

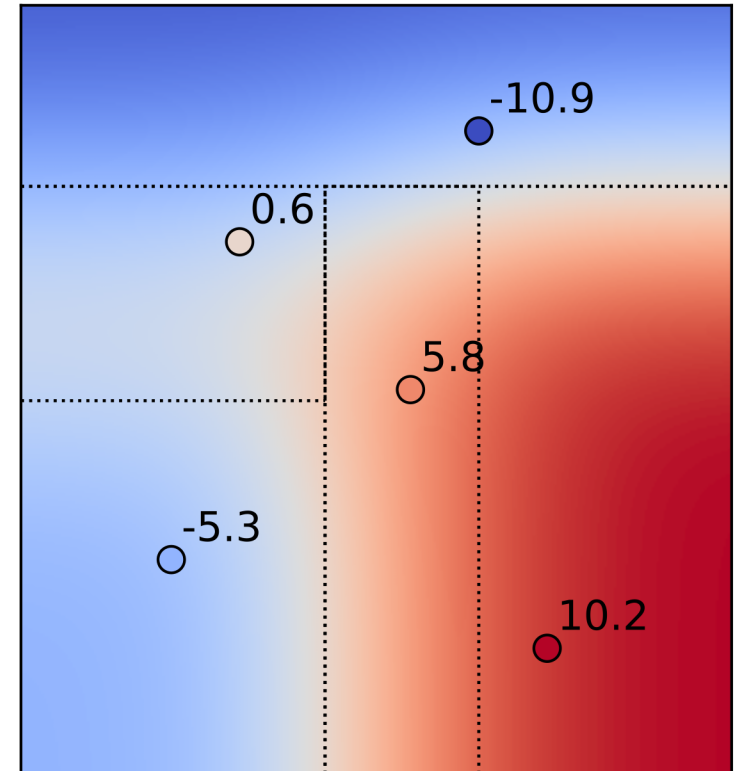
Improving Random Forests by Smoothing

Ziyi (Neil) Liu¹, Phuc Luong¹, **Mario Boley**^{2,1}, Daniel F. Schmidt¹

¹Department of Data Science and AI, Monash University

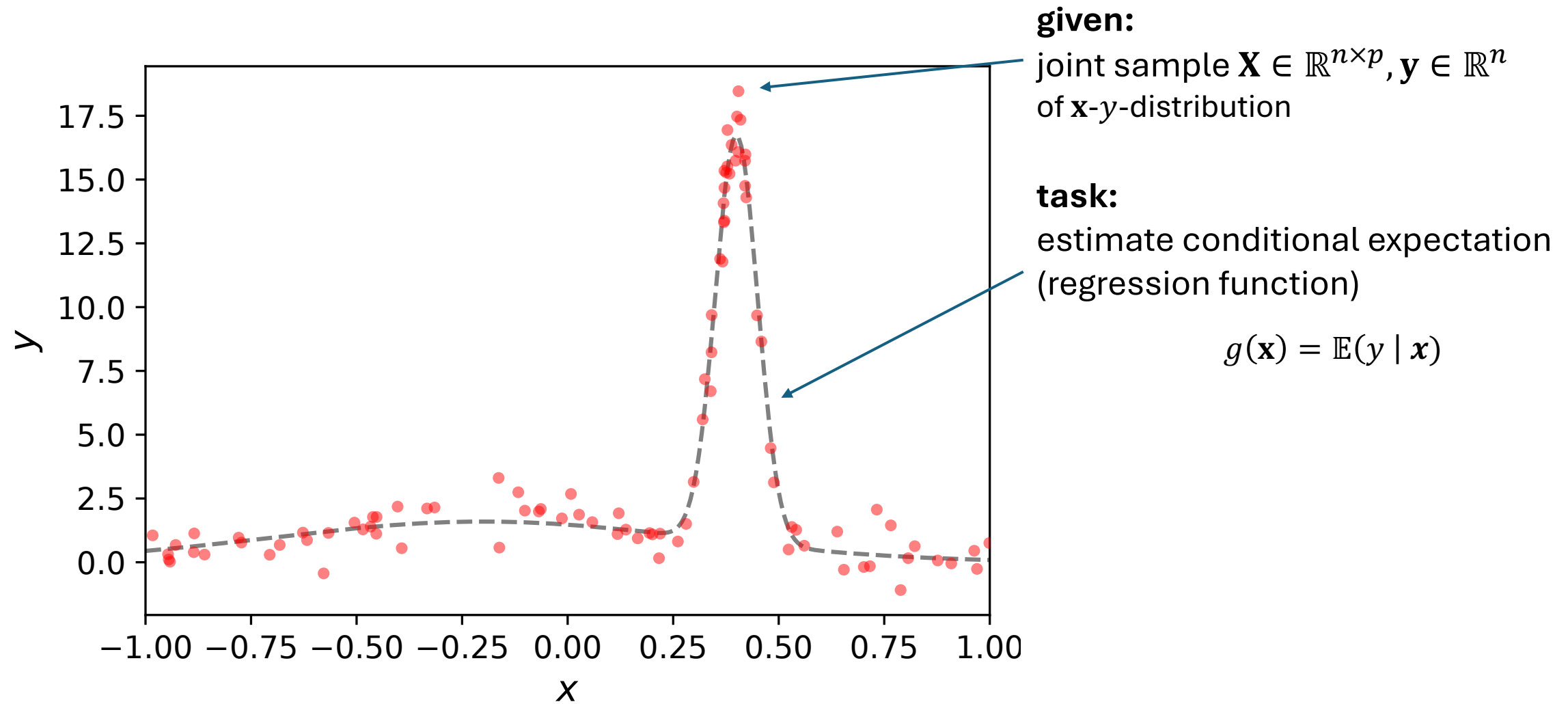
²Department for Information Systems, University of Haifa

`mboley@is.haifa.ac.il`



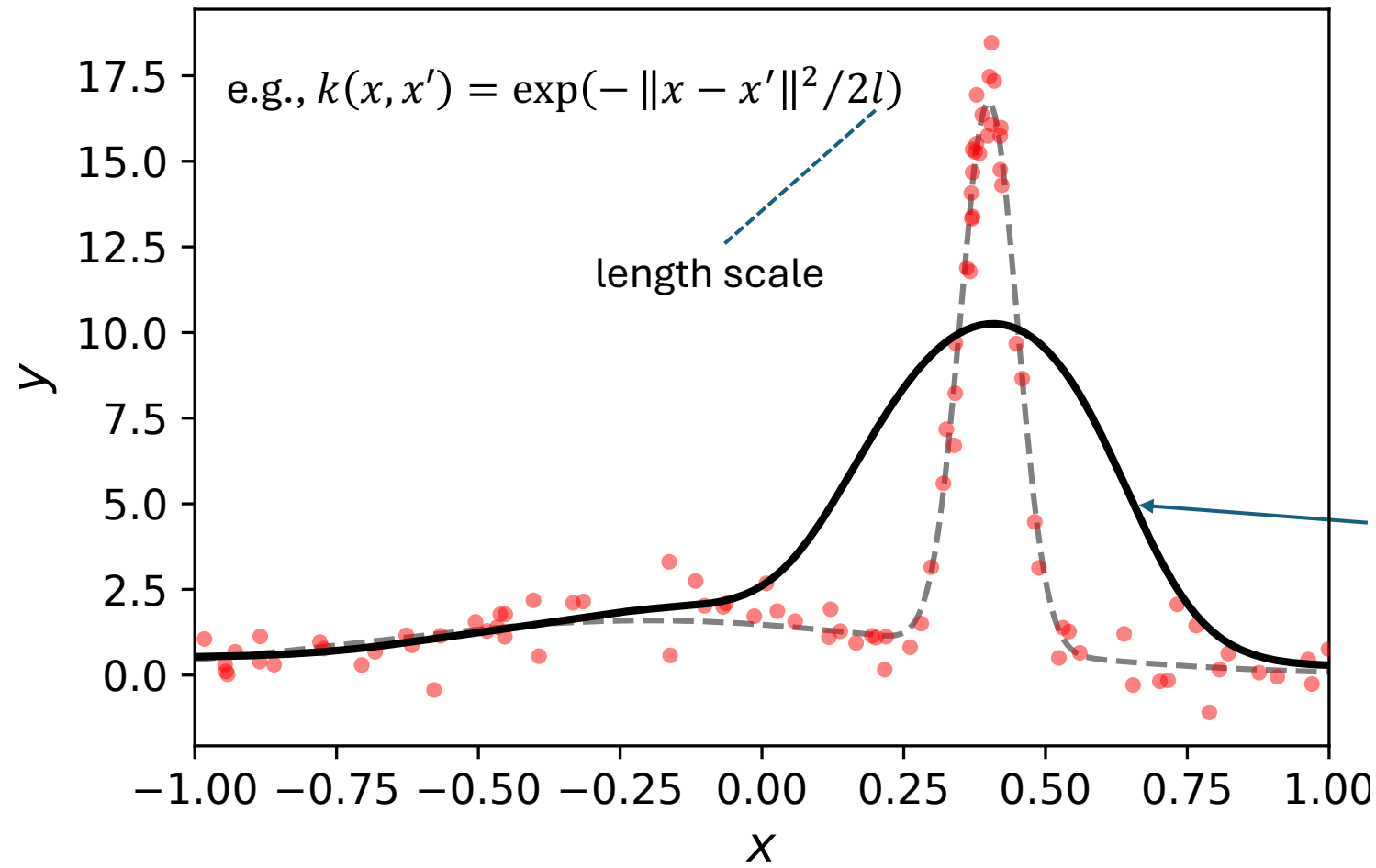
Non-Parametric Noisy Function Estimation

2



Issue of Kernel Smoothers: No Spatial Adaptivity

3



given:

joint sample $\mathbf{X} \in \mathbb{R}^{n \times p}$, $\mathbf{y} \in \mathbb{R}^n$
of $\mathbf{x}$ - $\mathbf{y}$ -distribution

task:

estimate conditional expectation
(regression function)

$$g(\mathbf{x}) = \mathbb{E}(y \mid \mathbf{x})$$

Gaussian process regression:

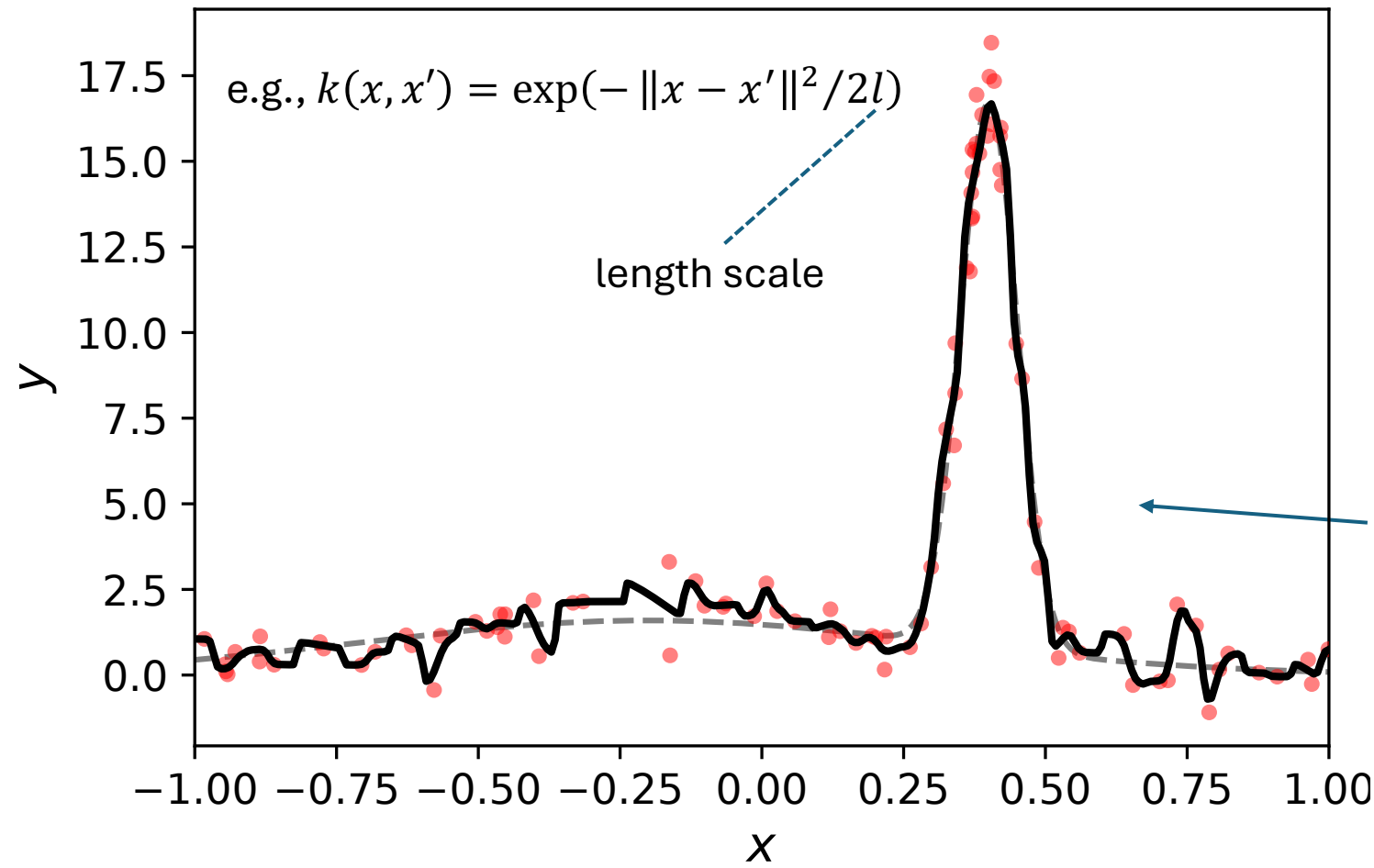
$$f_{\text{GPR}}(\mathbf{x}_0) = \mathbf{k}_0^\top (\mathbf{K} + \sigma^2 \mathbf{I})^{-1} \mathbf{y}$$

$$(k(\mathbf{x}_0, \mathbf{x}_1), \dots, k(\mathbf{x}_0, \mathbf{x}_n))^\top$$

kernel function

Issue of Kernel Smoothers: No Spatial Adaptivity

4



given:

joint sample $\mathbf{X} \in \mathbb{R}^{n \times p}$, $\mathbf{y} \in \mathbb{R}^n$
of $\mathbf{x}$ - $\mathbf{y}$ -distribution

task:

estimate conditional expectation
(regression function)

$$g(\mathbf{x}) = \mathbb{E}(y \mid \mathbf{x})$$

Gaussian process regression:

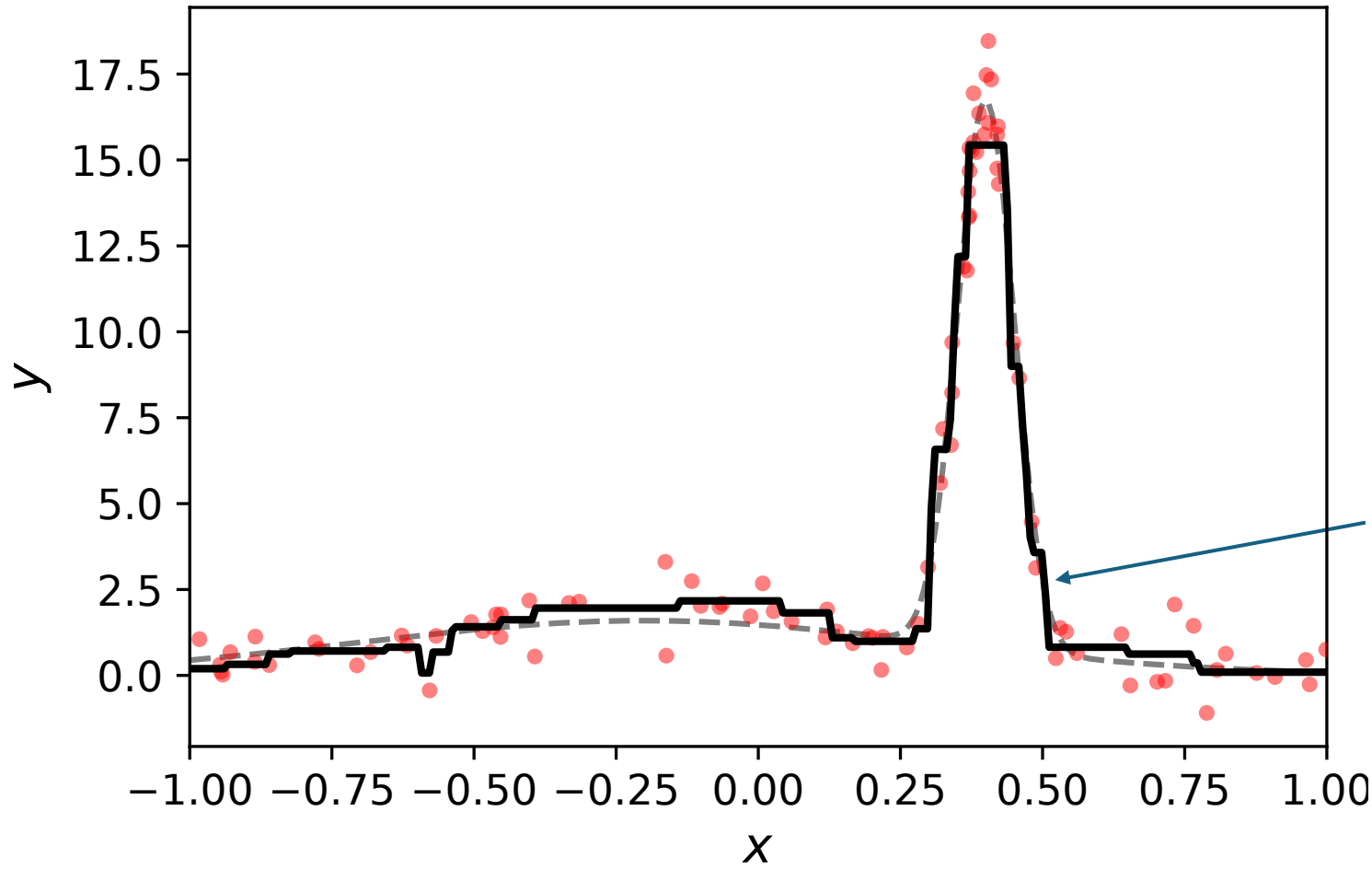
$$f_{\text{GPR}}(\mathbf{x}_0) = \mathbf{k}_0^\top (\mathbf{K} + \sigma^2 \mathbf{I})^{-1} \mathbf{y}$$

$$(k(\mathbf{x}_0, \mathbf{x}_1), \dots, k(\mathbf{x}_0, \mathbf{x}_n))^\top$$

kernel function

Issue of Tree-based Models: No Smoothness

5



given:

joint sample $\mathbf{X} \in \mathbb{R}^{n \times p}$, $\mathbf{y} \in \mathbb{R}^n$
of $\mathbf{x}$ - $\mathbf{y}$ -distribution

task:

estimate conditional expectation
(regression function)

$$g(\mathbf{x}) = \mathbb{E}(y \mid \mathbf{x})$$

Random forest regression:

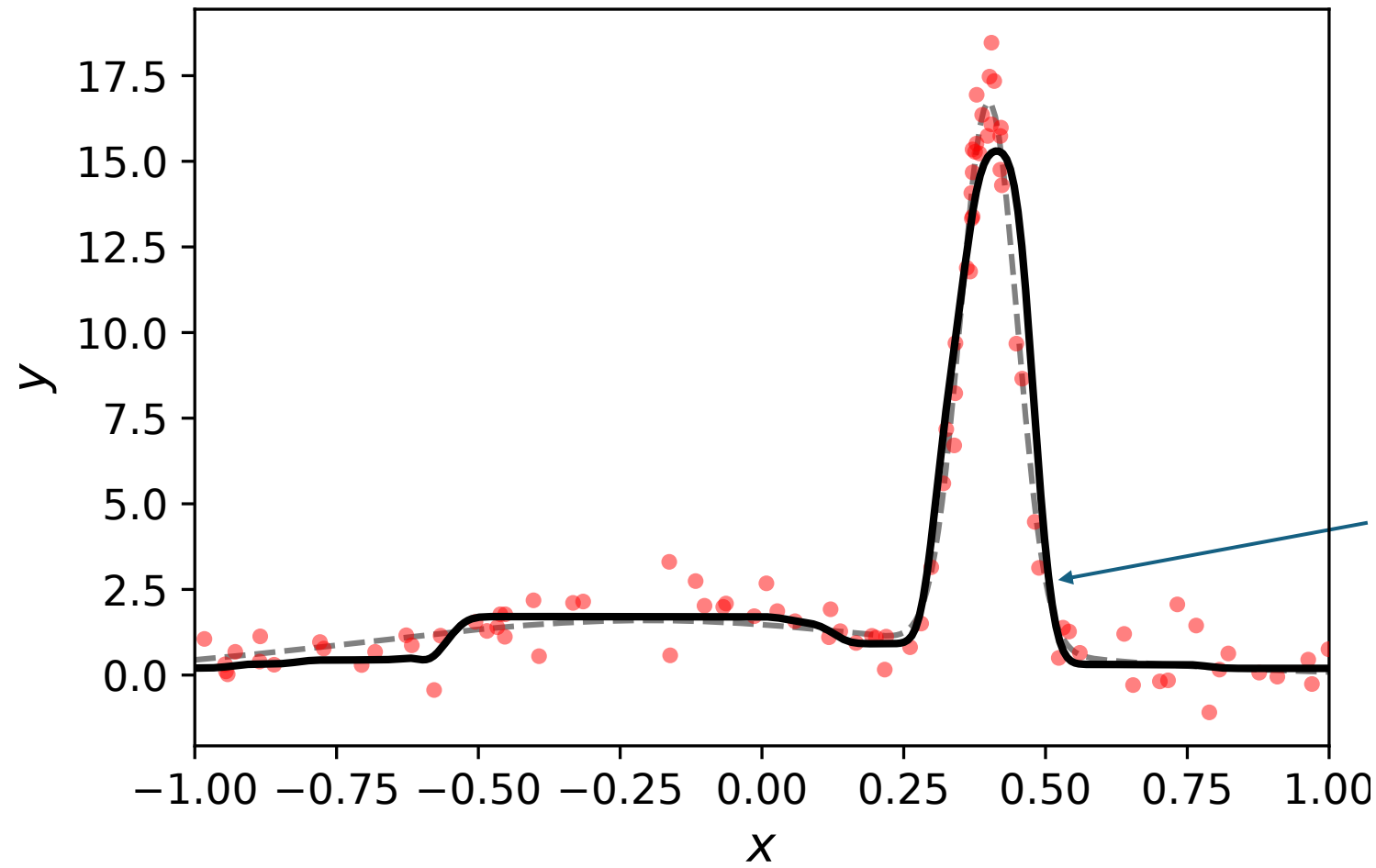
$$f_{\text{RF}}(\mathbf{x}_0) = \mathbf{w}(\mathbf{x}_0)^\top \mathbf{y}$$

$$w_i(\mathbf{x}_0) = \frac{1}{T} \sum_{t=1}^T \frac{m_{t,i} \mathbb{I}(\mathcal{T}_t(\mathbf{x}_i) = \mathcal{T}_t(\mathbf{x}_0))}{\sum_{j=1}^n m_{t,j} \mathbb{I}(\mathcal{T}_t(\mathbf{x}_j) = \mathcal{T}_t(\mathbf{x}_0))}$$

occurrences of $\mathbf{x}_j$ in bootstrap sample t

Goal: *Adaptivity and Smoothness*

6



given:

joint sample $\mathbf{X} \in \mathbb{R}^{n \times p}, \mathbf{y} \in \mathbb{R}^n$
of $\mathbf{x}$ - $\mathbf{y}$ -distribution

task:

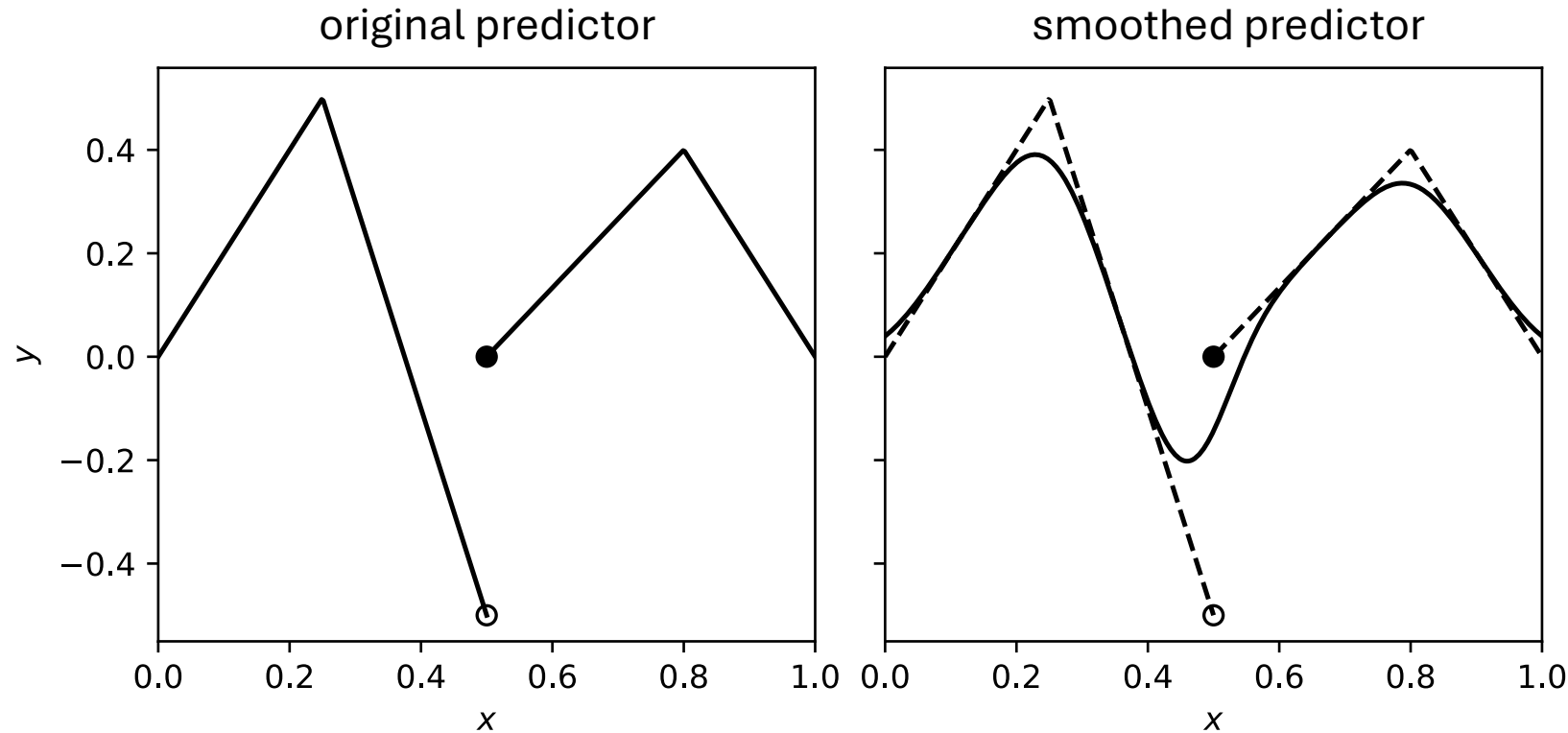
estimate conditional expectation
(regression function)

$$g(\mathbf{x}) = \mathbb{E}(y | \mathbf{x})$$

Smoothed Random forest

Idea: Post-hoc Probabilistic Smoothing

7

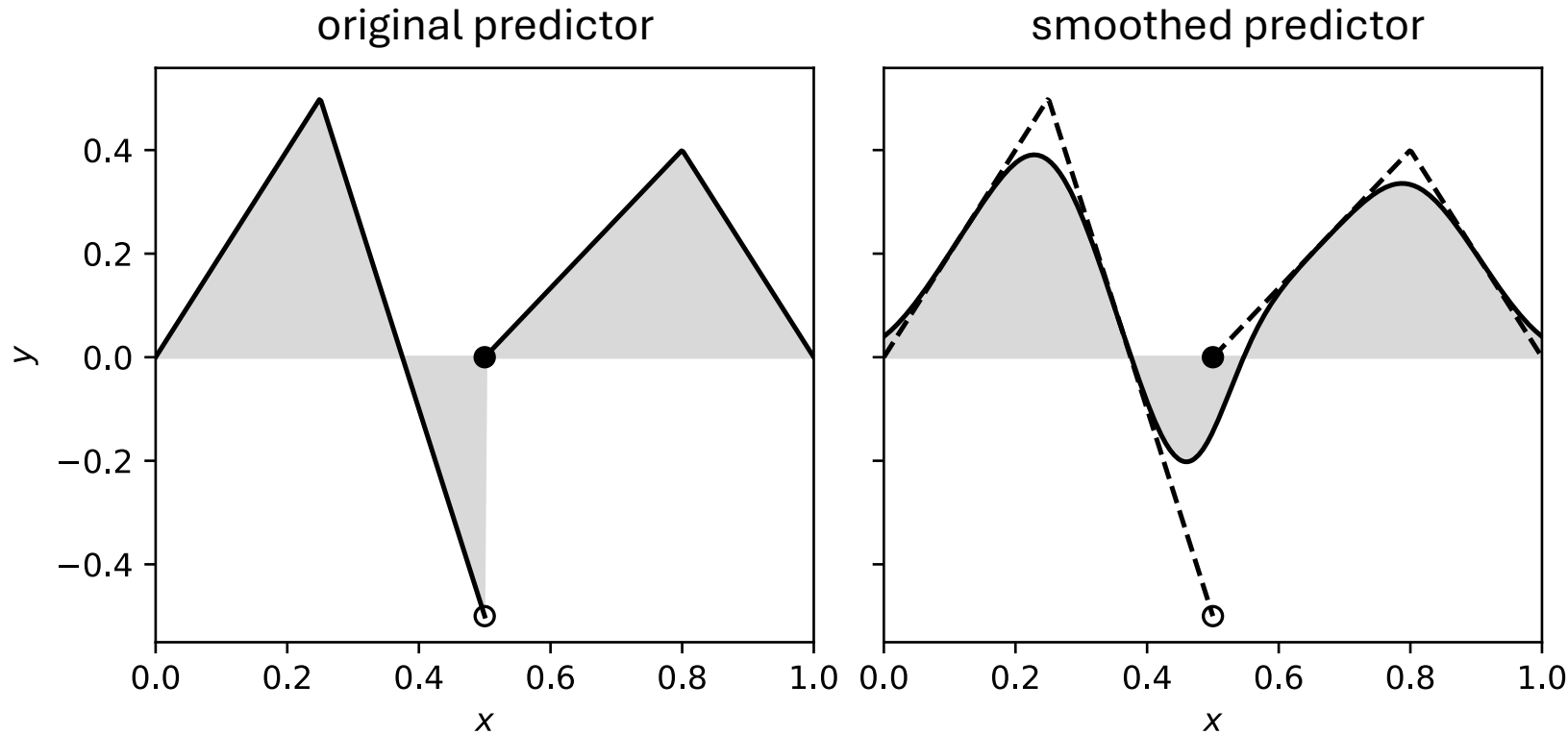


Smoothing transformation: $\tilde{f}(\mathbf{x}_0 \mid \lambda) = \mathbb{E}_{\mathbf{z} \sim k(\cdot \mid \mathbf{x}_0, \lambda)}[f(\mathbf{z})] = \int_{\mathbb{R}^p} f(\mathbf{z}) k(\mathbf{z} \mid \mathbf{x}_0, \lambda) d\mathbf{z}$

test point original predictor latent random test point probability kernel (density) of location/scale family

Idea: Post-hoc Probabilistic Smoothing

8

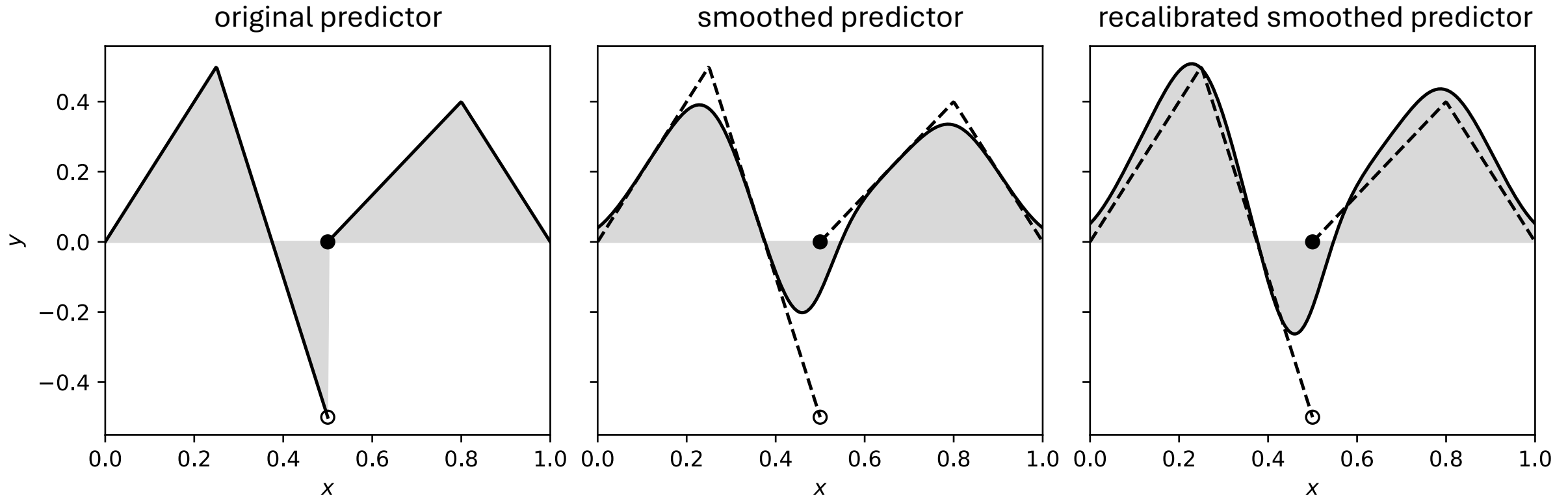


Smoothing transformation: $\tilde{f}(\mathbf{x}_0 \mid \lambda) = \mathbb{E}_{\mathbf{z} \sim k(\cdot \mid \mathbf{x}_0, \lambda)}[f(\mathbf{z})] = \int_{\mathbb{R}^p} f(\mathbf{z}) k(\mathbf{z} \mid \mathbf{x}_0, \lambda) d\mathbf{z}$

Shrinkage in L2 space: $\|\tilde{f}\|_2 \leq \|f\|_2$

Idea: Post-hoc Probabilistic Smoothing

9

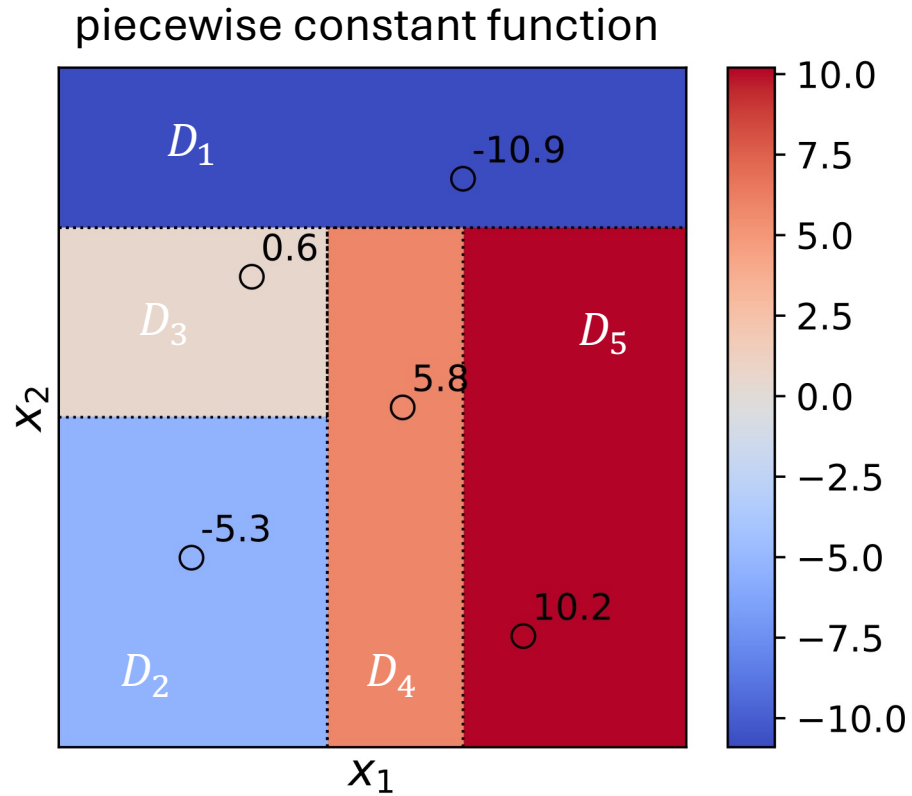


Smoothing transformation: $\tilde{f}(\mathbf{x}_0 \mid \boldsymbol{\lambda}) = \mathbb{E}_{\mathbf{z} \sim k(\cdot \mid \mathbf{x}_0, \boldsymbol{\lambda})}[f(\mathbf{z})] = \int_{\mathbb{R}^p} f(\mathbf{z}) k(\mathbf{z} \mid \mathbf{x}_0, \boldsymbol{\lambda}) d\mathbf{z}$

Recalibrated prediction function: $\hat{y}(\mathbf{x}_0 \mid \boldsymbol{\lambda}, \beta_0, \beta_1) = \beta_1 \tilde{f}(\mathbf{x}_0 \mid \boldsymbol{\lambda}) + \beta_0, \quad \beta_0 \in \mathbb{R}, \beta_1 \in \mathbb{R}_+$

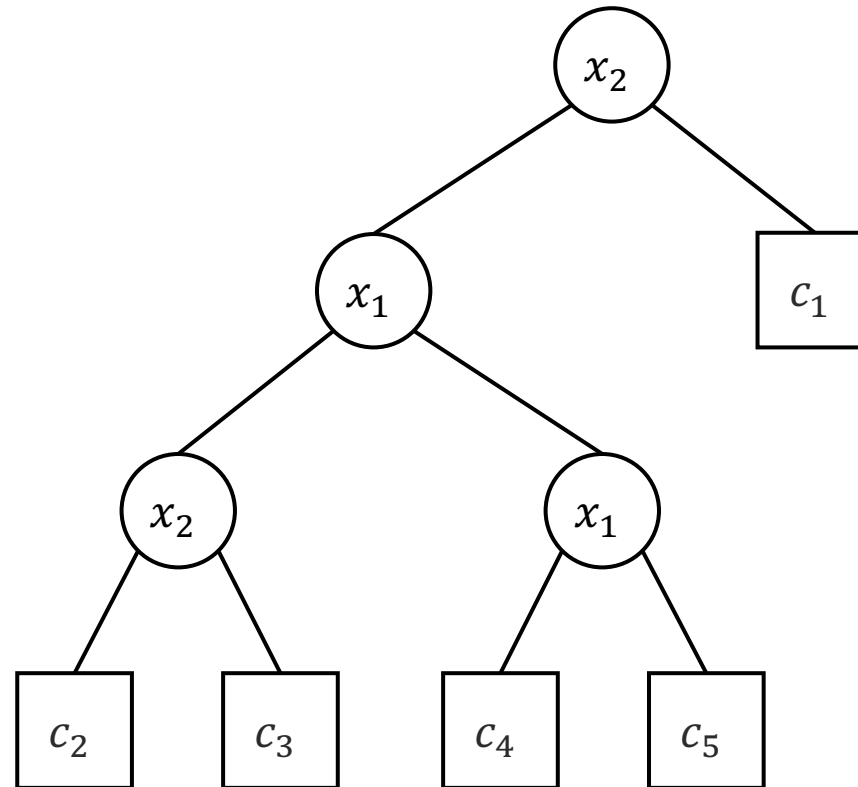
Application to Trees

10



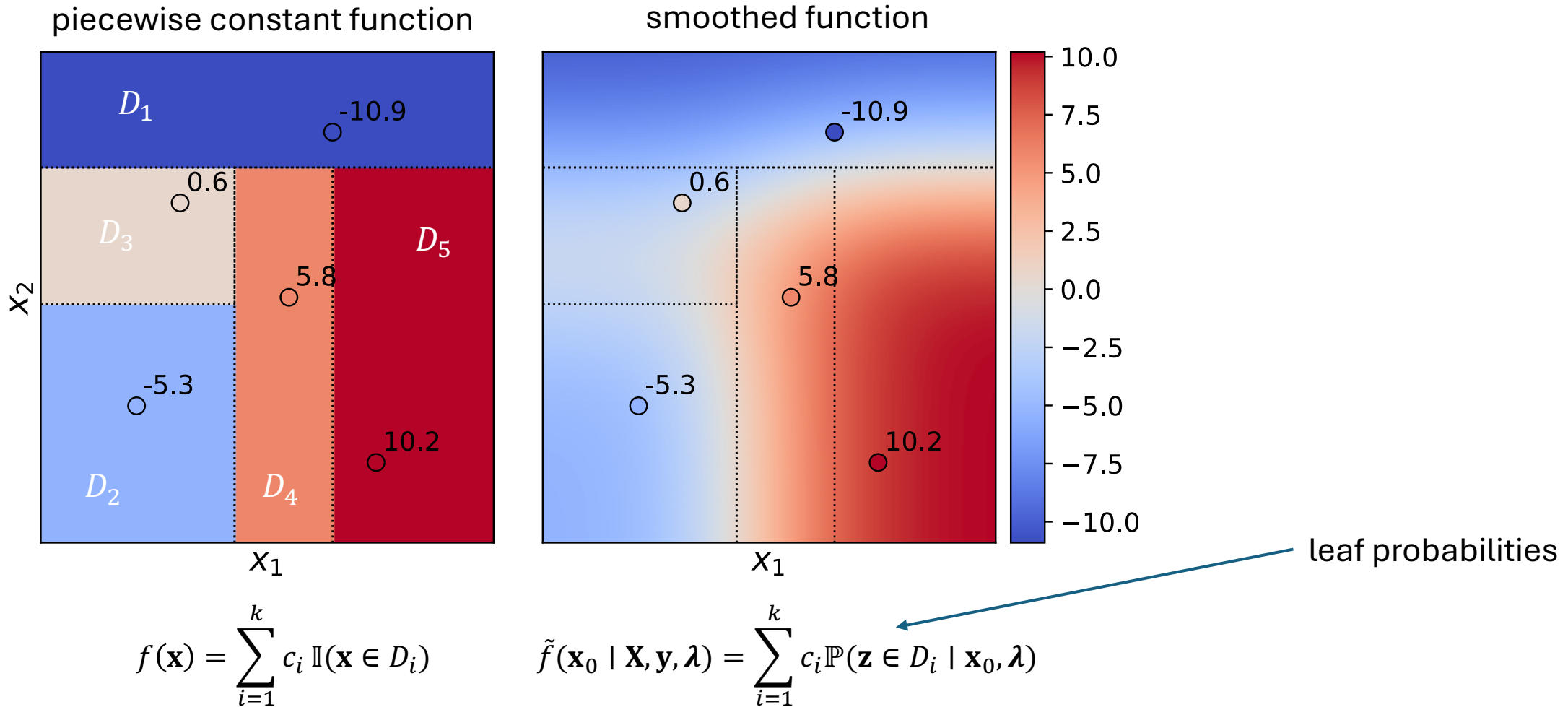
$$f(\mathbf{x}) = \sum_{i=1}^k c_i \mathbb{I}(\mathbf{x} \in D_i)$$

regression tree



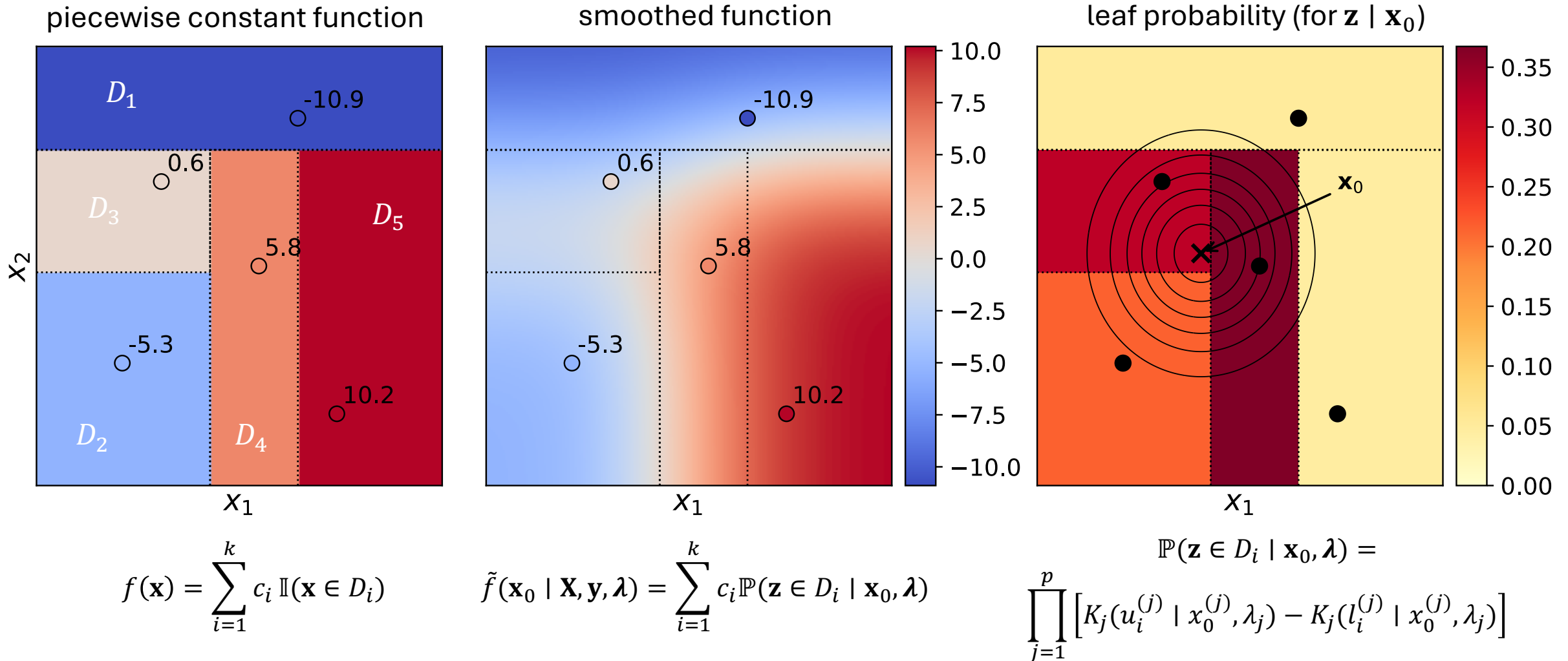
Application to Trees

11

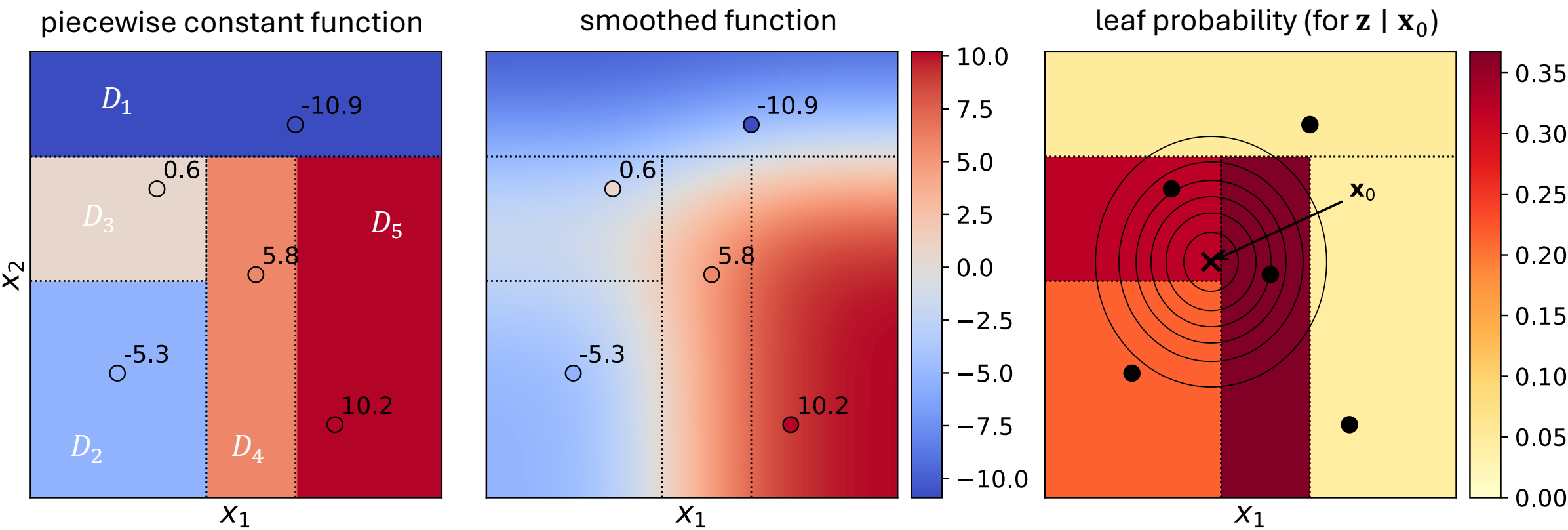


Application to Trees

12



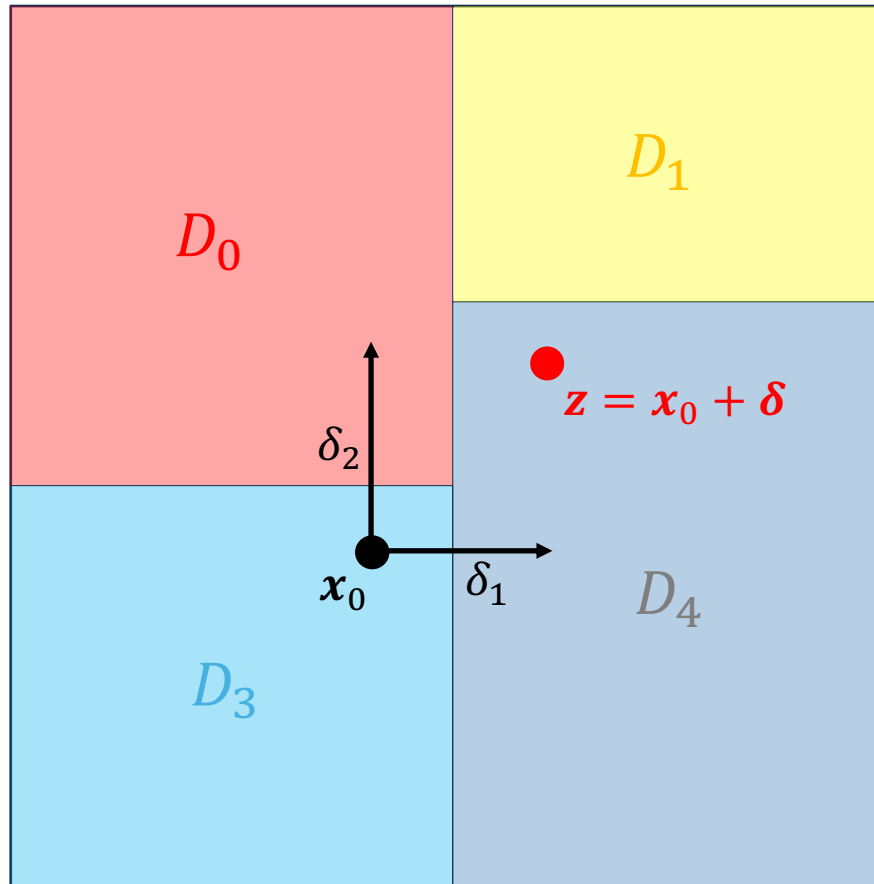
Application to Trees



computational complexity per prediction $O(kp)$

Latent test point approximates tree cut variability

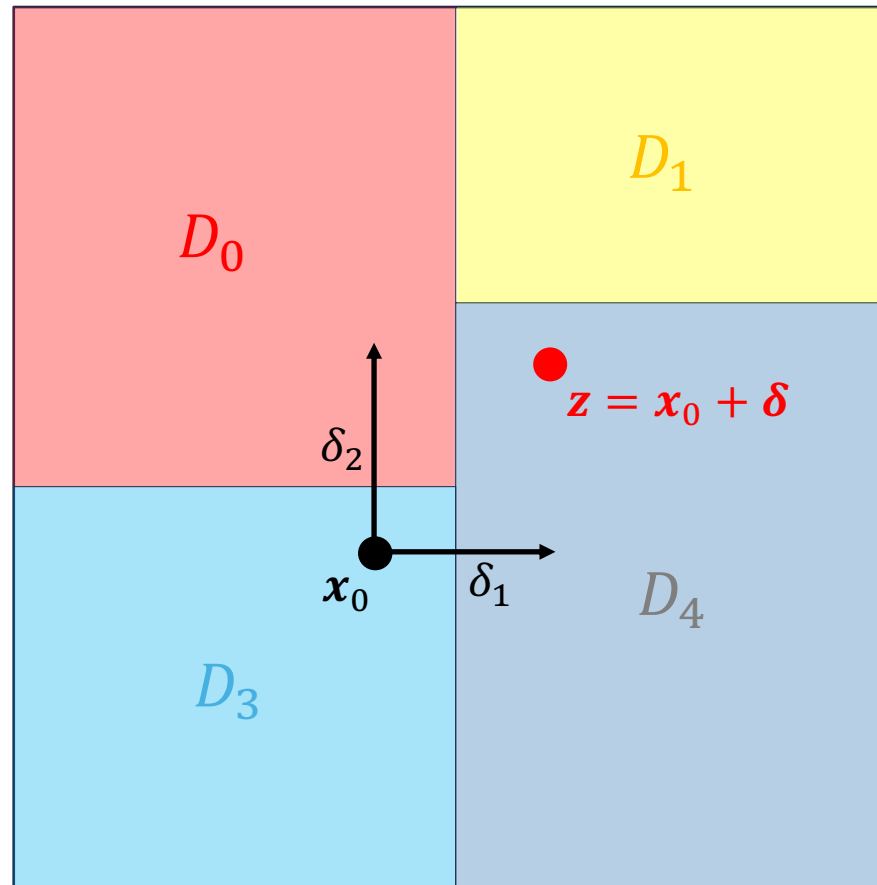
14



original tree

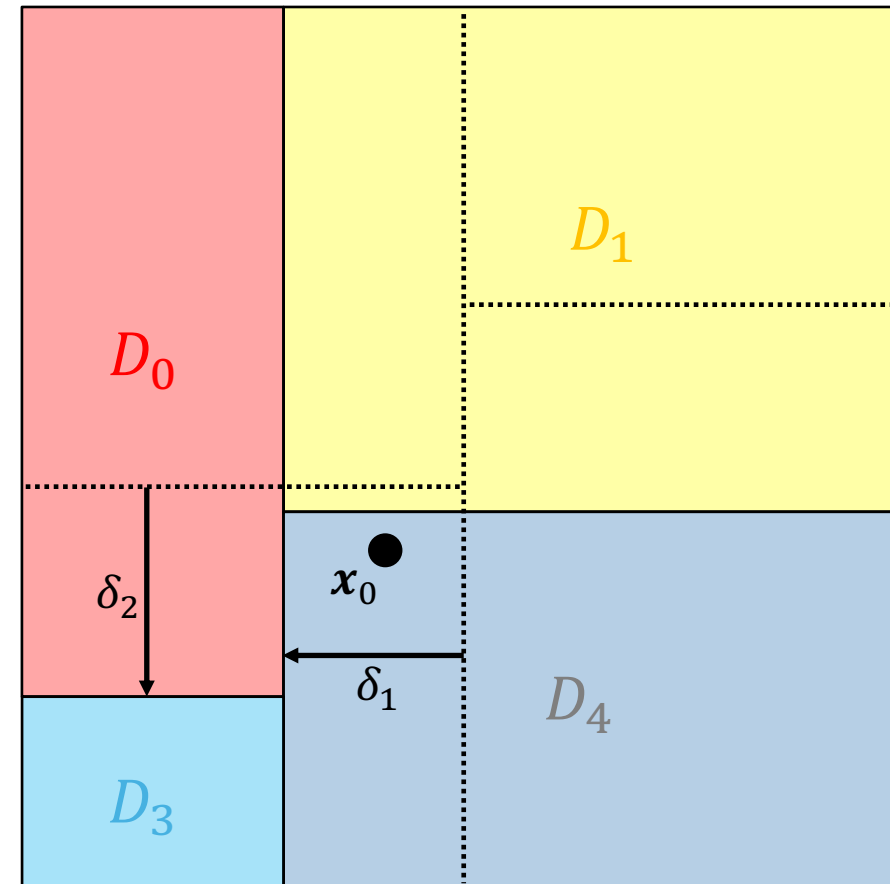
Latent test point approximates tree cut variability

15



original tree

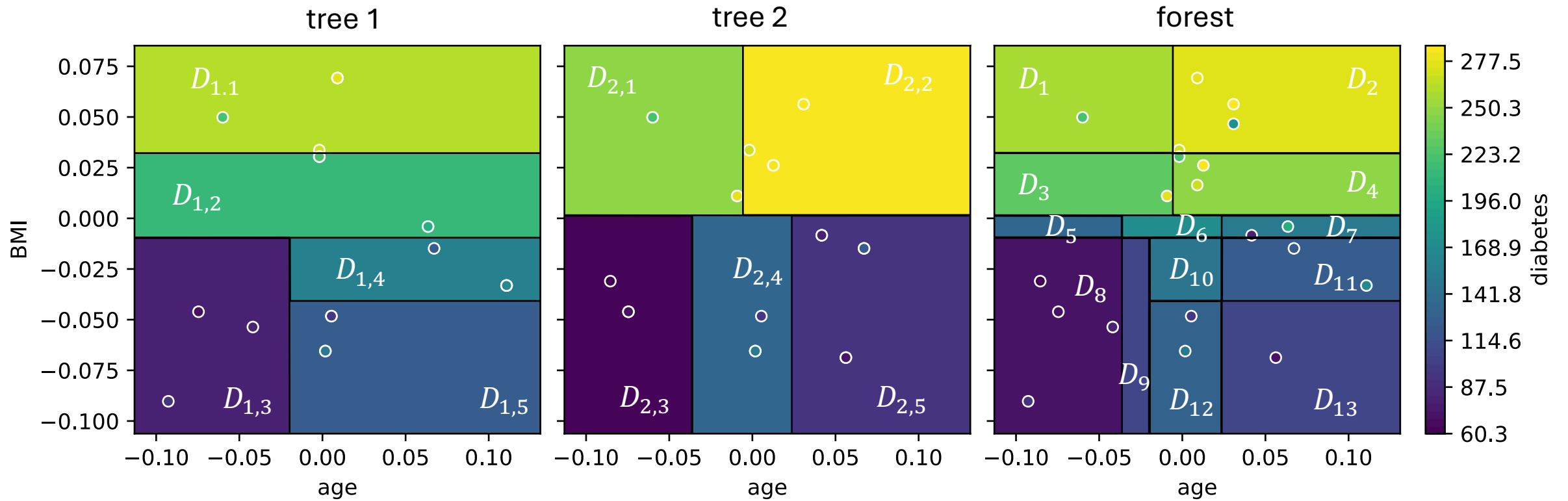
$\approx$



translated region boundaries

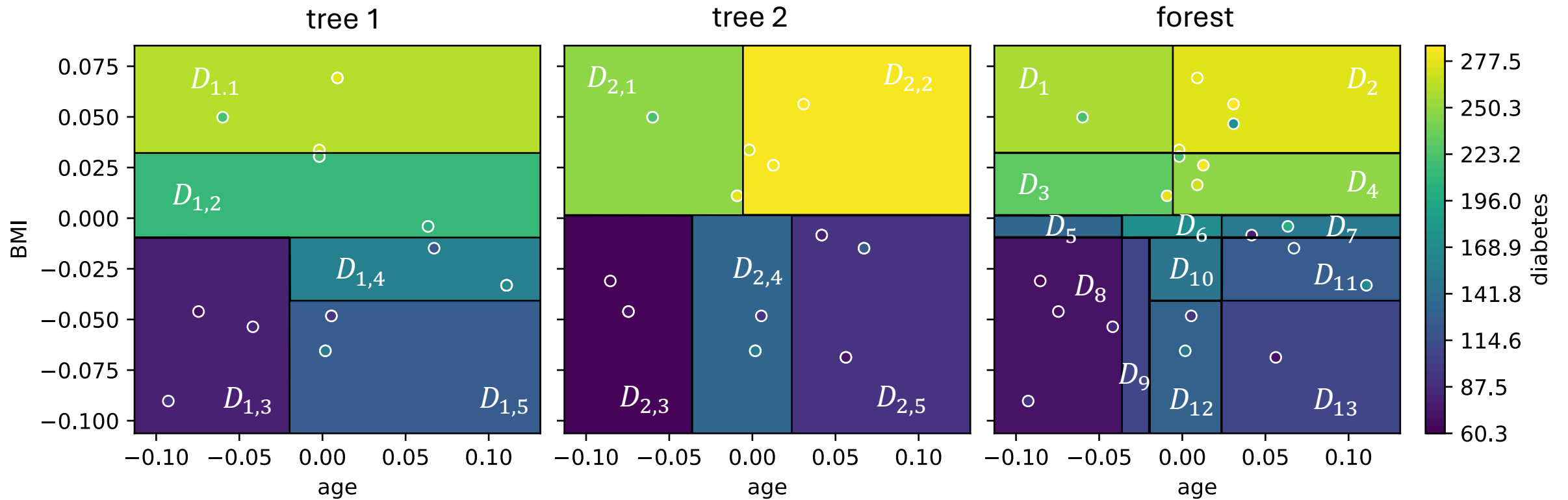
Application to Random Forests

16



Application to Random Forests

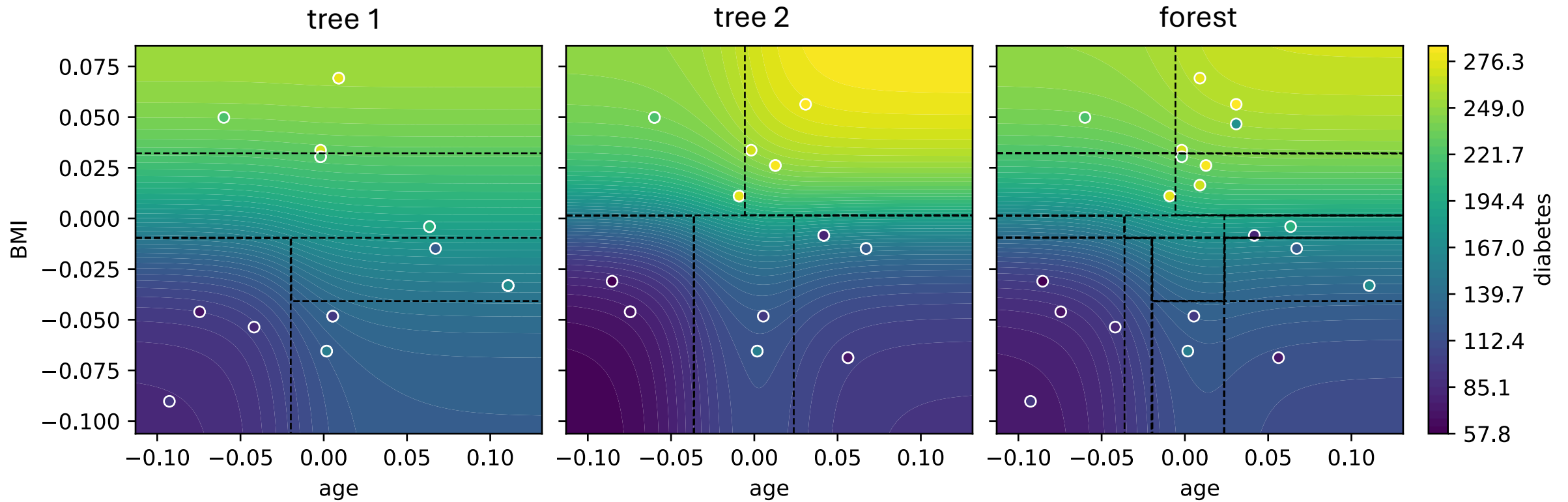
17



Linearity of smoothing transformation: if $f(\mathbf{x}) = \sum_{t=1}^T w_t f_t(\mathbf{x})$ then $\tilde{f}(\mathbf{x} | \lambda) = \sum_{t=1}^T w_t \tilde{f}_t(\mathbf{x} | \lambda)$

Application to Random Forests

18



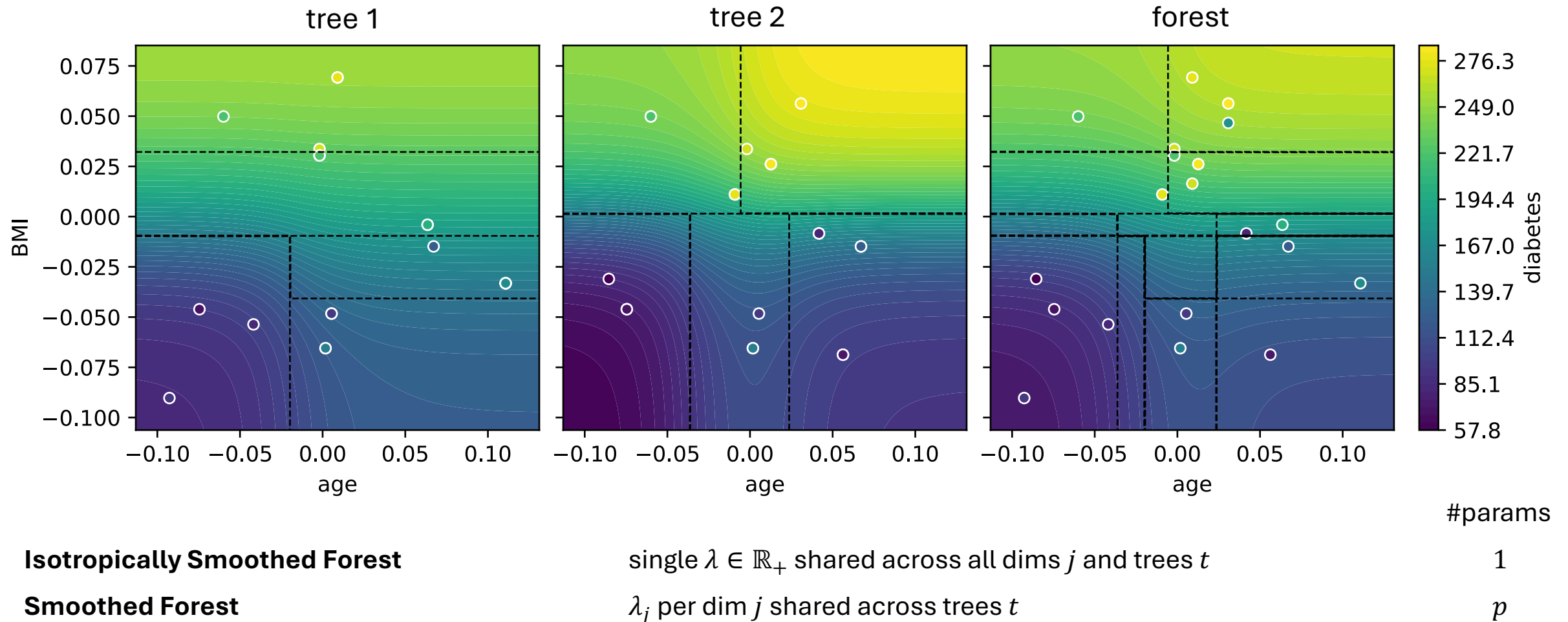
Linearity of smoothing transformation: if $f(\mathbf{x}) = \sum_{t=1}^T w_t f_t(\mathbf{x})$ then $\tilde{f}(\mathbf{x} | \boldsymbol{\lambda}) = \sum_{t=1}^T w_t \tilde{f}_t(\mathbf{x} | \boldsymbol{\lambda})$

Allows prediction complexity: $O(lp) \ll O(kp) = O(2^l p)$ in worst case where $l = l_1 + \dots + l_T$

(same as GPR, factor of p more than RF worst case complexity)

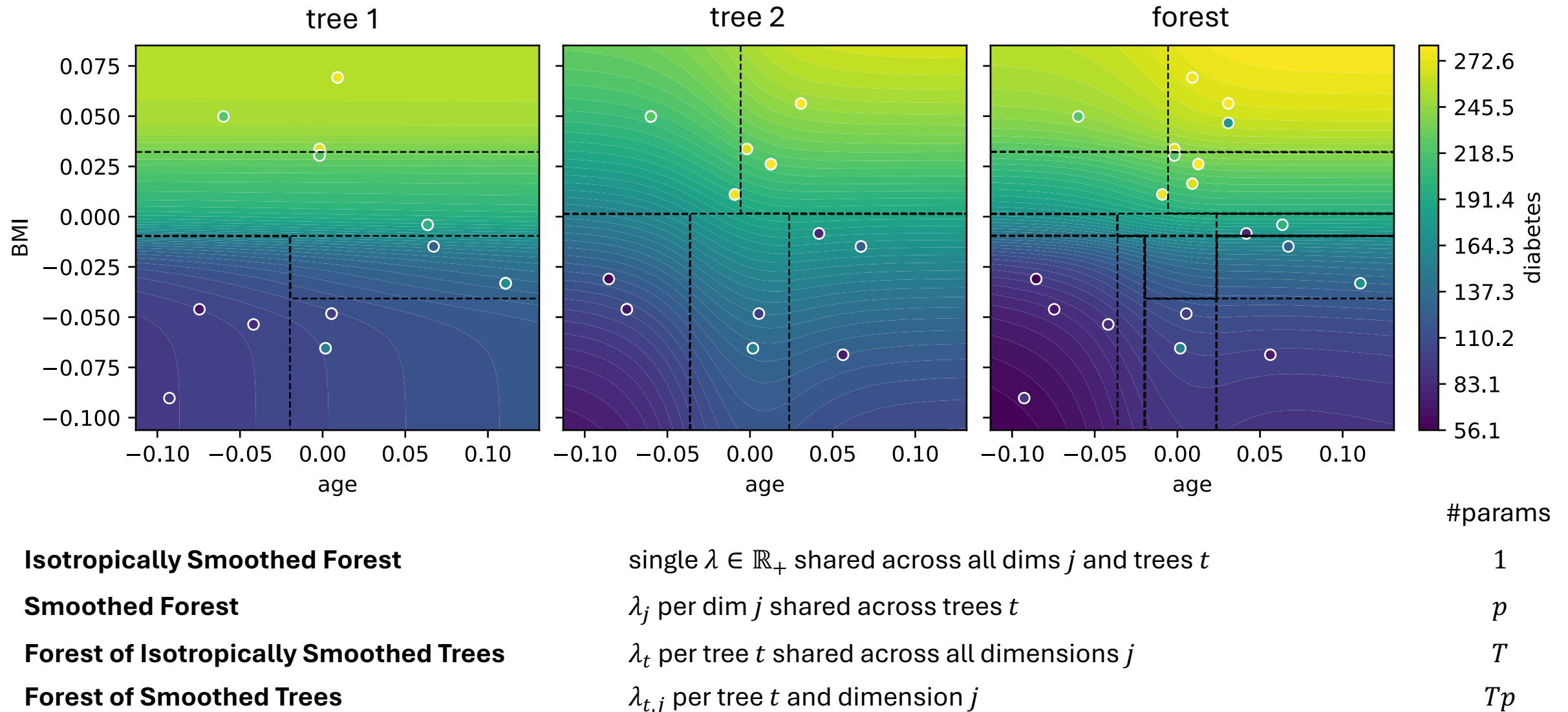
Different Model Variants Possible

19



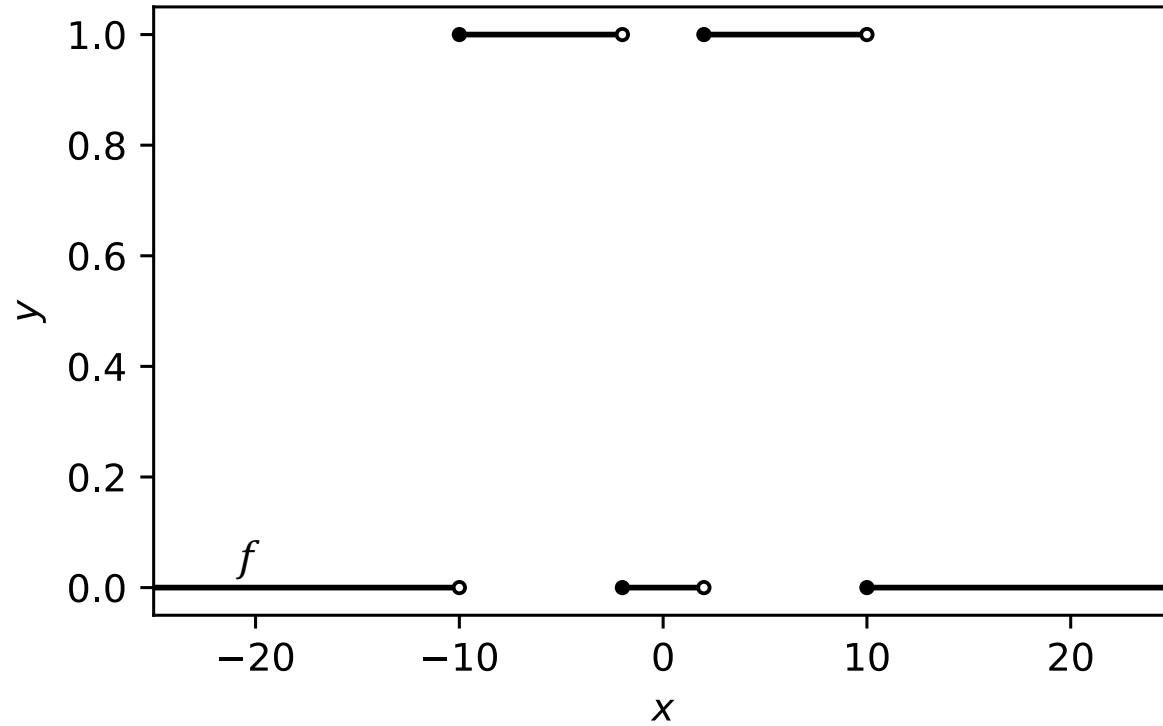
Different Model Variants Possible

20



Model Selection via Out-of-bag Loss

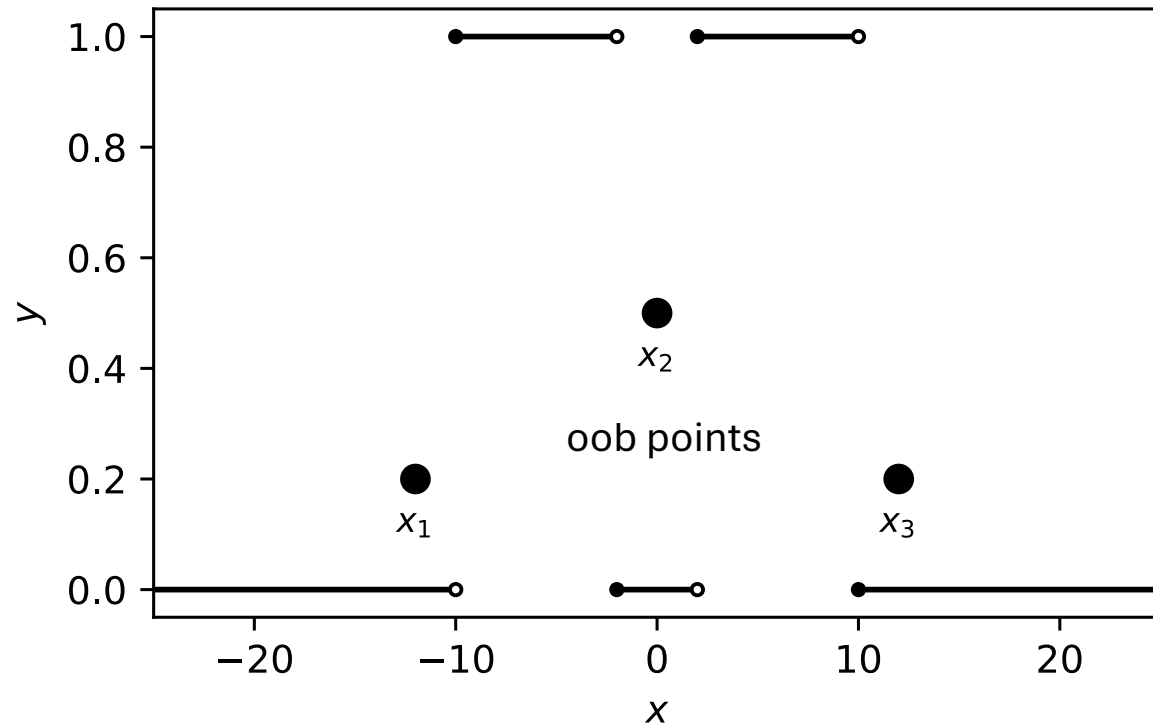
21



Training loss: under-smoothing

Model Selection via Out-of-bag Loss

22

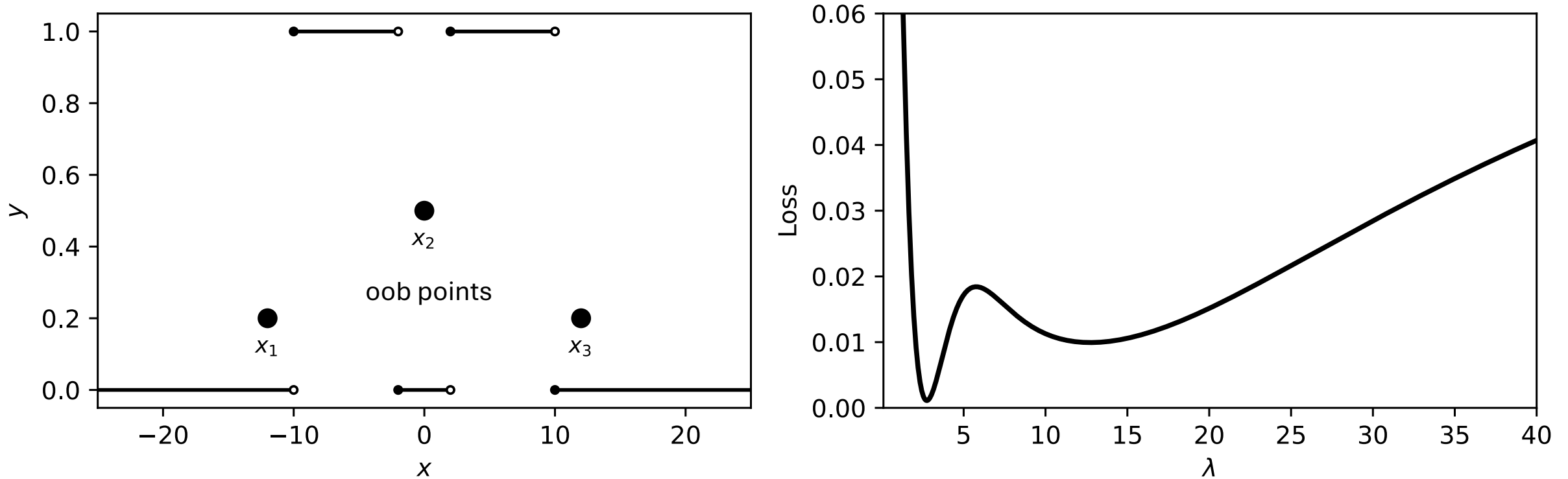


Training loss: under-smoothing

Instead use out-of-bag loss: $L_{\text{oob}}(\Lambda, \beta_0, \beta_1) = \sum_{t, i \in \mathcal{O}} (y_i - \beta_1 \tilde{f}(\mathbf{x}_i \mid \mathcal{T}_t, \lambda_t) - \beta_0)^2$

Model Selection via Out-of-bag Loss

23

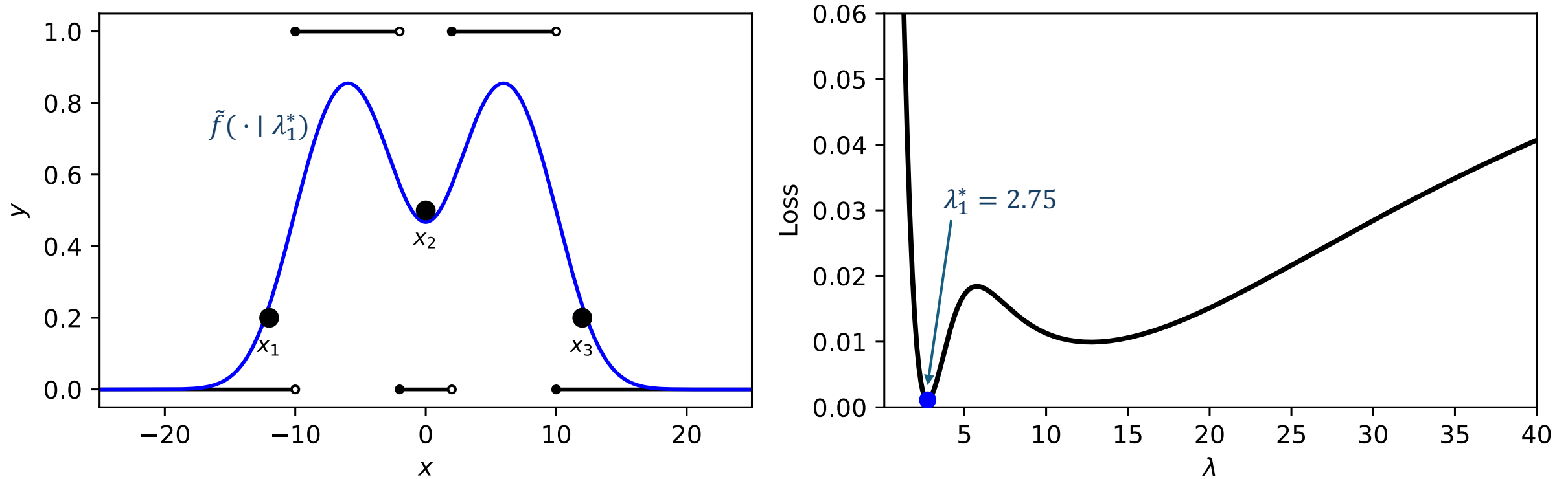


Training loss: under-smoothing

Instead use out-of-bag loss: $L_{\text{oob}}(\Lambda, \beta_0, \beta_1) = \sum_{t, i \in \mathcal{O}} (y_i - \beta_1 \tilde{f}(\mathbf{x}_i | \mathcal{T}_t, \lambda_t) - \beta_0)^2$

Model Selection via Out-of-bag Loss

24

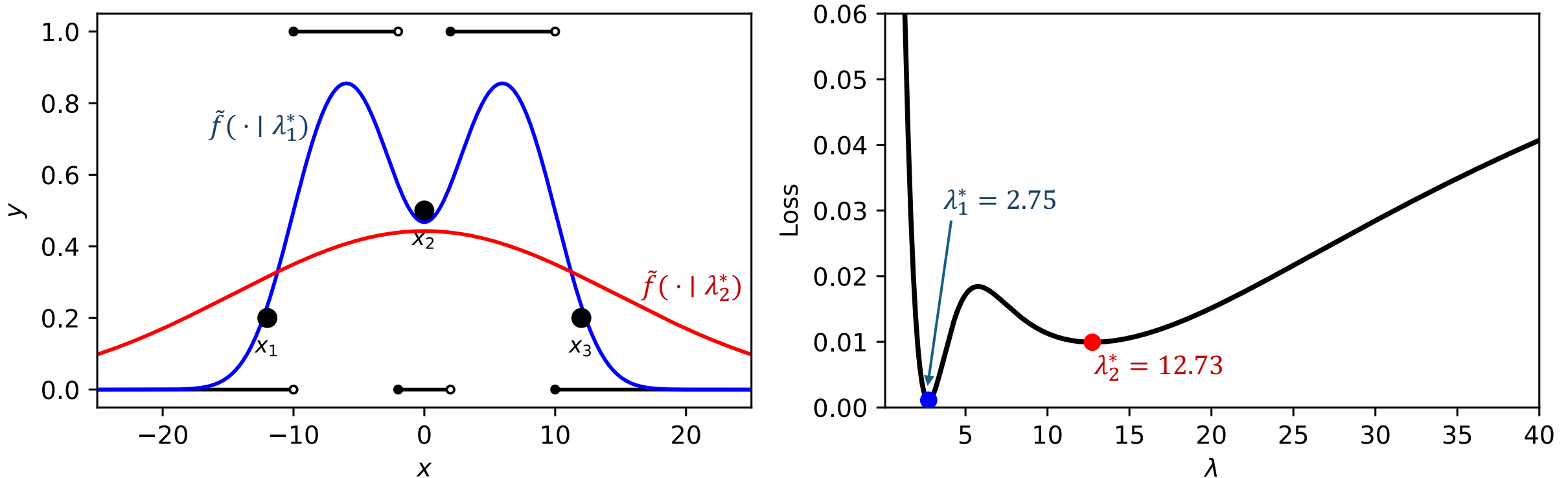


Training loss: under-smoothing

Instead use out-of-bag loss: $L_{\text{oob}}(\Lambda, \beta_0, \beta_1) = \sum_{t, i \in \mathcal{O}} (y_i - \beta_1 \tilde{f}(\mathbf{x}_i | \mathcal{T}_t, \lambda_t) - \beta_0)^2$

Model Selection via Out-of-bag Loss

25



Training loss: under-smoothing

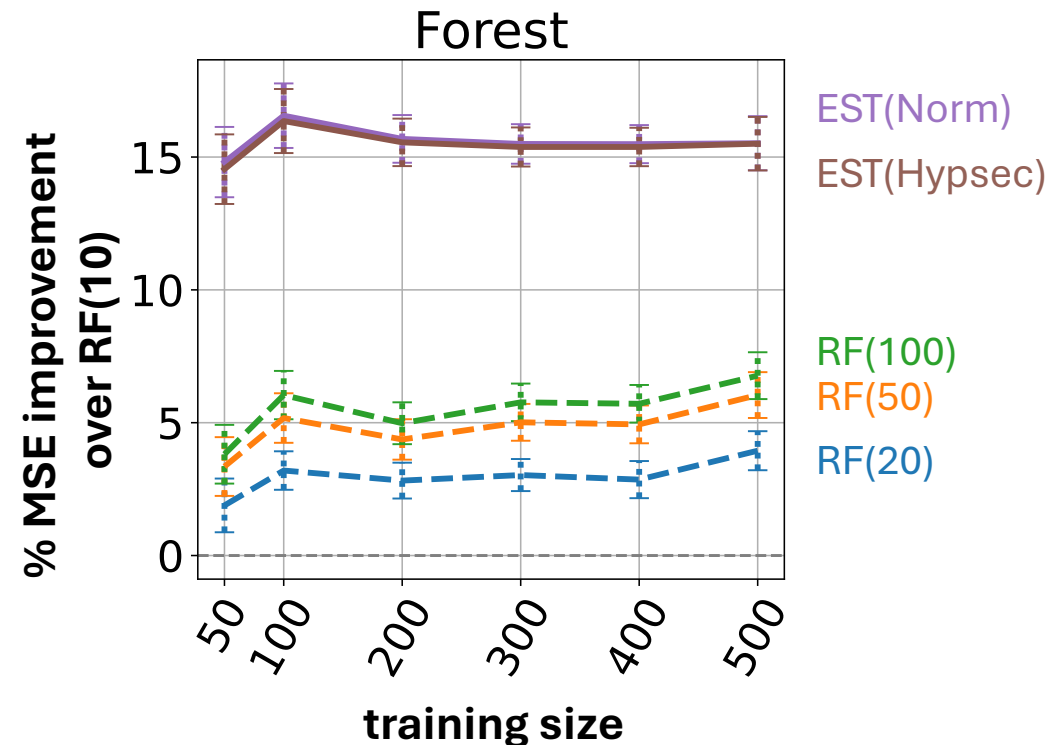
Instead use out-of-bag loss: $L_{\text{oob}}(\Lambda, \beta_0, \beta_1) = \sum_{t, i \in \mathcal{O}} (y_i - \beta_1 \tilde{f}(\mathbf{x}_i | \mathcal{T}_t, \lambda_t) - \beta_0)^2$

Solve via gradient method: several local minima, but differentiable

Effect of Smoothing in Practice

26

$$\text{PI}_{\text{MSE}}(\cdot) = \frac{\text{MSE}(\text{RF}(10)) - \text{MSE}(\cdot)}{\text{MSE}(\text{RF}(10))} \times 100\%$$

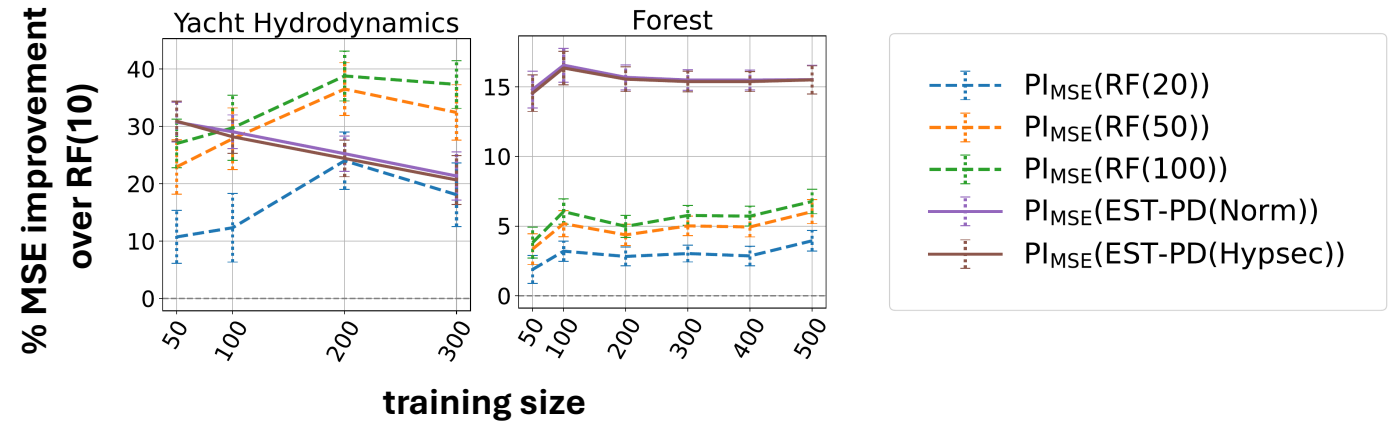


$$k_j(\mathbf{z}; \mathbf{x}_0, \lambda) = \frac{1}{\sqrt{2\pi\lambda^2}} \exp\left(-\frac{(\mathbf{z} - \mathbf{x}_0)^2}{2\lambda}\right)$$

$$k_j(\mathbf{z}; \mathbf{x}_0, \lambda) = \frac{1}{2\lambda} \operatorname{sech}\left(\frac{\pi(\mathbf{z} - \mathbf{x}_0)}{2\lambda}\right)$$

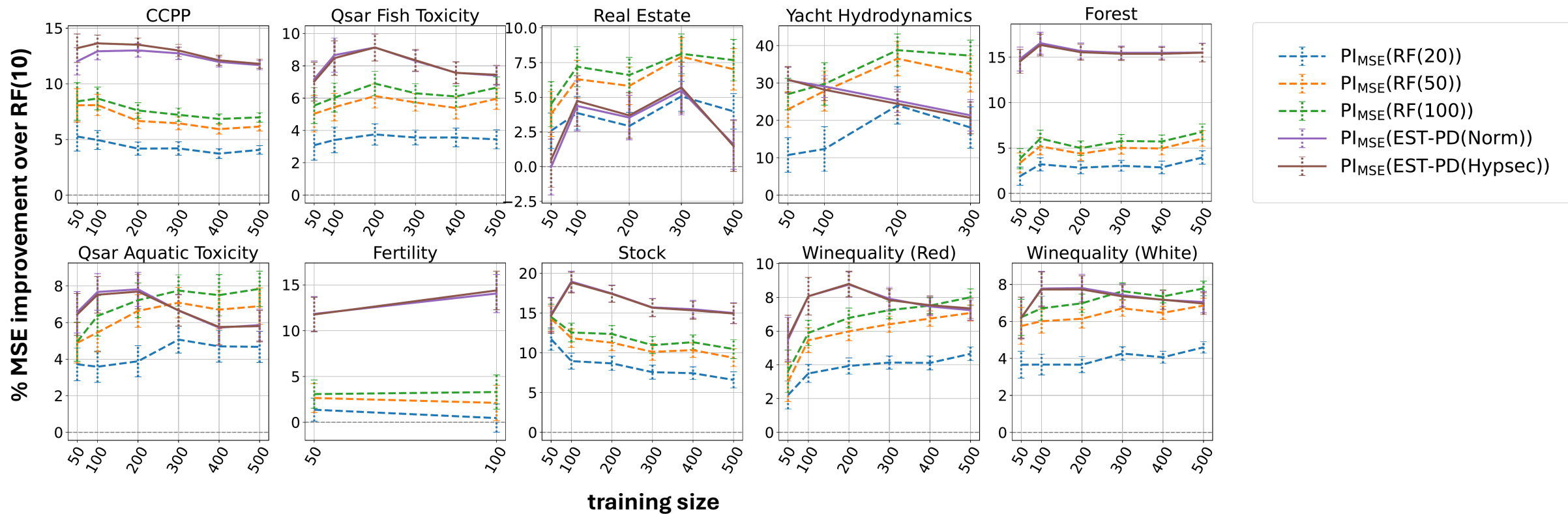
Effect of Smoothing in Practice

27



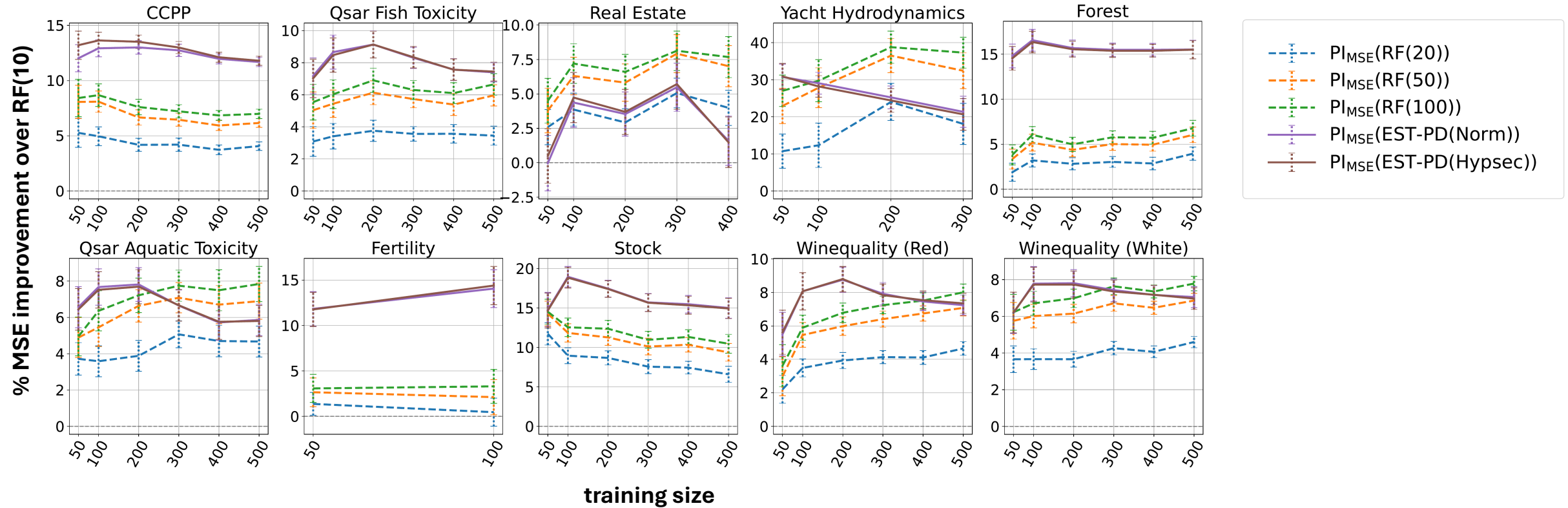
Effect of Smoothing in Practice

28

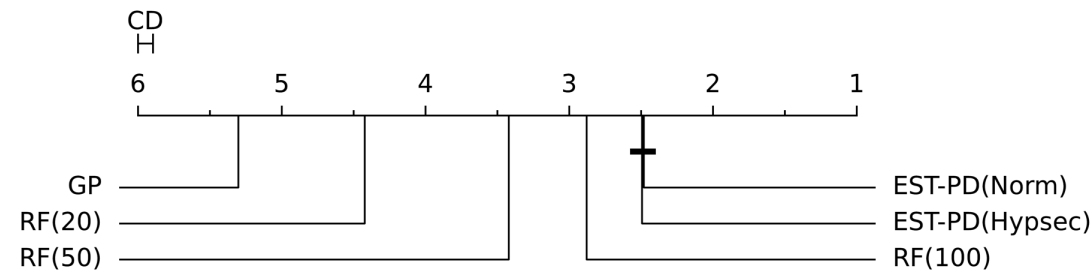


Effect of Smoothing in Practice

29



**Nemenyi post-hoc
critical-difference analysis**

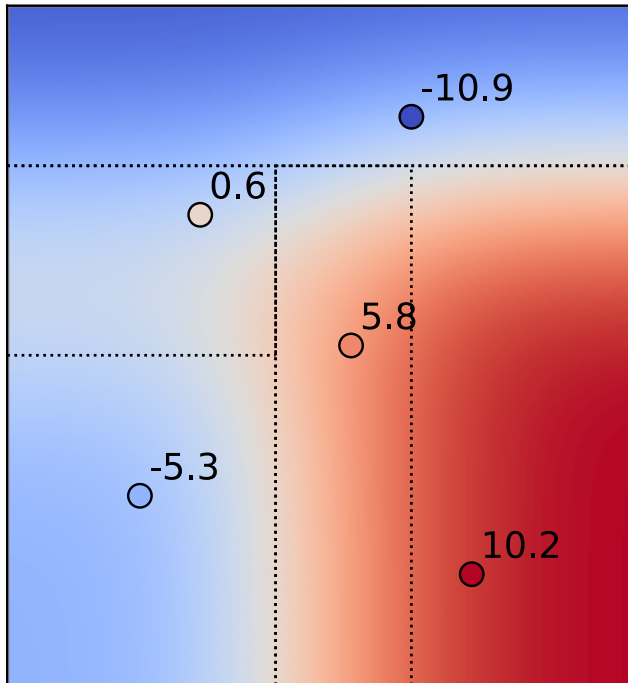


Summary

30

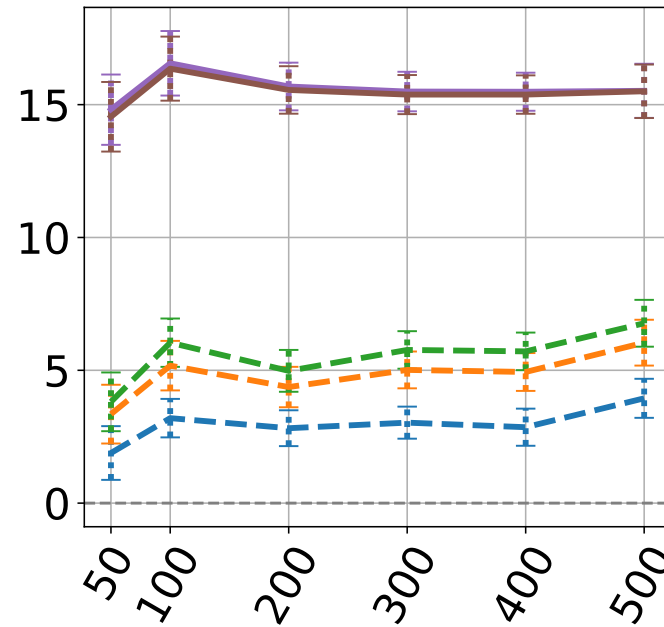
Idea:

achieve spatial adaptivity and smoothness by post-hoc smoothing of random forest



Finding:

can lead to notable MSE improvements in small data regime



Future:

- compare to *soft tree forests* and *RF as mean function* of GP
- integrate calibration with smoothing selection
- theoretical characterisation



github.com/Neal-Liu-Ziyi/SmoothedRandomForest

Try it out yourself!