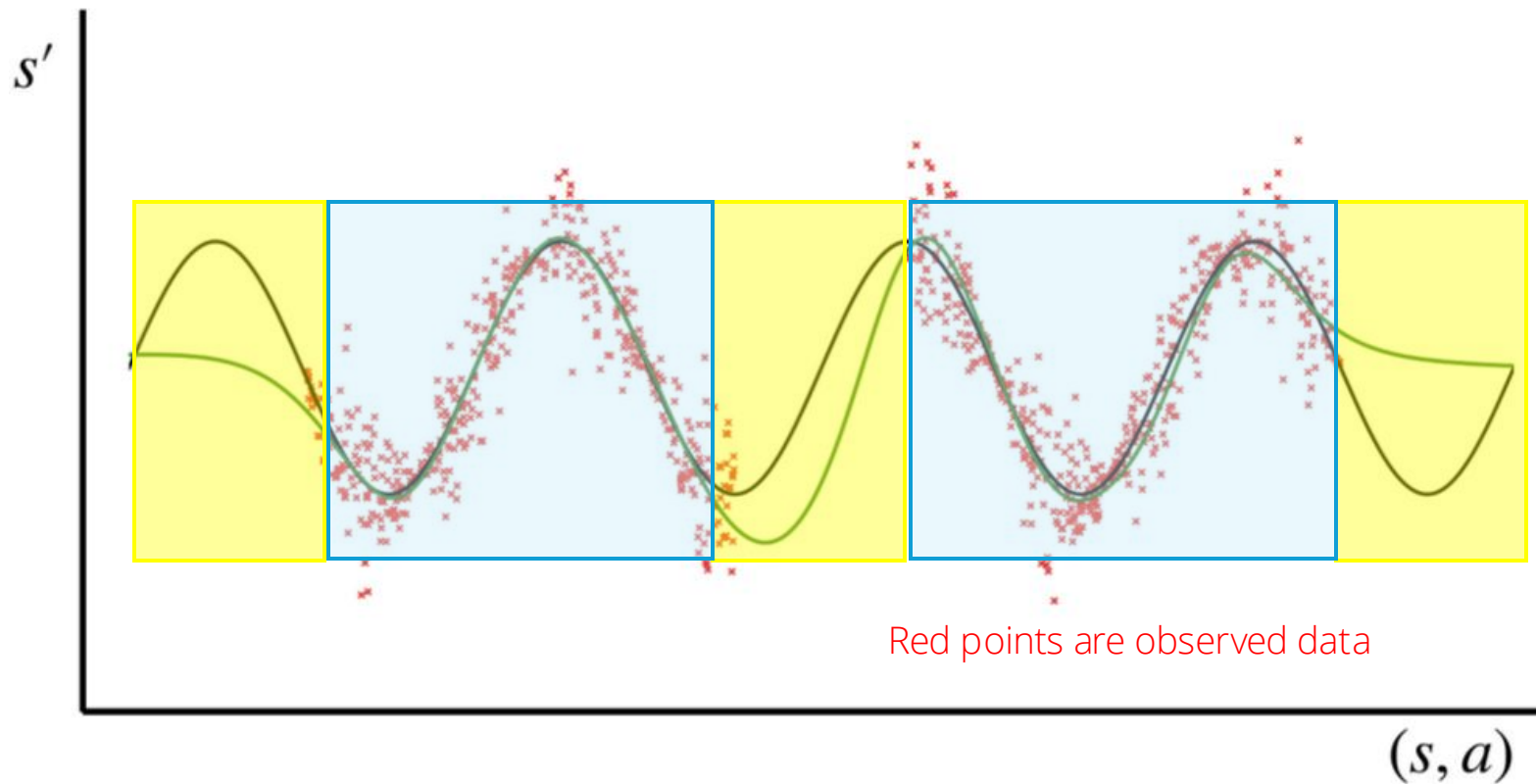
High-variance reward trajectory
induced by underfitting



Model-based RL attempts to go to high-variance reward trajectory,
which could have low expected reward instead.

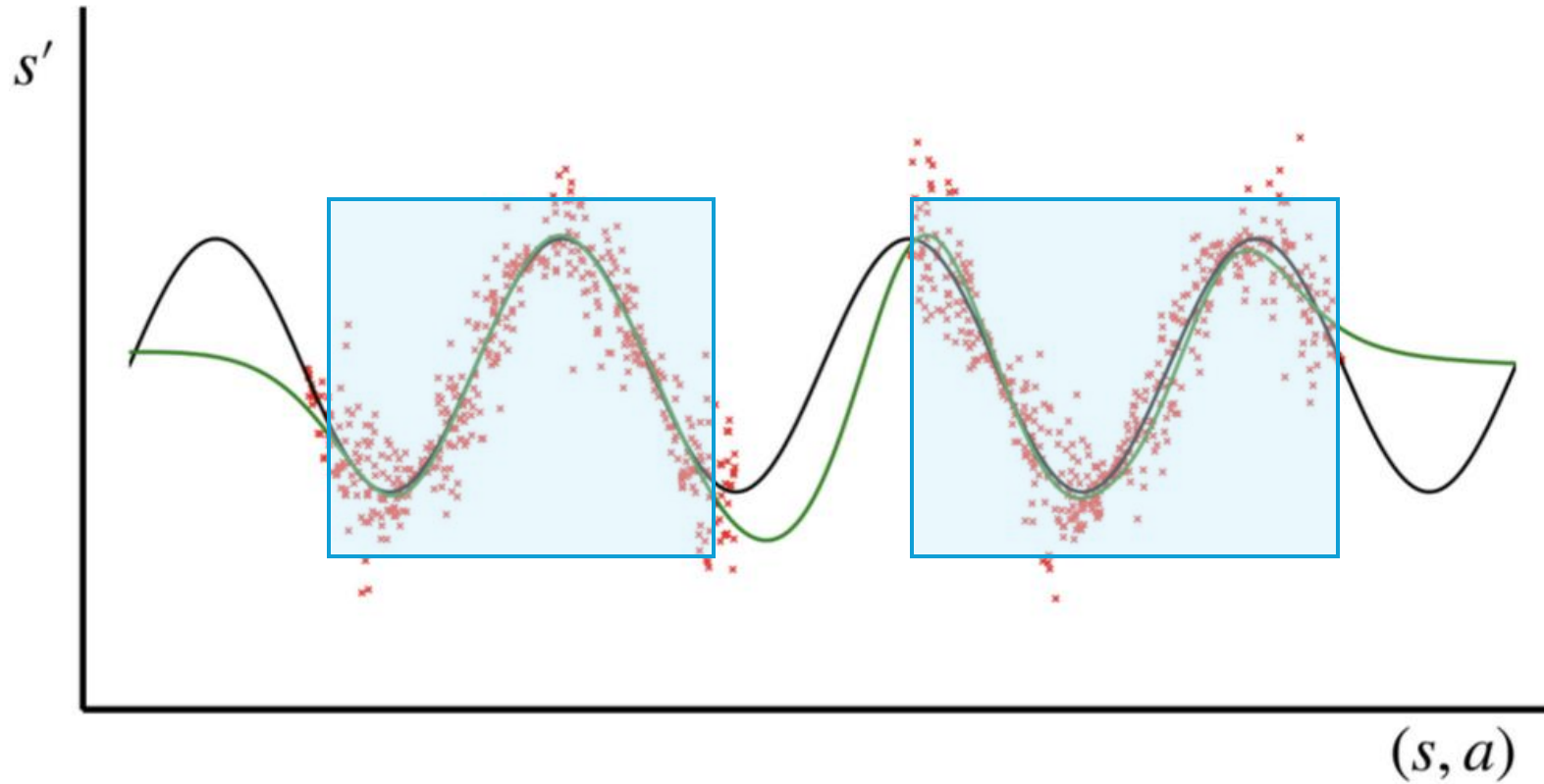
Idea: Account for Model Uncertainty

- Two types of uncertainty: (1) Epistemic uncertainty, (2) Aleatoric uncertainty



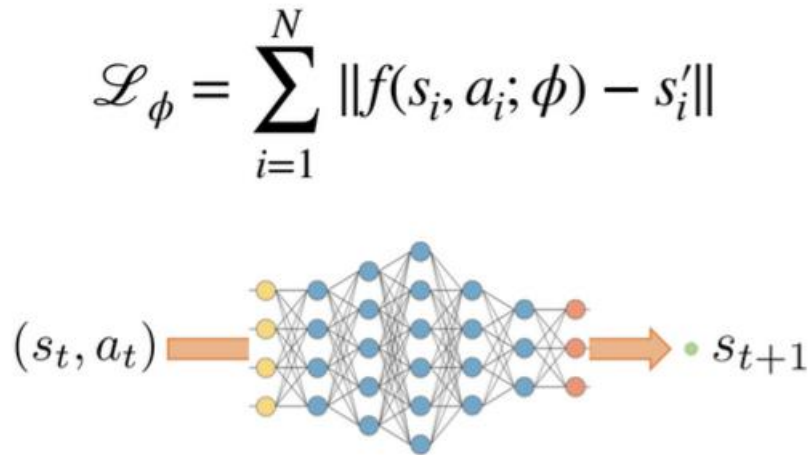
Aleatoric Uncertainty

- Aleatoric uncertainty: arises from inherent stochasticity of a system, e.g. observation noise, process noise or multi-modality



Aleatoric Uncertainty

- Aleatoric uncertainty: arises from inherent stochasticity of a system, e.g. observation noise, process noise or multi-modality
- How to model aleatoric uncertainty? If the environment is stochastic:
 - Deterministic regression models fails to capture the output distribution



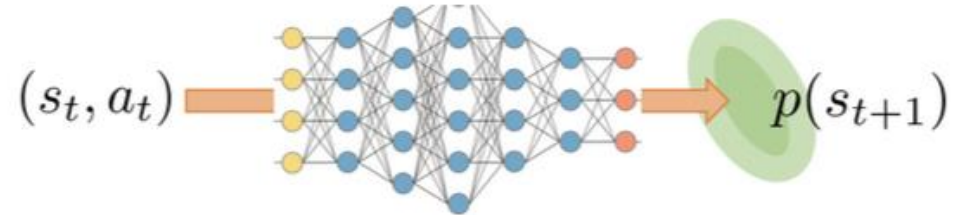
Aleatoric Uncertainty

- Aleatoric uncertainty: arises from inherent stochasticities of a system, e.g. observation noise, process noise or multi-modality
- How to model aleatoric uncertainty? If the environment is stochastic:
 - Deterministic regression models fails to capture the output distribution
 - Probabilistic models can capture the output distribution
- An example that captures Gaussian distribution:

$$p_{\phi}(s'|s, a) = \frac{\exp\left(-\frac{1}{2}(s' - \mu(s, a; \phi))^{\top}(\Sigma(s, a; \phi))^{-1}(s' - \mu(s, a; \phi))\right)}{\sqrt{(2\pi)^d \det \Sigma(s, a; \phi)}}$$

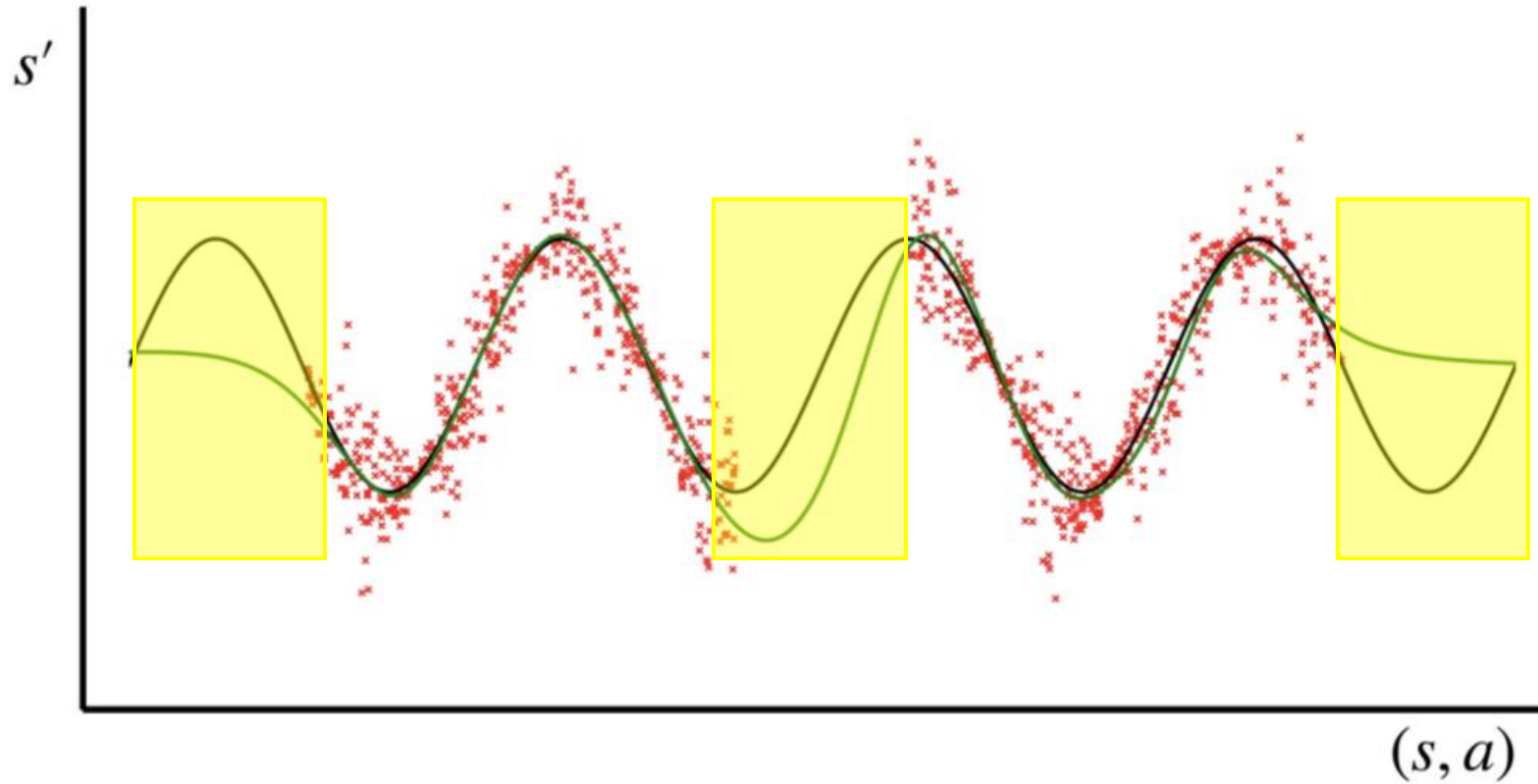
$$\mathcal{L}_{\phi} = -\frac{1}{N} \sum_{i=1}^N \log p_{\phi}(s'_i | s_i, a_i)$$

$$= \frac{1}{2}(s'_i - \mu(s_i, a_i; \phi))^{\top} \Sigma(s_i, a_i; \phi)^{-1} (s'_i - \mu(s_i, a_i; \phi)) + \frac{1}{2} \log(\det \Sigma(s_i, a_i; \phi)) + \text{const.}$$



Epistemic Uncertainty

- Epistemic uncertainty: subjective uncertainty about the dynamics function, due to a lack of sufficient data to uniquely determine the underlying system exactly
- The model is certain about the data, but we are not certain about the model



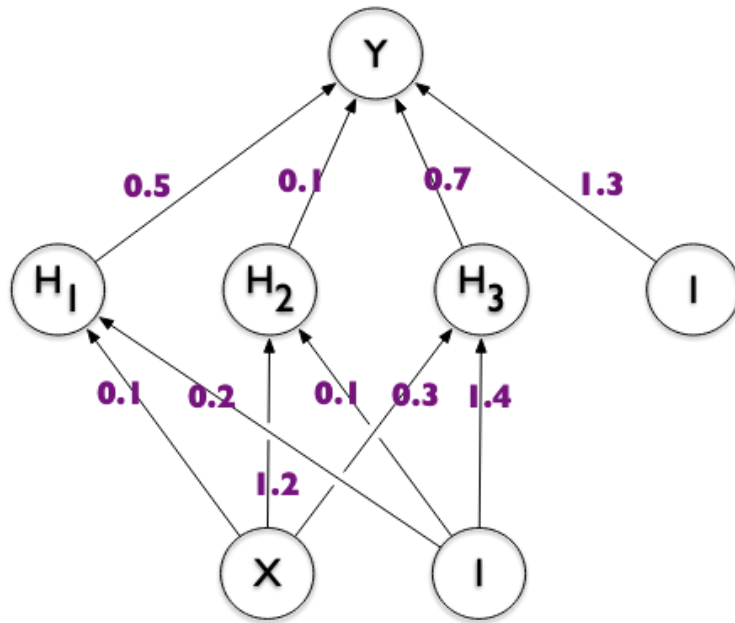
How to Model Epistemic Uncertainty? A Bayesian Inference Method

$$P(\text{hypothesis} \mid \text{data}) = \frac{P(\text{hypothesis})P(\text{data} \mid \text{hypothesis})}{\sum_h P(h)P(\text{data} \mid h)}$$

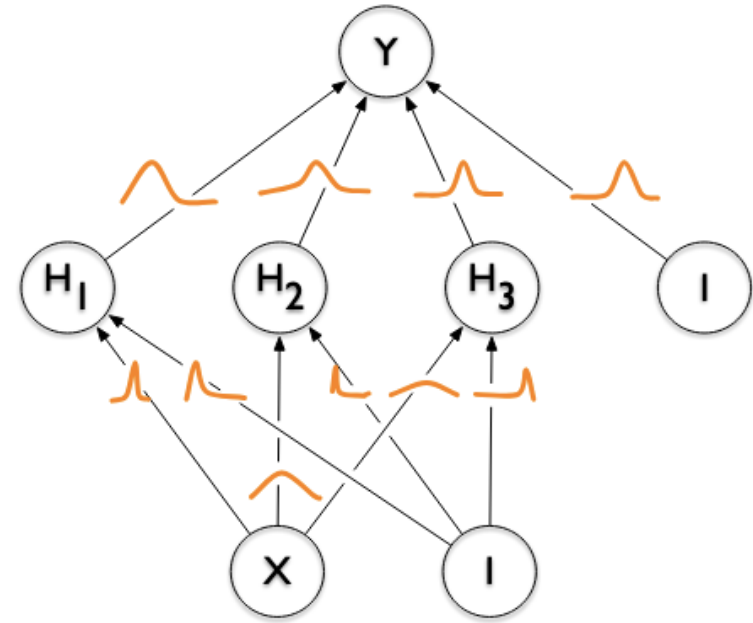
- **Hypotheses** here are weights for our learning model, i.e., weights of our neural networks that learns the transition dynamics
- The idea is to keep all the **hypotheses** that fit equally well the training set instead of committing to one, so that I can represent my uncertainty

How to Model Epistemic Uncertainty? A Bayesian Inference Method

Network with fixed weights



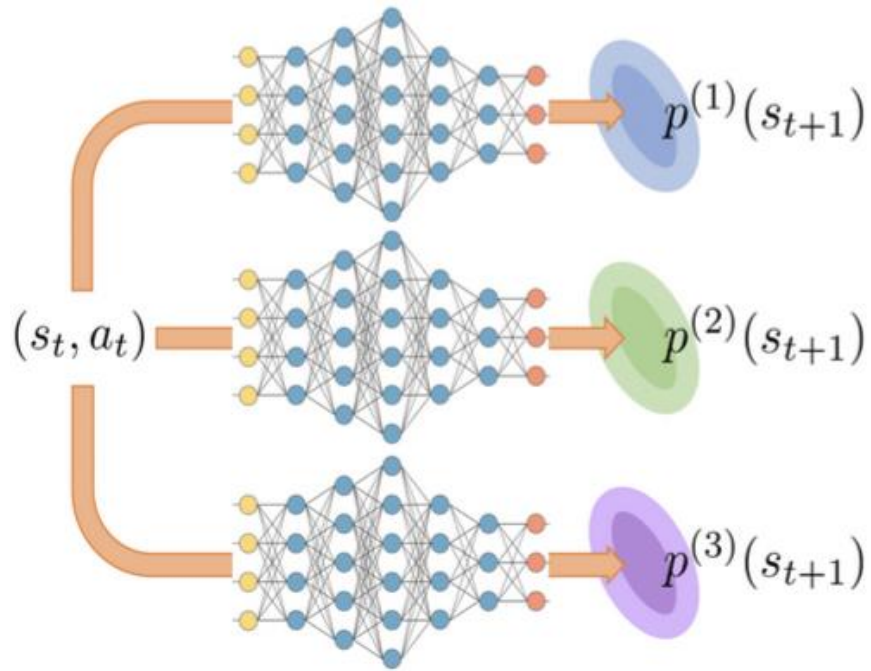
Network with a distribution of weights



predict according to:

$$\int p(\mathbf{s}_{t+1} | \mathbf{s}_t, \mathbf{a}_t, \theta) p(\theta | \mathcal{D}) d\theta$$

How to Model Epistemic Uncertainty? An Ensemble-of-Bootstraps Method

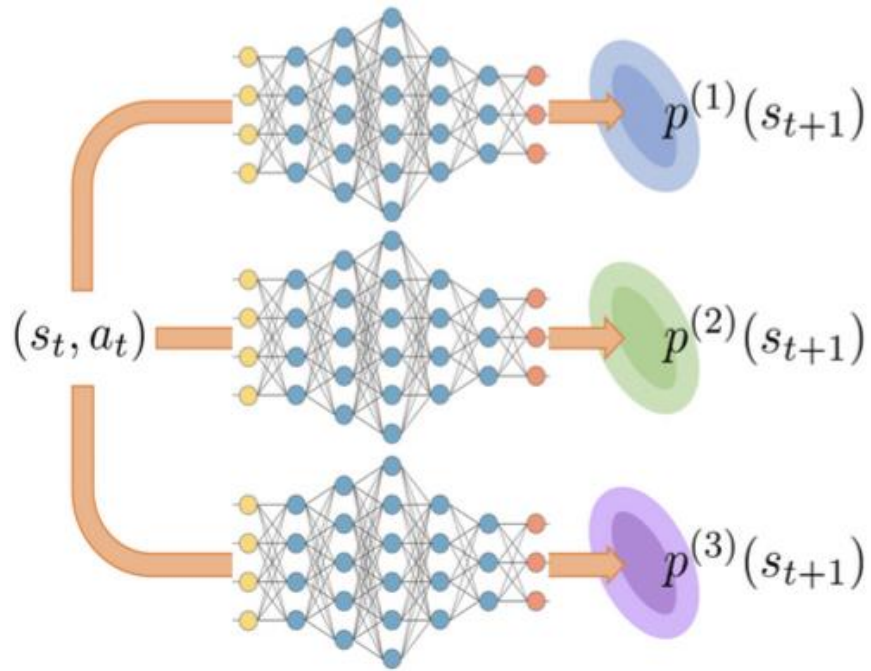


- Model ensembles are good approximation of Bayesian Nets:

$$p(\theta|D) \approx \frac{1}{N} \sum_i \sigma(\theta_i)$$

$$\int p(s_{t+1}|s_t, a_t, \theta) p(\theta|D) d\theta \approx \frac{1}{N} \sum_i p(s_{t+1}|s_t, a_t, \theta_i)$$

How to Model Epistemic Uncertainty? An Ensemble-of-Bootstraps Method



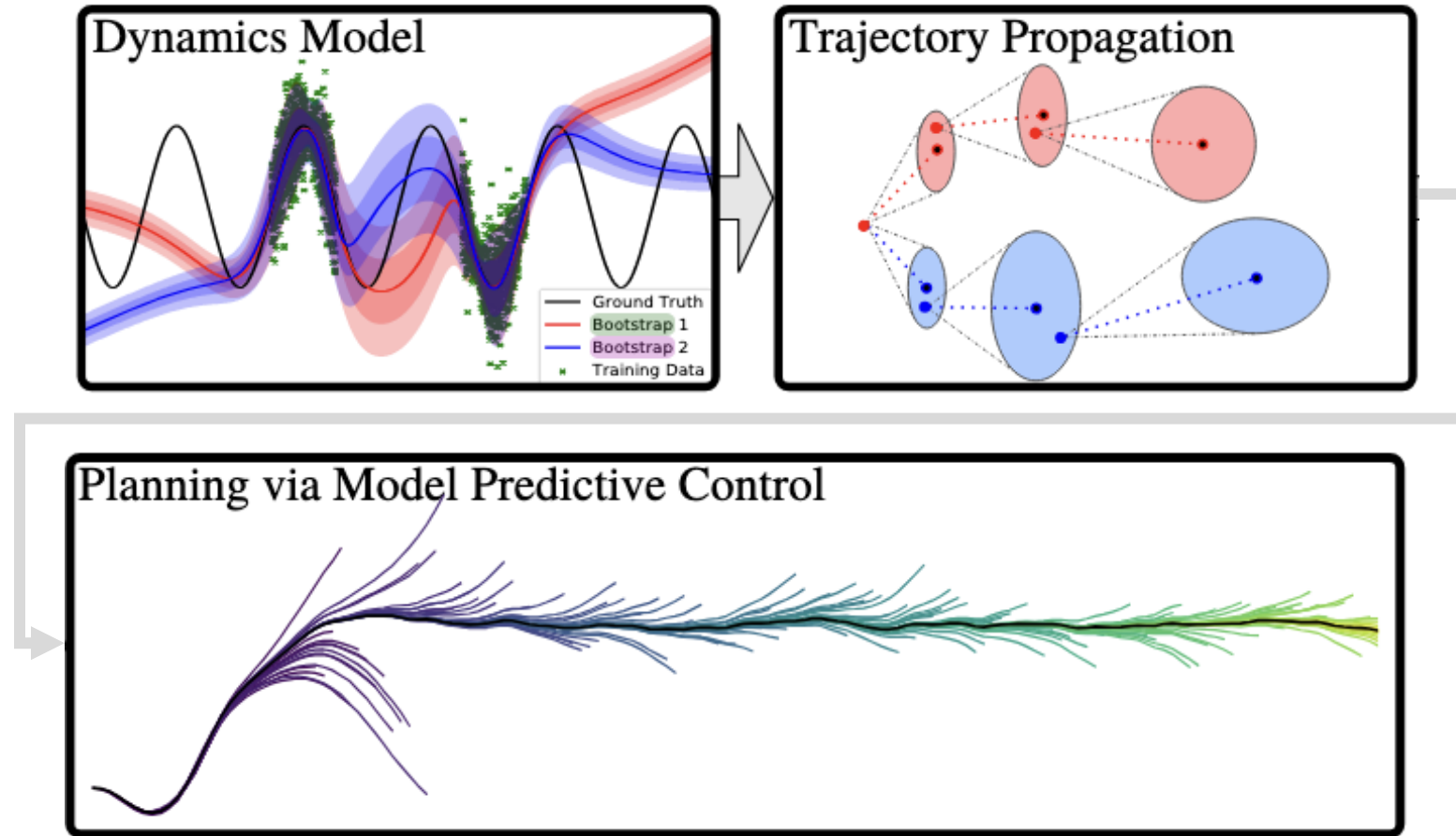
- Model ensembles are good approximation of Bayesian Nets:

$$p(\theta|D) \approx \frac{1}{N} \sum_i \sigma(\theta_i)$$

$$\int p(s_{t+1}|s_t, a_t, \theta) p(\theta|D) d\theta \approx \frac{1}{N} \sum_i p(s_{t+1}|s_t, a_t, \theta_i)$$

- Main idea: generate “independent” datasets for each “independent” model:
 - For seen data, they all agree (low-entropy prediction)
 - For unseen data, they disagree (high-entropy prediction)

Unrolling an Ensemble of Bootstraps



Unrolling an Ensemble of Bootstraps

Algorithm 1 Our model-based MPC algorithm ‘*PETS*’:

- 1: Initialize data $\mathbb{D}$ with a random controller for one trial.
 - 2: **for** Trial $k = 1$ to K **do**
 - 3: Train a *PE* dynamics model $\tilde{f}$ given $\mathbb{D}$. Trajectory sampling: the strategy to sample bootstrap as the dynamics model
 - 4: **for** Time $t = 0$ to TaskHorizon **do**
 - 5: **for** Actions sampled $\mathbf{a}_{t:t+T} \sim \text{CEM}(\cdot)$, 1 to NSamples **do**
 - 6: Propagate state particles $\mathbf{s}_\tau^p$ using *TS* and $\tilde{f}|\{\mathbb{D}, \mathbf{a}_{t:t+T}\}$.
 - 7: Evaluate actions as $\sum_{\tau=t}^{t+T} \frac{1}{P} \sum_{p=1}^P r(\mathbf{s}_\tau^p, \mathbf{a}_\tau)$
 - 8: Update $\text{CEM}(\cdot)$ distribution.
 - 9: Execute first action $\mathbf{a}_t^*$ (only) from optimal actions $\mathbf{a}_{t:t+T}^*$.
 - 10: Record outcome: $\mathbb{D} \leftarrow \mathbb{D} \cup \{\mathbf{s}_t, \mathbf{a}_t^*, \mathbf{s}_{t+1}\}$.
-

Results: Performant Model-Based RL

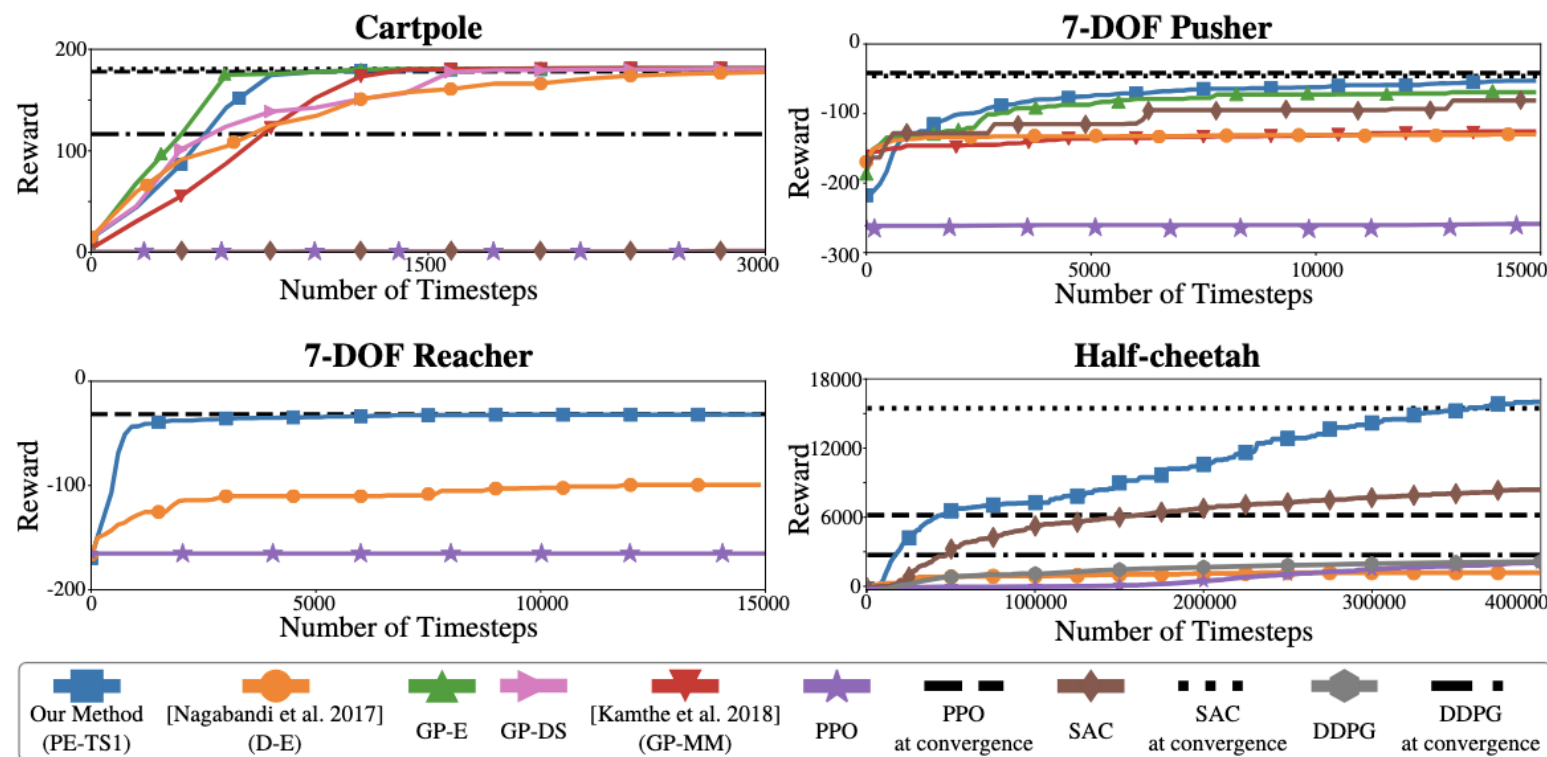
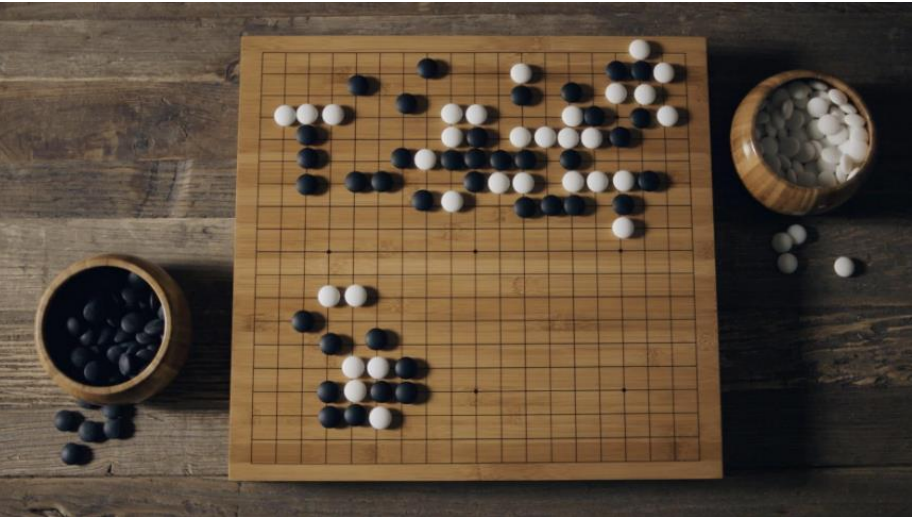


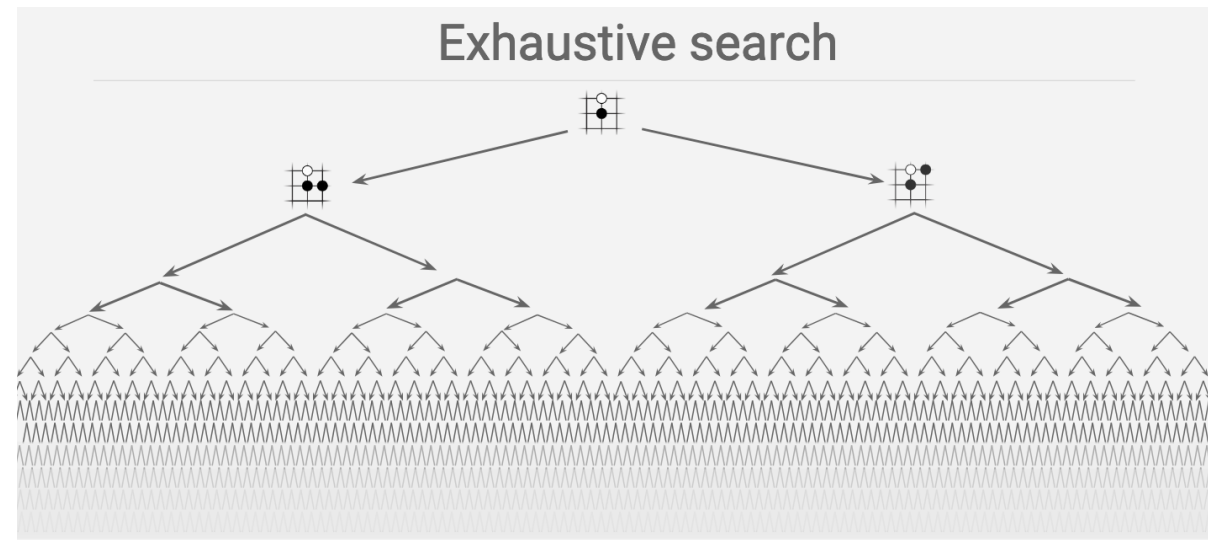
Figure 3: Learning curves for different tasks and algorithm. For all tasks, our algorithm learns in under 100K time steps or 100 trials. With the exception of Cartpole, which is sufficiently low-dimensional to efficiently learn a GP model, our proposed algorithm significantly outperform all other baselines. For each experiment, one time step equals 0.01 seconds, except Cartpole with 0.02 seconds. For visual clarity, we plot the average over 10 experiments of the maximum rewards seen so far.

From AlphaGO (Model-free RL) to
MuZero (Model-based RL)

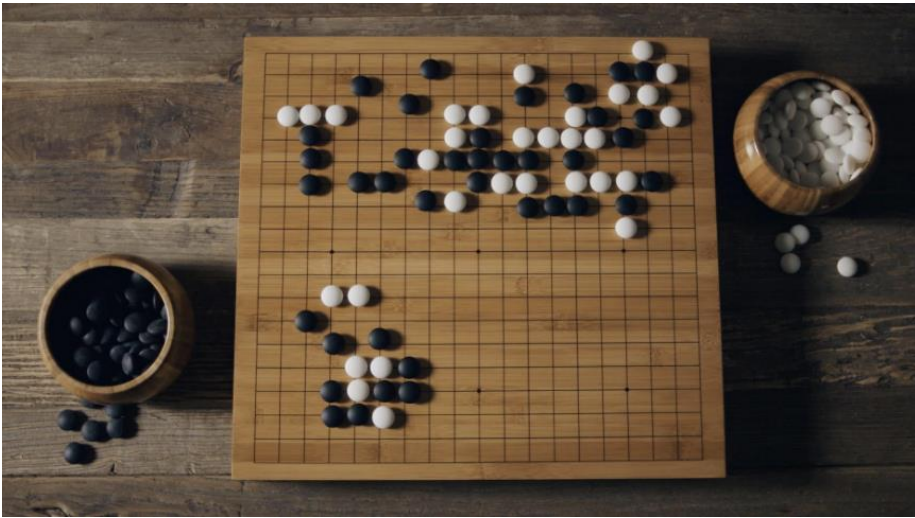
Go is Hard...



- Game tree complexity = b^d (b : game breadth, d : game depth)
 - Search space is huge
 - “Impossible” for computers to evaluate who is winning

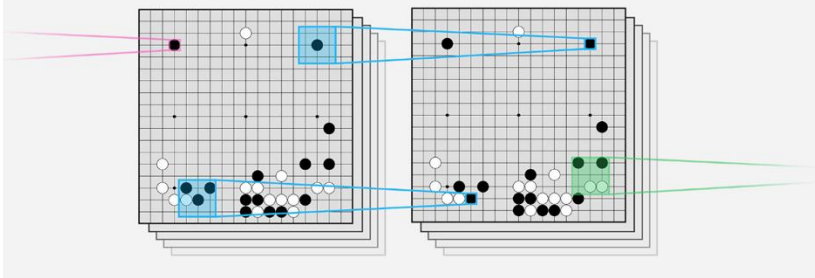


Go is Hard...

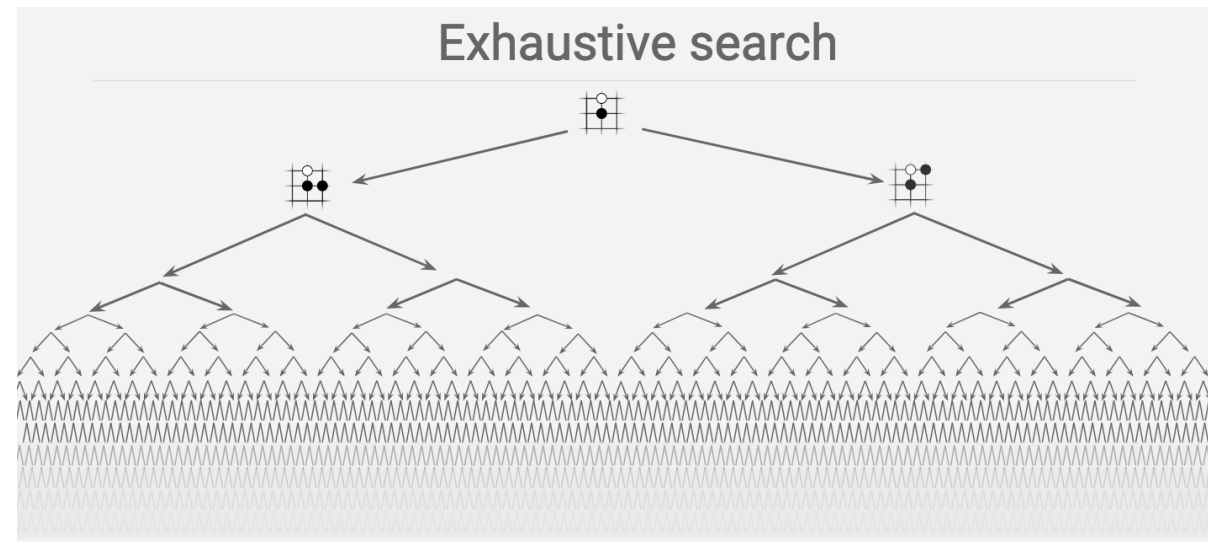


Visual encoder:

Convolutional neural network



- Game tree complexity = b^d (b : game breadth, d : game depth)
 - Search space is huge
 - “Impossible” for computers to evaluate who is winning



Imitation Learning with Expert Demonstrations

Policy network: 12 layer convolutional neural network

Training data: 30M positions from human expert games (KGS 5+ dan)

Training algorithm: maximise likelihood by stochastic gradient descent

$$\Delta\sigma \propto \frac{\partial \log p_{\sigma}(a|s)}{\partial \sigma}$$

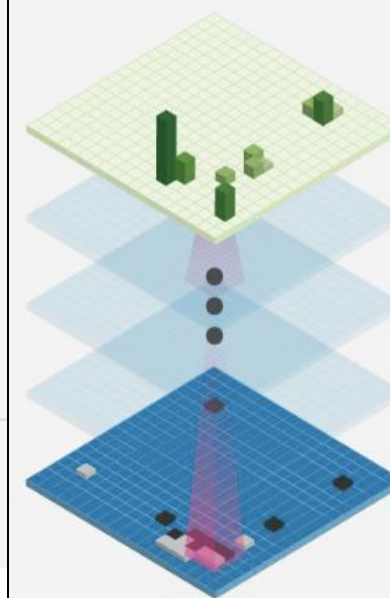
Training time: 4 weeks on 50 GPUs using Google Cloud

Results: 57% accuracy on held out test data (state-of-the art was 44%)



Policy network:

Move probabilities



Position

Reinforcement Learning with Policy Gradient

Policy network: 12 layer convolutional neural network

Training data: games of self-play between policy network

Training algorithm: maximise wins z by policy gradient reinforcement learning

$$\Delta\sigma \propto \frac{\partial \log p_{\sigma}(a|s)}{\partial \sigma} z$$

z denotes if the current state leads to a winning game

Training time: 1 week on 50 GPUs using Google Cloud

Results: 80% vs supervised learning. Raw network ~3 amateur dan.



Can We Do Better? Yes! MCTS + Policy/Value Networks

Value network: 12 layer convolutional neural network

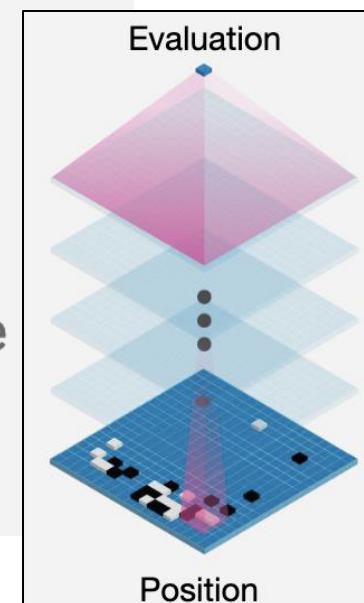
Training data: 30 million games of self-play

Training algorithm: minimise MSE by stochastic gradient descent

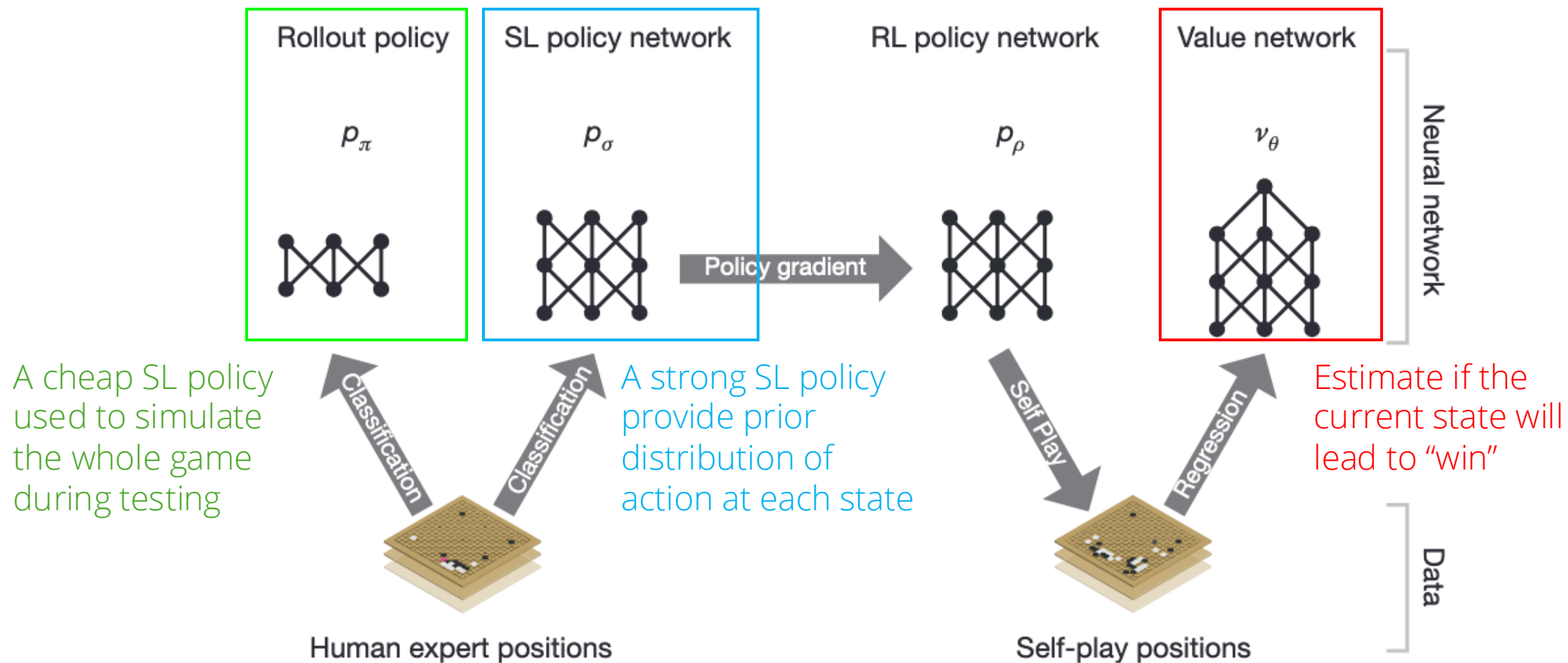
$$\Delta\theta \propto \frac{\partial v_{\theta}(s)}{\partial \theta} (z - v_{\theta}(s))$$

Training time: 1 week on 50 GPUs using Google Cloud

Results: First strong position evaluation function - previously thought impossible

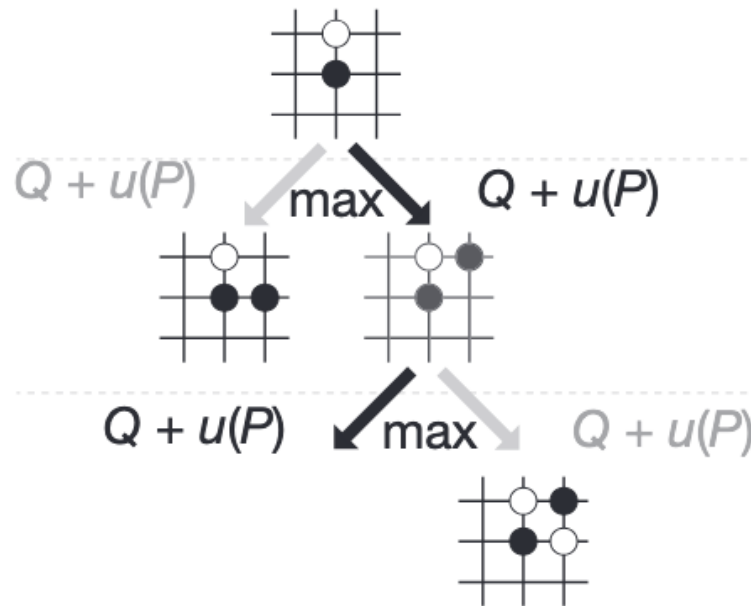


Can We Do Better? Yes! MCTS + Policy/Value Networks



Monte-Carlo Tree Search During Testing

Selection: selecting actions within the expanded tree

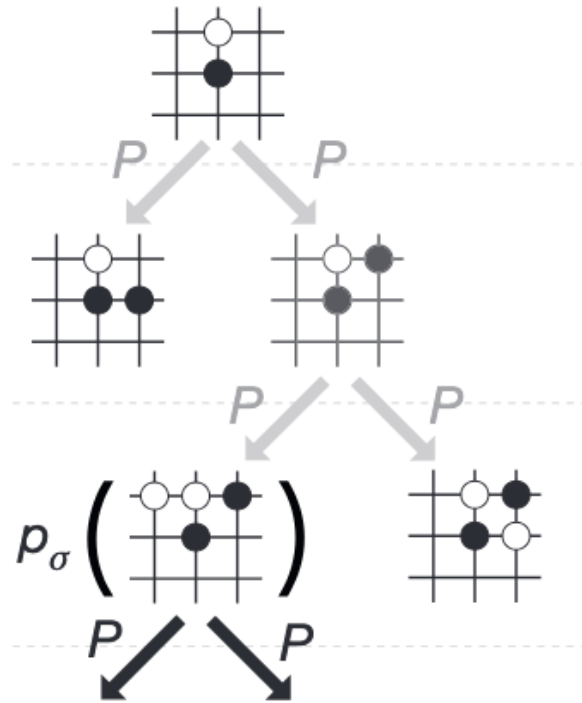


$$a_t = \operatorname{argmax}_a \underbrace{Q(s_t, a)}_{\text{exploitation}} + \underbrace{c\sqrt{N(s_t)} \frac{\pi_\theta(a_t | s_t)}{1 + N(s_t, a_t)}}_{\text{exploration}}$$

- a_t - action selected at time step t from state s_t
- $Q(s_t, a)$ - average reward collected so far from MC simulations
- $\pi_\theta(a | s)$ - prior expert probability provided by the SL policy p_σ
- $N(s, a)$ - number of times we have taken action a from state s during MC simulations

Monte-Carlo Tree Search During Testing

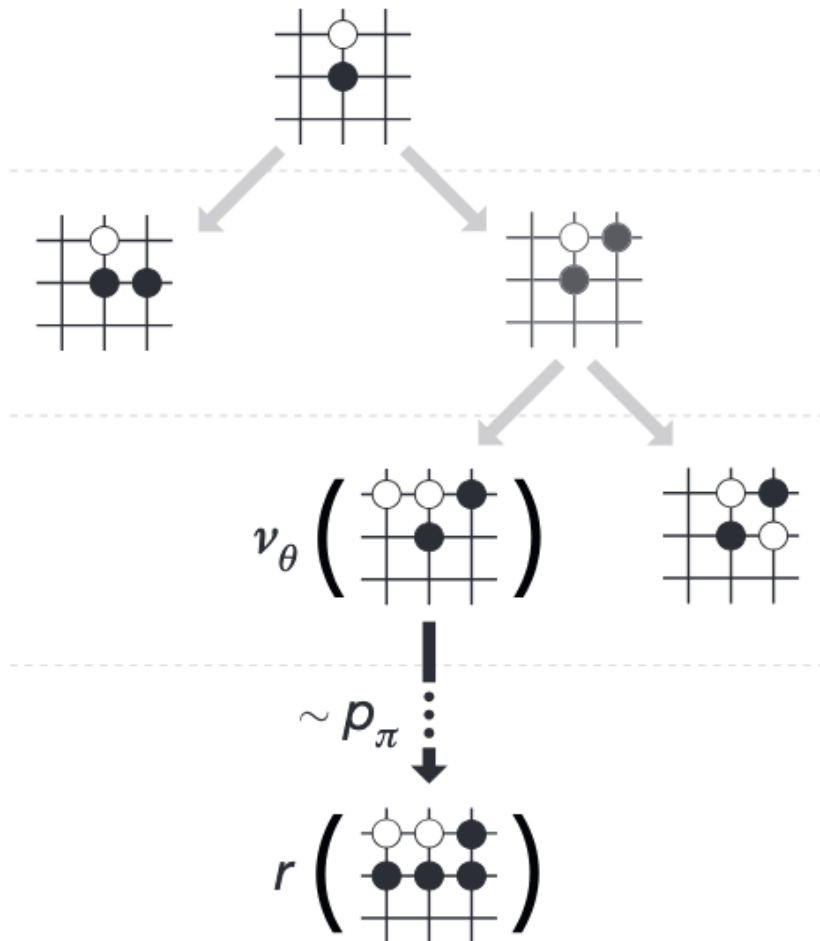
Expansion: when reaching a leaf, play the action with highest score from p_σ



- When leaf node is reached, it has a chance to be expanded
- Processed once by SL policy network (p_σ) and stored as prior probs $P(s, a)$
- Pick child node with highest prior prob

Monte-Carlo Tree Search During Testing

Simulation/Evaluation: use the rollout policy to reach to the end of the game



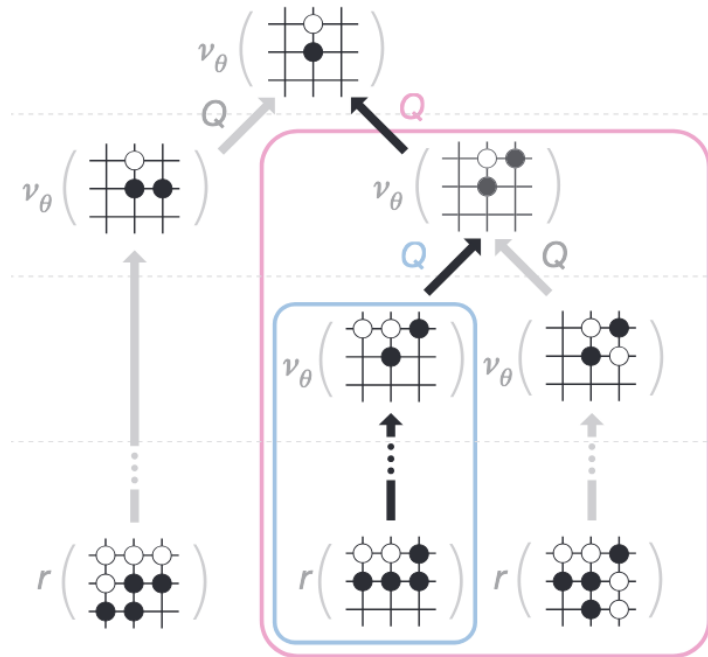
- From the selected leaf node, run multiple simulations in parallel using the rollout policy
- Evaluate the leaf node as:

$$V(s_L) = (1 - \lambda)v_\rho(s_L) + \lambda z_L$$

- v_ρ : value from the trained value function for board position s_L
- z_L : Reward from fast rollout p_x
 - Played until terminal step
- λ - mixing parameter

Monte-Carlo Tree Search During Testing

- Backup: update visitation counts and recorded rewards for the chosen path inside the tree



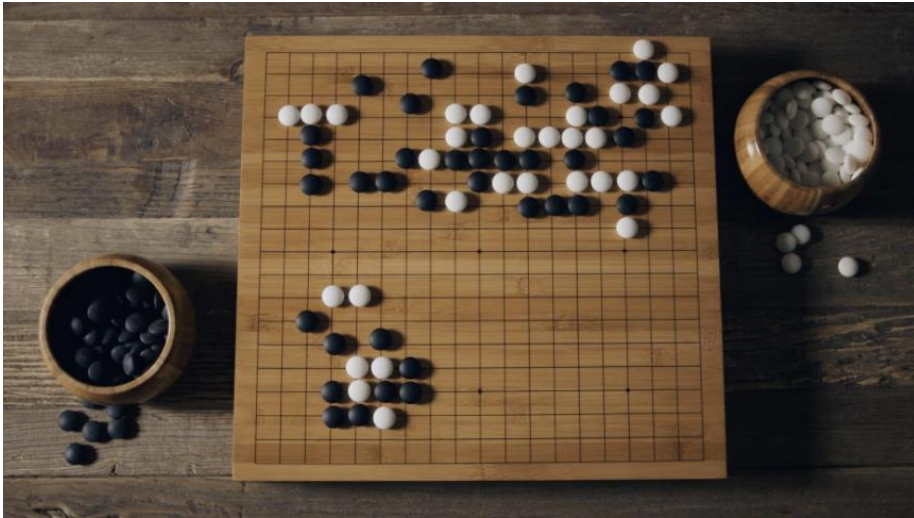
$$N(s, a) = \sum_{i=1}^n \mathbf{1}_{(s,a) \in \tau_i}$$

$$Q(s, a) = \frac{1}{N(s, a)} \sum_{i=1}^n \mathbf{1}_{(s,a) \in \tau_i} V(s_L^i)$$

- Extra index is to denote the i simulation, n total simulations
- Update visit count and mean reward of simulations passing through node
- Once MCTS completes, *the algorithm chooses the most visited move from the root position.*

To Do Monte-Carlo Tree Search During Test Time, We Need to Simulate

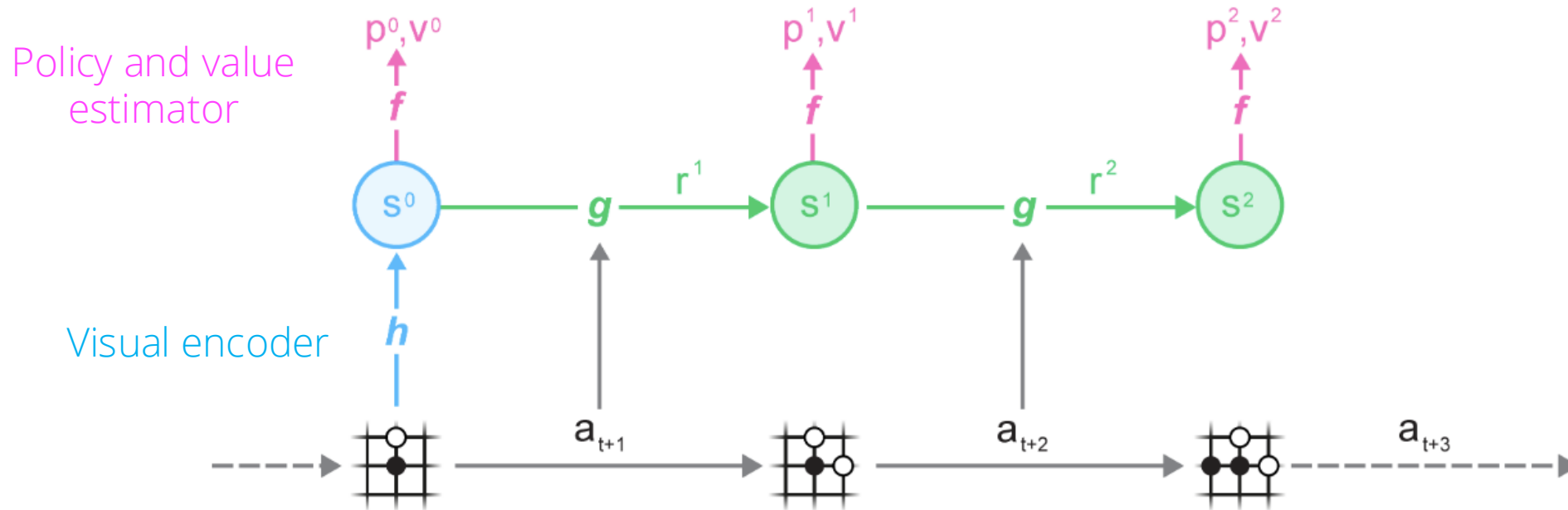
The rule is clear, trivial to build a simulator



The rule is complex, non-trivial to build a simulator

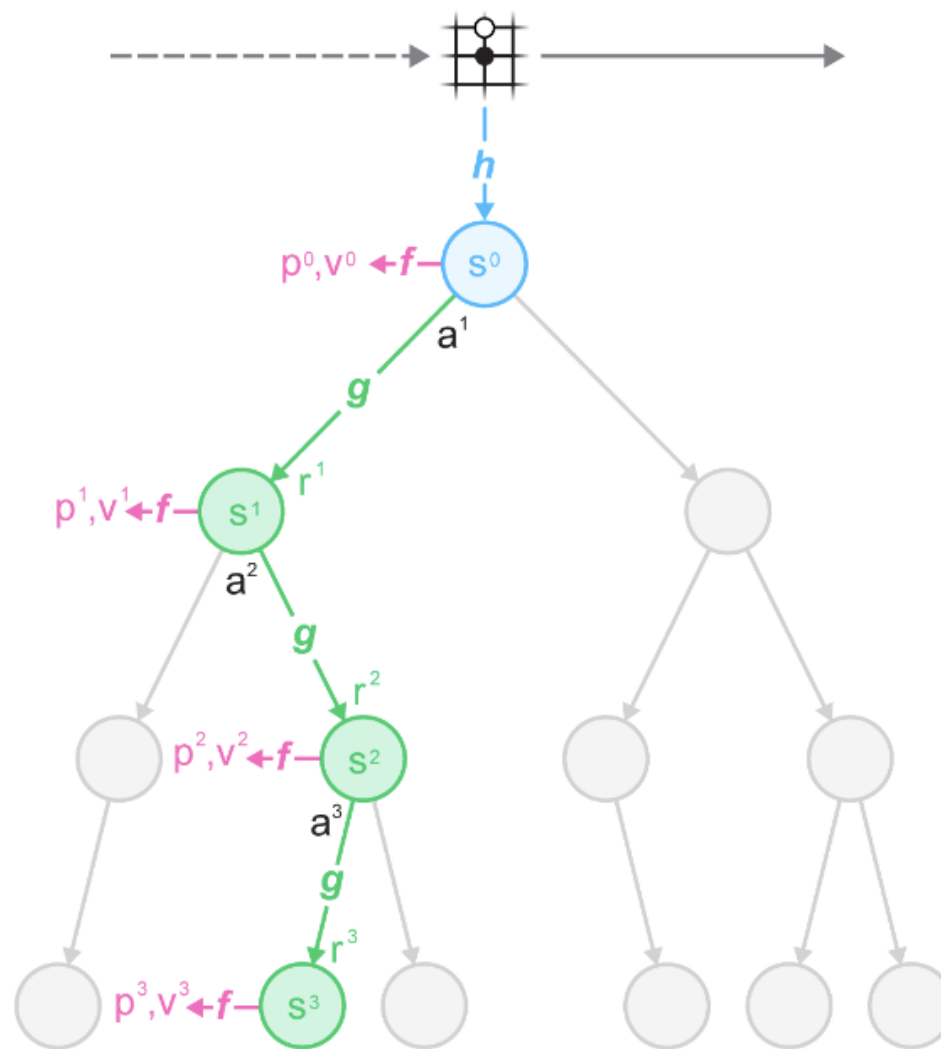


Learn a Dynamics Model for Different Games



- The dynamic model $r^t, s^t = g_\theta(s^{t-1}, a^t)$
- State s^t has no semantics of environment state attached to it, simply the hidden state of the overall model, and its sole purpose is to accurately predict relevant, future quantities: policies, values, and rewards

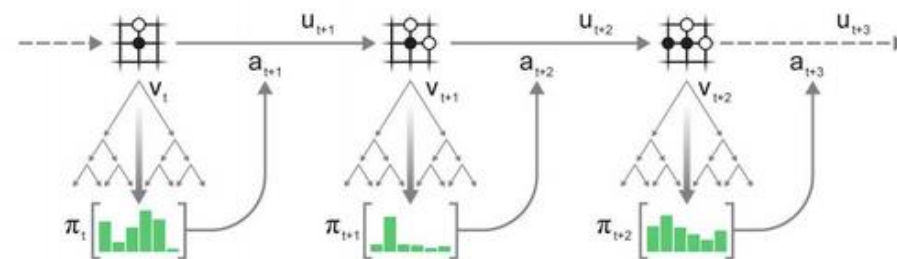
Use the Learned Dynamic Model for MCTS



Use MCTS to Improve Policy and Value Estimation

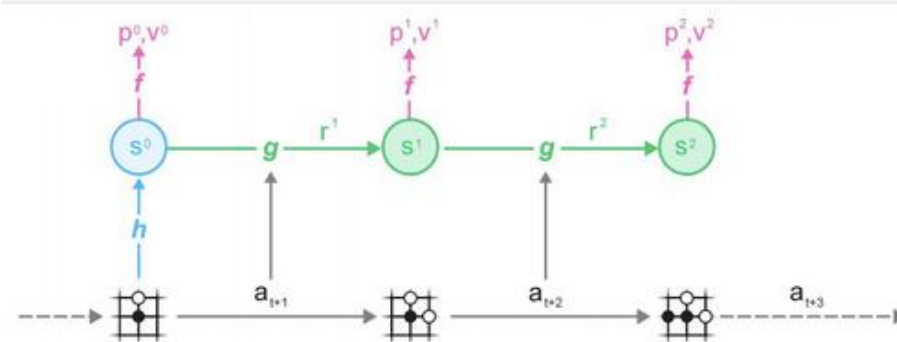
Policy MCTS Distillation

$$p_t \approx \pi_t$$



Reward

$$r_t \approx u_t$$

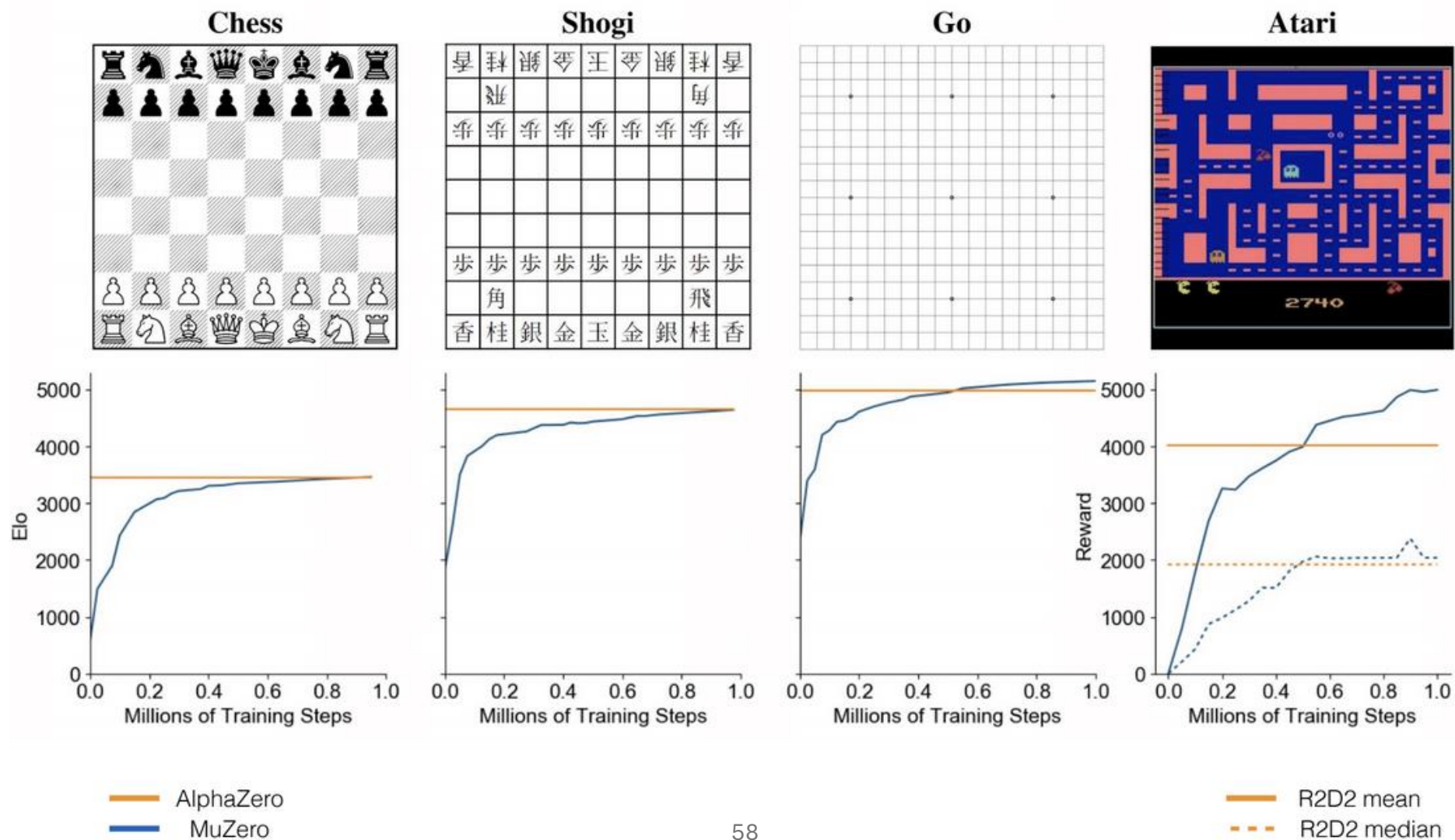


Value n-step bootstrapping

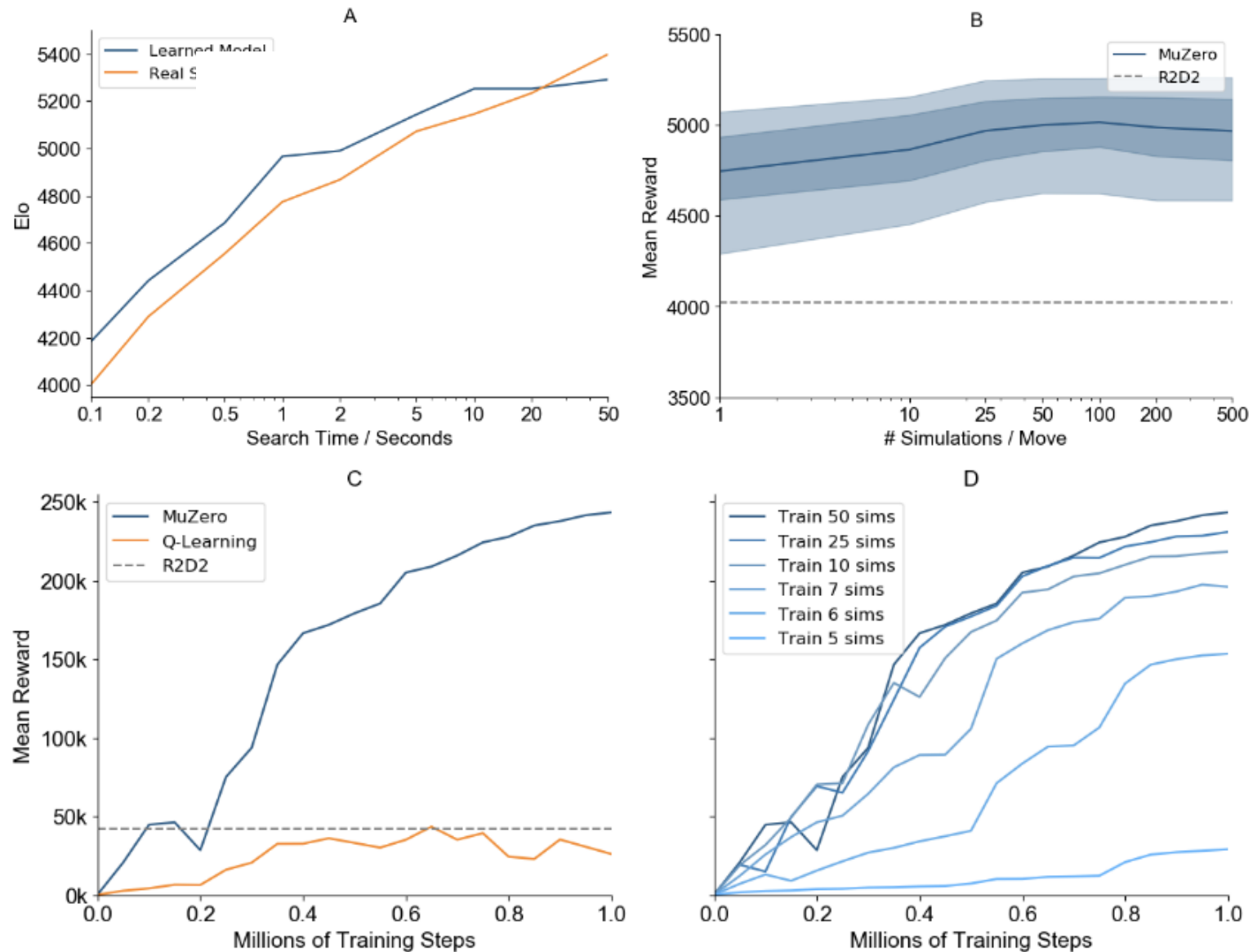
$$v_t \approx u_{t+1} + \gamma u_{t+2} + \dots + \gamma^{n-1} u_{t+n} + \gamma^n v_{t+n}$$



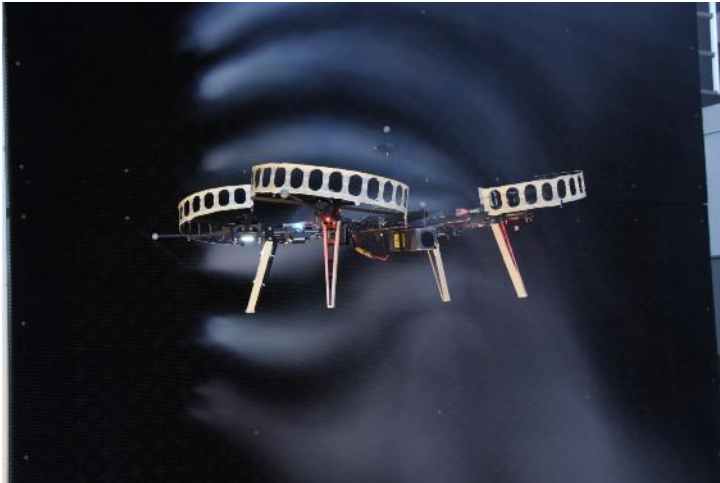
Results



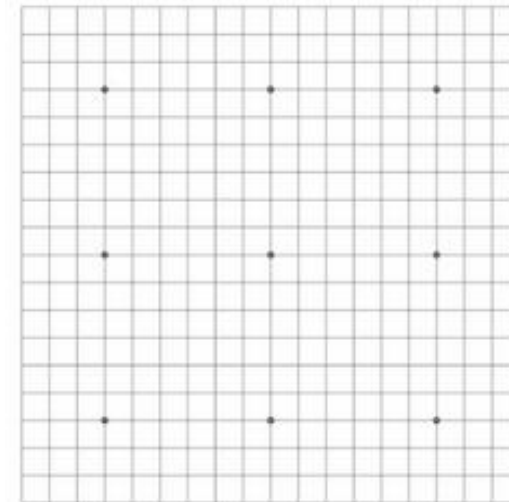
MuZero is Much More Sample Efficient



What if Real-world Model is not Stationary?



V.S



AnyCar to Anywhere Learning Universal Dynamics
Model for Agile and Adaptive Mobility. Xiao et al.

Adaptive Model-Based RL

- The formulation of the dynamics model:

$$x_{t+1:t+H} \approx f_{\theta}^{\text{AnyCar}} \left(\underbrace{\hat{x}_{t-K:t}, \hat{a}_{t-K:t-1}}_{\text{noisy state and action history}}, \underbrace{a_{t:t+H-1}}_{\text{future actions}} \right),$$

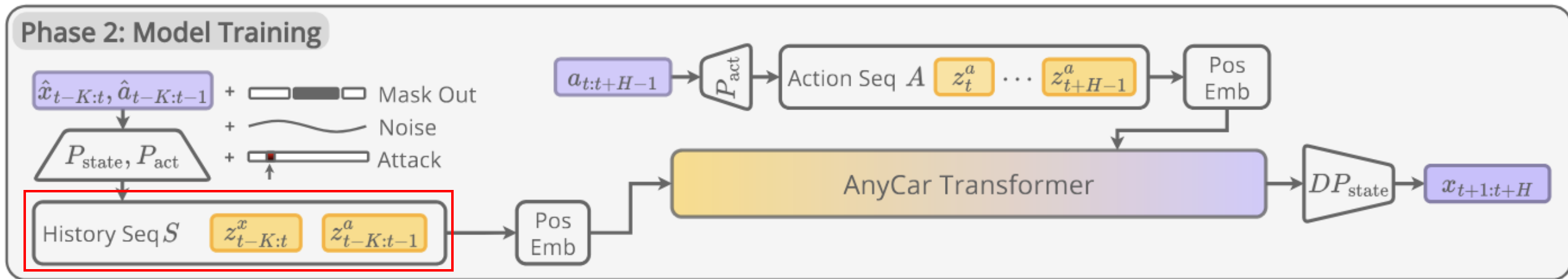
- The definition of states:

$$x_t \triangleq [p_t^x, p_t^y, \psi_t, \dot{p}_t^x, \dot{p}_t^y, \omega]$$

where (p_t^x, p_t^y) denotes the car position, ψ_t is the heading angle, $(\dot{p}_t^x, \dot{p}_t^y)$ denotes the velocity, and ω is the angular velocity

- The definition of actions a_t contain the throttle and steering angle

Adaptive Model-Based RL



Historical state-action pairs reveal the current dynamics of the environment



Learning the Dynamics Model from Existing Synthesized Demonstrations

Phase 1: Data Collection



- The simulation data generation has three sources of diversity:
 - Dynamics
 - Scenario
 - Controller
- Curriculum data generation:
 - Off-policy data: Pure pursuit controller for steering δ and a PD controller for throttle T , to track randomly synthesized reference track in simulation
 - On-policy data: NN-MPPI controller to track agile trajectories in simulation
 - Real-world data: NN-MPPI controller to track agile trajectories in the real world

Deployment

Phase 3: Deployment



1/16 Scale Car



20 cm

1/10 Scale Car



31 cm

w/ optional 3D-Printed Wheels

State Estimation Method

(a) MoCap



(b) ZED-VIO



(c) LiDAR-SLAM

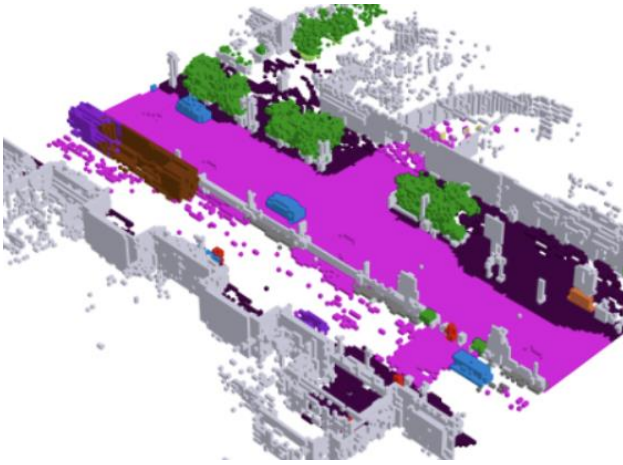


Adaptive Model-Based RL is Not Solved

- What if the state-action dynamic changes suddenly/accidentally (e.g. splitting water on the ground)?
- How about contact-rich manipulation?

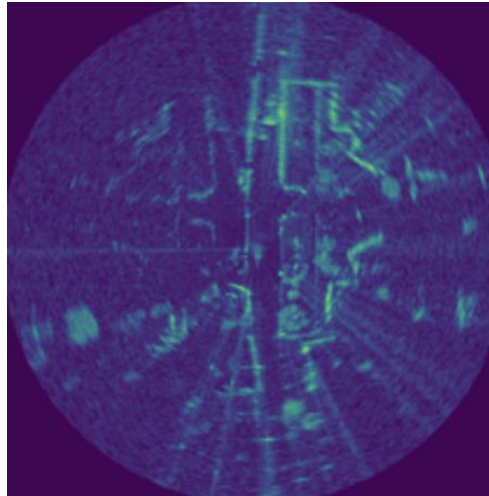
Usually We Have High-dimensional Sensory Input What is the State Representations of the Dynamics Models?

Occupancy



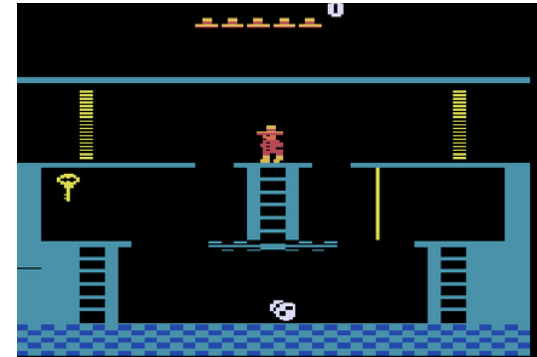
OccWorld: Learning a 3D Occupancy World Model for Autonomous Driving. Zhen et al.

Radar



CramNet: Camera-Radar Fusion with Ray-Constrained Cross-Attention for Robust 3D Object Detection. Hwang et al.

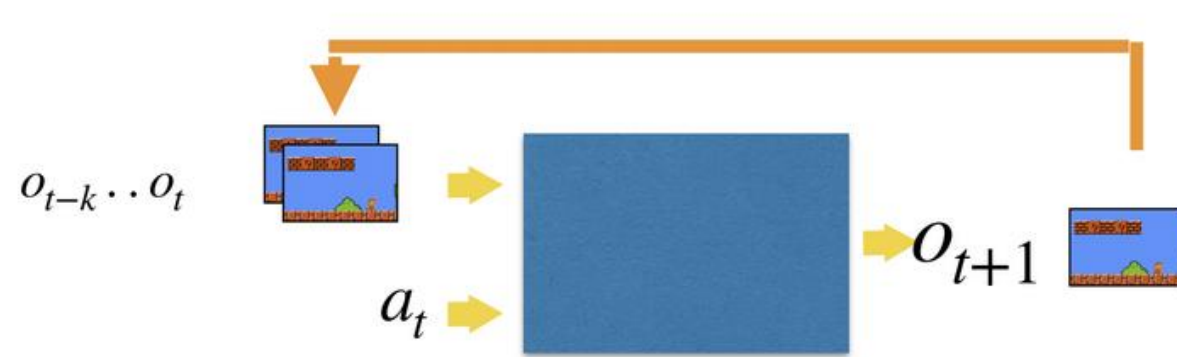
Image



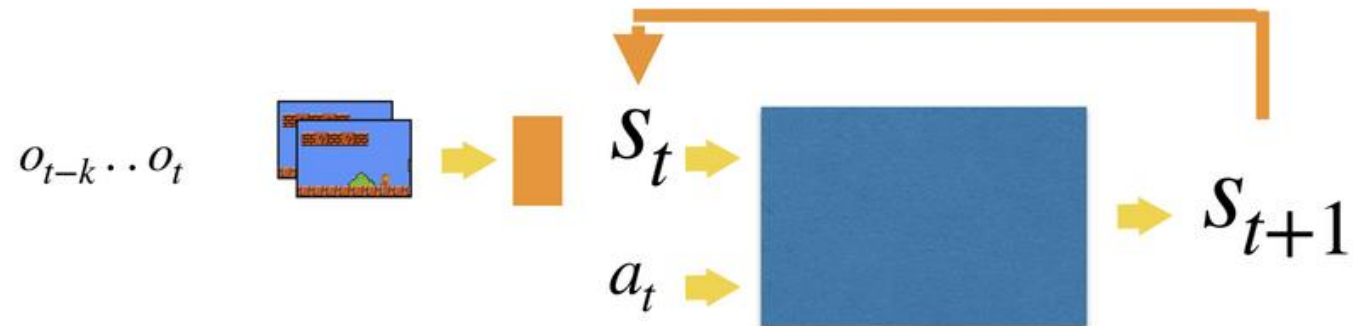
...

Usually We Have High-dimensional Sensory Input What is the State Representations of the Dynamics Models?

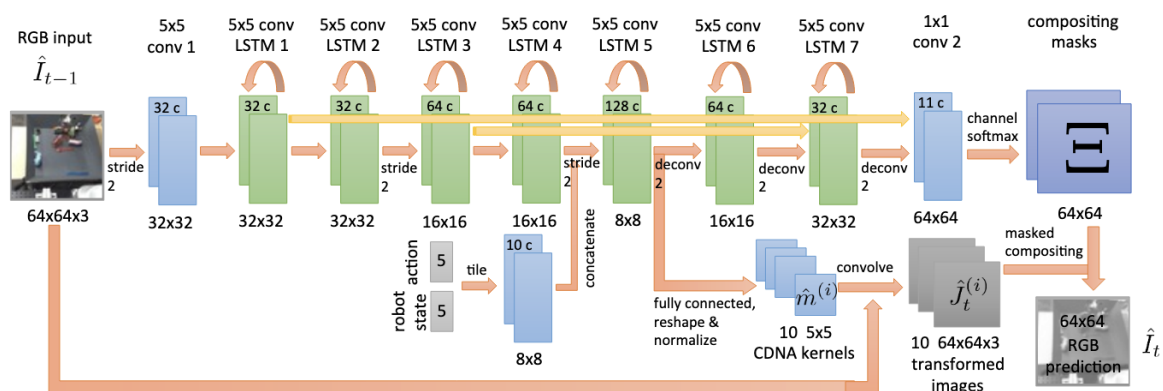
Unrolling in the observation space: the model is trained to predict future observations.



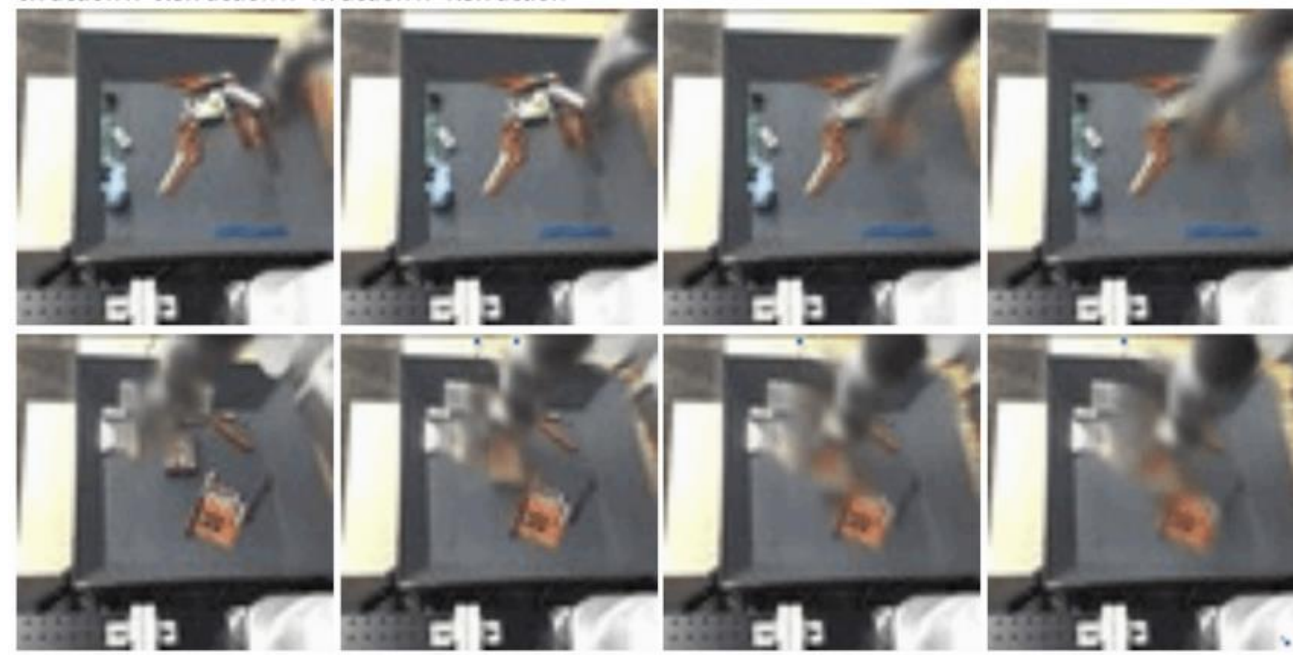
Unrolling in the latent space.



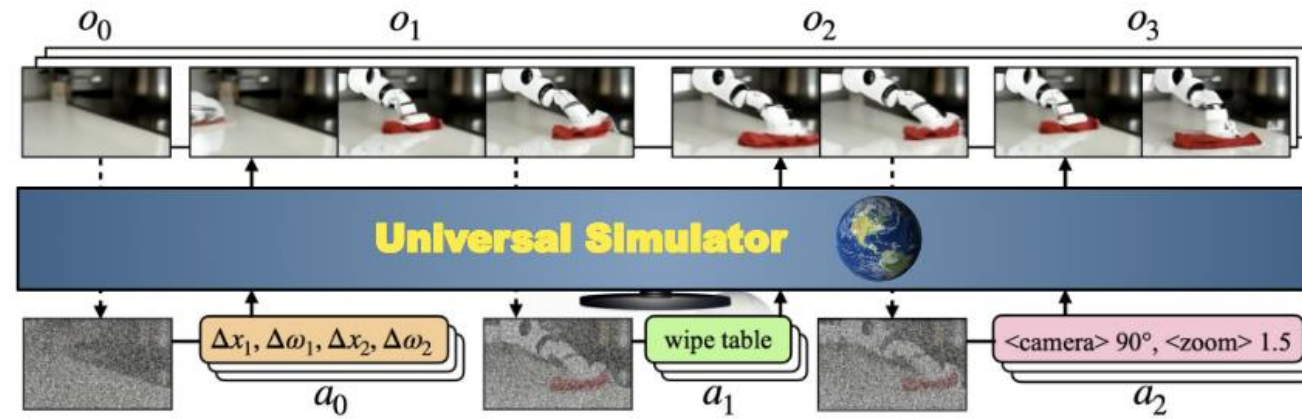
Unrolling in the Observation Space



0x action // 0.5x action // 1x action // 1.5x action



Unrolling in the Observation Space



Simulated rollout
from $\Delta x, \Delta y$ moving
left, right, down, up



Simulated rollout
from $\Delta x, \Delta y$
moving diagonally

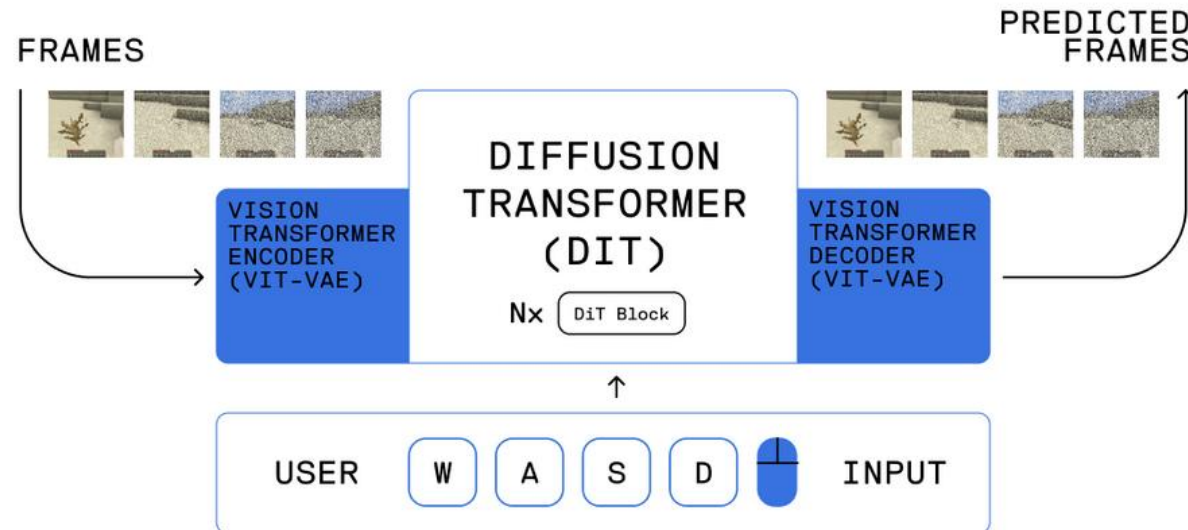


Real-robot execution
of “move blue cube to
green circle”



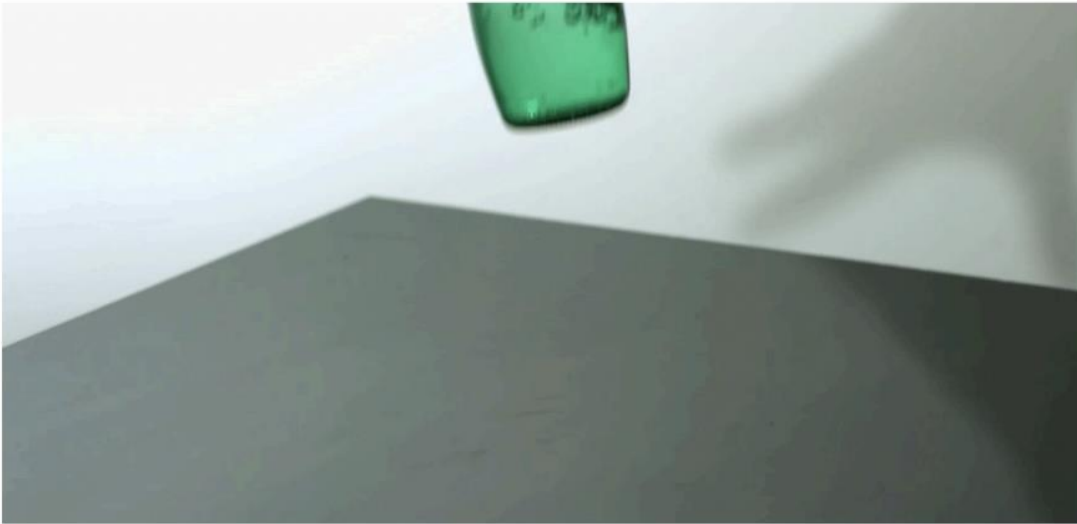
Unrolling in the Observation Space

Oasis 500M



What are the Problems?

What will happen on dropping the bottle?

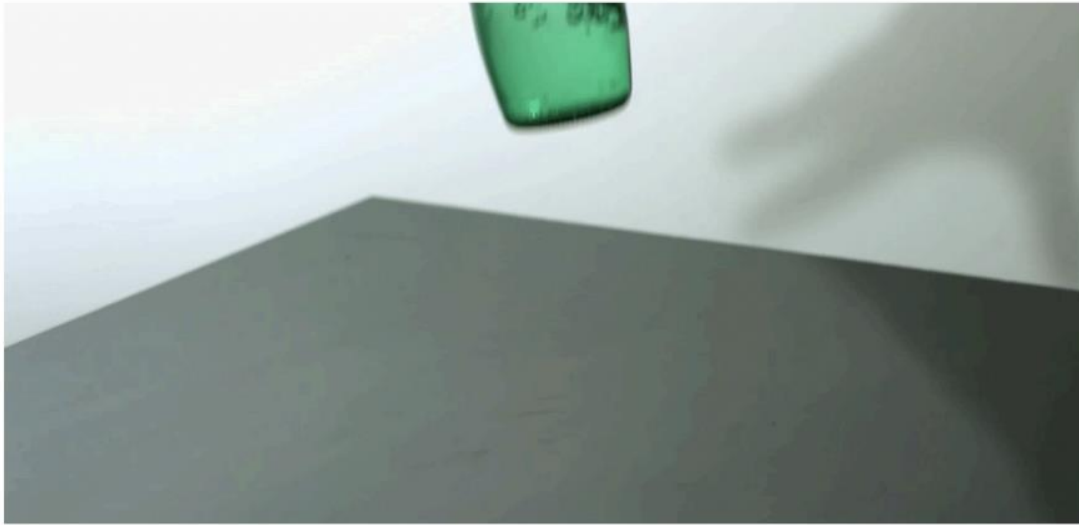


Simulating the water splash and breaking bottle is hard...
Do we need to model them?



What are the Problems?

What will happen on dropping the bottle?



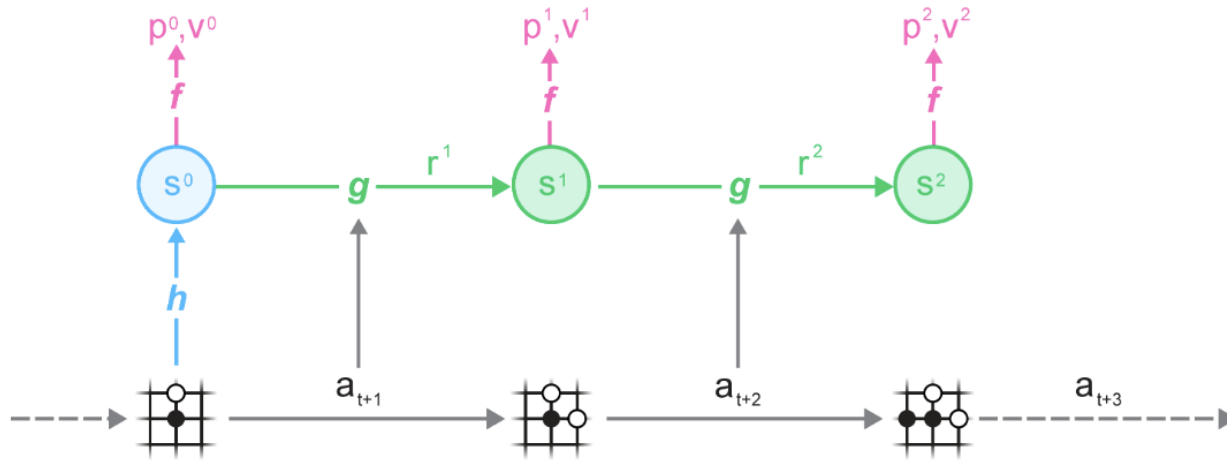
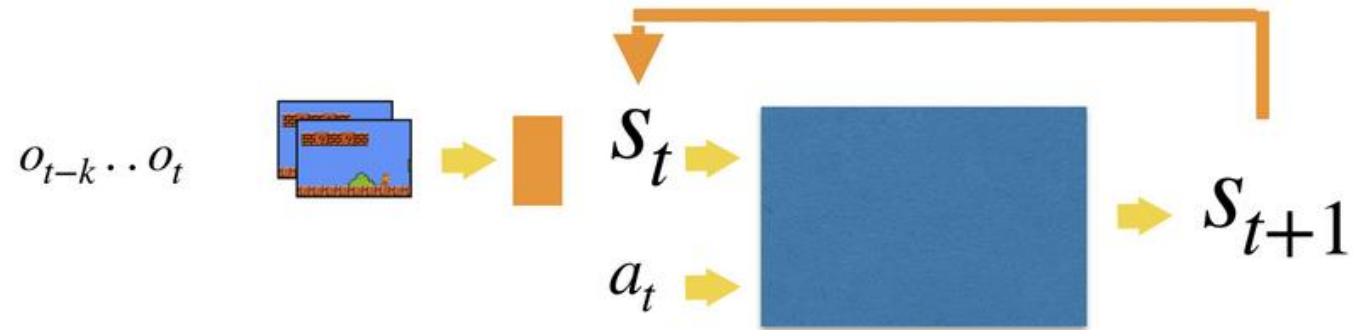
Simulating the water splash and breaking bottle is hard...
Do we need to model them?



Simulating every pixel takes time, can we afford to do test-time searching?



Unroll in the Latent Space



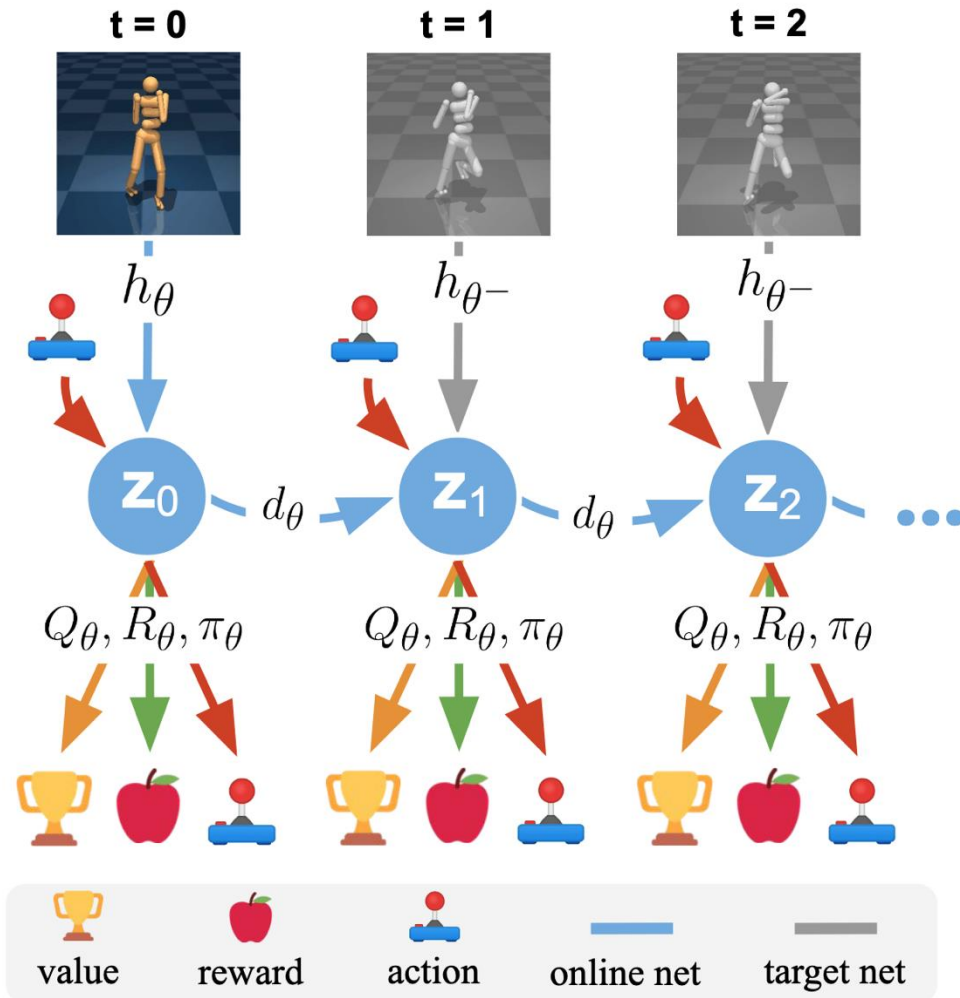
Remember in MuZero:

- The state representations s have no semantics of the environment, but encode information of policies, values and rewards
- Plan in the latent space

Task-Relevant Latent Dynamic Models

- Model-based RL typically **models everything** in the environment
- Model-free RL only **retains information predictive of reward**
- Idea:
 - Infer a latent space under all task-relevant losses: action prediction, reward prediction, Q value prediction
 - Learn jointly the dynamics model and the value estimation function

Task-Relevant Latent Dynamic Models



Representation:
 Latent dynamics:
 Reward:
 Value:
 Policy:

$$\begin{aligned} \mathbf{z}_t &= h_\theta(\mathbf{s}_t) \\ \mathbf{z}_{t+1} &= d_\theta(\mathbf{z}_t, \mathbf{a}_t) \\ \hat{r}_t &= R_\theta(\mathbf{z}_t, \mathbf{a}_t) \\ \hat{q}_t &= Q_\theta(\mathbf{z}_t, \mathbf{a}_t) \\ \hat{\mathbf{a}}_t &\sim \pi_\theta(\mathbf{z}_t) \end{aligned}$$

In Vanilla MPC

Algorithm:

- 
1. (Testing) Plan an optimal sequence of actions $a_{0:T}^*$:
$$a_{0:T}^* = \max_{a_{0:T}} r(s_0, a_0) + r(f(s_0, a_0), a_1) + r(f(f(s_0, a_0), a_1), a_2) + \dots$$
 2. Execute the first planned action a_0^* and observe resulting state s_1

- Major challenges of MPC:
 - MPC optimizes a trajectory over a shorter, finite horizon, which yields only temporally local optimal solutions
 - Long-horizon planning has compounding errors of dynamics models and is computationally expensive

In Vanilla MPC

Algorithm:

- 
1. (Testing) Plan an optimal sequence of actions $a_{0:T}^*$:
$$a_{0:T}^* = \max_{a_{0:T}} r(s_0, a_0) + r(f(s_0, a_0), a_1) + r(f(f(s_0, a_0), a_1), a_2) + \dots$$
 2. Execute the first planned action a_0^* and observe resulting state s_1

- Major challenges of MPC:
 - MPC optimizes a trajectory over a shorter, finite horizon, which yields only temporally local optimal solutions
 - Long-horizon planning has compounding errors of dynamics models and is computationally expensive

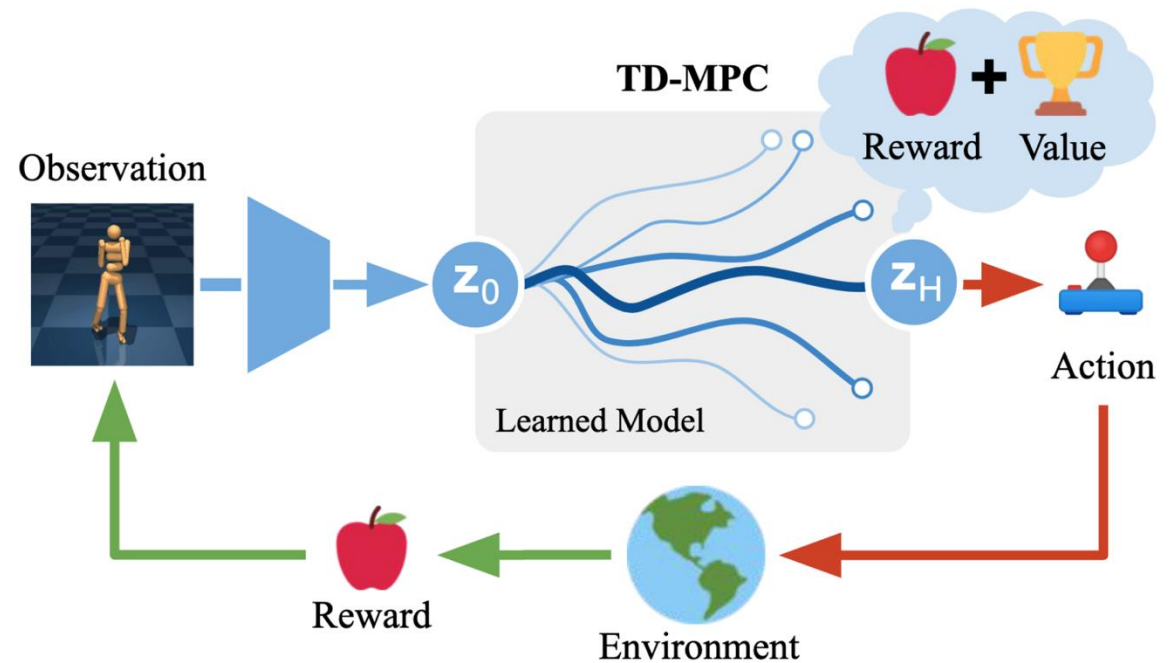
Instead of unrolling infinite-horizon trajectories, use a value function to approximate the globally optimal solution!

TD Learning for MPC

Inference (planning)

- Planning in latent space
- Return estimate:

$$\mathbb{E}_{\Gamma} \left[\underbrace{\gamma^H Q_{\theta}(\mathbf{z}_H, \mathbf{a}_H)}_{\text{Value}} + \underbrace{\sum_{t=0}^{H-1} \gamma^t R_{\theta}(\mathbf{z}_t, \mathbf{a}_t)}_{\text{Rewards}} \right]$$



Model Predictive Path Integral Control for Planning

Algorithm 1 TD-MPC (*inference*)

Require: θ : learned network parameters

μ^0, σ^0 : initial parameters for $\mathcal{N}$

N, N_π : num sample/policy trajectories

$\mathbf{s}_t, H$: current state, rollout horizon

- 1: Encode state $\mathbf{z}_t \leftarrow h_\theta(\mathbf{s}_t)$ $\triangleleft$ *Assuming TOLD model*
 - 2: **for** each iteration $j = 1..J$ **do**
 - 3: Sample N traj. of len. H from $\mathcal{N}(\mu^{j-1}, (\sigma^{j-1})^2 \mathbf{I})$
 - 4: **Sample** N_π traj. of length H using π_θ, d_θ
 *// Estimate trajectory returns ϕ_Γ using $d_\theta, R_\theta, Q_\theta$,
 starting from $\mathbf{z}_t$ and initially letting $\phi_\Gamma = 0$:*
 - 5: **for** all $N + N_\pi$ trajectories $(\mathbf{a}_t, \mathbf{a}_{t+1}, \dots, \mathbf{a}_{t+H})$ **do**
 - 6: **for** step $t = 0..H - 1$ **do**
 - 7: $\phi_\Gamma = \phi_\Gamma + \gamma^t R_\theta(\mathbf{z}_t, \mathbf{a}_t)$ $\triangleleft$ *Reward*
 - 8: $\mathbf{z}_{t+1} \leftarrow d_\theta(\mathbf{z}_t, \mathbf{a}_t)$ $\triangleleft$ *Latent transition*
 - 9: $\phi_\Gamma = \phi_\Gamma + \gamma^H Q_\theta(\mathbf{z}_H, \mathbf{a}_H)$ $\triangleleft$ *Terminal value*
 // Update parameters μ, σ for next iteration:
 - 10: $\mu^j, \sigma^j = \text{Equation 4 (and Equation 5)}$
 - 11: **return** $\mathbf{a} \sim \mathcal{N}(\mu^J, (\sigma^J)^2 \mathbf{I})$
-

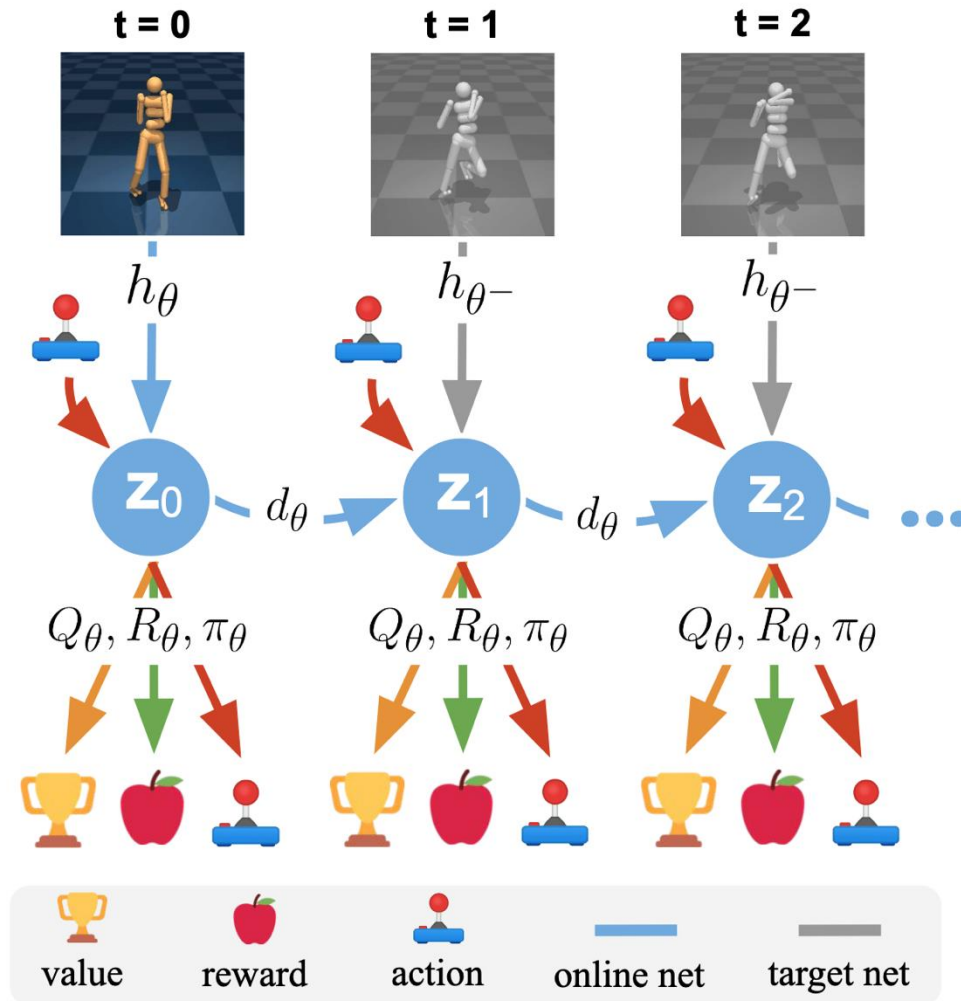
Model Predictive Path Integral Control for Planning

- Update the distribution of trajectories towards high-value samples:

$$\mu^j = \frac{\sum_{i=1}^k \Omega_i \Gamma_i^*}{\sum_{i=1}^k \Omega_i}, \quad \sigma^j = \sqrt{\frac{\sum_{i=1}^k \Omega_i (\Gamma_i^* - \mu^j)^2}{\sum_{i=1}^k \Omega_i}},$$

$$\Omega_i = e^{\tau(\phi_{\Gamma,i}^*)}$$

Learning Objectives



TOLD minimizes the objective

$$\mathcal{J}(\theta; \Gamma) = \sum_{i=t}^{t+H} \lambda^{i-t} \mathcal{L}(\theta; \Gamma_i), \quad (7)$$

where

$$\mathcal{L}(\theta; \Gamma_i) = c_1 \underbrace{\|R_\theta(\mathbf{z}_i, \mathbf{a}_i) - r_i\|_2^2}_{\text{reward}} \quad (8)$$

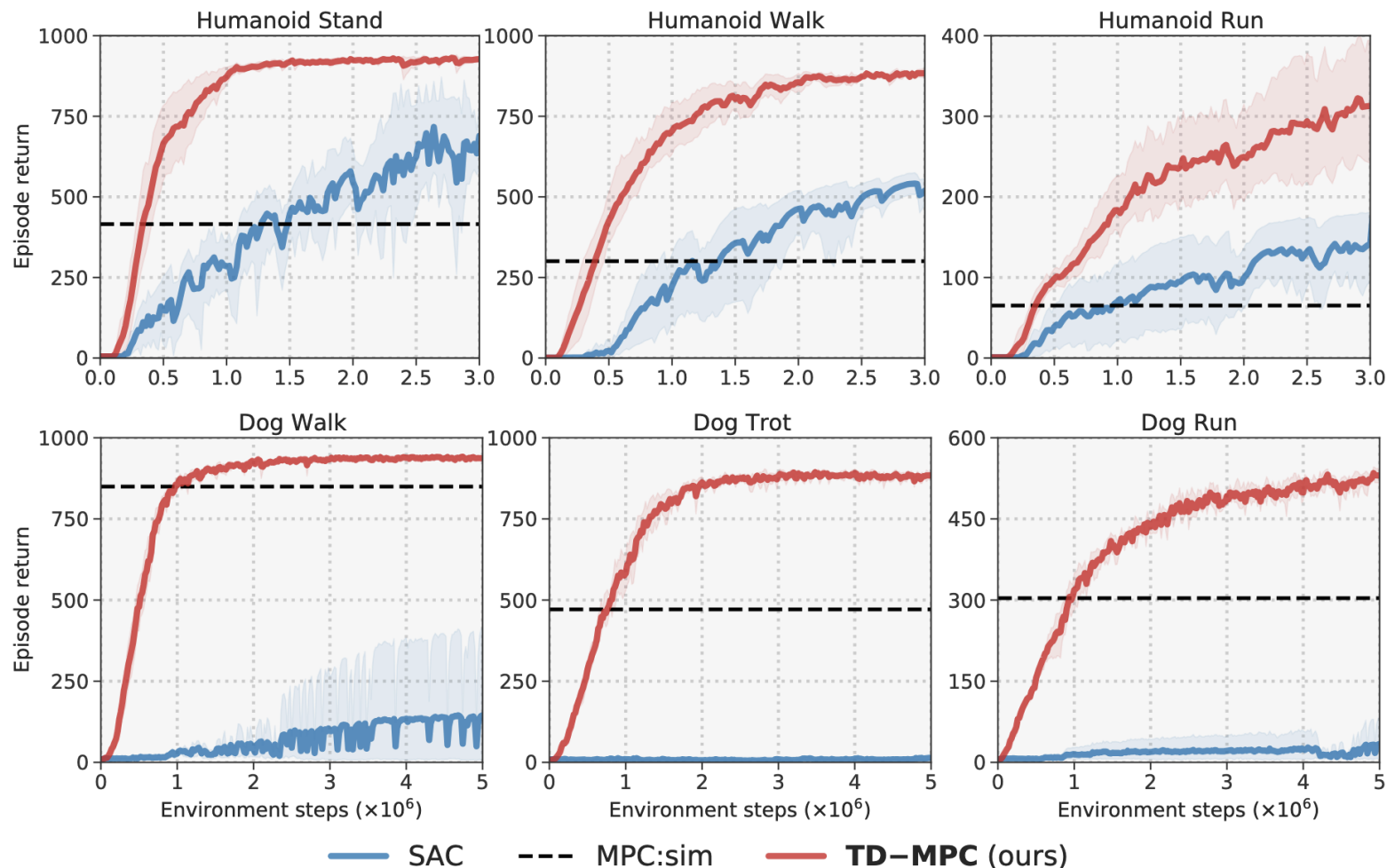
$$+ c_2 \underbrace{\|Q_\theta(\mathbf{z}_i, \mathbf{a}_i) - (r_i + \gamma Q_{\theta^-}(\mathbf{z}_{i+1}, \pi_\theta(\mathbf{z}_{i+1})))\|_2^2}_{\text{value}} \quad (9)$$

$$+ c_3 \underbrace{\|d_\theta(\mathbf{z}_i, \mathbf{a}_i) - h_{\theta^-}(\mathbf{s}_{i+1})\|_2^2}_{\text{latent state consistency}} \quad (10)$$

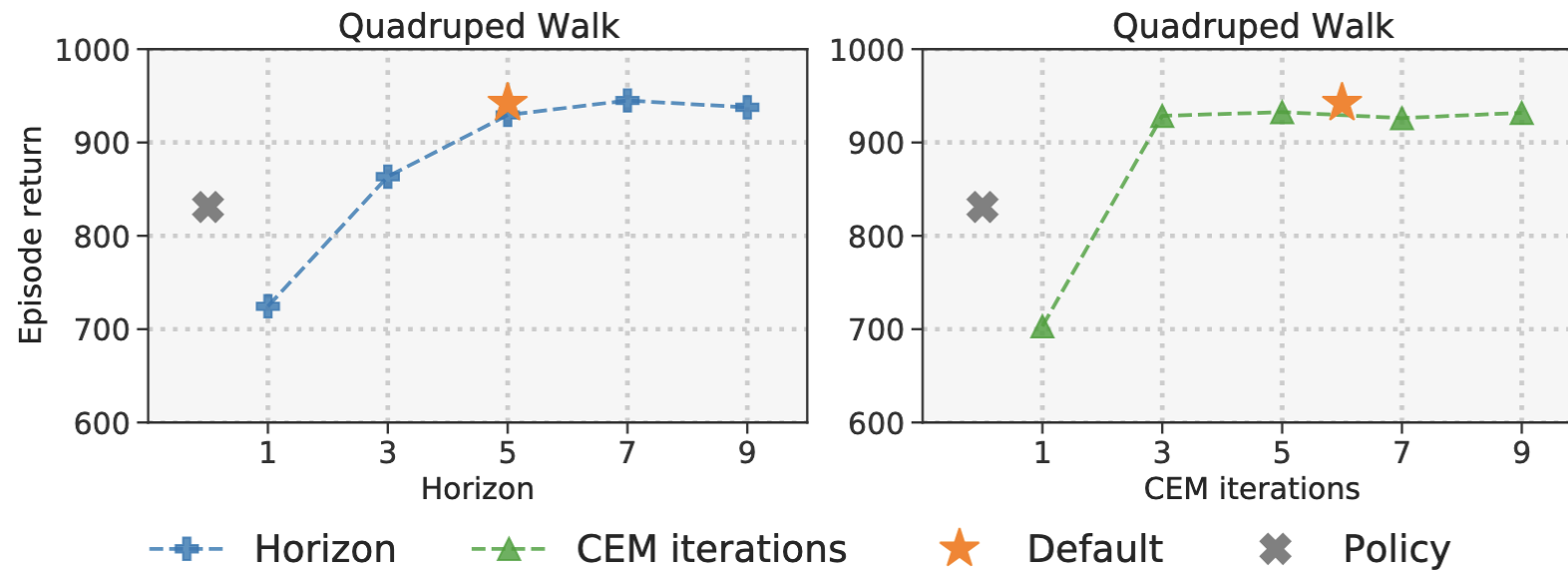
and the policy minimizes

$$\mathcal{J}_\pi(\theta; \Gamma) = - \sum_{i=t}^{t+H} \lambda^{i-t} Q_\theta(\mathbf{z}_i, \pi_\theta(\text{sg}(\mathbf{z}_i))), \quad (11)$$

TD-MPC is More Sample Efficient



Likewise, Planning Helps



A Quick Summary of Model-Based RL

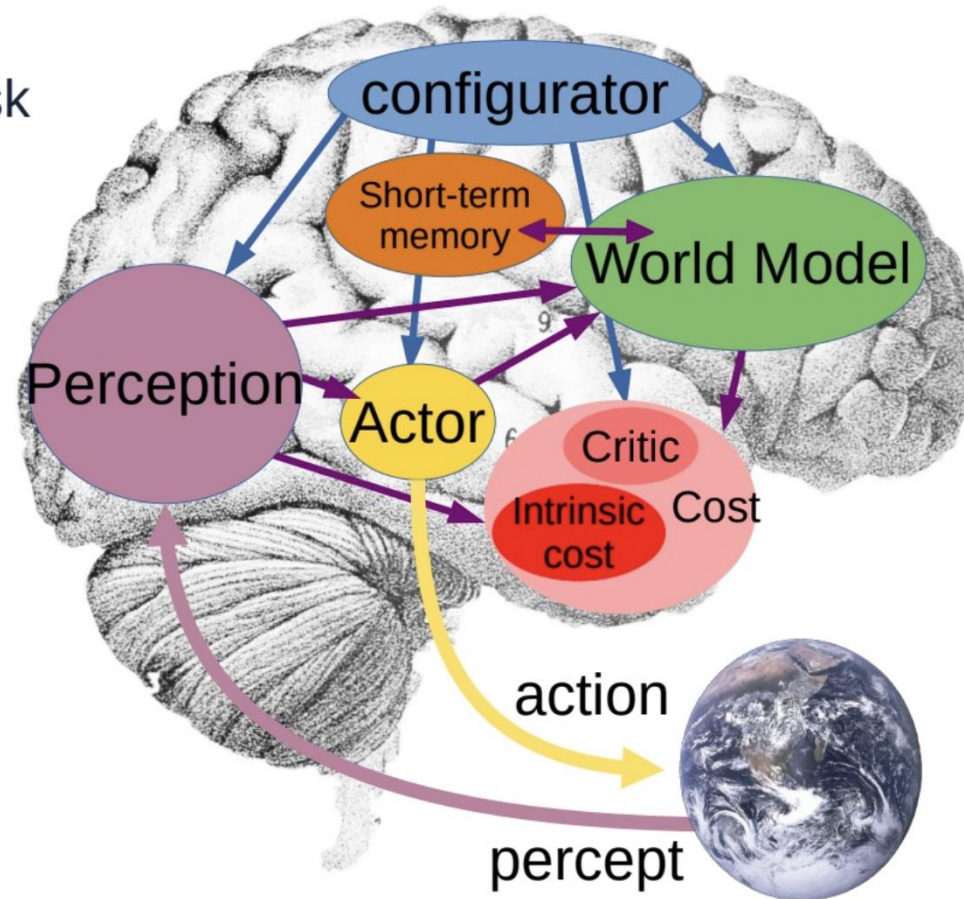
- Goodness of MBRL:
 - MBRL is more sample efficient
 - MBRL enables planning in simulation
 - MBRL allows evaluation of an action
 - MBRL re-uses interaction episodes, even those with low rewards
 - With careful architectural design, MBRL can adapt to non-stationary dynamics
- Challenges of MBRL:
 - Compounding errors of long-horizon rollouts
 - Inaccurate modeling of dynamics due to domain shift, overfitting, underfitting...
 - What is the right state representations for the dynamics models?

Learning the Dynamics Models Enable Us to Predict the Outcome of Our Actions

Y. LeCun

Modular Architecture for Autonomous AI

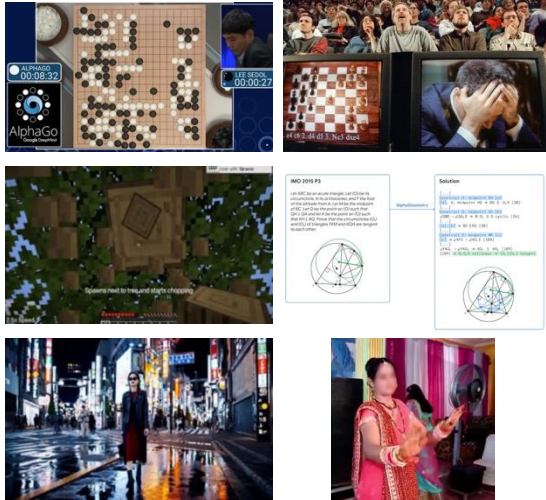
- ▶ **Configurator**
 - ▶ Configures other modules for task
- ▶ **Perception**
 - ▶ Estimates state of the world
- ▶ **World Model**
 - ▶ Predicts future world states
- ▶ **Cost**
 - ▶ Compute “discomfort”
- ▶ **Actor**
 - ▶ Find optimal action sequences
- ▶ **Short-Term Memory**
 - ▶ Stores state-cost episodes



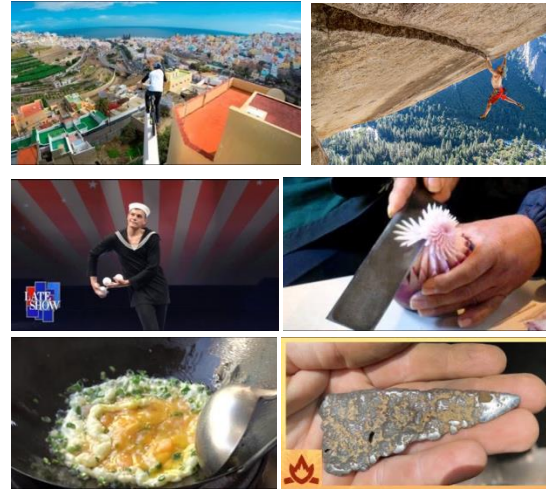
Remember Moravec's Paradox?

What makes the difference?
The key is that we are in a physical world

Easy universe



Hard universe



"We are all prodigious Olympians in perceptual and motor areas, so good that we make the difficult look easy. Abstract thought, though, is a new trick, perhaps less than 100 thousand years old. We have not yet mastered it. It is not all that intrinsically difficult; it just seems so when we do it"

Hans Moravec

"The main lessons of thirty-five years of AI research is that the hard problems are easy and the easy problems are hard. The mental abilities of a four-year-old that we take for granted-recognizing a face, lifting a pencil, walking across a room, answering a question-in fact solve some of the hardest engineering problems ever conceived."

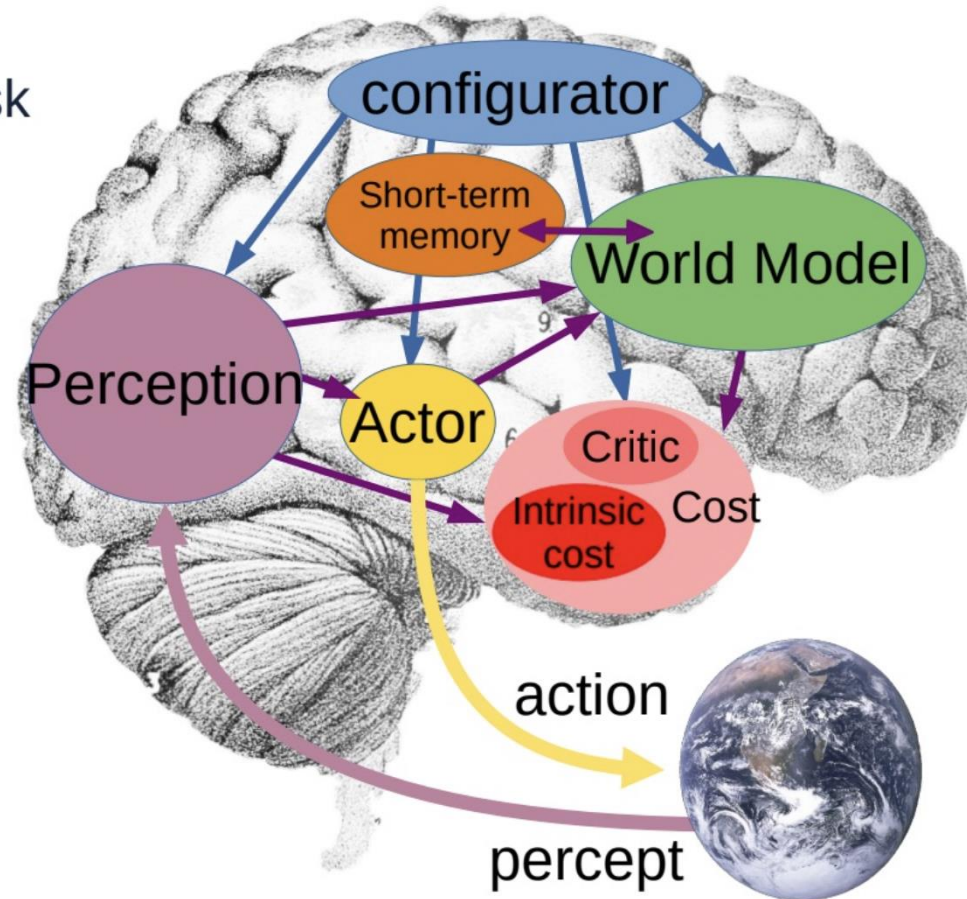
Steven Pinker

Is world model the next frontier in our path to AGI?

Y. LeCun

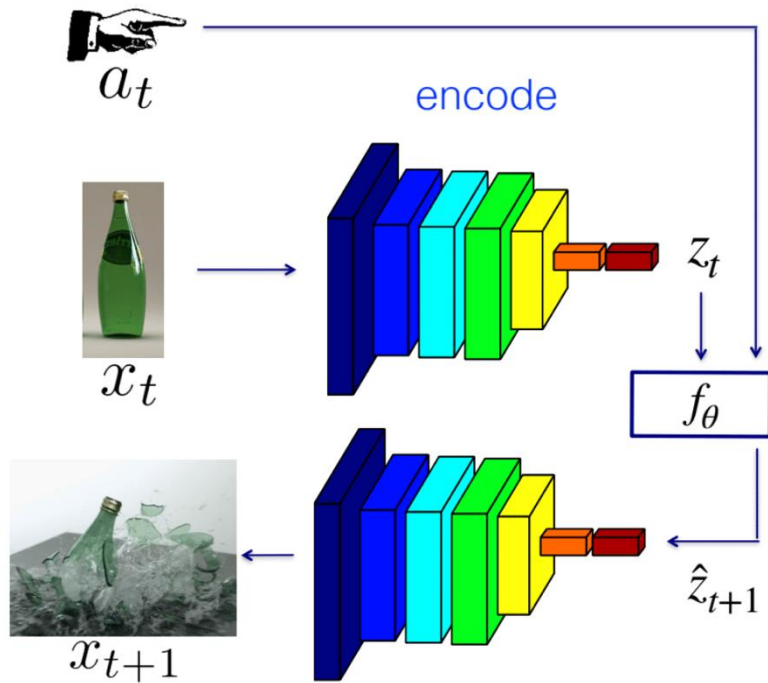
Modular Architecture for Autonomous AI

- ▶ **Configurator**
 - ▶ Configures other modules for task
- ▶ **Perception**
 - ▶ Estimates state of the world
- ▶ **World Model**
 - ▶ Predicts future world states
- ▶ **Cost**
 - ▶ Compute “discomfort”
- ▶ **Actor**
 - ▶ Find optimal action sequences
- ▶ **Short-Term Memory**
 - ▶ Stores state-cost episodes

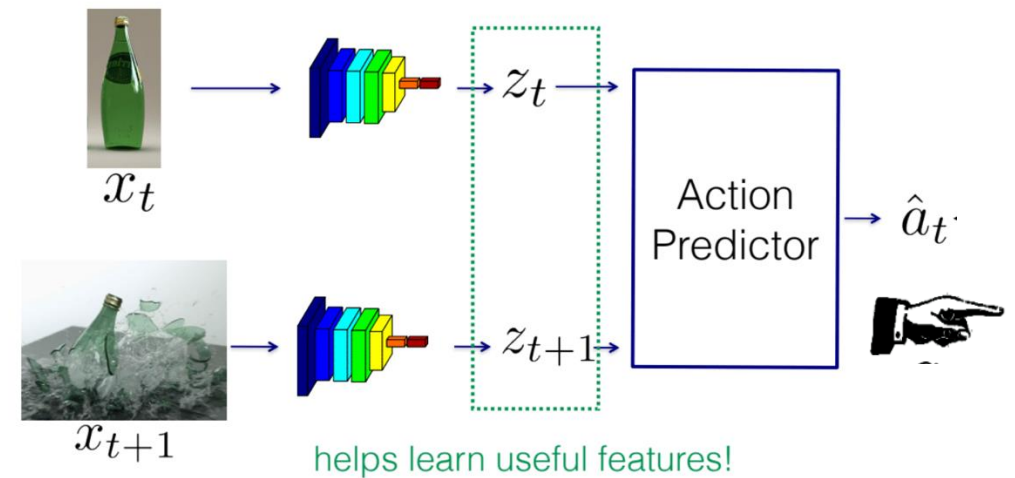


An Alternative to Model-Based RL: Sub-goal Prediction + Inverse Models

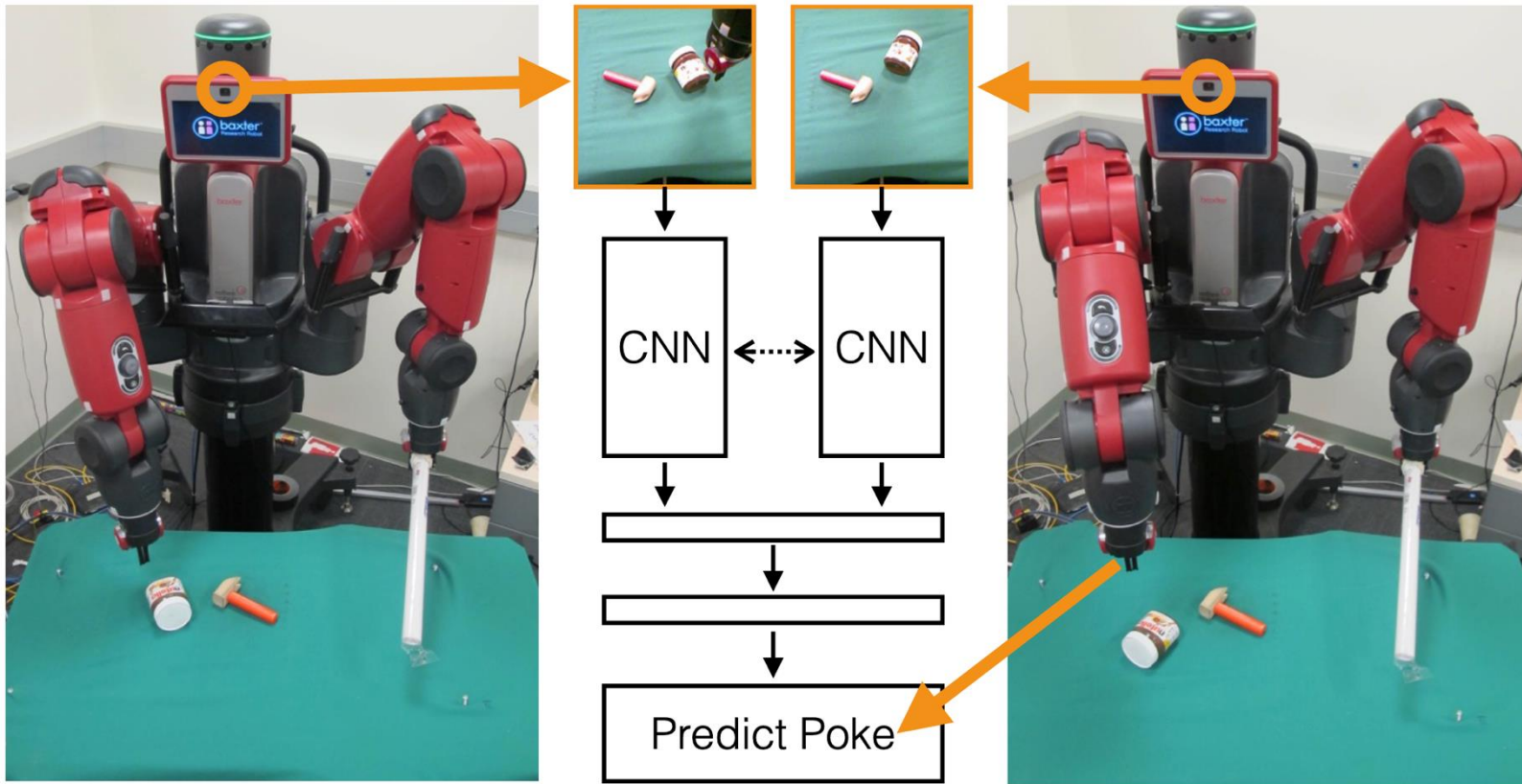
Forward Dynamic Models



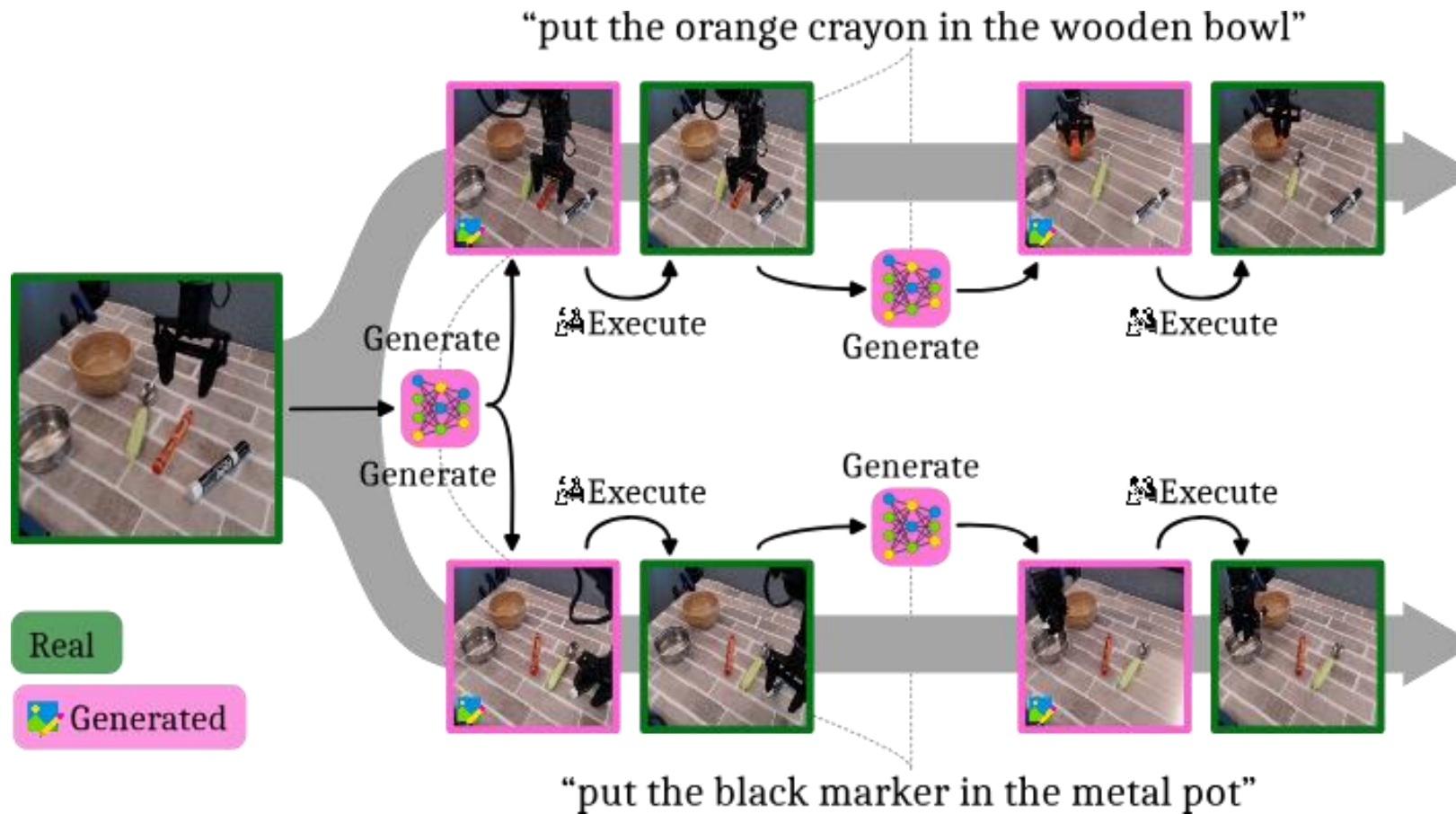
Inverse Models with Subgoals



An Alternative to Model-Based RL: Sub-goal Prediction + Inverse Models



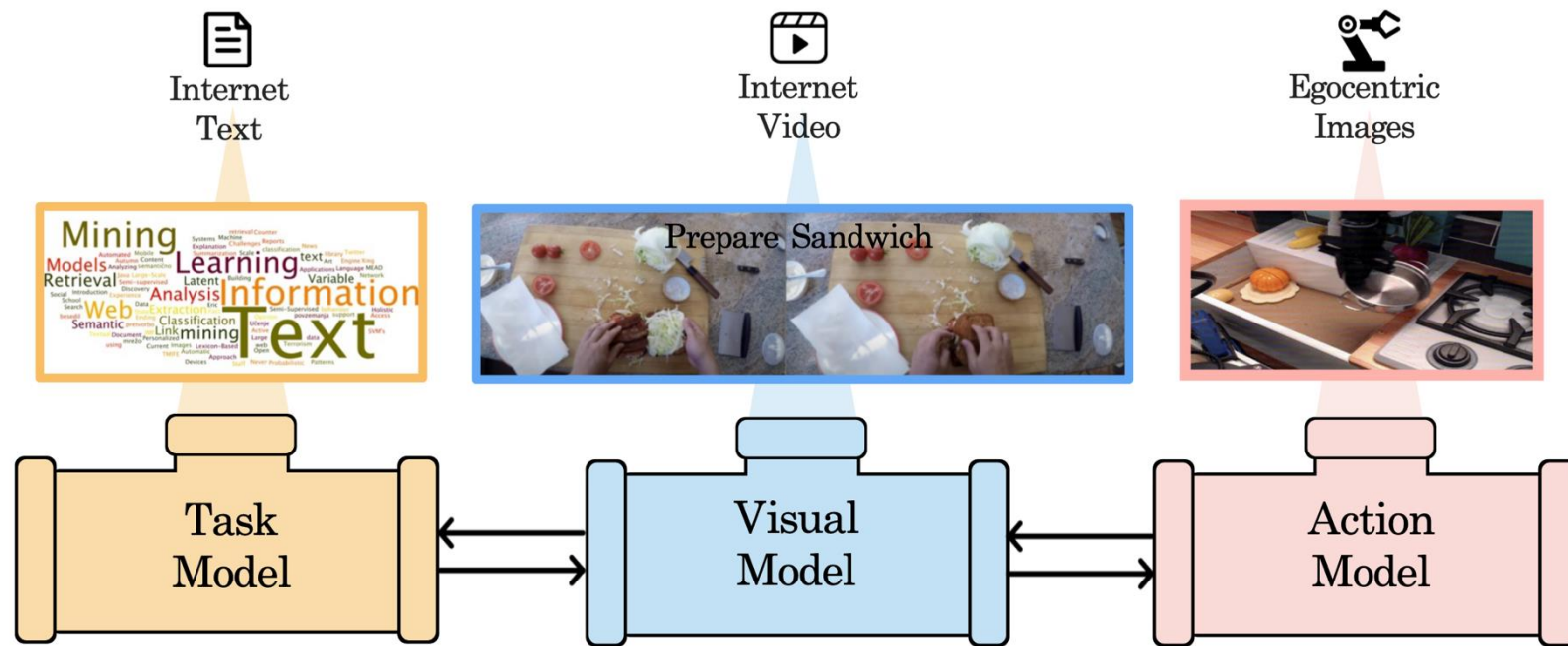
SuSIE: Predict Subgoals with an Image Generative Model, and Actions with a Goal-conditioned Policy



SuSIE: Predict Subgoals with an Image Generative Model, and Actions with a Goal-conditioned Policy



An Alternative to Model-Based RL: LLM + Sub-goal Prediction + Inverse Models



An Alternative to Model-Based RL: LLM + Sub-goal Prediction + Inverse Models

