

Offline Reinforcement Learning

Deep Reinforcement Learning

Mohammad Hossein Rohban, Ph.D.

Courtesy: Some slides are adopted from CS 285 Berkeley, and CS 234 Stanford, and Pieter Abbeel's compact series on RL.



RIML Lab, CE Department, Sharif
University of Technology

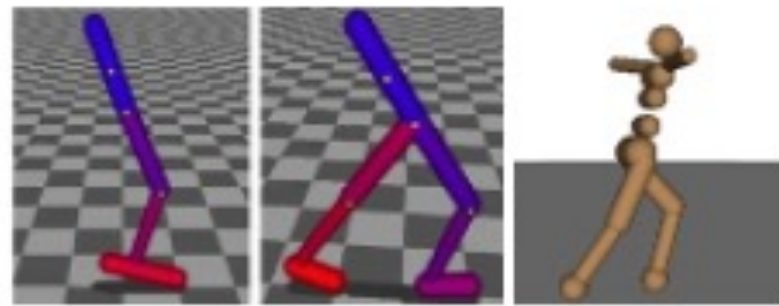
Spring 2026

The generalization gap

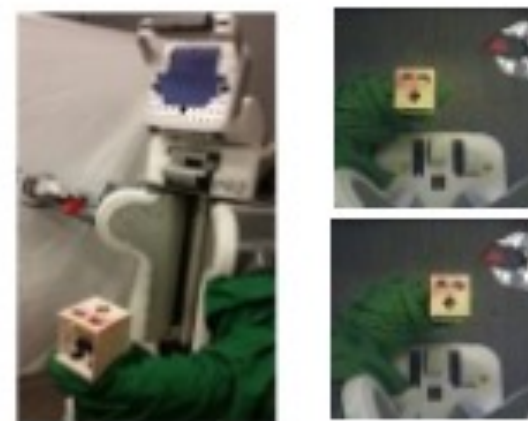
Tehran
Isfahan



Mnih et al. '13



Schulman et al. '14 & '15



Levine*, Finn*, et al. '16



Closed World

Offline RL

RL on Data

RL on Simulator

VS.

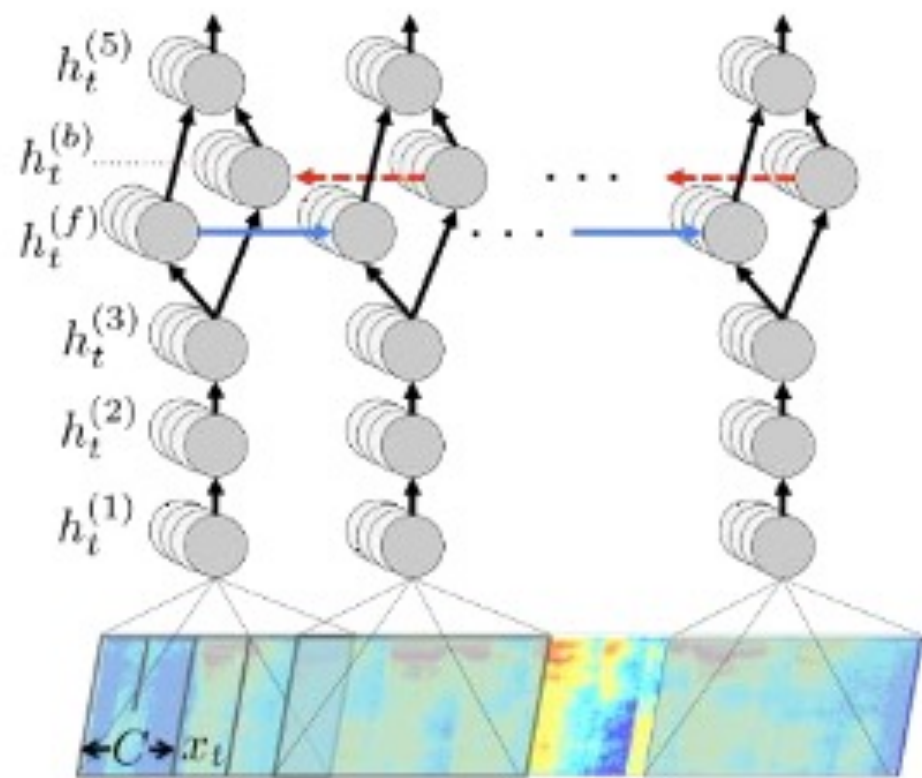
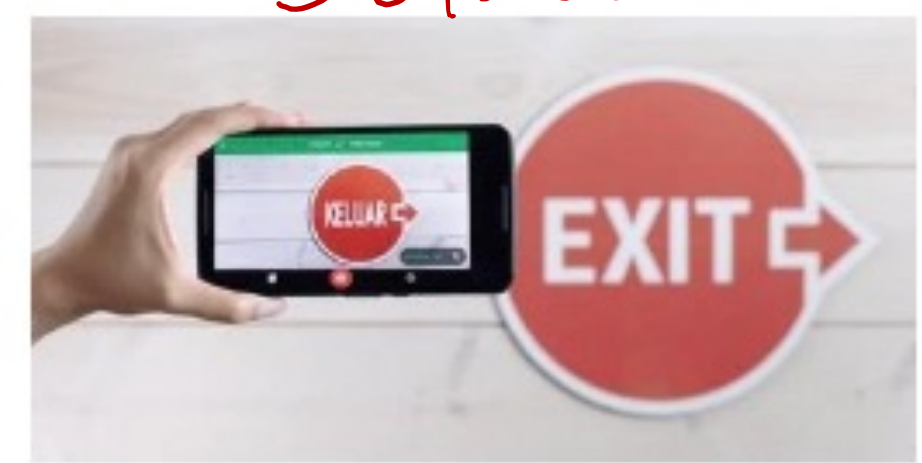
this is done many times

enormous gulf

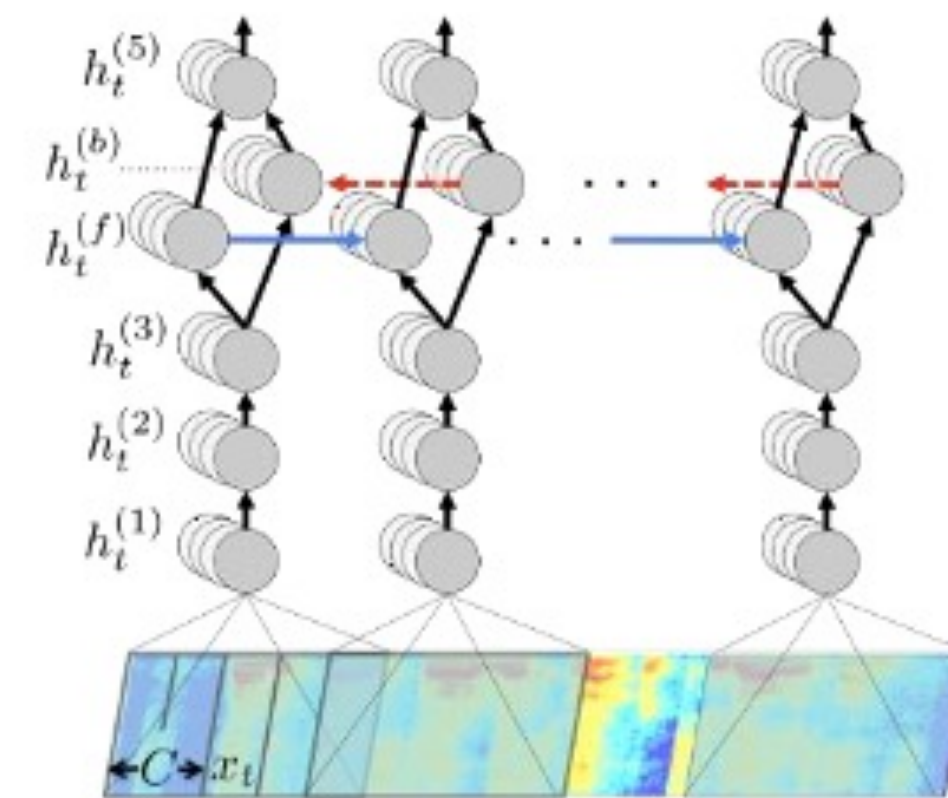
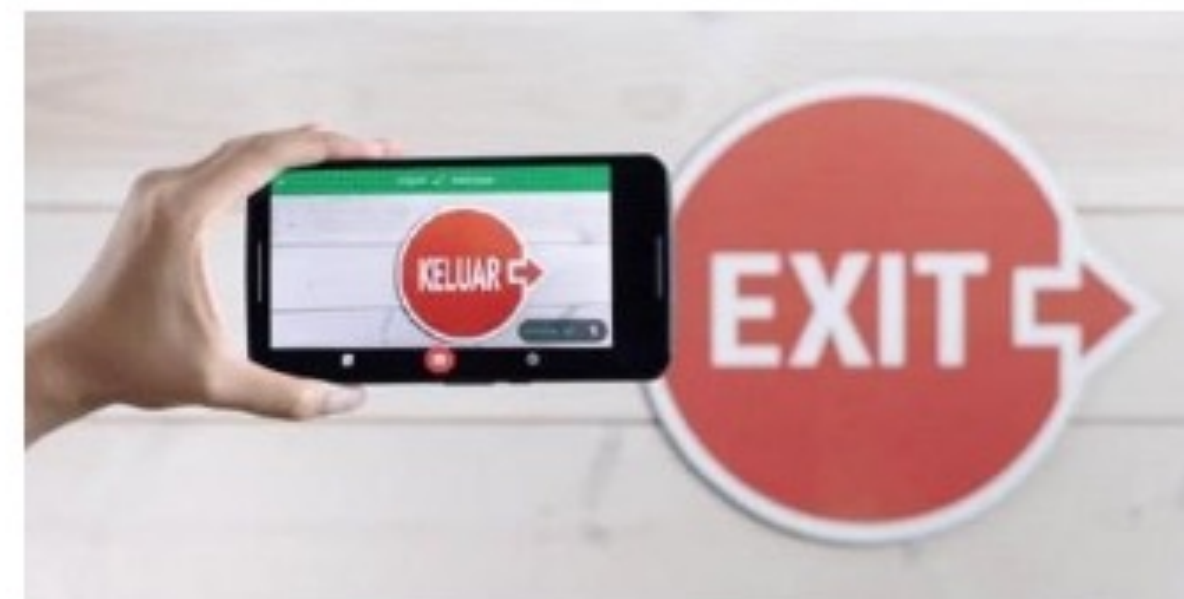
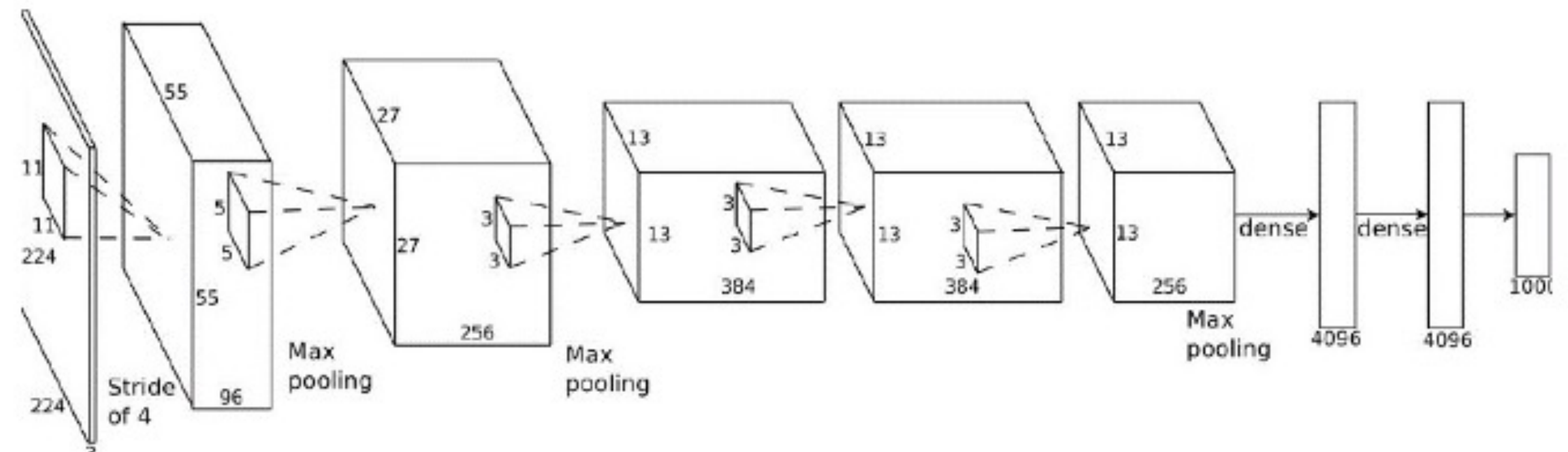
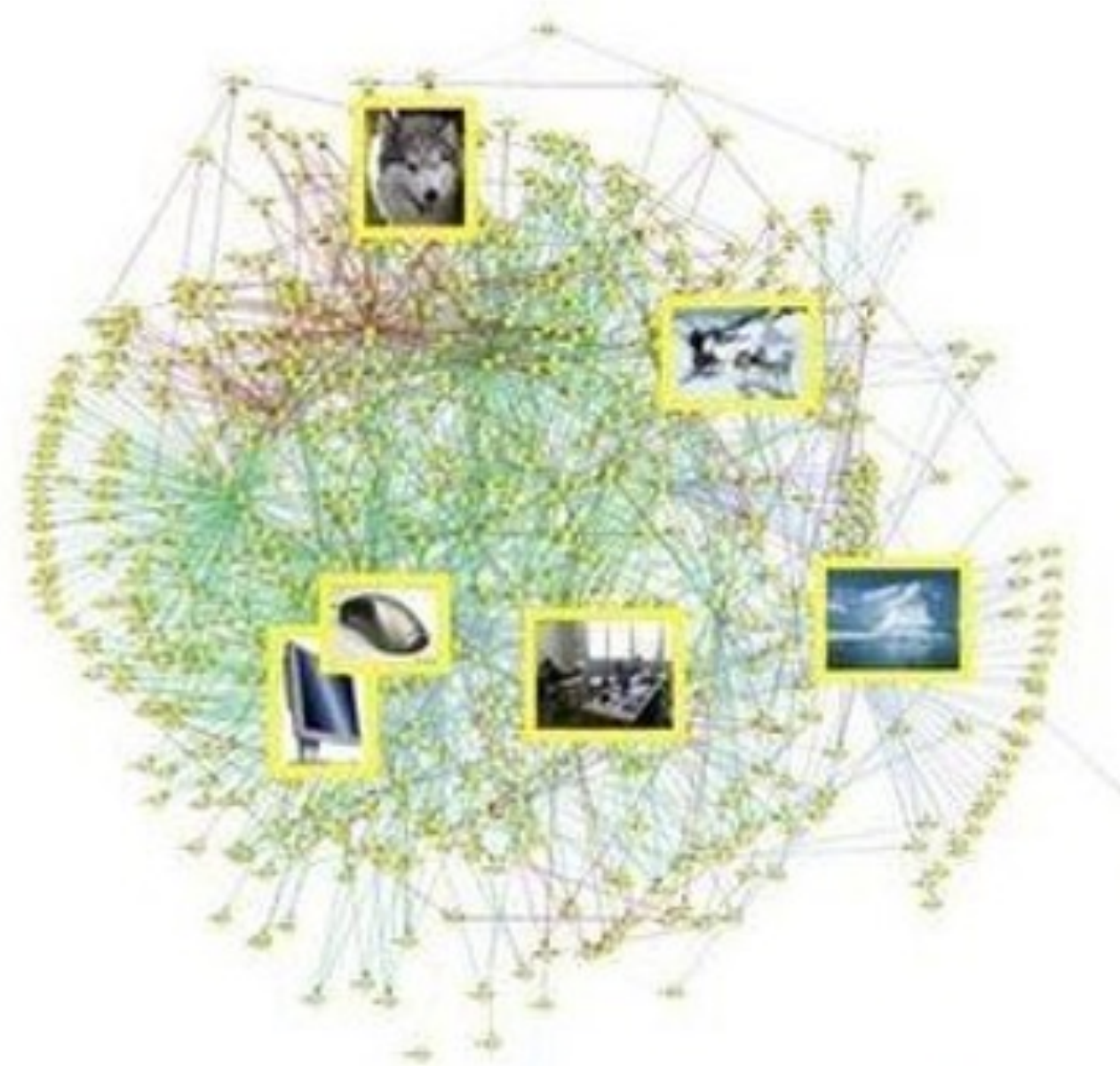
SL

Open World

Dataset

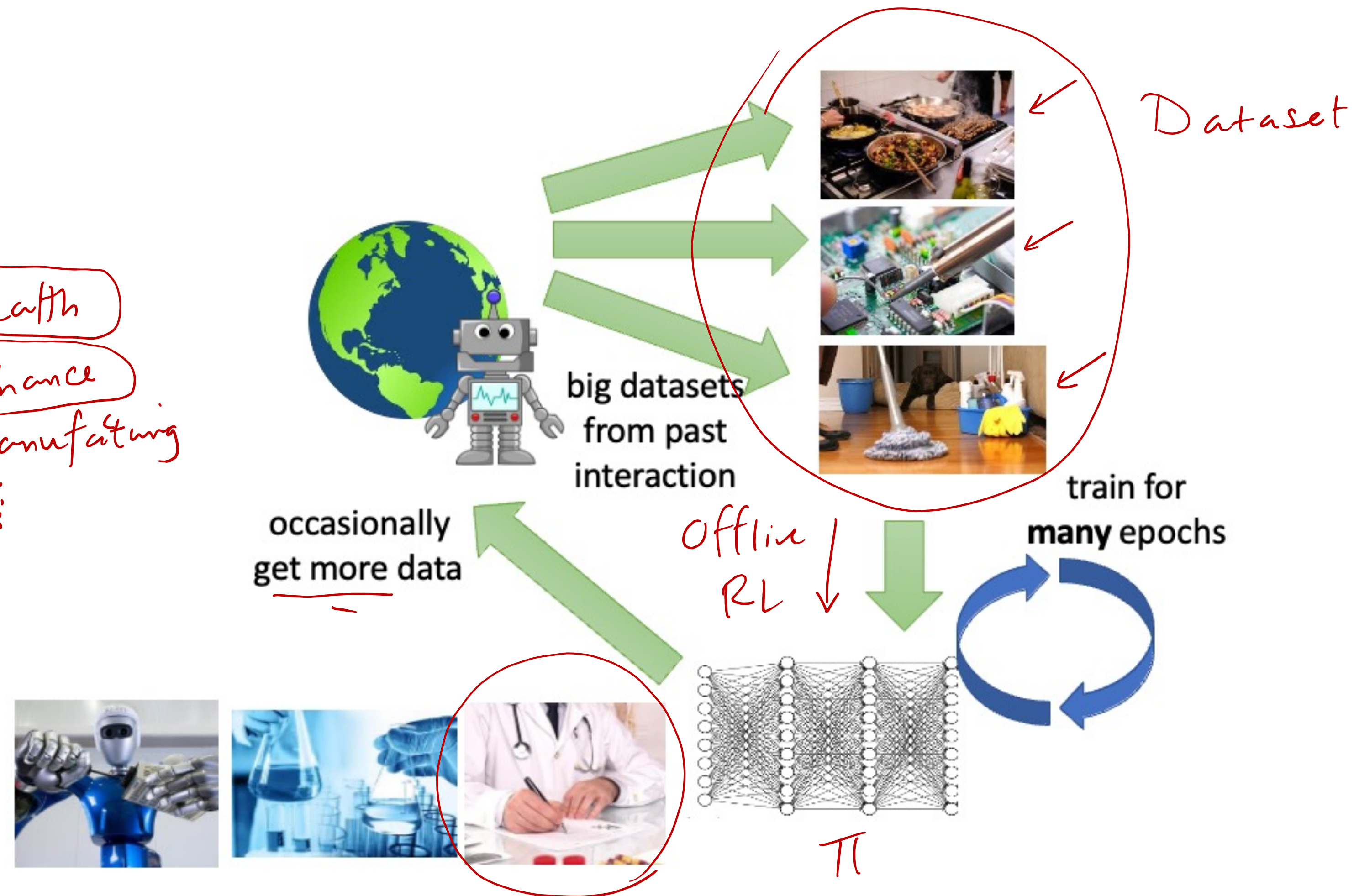


What makes modern machine learning work?



Can we develop data-driven RL methods?

- ① Scalability
- ② Safety
 - Health
 - * Finance
 - Manufacturing
 - ⋮



What does offline RL mean?

$$D = \{(s_i, a_i, s'_i, r_i)\}$$

$$s \sim d^{\pi_\beta}(s)$$

$$a \sim \pi_\beta(a|s) \text{ (generally not known)}$$

$$s' \sim p(s'|s, a)$$

$$r = r(s, a)$$

$$\max_{\pi} \sum_{t=0}^T E_{s_t \sim d^{\pi}(s), a_t \sim \pi(a|s)} [\gamma^t r(s_t, a_t)]$$

go beyond
the perf.
achieved by π_β

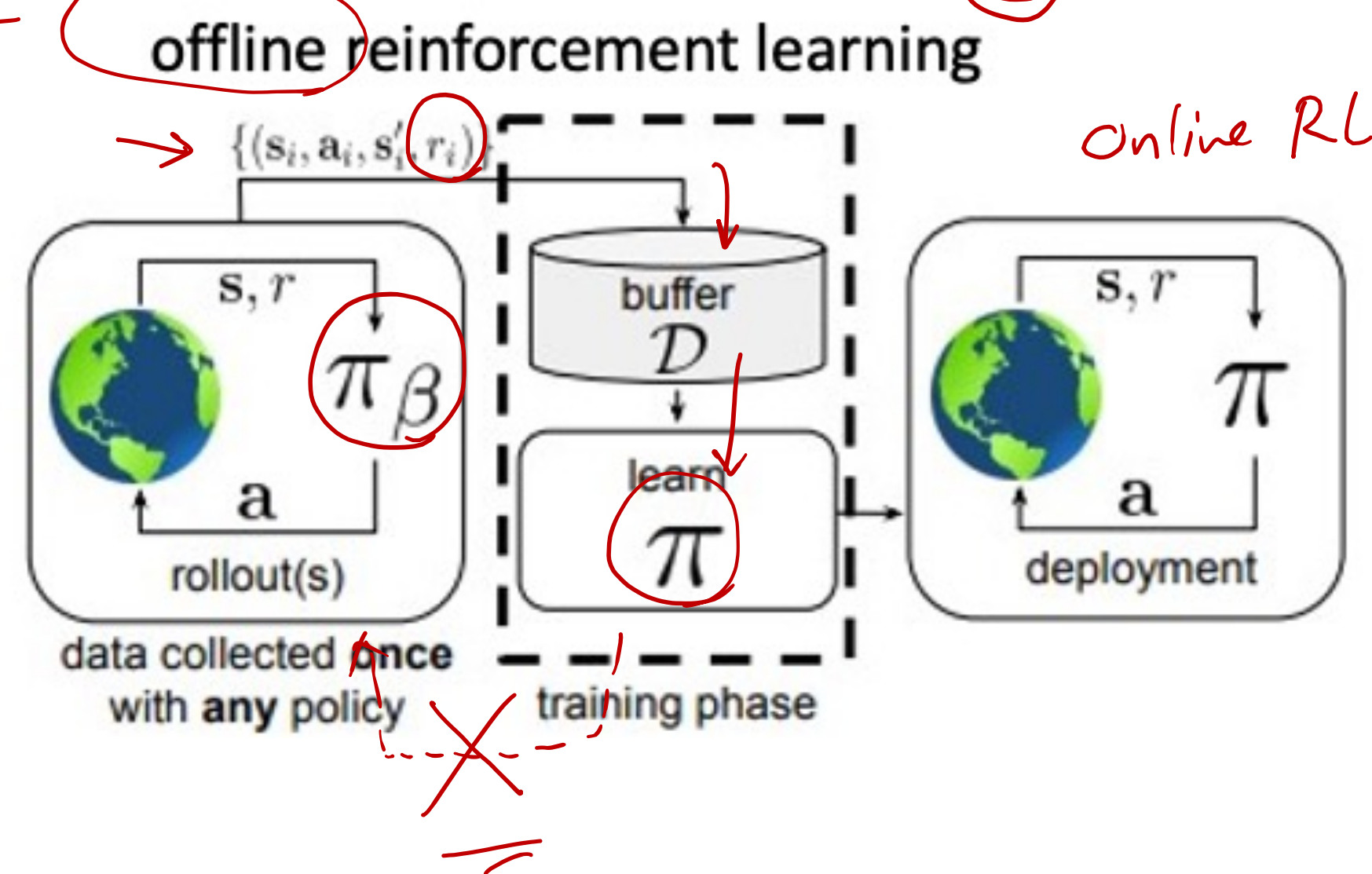
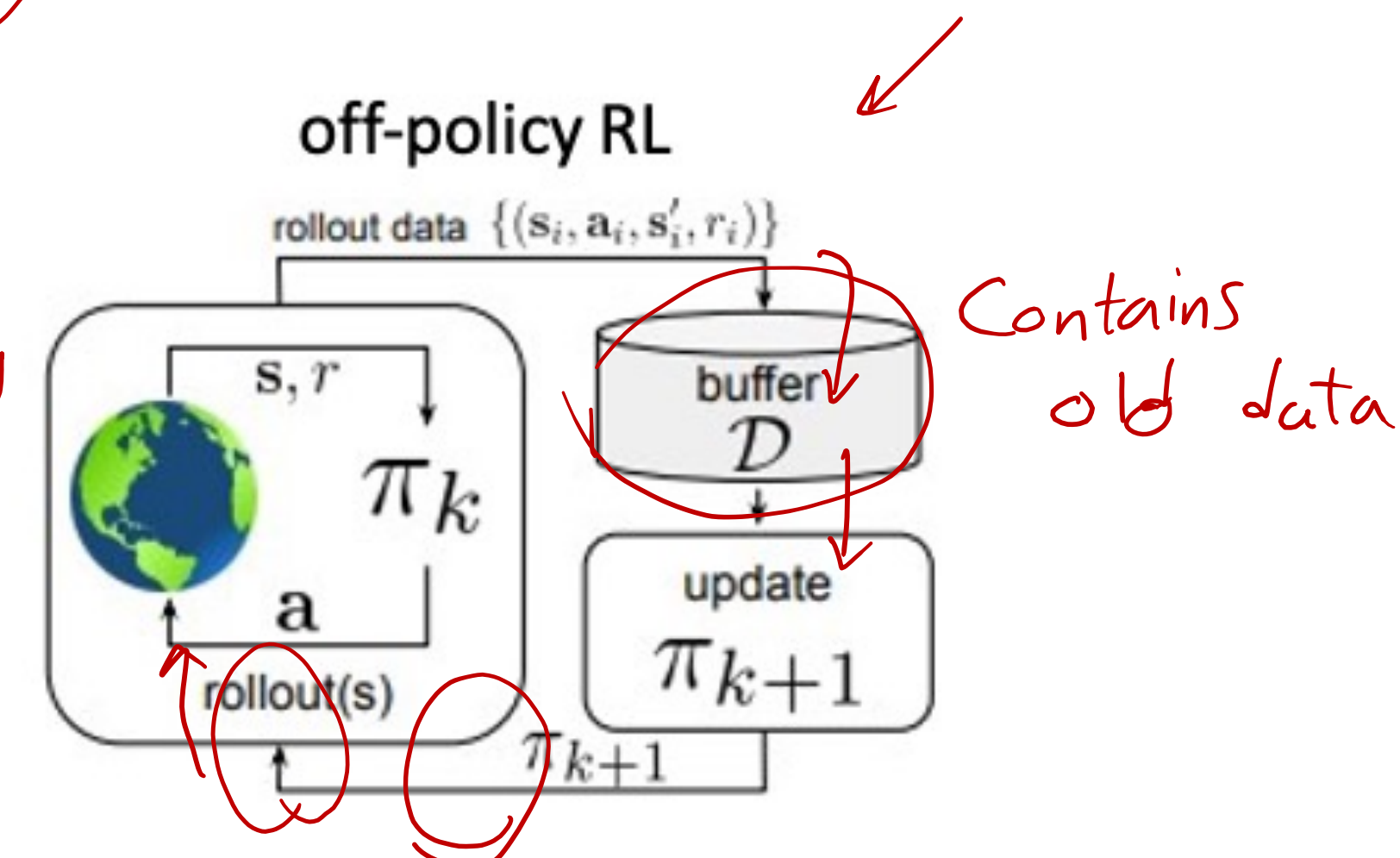
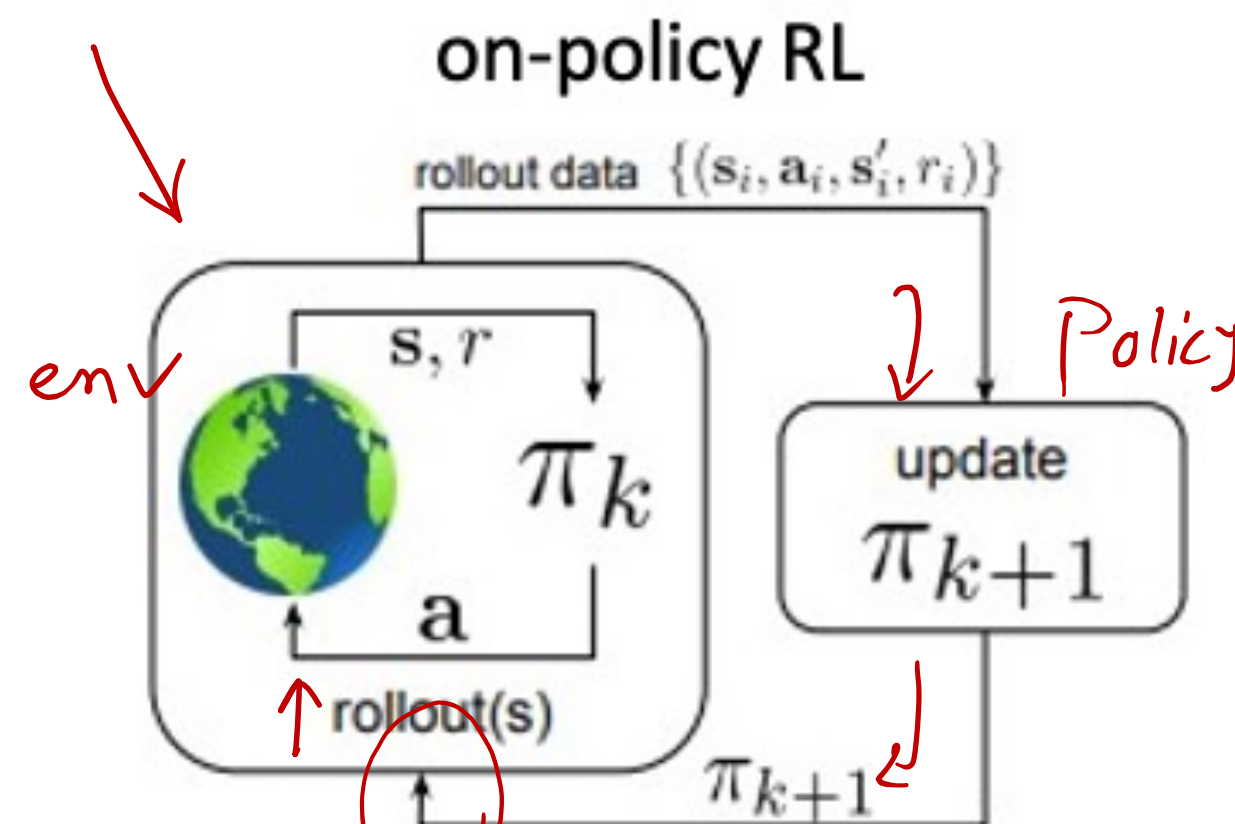
not necessarily
expert data

vs.
imitation learning
(expert data)

vs.

SL

(e.g. policy grad)



Types of offline RL problems

→ off-policy evaluation (OPE):

①

$$J(\pi) = E_{\pi} \left[\sum_{t=1}^T r(s_t, a_t) \right]$$

given $\mathcal{D}$, ^{$\sim \pi_{\beta}$} estimate $J(\pi)$

$$\pi_{\beta} \neq \pi$$

$$\mathcal{D} = \{(s_i, a_i, s'_i, r_i)\}$$

$$s \sim d^{\pi_{\beta}}(s)$$

$$a \sim \pi_{\beta}(a|s)$$

$$s' \sim p(s'|s, a)$$

$$r \leftarrow r(s, a)$$

②

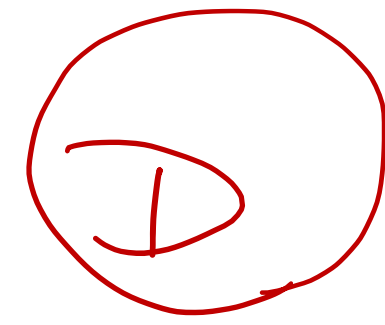
offline reinforcement learning: (a.k.a. batch RL, sometimes fully off-policy RL)

given $\mathcal{D}$, learn the best possible policy π_{θ}

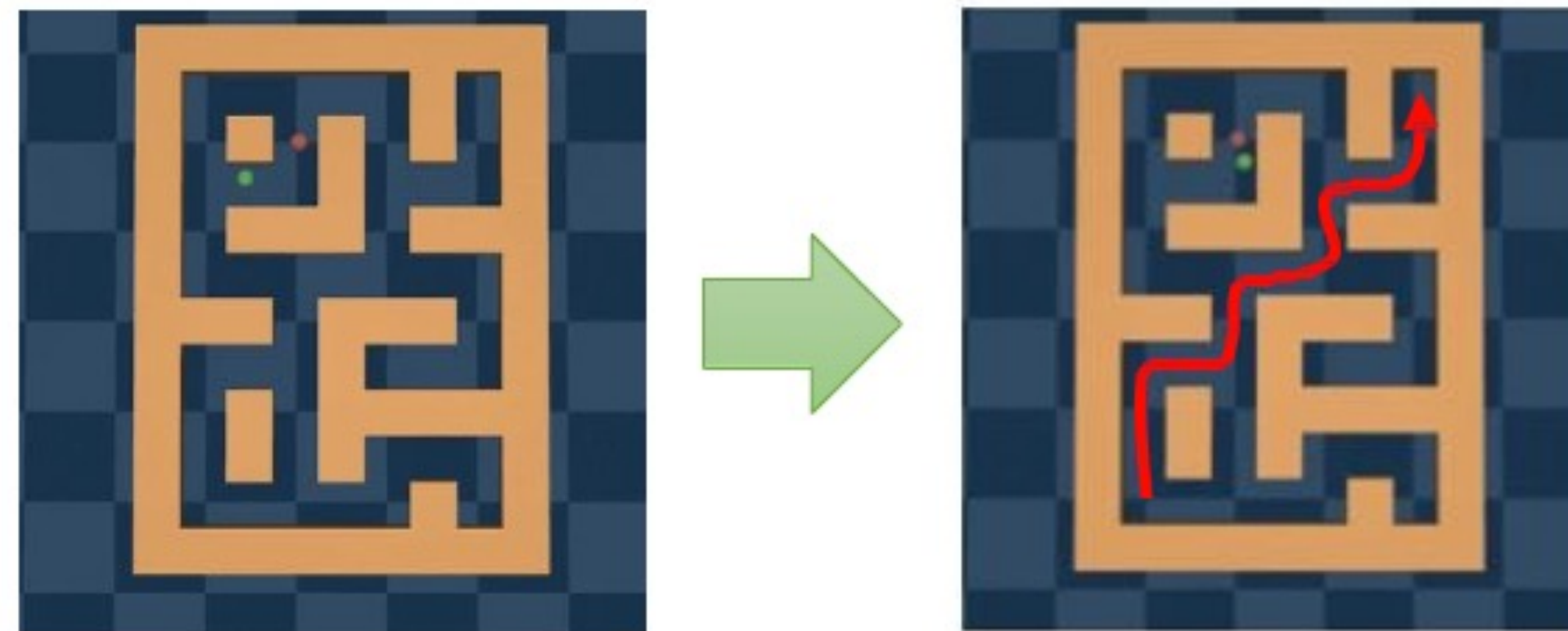
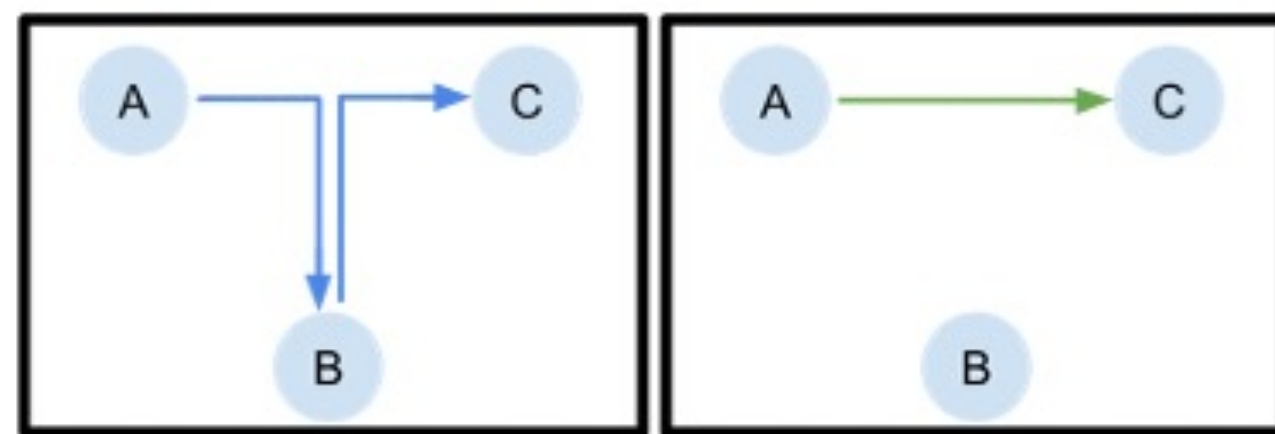
$\max_{\pi} J(\pi)$
 π based on $\mathcal{D}$.

not necessarily obvious what this means

How is this even possible?



1. Find the “good stuff” in a dataset full of good and bad behaviors
2. Generalization: good behavior in one place may suggest good behavior in another place
a.k.a state
3. “Stitching”: parts of good behaviors can be recombined



What do we expect offline RL methods to do?

- ♦ Bad intuition: it's like imitation learning

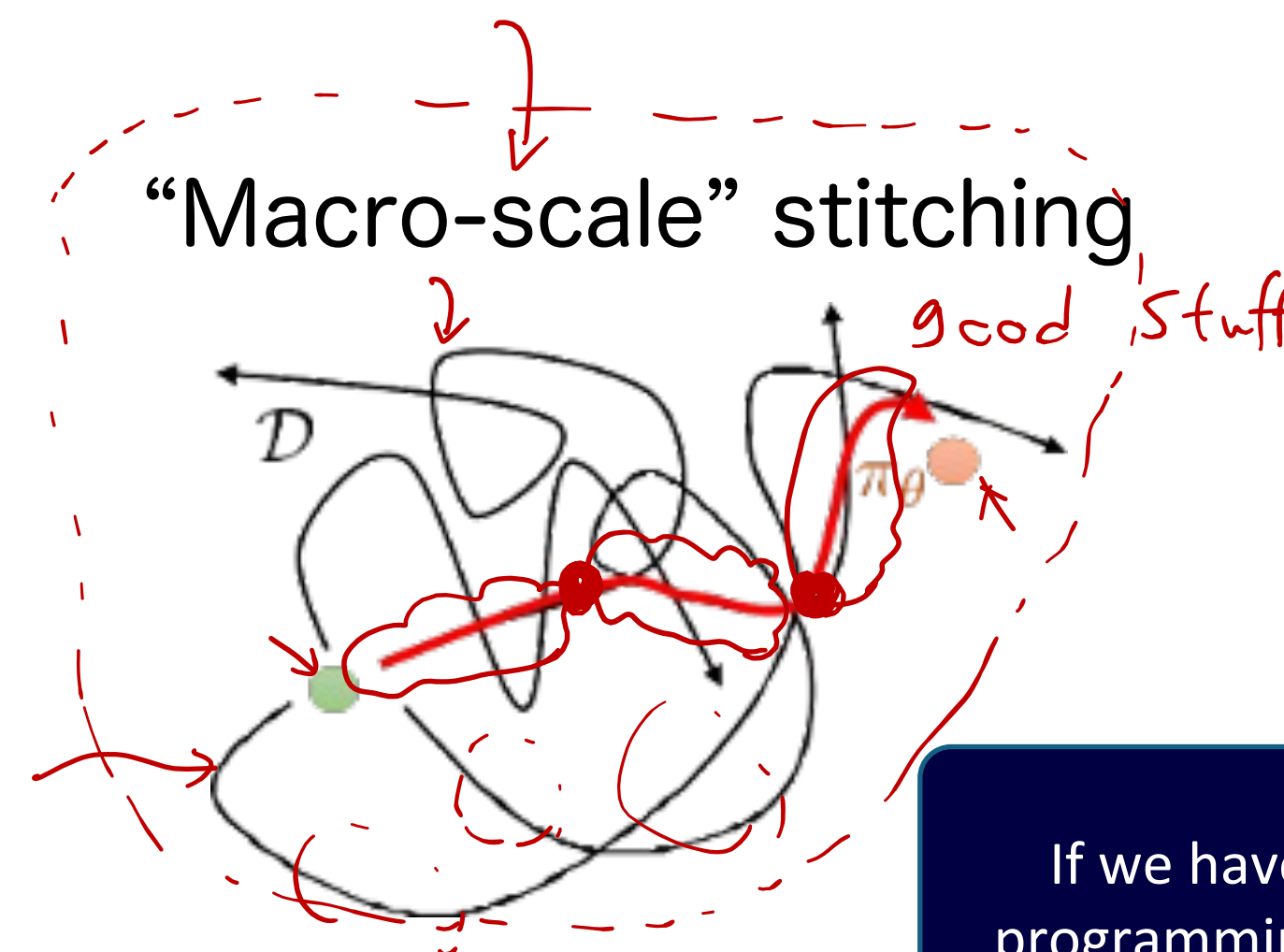
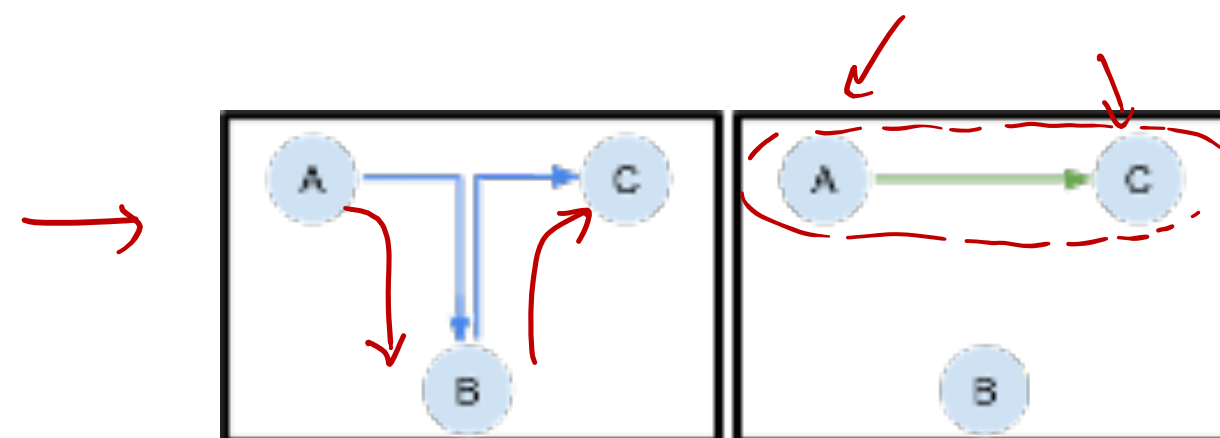
Though it can be shown to be provably better than imitation learning even with optimal data, under some structural assumptions!

See: Kumar, Hong, Singh, Levine. Should I Run Offline Reinforcement Learning or Behavioral Cloning?

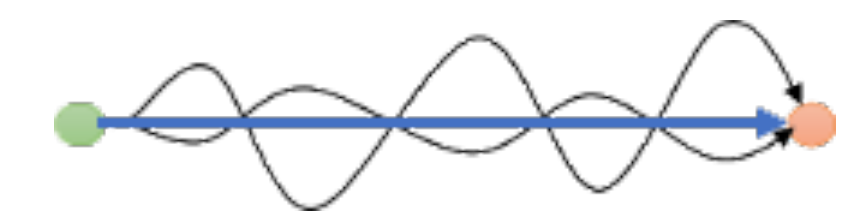


- ♦ Better intuition: get order from chaos

But this is just the clearest example!

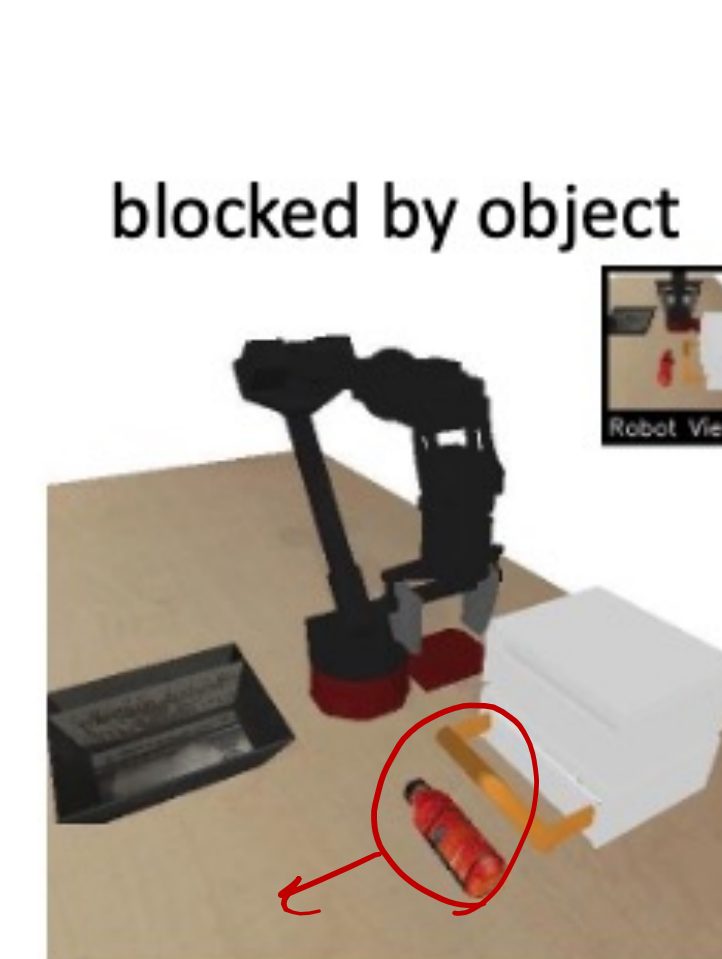
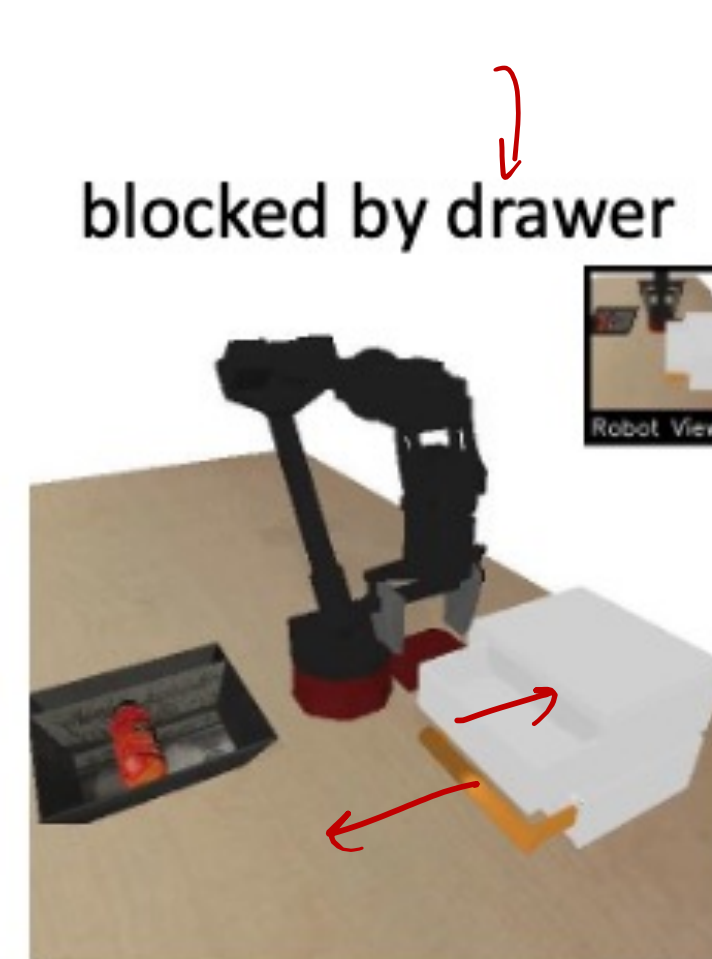
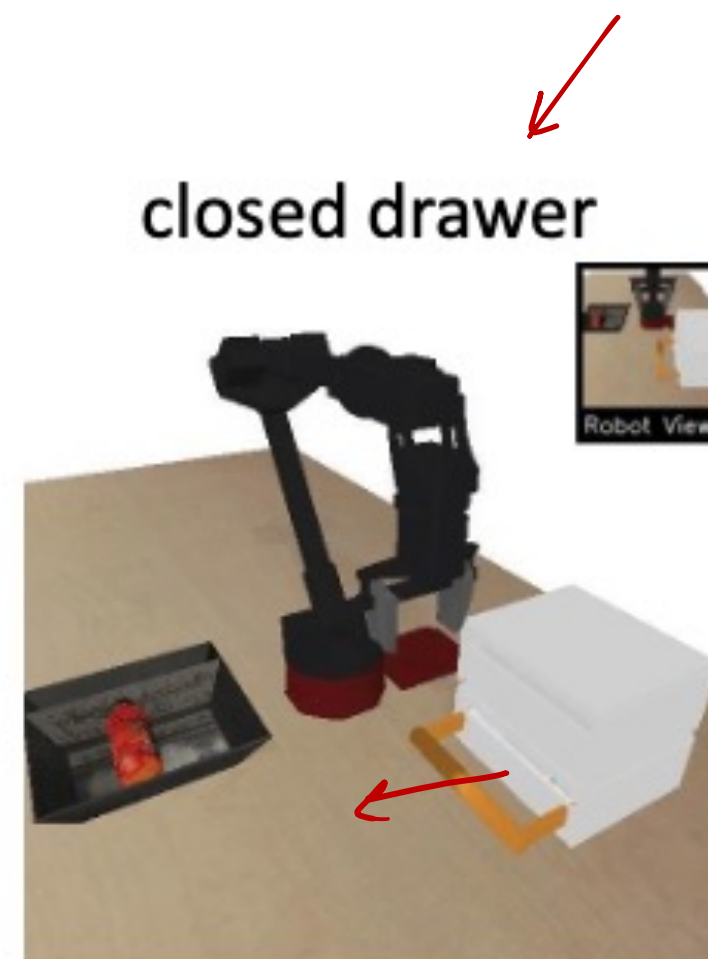
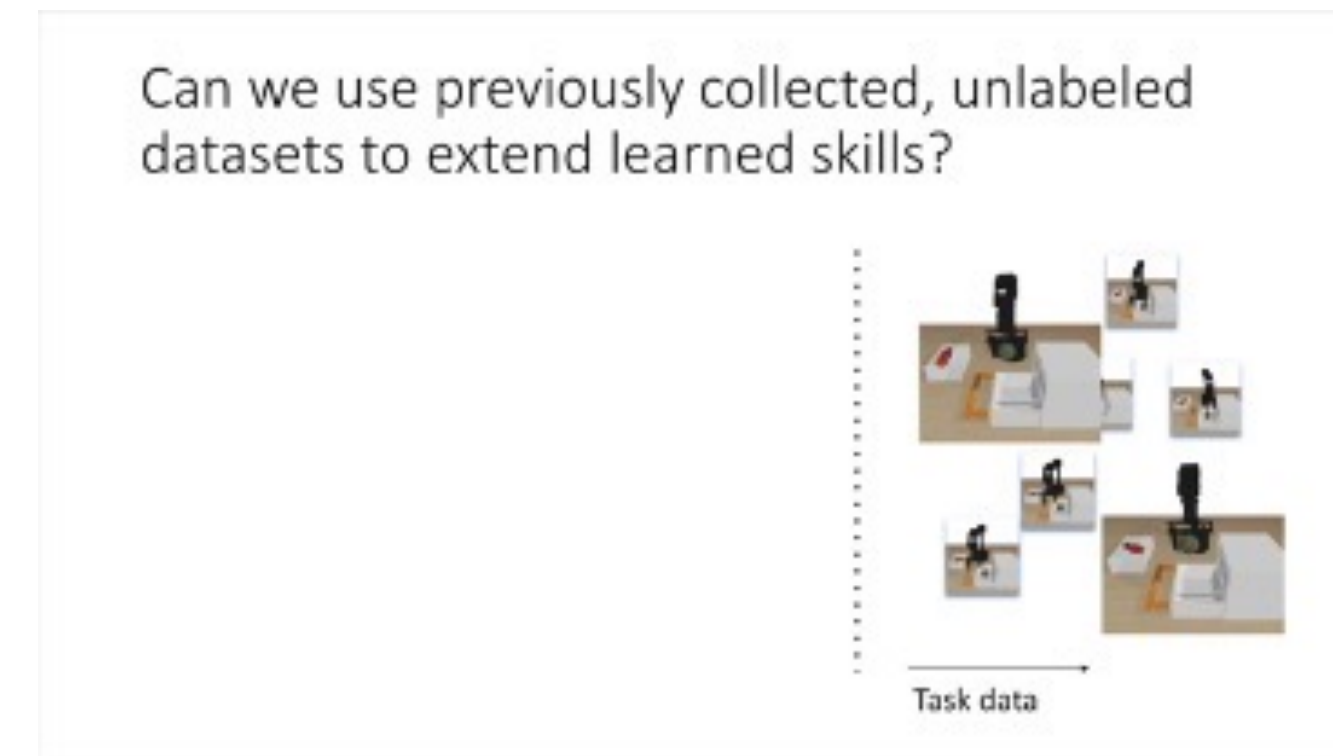
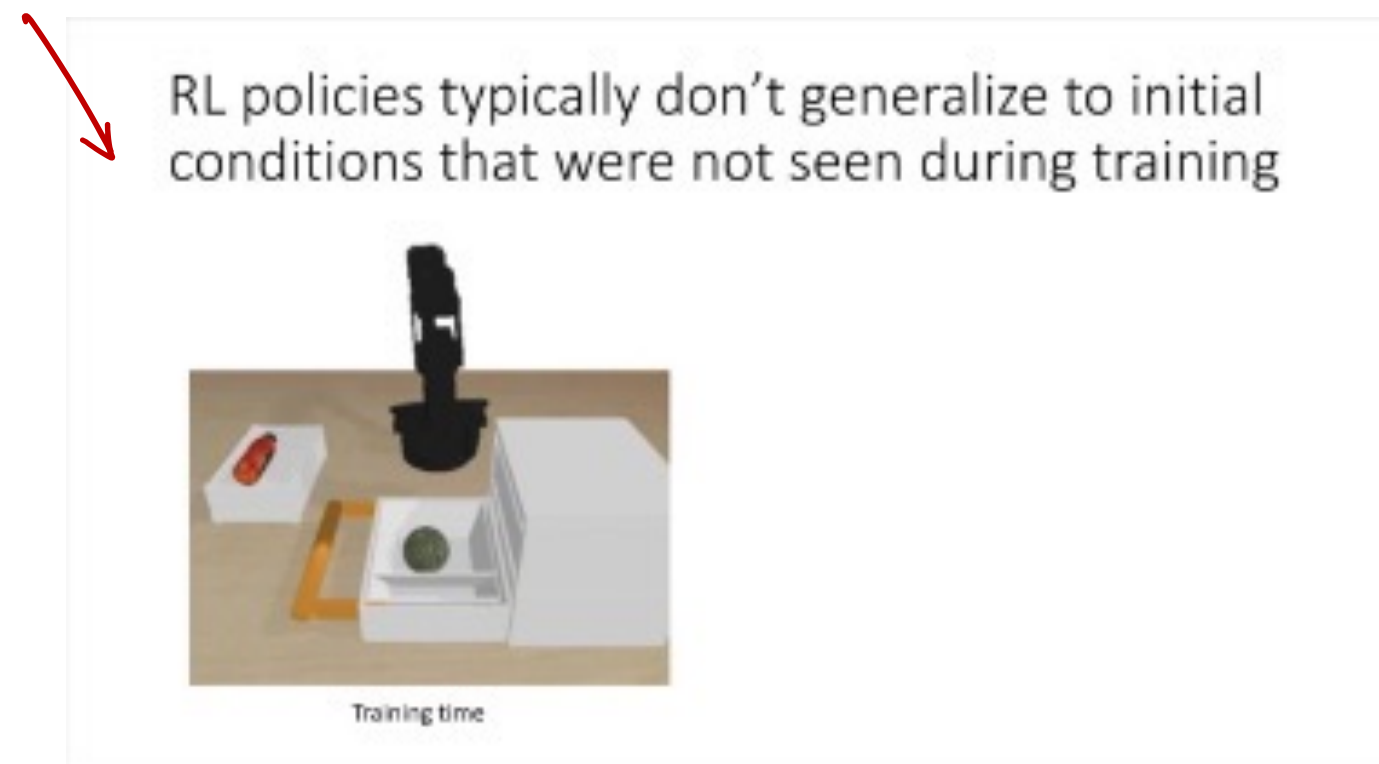


“Micro-scale” stitching:

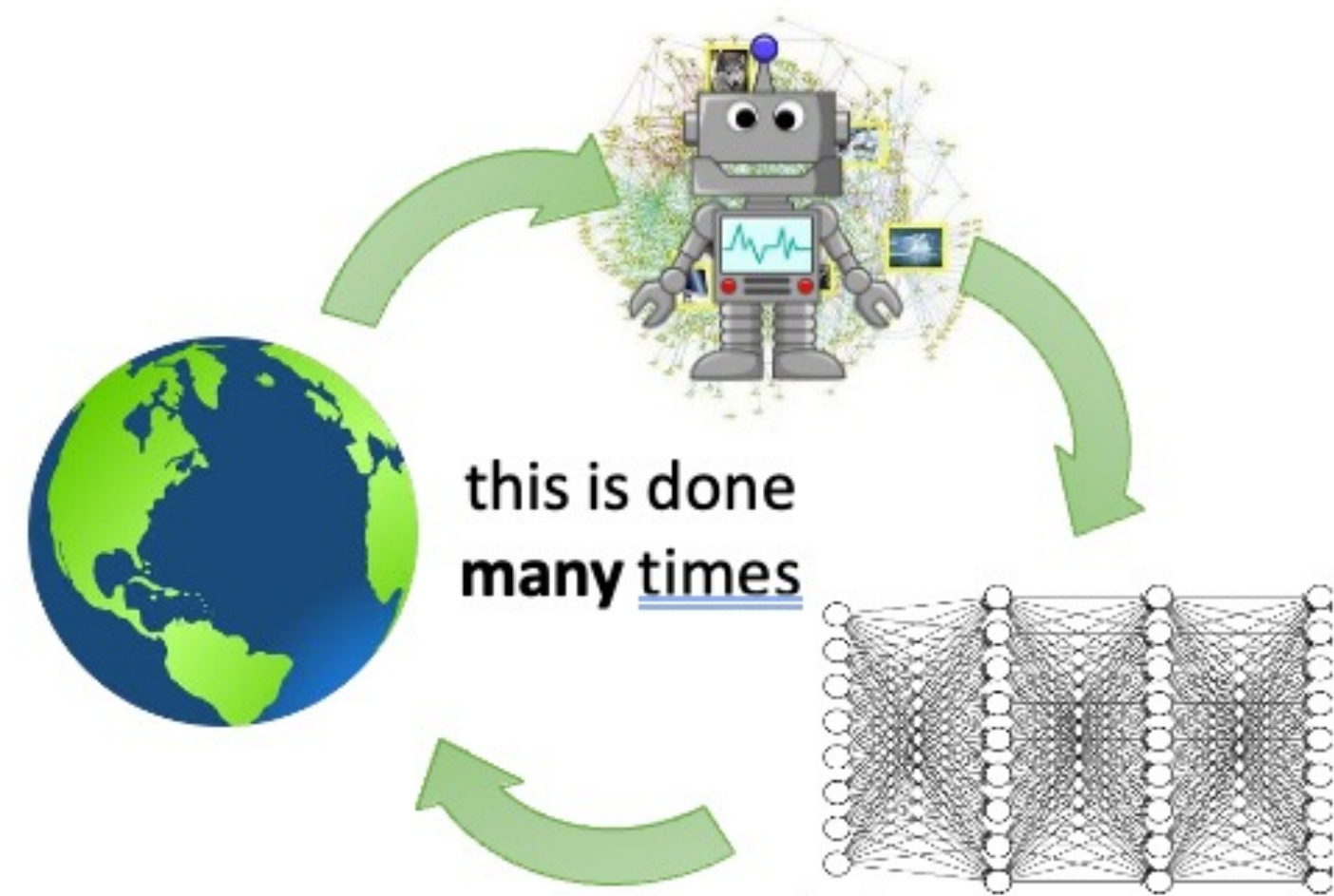


If we have algorithms that properly perform dynamic programming, we can take this idea much further and get near-optimal policies from highly suboptimal data

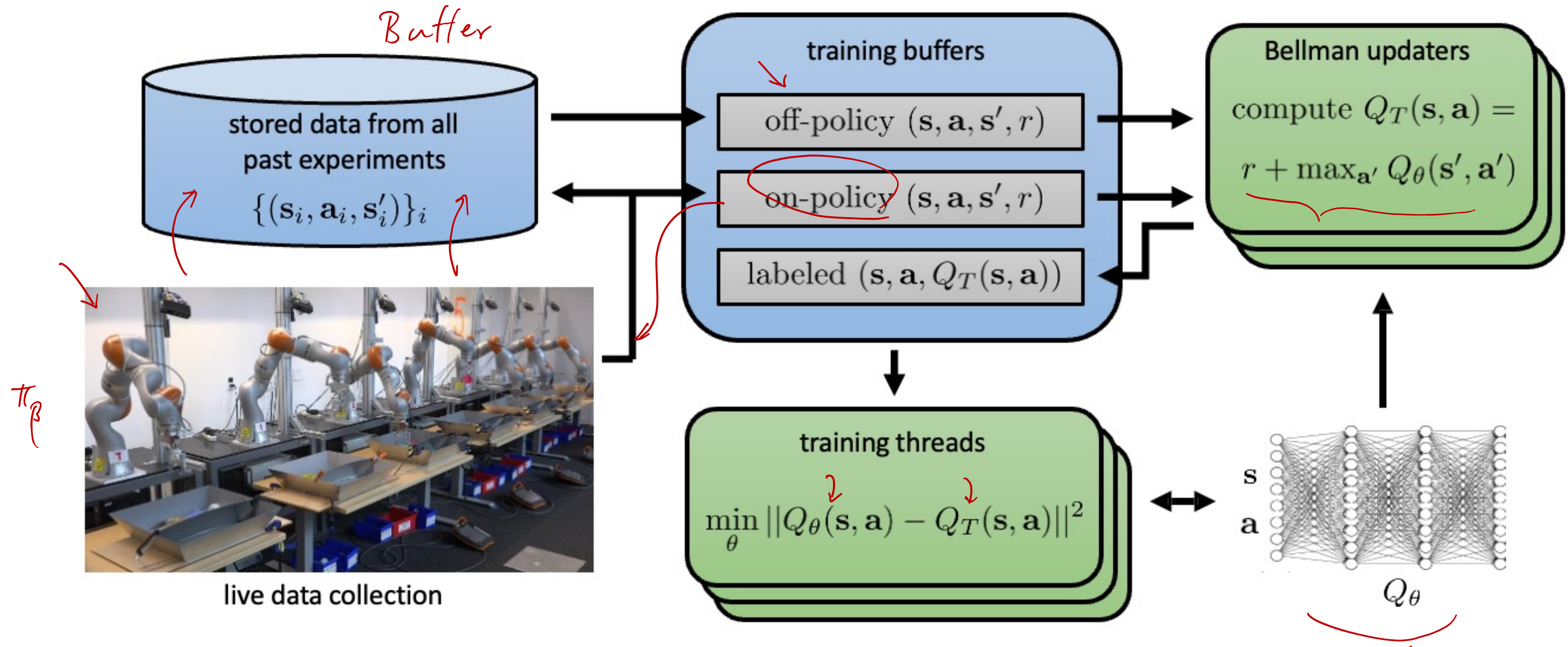
A vivid example



Why should we care?



Does it work?



Kalashnikov, Irpan, Pastor, Ibarz, Herzog, Jang, Quillen, Holly, Kalakrishnan, Vanhoucke, Levine. QT-Opt: Scalable Deep Reinforcement Learning of Vision-Based Robotic Manipulation Skills

Does it work? (Cont.)

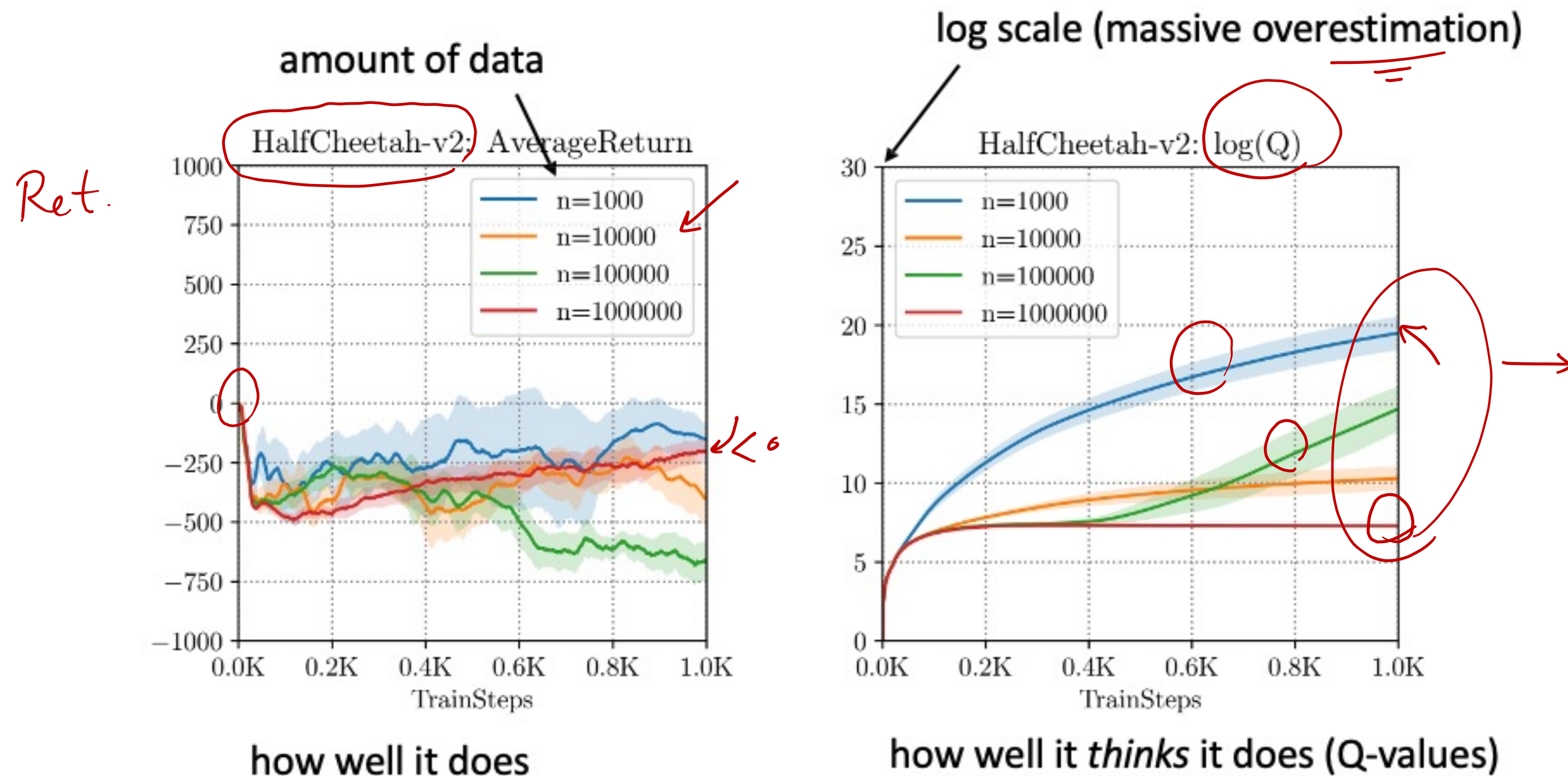


<u>Method</u>	<u>Dataset</u>	<u>Success</u>	<u>Failure</u>
Offline QT-Opt	580k offline	87%	13%
Finetuned QT-Opt	580k offline + 28k online	96%	4%

Kalashnikov, Irpan, Pastor, Ibarz, Herzog, Jang, Quillen, Holly, Kalakrishnan, Vanhoucke, Levine. QT-Opt: Scalable Deep Reinforcement Learning of Vision-Based Robotic Manipulation Skills

Why is offline RL hard?

~~Overfitting~~

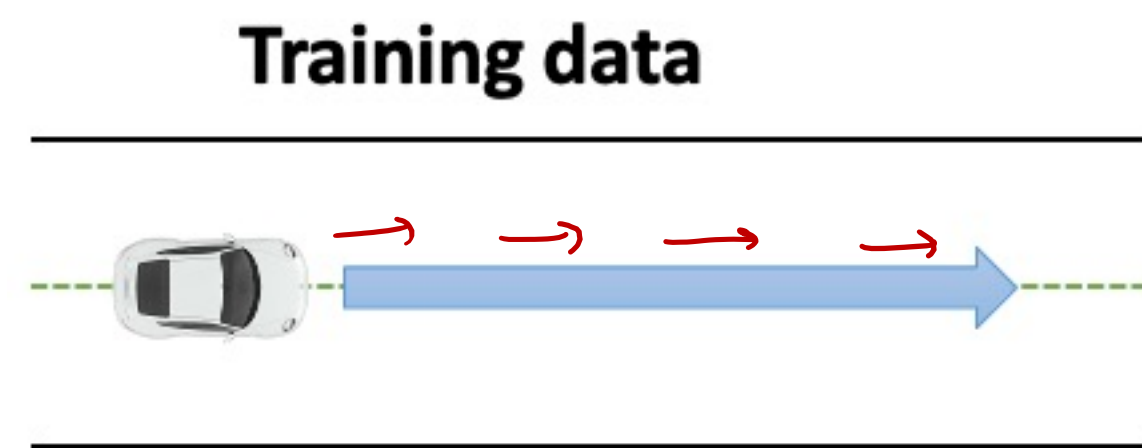


Kumar, Fu, Tucker, Levine. Stabilizing Off-Policy Q-Learning via Bootstrapping Error Reduction. NeurIPS '19

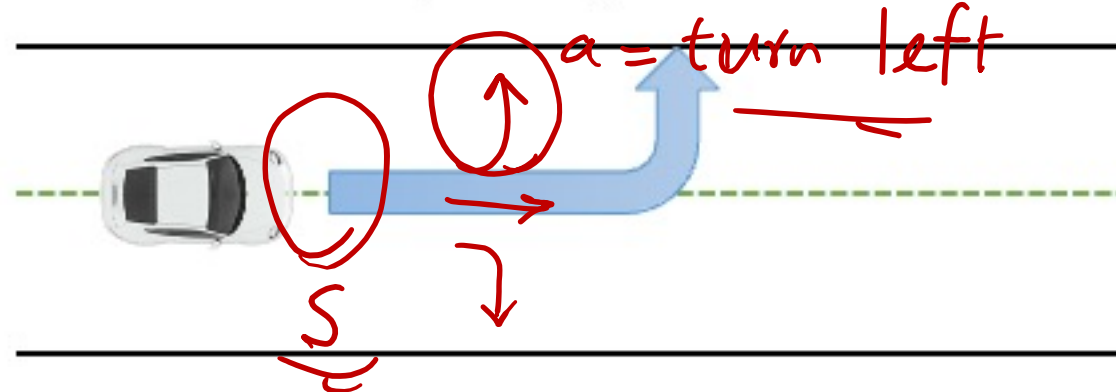
Why is offline RL hard? (Cont.)

- ♦ Fundamental problem: counterfactual queries

D



What the policy wants to do



$$\max_a Q(s, a) = \max \{ Q(s, \text{left}), Q(s, \text{right}), Q(s, \text{straight}) \}$$

Is this good? Bad?
How do we know if we didn't see it in the data?

Buffer sampling $\rightarrow Q(s, \text{left})$

Overshoot
cvf.

- ♦ Online RL algorithms don't have to handle this, because they can simply try this action and see what happens
- ♦ Offline RL methods must somehow account for these unseen ("out-of-distribution") actions, ideally in a safe way
- ♦ ...while still making use of generalization to come up with behaviors
- ♦ that are better than the best thing seen in the data!

Distribution shift in a nutshell

Example empirical risk minimization (ERM) problem:

$$\theta \leftarrow \operatorname{argmin}_{\theta} E_{x \sim p(x), y \sim p(y|x)} [(f_{\theta}(x) - y)^2]$$

given some x^* , is $f_{\theta}(x^*)$ correct?

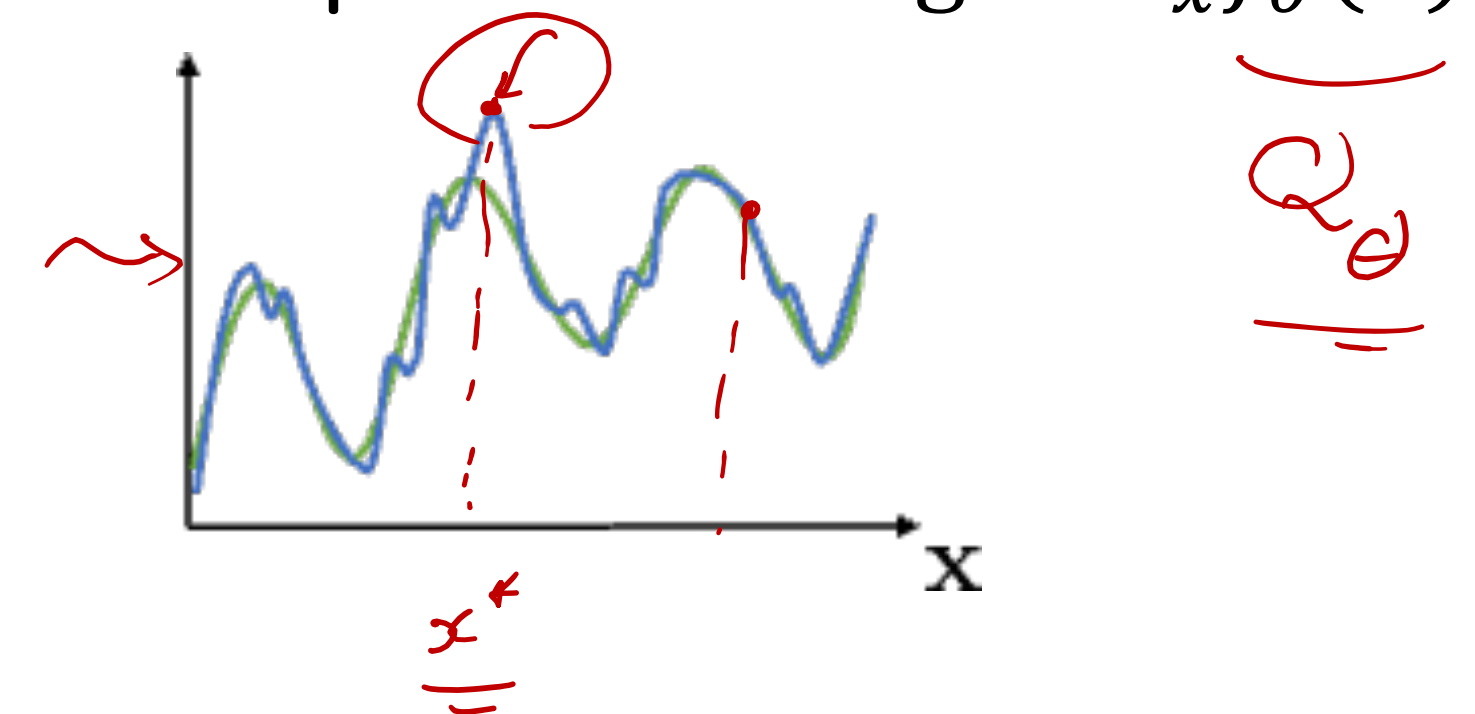
$E_{x \sim p(x), y \sim p(y|x)} [(f_{\theta}(x) - y)^2]$ is low

$E_{x \sim \bar{p}(x), y \sim p(y|x)} [(f_{\theta}(x) - y)^2]$ is not, for general $\bar{p}(x) \neq p(x)$

what if $x^* \sim p(x)$? not necessarily...

usually we are not worried – neural nets generalize well!

what if we pick $x^* \leftarrow \operatorname{argmax}_x f_{\theta}(x)$?



Where do we suffer from distribution shift?

$$Q(s, a) \leftarrow r(s, a) + \max_{a'} Q(s', a')$$

actual (pointing to $r(s, a)$)

OOD action w.r.t π_β (pointing to a' in $\max_{a'}$)

$$Q(s, a) \leftarrow r(s, a) + \underbrace{E_{a' \sim \pi_{new}}[Q(s', a')]}_{y(s, a)}$$

under $\pi_\beta = \pi_{new}$ (pointing to π_{new})

precise (pointing to the expectation term)

what is the objective?

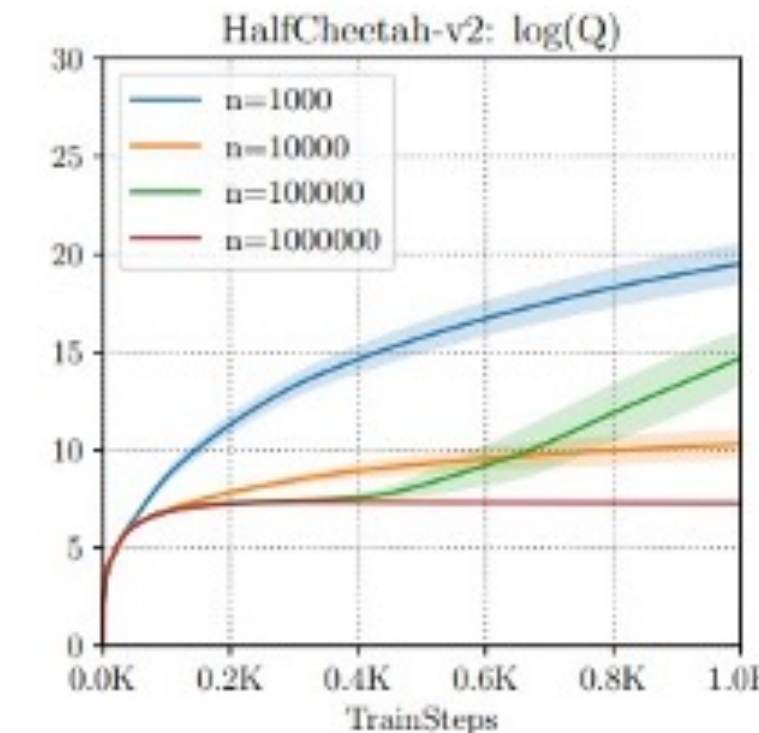
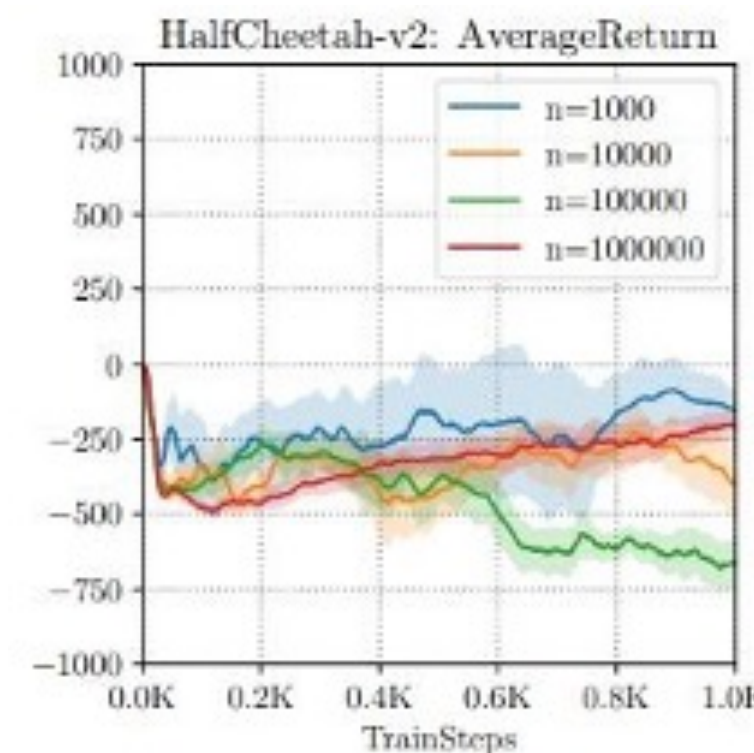
$$\min_Q E_{(s,a) \sim \pi_\beta(s,a)} [(Q(s, a) - \underbrace{y(s, a)}_{\text{target value}})^2]$$

behavior policy

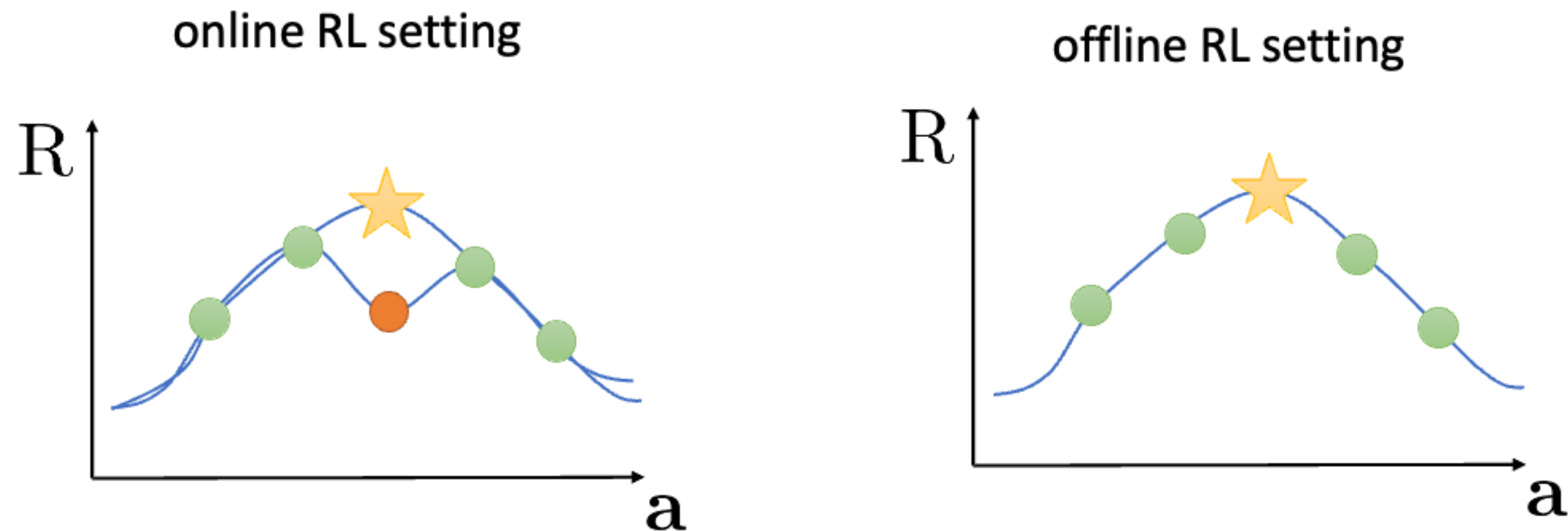
even worse: $\pi_{new} = \operatorname{argmax}_\pi E_{a \sim \pi(a|s)} [Q(s, a)]$

expect good accuracy when $\pi_\beta(a|s) = \pi_{new}(a|s)$ ✓
 how often does that happen?

(what if we pick $x^* \leftarrow \operatorname{argmax}_x f_\theta(x)$?)



Issues with generalization are not corrected



Existing challenges with sampling error and function approximation error in standard RL become much more severe in offline RL

Policy Constraint Methods



How do prior methods address this?

Actor-Critic

$$Q(s, a) \leftarrow r(s, a) + E_{a' \sim \pi_{new}}[Q(s', a')]$$

$$\pi_{new}(a|s) = \operatorname{argmax}_{\pi} E_{a \sim \pi(a|s)}[Q(s, a)] \text{ s.t. } D_{KL}(\pi \parallel \pi_{\beta}) \leq \epsilon$$

- ◆ This solves distribution shift, right?
- ◆ No more erroneous values?
- ◆ Issue 1: we usually don't know the behavior policy

$$\pi_{\beta}(a|s)$$

- ✓ human-provided data
 - ✓ data from hand-designed controller
 - ✓ data from many past RL runs
 - ✓ all of the above

- ◆ Issue 2: this is both too pessimistic and not pessimistic enough

$$\pi_{\beta} \approx \pi_{new} \Rightarrow$$

$$E_{a' \sim \pi_{new}}[Q(s', a')] \text{ precise}$$

- ◆ “policy constraint” method
- ◆ very old idea (but it had no single name?)
- ◆ Todorov et al. [passive dynamics in linearly-solvable]
- ◆ Kappen et al. [KL-divergence control, etc.]
- ◆ trust regions, covariant policy gradients, natural po
- ◆ used in some form in recent papers:
- ◆ Fox et al. ‘15 (“Taming the Noise...”)
- ◆ Fujimoto et al. ‘18 (“Off Policy...”)
- ◆ Jaques et al. ‘19 (“Way Off Policy...”)
- ◆ Kumar et al. ‘19 (“Stabilizing...”)
- ◆ Wu et al. ‘19 (“Behavior Regularized...”)

$$D_{KL}(\pi \parallel \pi_{\beta})$$

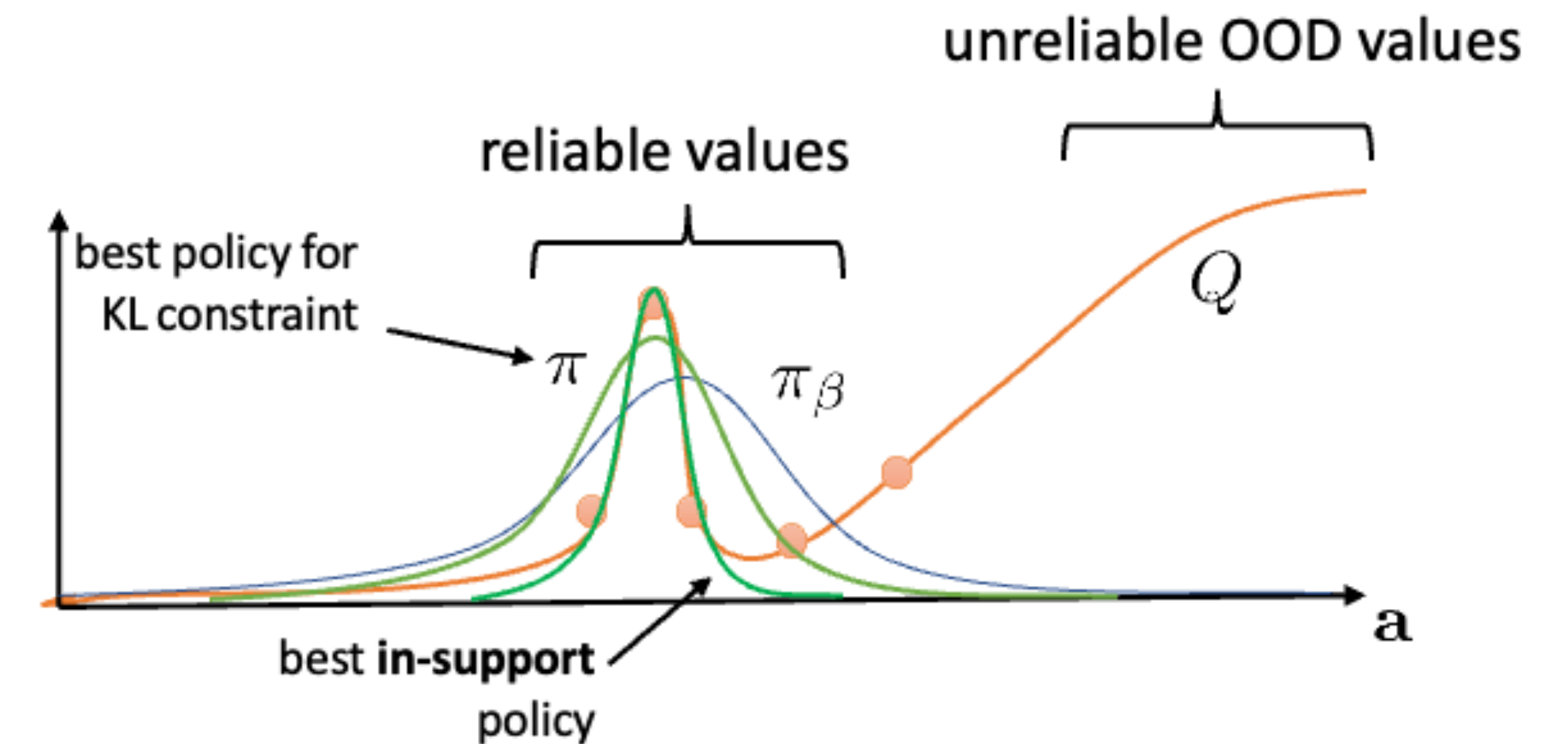
Explicit policy constraint methods

① ♦ What kinds of constraints can we use?

KL-divergence: $D_{KL}(\pi \parallel \pi_\beta)$

+ easy to implement (more on this later)

- not necessarily what we want



② ♦ support constraint: $\pi(a|s) \geq 0$ only if $\pi_\beta(a|s) \geq \epsilon$
can approximate with MMD

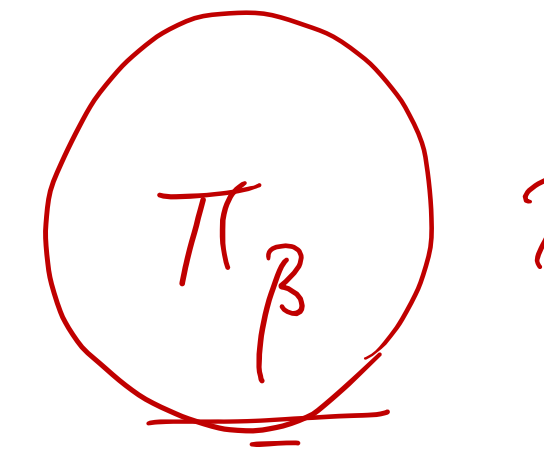
- significantly more complex to implement

+ much closer to what we really want

♦ For more information, see:

- ✓ Levine, Kumar, Tucker, Fu. Offline Reinforcement Learning: Tutorial, Review, and Perspectives on Open Problems. '20
- ✓ Kumar, Fu, Tucker, Levine. Stabilizing Off-Policy Q-Learning via Bootstrapping Error Reduction. '19
- ✓ Wu, Tucker, Nachum. Behavioral Regularized Offline Reinforcement Learning. '19

Explicit policy constraint methods (Cont.)



♦ How do we implement constraints?

✓ Modify the actor objective

$$\theta \leftarrow \operatorname{argmax}_{\theta} E_{s \sim D} [E_{a \sim \pi_{\theta}(a|s)} [Q(s, a)]]$$

Lagrange multiplier

easy to compute and differentiate for Gaussian or categorical po

$$\theta \leftarrow \operatorname{argmax}_{\theta} E_{s \sim D} [E_{a \sim \pi_{\theta}(a|s)} [Q(s, a) + \lambda \log \pi_{\beta}(a|s)]] + \lambda H(\pi(a|s))$$

$$D_{KL}(\pi \parallel \pi_{\beta}) = E_{\pi} [\log \pi(a|s) - \log \pi_{\beta}(a|s)] = -E_{\pi} [\log \pi_{\beta}(a|s)] - H(\pi)$$

✓ Modify the reward function

$$\bar{r}(s, a) = r(s, a) - D(\pi, \pi_{\beta})$$

simple modification to directly penalize

divergence also accounts for future divergence

See: Wu, Tucker, Nachum. Behavior Regularized Offline Reinforcement Learning. `19

generally, the best modern offline RL methods do not do either of these things

Implicit policy constraint methods

$$\frac{\partial}{\partial \pi(s)} \left(\mathbb{E}_{a \sim \pi} [Q(s, a)] - \lambda D_{KL}(\pi \parallel \pi_\beta) \right) = 0$$

$$\pi_{new}(\mathbf{a}|\mathbf{s}) = \operatorname{argmax}_{\pi} \mathbb{E}_{\mathbf{a} \sim \pi(\mathbf{a}|\mathbf{s})} [Q(\mathbf{s}, \mathbf{a})] \text{ s.t. } D_{KL}(\pi \parallel \pi_\beta) \leq \epsilon$$

closed form sol.

$$\rightarrow \pi^*(\mathbf{a}|\mathbf{s}) = \frac{1}{Z(\mathbf{s})} \pi_\beta(\mathbf{a}|\mathbf{s}) \exp \left(\frac{1}{\lambda} A^\pi(\mathbf{s}, \mathbf{a}) \right)$$

$$\pi_\Theta \leftrightarrow \pi^*$$

straightforward to show via duality

See also: Peters et al. (REPS)

approximate via weighted max likelihood!

$$\pi_{new}(\mathbf{a}|\mathbf{s}) = \operatorname{argmax}_{\pi} \mathbb{E}_{(\mathbf{s}, \mathbf{a}) \sim \pi_\beta} \left[\log \pi(\mathbf{a}|\mathbf{s}) \frac{1}{Z(\mathbf{s})} \exp \left(\frac{1}{\lambda} A^{\pi_{old}}(\mathbf{s}, \mathbf{a}) \right) \right]$$

samples from dataset $\mathbf{a} \sim \pi_\beta(\mathbf{a}|\mathbf{s})$

critic can be used to give us this

$$\begin{aligned} & \max_{\Theta} \mathbb{E}_{\pi^*} [\log \pi_\Theta] \\ & \int \log \pi_\Theta \frac{1}{Z(\mathbf{s})} \pi_\beta(\mathbf{a}|\mathbf{s}) d\mathbf{a} \\ & \mathbb{E}_{(\mathbf{s}, \mathbf{a}) \sim \pi_\beta} \left\{ \log \pi_\Theta(\mathbf{s}, \mathbf{a}) \right\} \\ & = \mathbb{E}_{(\mathbf{s}, \mathbf{a}) \sim \pi_\beta} \left\{ g_\Theta(\mathbf{s}, \mathbf{a}) \right\} \end{aligned}$$

Implicit policy constraint methods (Cont.)

$$\mathcal{L}_C(\phi) = E_{(s,a,s') \sim D} [(Q_\phi(s,a) - (r(s,a) + \gamma E_{a' \sim \pi_\theta(a'|s')} [Q_\phi(s',a')]))^2]$$

it's the same policy that maximises the Q-func.

$$\mathcal{L}_A(\theta) = -E_{(s,a) \sim \pi_\beta} \left[\log \pi_\theta(a|s) \frac{1}{Z(s)} \exp \left(\frac{1}{\lambda} A^{\pi_{old}}(s,a) \right) \right]$$

Offline Data

implicitly enforces $D_{KL}(\pi_\theta \parallel \pi_\beta) \approx 0$

$$1. \phi \leftarrow \phi - \alpha \nabla_\phi \mathcal{L}_C(\phi)$$

$$2. \theta \leftarrow \theta - \alpha \nabla_\theta \mathcal{L}_A(\theta)$$

$$Q(s,a) \leftarrow r(s,a) + E_{a' \sim \pi_{new}} [Q(s',a')]$$

$$\pi_{new}(a|s) = \underset{\pi}{\operatorname{argmax}} E_{a \sim \pi(a|s)} [Q(s,a)] \text{ s.t. } D_{KL}(\pi \parallel \pi_\beta) \leq \epsilon$$

Can we also avoid all OOD actions in the Q update?

$$Q(s, a) \leftarrow r(s, a) + E_{a' \sim \pi_{new}}[Q(s', a')] \rightarrow V^*(s')$$

$V(s') \leftarrow$ just another neural network

$$V^* = \max_a Q(s', a)$$

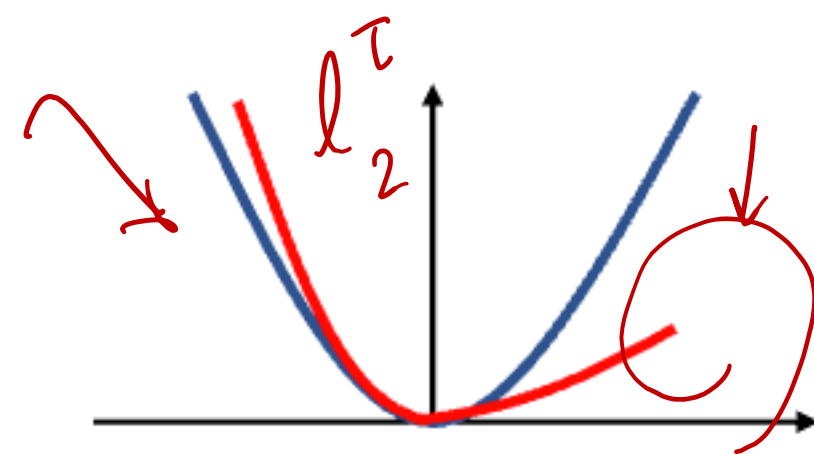
$$V \leftarrow \operatorname{argmin} \frac{1}{N} \sum_{i=1}^N \ell(V(s_i), Q(s_i, a_i)) \rightarrow V^*(s) =$$

this action comes from π_β not from π_{new}

e.g., MSE loss $(V(s_i) - Q(s_i, a_i))^2$

$$\frac{1}{N} \sum_i 2(V(s_i) - Q(s_i, a_i)) = 0$$

expectile: $\ell_2^\tau(x) = \begin{cases} (1-\tau)x^2 & \text{if } x > 0 \\ \tau x^2 & \text{else} \end{cases}$



$$V(s) \leftarrow \max_{a \in \Omega(s)} Q(s, a)$$

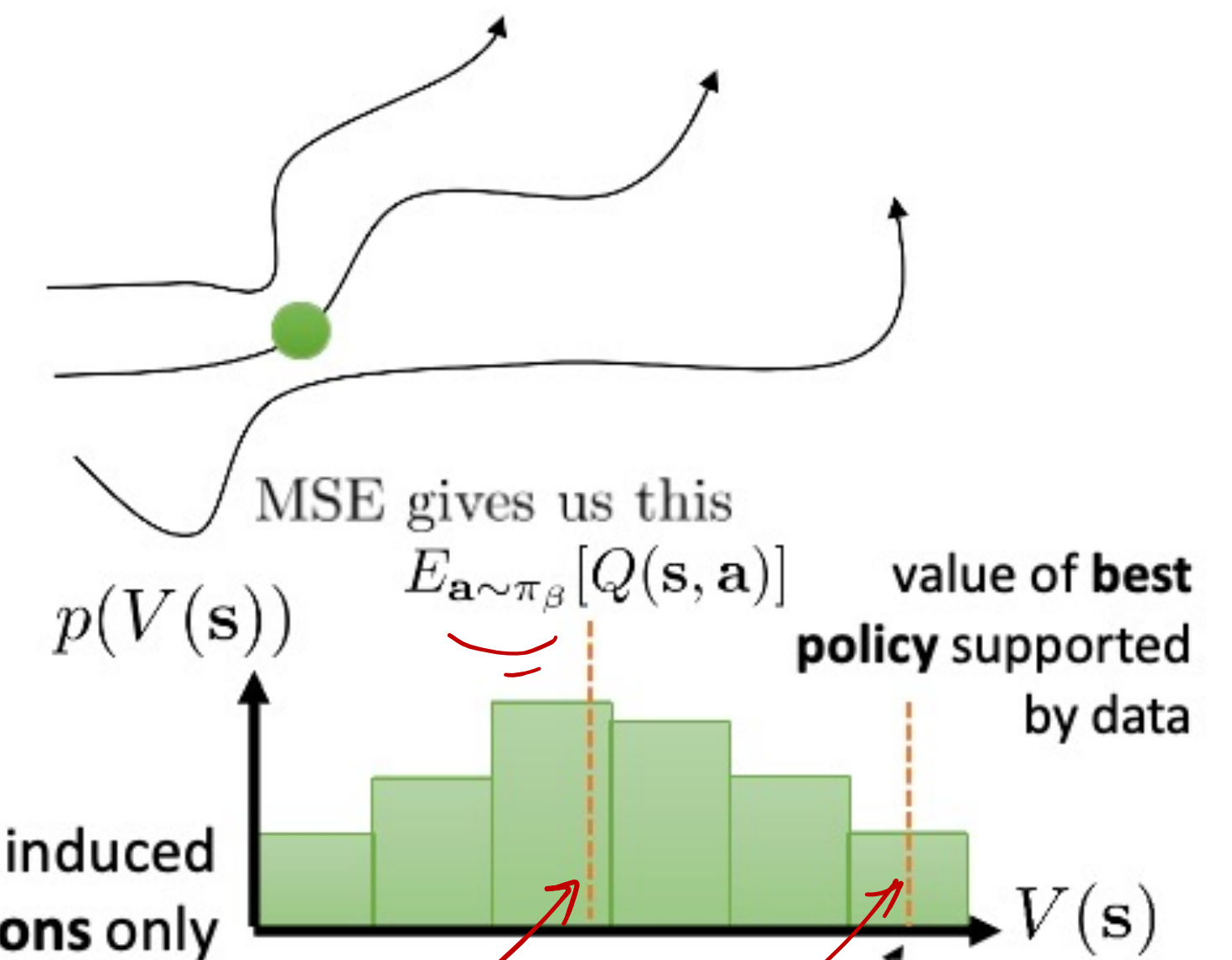
$$\Omega(s) = \{a: \pi_\beta(a|s) \geq \epsilon\}$$

if we use ℓ_2^τ for big τ

$$V(s) = \frac{1}{N} \sum_{i=1}^N Q(s, a_i) \rightarrow a_i \sim \underline{D} \sim \pi_\beta$$

distribution is induced by actions only

could another loss give us this?



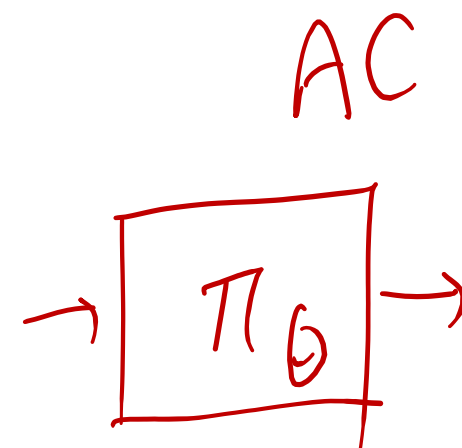
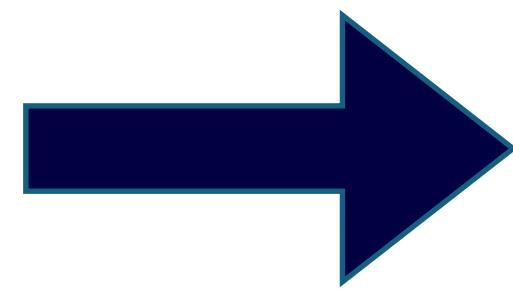
Implicit Q-learning (IQL)

$$Q(s, a) \leftarrow r(s, a) + V(s')$$

$$V(s) \leftarrow \max_{a \in \Omega(s)} Q(s, a)$$

$$\Omega(s) = \{a: \pi_{\beta}(a|s) \geq \epsilon\}$$

if we use ℓ_2^T for big τ



Q_{ϕ} $r(s, a) + \gamma V(s')$

V_{ψ}

ℓ_2^T $\tau \sim 95\%$

$V \leftarrow \underset{V}{\operatorname{argmin}} \frac{1}{N} \sum_{i=1}^N \ell_2^V(V(s_i), Q(s_i, a_i))$

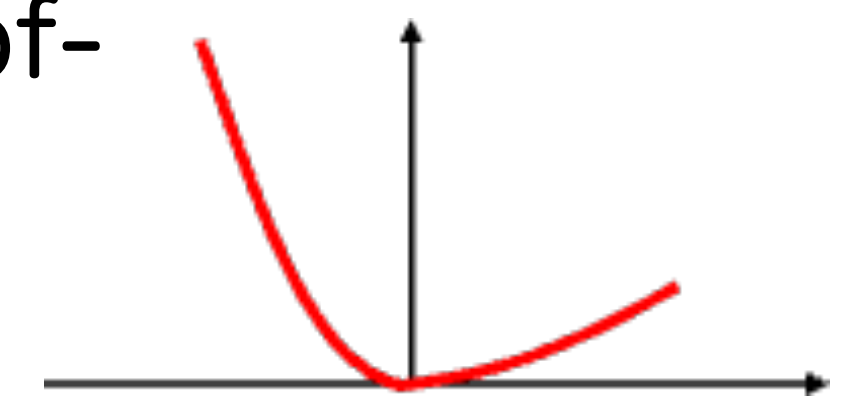
π

$$Q(s, a) \leftarrow r(s, a) + \max_{a' \in \Omega(s')} Q(s', a')$$

“implicit” policy

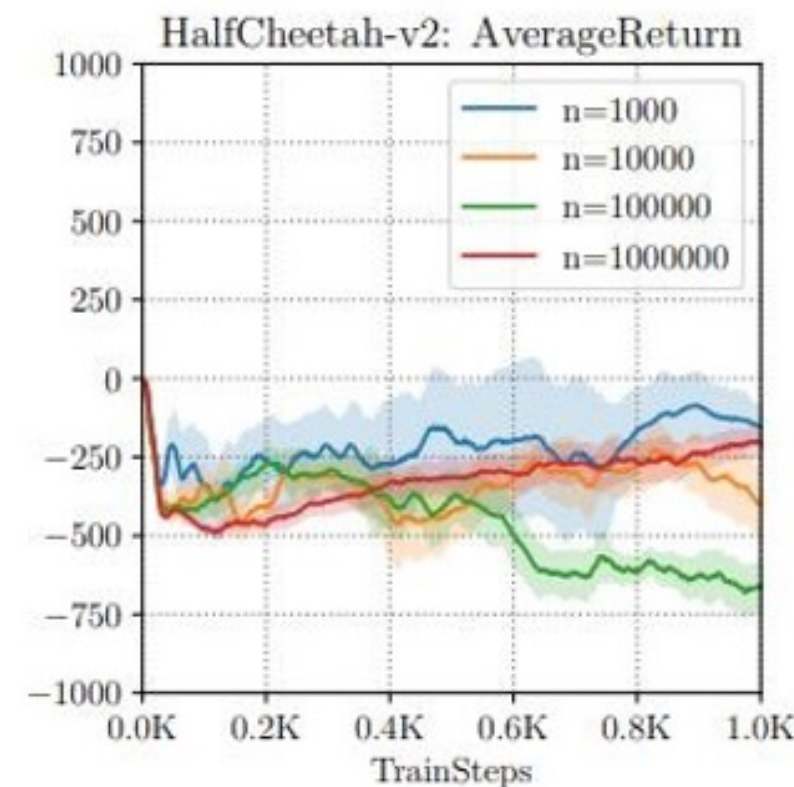
$$\pi_{new}(a|s) = \delta(a = \arg \max_{a \in \Omega(s)} Q(s, a))$$

✓ Now we can do value function updates without ever risking out-of-distribution actions!

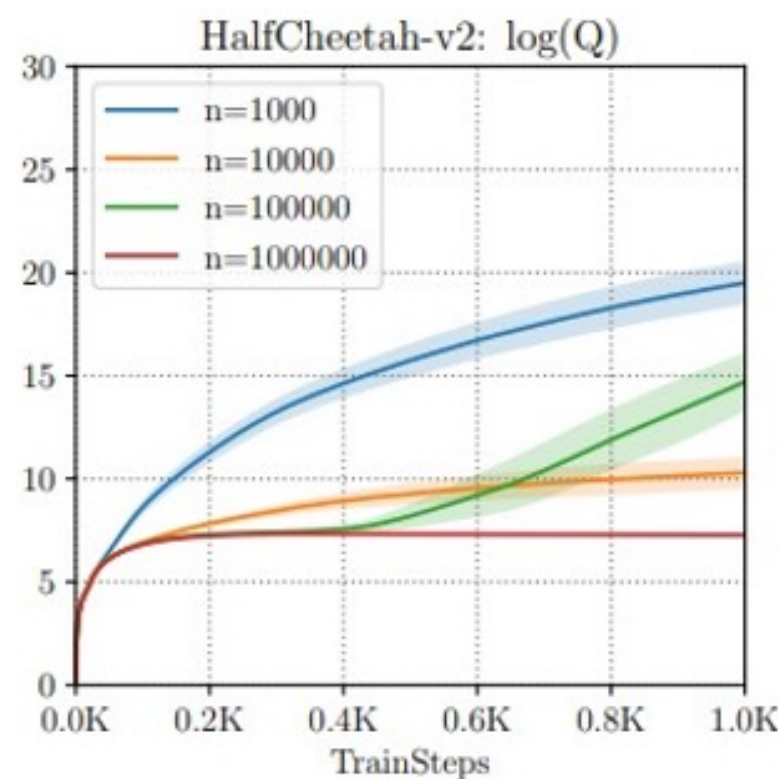


Conservative Q-Learning

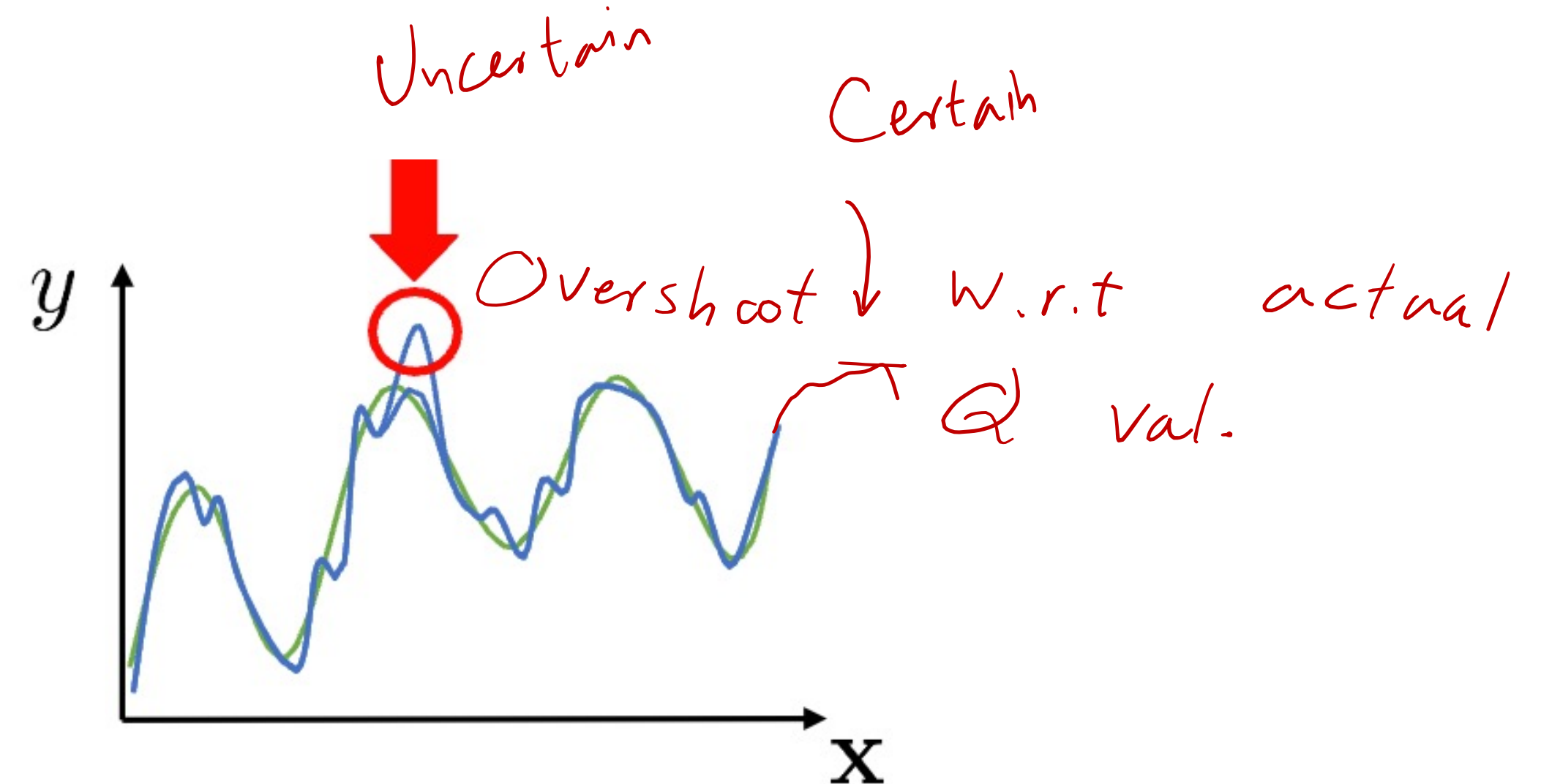
Conservative Q-learning (CQL)



how well it does



how well it *thinks* it does (Q-values)



$$\hat{Q}^\pi = \arg \min_Q$$

$$\max_{\mu} \alpha E_{s \sim D, a \sim \mu(a|s)} [Q(s, a)]$$

$$E \{ U_{(s,a)} \}$$

$a \sim \mu$

term to push down big Q-values

regular objective: $+E_{(s,a,s') \sim D} [(Q(s, a) - (r(s, a) + E_\pi[Q(s', a')]))^2]$

can show that $\hat{Q}^\pi \leq \underline{Q}^\pi$ for large enough α

true Q-function

π

θ

$E[Q(s, a)]$

$a \sim \pi_\theta$

Conservative Q-learning (CQL) - (Cont.)

A better bound:

always pushes Q-values down

push up on (s, a) samples in data

$$\hat{Q}^\pi = \operatorname{argmin}_Q \max_\mu \alpha E_{s \sim D, a \sim \mu(a|s)} [Q(s, a)] - \alpha E_{(s, a) \sim D} [Q(s, a)] + E_{(s, a, s') \sim D} [(Q(s, a) - (r(s, a) + E_\pi[Q(s', a')]))^2]$$

$\mathcal{L}_{CQL} \hat{Q}^\pi$

no longer guaranteed that $\hat{Q}^\pi(s, a) \leq Q^\pi(s, a)$ for all (s, a)

but guaranteed that $E_{\pi(a|s)} [\hat{Q}^\pi(s, a)] \leq E_{\pi(a|s)} [Q^\pi(s, a)]$ for all $s \in D$

Conservative Q-learning (CQL) - (Cont.)

- 1. Update $\hat{Q}^\pi$ w.r.t. $\mathcal{L}_{CQL}(\hat{Q}^\pi)$ using $\mathcal{D}$
- 2. Update policy π

if actions are discrete:

$$\pi(\mathbf{a} | \mathbf{s}) = \begin{cases} 1 & \text{if } \mathbf{a} = \arg \max_{\mathbf{a}} \hat{Q}(\mathbf{s}, \mathbf{a}) \\ 0 & \text{otherwise} \end{cases}$$

$$\hat{Q} \leq Q$$

if actions are continuous:

$$\theta \leftarrow \theta + \alpha \nabla_{\theta} \sum_i E_{\mathbf{a} \sim \pi_{\theta}(\mathbf{a} | \mathbf{s}_i)} [\hat{Q}(\mathbf{s}_i, \mathbf{a})]$$

Conservative Q-learning (CQL) - (Cont.)

bi-level

$$\hat{Q}^\pi = \underset{Q}{\operatorname{argmin}} \underset{\mu}{\operatorname{max}} \alpha E_{s \sim D, a \sim \mu(a|s)}[Q(s, a)] - \alpha E_{(s, a) \sim D}[Q(s, a)] + \mathcal{R}(\mu) + E_{(s, a, s') \sim D}[(Q(s, a) - (r(s, a) + E_\pi[Q(s', a')]))^2]$$

$\mathcal{L}_{CQL} \hat{Q}^\pi$

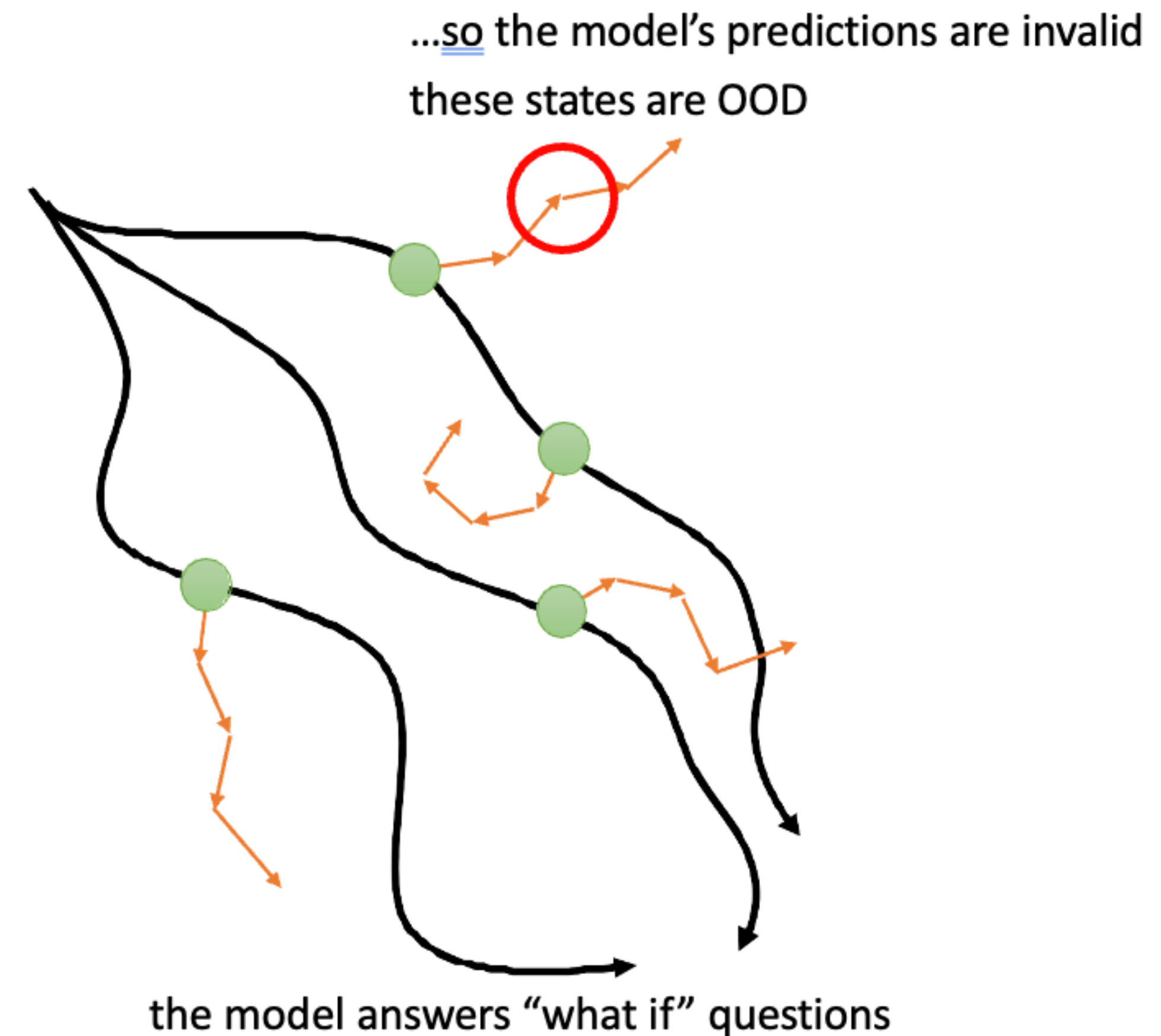
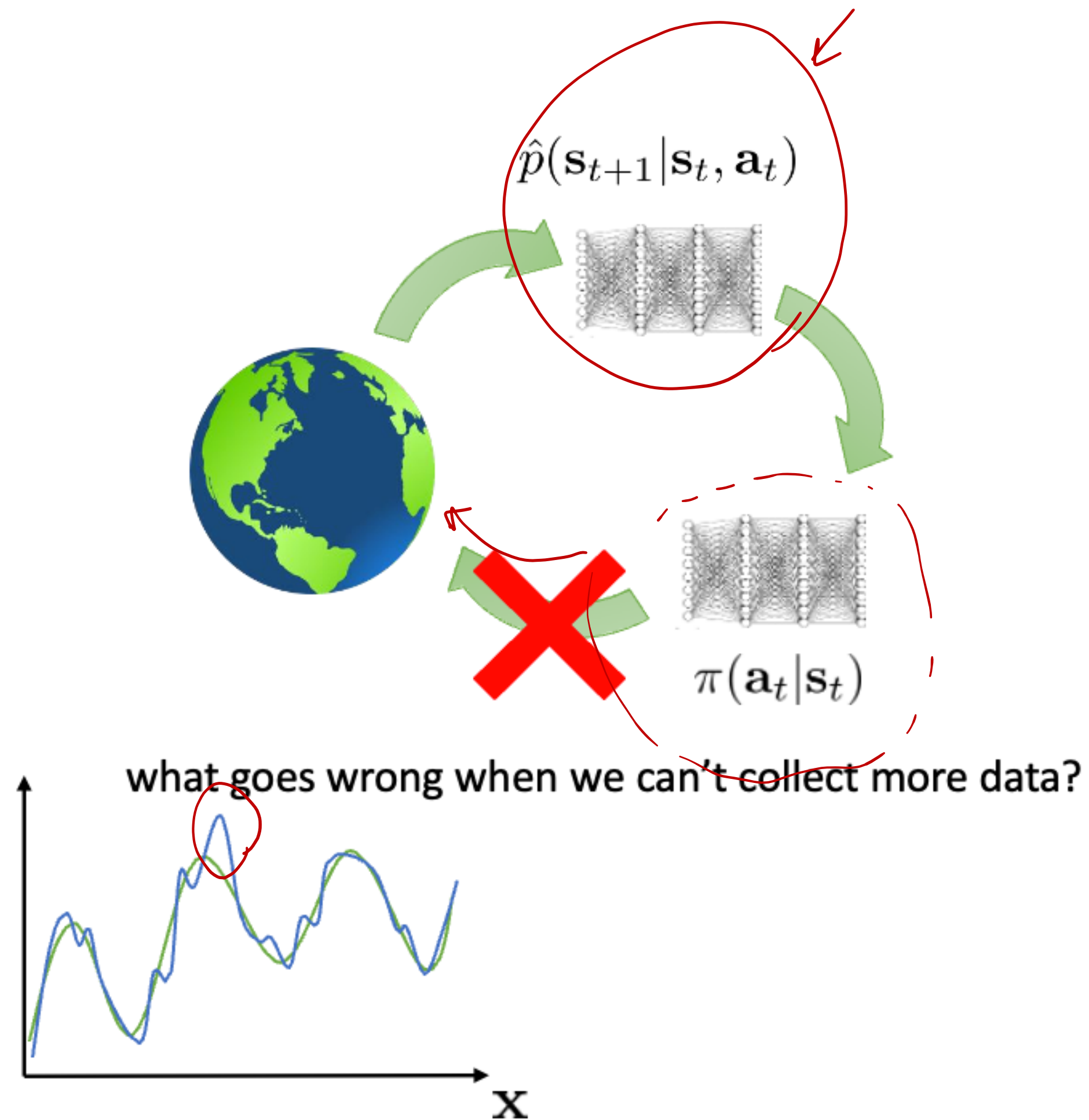
$D_{KL}(\mu \parallel \rho)$
 $\mathcal{H}(\mu)$
 $\mu(a|s) = \frac{e^{\beta Q(s, a)}}{\sum_{a'} e^{\beta Q(s, a)}}$

common choice: $\mathcal{R} = E_{s \sim D}[\mathcal{H}(\mu(\cdot | s))]$ maximum entropy regularization

Model-Based Offline RL



How does model-based RL work?



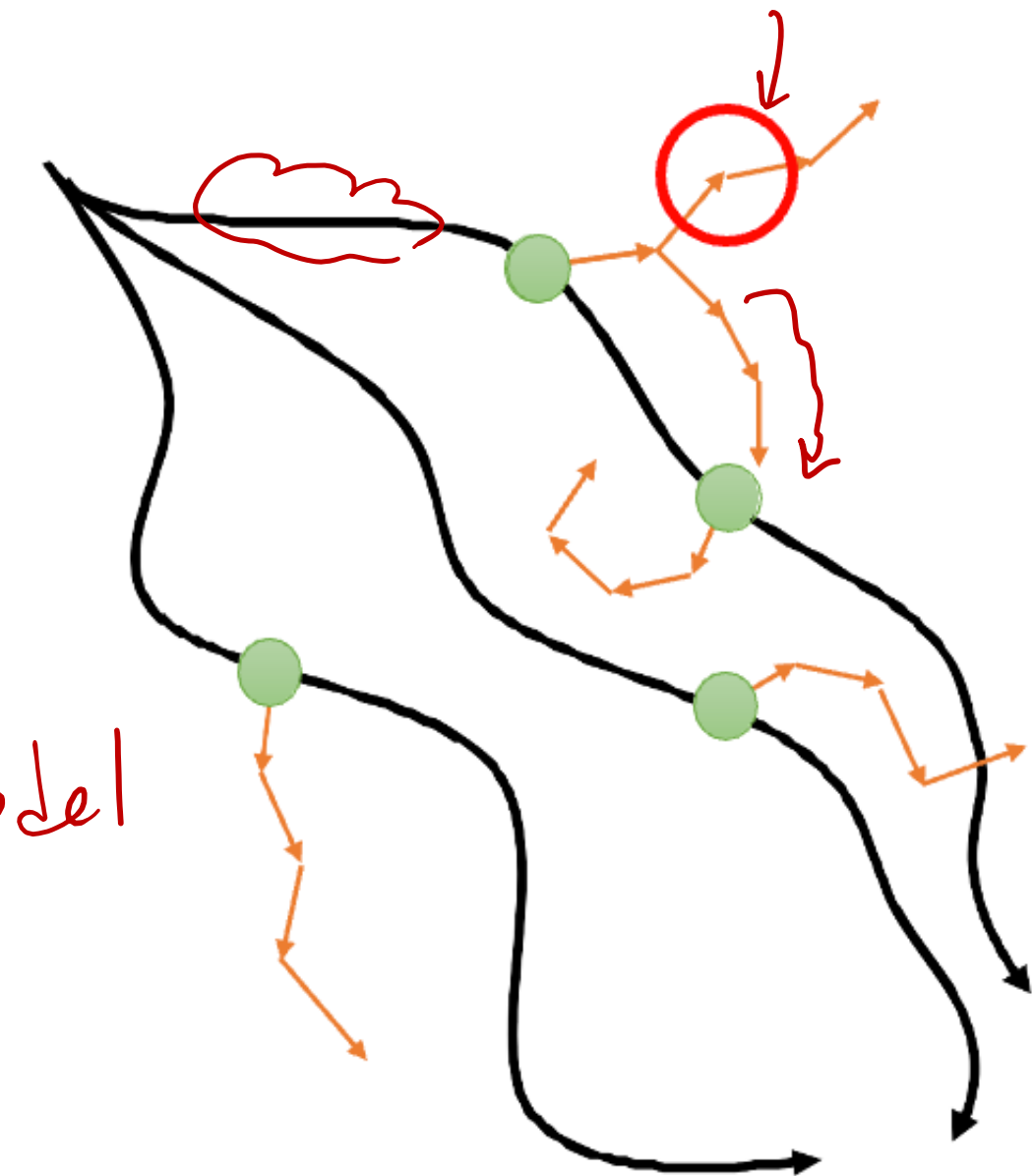
Model-Based Offline RL

solution: “punish” the policy for exploiting

$$\tilde{r}(s, a) = r(s, a) - \lambda u(s, a)$$

Ensemble of World Model

uncertainty penalty



...and then use any existing model-based RL algorithm

Yu, Thomas, Yu, Ermon, Zou, Levine, Finn, Ma. MOPO: Model-Based Offline Policy Optimization. '20

See also: Kidambi et al., MOREL: Model-Based Offline Reinforcement Learning. '20 (concurrent)

Conservative Model-Based RL

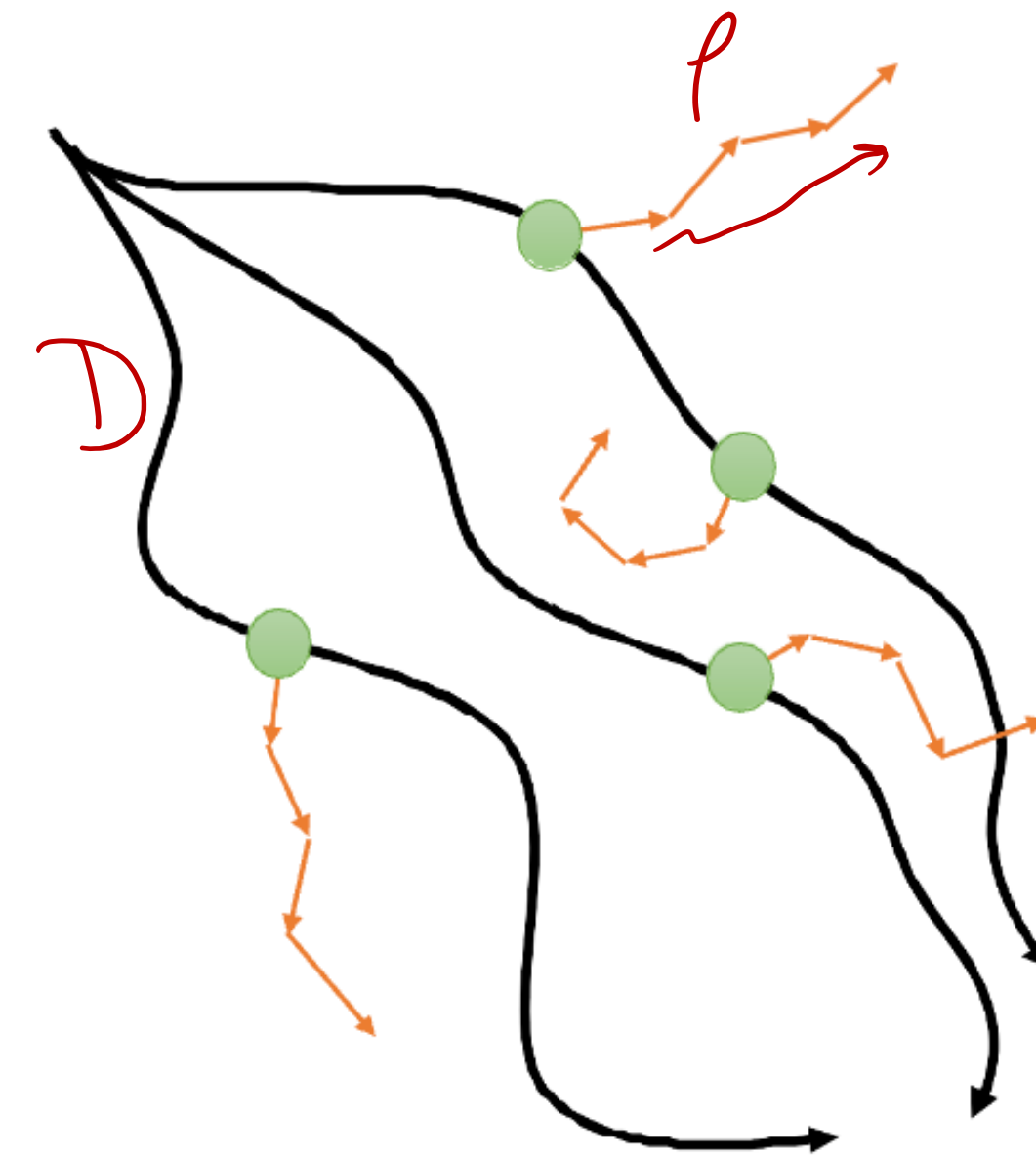
Basic idea: just like CQL minimizes Q-value of policy actions, we can minimize Q-value of model state-action tuples

$$\underline{\rho} = \underline{\hat{\mathcal{M}}} \cdot \underline{\pi}$$

state-action tuples from the model

$$Q^{k+1} \leftarrow \operatorname{argmin}_Q \beta \left(\mathbb{E}_{s,a \sim \underline{\rho}(s,a)} [Q(s,a)] - \mathbb{E}_{s,a \sim \underline{D}} [Q(s,a)] \right)$$

$$+ \frac{1}{2} \mathbb{E}_{s,a,s' \sim d_f} \left[\left(Q(s,a) - B^{\pi} \hat{Q}^k(s,a) \right)^2 \right].$$



Intuition: if the model produces something that looks clearly different from real data, it's easy for the Q-function to make it look bad

Summary, Applications, Open Questions

Which offline RL algorithm do I use?

If you want to only train offline...

✓ Conservative Q-learning + just one hyperparameter + well understood and widely tested Implicit Q-learning + more flexible (offline + online) more hyperparameters τ

If you want to only train offline and finetune online

Advantage-weighted actor-critic (AWAC) + widely used and well tested Implicit Q-learning + seems to perform much better!

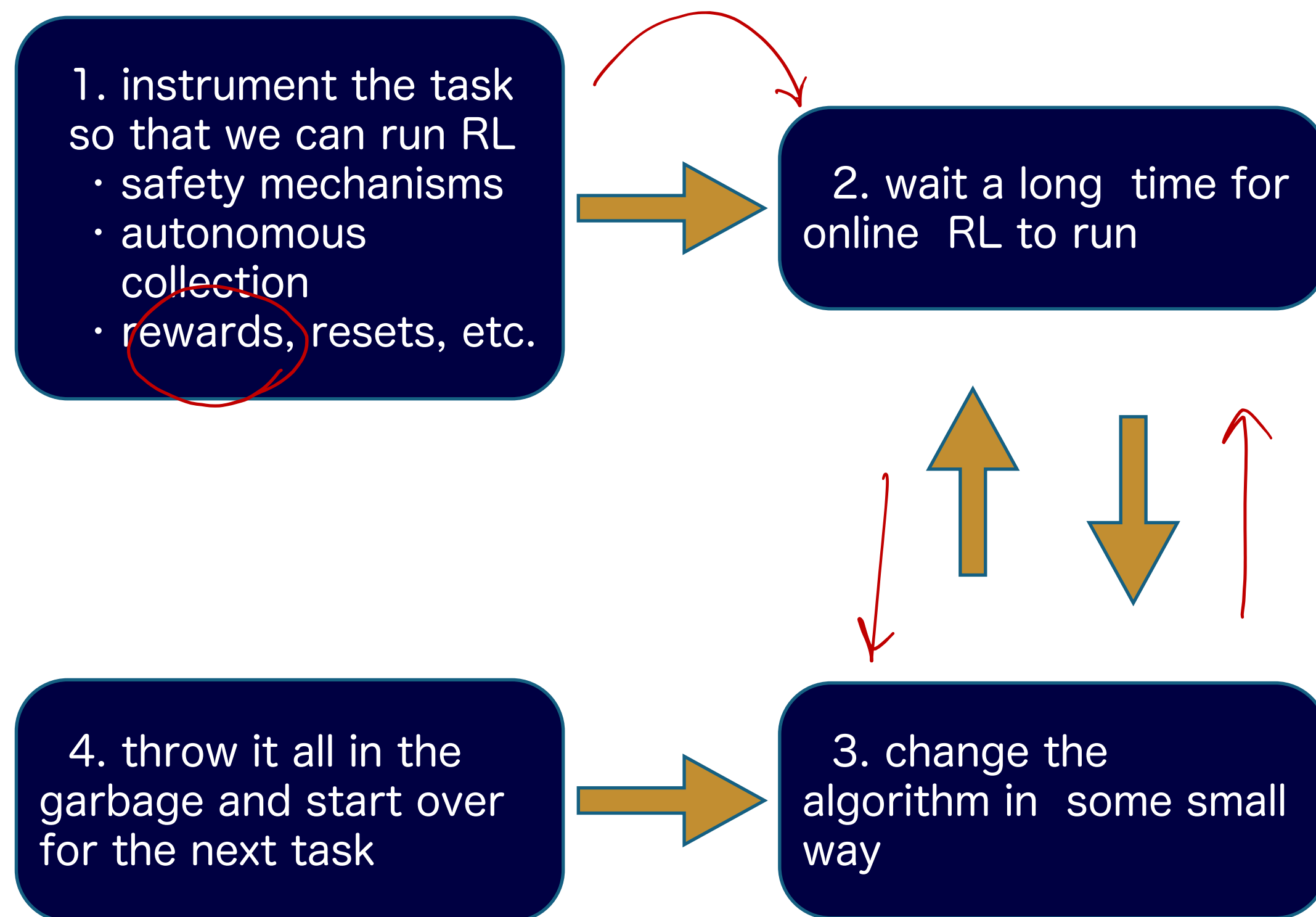
If you have a good way to train models in your domain

COMBO + similar properties as CQL, but benefits from models

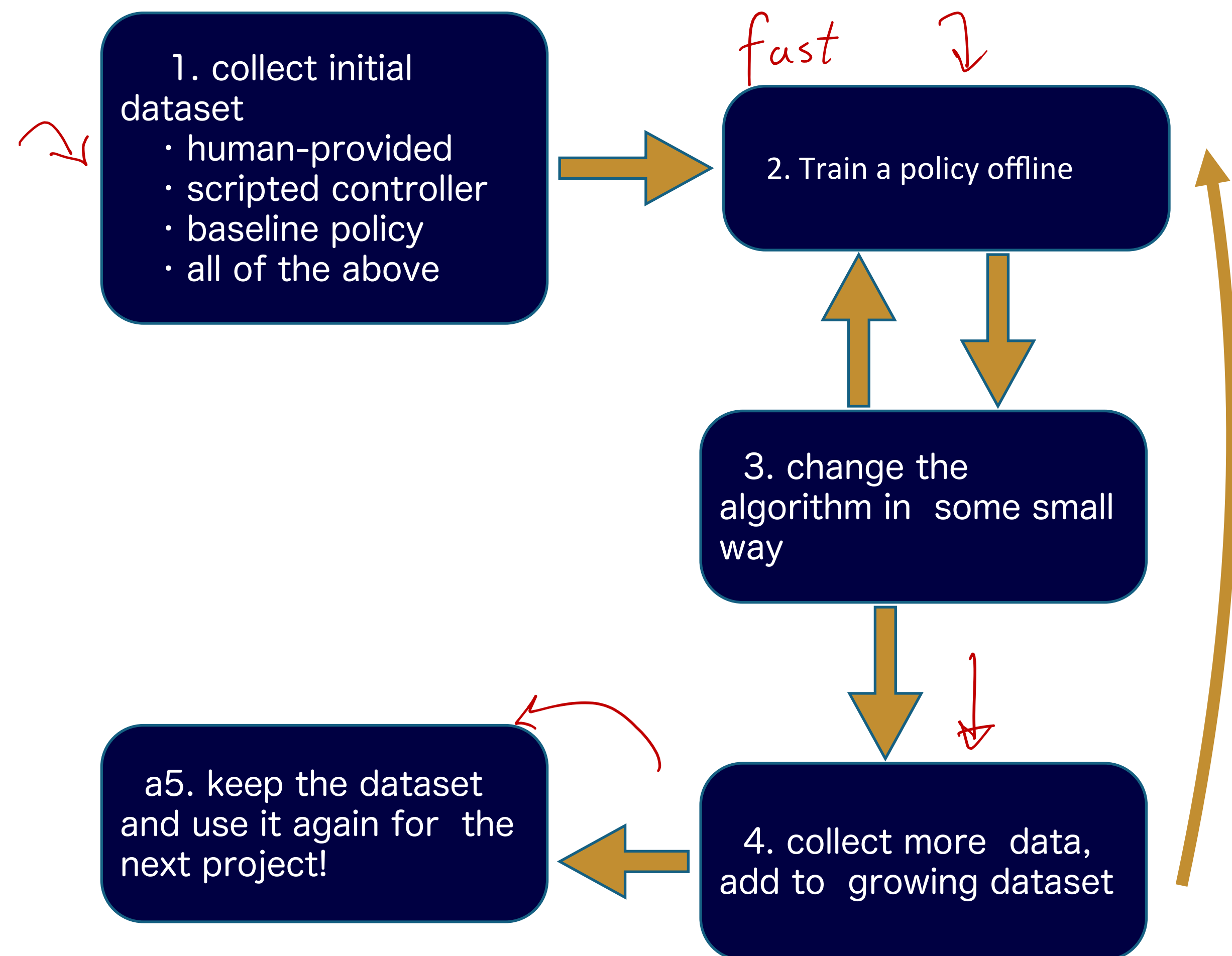
- not always easy to train a good model in your domain!

The power of offline RL

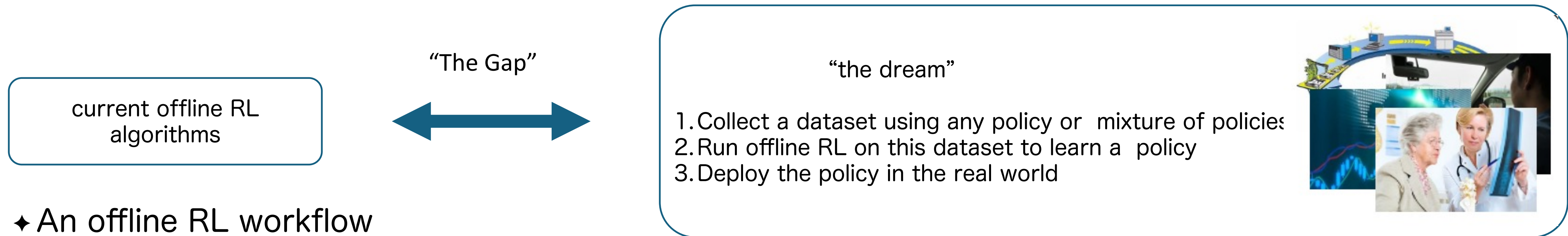
standard real-world RL process



offline RL process



Takeaways, conclusions, future directions



♦ An offline RL workflow

- ✓ Supervised learning workflow: train/test split
- ✓ Offline RL workflow: ???

♦ Statistical guarantees

- ✓ Biggest challenge: distributional shift/counterfactuals
- ✓ Can we make any guarantees?

♦ Scalable methods, large-scale applications

- ✓ Dialogue systems
- ✓ Data-driven navigation and driving

$$\hat{Q} \leq Q$$