

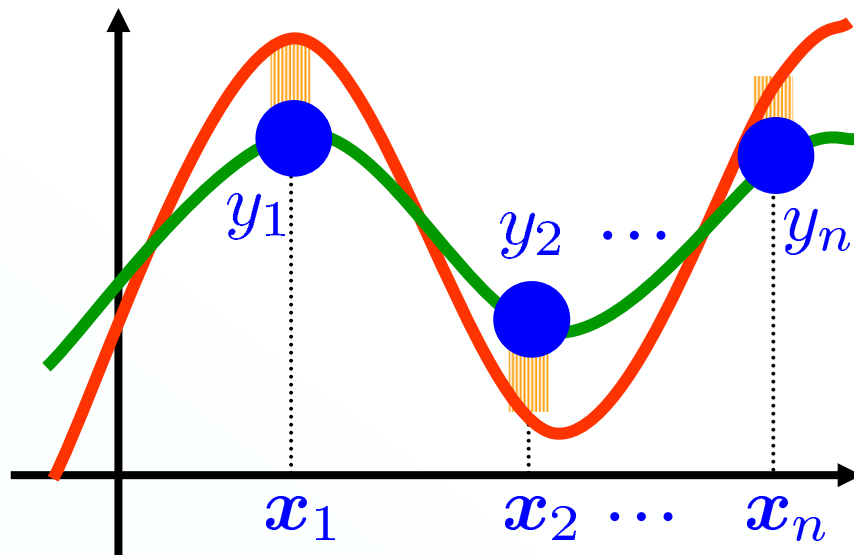
Estimating the Error at Given Test Input Points for Linear Regression



Masashi Sugiyama

Fraunhofer FIRST-IDA, Berlin, Germany
Tokyo Institute of Technology, Tokyo, Japan

Regression Problem



$f(\mathbf{x})$: Underlying function

$\hat{f}(\mathbf{x})$: Learned function

$\{(\mathbf{x}_i, y_i)\}_{i=1}^n$: Training examples

$$y_i = f(\mathbf{x}_i) + \epsilon_i$$

(noise)

$\epsilon_i \stackrel{i.i.d.}{\sim}$ mean 0, variance σ^2

From $\{(\mathbf{x}_i, y_i)\}_{i=1}^n$, obtain a good approximation $\hat{f}(\mathbf{x})$ to $f(\mathbf{x})$

Typical Method of Learning

3

■ Linear regression model

$$\hat{f}(\mathbf{x}) = \sum_{i=1}^p \alpha_i \varphi_i(\mathbf{x})$$

α_i :Parameters

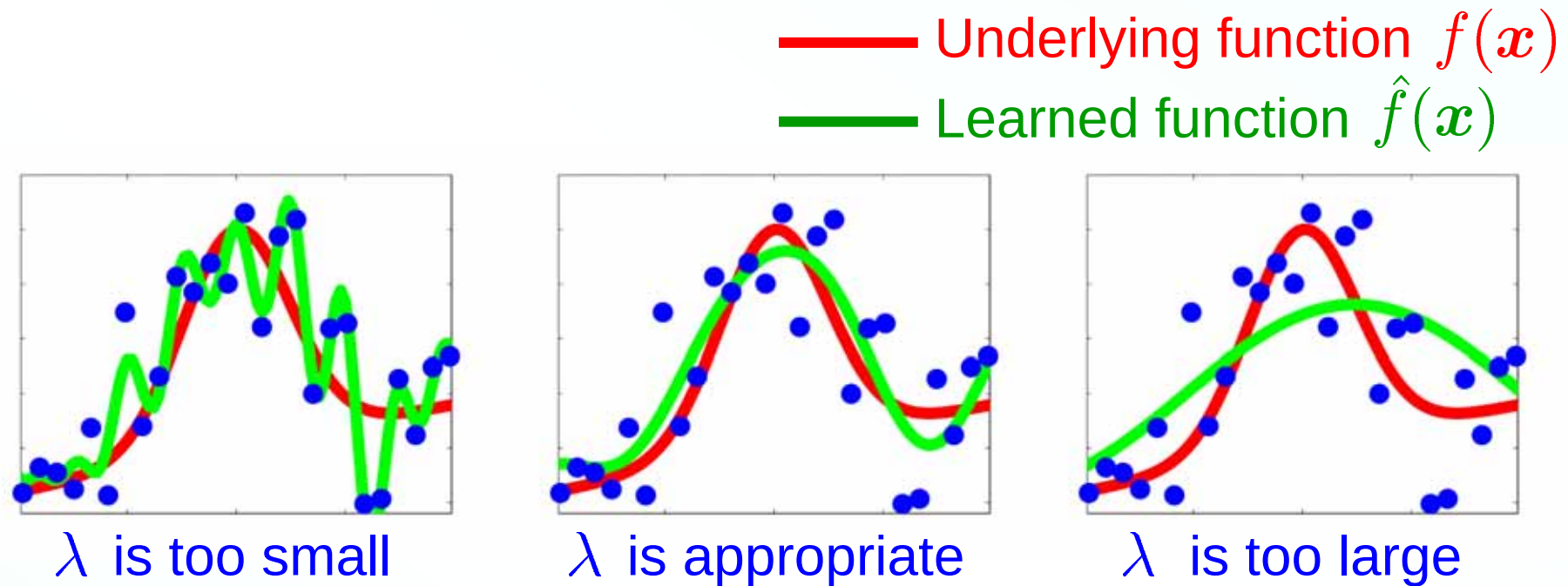
$\{\varphi_i(\mathbf{x})\}_{i=1}^p$:Fixed basis functions

■ Ridge estimation

$$\min_{\{\alpha_i\}} \left[\sum_{i=1}^n \left(\hat{f}(\mathbf{x}_i) - y_i \right)^2 + \lambda \sum_{i=1}^p \alpha_i^2 \right]$$

λ :Ridge parameter (model parameter)

Model Selection



Choice of the model is crucial
for obtaining good learned function $\hat{f}(x)$!

Generalization Error

For model selection, we need a criterion that measures 'closeness' between $\hat{f}(x)$ and $f(x)$:

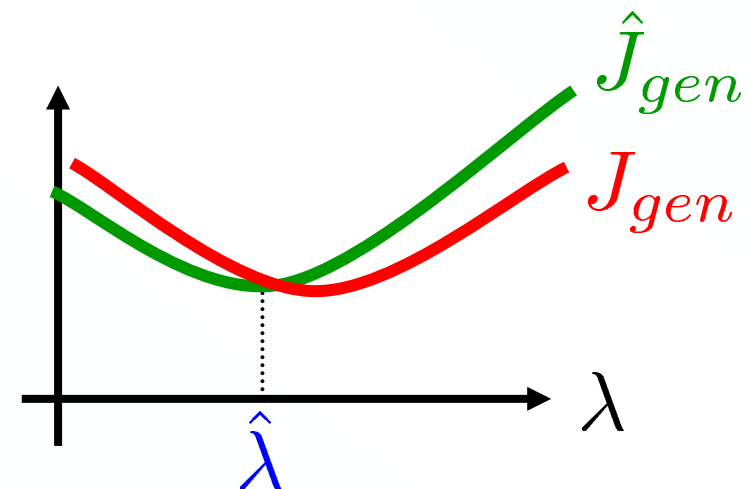
➔ Generalization error, e.g.,

$$J_{gen}(\lambda) = \int \left(\hat{f}(x) - f(x) \right)^2 p(x) dx$$

Determine the model λ so that an estimator $\hat{J}_{gen}$ of the unknown generalization error J_{gen} is minimized.

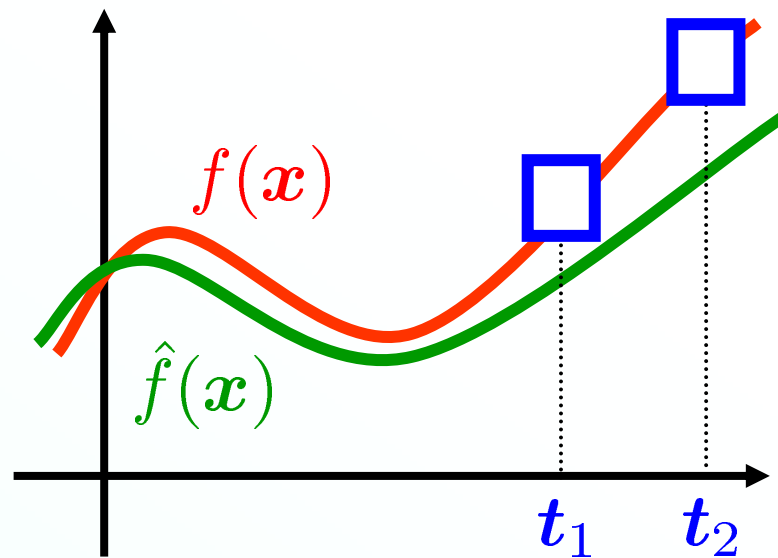
$$\hat{\lambda} = \underset{\lambda}{\operatorname{argmin}} \hat{J}_{gen}(\lambda)$$

$p(x)$: Probability density of test input points



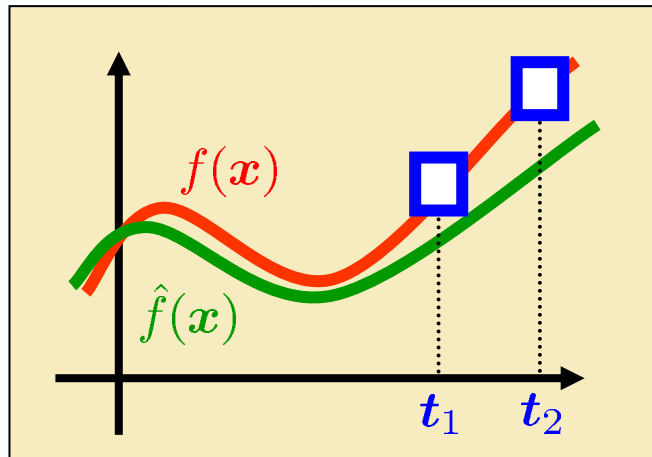
Transductive Inference

- Test input points are specified in advance.
- We do not have to estimate the entire function $f(x)$, but just estimate the values of the function at the test input points $\{t_i\}_{i=1}^{n_t}$.

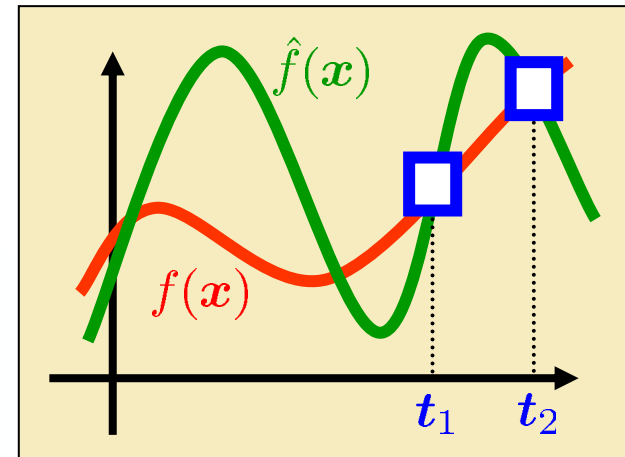


Model Selection for Transductive Inference

- Test error at given test input points is different from the generalization error.
- Model should be chosen so that the test error only at $\{t_i\}_{i=1}^{n_t}$ is minimized.



Small generalization error
Large test error



Large generalization error
Small test error

Goal of Our Research

- We want to estimate the test error at the given test input points!

$$J_{test} = \mathbb{E} \sum_{i=1}^{n_t} \left(\hat{f}(\mathbf{t}_i) - f(\mathbf{t}_i) \right)^2$$

$\mathbb{E}$: Expectation over noise



Setting

9

■ Linear regression model

$$\hat{f}(\mathbf{x}) = \sum_{i=1}^p \alpha_i \varphi_i(\mathbf{x})$$

α_i :Parameters
 $\{\varphi_i(\mathbf{x})\}_{i=1}^p$:Fixed basis functions

■ Linear estimation

$$\hat{\alpha} = Xy$$

$$\hat{\alpha} = (\hat{\alpha}_1, \hat{\alpha}_2, \dots, \hat{\alpha}_p)^\top$$

X :A matrix

$$y = (y_1, y_2, \dots, y_n)^\top$$

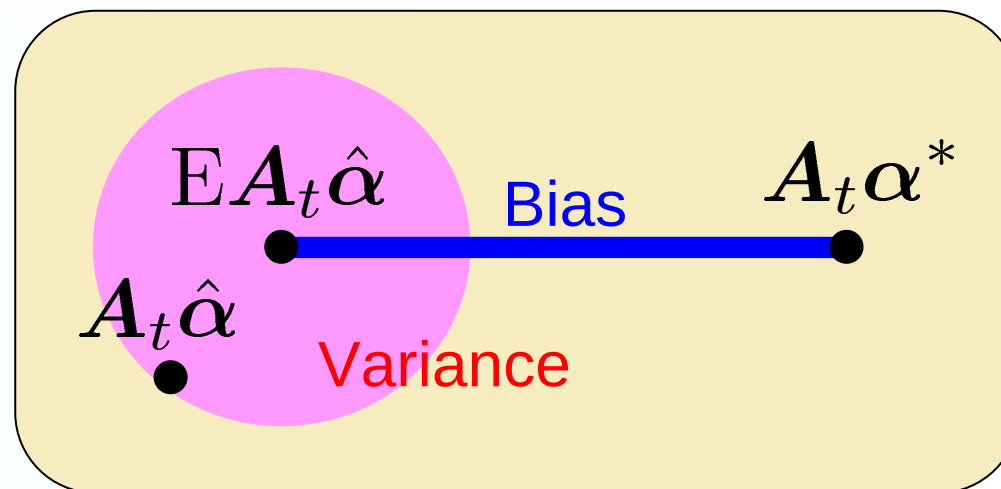
■ Realizability

$$f(\mathbf{x}) = \sum_{i=1}^p \alpha_i^* \varphi_i(\mathbf{x})$$

α_i^* :Unknown true parameters

Bias / Variance Decomposition ¹⁰

$$\begin{aligned} J_{test} &= \mathbb{E} \sum_{i=1}^{n_t} \left(\hat{f}(\mathbf{t}_i) - f(\mathbf{t}_i) \right)^2 & [\mathbf{A}_t]_{i,j} &= \varphi_j(\mathbf{t}_i) \\ &= \mathbb{E} \|\mathbf{A}_t \hat{\boldsymbol{\alpha}} - \mathbf{A}_t \boldsymbol{\alpha}^*\|^2 & \mathbf{A}_t \hat{\boldsymbol{\alpha}} &= (\hat{f}(\mathbf{t}_1), \hat{f}(\mathbf{t}_2), \dots, \hat{f}(\mathbf{t}_{n_t}))^\top \\ &= \underbrace{\|\mathbb{E} \mathbf{A}_t \hat{\boldsymbol{\alpha}} - \mathbf{A}_t \boldsymbol{\alpha}^*\|^2}_{\text{Bias}} + \underbrace{\mathbb{E} \|\mathbf{A}_t \hat{\boldsymbol{\alpha}} - \mathbb{E} \mathbf{A}_t \hat{\boldsymbol{\alpha}}\|^2}_{\text{Variance}} & \boldsymbol{\alpha}^* &= (\alpha_1^*, \alpha_2^*, \dots, \alpha_p^*)^\top \end{aligned}$$





Tricks for Estimating Bias

11

Sugiyama & Ogawa (Neural Comp., 2001)

Sugiyama & Müller (JMLR, 2002)

$$Bias = \|E\mathbf{A}_t\hat{\boldsymbol{\alpha}} - \mathbf{A}_t\boldsymbol{\alpha}^*\|^2$$

- True parameter $\boldsymbol{\alpha}^*$ is unknown.
- We utilize an unbiased estimator of the true parameter for estimating the bias.

$$\hat{\boldsymbol{\alpha}}_u = \mathbf{A}^\dagger \mathbf{y}$$

$$E\hat{\boldsymbol{\alpha}}_u = \boldsymbol{\alpha}^*$$

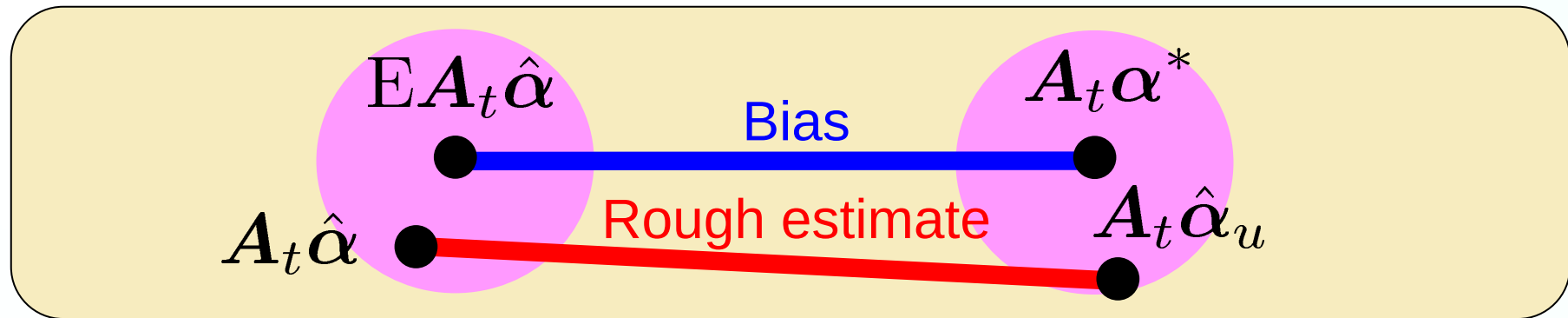
$\mathbf{A}$: Design matrix

† : Generalized inverse

$$\mathbf{A}_{i,j} = \varphi_j(\mathbf{x}_i)$$

$$\mathbf{A}\hat{\boldsymbol{\alpha}} = (\hat{f}(\mathbf{x}_1), \hat{f}(\mathbf{x}_2), \dots, \hat{f}(\mathbf{x}_n))^\top$$

Unbiased Estimator of Bias



$$\begin{aligned}
 Bias &= \|E A_t \hat{\alpha} - A_t \alpha^*\|^2 & \epsilon &= (\epsilon_1, \epsilon_2, \dots, \epsilon_n)^\top \\
 & & z &= (f(x_1), f(x_2), \dots, f(x_n))^\top \\
 &= \|A_t \hat{\alpha} - A_t \hat{\alpha}_u\|^2 - 2\langle A_t X z, A_t A^\dagger \epsilon \rangle - \|A_t (X - A^\dagger) \epsilon\|^2 \\
 &\quad \downarrow \text{E} \qquad \qquad \qquad \downarrow \text{E} \\
 \widehat{Bias} &= \|A_t \hat{\alpha} - A_t \hat{\alpha}_u\|^2 - 0 - \sigma^2 \text{tr} \left(A_t (X - A^\dagger) [A_t (X - A^\dagger)]^\top \right)
 \end{aligned}$$

$$E \widehat{Bias} = Bias$$

Unbiased Estimator of Variance¹³

$$\begin{aligned}\blacksquare \text{Var} &= \mathbb{E} \|\mathbf{A}_t \hat{\boldsymbol{\alpha}} - \mathbb{E} \mathbf{A}_t \hat{\boldsymbol{\alpha}}\|^2 \\ &= \sigma^2 \text{tr} (\mathbf{A}_t \mathbf{X} (\mathbf{A}_t \mathbf{X})^\top)\end{aligned}$$

σ^2 : Noise variance

An unbiased estimator of noise variance:

$$\hat{\sigma}^2 = \frac{\|\mathbf{y} - \mathbf{A} \mathbf{A}^\dagger \mathbf{y}\|^2}{n - p}$$

$$\blacksquare \widehat{\text{Var}} = \hat{\sigma}^2 \text{tr} (\mathbf{A}_t \mathbf{X} (\mathbf{A}_t \mathbf{X})^\top)$$

$$\mathbb{E} \widehat{\text{Var}} = \text{Var}$$

Unbiased Estimator of Test Error¹⁴

- Adding bias and variance estimators, we have an unbiased estimator of test error.

$$\hat{J}_{test} = \widehat{Bias} + \widehat{Var}$$

- For simplicity, we ignore constant terms

$$\begin{aligned}\hat{J}_t = & \|A_t X y\|^2 - 2\langle A_t X y, A_t A^\dagger y \rangle \\ & + 2\hat{\sigma}^2 \operatorname{tr} \left(A_t X (A_t A^\dagger)^\top \right)\end{aligned}$$

$$\mathbb{E} \hat{J}_t = J_t$$

$$J_t = \mathbb{E} \sum_{i=1}^{n_t} \left(\hat{f}(\mathbf{t}_i) - f(\mathbf{t}_i) \right)^2 - \sum_{i=1}^{n_t} f(\mathbf{t}_i)^2$$

Unrealizable Cases

- So far, we assumed that the model includes the underlying function.

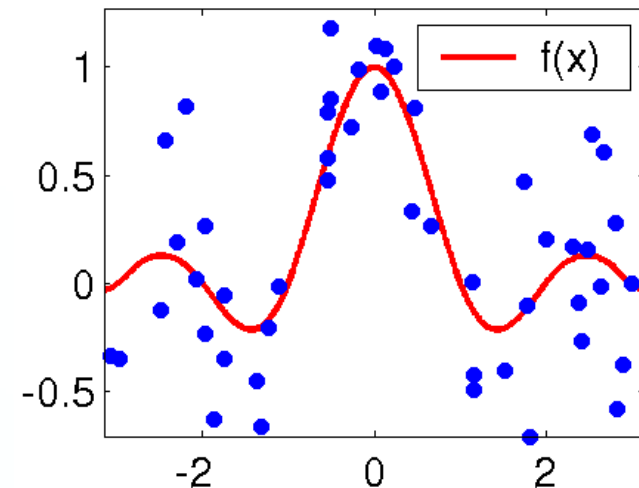
$$f(\mathbf{x}) = \sum_{i=1}^p \alpha_i^* \varphi_i(\mathbf{x}) \quad \alpha_i^* : \text{Unknown true parameters}$$

- We can prove that even when the above assumption is not rigorously fulfilled, $\hat{J}_t$ is still almost unbiased.

$$\mathbb{E} \hat{J}_t \approx J_t$$

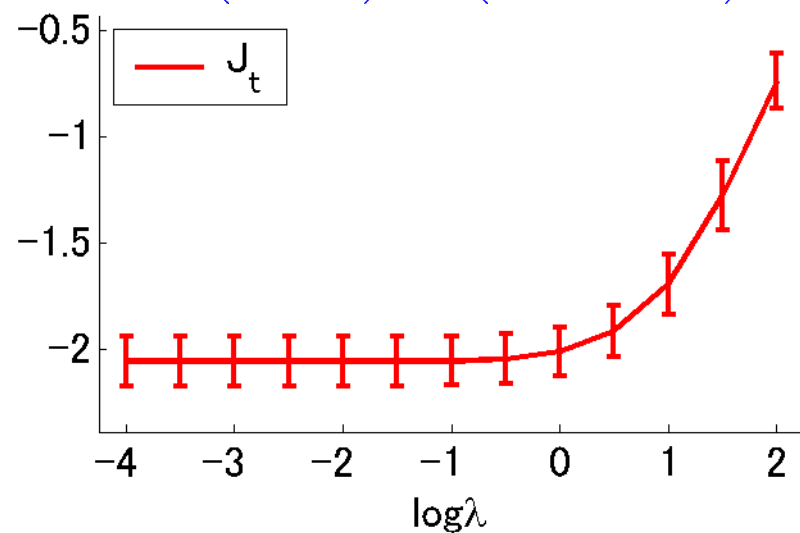
Simulation: Toy Data Sets

- **Basis functions:** 10 Gaussian functions centered at equally located points in $[-\pi, \pi]$.
- **Target function:** sinc-like function (realizable).
- **Training examples** $\{(x_i, y_i)\}_{i=1}^n$:
$$x_i \stackrel{i.i.d.}{\sim} U(-\pi, \pi)$$
$$y_i = f(x_i) + \epsilon_i, \quad \epsilon_i \stackrel{i.i.d.}{\sim} N(0, \sigma^2)$$
- **Test input points** $\{t_i\}_{i=1}^{50}$:
$$t_i \stackrel{i.i.d.}{\sim} U(-\pi, \pi)$$
- **Ridge estimation** is used for learning.

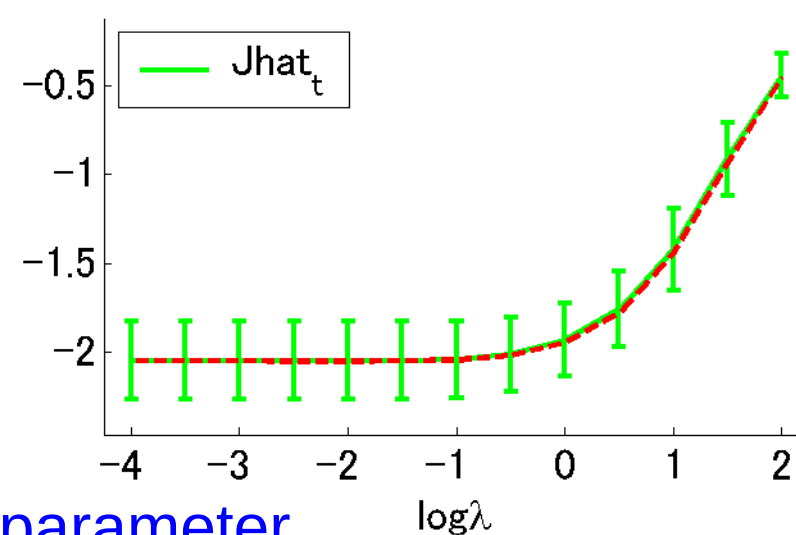
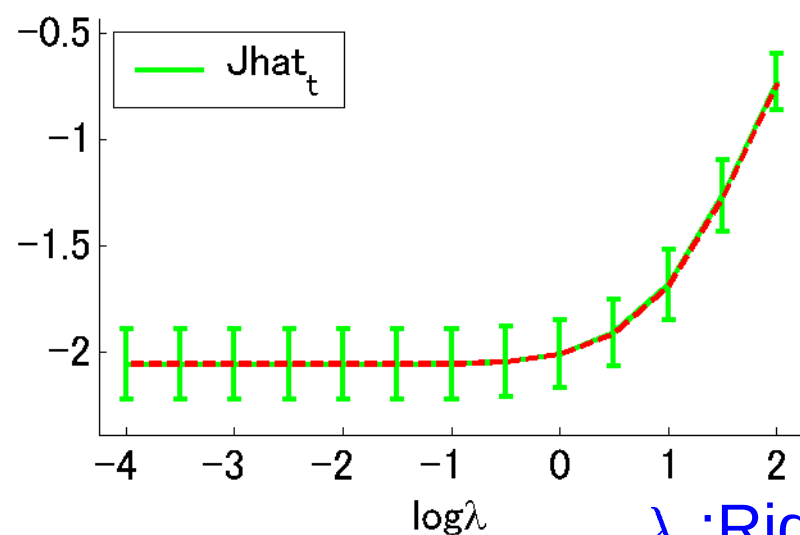
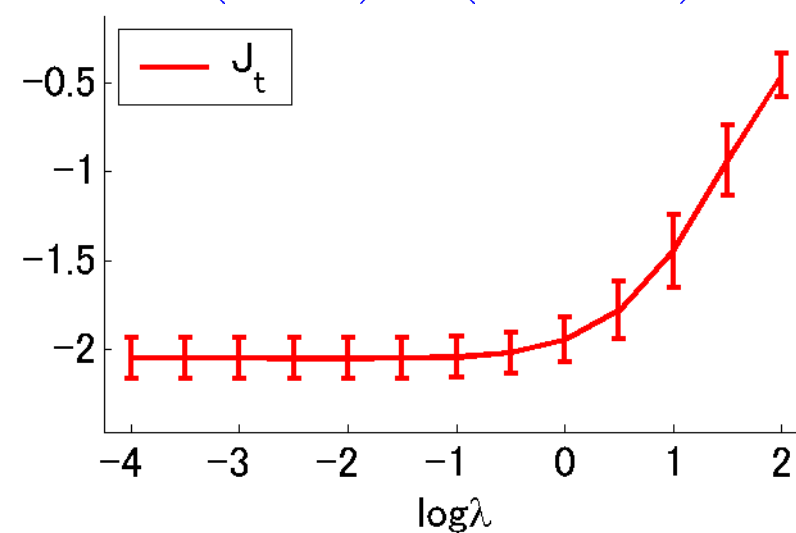


Results (1)

$$(n, \sigma^2) = (100, 0.01)$$



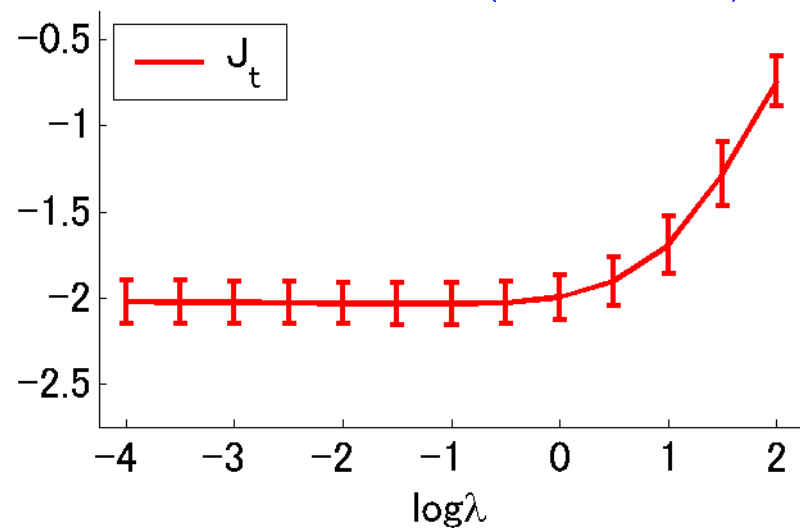
$$(n, \sigma^2) = (50, 0.01)$$



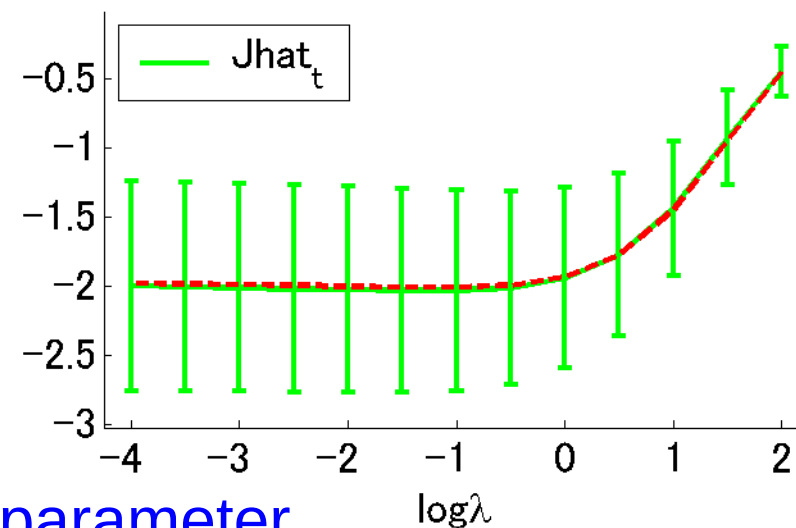
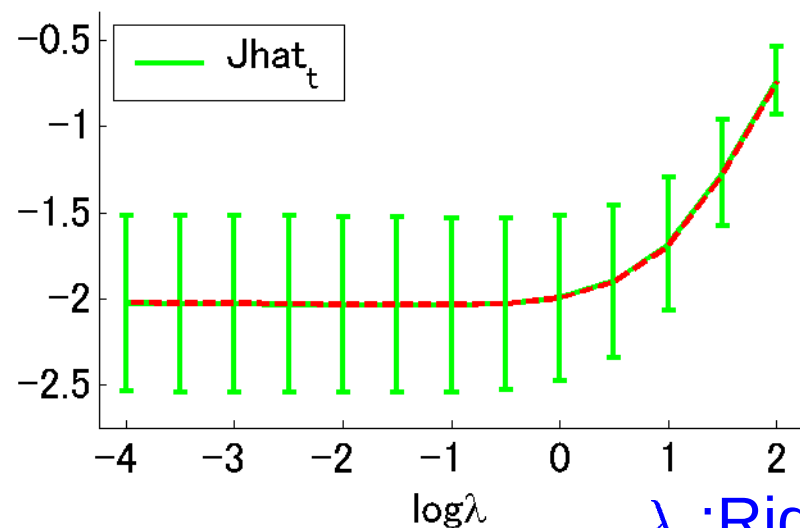
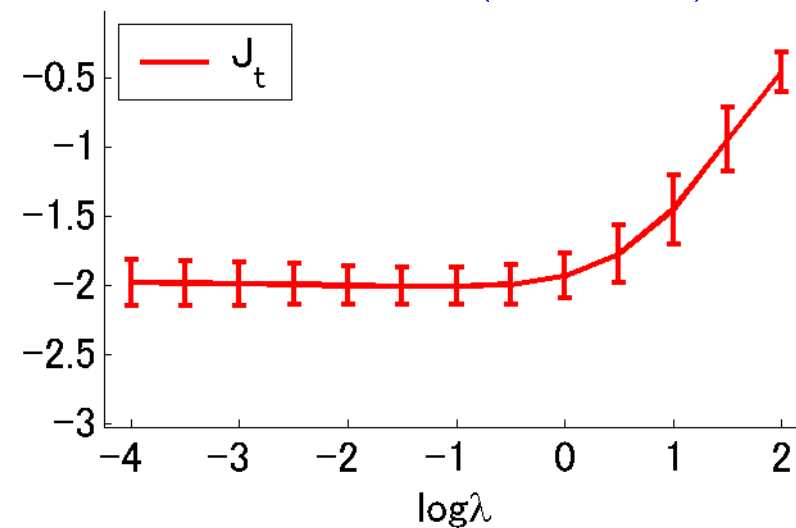
λ : Ridge parameter

Results (2)

$$(n, \sigma^2) = (100, 0.09)$$



$$(n, \sigma^2) = (50, 0.09)$$



λ : Ridge parameter

Simulation: DELVE Data Sets

19

- **Training set**: 100 randomly selected samples.
- **Test set**: 50 randomly selected samples.
- **Basis functions**: Gaussian function centered at first 50 training input points.
- **Ridge estimation** is used for learning.
- Ridge parameter is selected by **the proposed method, leave-one-out cross-validation, or an empirical Bayesian method.**

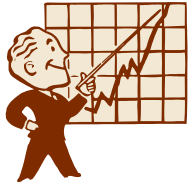
Normalized Test Errors

Mean (Standard deviation)

Data set	Proposed method	LOO cross-validation	Empirical Bayes
Boston	1.17 (0.54)	1.26 (0.58)	1.39 (0.59)
Bank-8fm	1.07 (0.29)	1.11 (0.32)	1.09 (0.31)
Bank-8nm	1.09 (0.51)	1.12 (0.56)	1.18 (0.60)
Kin-8fm	1.06 (0.32)	1.17 (0.36)	1.68 (0.48)
Kin-8nm	1.11 (0.27)	1.09 (0.24)	1.15 (0.24)

Red: Best and others with no significant difference by 99% t-test

Proposed method can be successfully applied to transductive model selection!



Conclusions

- Model selection is usually carried out so that estimated generalization error is minimized.
- When test input points are specified in advance (**transductive inference**), it is natural to choose a model so that **the test error only at the test input points is minimized**.
- We derived **an unbiased estimator of the test error at given test input points**.
- Simulation showed **the proposed method works well in practical situations**.