

Interactive Imitation Learning

Recap

The Behavior Cloning algorithm:

Choose regression (for continuous action) or classification loss $\ell(\pi(s), a)$, and perform SL:

$$\hat{\pi} = \min_{\pi \in \Pi} \sum_{i=1}^M \ell(\pi(s_i^*), a_i^*)$$

ERM

(s^*, a^*)

$s^* \sim d^{\pi^*}$

$a^* \sim \pi^*(\cdot | s^*)$

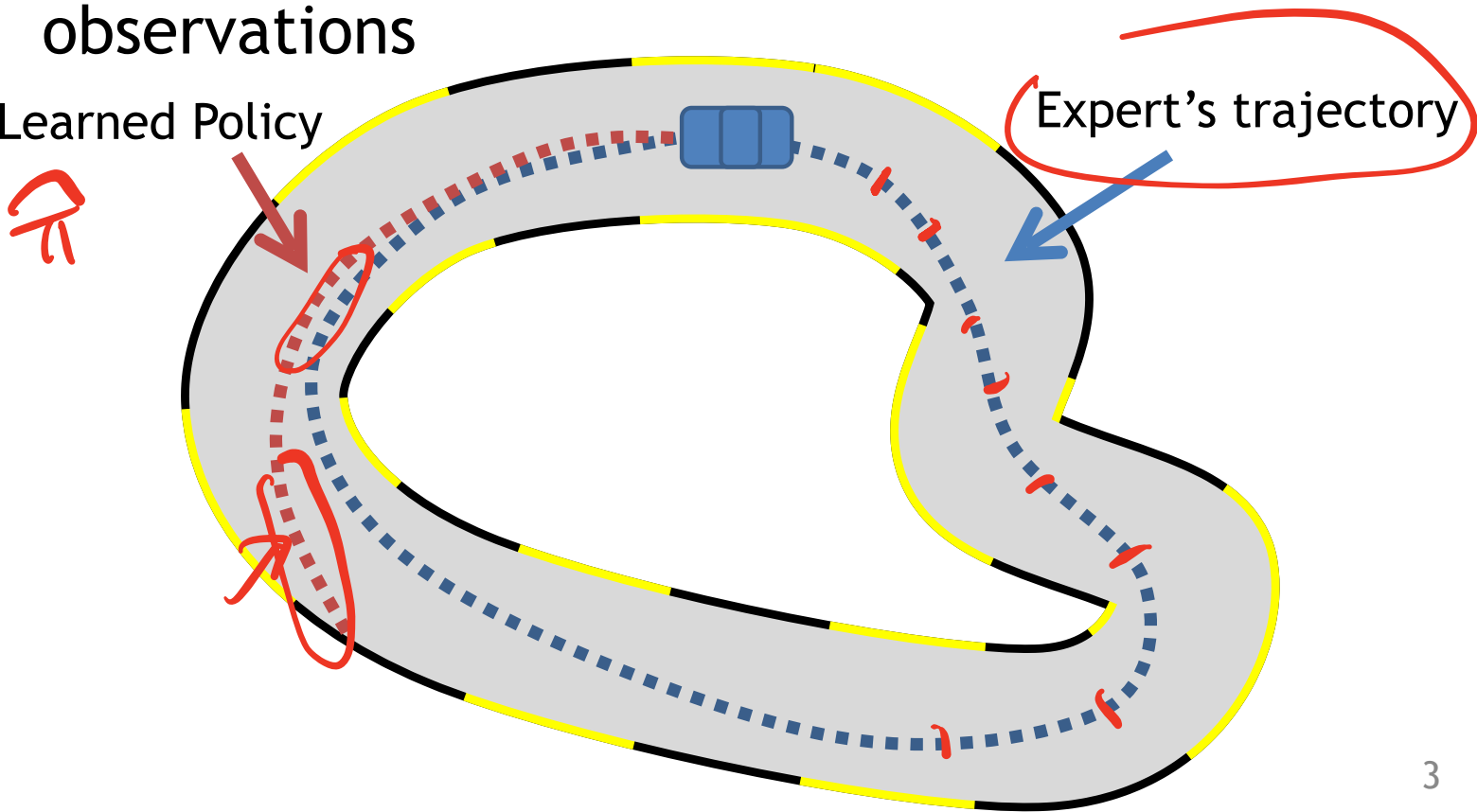
What could go wrong?

[Pomerleau89, Daume09]

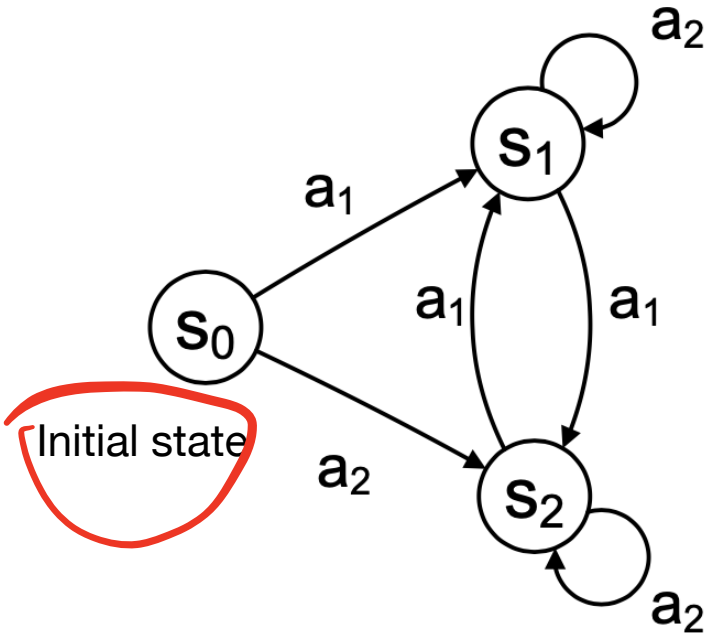
- Predictions affect future inputs/observations

Learned Policy

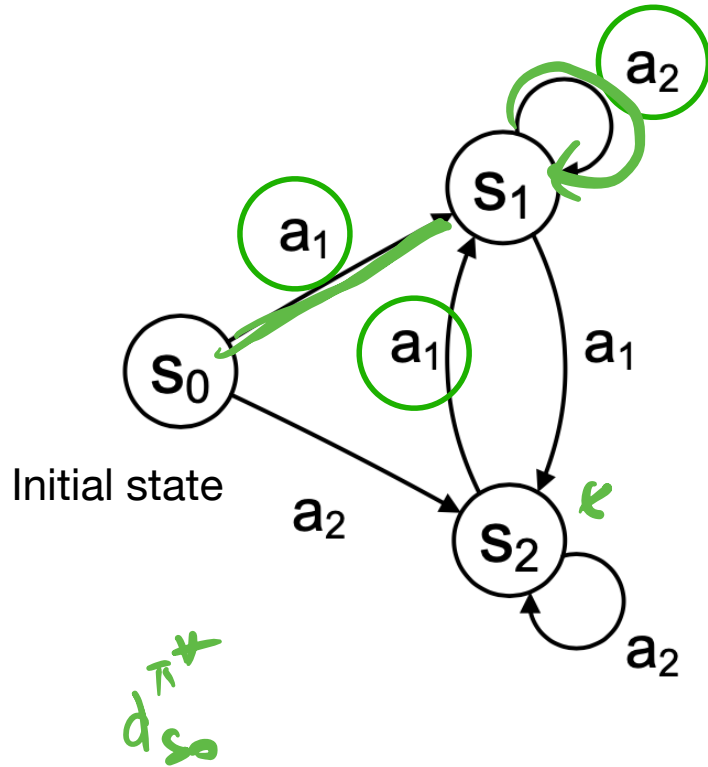
Expert's trajectory



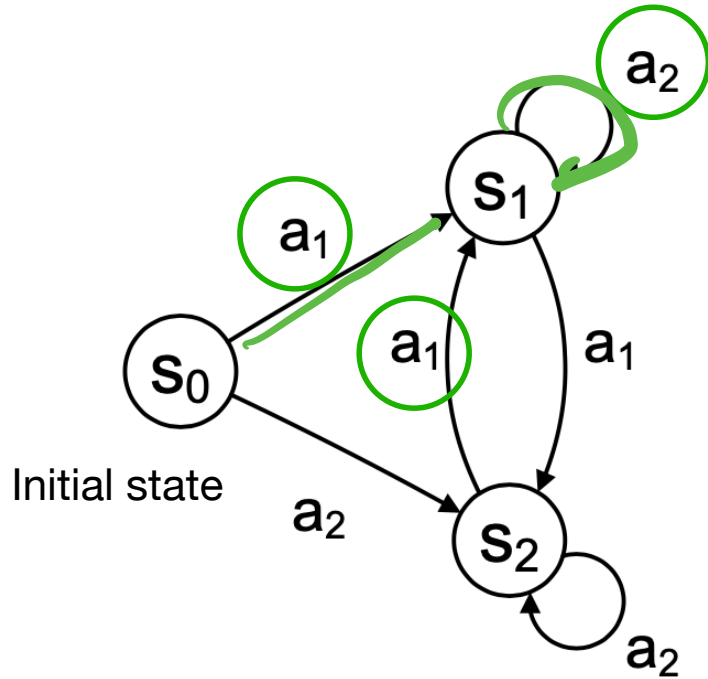
Distribution Shift: Example



Distribution Shift: Example

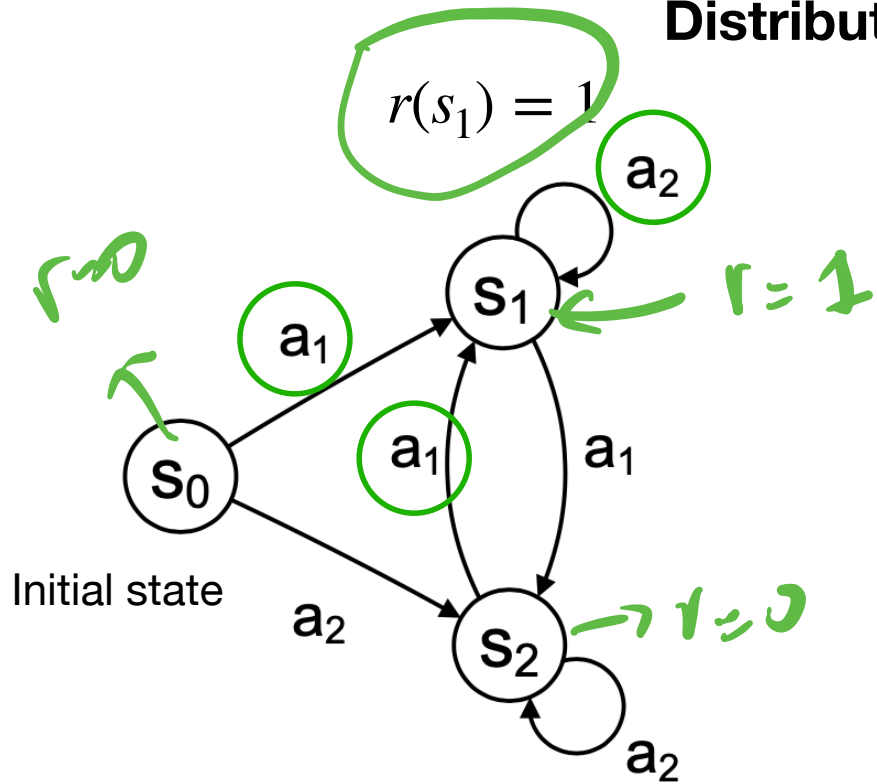


Distribution Shift: Example



$$\underline{d_{s_0}^{\pi^*}(s_0) = 1 - \gamma}, \quad \underline{d_{s_0}^{\pi^*}(s_1) = \gamma}, \quad \underline{d_{s_0}^{\pi^*}(s_2) = 0}$$

Distribution Shift: Example

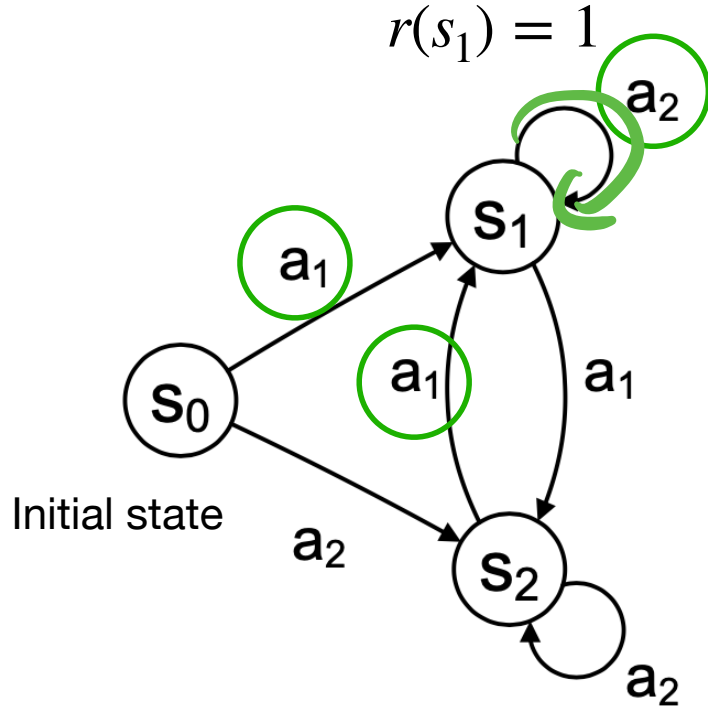


$$d_{s_0}^{\pi^*}(s_0) = 1 - \gamma, d_{s_0}^{\pi^*}(s_1) = \gamma, d_{s_0}^{\pi^*}(s_2) = 0$$

$$s \sim d_{s_0}^{\pi^*}$$

$$a \sim \pi^*(\cdot | s)$$

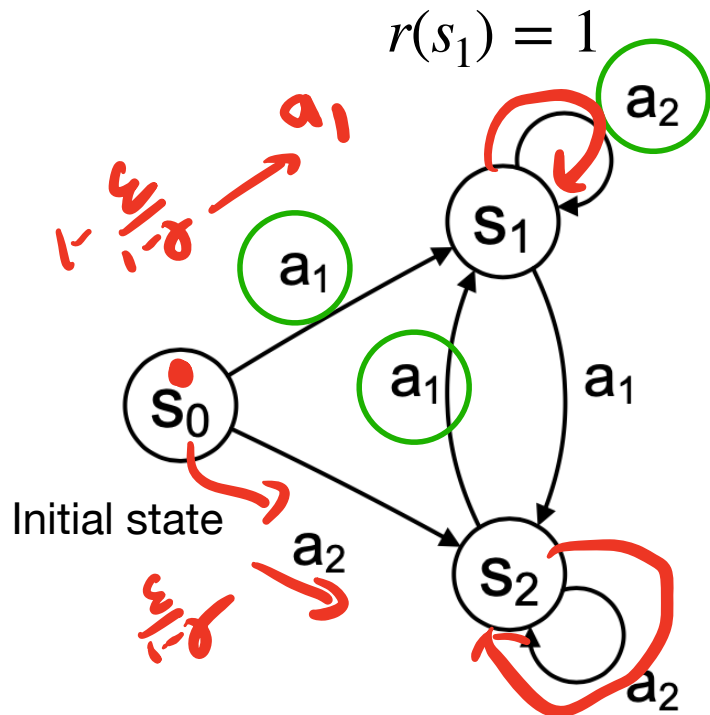
Distribution Shift: Example



$$d_{s_0}^{\pi^*}(s_0) = 1 - \gamma, d_{s_0}^{\pi^*}(s_1) = \gamma, d_{s_0}^{\pi^*}(s_2) = 0$$

$$V_{s_0}^{\pi^*} = \frac{\gamma}{1 - \gamma}$$

Distribution Shift: Example



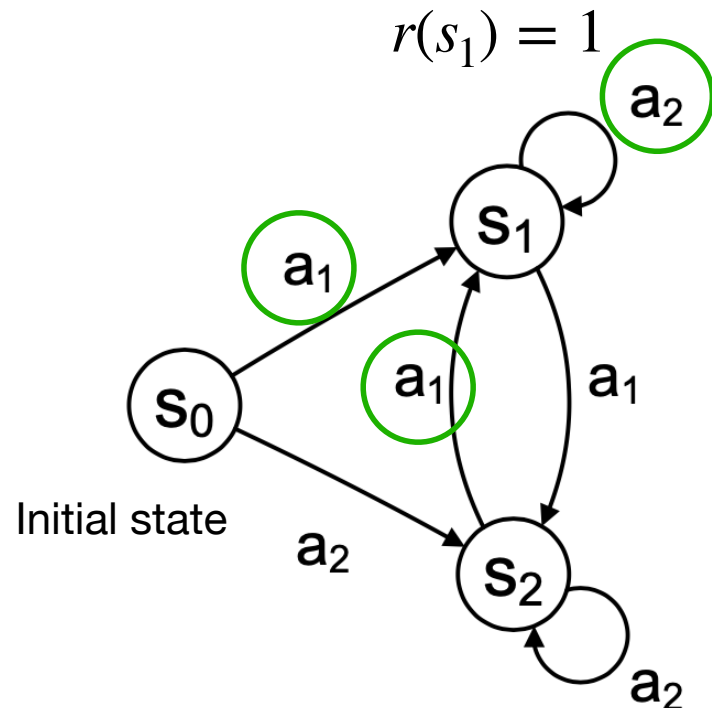
$$d_{s_0}^{\pi^*}(s_0) = 1 - \gamma, d_{s_0}^{\pi^*}(s_1) = \gamma, d_{s_0}^{\pi^*}(s_2) = 0$$

$$V_{s_0}^{\pi^*} = \frac{\gamma}{1 - \gamma}$$

Assume SL returned such policy $\hat{\pi}$

$$\hat{\pi}(s_0) = \begin{cases} a_1 & \text{w/ prob } 1 - \epsilon / (1 - \gamma) \\ a_2 & \text{w/ prob } \epsilon / (1 - \gamma) \end{cases}, \quad \hat{\pi}(s_1) = a_2, \hat{\pi}(s_2) = a_2$$

Distribution Shift: Example



$$d_{s_0}^{\pi^*}(s_0) = 1 - \gamma, d_{s_0}^{\pi^*}(s_1) = \gamma, d_{s_0}^{\pi^*}(s_2) = 0$$

$$V_{s_0}^{\pi^*} = \frac{\gamma}{1 - \gamma}$$

Assume SL returned such policy $\hat{\pi}$

$$\hat{\pi}(s_0) = \begin{cases} a_1 & \text{w/ prob } 1 - \epsilon/(1 - \gamma) \\ a_2 & \text{w/ prob } \epsilon/(1 - \gamma) \end{cases}, \quad \hat{\pi}(s_1) = a_2, \hat{\pi}(s_2) = a_2$$

We will have good supervised learning error:

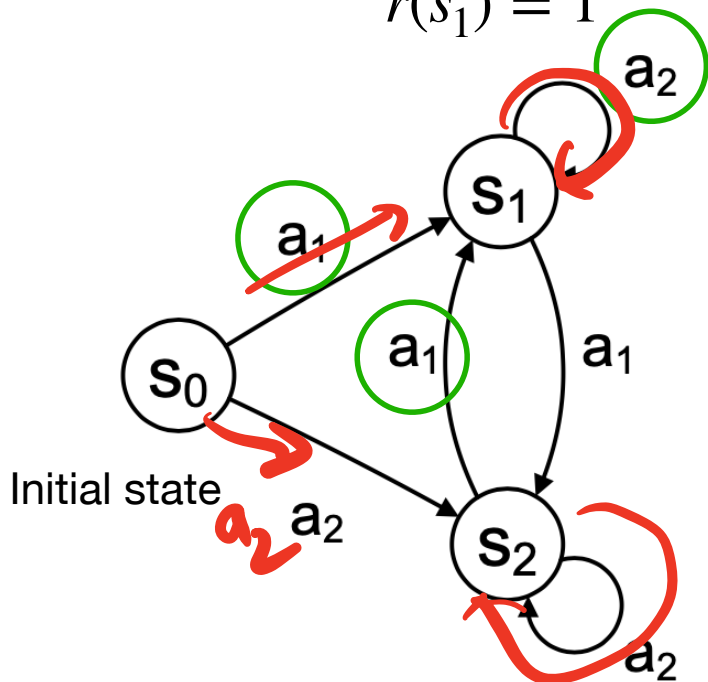
$$\hat{\pi} \leftrightarrow V^{\pi^*}$$

$$\mathbb{E}_{s \sim d_{s_0}^{\pi^*}} \mathbb{E}_{a \sim \hat{\pi}(\cdot|s)} \mathbf{1}(a \neq \pi^*(s)) = \epsilon$$

Training Dis $d_{s_0}^{\pi^*}$

Distribution Shift: Example

$$r(s_1) = 1$$



$$d_{s_0}^{\pi^*}(s_0) = 1 - \gamma, d_{s_0}^{\pi^*}(s_1) = \gamma, d_{s_0}^{\pi^*}(s_2) = 0$$

$$V_{s_0}^{\pi^*} = \frac{\gamma}{1 - \gamma}$$

V^{π^*}

Assume SL returned such policy $\hat{\pi}$

$$\hat{\pi}(s_0) = \begin{cases} a_1 & \text{w/ prob } 1 - \epsilon/(1 - \gamma) \\ a_2 & \text{w/ prob } \epsilon/(1 - \gamma) \end{cases}, \quad \hat{\pi}(s_1) = a_2, \hat{\pi}(s_2) = a_2$$

We will have good supervised learning error:

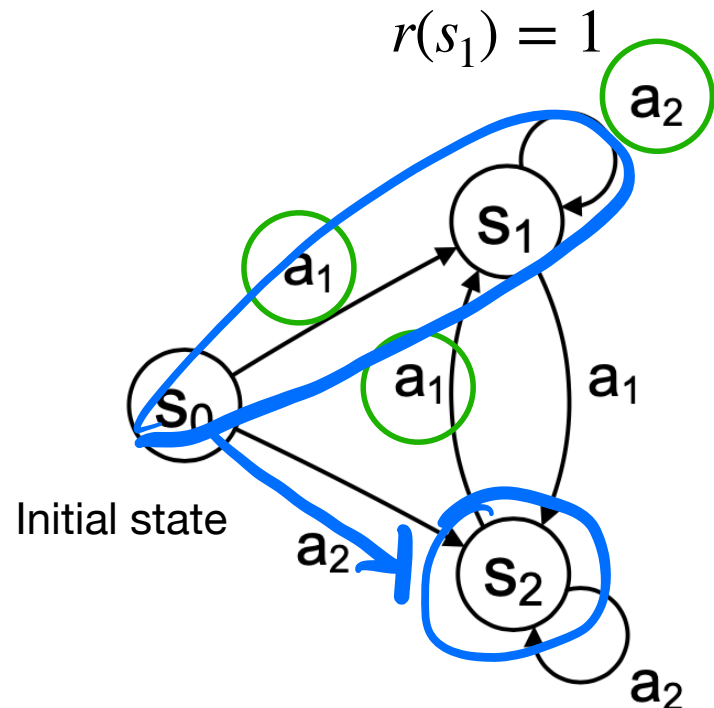
$$\mathbb{E}_{s \sim d_{s_0}^{\pi^*}} \mathbb{E}_{a \sim \hat{\pi}(\cdot|s)} \mathbf{1}(a \neq \pi^*(s)) = \epsilon$$

But we have quadratic error in performance:

$$V_{s_0}^{\hat{\pi}} = \frac{\gamma}{1 - \gamma} - \frac{\epsilon\gamma}{(1 - \gamma)^2} = V_{s_0}^{\pi^*} - \frac{\epsilon\gamma}{(1 - \gamma)^2}$$

$V^{\hat{\pi}} = \left(1 - \frac{\epsilon}{1 - \gamma}\right) \frac{\gamma}{1 - \gamma}$

Distribution Shift: Example



$$d_{s_0}^{\pi^*}(s_0) = 1 - \gamma, d_{s_0}^{\pi^*}(s_1) = \gamma, d_{s_0}^{\pi^*}(s_2) = 0$$

$$V_{s_0}^{\pi^*} = \frac{\gamma}{1 - \gamma}$$

Assume SL returned such policy $\hat{\pi}$

$$\hat{\pi}(s_0) = \begin{cases} a_1 & \text{w/ prob } 1 - \epsilon/(1 - \gamma) \\ a_2 & \text{w/ prob } \epsilon/(1 - \gamma) \end{cases}, \quad \hat{\pi}(s_1) = a_2, \hat{\pi}(s_2) = a_2$$

We will have good supervised learning error:

$$\mathbb{E}_{s \sim d_{s_0}^{\pi^*}} \mathbb{E}_{a \sim \hat{\pi}(\cdot|s)} \mathbf{1}(a \neq \pi^*(s)) = \epsilon$$

But we have quadratic error in performance:

$$V_{s_0}^{\hat{\pi}} = \frac{\gamma}{1 - \gamma} - \frac{\epsilon\gamma}{(1 - \gamma)^2} = V_{s_0}^{\pi^*} - \frac{\epsilon\gamma}{(1 - \gamma)^2}$$

Issue: once we make a mistake at s_0 , we end up in s_2 which is not in the training data!

$$\frac{\epsilon \cdot \gamma}{(1 - \gamma)^2}$$

An Autonomous Land Vehicle In A Neural Network *[Pomerleau, NIPS '88]*



“If the network is not presented with sufficient variability in its training exemplars to cover the conditions it is likely to encounter...[it] will perform poorly”

An Autonomous Land Vehicle In A Neural Network *[Pomerleau, NIPS '88]*



“If the network is not presented with sufficient variability in its training exemplars to cover the conditions it is likely to encounter...[it] will perform poorly”

Question for today:

How to mitigate the distribution shift issue?

Solution:

Interactive Imitation Learning Setting

Solution:

Interactive Imitation Learning Setting

Key assumption:

we can query expert $\pi^\star$ at any time and any state during training

Solution:

Interactive Imitation Learning Setting

Key assumption:

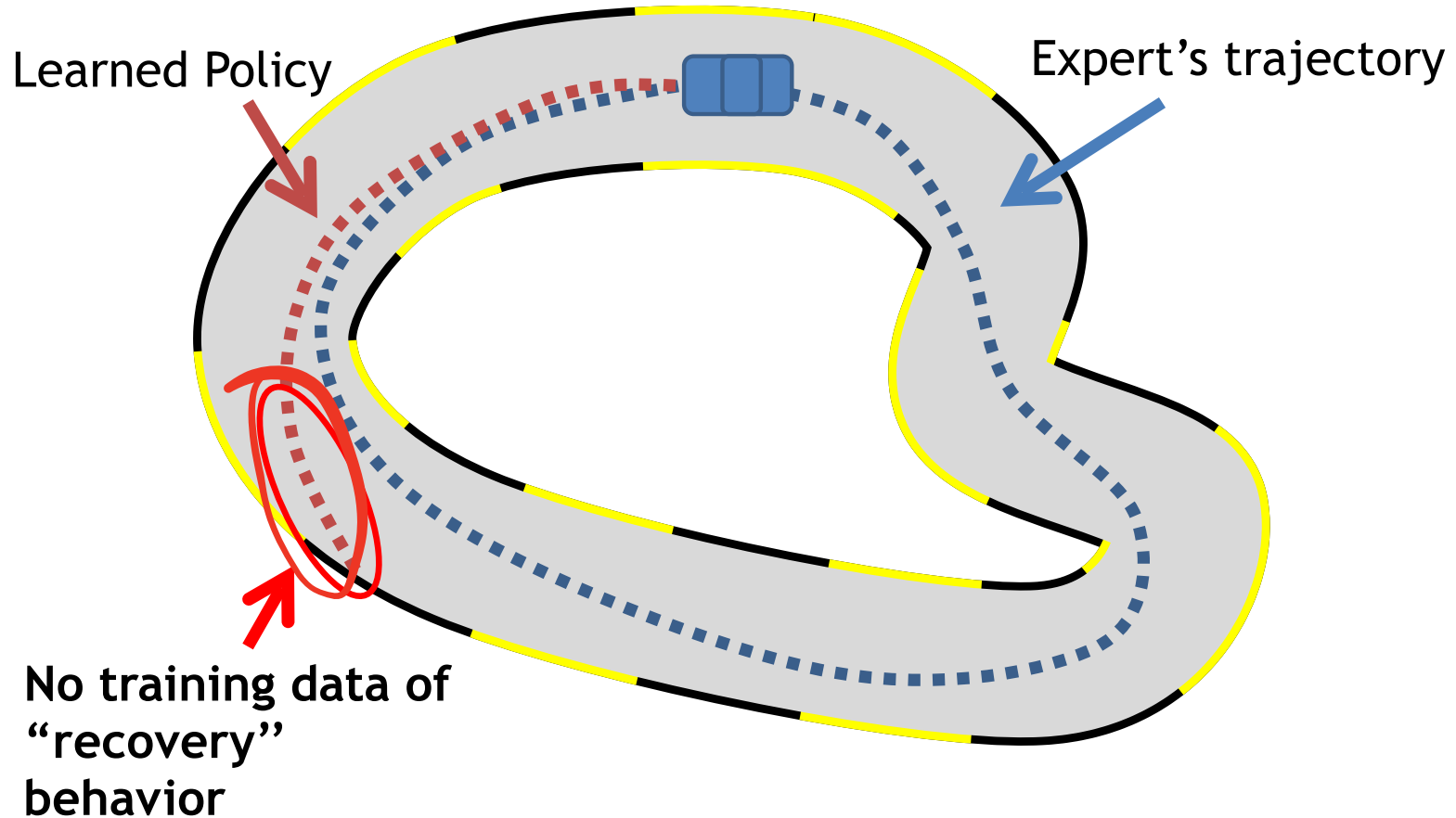
we can query expert $\pi^\star$ at any time and any state during training

(Recall that previously we only had an offline dataset $\mathcal{D} = (s_i^\star, a_i^\star)_{i=1}^M \sim d_\mu^{\pi^\star}$)

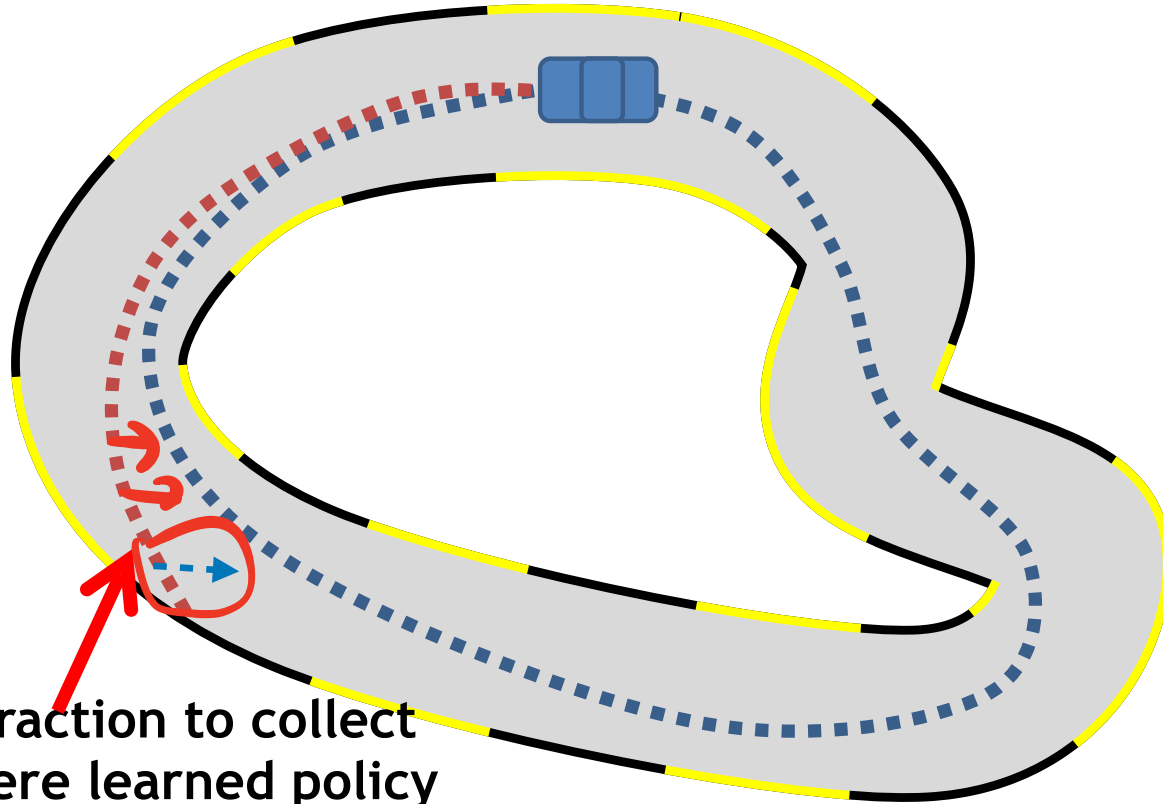
Outline for today:

1. The DAgger (Data Aggregation) Algorithm
2. Analysis of DAgger: DAgger as online learning

Recall the Main Problem from Behavior Cloning:

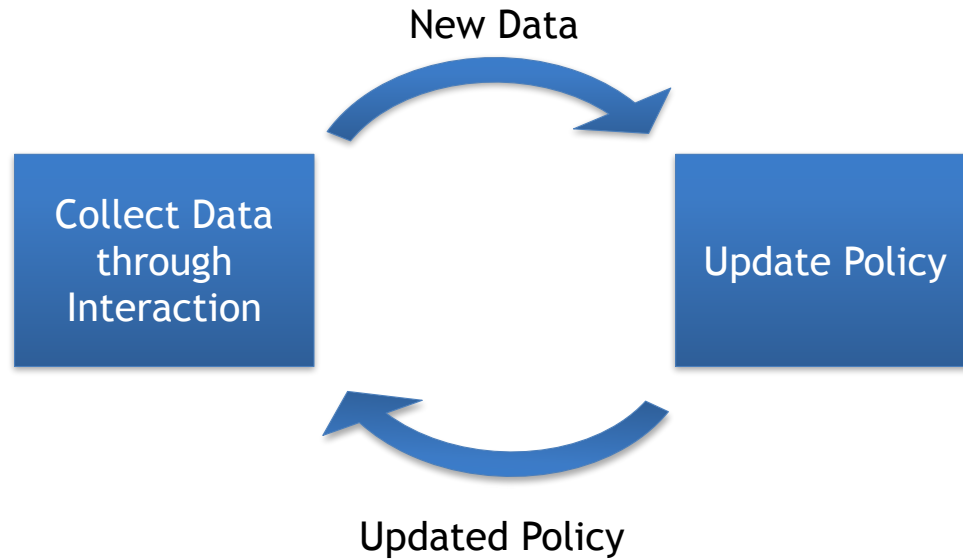


Intuitive solution: Interaction



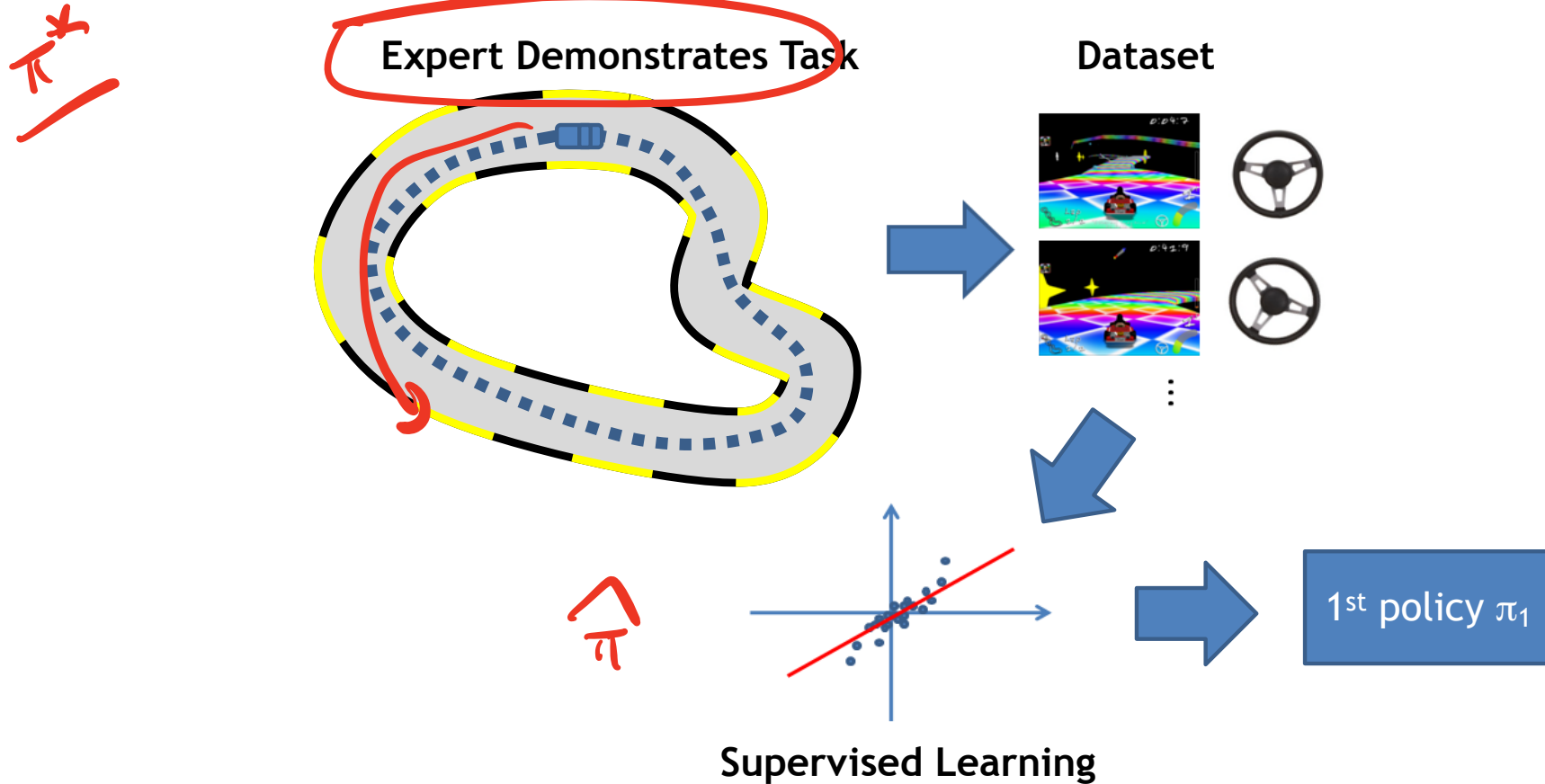
Use interaction to collect
data where learned policy
goes

General Idea: Iterative Interactive Approach



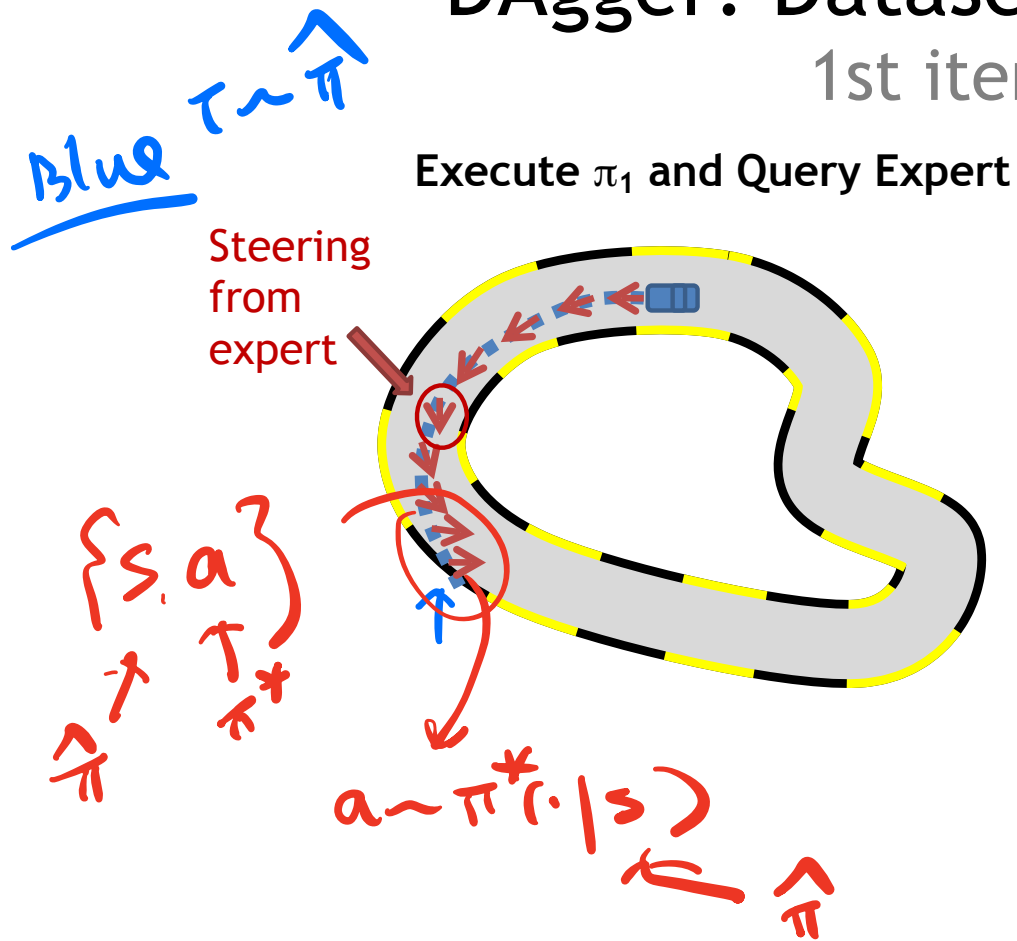
DAgger: Dataset Aggregation ^[Ross11a]

0th iteration



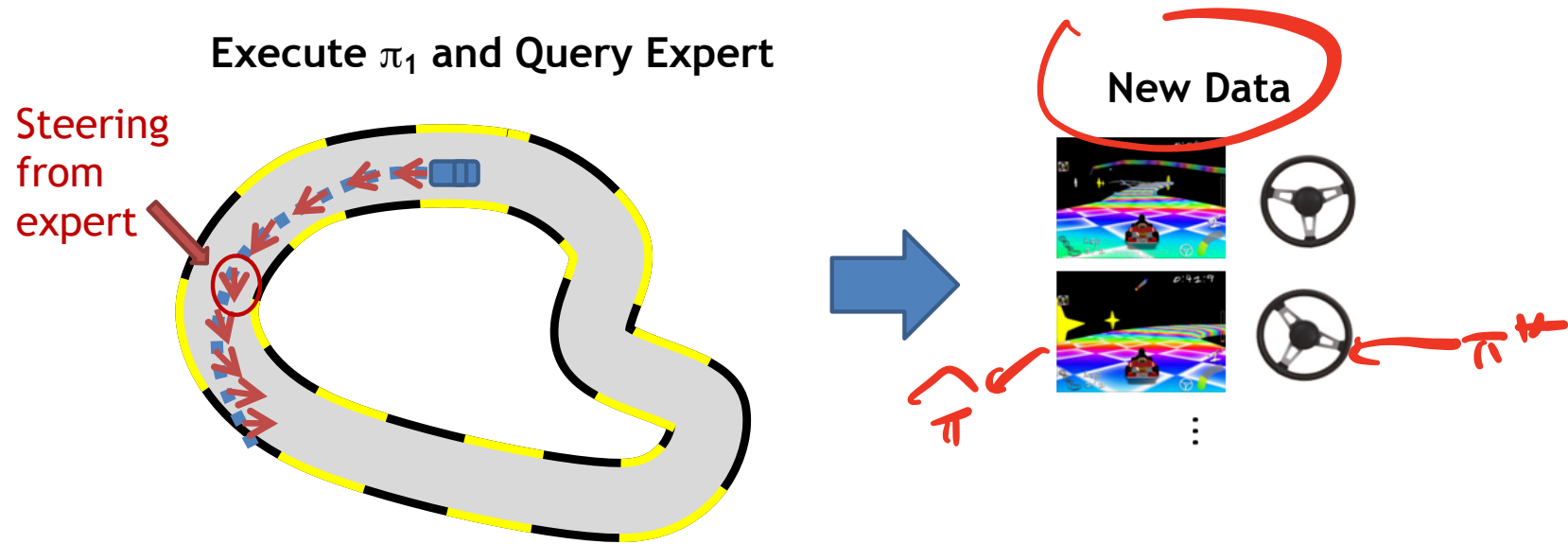
DAgger: Dataset Aggregation ^[Ross11a]

1st iteration



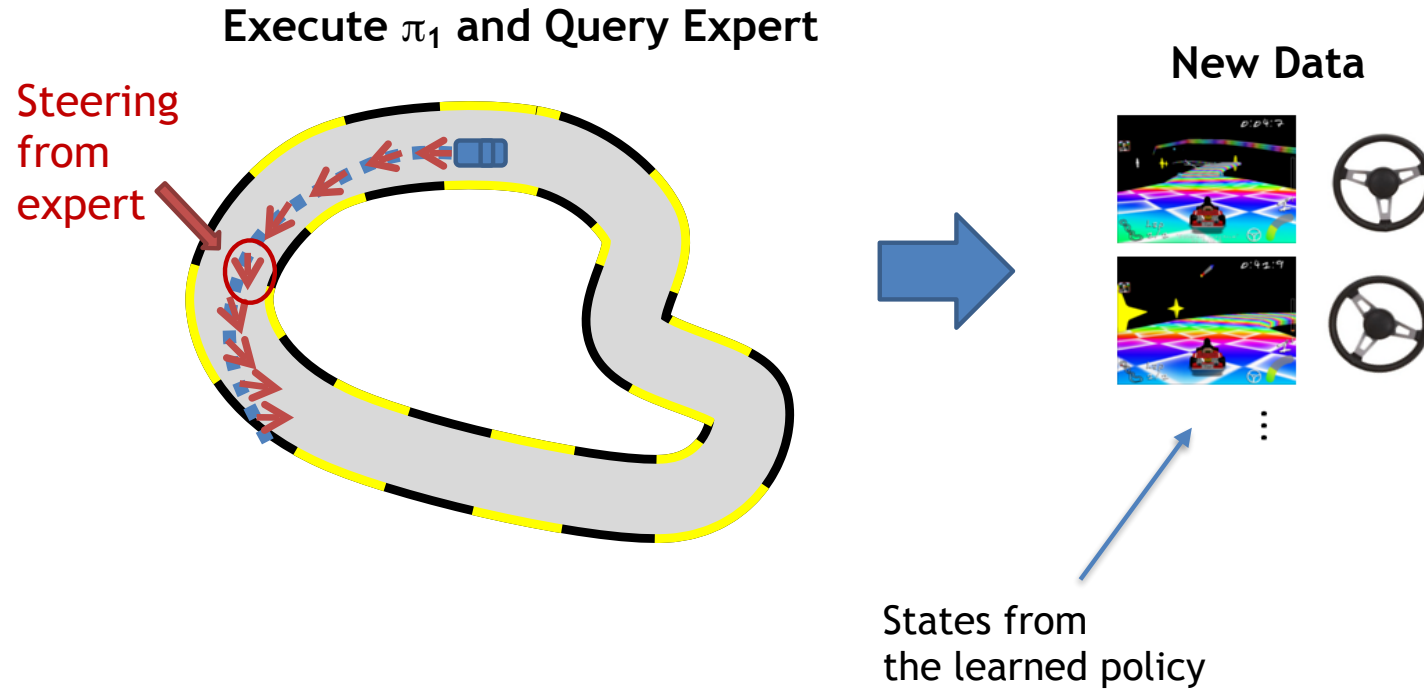
Dagger: Dataset Aggregation ^[Ross11a]

1st iteration



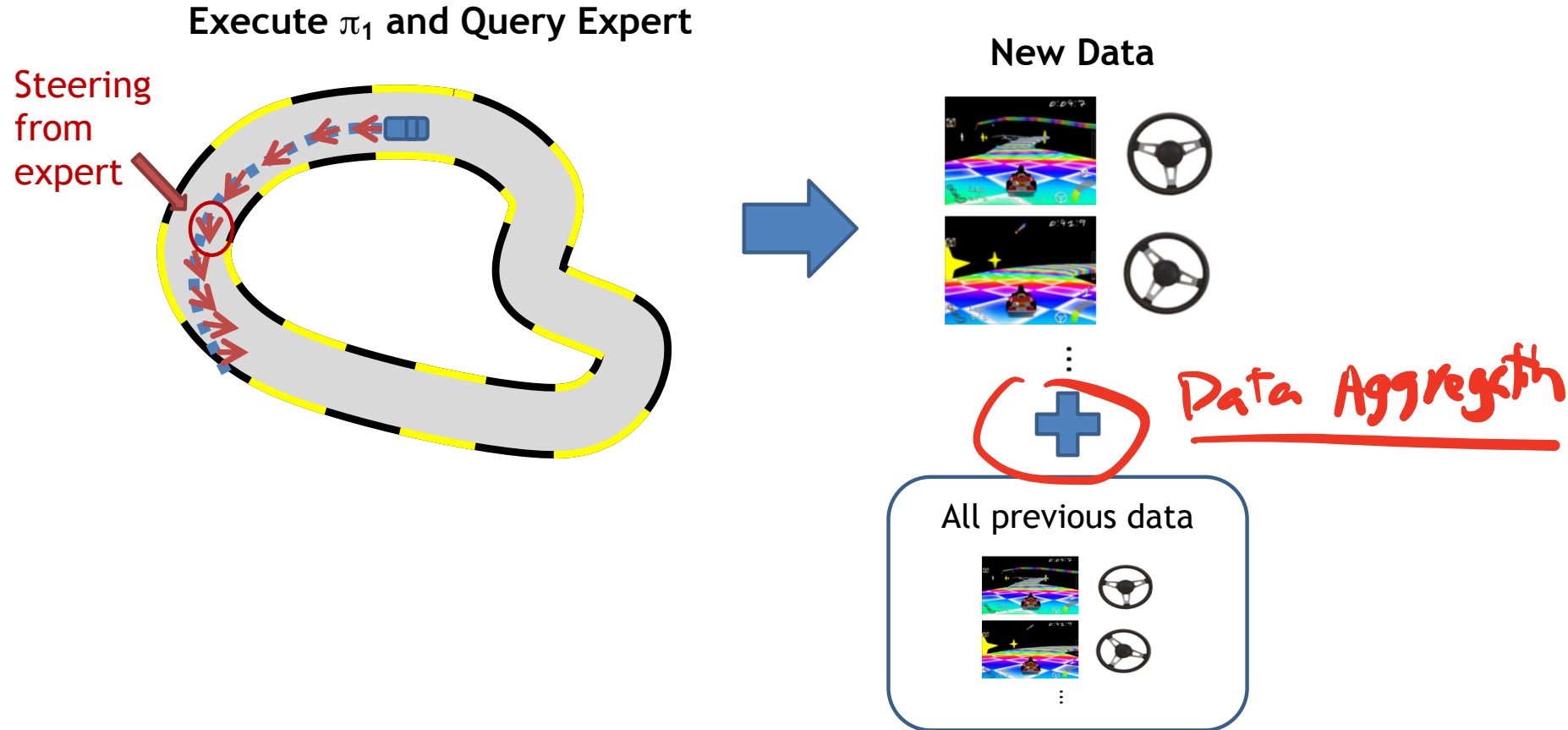
Dagger: Dataset Aggregation ^[Ross11a]

1st iteration



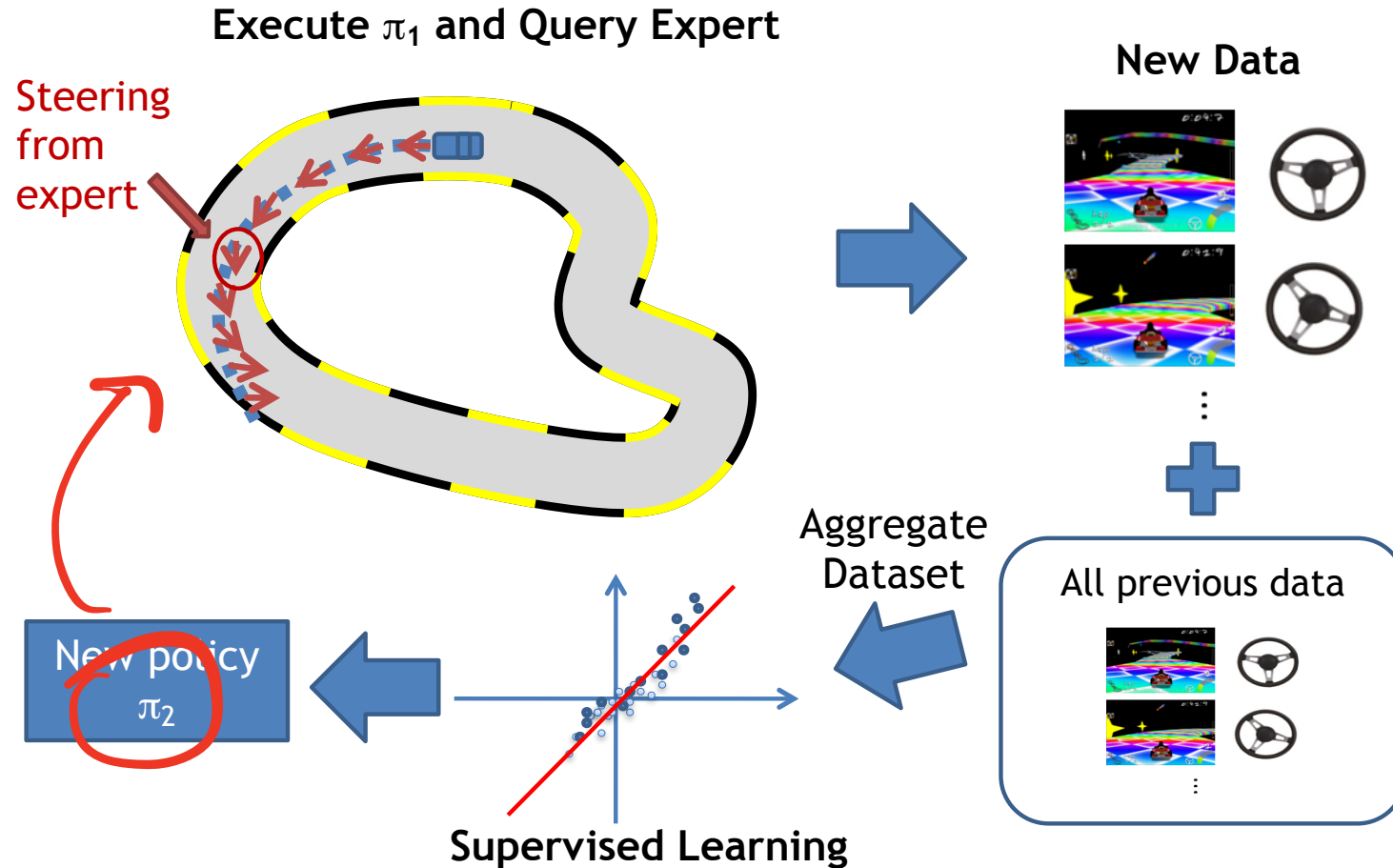
Dagger: Dataset Aggregation ^[Ross11a]

1st iteration



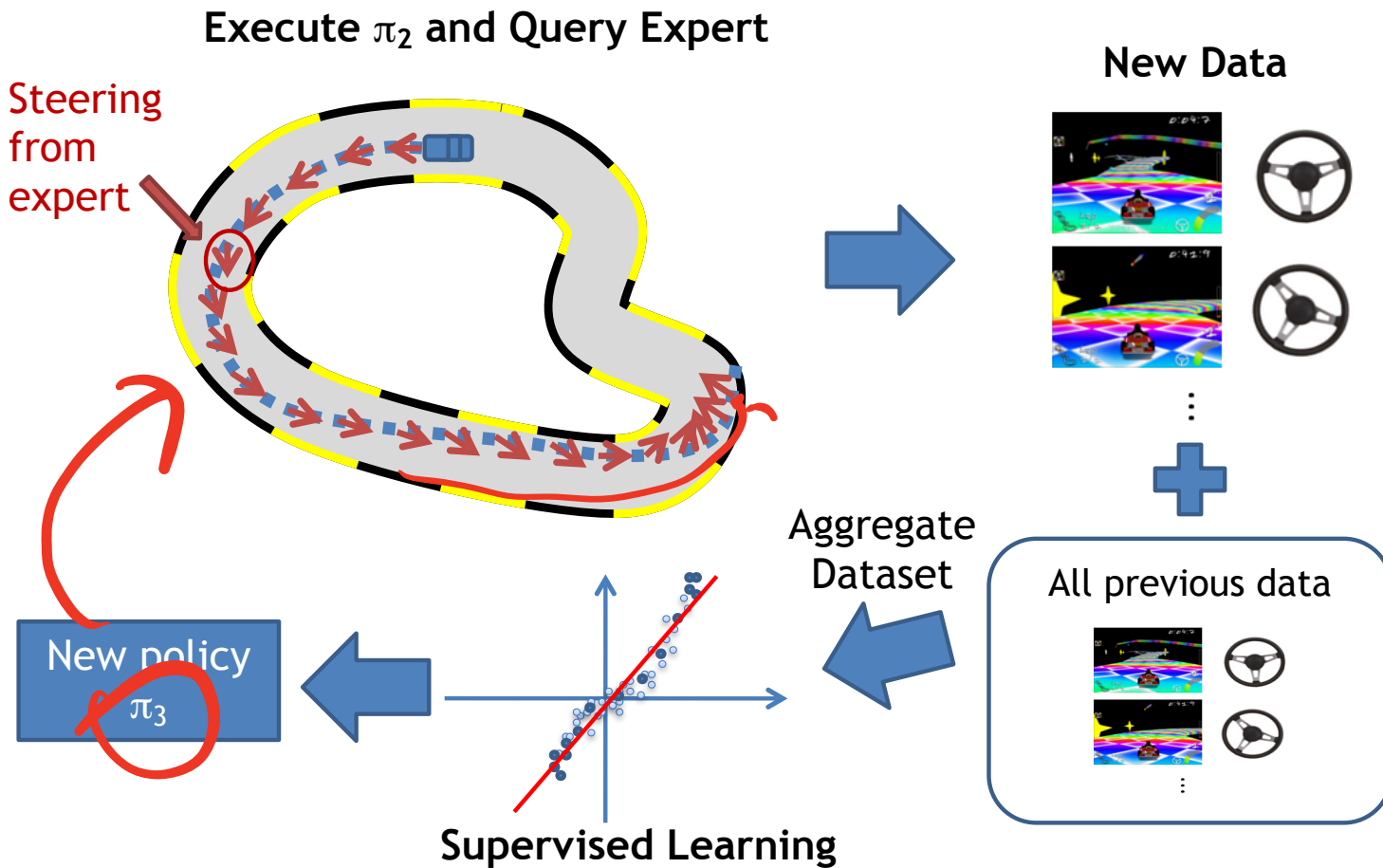
Dagger: Dataset Aggregation ^[Ross11a]

1st iteration



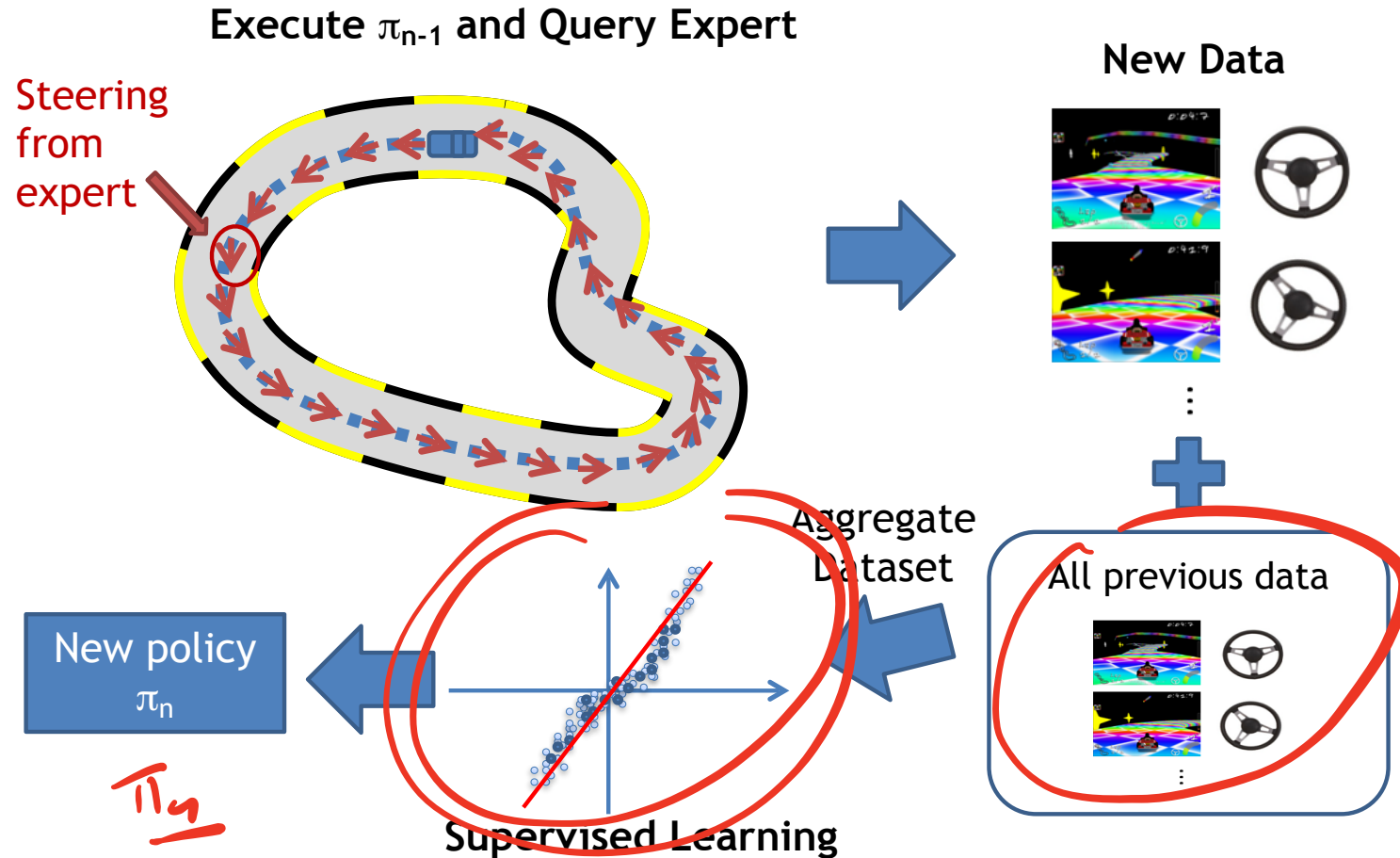
DAgger: Dataset Aggregation ^[Ross11a]

2nd iteration



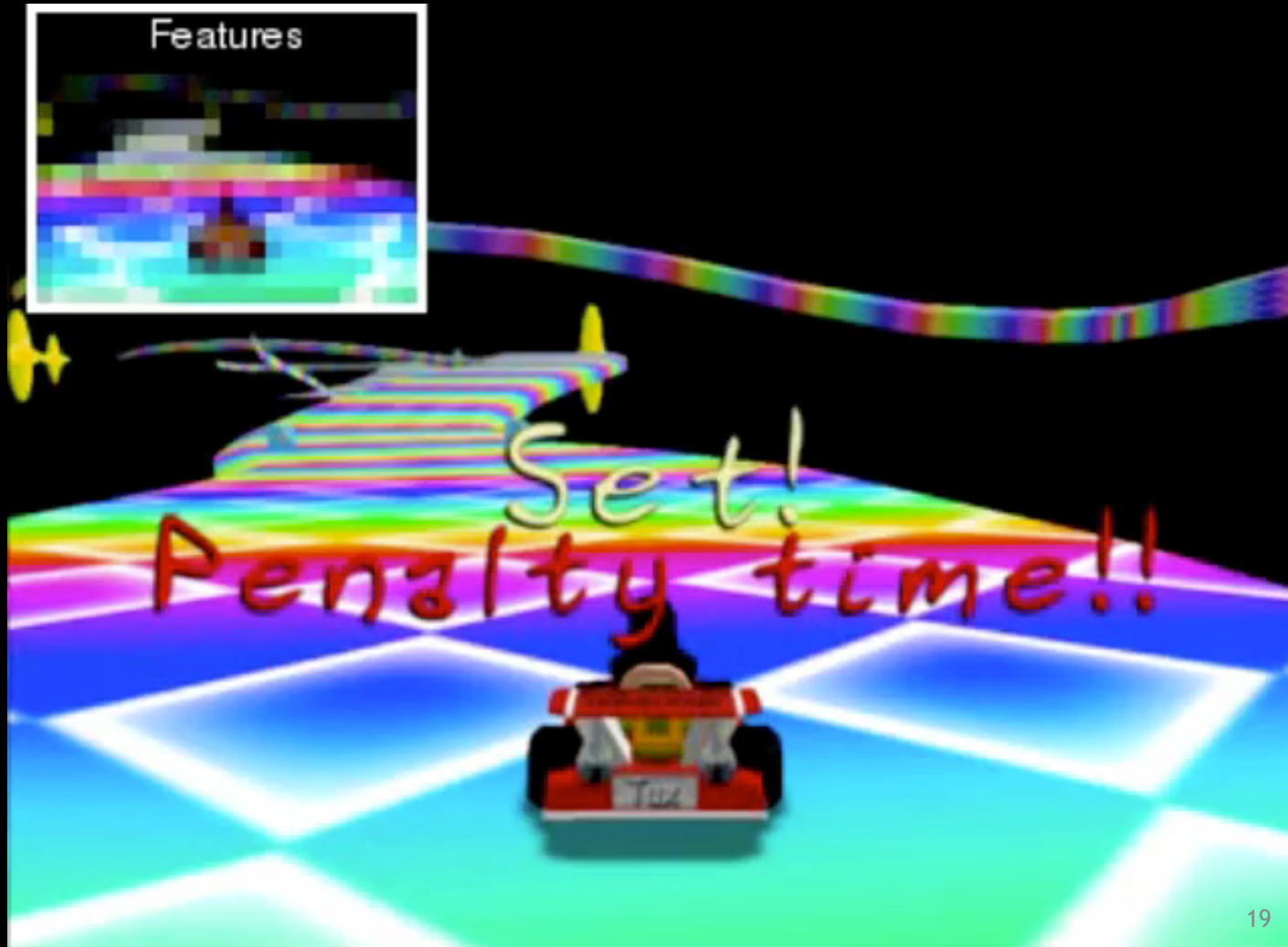
Dagger: Dataset Aggregation ^[Ross11a]

n^{th} iteration



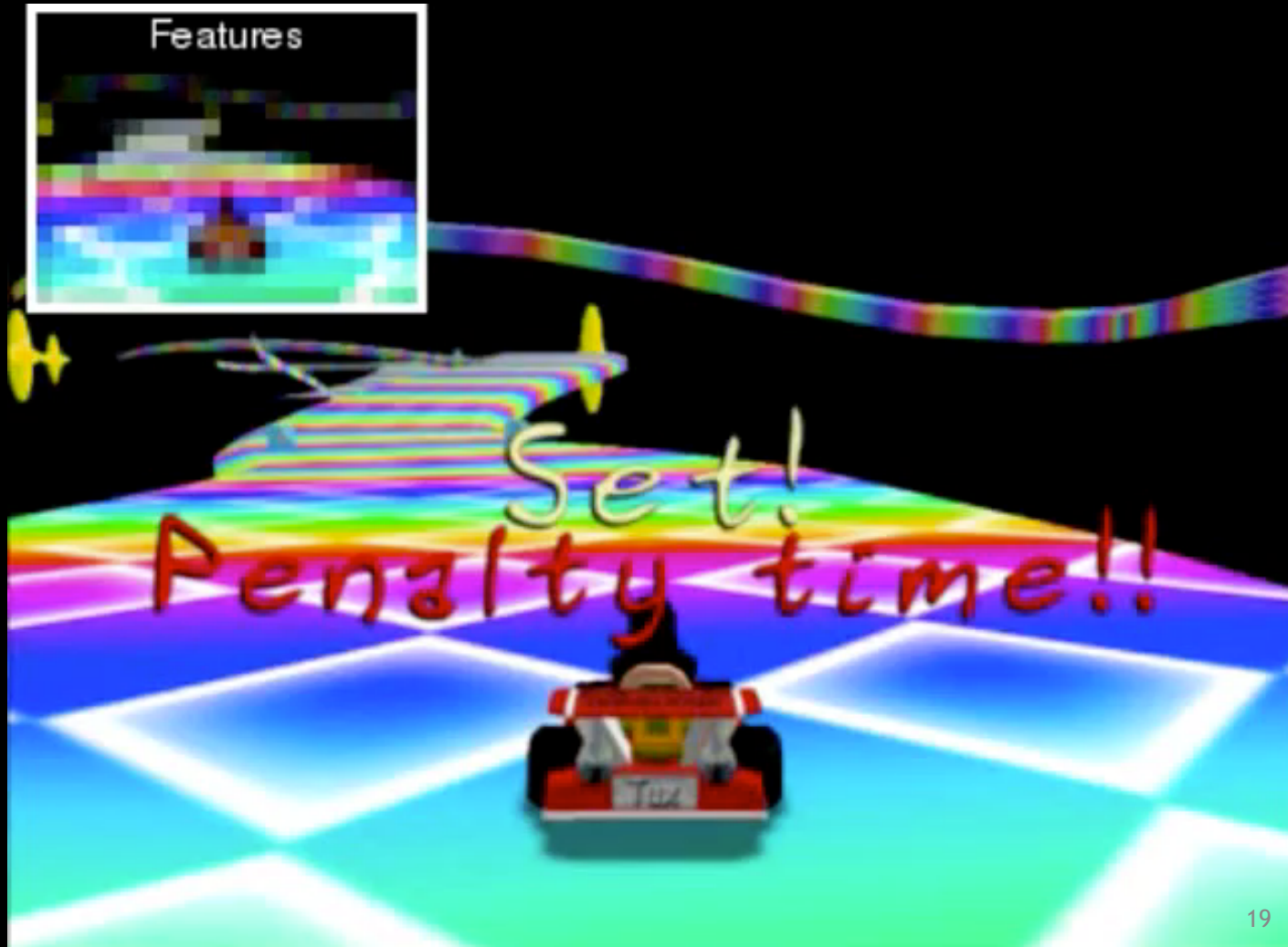
Success!

[Ross AISTATS 2011]



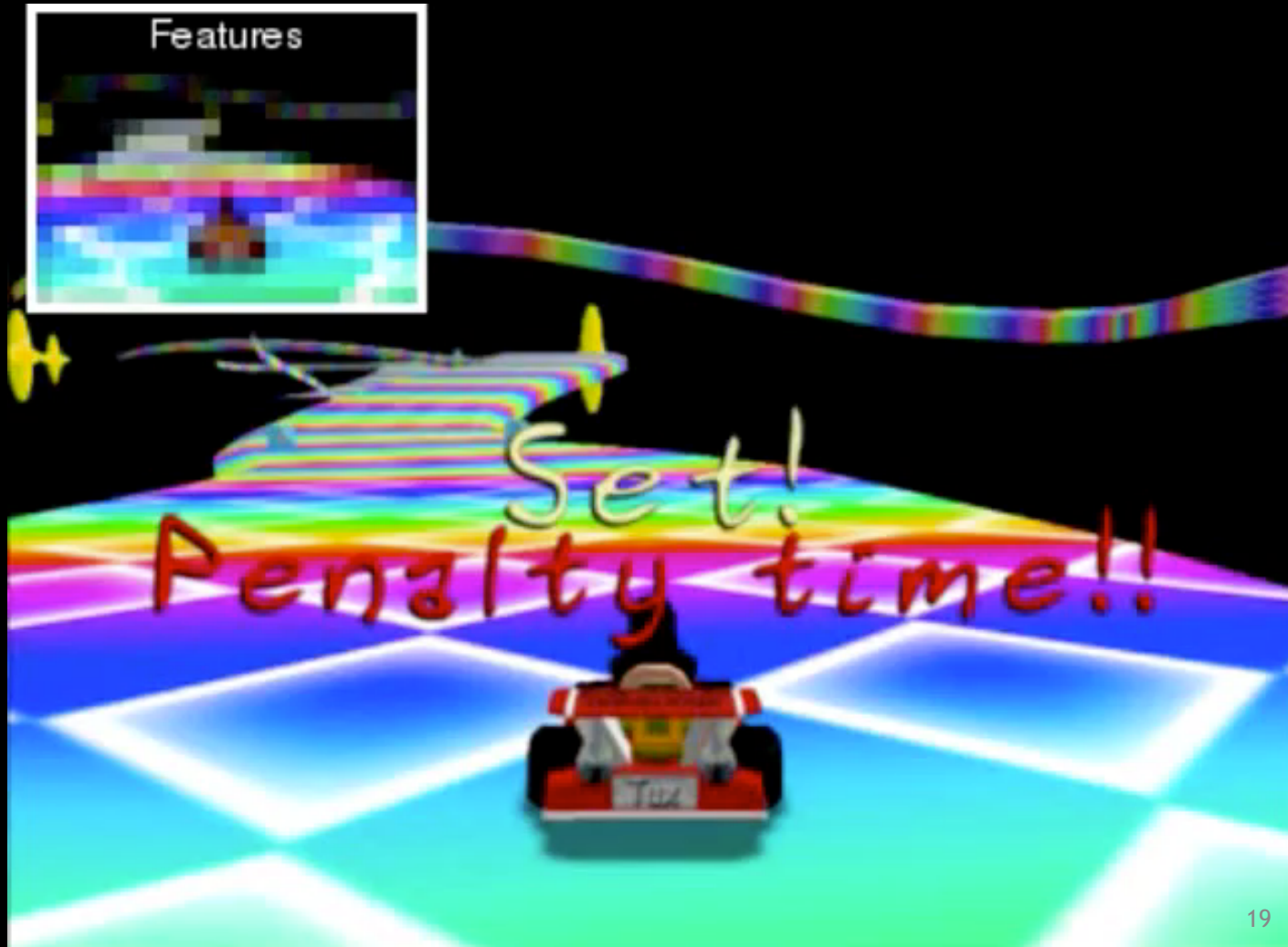
Success!

[Ross AISTATS 2011]

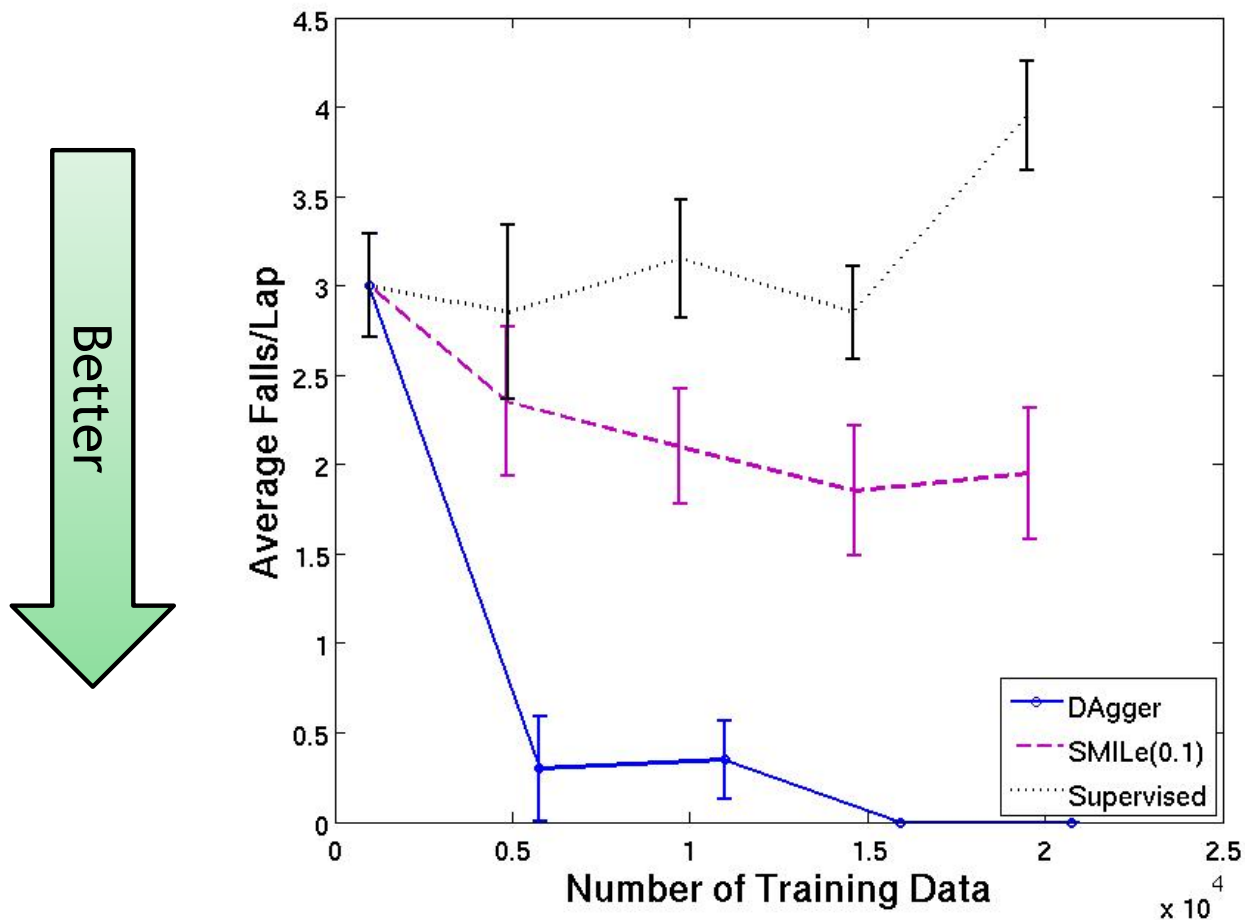


Success!

[Ross AISTATS 2011]



Average Falls/Lap



FPS: 24

Attempt: 1 of 1

AgentLinear

Selected Actions:

RIGHT

SPEED

FPS: 24

Attempt: 1 of 1

AgentLinear

Selected Actions:

RIGHT

SPEED

FPS: 24

Attempt: 1 of 1

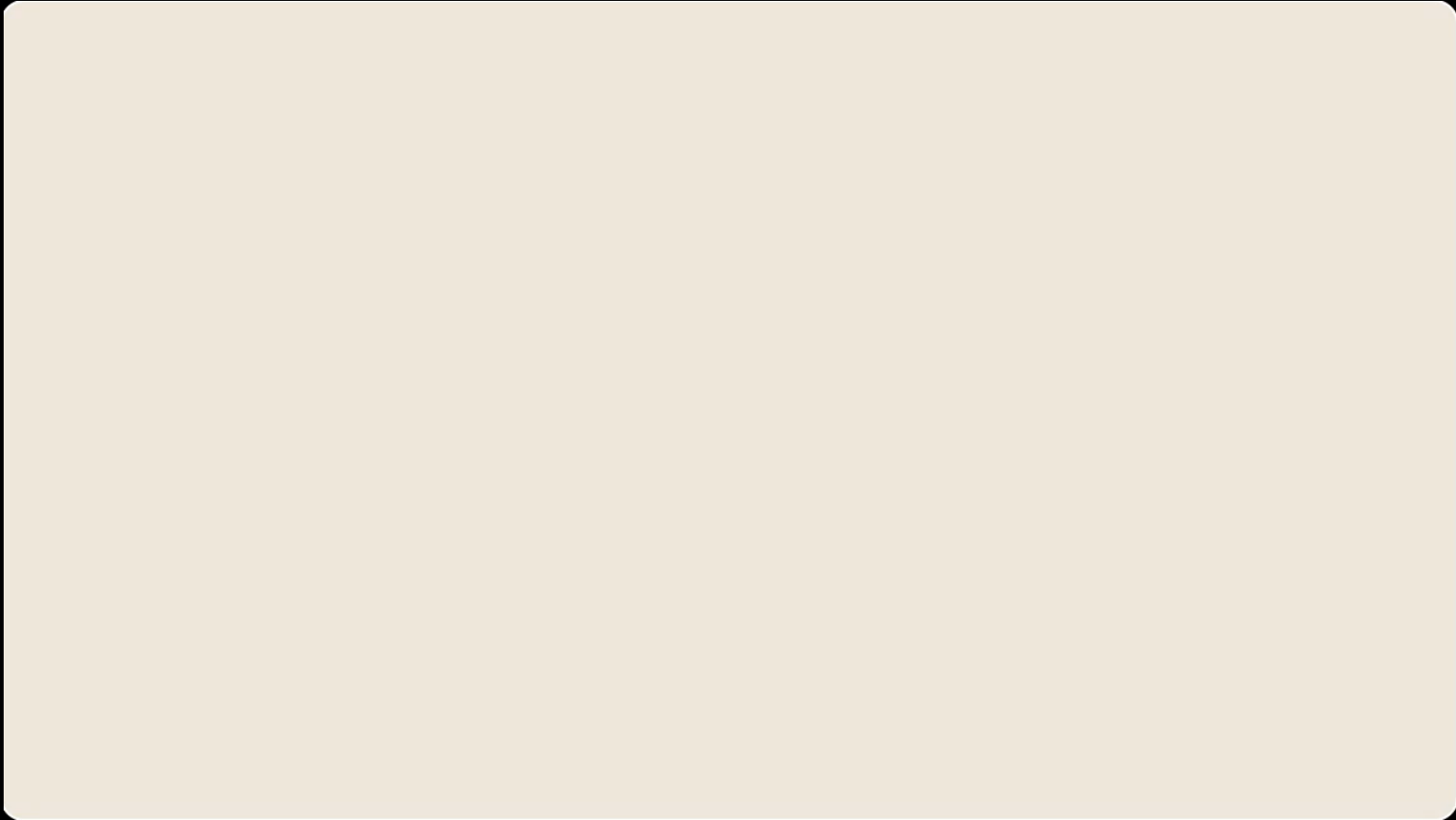
AgentLinear

Selected Actions:

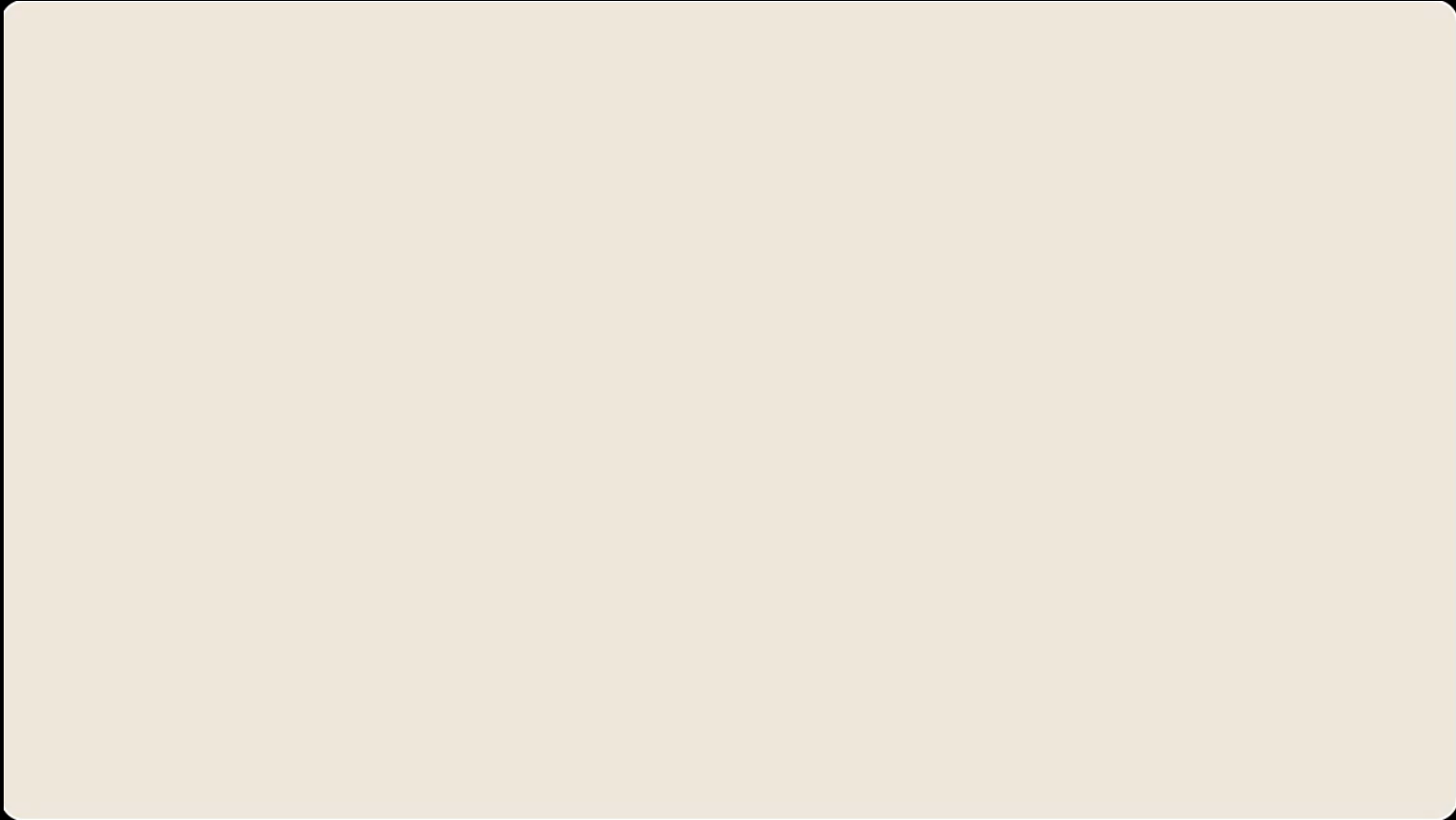
RIGHT

SPEED

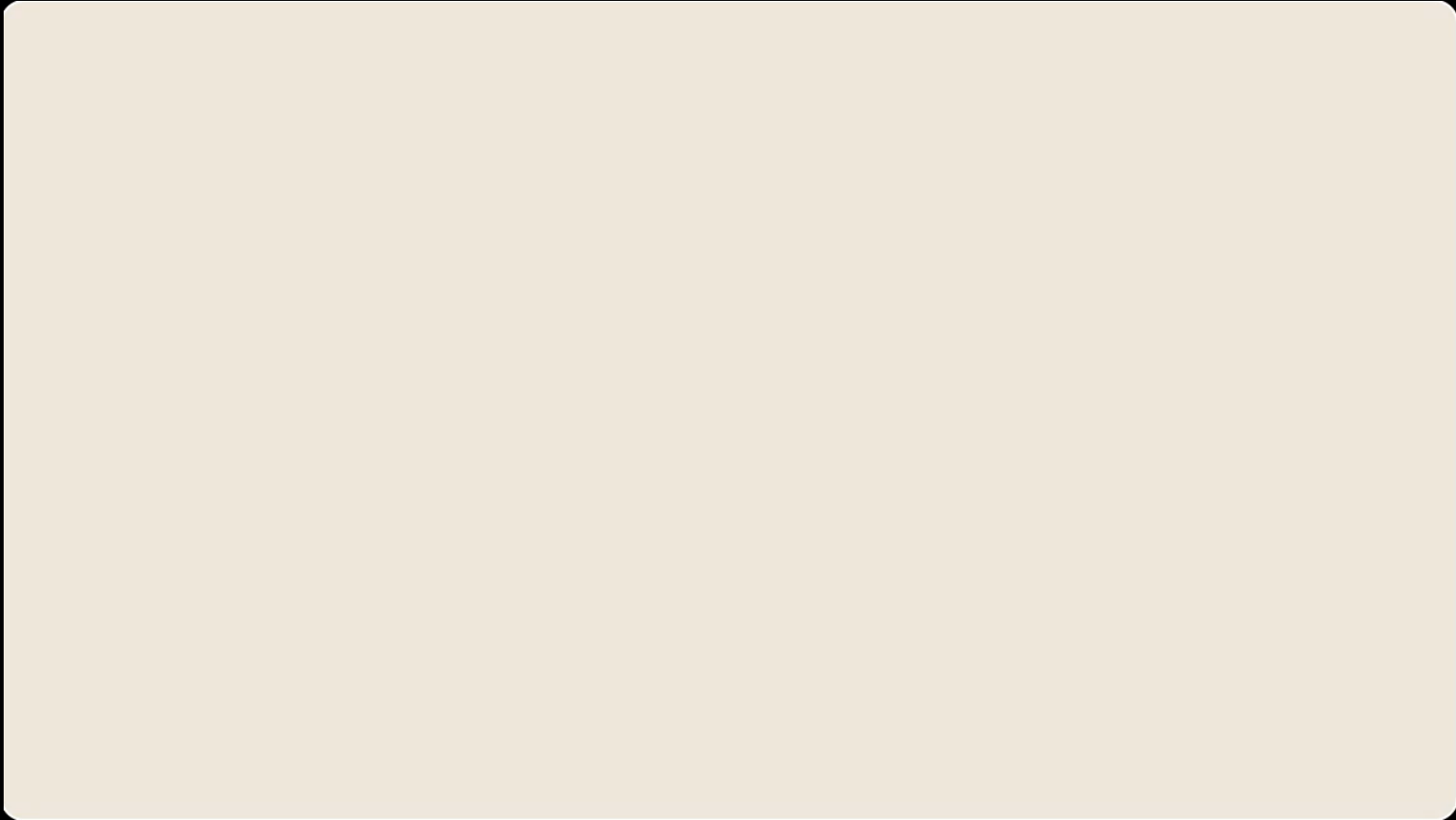
More fun than Video Games...



More fun than Video Games...



More fun than Video Games...



Forms of the Interactive Experts

Interactive Expert is expensive, especially when the expert is human...

But expert does not have to be human...

Forms of the Interactive Experts

Interactive Expert is expensive, especially when the expert is human...

But expert does not have to be human...

Example: high-speed off-road driving

[Pan et al, RSS 18, Best System Paper]



Fig. 4: The AutoRally car and the test track.

Forms of the Interactive Experts

Interactive Expert is expensive, especially when the expert is human...

But expert does not have to be human...

Example: high-speed off-road driving

[Pan et al, RSS 18, Best System Paper]



Goal: learn a racing control policy that maps from data on cheap on-board sensors (raw-pixel image) to low-level control (steer and throttle)

Fig. 4: The AutoRally car and the test track.

Forms of the Interactive Experts

Interactive Expert is expensive, especially when the expert is human...

But expert does not have to be human...

Example: high-speed off-road driving

[Pan et al, RSS 18, Best System Paper]



Fig. 4: The AutoRally car and the test track.

Goal: learn a racing control policy that maps from data on cheap on-board sensors (raw-pixel image) to low-level control (steer and throttle)



(a) raw image

→ Steering + throttle

Forms of the Interactive Experts

Example: high-speed off-road driving

[Pan et al, RSS 18, Best System Paper]

Forms of the Interactive Experts

Example: high-speed off-road driving

[Pan et al, RSS 18, Best System Paper]

Their Setup:

At Training, we have expensive sensors for accurate state estimation
and we have computation resources for **MPC** (i.e., high-frequency replanning)

Forms of the Interactive Experts

Example: high-speed off-road driving

[Pan et al, RSS 18, Best System Paper]

Their Setup:

At Training, we have expensive sensors for accurate state estimation
and we have computation resources for **MPC** (i.e., high-frequency replanning)

The MPC is the expert in this case!

Forms of the Interactive Experts

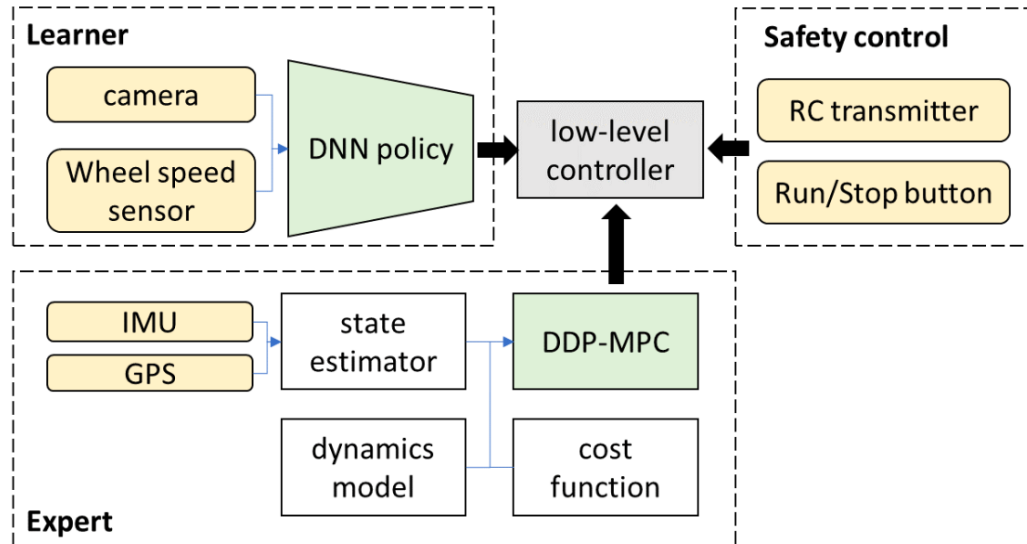
Example: high-speed off-road driving

[Pan et al, RSS 18, Best System Paper]

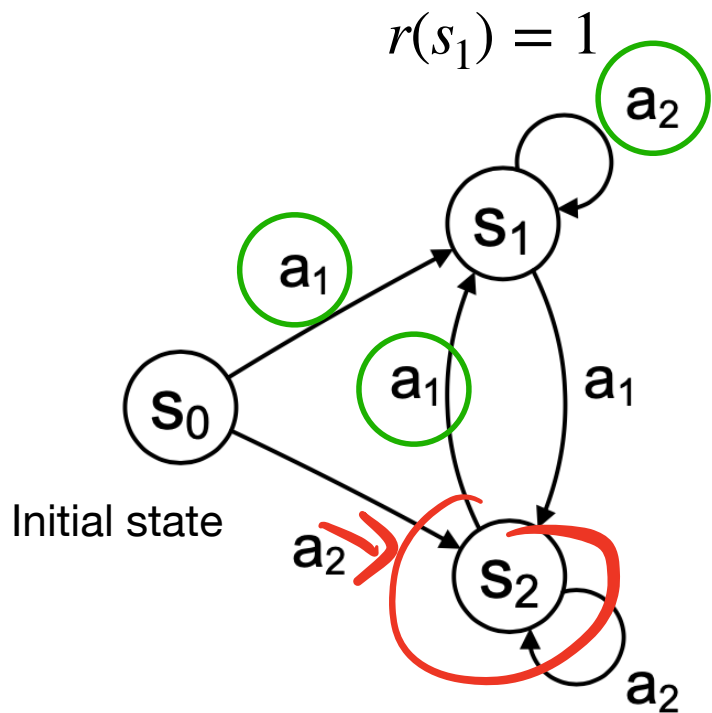
Their Setup:

At Training, we have expensive sensors for accurate state estimation and we have computation resources for **MPC** (i.e., high-frequency replanning)

The MPC is the expert in this case!



Distribution Shift: Example



$$d_{s_0}^{\pi^*}(s_0) = 1 - \gamma, d_{s_0}^{\pi^*}(s_1) = \gamma, d_{s_0}^{\pi^*}(s_2) = 0$$

$$V_{s_0}^{\pi^*} = \frac{\gamma}{1 - \gamma}$$

Assume SL returned such policy $\hat{\pi}$

$$\hat{\pi}(s_0) = \begin{cases} a_1 & \text{w/ prob } 1 - \epsilon/(1 - \gamma) \\ a_2 & \text{w/ prob } \epsilon/(1 - \gamma) \end{cases}, \quad \hat{\pi}(s_1) = a_2, \hat{\pi}(s_2) = a_2$$

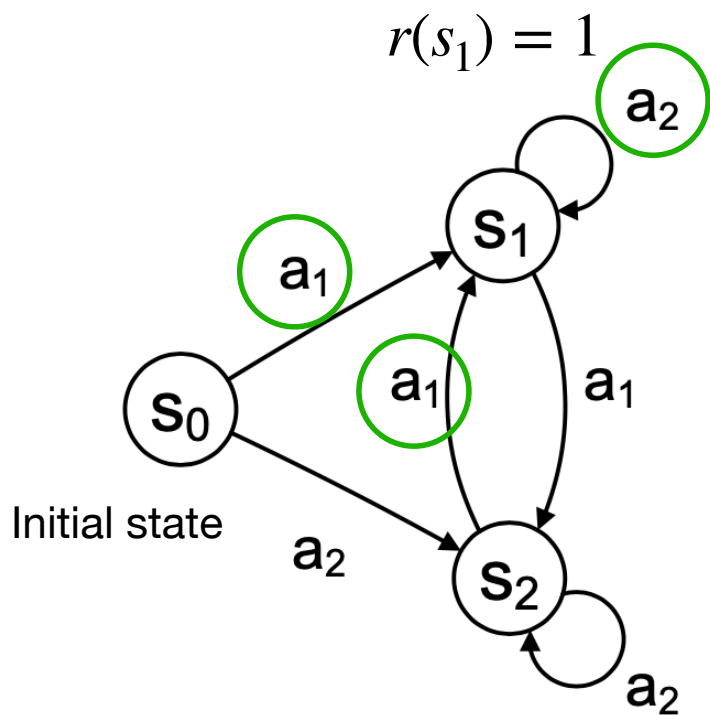
We will have good supervised learning error:

$$\mathbb{E}_{s \sim d_{s_0}^{\pi^*}} \mathbb{E}_{a \sim \hat{\pi}(\cdot|s)} \mathbf{1}(a \neq \pi^*(s)) = \epsilon$$

But we have quadratic error in performance:

$$V_{s_0}^{\hat{\pi}} = V_{s_0}^{\pi^*} - \frac{\epsilon\gamma}{(1 - \gamma)^2}$$

Distribution Shift: Example



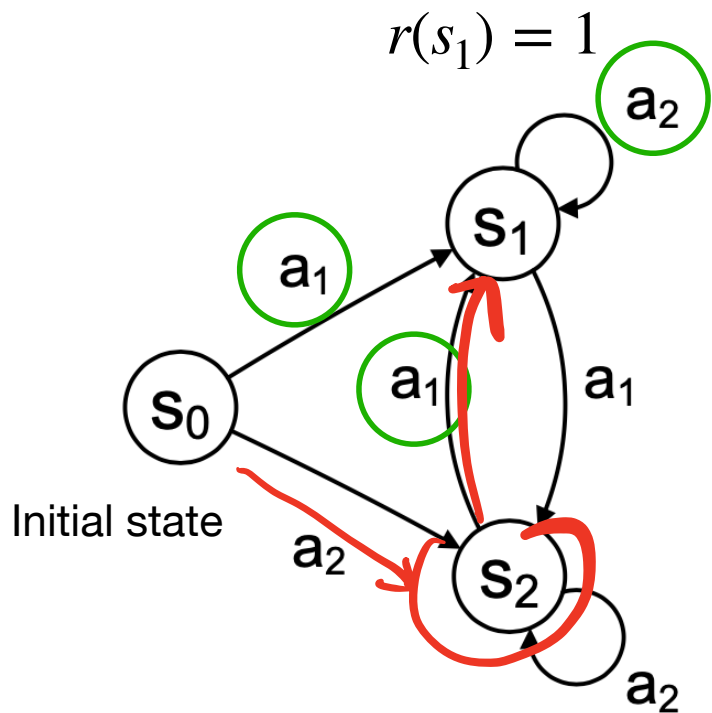
Assume SL returned such policy $\hat{\pi}$

$$\hat{\pi}(s_0) = \begin{cases} a_1 & \text{w/ prob } 1 - \epsilon/(1 - \gamma) \\ a_2 & \text{w/ prob } \epsilon/(1 - \gamma) \end{cases}, \quad \hat{\pi}(s_1) = a_2, \hat{\pi}(s_2) = a_2$$

$$d_{s_0}^{\pi^*}(s_0) = 1 - \gamma, d_{s_0}^{\pi^*}(s_1) = \gamma, d_{s_0}^{\pi^*}(s_2) = 0$$

$$V_{s_0}^{\pi^*} = \frac{\gamma}{1 - \gamma}$$

Distribution Shift: Example



Assume SL returned such policy $\hat{\pi}$

$$\hat{\pi}(s_0) = \begin{cases} a_1 & \text{w/ prob } 1 - \epsilon/(1 - \gamma) \\ a_2 & \text{w/ prob } \epsilon/(1 - \gamma) \end{cases}, \quad \hat{\pi}(s_1) = a_2, \hat{\pi}(s_2) = a_2$$

Why DAgger can fix this problem?

$$d_{s_0}^{\pi^*}(s_0) = 1 - \gamma, d_{s_0}^{\pi^*}(s_1) = \gamma, d_{s_0}^{\pi^*}(s_2) = 0$$

$$V_{s_0}^{\pi^*} = \frac{\gamma}{1 - \gamma}$$

Outline for today:

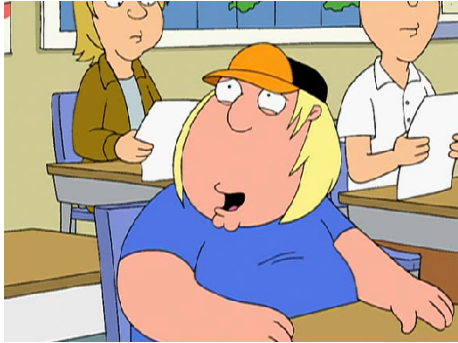


1. The DAgger (Data Aggregation) Algorithm

2. Analysis of DAgger: DAgger as online learning

Online Learning

Learner



convex Decision set Θ

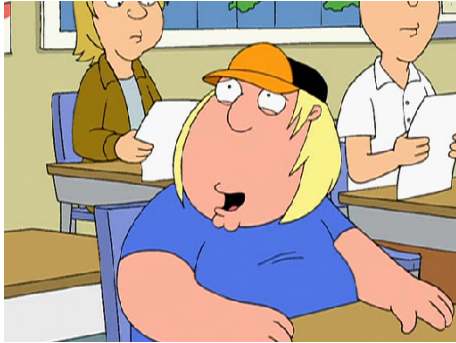
Adversary



...

Online Learning

Learner



convex Decision set Θ

Learner picks a decision θ_0



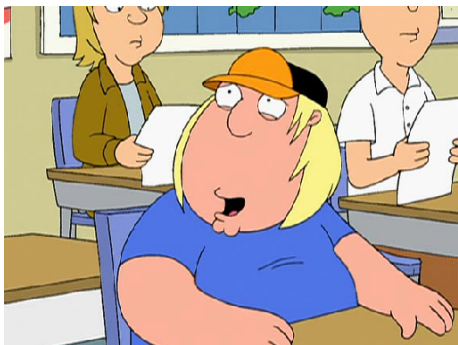
(x, y) **Adversary**



...

Online Learning

Learner



convex Decision set Θ

Learner picks a decision θ_0



Adversary picks a loss $\ell_0 : \Theta \rightarrow \mathbb{R}$



$$\ell_0(\theta_0) = (\theta_0^T x_0 - y_0)^2$$

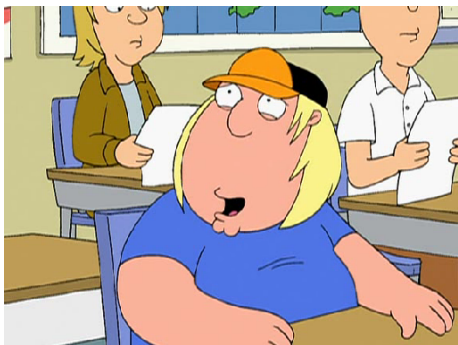
Adversary



...

Online Learning

Learner



convex Decision set Θ

Learner picks a decision θ_0



Adversary picks a loss $\ell_0 : \Theta \rightarrow \mathbb{R}$



Learner picks a new decision θ_1



Adversary

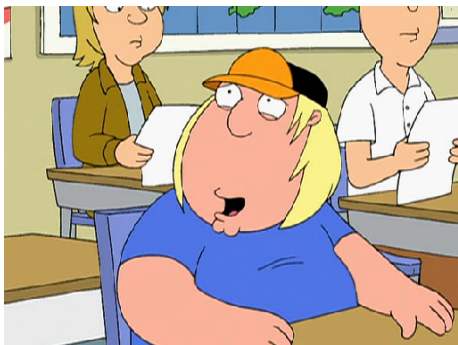


(x_1, y_1)

$\dots \ell_1(\theta_1) = (\theta_1^T x_1 - y_1)^2$

Online Learning

Learner



convex Decision set Θ

Learner picks a decision θ_0



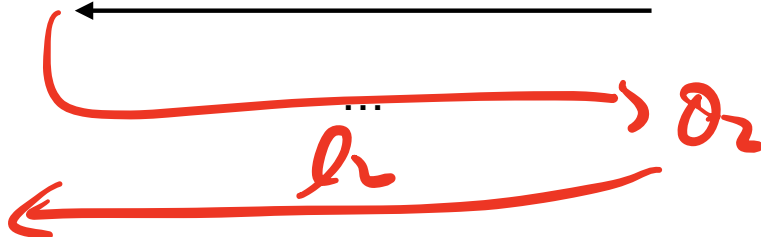
Adversary picks a loss $\ell_0 : \Theta \rightarrow \mathbb{R}$



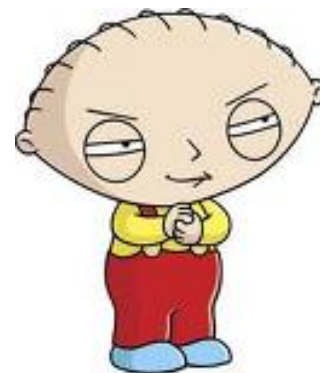
Learner picks a new decision θ_1



Adversary picks a loss $\ell_1 : \Theta \rightarrow \mathbb{R}$

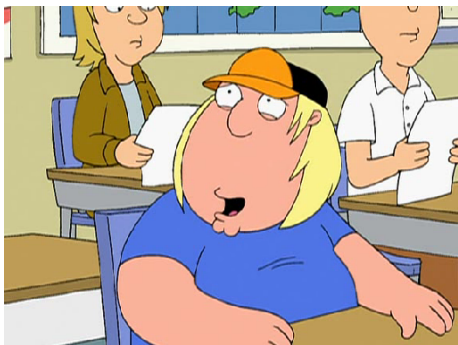


Adversary



Online Learning

Learner



convex Decision set Θ

Learner picks a decision θ_0



Adversary picks a loss $\ell_0 : \Theta \rightarrow \mathbb{R}$



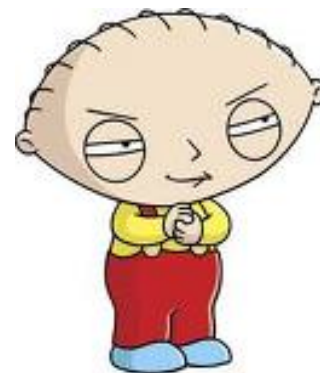
Learner picks a new decision θ_1



Adversary picks a loss $\ell_1 : \Theta \rightarrow \mathbb{R}$



Adversary



$$\text{Regret} = \sum_{t=0}^{T-1} \ell_t(\theta_t) - \min_{\theta \in \Theta} \sum_{t=0}^{T-1} \ell_t(\theta) \rightarrow \frac{\sqrt{T}}{T} \rightarrow 0 \quad T \rightarrow \infty$$

Example: online linear regression

Can we perform linear regression in online fashion with non i.i.d (or even adversary) data?

Example: online linear regression

Can we perform linear regression in online fashion with non i.i.d (or even adversary) data?

Every iteration t :

Example: online linear regression

Can we perform linear regression in online fashion with non i.i.d (or even adversary) data?

Every iteration t :

1. Learner first picks $\theta_t \in \text{Ball} \subset \mathbb{R}^d$



Example: online linear regression

Can we perform linear regression in online fashion with non i.i.d (or even adversary) data?

Every iteration t :

1. Learner first picks $\theta_t \in \text{Ball} \subset \mathbb{R}^d$
2. Adversary **then** picks $\underline{x_t} \in \mathcal{X} \subset \mathbb{R}^d, \underline{y_t} \in [a, b]$

(see θ_t)

Example: online linear regression

Can we perform linear regression in online fashion with non i.i.d (or even adversary) data?

Every iteration t :

1. Learner first picks $\theta_t \in \text{Ball} \subset \mathbb{R}^d$
2. Adversary **then** picks $x_t \in \mathcal{X} \subset \mathbb{R}^d, y_t \in [a, b]$
3. Learner suffers loss $\ell_t(\theta_t) = (\theta_t^\top x_t - y_t)^2$

Example: online linear regression

Can we perform linear regression in online fashion with non i.i.d (or even adversary) data?

Every iteration t :

1. Learner first picks $\theta_t \in \text{Ball} \subset \mathbb{R}^d$
2. Adversary **then** picks $x_t \in \mathcal{X} \subset \mathbb{R}^d, y_t \in [a, b]$
3. Learner suffers loss $\ell_t(\theta_t) = (\theta_t^\top x_t - y_t)^2$

Learner has to make decision θ_t based on history up to $t - 1$,
while adversary could pick (x_t, y_t) even after seeing θ_t

Example: online linear regression

Can we perform linear regression in online fashion with non i.i.d (or even adversary) data?

Every iteration t :

1. Learner first picks $\theta_t \in \text{Ball} \subset \mathbb{R}^d$
2. Adversary **then** picks $x_t \in \mathcal{X} \subset \mathbb{R}^d, y_t \in [a, b]$
3. Learner suffers loss $\ell_t(\theta_t) = (\theta_t^\top x_t - y_t)^2$

Learner has to make decision θ_t based on history up to $t - 1$,
while adversary could pick (x_t, y_t) even after seeing θ_t

Adversary seems too powerful...

Example: online linear regression

BUT, a very intuitive algorithm actually achieves no-regret property:

Example: online linear regression

BUT, a very intuitive algorithm actually achieves no-regret property:

Every iteration t :

1. Learner first picks θ_t that minimizes the aggregated loss

$$\theta_t = \arg \min_{\theta \in \text{Ball}} \sum_{i=0}^{t-1} (\theta^\top x_i - y_i)^2 + \lambda \|\theta\|_2^2$$

Example: online linear regression

BUT, a very intuitive algorithm actually achieves no-regret property:

Every iteration t :

1. Learner first picks θ_t that minimizes the aggregated loss

$$\theta_t = \arg \min_{\theta \in \text{Ball}} \sum_{i=0}^{t-1} (\theta^\top x_i - y_i)^2 + \lambda \|\theta\|_2^2$$

This is called Follow-the-Regularized-Leader (FTRL), and it achieves no-regret property:

Example: online linear regression

BUT, a very intuitive algorithm actually achieves no-regret property:

Every iteration t :

1. Learner first picks θ_t that minimizes the aggregated loss

$$\theta_t = \arg \min_{\theta \in \text{Ball}} \left(\sum_{i=0}^{t-1} (\theta^\top x_i - y_i)^2 + \lambda \|\theta\|_2^2 \right)$$

This is called Follow-the-Regularized-Leader (FTRL), and it achieves no-regret property:

$$\sum_{i=0}^{T-1} \ell_i(\theta_i) - \min_{\theta \in \text{Ball}} \sum_{i=0}^{T-1} \ell_i(\theta) = O(\sqrt{T})$$

Avg-Reg
 $= \frac{\sqrt{T}}{T} \rightarrow 0$

Generally, Follow-the-Regularized-Leader is no-regret

At time step t , learner has seen $\ell_0, \dots, \ell_{t-1}$, which new decision she could pick?

$$\text{FTRL: } \theta_t = \min_{\theta \in \Theta} \sum_{i=0}^{t-1} \ell_i(\theta) + \lambda R(\theta)$$

Ex $\| \theta \|_2^2$

Generally, Follow-the-Regularized-Leader is no-regret

At time step t , learner has seen $\ell_0, \dots, \ell_{t-1}$, which new decision she could pick?

$$\text{FTRL: } \theta_t = \min_{\theta \in \Theta} \sum_{i=0}^{t-1} \ell_i(\theta) + \lambda R(\theta)$$

Informal Theorem (FTRL): when things are convex, FTRL is no-regret, i.e.,

$$\frac{1}{T} \left[\sum_{t=0}^{T-1} \ell_t(\theta_t) - \min_{\theta \in \Theta} \sum_{t=0}^{T-1} \ell_t(\theta) \right] = O\left(1/\sqrt{T}\right)$$

Dagger Revisit

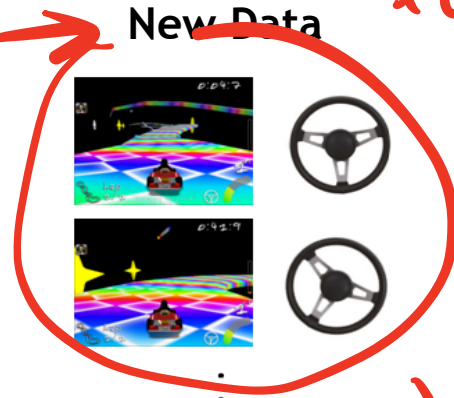
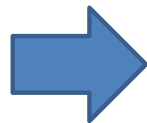
Steering
from
expert

At iteration t:

New Data

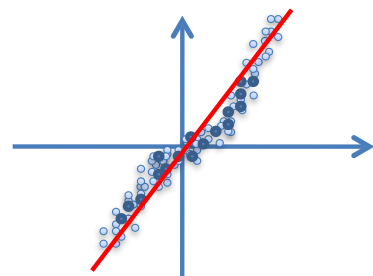
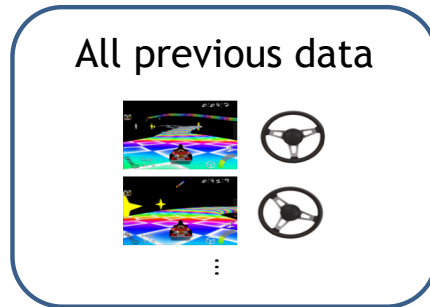
l_t

π_e



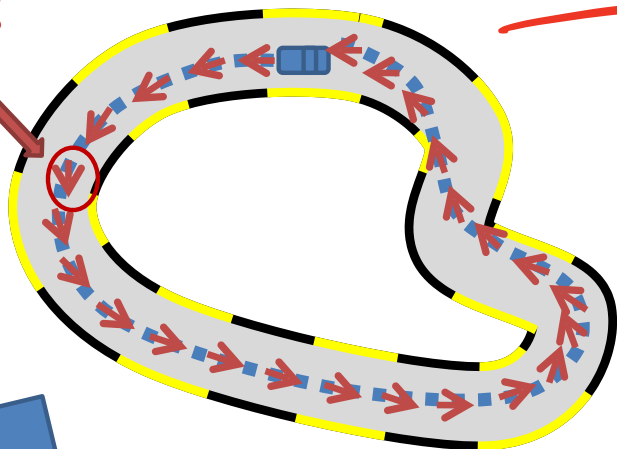
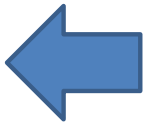
Data
Agg

Aggregate
Dataset

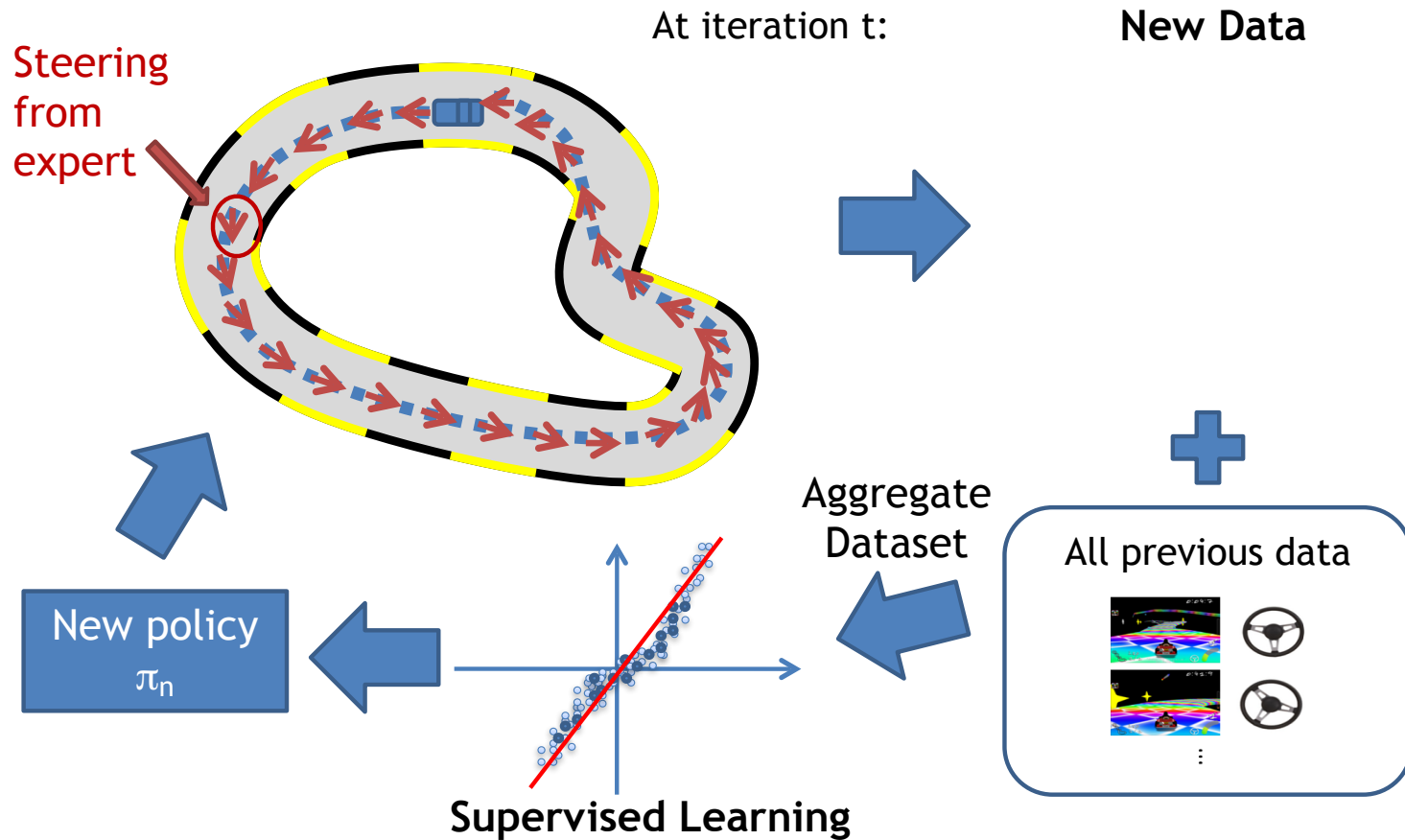


Supervised Learning

New policy
 π_n

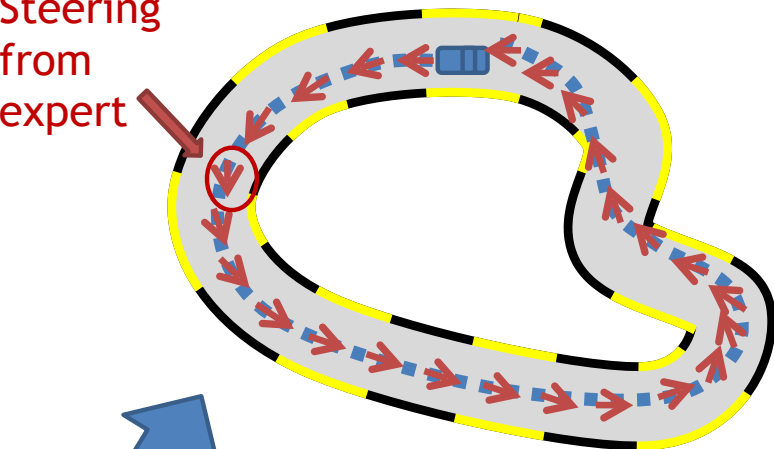


Dagger Revisit



Dagger Revisit

Steering
from
expert



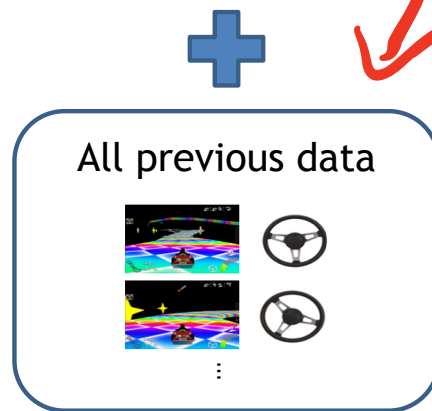
At iteration t:

New Data

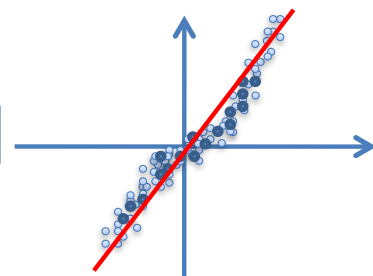
a is continuous

$$\ell_t(\pi) = \sum_{i=1}^m \|\pi(s^i) - \pi^*(s^i)\|_2^2$$

Aggregate
Dataset

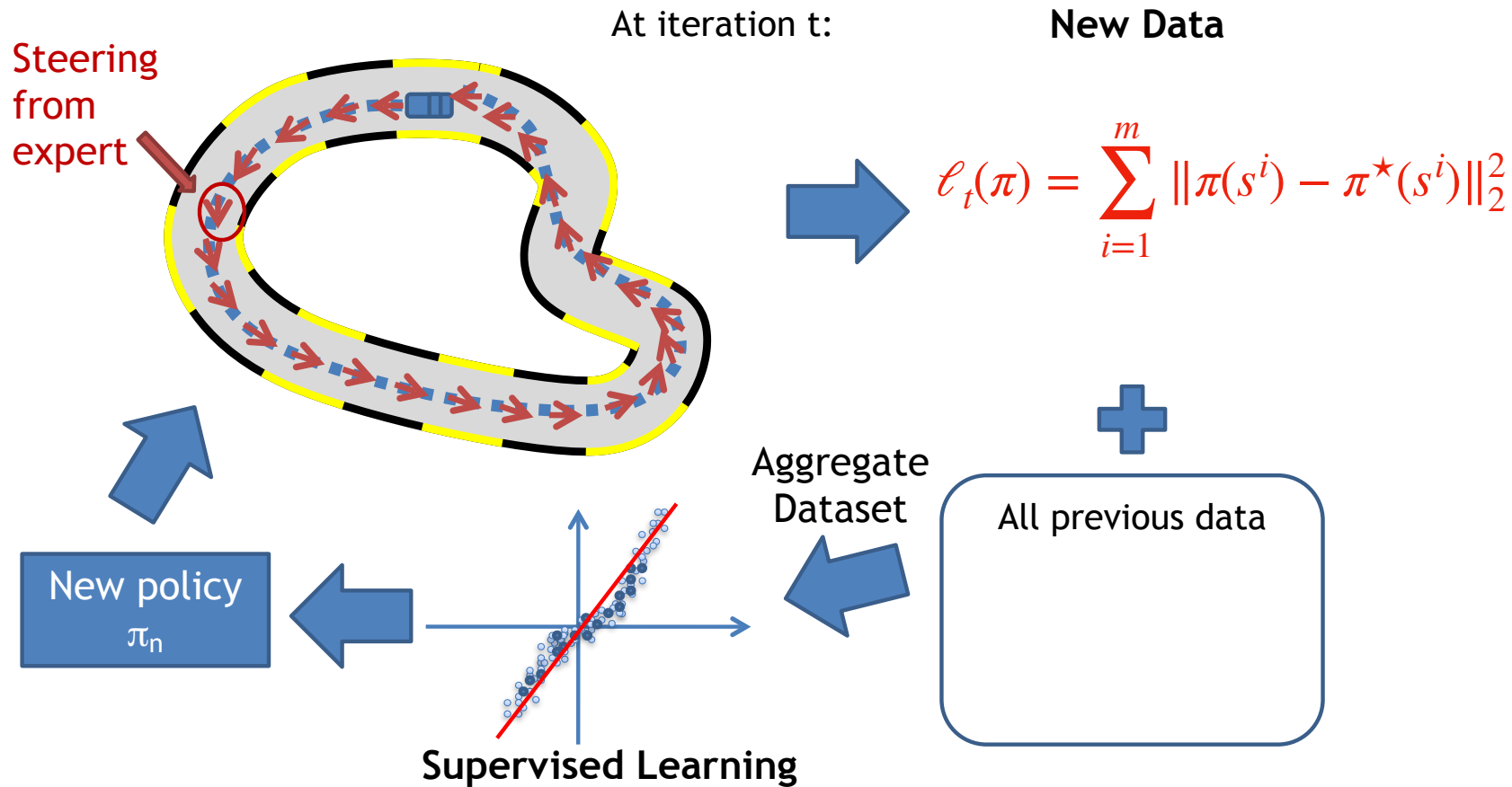


New policy
 π_n

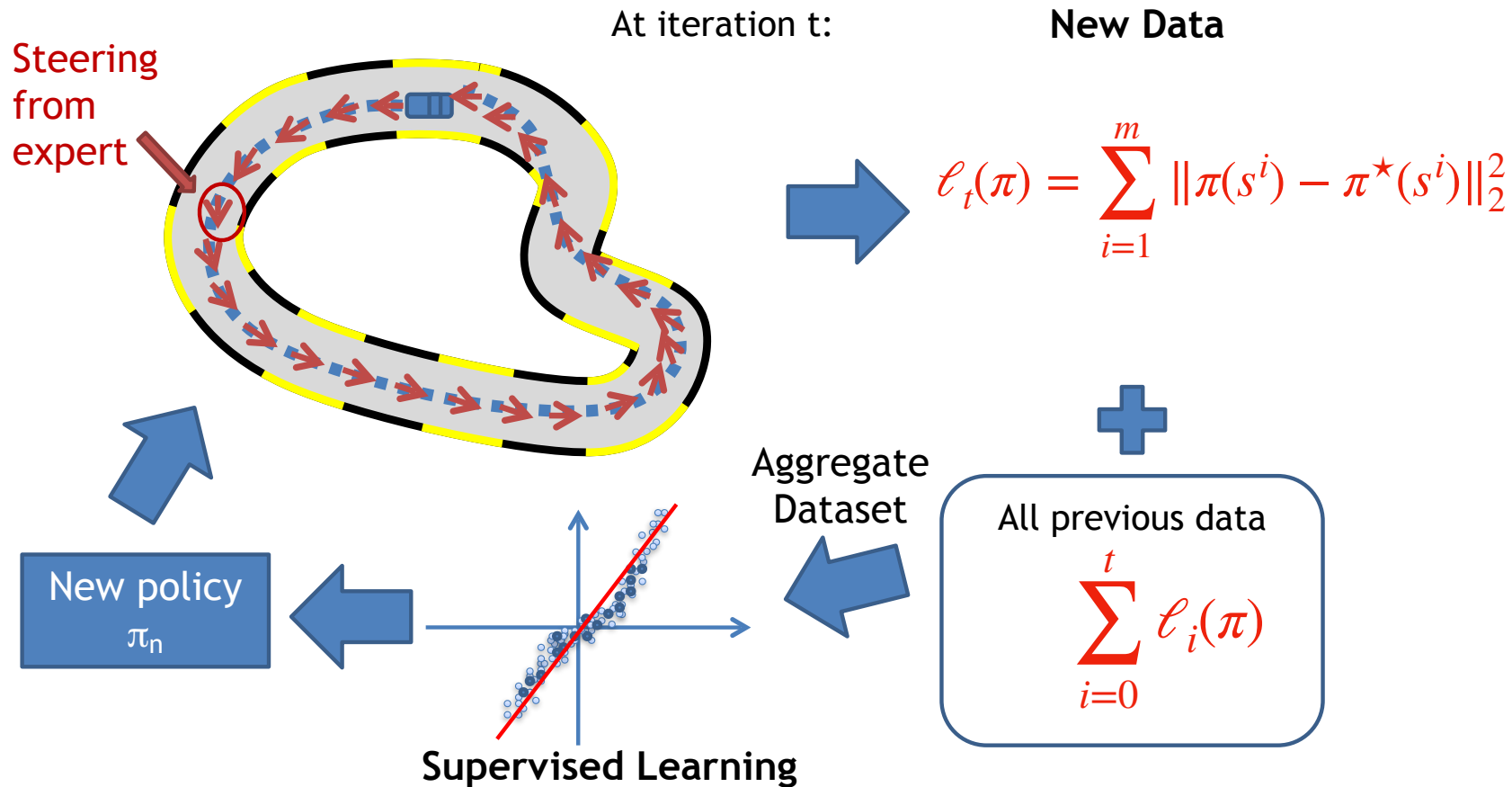


Supervised Learning

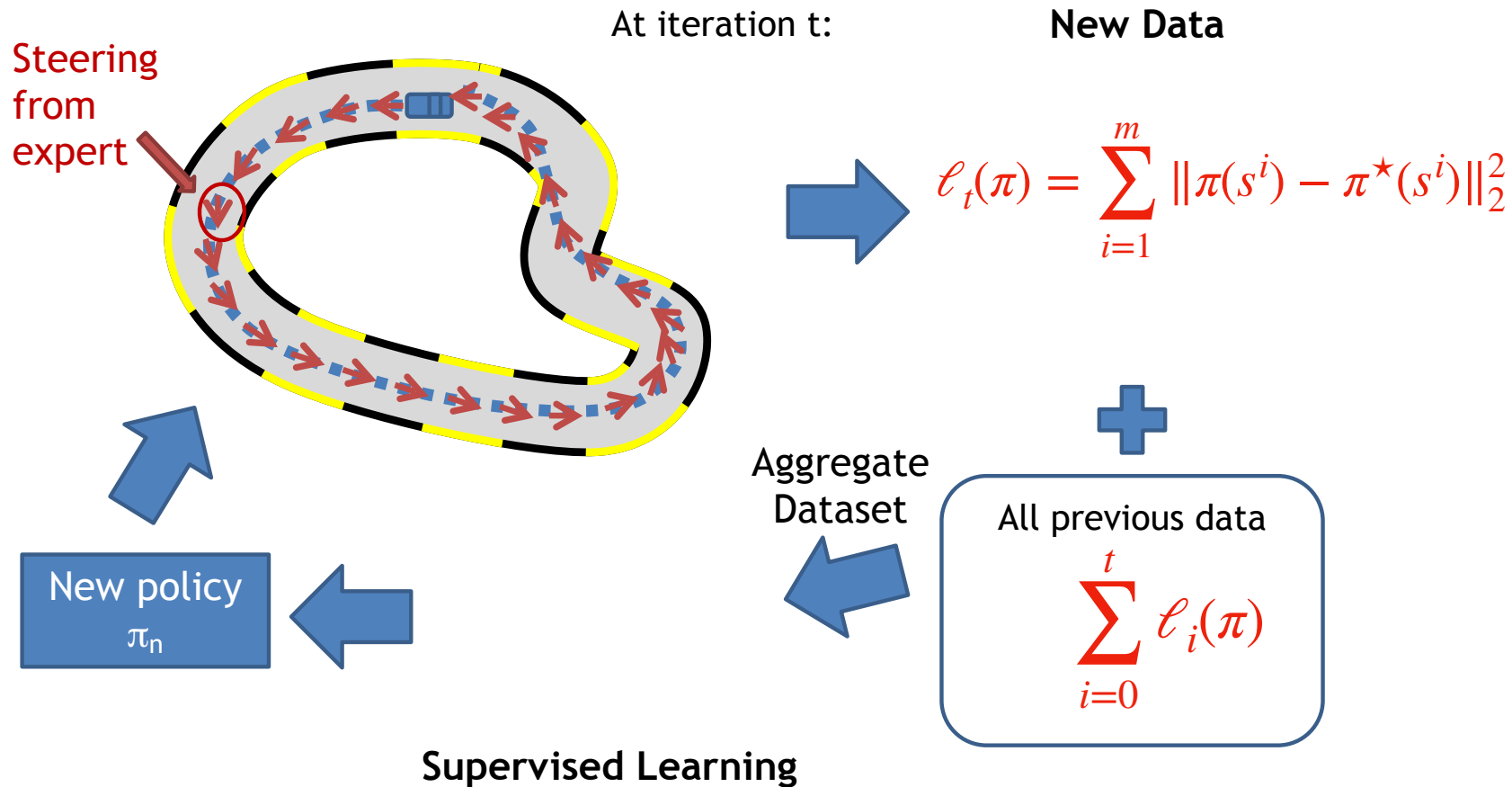
Dagger Revisit



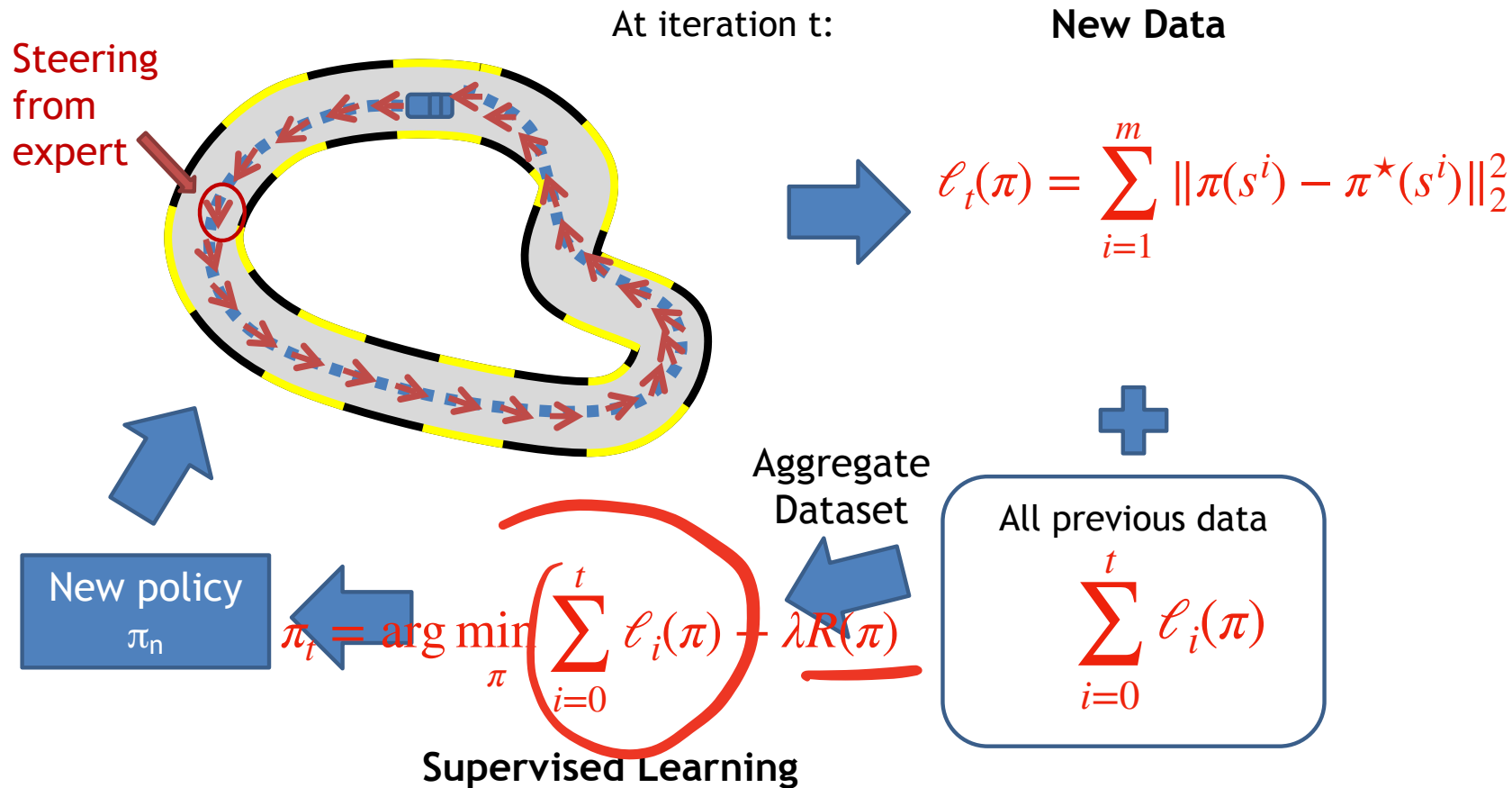
Dagger Revisit



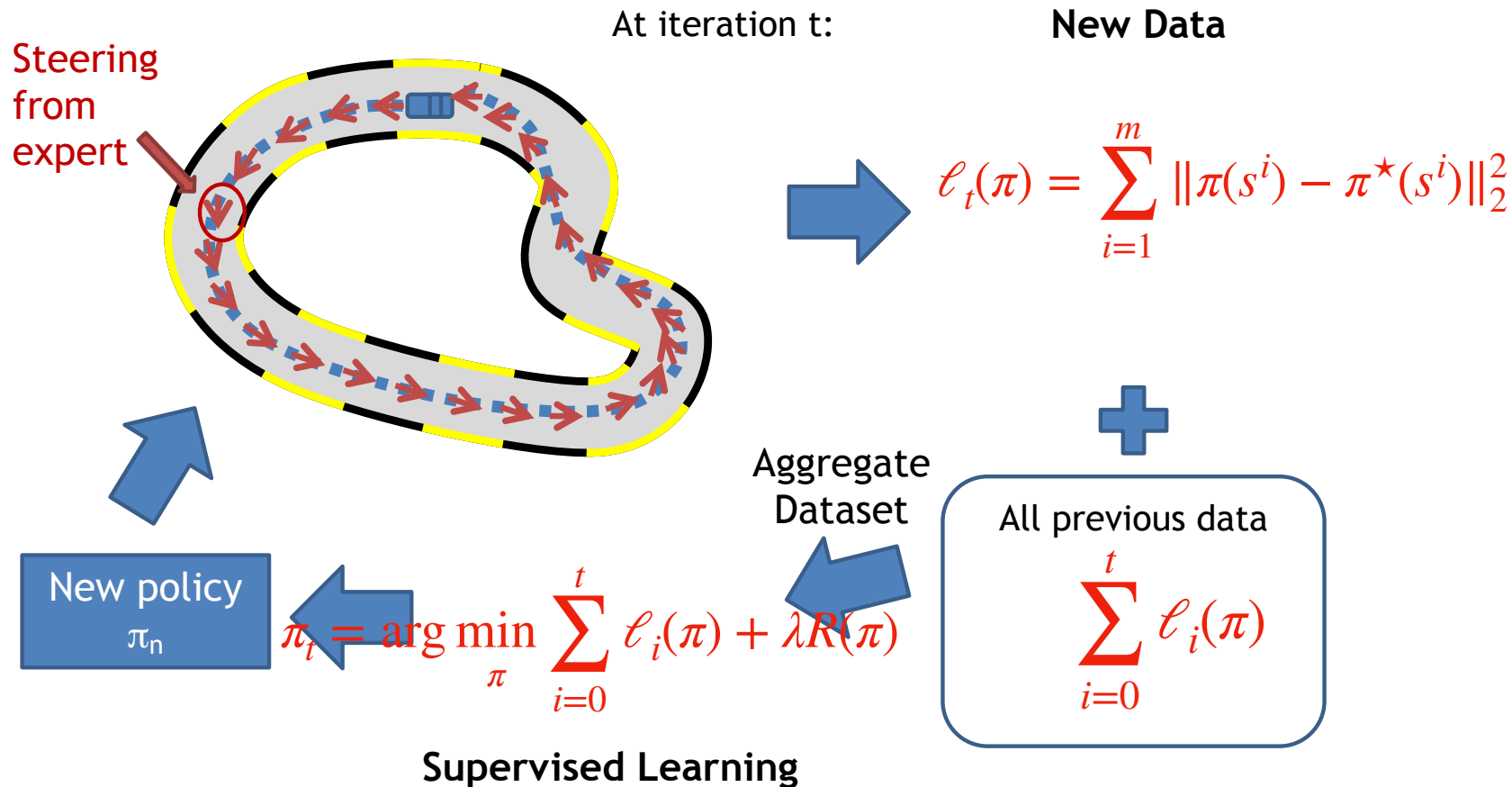
Dagger Revisit



Dagger Revisit



Dagger Revisit



Summary for Today

1. The DAgger algorithm

Initialize π^0 , and dataset $\mathcal{D} = \emptyset$

For $t = 0 \rightarrow T - 1$:

1. W/ π^t , generate dataset $\mathcal{D}^t = \{s_i, a_i^\star\}, s_i \sim d_\mu^{\pi^t}, a_i^\star = \pi^\star(s_i)$
2. **Data aggregation:** $\mathcal{D} = \mathcal{D} + \mathcal{D}^t$
3. **Update policy via Supervised-Learning:** $\pi^{t+1} = \text{SL}(\mathcal{D})$

Summary for Today

1. The DAgger algorithm

Initialize π^0 , and dataset $\mathcal{D} = \emptyset$

For $t = 0 \rightarrow T - 1$:

1. W/ π^t , generate dataset $\mathcal{D}^t = \{s_i, a_i^\star\}, s_i \sim d_\mu^{\pi^t}, a_i^\star = \pi^\star(s_i)$
2. **Data aggregation:** $\mathcal{D} = \mathcal{D} + \mathcal{D}^t$
3. **Update policy via Supervised-Learning:** $\pi^{t+1} = \text{SL}(\mathcal{D})$

2. We can see that DAgger is essentially an online-learning algorithm (FTRL)