
EXPLAINABLE LEARNING WITH HIERARCHICAL ONLINE DETERMINISTIC ANNEALING

Christos N. Mavridis and John S. Baras

Department of Electrical and Computer Engineering
Institute for Systems Research
University of Maryland

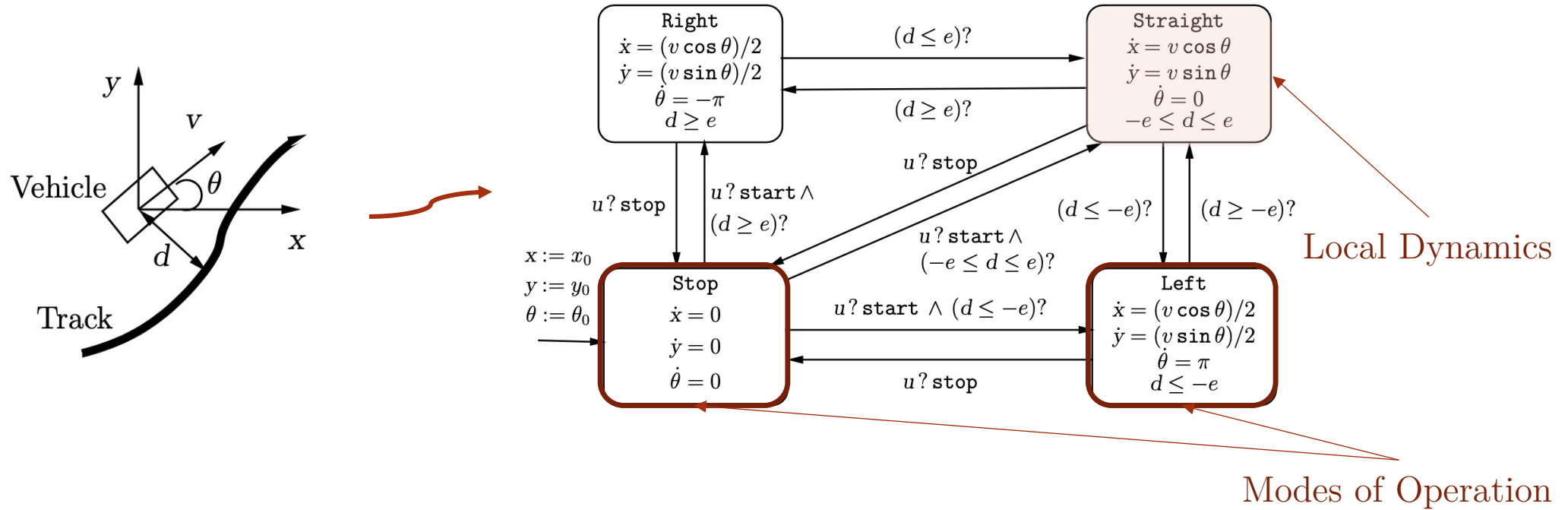


ECML PKDD 2023
Uncertainty meets Explainability

The
Institute for
Systems
Research

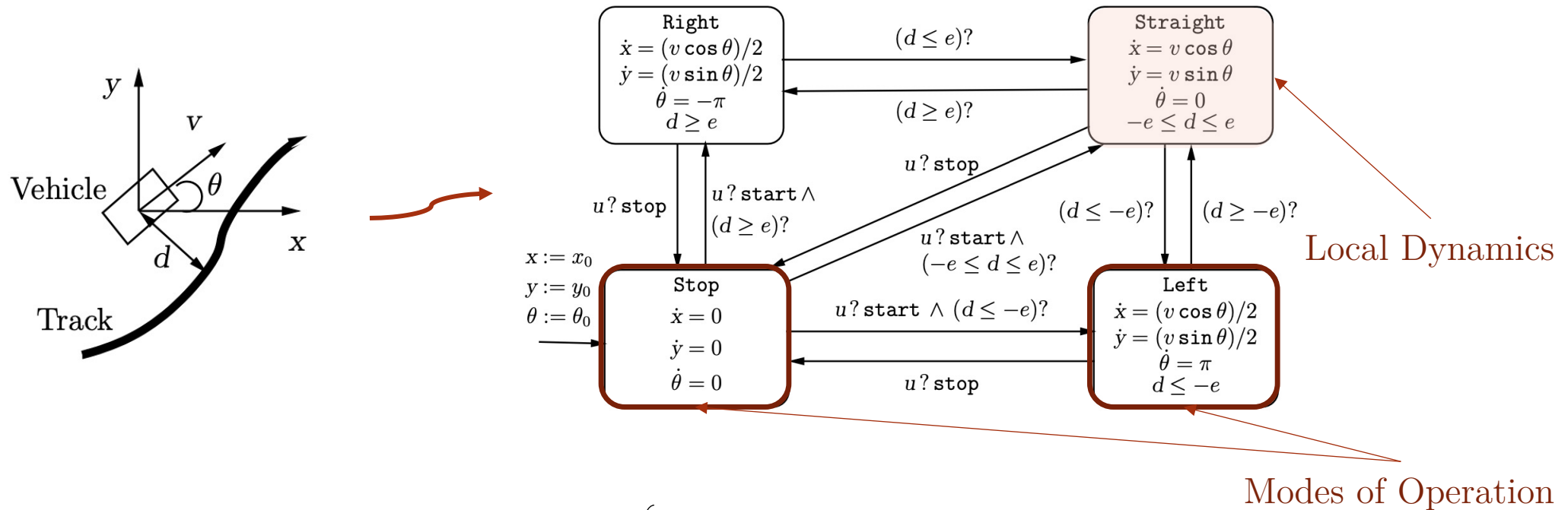
Explainable Learning – The Control-Theoretic Perspective

➤ Autonomous Vehicle Control



Explainable Learning – The Control-Theoretic Perspective

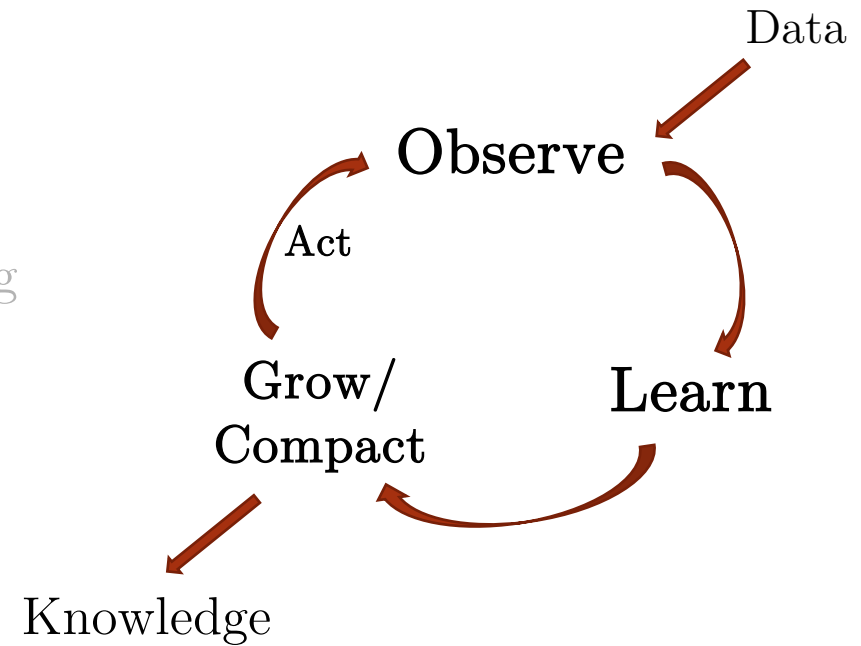
➤ Autonomous Vehicle Control



- Intelligent Autonomous Systems:
- How many modes?
 - Local System Identification?
 - Simultaneous, real-time learning?

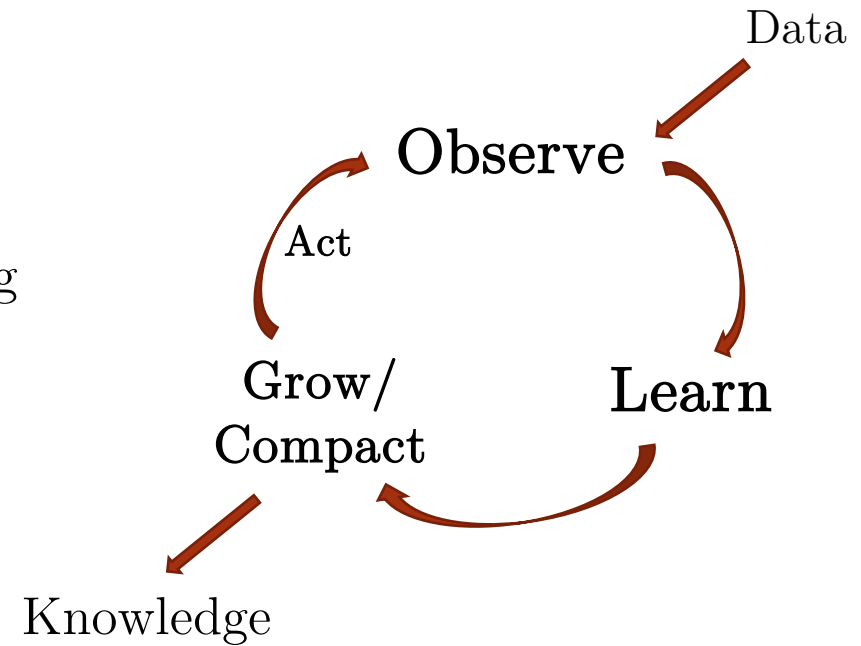
Learning Properties in Cyber-Physical Systems

- **Continuous/Dynamic/Adaptive Process**
- Interpretation
 - Why and when doesn't it work?
 - Knowledge Representation and Reasoning
- Robustness
 - Model uncertainty, overfitting, etc.
 - Transfer to real system?
- **Time and Memory Efficiency**
 - Real-time?
 - Processing/Communication bandwidth
 - Hyperparameter-tuning
 - Performance-Complexity Trade-off
 - Hierarchical Learning?



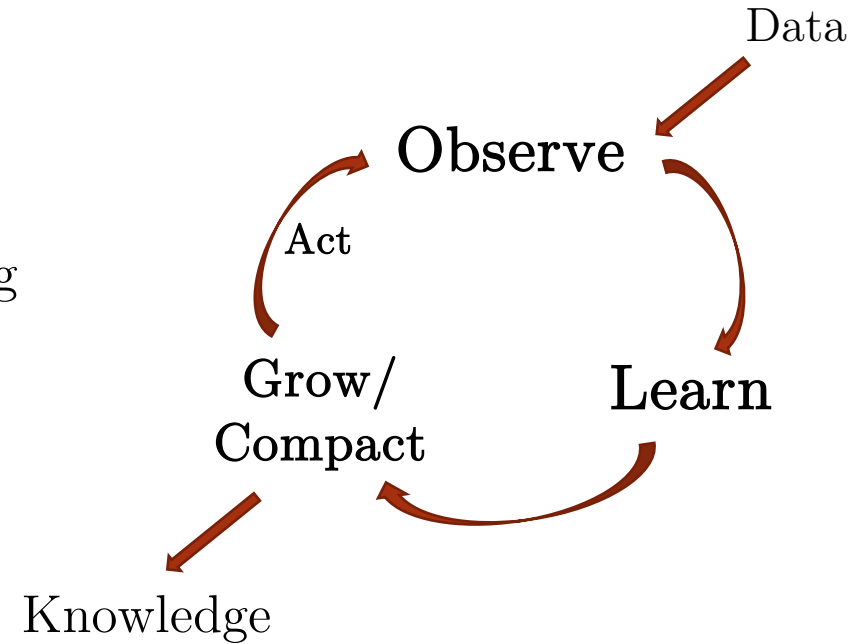
Learning Properties in Cyber-Physical Systems

- **Continuous/Dynamic/Adaptive Process**
- **Interpretation**
 - Why and when doesn't it work?
 - Knowledge Representation and Reasoning
- **Robustness**
 - Model uncertainty, overfitting, etc.
 - Transfer to real system?
- **Time and Memory Efficiency**
 - Real-time?
 - Processing/Communication bandwidth
 - Hyperparameter-tuning
 - Performance-Complexity Trade-off
 - Hierarchical Learning?



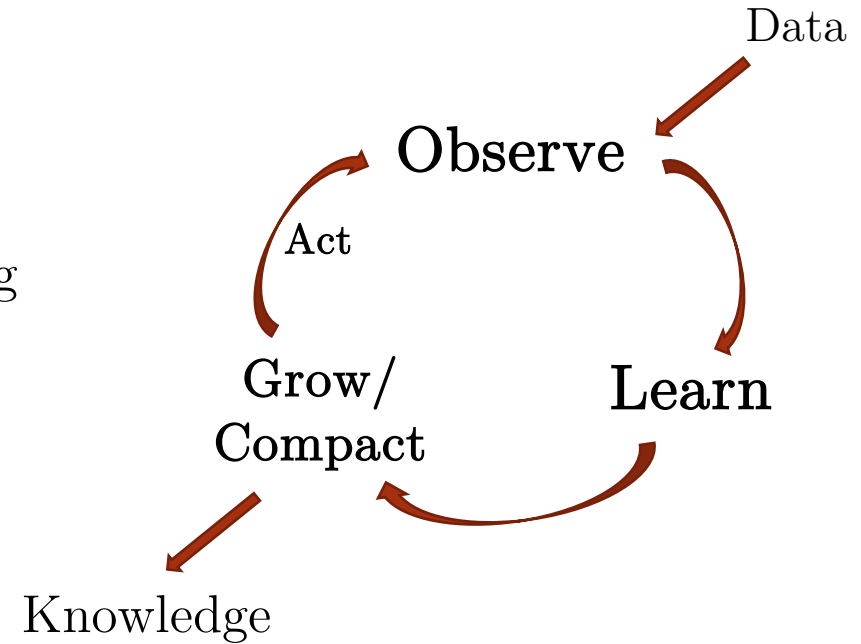
Learning Properties in Cyber-Physical Systems

- **Continuous/Dynamic/Adaptive Process**
- **Interpretation**
 - Why and when doesn't it work?
 - Knowledge Representation and Reasoning
- **Robustness**
 - Model uncertainty, overfitting, etc.
 - Transfer to real system?
- **Time and Memory Efficiency**
 - Real-time?
 - Processing/Communication bandwidth
 - Hyperparameter-tuning
 - Performance-Complexity Trade-off
 - Hierarchical Learning?



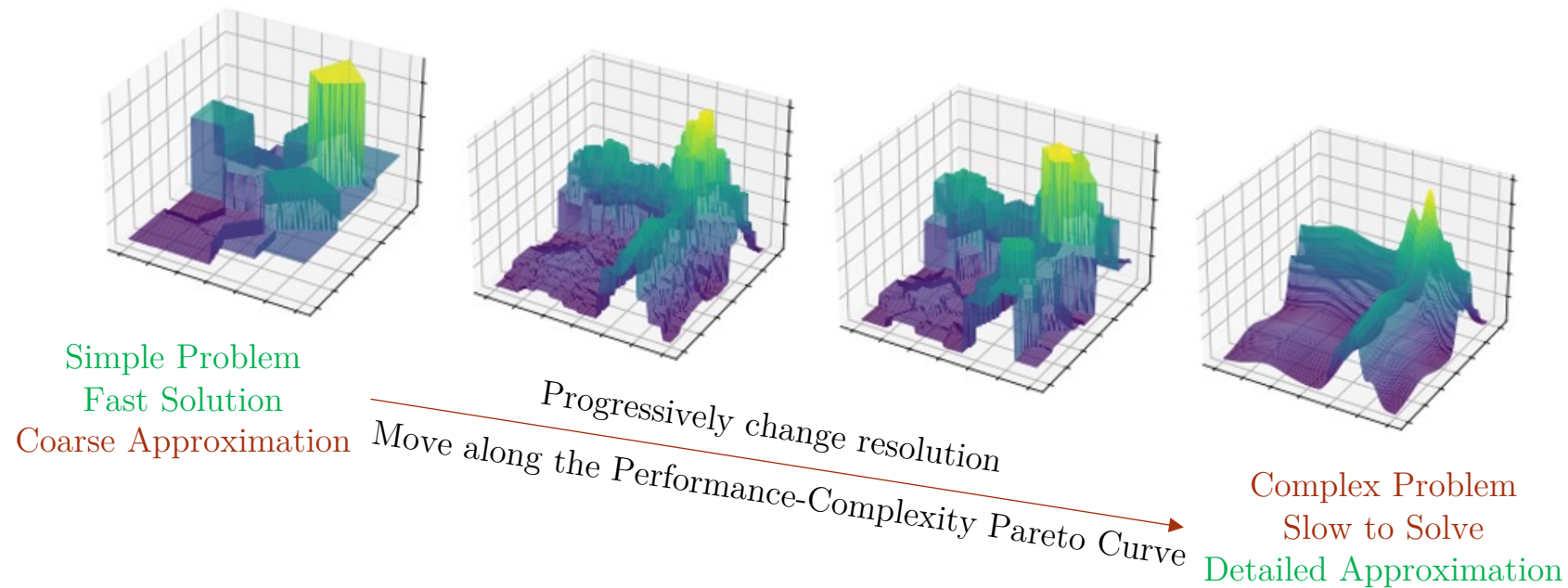
Learning Properties in Cyber-Physical Systems

- **Continuous/Dynamic/Adaptive Process**
- **Interpretation**
 - Why and when doesn't it work?
 - Knowledge Representation and Reasoning
- **Robustness**
 - Model uncertainty, overfitting, etc.
 - Transfer to real system?
- **Time and Memory Efficiency**
 - Real-time?
 - Processing/Communication bandwidth
 - Hyperparameter-tuning
 - Performance-Complexity Trade-off
 - Hierarchical Learning?



Towards Explainable Hierarchical Learning

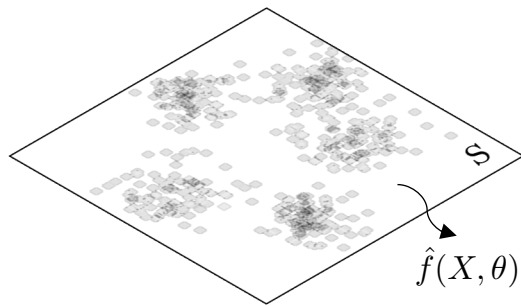
- Goal: Hierarchically Approximate Optimal Solutions*



* function approximation, reinforcement learning, game policies, system identification, clustering/classification

Towards Explainable Hierarchical Learning (II)

- **Divide and Conquer:** Partition the space and use local models



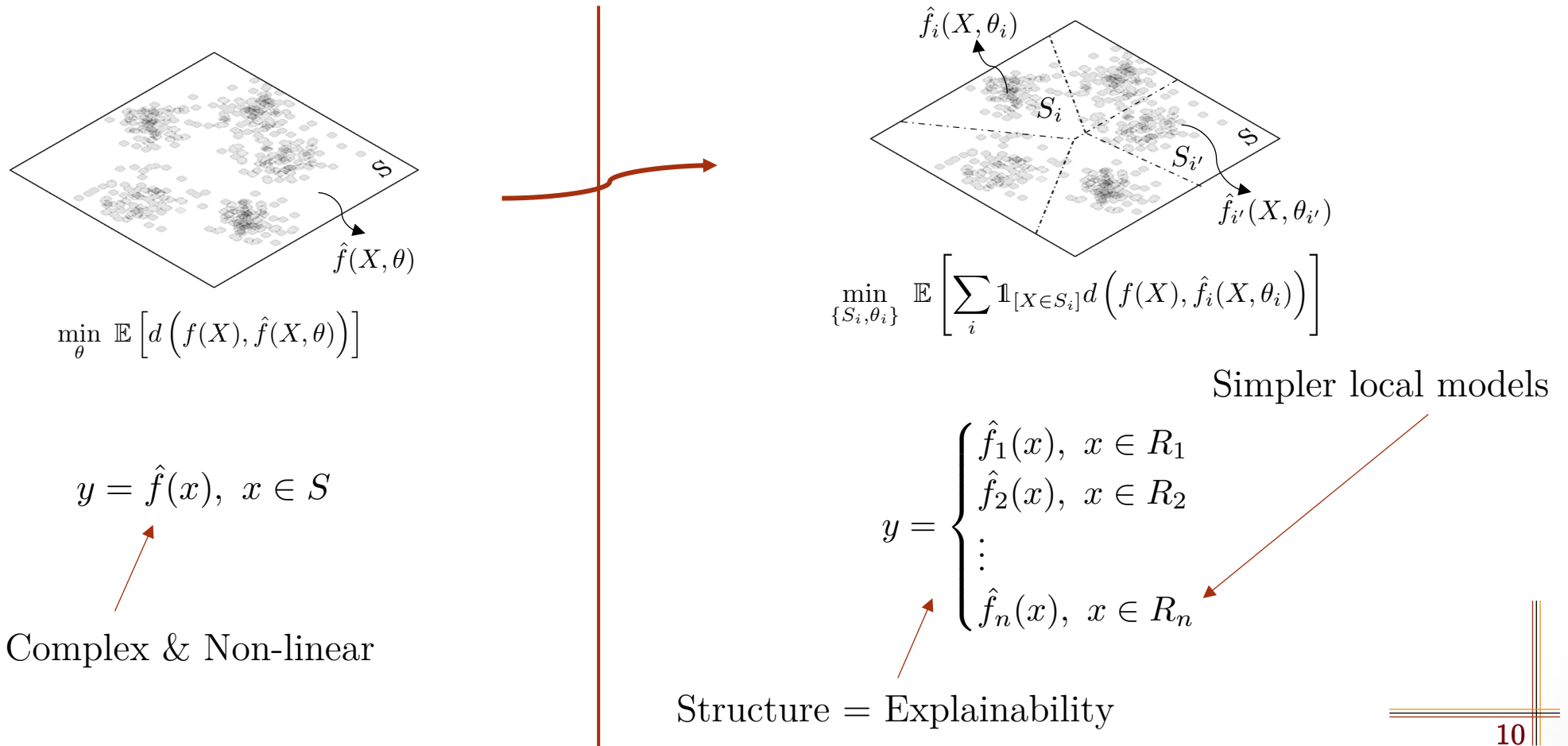
$$\min_{\theta} \mathbb{E} \left[d \left(f(X), \hat{f}(X, \theta) \right) \right]$$

$$y = \hat{f}(x), \quad x \in S$$

Highly Complex & Non-linear

Towards Explainable Hierarchical Learning (II)

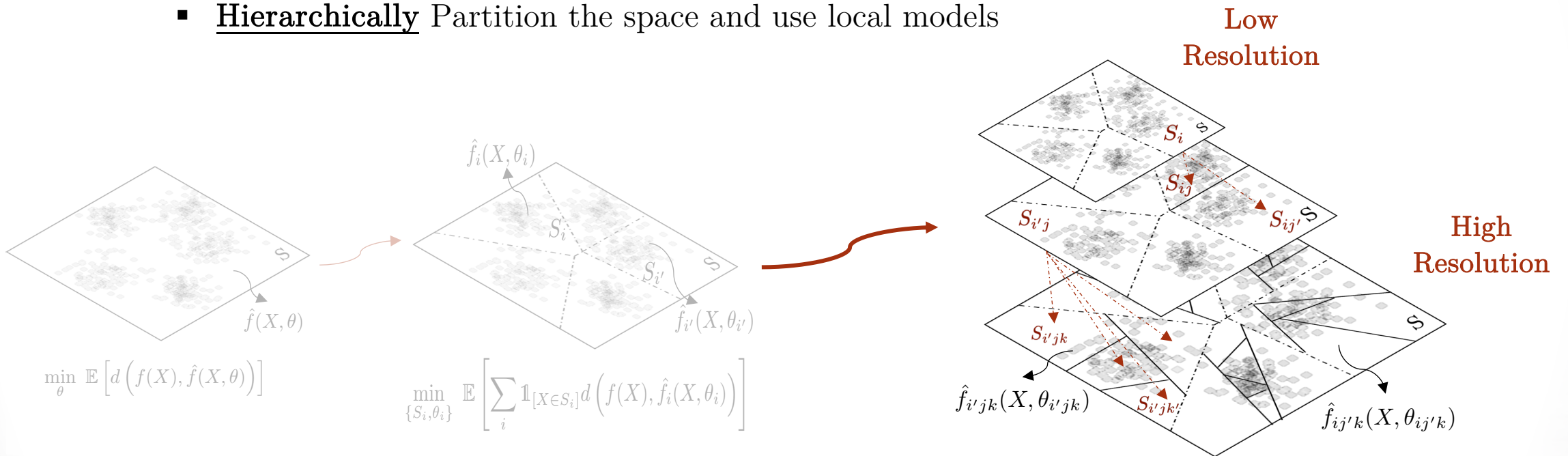
- **Divide and Conquer:** Partition the space and use local models



Towards Explainable Hierarchical Learning (II)

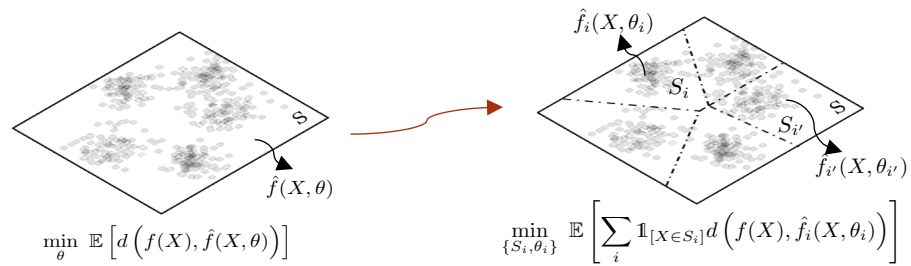
➤ Divide and Conquer

- Hierarchically Partition the space and use local models



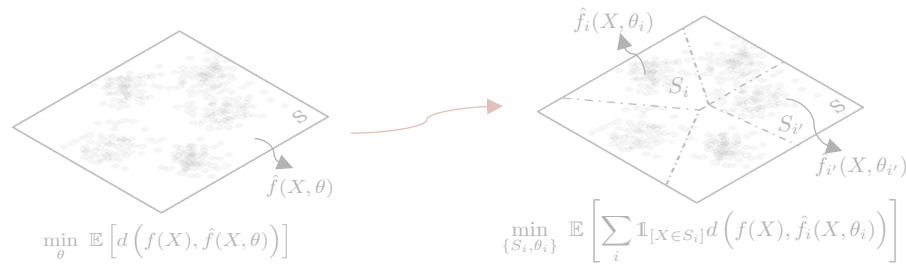
Towards Explainable Hierarchical Learning (III)

➤ Problems with Simultaneous Partitioning and Local Learning?



Towards Explainable Hierarchical Learning (III)

➤ Problems with Simultaneous Partitioning and Local Learning?



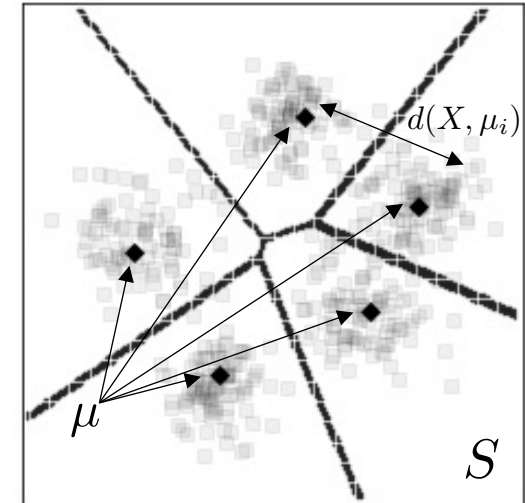
➤ Problems:

- How many regions?
 - Start with few and add as needed?
- Optimal parameters?
 - Local minima? Gradients?
 - Robustness?
- Simultaneously learn local models?

➔ Online Deterministic Annealing

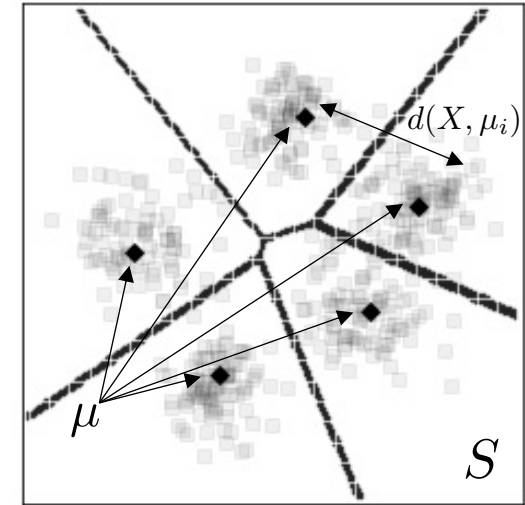
Online Deterministic Annealing

- *Observations:* $X^N := \{x_i\}_{i=1}^N$, $x_i \in S$ realizations of a r.v. $X \in S$
- *Codevectors:* $\mu = \{\mu_i\}_{i=1}^M$, $\mu_i \in S$ domain of a r.v. $Q \in S$
defined by: $p(\mu_i|x) = \mathbb{P}[Q = \mu_i|X = x]$
- *Dissimilarity:* $d : S \times S \rightarrow [0, \infty)$



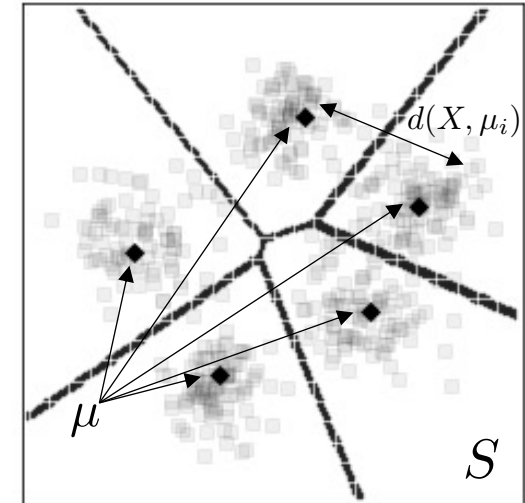
Online Deterministic Annealing

- *Observations:* $X^N := \{x_i\}_{i=1}^N$, $x_i \in S$ realizations of a r.v. $X \in S$
- *Codevectors:* $\mu = \{\mu_i\}_{i=1}^M$, $\mu_i \in S$ domain of a r.v. $Q \in S$
defined by: $p(\mu_i|x) = \mathbb{P}[Q = \mu_i|X = x]$
- *Dissimilarity:* $d : S \times S \rightarrow [0, \infty)$



Online Deterministic Annealing

- *Observations:* $X^N := \{x_i\}_{i=1}^N$, $x_i \in S$ realizations of a r.v. $X \in S$
- *Codevectors:* $\mu = \{\mu_i\}_{i=1}^M$, $\mu_i \in S$ domain of a r.v. $Q \in S$
defined by: $p(\mu_i|x) = \mathbb{P}[Q = \mu_i|X = x]$
- *Dissimilarity:* $d : S \times S \rightarrow [0, \infty)$

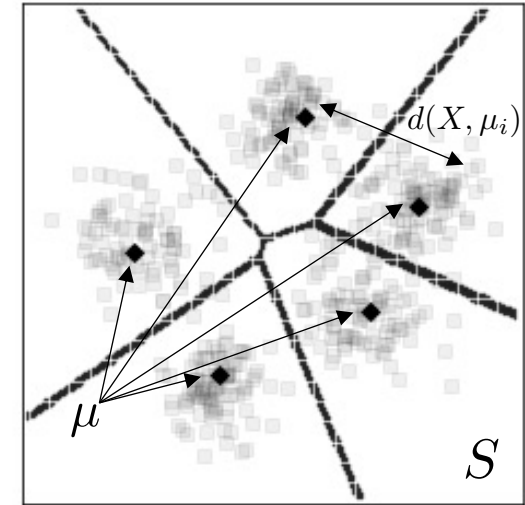


Online Deterministic Annealing

- *Observations:* $X^N := \{x_i\}_{i=1}^N$, $x_i \in S$ realizations of a r.v. $X \in S$
- *Codevectors:* $\mu = \{\mu_i\}_{i=1}^M$, $\mu_i \in S$ domain of a r.v. $Q \in S$
defined by: $p(\mu_i|x) = \mathbb{P}[Q = \mu_i|X = x]$
- *Dissimilarity:* $d : S \times S \rightarrow [0, \infty)$

Clustering?

$$\min_{\mu} D(X, Q) := \mathbb{E}[d(X, Q)] = \int p(x) \sum_i p(\mu_i|x) d(x, \mu_i) dx$$



Online Deterministic Annealing

- *Observations:* $X^N := \{x_i\}_{i=1}^N$, $x_i \in S$ realizations of a r.v. $X \in S$
- *Codevectors:* $\mu = \{\mu_i\}_{i=1}^M$, $\mu_i \in S$ domain of a r.v. $Q \in S$
defined by: $p(\mu_i|x) = \mathbb{P}[Q = \mu_i|X = x]$
- *Dissimilarity:* $d : S \times S \rightarrow [0, \infty)$

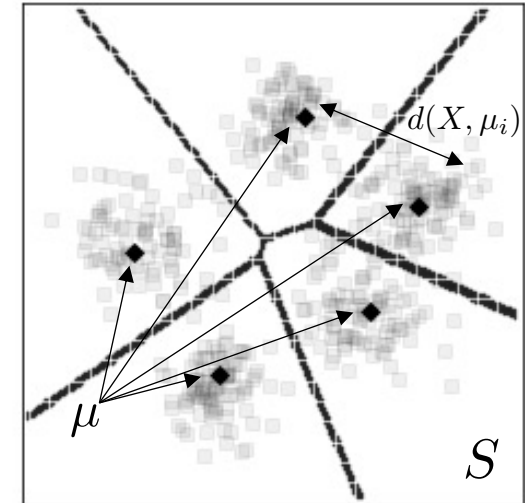
Clustering?

$$\min_{\mu} D(X, Q) := \mathbb{E}[d(X, Q)] = \int p(x) \sum_i p(\mu_i|x) d(x, \mu_i) dx$$

Online Deterministic Annealing

$$\min_{\mu} F_T := D - TH \quad \text{for decreasing values of } T.$$

where $\underbrace{H(X, Q)}_{\text{Entropy}} := \mathbb{E}[-\log P(X, Q)] = H(X) - \int p(x) \sum_i p(\mu_i|x) \log p(\mu_i|x) dx$



Adaptive
Robust
Progressive

Why Maximum Entropy?

➤ Jayne's Maximum Entropy Principle

- Most “Unbiased” estimator: each sub-problem induces “good” initial conditions for the next
- Duality (Legendre-type) and Regularization*:

$$\frac{1}{\beta} \log \mathbb{E}_{P_\mu} [e^{\beta Z}] = \inf_{P_\nu \in \mathcal{P}(\Omega)} \left\{ \mathbb{E}_{P_\nu} [Z] - \frac{1}{\beta} D_{KL}(P_\nu, P_\mu) \right\}, \beta < 0$$


$$\min F_T \simeq \min \frac{1}{\beta} \log \mathbb{E} [e^{\beta D}], \beta = -\frac{1}{T}$$

Risk-Sensitivity


$$\frac{1}{\beta} \log \mathbb{E} [e^{\beta J}] = \mathbb{E} [J] + \frac{\beta}{2} \text{Var} [J] + O(\beta^2)$$

➤ **Robustness** w.r.t. initial conditions, input perturbations.

➤ **Bifurcation:** Progressively grow set of models

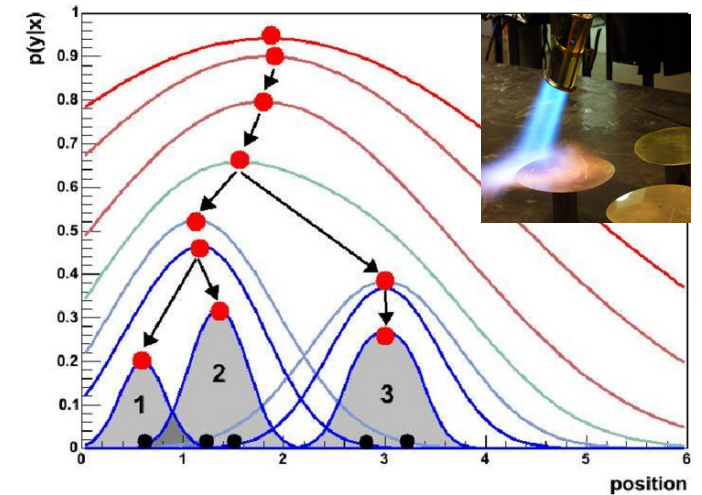
*Mavridis et al., Risk Sensitivity and Entropy Regularization in Prototype-based Learning, IEEE MED 2022.

Online Deterministic Annealing

Online Deterministic Annealing

Solve: $\min_{\mu} F_T := D - TH$ for decreasing values of T .

$\begin{cases} D(X, Q) : \text{Distortion} \\ H(X, Q) : \text{Entropy} \end{cases}$

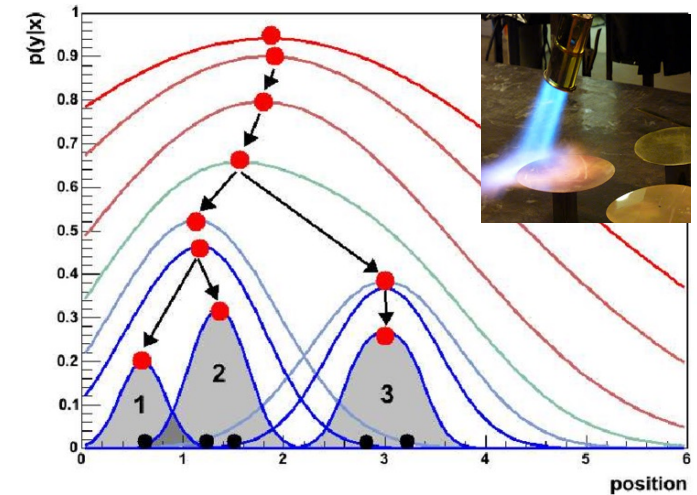


Online Deterministic Annealing

Online Deterministic Annealing

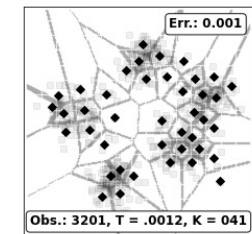
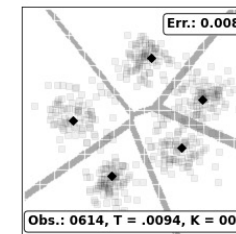
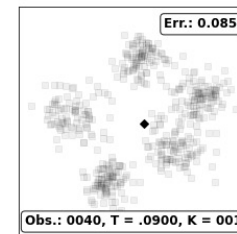
Solve: $\min_{\mu} F_T := D - TH$ for decreasing values of T .

$\begin{cases} D(X, Q) : \text{Distortion} \\ H(X, Q) : \text{Entropy} \end{cases}$



■ Lagrange (Temperature) Coefficient T

- Controls Performance/Complexity Tradeoff
- Simulates Annealing Optimization (Temperature)
- Stochastic Approximation
 - Simultaneous local system identification
- Triggers Bifurcation
 - Progressively adjust number of regions/codevectors



Online Deterministic Annealing (II)

Solving the Optimization Problem $\min F_T := D - TH$

- **Lemma.** The solution to $F^*(\mu) := \min_{\{p(\mu_i|x)\}} F(\mu)$
s.t. $\sum_i p(\mu_i|x) = 1$, is given by the Gibbs distributions
$$p^*(\mu_i|x) = \frac{e^{-\frac{d(x, \mu_i)}{T}}}{\sum_j e^{-\frac{d(x, \mu_j)}{T}}}, \quad \forall x \in S.$$

- **Theorem.** The solution to $\min_{\mu} F^*(\mu)$ is given by

$$\mu_i^* = \mathbb{E}[X|\mu_i] = \frac{\int x p(x) p^*(\mu_i|x) dx}{p^*(\mu_i)}$$

centroid form

if $d := d_\phi$ is a Bregman divergence. (sufficient condition)

e.g., squared Euclidean distance, KL divergence, ...



Online Deterministic Annealing (III)

Solving the Optimization Problem $\min F_T := D - TH$

► **Theorem.** *The dynamic stochastic process created by the recursive updates*

$$\mu_i(n+1) = \frac{\beta(n)}{\rho_i(n)} \left[\frac{\sigma_i(n+1)}{\rho_i(n+1)} (\rho_i(n) - \hat{p}(\mu_i|x_n)) + (x_n \hat{p}(\mu_i|x_n) - \sigma_i(n)) \right]$$

where the quantities $\rho_i(n)$, $\sigma_i(n)$, and $\hat{p}(\mu_i|x_n)$ are recursively updated by:

$$\begin{cases} \rho_i(n+1) &= \rho_i(n) + \alpha(n) [\hat{p}(\mu_i|x_n) - \rho_i(n)] \\ \sigma_i(n+1) &= \sigma_i(n) + \alpha(n) [x_n \hat{p}(\mu_i|x_n) - \sigma_i(n)] \end{cases}$$

$$\hat{p}(\mu_i|x_n) = \frac{\rho_i(n) e^{-\frac{d(x_n, \mu_i(n))}{T}}}{\sum_i \rho_i(n) e^{-\frac{d(x_n, \mu_i(n))}{T}}}$$

converges almost surely to a possibly sample path dependent solution of the optimization $\min_{\mu} F^*(\mu)$, as $n \rightarrow \infty$.

$$\mu_i(n) = \frac{\sigma_i(n) \rightarrow \mathbb{E}[\mathbf{1}_{[\mu]} X]}{\rho_i(n) \rightarrow \mathbb{P}[\mu]}$$

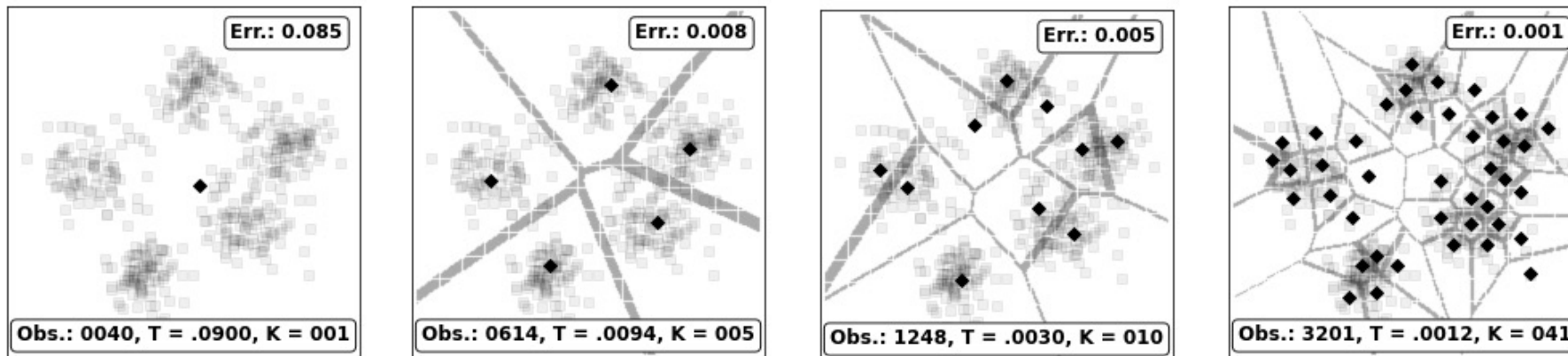
Stochastic Approximation: Gradient-Free !



Online Deterministic Annealing (IV)

Bifurcation and the number of codevectors

- Sequentially solve: $\min F_{T_\infty} := D - T_\infty H$
...
 $\min F_{T_0} := D - T_0 H$, $T_i < T_{i+1}$: Decreasing Temperature
- **Remark.** As $T \rightarrow \infty$, we get $\mu_i = \mathbb{E}[f(X)]$, $\forall i$, i.e., one unique pseudo-input.
- **Remark.** As T is lowered below a critical value, a bifurcation phenomenon occurs, and the number of pseudo-inputs increases.



Performance-Complexity Trade-off

Online Deterministic Annealing (V)

Training Local Models: Two-Timescale Stochastic Approximation

Algorithm 1 Online Deterministic Annealing

Initialize

while Termination Criterion **do**

Perturb $\mu^i \leftarrow \{\mu^i + \delta, \mu^i - \delta\}, \forall i$

repeat

Observe (x, c)

for $i = 1, \dots, K$ **do**

$s^i = \mathbb{1}_{[c_{\mu^i}=c]}$

Update:

$$p(\mu^i|x) \leftarrow \frac{p(\mu^i)e^{-\frac{d_\phi(x, \mu^i)}{T}}}{\sum_i p(\mu^i)e^{-\frac{d_\phi(x, \mu^i)}{T}}}$$

$$p(\mu^i) \leftarrow p(\mu^i) + \beta_t [s^i p(\mu^i|x) - p(\mu^i)]$$

$$\sigma(\mu^i) \leftarrow \sigma(\mu^i) + \beta_t [s^i x p(\mu^i|x) - \sigma(\mu^i)]$$

$$\mu^i \leftarrow \frac{\sigma(\mu^i)}{p(\mu^i)}$$

end for

until Convergence

Keep effective codevectors

Remove idle codevectors

Lower temperature $T \leftarrow \gamma T$

end while

$$\mu_{t+1} = \mu_t + \beta(t) [g(\theta_t, \mu_t) + M_{t+1}^{(\mu)}]$$

Slow SA

Partition

Online Deterministic Annealing (VI)

Training Local Models: Two-Timescale Stochastic Approximation

Algorithm 1 Online Deterministic Annealing

Initialize

while Termination Criterion **do**

Perturb $\mu^i \leftarrow \{\mu^i + \delta, \mu^i - \delta\}, \forall i$

repeat

Observe (x, c)

for $i = 1, \dots, K$ **do**

$s^i = \mathbb{1}_{[c_{\mu^i}=c]}$

Update:

$$p(\mu^i|x) \leftarrow \frac{p(\mu^i)e^{-\frac{d_\phi(x, \mu^i)}{T}}}{\sum_i p(\mu^i)e^{-\frac{d_\phi(x, \mu^i)}{T}}}$$

$$p(\mu^i) \leftarrow p(\mu^i) + \beta_t [s^i p(\mu^i|x) - p(\mu^i)]$$

$$\sigma(\mu^i) \leftarrow \sigma(\mu^i) + \beta_t [s^i x p(\mu^i|x) - \sigma(\mu^i)]$$

$$\mu^i \leftarrow \frac{\sigma(\mu^i)}{p(\mu^i)}$$

end for

until Convergence

Keep effective codevectors

Remove idle codevectors

Lower temperature $T \leftarrow \gamma T$

end while

$$\mu_{t+1} = \mu_t + \beta(t) [g(\theta_t, \mu_t) + M_{t+1}^{(\mu)}]$$

$$\theta_{t+1} = \theta_t + \alpha(t) [f(\theta_t, \mu_t) + M_{t+1}^{(\theta)}]$$

$$\frac{\beta(t)}{\alpha(t)} \rightarrow 0$$

Slow SA

Fast SA

Gradient Descent
Q-Learning
...

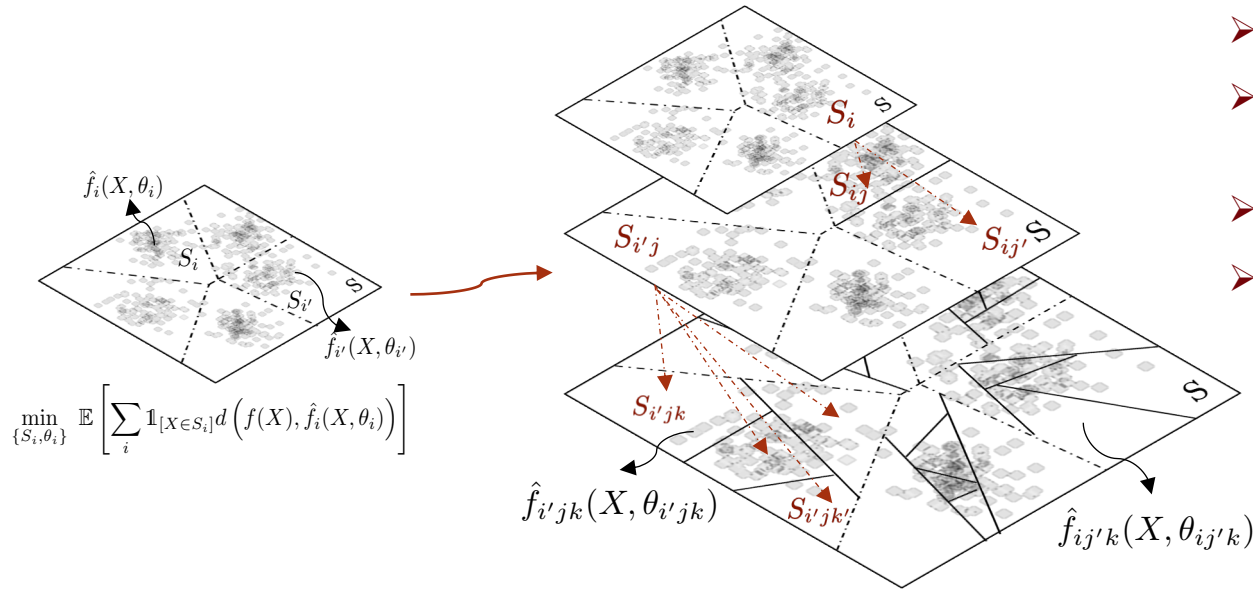
Partition

Local models

$$\min_{\{S_i, \theta_i\}} \mathbb{E} \left[\sum_i \mathbb{1}_{[X \in S_i]} d \left(f(X), \hat{f}_i(X, \theta_i) \right) \right]$$

Hierarchical Online Deterministic Annealing

Tree-Structured Hierarchical Learning



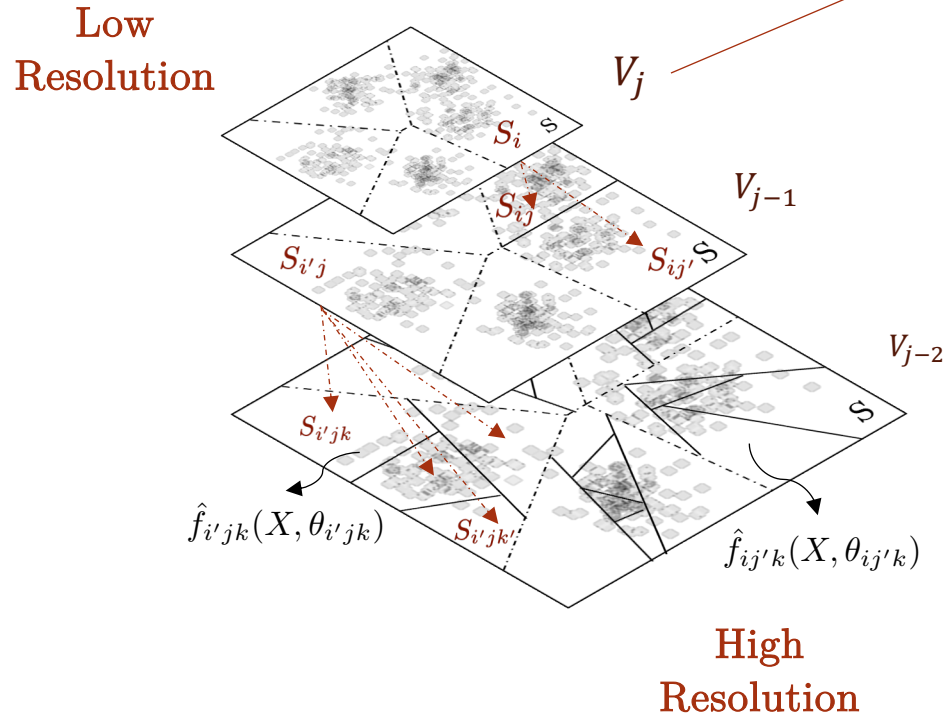
- Constructive (Structured Representation)
- Provably Consistent
- Localization
 - Emphasis on regions with high error
- Asynchronous/Parallel Computation
- Reduced Complexity

$$O\left(\frac{k^{\bar{l}} - 1}{k(k-1)} N_c (2\bar{k})^2 d\right)$$

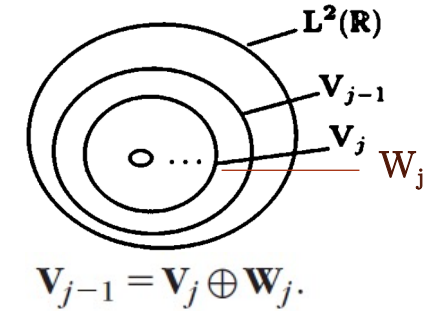
$$\bar{k} = \sum_{n=0}^{1/\bar{l} \log_2 K_{max}} 2^n$$

Hierarchical Online Deterministic Annealing

Multi-Resolution Hierarchical Learning



Example: Group-convolutional Wavelets



- Constructive (Structured Representation)
- Provably Consistent
- Localization
 - Emphasis on regions with high error
- Asynchronous/Parallel Computation
- Reduced Complexity

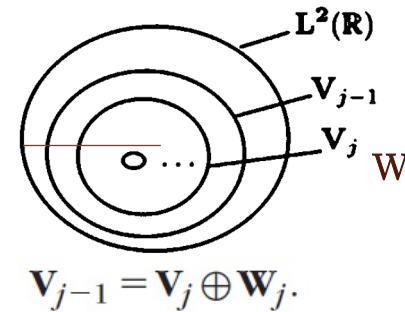
$$O\left(\frac{k^{\bar{l}} - 1}{k(k-1)} N_c (2\bar{k})^2 d\right)$$

$$\bar{k} = \sum_{n=0}^{1/\bar{l} \log_2 K_{max}} 2^n$$

Group-Convolutional Wavelets

- Wavelet Transform

- Multi-Resolution Analysis
- Sparse, Stable, Translation Covariant

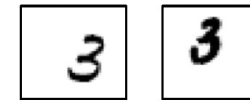


- Convolution on Groups

$$(f * g)(x) = \int_G f(y)g(y^{-1}x)d\lambda(y)$$

where for a Lie Group G : $g \in G \rightarrow g.f(x) := f(g^{-1}x)$

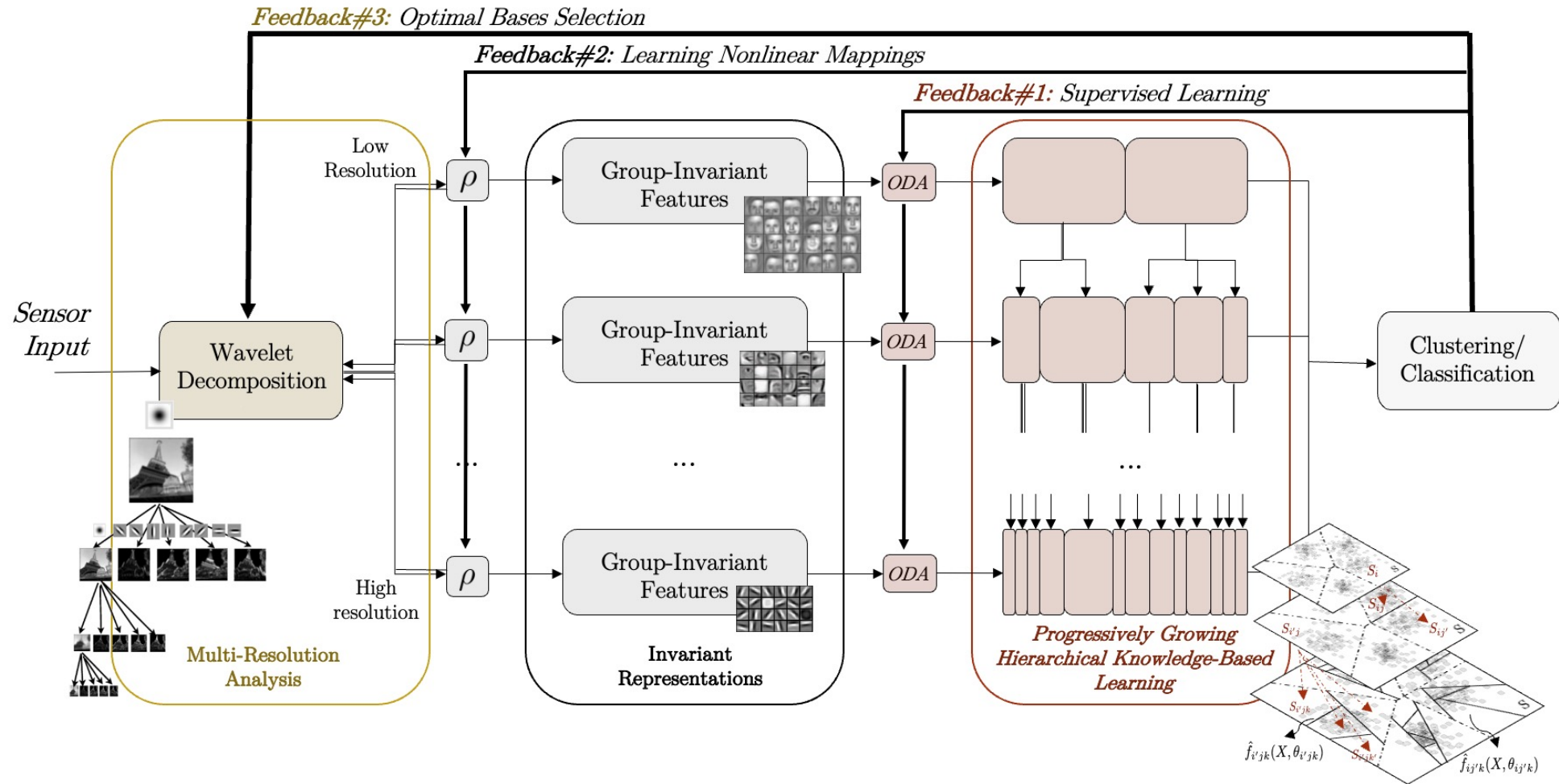
- Locally Group-Invariant Representations



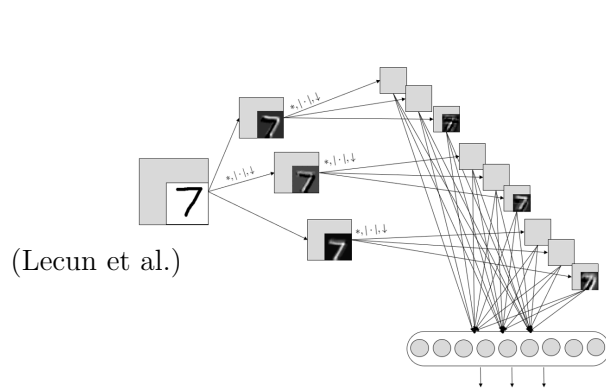
Repeat

- Build group-covariant representations (**wavelets**)
- Make them locally invariant (**non-linearity + averaging**)

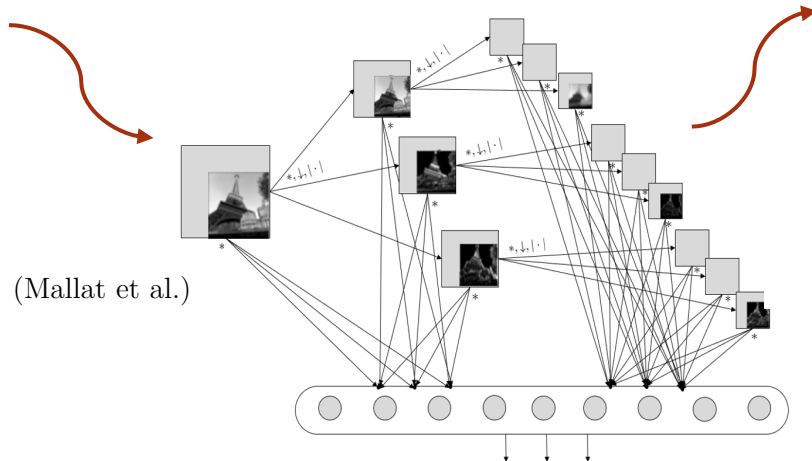
Closed-Loop Hierarchical Learning Architecture



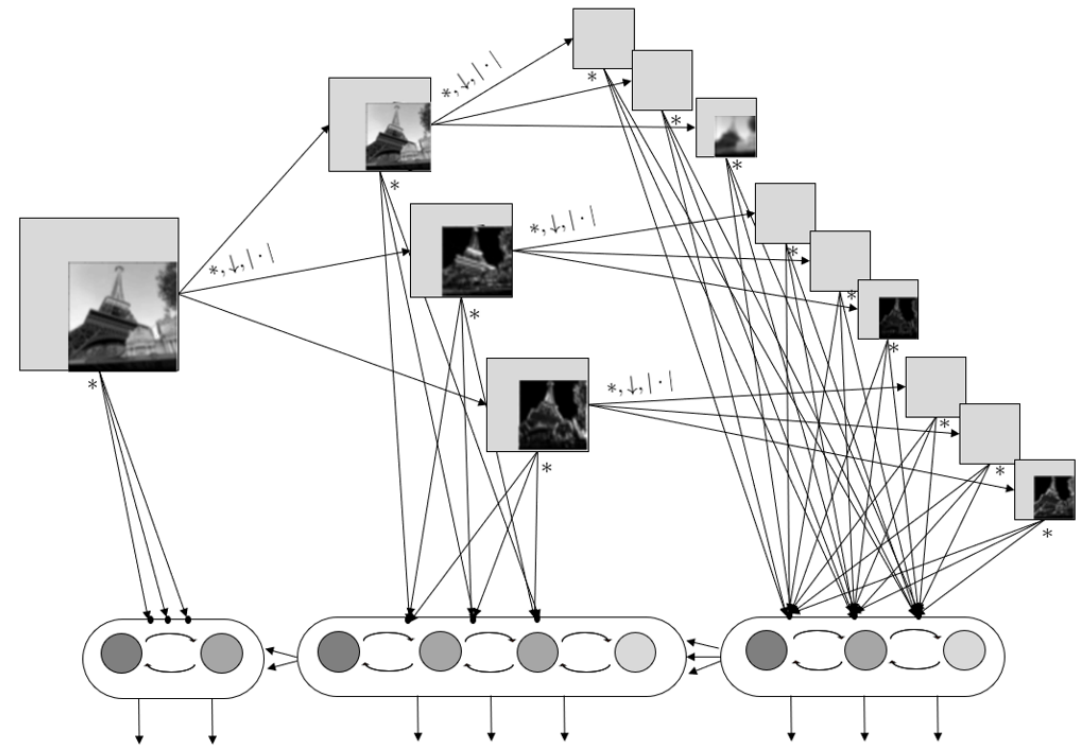
A Deep Learning Architecture



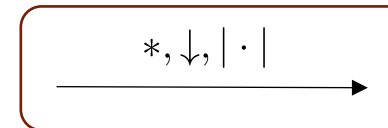
Deep Convolutional Network



Scattering Convolutional Network



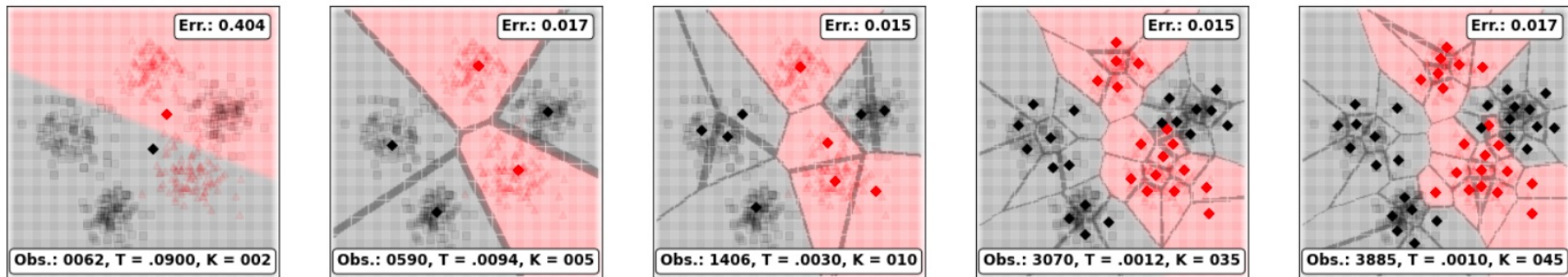
Our Approach



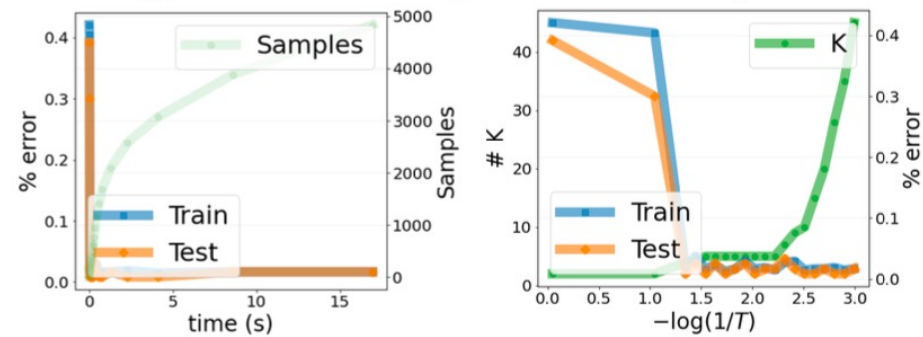
Simulation Results

➤ Single Resolution. Binary Classification on Mixture of Gaussians.

Performance-Complexity Trade-off



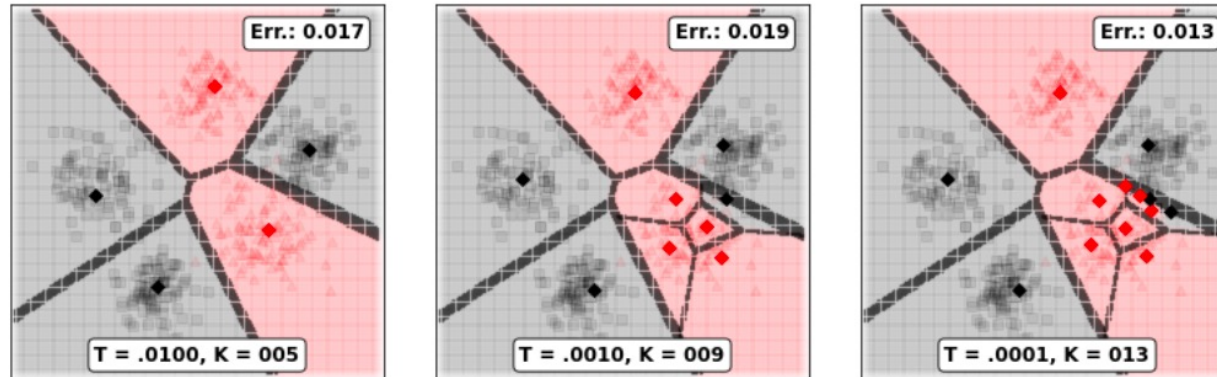
(a) Evolution of the algorithm in the data space.



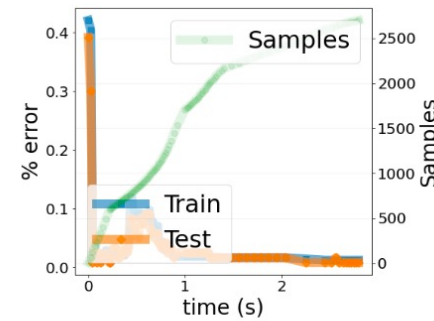
(b) Performance curves.

Simulation Results (II)

- **Single Resolution – Tree-Structured.** Binary Classification on Mixture of Gaussians.



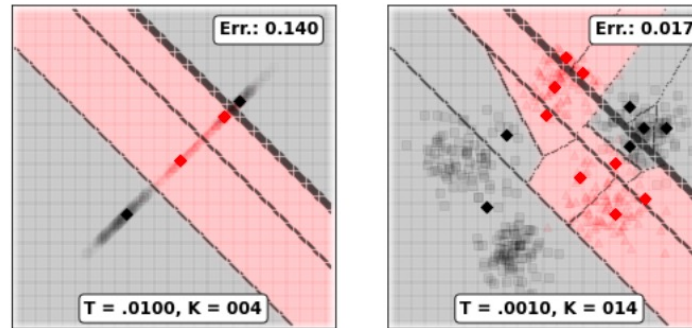
(a) Evolution of the algorithm in the data space.



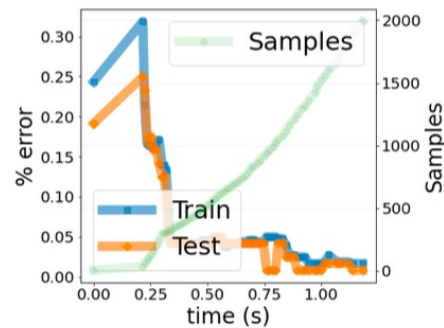
(b) Performance curves.

Simulation Results (III)

- **Multiple Resolutions w/ PCA.** Binary Classification on Mixture of Gaussians.



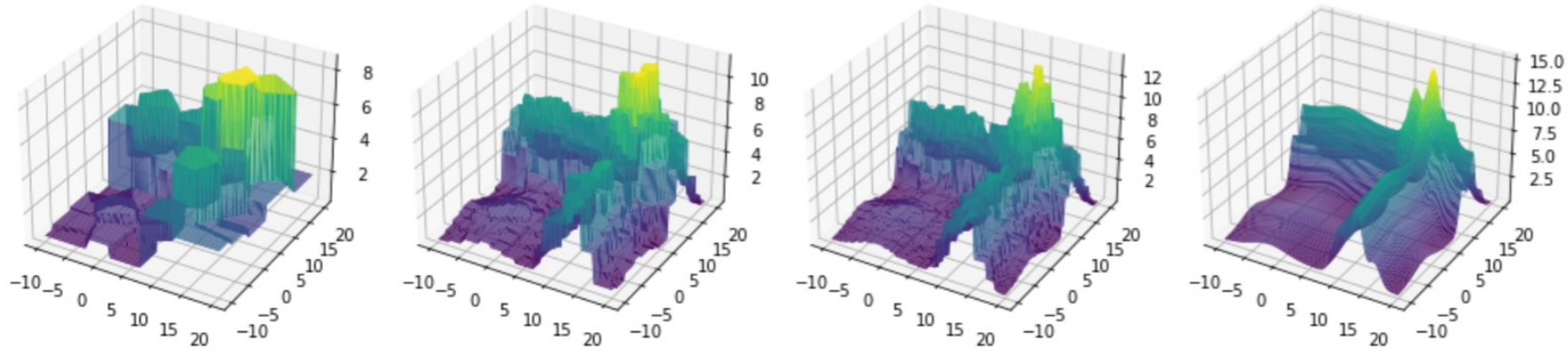
(a) Convergence of first layer with low-resolution features. (b) Convergence of second layer with high-resolution features.



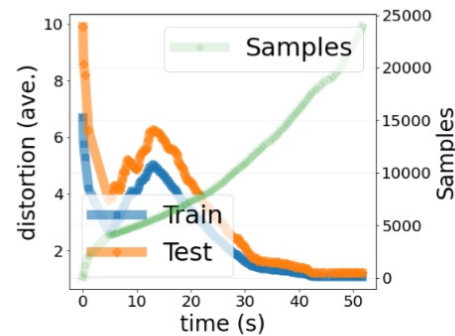
(c) Performance curve.

Simulation Results (IV)

➤ **Single Resolution.** 2D Regression with Constant Local Models.



(a) Evolution of the algorithm in the data space (original function on the right).

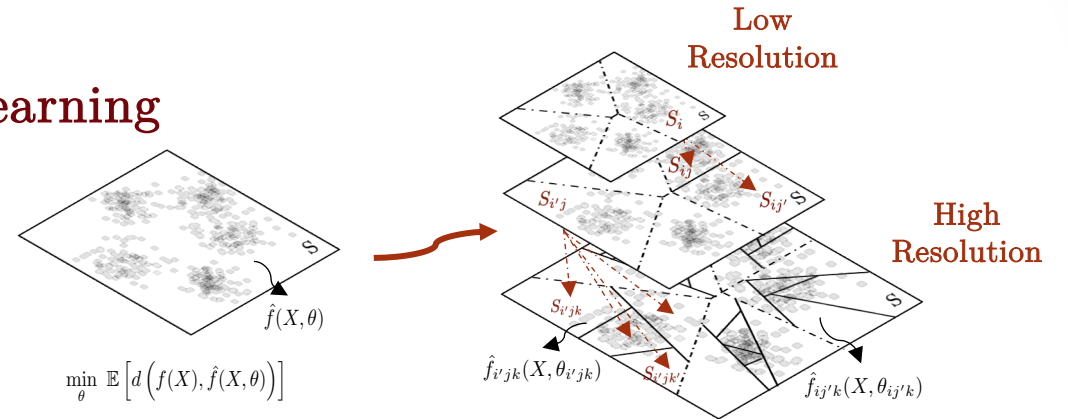


(b) Performance curves.

Thank you!

➤ Simultaneous Partitioning and Local Learning

- Explainability
- Robustness w.r.t. Init. & Noise



➤ Hierarchical Online Deterministic Annealing

- Multi-Resolution Partitioning
- Online, Adaptive, Gradient-Free
- Simultaneous local model training



Questions?

Christos Mavridis

mavridis@umd.edu

mavridis@kth.se

<https://mavridischristos.github.io/>

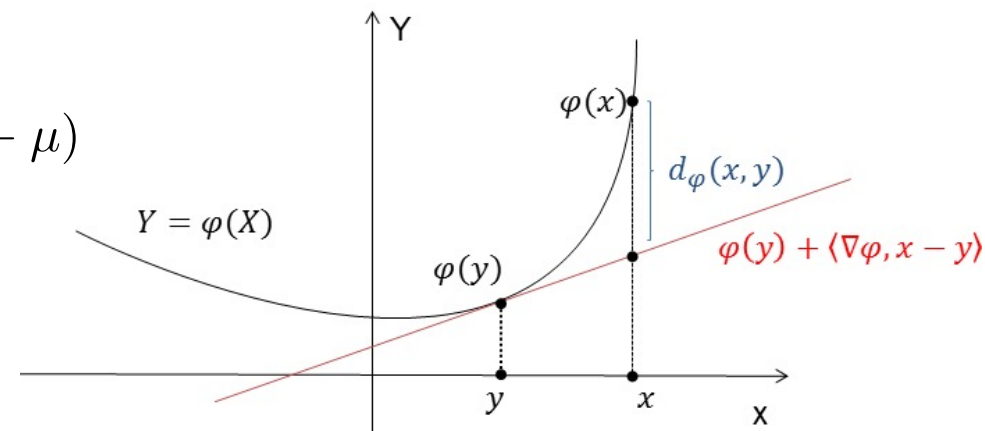


Bregman Divergences



►
$$d_{\phi}(x, \mu) = \phi(x) - \phi(\mu) - \frac{\partial \phi}{\partial \mu}(\mu)(x - \mu)$$

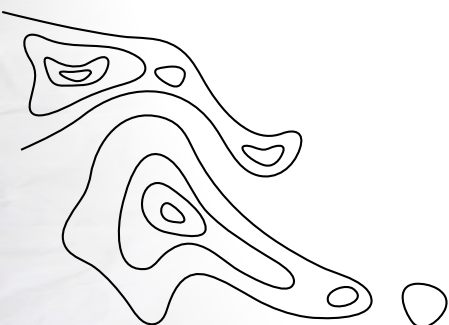
- Euclidean distance, KL divergence, ...

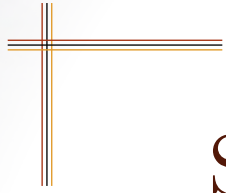


- **Theorem.** Let $X : \Omega \rightarrow S$ be a random variable defined in the probability space $(\Omega, \mathcal{F}, \mathbb{P})$ such that $\mathbb{E}[X] \in \text{ri}(S)$, and let a distortion measure $d : S \times \text{ri}(S) \rightarrow [0, \infty)$, where $\text{ri}(S)$ denotes the relative interior of S . Then

$$\mu := \mathbb{E}[X] \in \arg \min_{s \in \text{ri}(S)} \mathbb{E}[d(X, s)]$$

is the unique minimizer of $\mathbb{E}[d(X, s)]$ in $\text{ri}(S)$, if and only if d is a Bregman divergence for any function ϕ that satisfies the definition.





Stochastic Approximation



Theorem. *Almost surely, the sequence:*

$$x_{n+1} = x_n + \alpha(n) [h(x_n) + M_{n+1}], \quad n \geq 0 \quad (1)$$

converges to a (possibly sample path dependent) compact, connected, internally chain transitive, invariant set of the o.d.e:

$$\dot{x}(t) = h(x(t)), \quad t \geq 0, \quad (2)$$

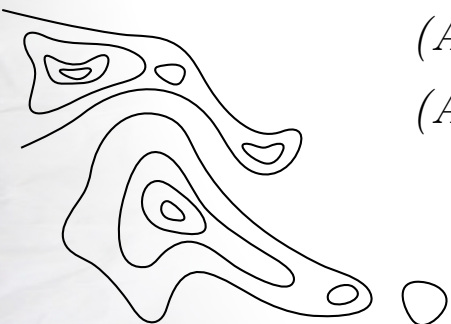
provided that:

- (A1) $h : \mathbb{R}^d \rightarrow \mathbb{R}^d$ is Lipschitz.
- (A2) $\sum_n \alpha(n) = \infty$, and $\sum_n \alpha^2(n) < \infty$
- (A3) $\{M_n\}$ is a martingale difference sequence
- (A4) $\{x_n\}$ remain bounded a.s.

Examples:

$$h(x) = \begin{cases} -\nabla J(x), & \text{SGD} \\ F(x) - x, & \text{Fixed-Point Iter.} \end{cases}$$

*Borkar, Stochastic approximation: a dynamical systems viewpoint, Springer, 2009



Bifurcation and the number of Codevectors



► **Theorem.** *Bifurcation occurs under the following condition*

$$\exists y_n \text{ s.t. } p(y_n) > 0 \text{ and } \det \left[I - T \frac{\partial^2 \phi(y_n)}{\partial y_n^2} C_{x|y_n} \right] = 0$$

where $C_{x|y_n} := \mathbb{E} [(x - y_n)(x - y_n)^T | y_n]$.

Proof. From variational calculus and the second order condition:

$$\frac{d^2}{d\epsilon^2} F^*(\{\mu + \epsilon\psi\})|_{\epsilon=0} \geq 0$$

- T_c depends on:
- The Bregman divergence
 - The data space

□

