

# META REINFORCEMENT LEARNING

- Through previous experience (motivation of Meta RL)
- Data efficiency (with a few examples only)

• How might you get a machine to accomplish this task?

- Provide model of image formation, etc.
- Fewer human priors, more data-driven priors: Greater success.

## Meta-RL Problem Formulation and Examples

META-LEARNING PROBLEM:

• In Supervised Learning:

Inputs:  $x$       Outputs:  $y$   
 $\hookrightarrow y = f(x; \theta)$

DATA:  $\{(x, y)_i\}$

• In Meta Supervised learning:

$D_{train}: \{(x, y)_{i:k}\}$

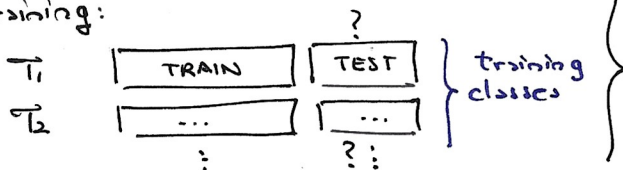
Inputs:  $D_{train}$  and  $x_{test}$       Outputs:  $y_{test}$   
 $\hookrightarrow y_{test} = f(D_{train}, x_{test}; \theta)$

DATA:  $\{D_i\}$   
 $\hookrightarrow$  diff. datasets (tasks)  
 $D_i: \{(x, y)_j\}$

### EXAMPLE: FEW-SHOT CLASSIFICATION

• Given 1 example of 5 classes: we want to classify new examples. ( $x_{test}$ )

• Meta-Training:

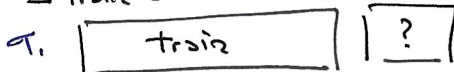


$\Uparrow$  META-TESTING  
 $\rightarrow$  META-TRAINING.

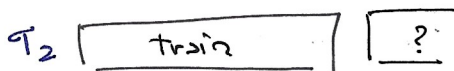
$\hookrightarrow$  which updates  $f$  of function (the parameters)

#### META-TRAINING

- Train dataset #1: Cat-Bird -

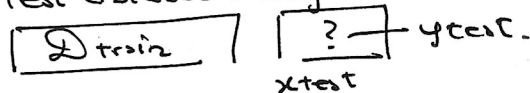


- Train dataset #2: flower-bike -



#### META-TESTING

- Test dataset: dog-butter -



In one gradient step

• In Reinforcement Learning:

Inputs:  $s_t$       Outputs:  $a_t$   
 $\hookrightarrow a_t = \pi(s_t; \theta)$

DATA:  $\{(s_t, a_t, r_t, s_{t+1})\}$

• In Meta RL:

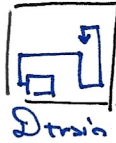
Inputs:  $D_{train} \oplus s_t$       Outputs:  $a_t$   
 $\hookrightarrow a_t = f(D_{train}, s_t; \theta)$   
 $\leftarrow$   $k$ -rollouts from  $\pi$

DATA:  $\{D_i\}$   
 dataset of datasets collected for each task.

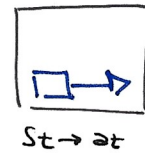
- ⇒ Design and Optimization of the  $f$  function
- ⇒ Collecting Appropriate Data (Learning to Explore - collect good  $\mathcal{D}_{train}$ )

### EXAMPLE: MAZE NAVIGATION

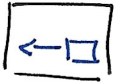
- Given a small amount of experience



⇒ learn to solve the task



By learning how to solve many other tasks... (other diff. mazes)



- For  $\mathcal{D}_{train}$ : (episodes)
  - ↳  $k$  rollouts from  $\pi$  ⇒ EPISODIC VARIANT
  - ↳  $1, \dots, k$  timesteps from  $\pi$  ⇒ ONLINE VARIANT.

### QUESTIONS

#### 1. Meta-Learning vs Transfer Learning

- Optimize for good transfer performance (new task learnt differently)
  - ↳ Meta-Learning.
- Similar Motivations: both want to understand how to learn and transfer a representation that will help with learning new skills.
  - ↳ less formal for generalization process.
- TL: Accomplish this goal without any additional training.
- ML: Explicitly says that you can gather little more data on the new task to update the model (quickly)

### Method Classes

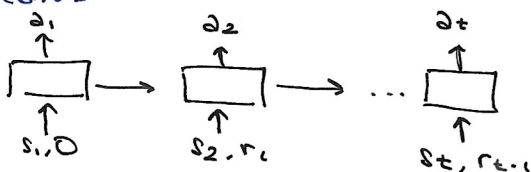
#### BLACK-BOX METHOD

Design of the  $f$  function:

$\mathcal{D}_{train}, s_t \rightarrow a_t$

#### Recurrent Network (LSTM, NTN, Conv)

$$a_t = f(\mathcal{D}_{train}, s_t; \theta)$$



⇒ diff compared to recurrent policy?

⇒ Hidden state maintained across episodes within a task.

- + : general + expressive
- + : variety of design choices in Architecture
- : complex model for complex task of learning
- : Impractical data requirements.

## PAPER: Simple Neural Attentive Meta-Learner

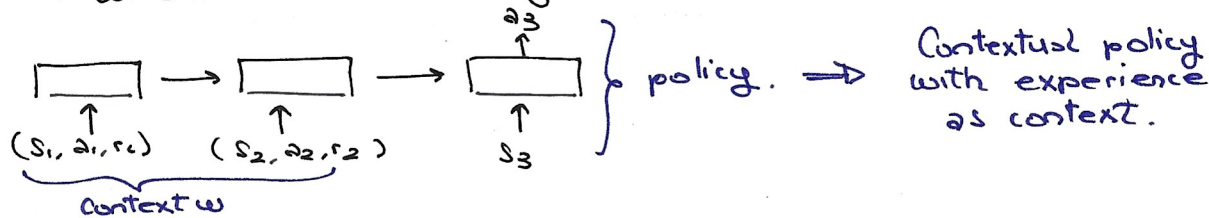
• Experiment: learning to visually navigate a maze

- train on 1000 small mazes
- test on held-out small mazes and large mazes.

• Digression: Connection to Contextual Policies.

• Contextual policy:  $\pi_{\theta}(a|s, w)$

↳  $w$ : stack location, walking direction, etc. (context)



What about goal-conditioned policies / value functions?

- Rewards: strict generalization of goals.

- Meta-RL: Adapt new tasks vs. generalize new tasks  
META-LEARNING (k-shot) CLASSIC LEARNING (0-shot)

## OPTIMIZATION-BASED APPROACHES

Line-Tuning  
[test-time]

$$\Phi_j \leftarrow \underset{\substack{\uparrow \\ \text{pretrained} \\ \text{parameters}}}{\Theta} - \alpha \nabla_{\Theta} \underset{\substack{\uparrow \\ \text{training data} \\ \text{for new task.}}}{\mathcal{L}_{\text{train}}^j(\Theta)}$$

OUR METHOD

$$\min_{\Theta} \sum_{\text{task } i} \mathcal{L}_{\text{test}}^i(\Theta - \alpha \nabla_{\Theta} \mathcal{L}_{\text{train}}^i(\Theta))$$

↳ Embedded gradient descent in Meta-Learning

↳ ∴ Train over many tasks, to learn parameter vector  $\Theta$  that transfers.

Suppose that:

- $\Theta$ : parameter vector being meta-learned
- $\Phi_i^*$ : optimal parameter vector for task  $i$

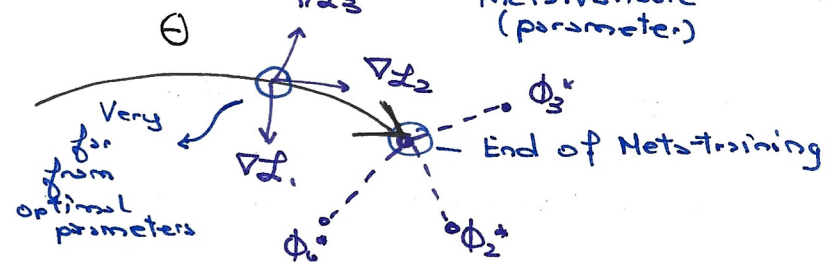
$\square$ : Meta-Learning  
 $--$ : Learning/Adaptation

Meta-Learning

$$\min_{\Theta} \sum_{\text{task } i} \mathcal{L}_{\text{test}}^i \quad \left( \square \right) \quad \left( \nabla_{\Theta} \mathcal{L}_{\text{train}}^i(\Theta) \right)$$

↑ Training/Adaptation

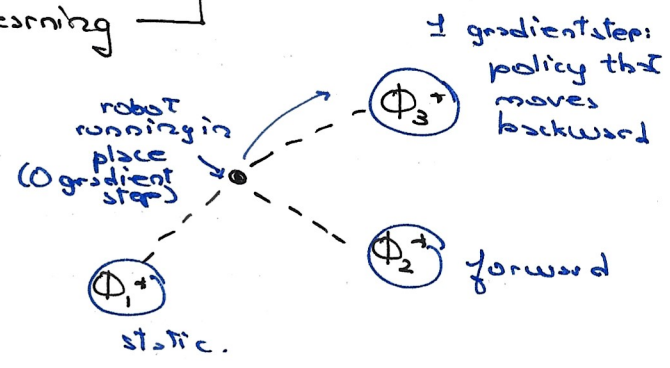
Optimize this metavariable (parameter)



Model-Agnostic Meta-Learning

MAML

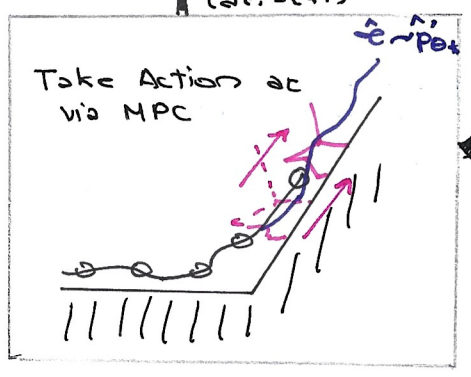
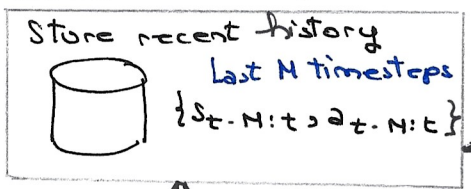
Example: At the end of meta-training



MAML and Model-Based RL

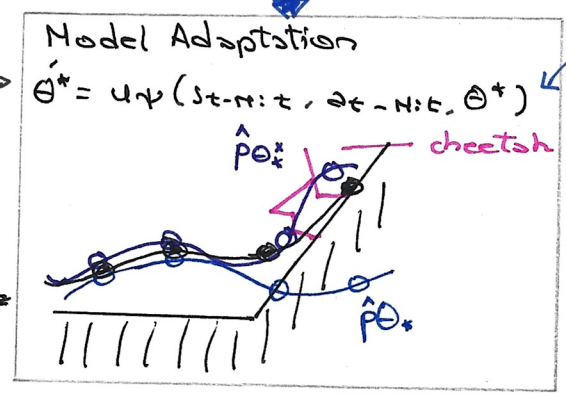
- Adapt to diff. dynamics in environments.

ONLINE ADAPTATION



Ex: On variable terrains

META-TRAIN A PRIOR  $\Theta^*$



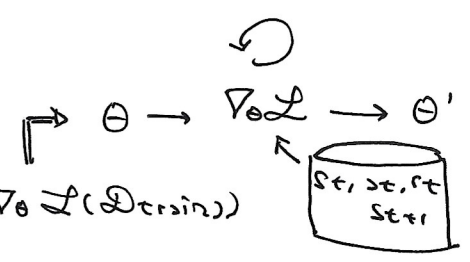
One step of gradient descent on learned dynamics model

Ex: META-TRAIN on variable terrains

META-TEST with slope, missing leg, payload, calibration errors.

To summarize...

- Black-Box Method:  $a_t = f(\mathcal{D}_{\text{train}}, s_t; \Theta)$
- Optimization-Based Method:  $a_t = f(s_t; \Theta - \alpha \nabla_{\Theta} \mathcal{L}(\mathcal{D}_{\text{train}}))$ 
  - + inductive bias of SGD built in
  - + model-agnostic (plug and play with existing archi)
  - RL gradients & informative (policy gradients, Bellman errors)



## Learning to Explore

How should we collect  $\mathcal{D}_{\text{train}}$ ?

Coupling problem  $\rightarrow$  Execution  $\rightarrow$  Exploration

#0 Optimize for Exploration + Execution End-to-End w.r.t Reward

⊕ simple  
leads to optimal strategy in principle

⊖ Challenging optimization when exploration is hard

Example of hard Meta-RL Problem (Exploration):

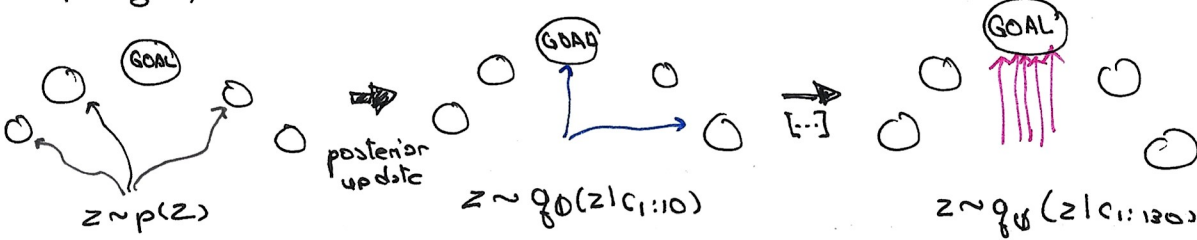
- Learned cooking tasks in previous kitchens (META-TRAINING)
- Goal: quickly learn tasks in a new kitchen. (META-TESTING)
  - 1 exploration episode
  - 1 execution episode

#1 Leverage Alternative Exploration Strategies.

a) Posterior sampling (Thompson sampling)

dist. you've collected so far.

- Learn dist. over latent variable  $p(z)$ ,  $q(z|\mathcal{D}_{\text{train}})$  and corresponding task policies  $\pi(a|s, z)$ .
- Sample  $z$  from current «posterior» and sample from policy  $\pi(a|s, z)$



$\Rightarrow$  Posterior sampling is bad when...

⊖ Goals far away and signs on wall that tells you the correct goal.

b) Use intrinsic rewards

c) Tasks dynamism and reward prediction

- Train model  $f(s', r | s, a, \mathcal{D}_{\text{train}})$
- Collect  $\mathcal{D}_{\text{train}}$  so that model is accurate.

Bad when...

⊖ Lots of distractors, or complex, high-dim state dynamics.

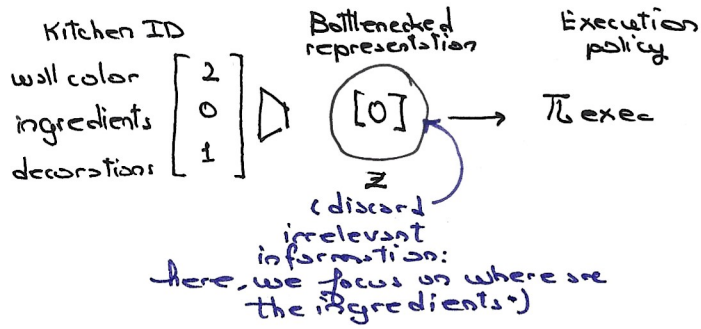
⊕ Easy to optimize  
Many based on principled strategies

⊖ Suboptimal by arbitrarily large amount in some environments.

#2: Decouple by acquiring representation of task relevant information

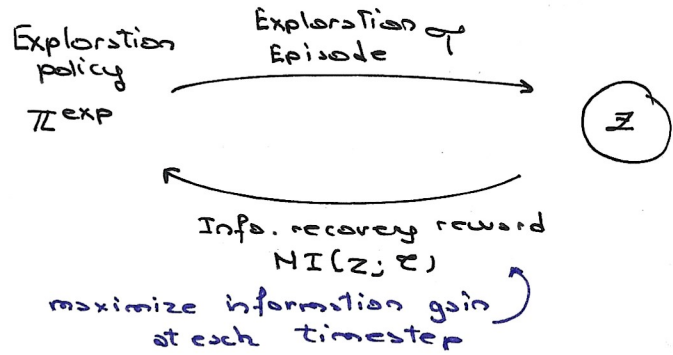
### META-TRAINING:

1) Learn execution  $\oplus$  identify key information.

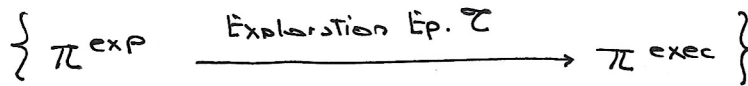


- $\oplus$ : Leads to optimal strategy in principle.  
• Easy to optimize in practice
- $\ominus$ : Requires task identifier (sign)

2) Learn to explore by recovering that information.



### META-TESTING:



Decoupled Reward-free Exploration and Execution in Meta-RL (DREAM)

$\Rightarrow$  Consistent with End-to-End Approach.

Example: Sparse Reward 3D Visual Navigation Problem

- $\rightarrow$  Task: go to the (key/block/ball), color specified by the sign
- $\rightarrow$  Agent starts on other side of barrier, must walk around to read the sign
- $\rightarrow$  Pixel observations (80x60 RGB)
- $\rightarrow$  Sparse Binary Reward

To summarize...

- $\rightarrow$  End-to-End
- $\rightarrow$  Alternative Strategies
- $\rightarrow$  Decoupled Exp. and Exec (DREAM)

## Challenges and Latest Developments.

1. Adopting to entirely new tasks. (meta-world benchmark?)
2. Robustness to out of distribution tasks.
3. Where do the tasks even come from?
4. Better RL Algorithms.