

Subspace Learning for Effective Meta-Learning

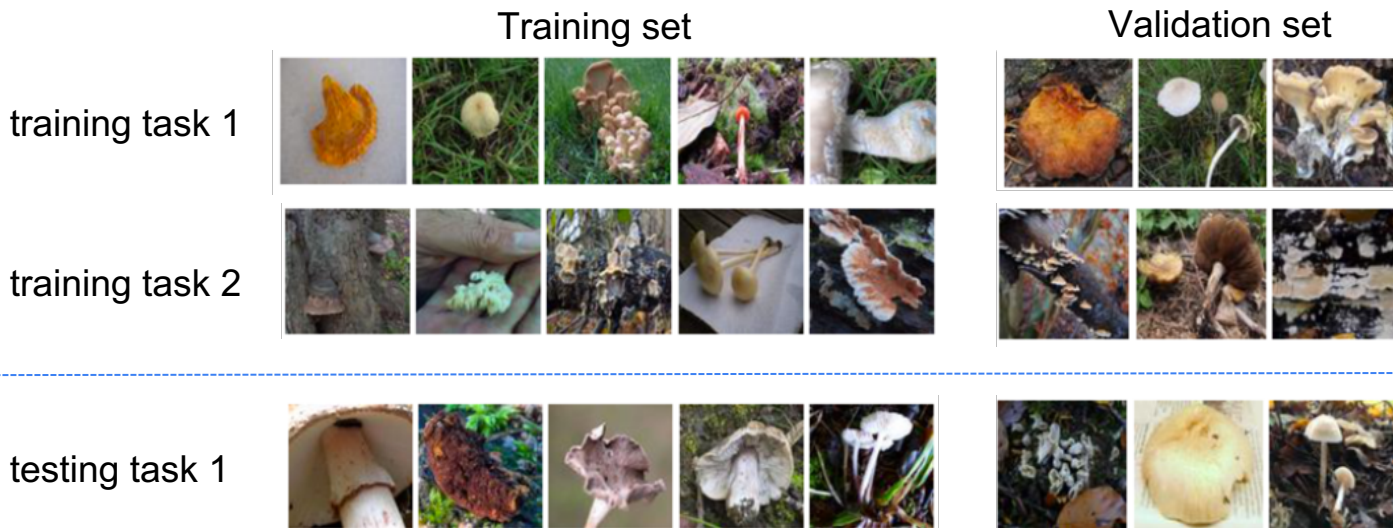
Weisen Jiang^{1 2} James T. Kwok² Yu Zhang^{1 3}

¹Southern University of Science and Technology

²Hong Kong University of Science and Technology

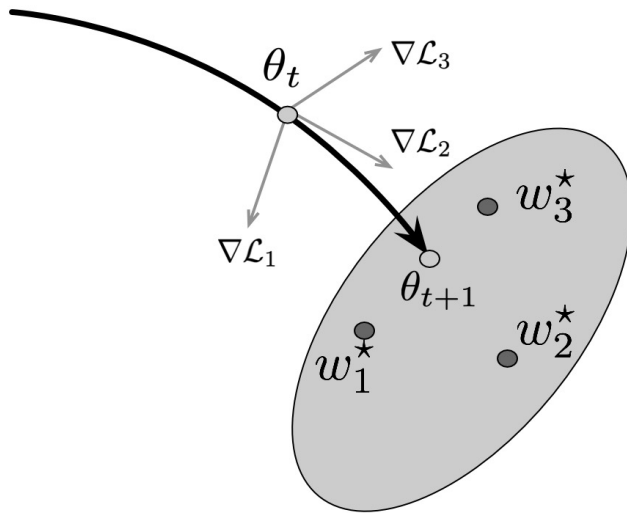
³Peng Cheng Laboratory

Background: Meta-Learning



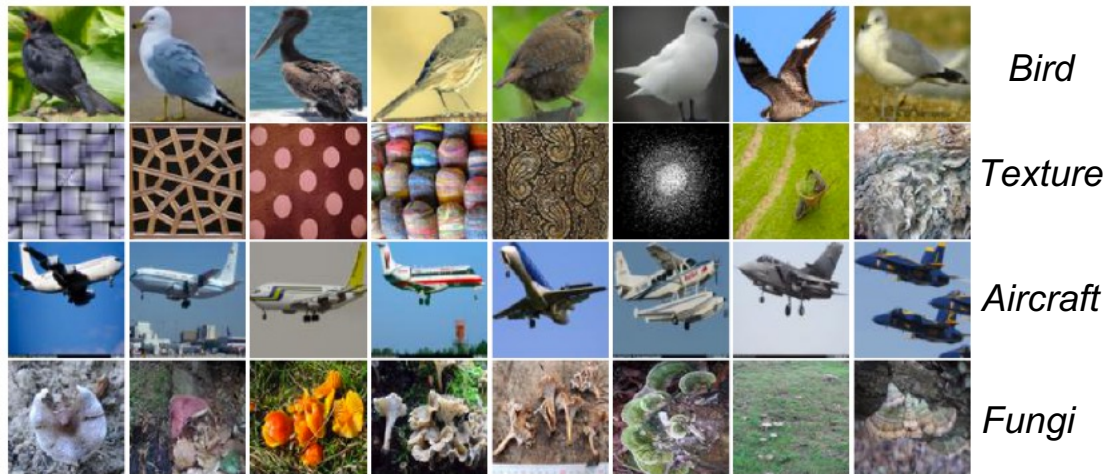
- Humans can learn new concepts with limited examples.
- However, deep networks are **data-hungry**.
- **Meta-learning** aims to extract meta-knowledge from seen tasks to accelerate learning unseen tasks.

Background: MAML



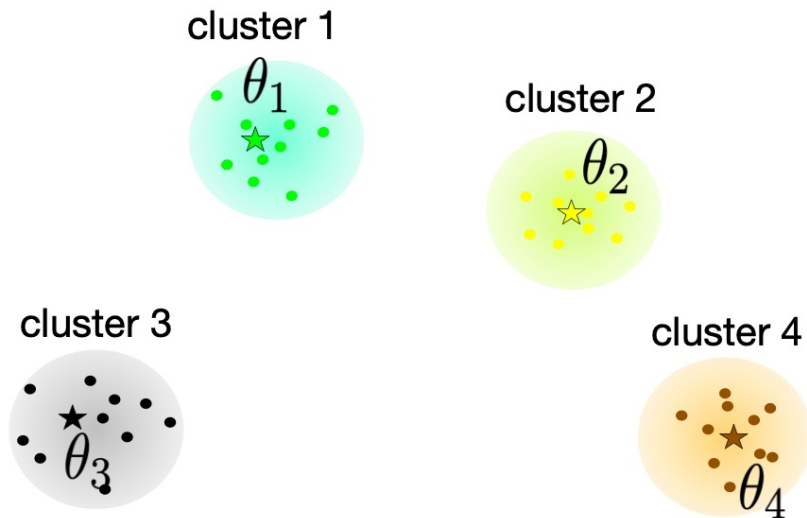
MAML learns a globally-shared initialization for all tasks.

Background: Complex Environment



A single meta-model may not be sufficient to capture all the meta-knowledge

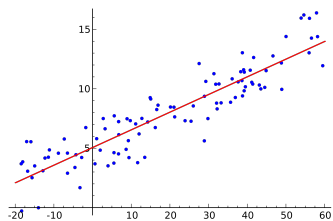
Background: Structured Meta-Learning



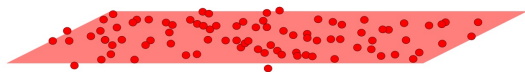
- Recent structured meta-learning methods (DPMM and TSA-MAML) cluster tasks into multiple groups .
- **Cluster centroids** form group-specific initializations.
- However, task model parameters may lie in **a subspace mixture**.

Background: Research Gap

Prior work

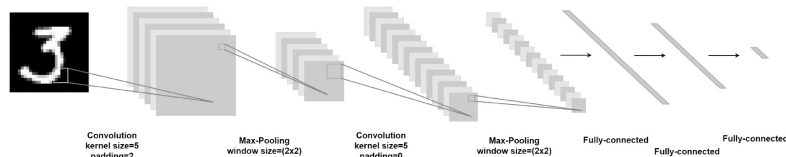


linear regression

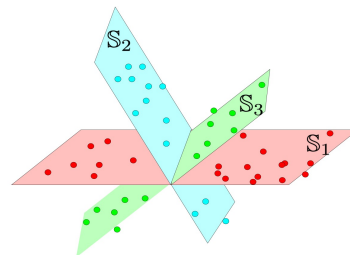


single subspace

This work



nonlinear models



a subspace mixture

MUSML: A Subspace Mixture

- task parameter w_τ 's are assumed to lie in K subspaces $\{\mathbb{S}_1, \dots, \mathbb{S}_K\}$
- $S_k \in \mathbb{R}^{d \times m}$ is a basis of $\mathbb{S}_k$
- meta-parameters: $\{S_1, \dots, S_K\}$

MUSML: Base Learner

- in each S_k , we search for a linear combination to form task model $w_\tau = S_k v_{\tau,k}^*$:

$$v_{\tau,k}^* = \arg \min_{v_\tau \in \mathbb{R}^m} \mathcal{L}(\mathcal{D}_\tau^{tr}; S_k v_\tau)$$
- use convex program when $\mathcal{L}(\mathcal{D}; w)$ is convex
- in nonconvex case, we seek an **approximate minimizer**:
for $t' = 1, \dots, T_{in}$,

$$v_{\tau,k}^{(t')} = v_{\tau,k}^{(t'-1)} - \alpha \nabla_{v_{\tau,k}^{(t'-1)}} \mathcal{L}(\mathcal{D}_\tau^{tr}; S_k v_{\tau,k}^{(t'-1)})$$

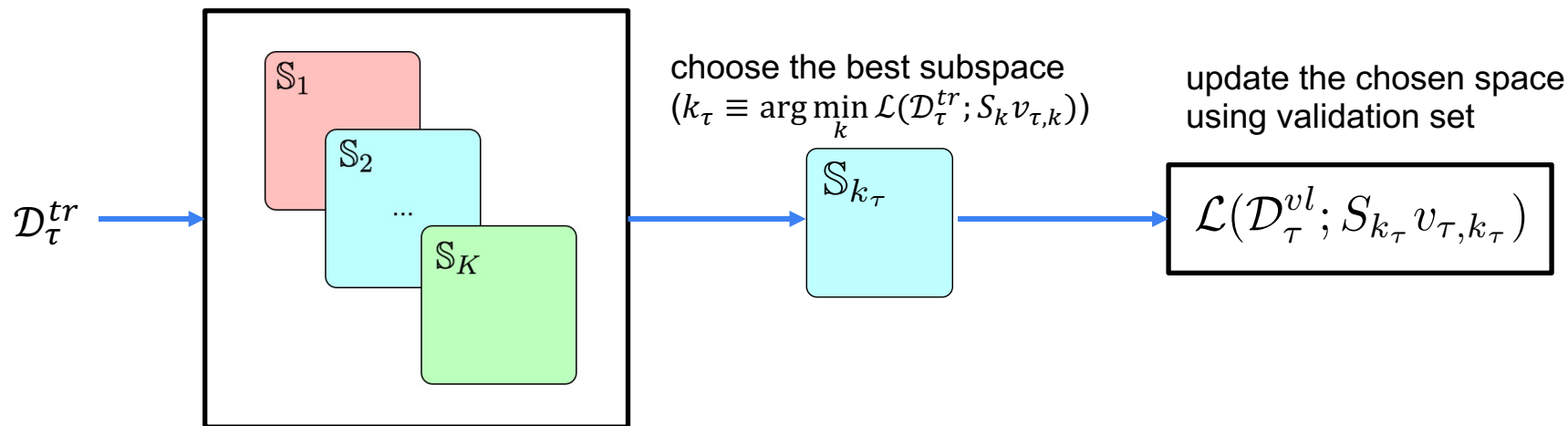
and then $v_{\tau,k} = v_{\tau,k}^{(T_{in})}$

Algorithm 1 Multiple Subspaces for Meta-Learning (MUSML).

Require: stepsize α , $\{\eta_t\}$; number of inner gradient steps T_{in} , number of subspaces K , subspace dimension m , temperature $\{\gamma_t\}$; initialization $\mathbf{v}^{(0)}$;

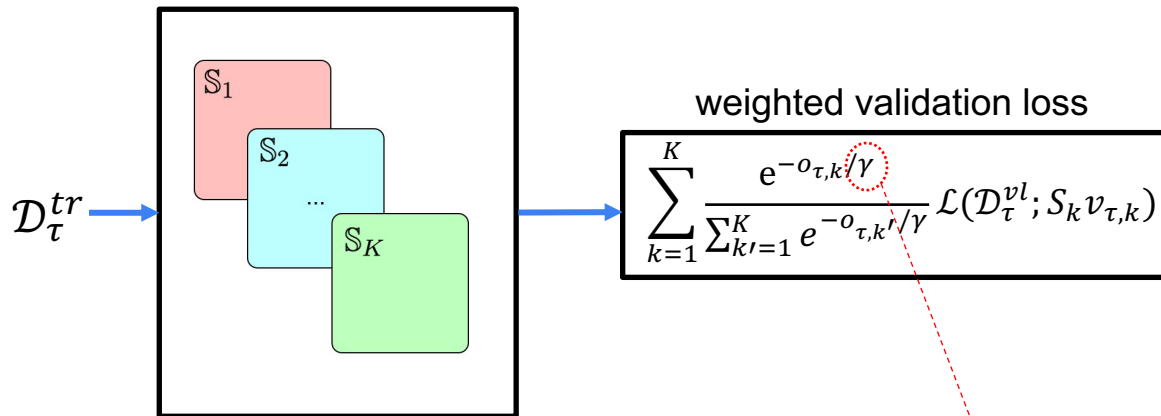
- 1: **for** $t = 0, 1, \dots, T - 1$ **do**
- 2: sample a task τ with $\mathcal{D}_\tau^{tr}$ and $\mathcal{D}_\tau^{vl}$;
- 3: *base learner*:
- 4: **for** $k = 1, \dots, K$ **do**
- 5: initialize $\mathbf{v}_{\tau,k}^{(0)} = \mathbf{v}^{(0)}$;
- 6: **for** $t' = 0, 1, \dots, T_{in} - 1$ **do**
- 7: $\mathbf{v}_{\tau,k}^{(t'+1)} = \mathbf{v}_{\tau,k}^{(t')} - \alpha \nabla_{\mathbf{v}_{\tau,k}^{(t')}} \mathcal{L}(\mathcal{D}_\tau^{tr}; S_{k,t} \mathbf{v}_{\tau,k}^{(t')})$;
- 8: **end for**
- 9: $\mathbf{v}_{\tau,k} \equiv \mathbf{v}_{\tau,k}^{(T_{in})}$;
- 10: $o_{\tau,k} = \mathcal{L}(\mathcal{D}_\tau^{tr}; S_{k,t} \mathbf{v}_{\tau,k})$;
- 11: **end for**
- 12: *meta-learner*:
- 13: $\mathcal{L}_{vl} = \sum_{k=1}^K \frac{\exp(-o_{\tau,k}/\gamma_t)}{\sum_{k'=1}^K \exp(-o_{\tau,k'}/\gamma_t)} \mathcal{L}(\mathcal{D}_\tau^{vl}; S_{k,t} \mathbf{v}_{\tau,k})$;
- 14: $\{\mathbf{S}_{1,t+1}, \dots, \mathbf{S}_{K,t+1}\} = \{\mathbf{S}_{1,t}, \dots, \mathbf{S}_{K,t}\} - \eta_t \nabla_{\{\mathbf{S}_{1,t}, \dots, \mathbf{S}_{K,t}\}} \mathcal{L}_{vl}$;
- 15: **end for**
- 16: **Return** $\mathbf{S}_{1,T}, \dots, \mathbf{S}_{K,T}$.

MUSML: Meta-Learner (one-hot selection)



inefficient: only one subspace is updated at each step.

MUSML: Meta-Learner (soft selection)



All spaces are updated simultaneously.

Algorithm 1 Multiple Subspaces for Meta-Learning (MUSML).

Require: stepsize α , $\{\eta_t\}$; number of inner gradient steps T_{in} , number of subspaces K , subspace dimension m , temperature $\{\gamma_t\}$; initialization $\mathbf{v}^{(0)}$;

```

1: for  $t = 0, 1, \dots, T - 1$  do
2:   sample a task  $\tau$  with  $\mathcal{D}_{\tau}^{tr}$  and  $\mathcal{D}_{\tau}^{vl}$ ;
3:   base learner:
4:   for  $k = 1, \dots, K$  do
5:     initialize  $\mathbf{v}_{\tau,k}^{(0)} = \mathbf{v}^{(0)}$ ;
6:     for  $t' = 0, 1, \dots, T_{in} - 1$  do
7:        $\mathbf{v}_{\tau,k}^{(t'+1)} = \mathbf{v}_{\tau,k}^{(t')}$   $- \alpha \nabla_{\mathbf{v}_{\tau,k}^{(t')}} \mathcal{L}(\mathcal{D}_{\tau}^{tr}; \mathcal{S}_{k,t} \mathbf{v}_{\tau,k}^{(t')})$ ;
8:     end for
9:      $\mathbf{v}_{\tau,k} \equiv \mathbf{v}_{\tau,k}^{(T_{in})}$ ;
10:     $o_{\tau,k} = \mathcal{L}(\mathcal{D}_{\tau}^{tr}; \mathcal{S}_{k,t} \mathbf{v}_{\tau,k})$ ;
11:  end for
12:  meta-learner:
13:   $\mathcal{L}_{vl} = \sum_{k=1}^K \frac{\exp(-o_{\tau,k}/\gamma_t)}{\sum_{k'=1}^K \exp(-o_{\tau,k'}/\gamma_t)} \mathcal{L}(\mathcal{D}_{\tau}^{vl}; \mathcal{S}_{k,t} \mathbf{v}_{\tau,k})$ ;
14:   $\{\mathcal{S}_{1,t+1}, \dots, \mathcal{S}_{K,t+1}\} = \{\mathcal{S}_{1,t}, \dots, \mathcal{S}_{K,t}\} - \eta_t \nabla_{\{\mathcal{S}_{1,t}, \dots, \mathcal{S}_{K,t}\}} \mathcal{L}_{vl}$ ;
15: end for
16: Return  $\mathcal{S}_{1,T}, \dots, \mathcal{S}_{K,T}$ .
  
```

- $\gamma \rightarrow 0$, soft selection becomes one-hot
- $\gamma \rightarrow \infty$, soft selection becomes uniform
- In practice, a linear annealing schedule

Analysis

➤ Under some conditions, we obtain an upper bound for the population risk

$$\mathcal{R}(\mathcal{S}) \leq \mathcal{R}^* + \rho \sqrt{m} \mathbb{E}_{\tau'} \mathbb{E}_{\mathcal{D}_{\tau'}^{tr}} \left\| v_{\tau', k_{\tau'}} - v_{\tau', k_{\tau'}}^* \right\|^2 + \varrho \mathbb{E}_{\tau'} \mathbb{E}_{\mathcal{D}_{\tau'}^{tr}} \text{dist} \left(w_{\tau'}^*, \mathbb{S}_{k_{\tau'}} \right) + K \sqrt{\frac{v^2 + 12 \varrho v (1 + m \alpha \delta)^{T_{in}}}{2 N_{tr}}}$$

(1) minimum risk

(2) distance between approximate minimizer and exact minimizer

(3) approximation error using learned subspaces

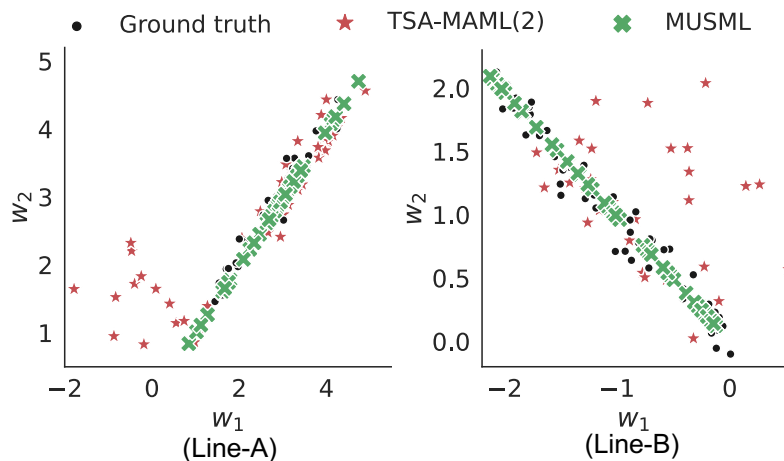
(4) complexity of subspaces (K and m)

Experiments: Few-Shot Regression (Synthetic Data : Setup)

- a **nonlinear** model: $f(x; w_\tau) = e^{0.1w_{\tau,1}x} + w_{\tau,2} |\sin(x)|$;
- $w_\tau = [w_{\tau,1}; w_{\tau,2}]$ is sampled from one of the **two** subspaces:
 - *Line-A*: $w_\tau = S_1 a_\tau + 0.1 \xi_\tau$, where $S_1 = [1; 1]$, $a_\tau \sim \mathcal{U}(1,5)$, $\xi_\tau \sim \mathcal{N}(0, I)$;
 - *Line-B*: $w_\tau = S_2 a_\tau + 0.1 \xi_\tau$, where $S_2 = [-1; 1]$, $a_\tau \sim \mathcal{U}(0,2)$, $\xi_\tau \sim \mathcal{N}(0, I)$;
- samples: $y = f(x; w_\tau) + 0.05 \xi$, where $x \sim \mathcal{U}(-5,5)$ and $\xi \sim \mathcal{N}(0,1)$;
- $K = 2, m = 1$.
- γ_t : a linear annealing schedule

Experiments: Few-Shot Regression (Synthetic Data : Results)

visualization of task model parameters



➤ MUSML discovers underlying subspaces

meta-testing MSE

MAML	0.74
BMG	0.67
DPMM	0.56
HSML	0.49
ARML	0.60
TSA-MAML	0.58
MUSML	0.07

➤ MUSML performs the best.

Experiments: Few-shot Classification (Setup)

- meta-datasets:
 1. *Meta-Dataset-BTAF*
 2. *Meta-Dataset-ABF*
 3. *Meta-Dataset-CIO*
- $f(\cdot; w)$: Conv4 backbone + Prototype classifier
- K and m are chosen from 1 to 5

	#classes (meta-train/meta-valid/meta-test)
<i>Bird</i>	64/16/20
<i>Texture</i>	30/7/10
<i>Aircraft</i>	64/16/20
<i>Fungi</i>	64/16/20
<i>CIFAR-FS</i>	64/16/20
<i>Mini-ImageNet</i>	64/16/20
<i>Omniglot</i>	71/15/16

Experiments: Few-shot Classification (Results)

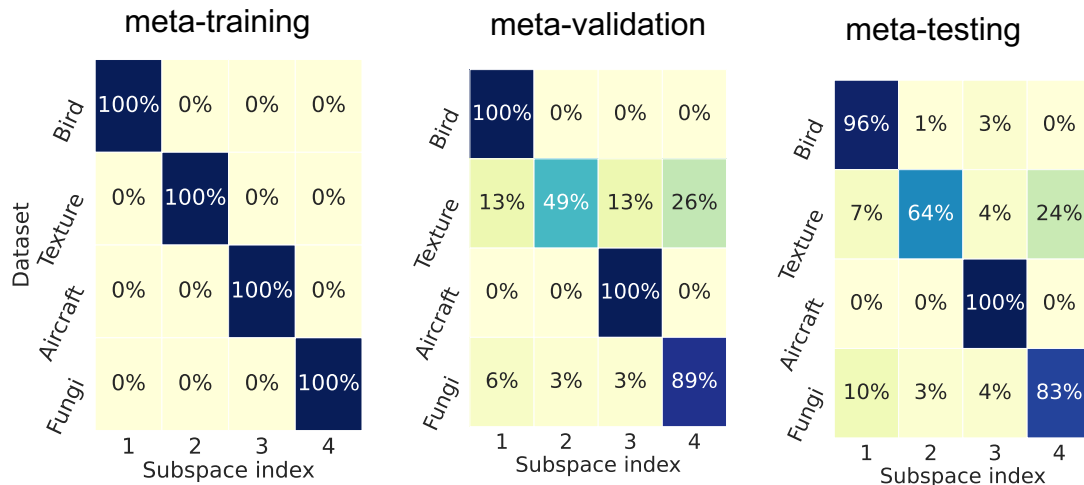
Accuracy of 5-way 5-shot classification

	<i>Meta-Dataset-BTAF</i>	<i>Meta-Dataset-ABF</i>	<i>Meta-Dataset-CIO</i>
MAML	57.78	63.86	74.46
ProtoNet	62.29	65.62	76.51
ANIL	58.57	64.43	74.61
BMG	60.10	65.80	77.46
DPMM	63.00	66.26	76.63
TSA-MAML	63.20	68.17	76.89
HSML	62.39	64.17	75.54
ARML	63.95	64.52	76.12
TSA-ProtoNet	63.57	68.77	77.27
MUSML	66.18	71.10	77.83

- MUSML is more accurate than both **structured** and **unstructured** meta-learning methods.

Experiments: Few-shot Classification (Results)

task assignment to the learned subspaces
(5-way 5-shot setting on *Meta-Dataset-BTAF*)



➤ MUSML discovers task structure.

Experiments: Cross-Domain Few-shot Classification

accuracy of cross-domain 5-way 5-shot setting
(*Meta-Dataset-BTAF* -> *Meta-Dataset-CIO*)

MAML	64.25
ProtoNet	66.13
ANIL	65.19
BMG	66.98
DPMM	66.73
TSA-MAML	66.85
HSML	65.18
ARML	65.37
TSA-ProtoNet	66.92
MUSML	67.41

- MUSML is also effective on unseen domains.

Experiments: Improving Existing Meta-learning Algorithms

accuracy of 5-way 5-shot classification

	<i>Meta-Dataset-BTAF</i>	<i>Meta-Dataset-ABF</i>	<i>Meta-Dataset-CIO</i>
Meta-SGD	58.93	64.19	75.95
MUSML-SGD	65.72	69.15	77.48
Meta-Curvature	60.02	64.51	76.13
MUSML-Curvature	66.10	69.23	77.96

- A subspace mixture is beneficial for both Meta-SGD and Meta-Curvature.
- MUSML can be used for any meta-learning algorithms.

Summary

- Problem: meta-learning in complex environment (task models are diverse)
- MUSML: learning a subspace mixture for building task models
- Each subspace can be viewed as a type of meta-knowledge
- Experimental results verify the effectiveness of MUSML