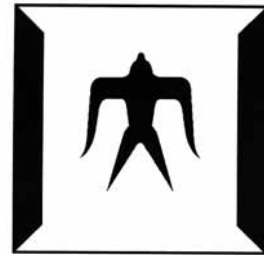


NIPS2005

Dec 5-8, 2005

Active Learning for Misspecified Models



Masashi Sugiyama
Tokyo Institute of Technology

Abstract

The goal of active learning is to determine the locations of training input points so that the generalization error is minimized. We discuss the problem of active learning in linear regression scenarios. Traditional active learning methods using least-squares learning often assume that the model used for learning is correctly specified. In many practical situations, however, this assumption may not be fulfilled. Recently, active learning methods using "importance"-weighted least-squares learning have been proposed, which are shown to be robust against misspecification of models. In this paper, we propose a new active learning method also using the weighted least-squares learning, which we call **ALICE** (Active Learning using the Importance-weighted least-squares learning based on Conditional Expectation of the generalization error). An important difference from existing methods is that we predict the **conditional expectation** of the generalization error given training input points, while existing methods predict the full expectation of the generalization error. By this difference, the training input design can be fine-tuned depending on the realization of training input points. Theoretically, we prove that the proposed active learning criterion is a more accurate predictor of the **single-trial** generalization error than the existing criterion in some sense. Numerical studies with toy and benchmark data sets show that the proposed method compares favorably to existing methods.

Linear Regression Problem

Learning target:

$$f(\mathbf{x}) : \mathbb{R}^d \rightarrow \mathbb{R}$$

Training examples:

$$\{\mathbf{x}_i, y_i\}_{i=1}^n$$

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

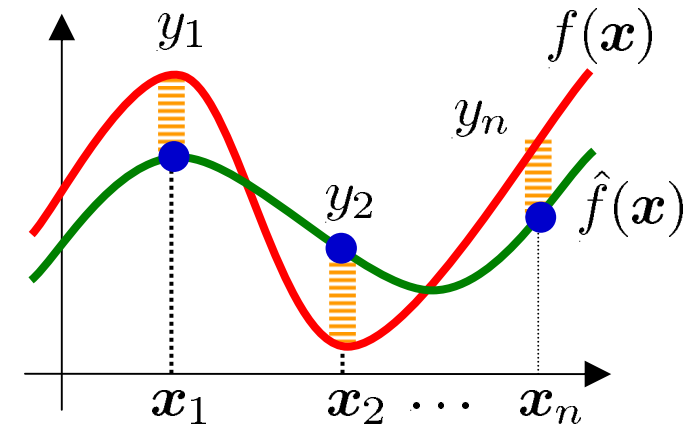
$$\{\epsilon_i\}_{i=1}^n : \text{iid, mean 0, variance } \sigma^2$$

Linear regression model:

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

α_i : Parameter

$\varphi_i(\mathbf{x})$: Basis function



Generalization Error

- Goal of regression :

Obtain a learned function $\hat{f}(x)$ that minimizes **generalization error** (expected error for unseen test input points).

- When test input points follows a density $q(x)$, generalization error is

$$\begin{aligned} G &= \int \left(\hat{f}(x) - f(x) \right)^2 q(x) dx \\ &= \|\hat{f} - f\|^2 \end{aligned}$$

Distribution of Training Input Points⁵

■ Passive learning (ordinary setting):

Draw training input points from the same density as test input points

$$\mathbf{x}_i \stackrel{i.i.d.}{\sim} q(\mathbf{x})$$

■ Active learning (experimental design):

Learner design training input density $p(\mathbf{x})$

$$\mathbf{x}_i \stackrel{i.i.d.}{\sim} p(\mathbf{x})$$

$$p(\mathbf{x}) \neq q(\mathbf{x})$$

Active Learning

- **Goal:** Design training input density $p(\mathbf{x})$ so that the generalization error G is minimized

$$G(p) = \|\hat{f}_p - f\|^2$$

$$\hat{f}_p(\mathbf{x}) = \sum_{i=1}^b \alpha_i \varphi_i(\mathbf{x}) \quad \mathbf{x}_i \stackrel{i.i.d.}{\sim} p(\mathbf{x})$$

Need to predict generalization error
before observing output values!

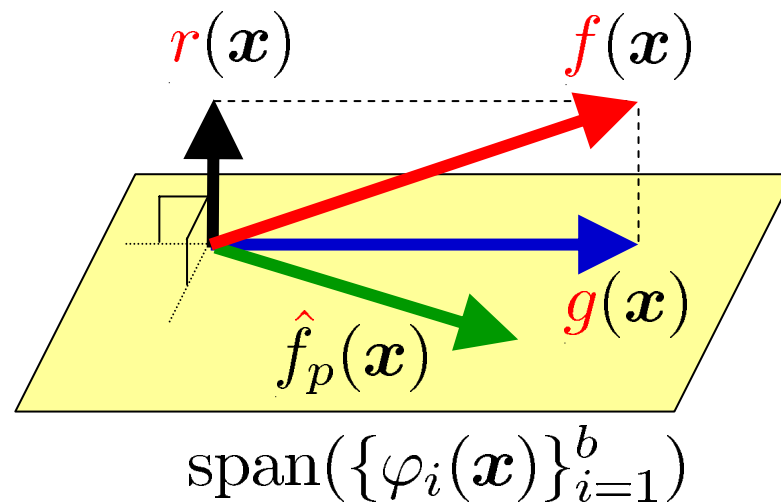
Decomposition of Learning Target⁷

$$f(\mathbf{x}) = g(\mathbf{x}) + r(\mathbf{x})$$

■ Linear part: $g(\mathbf{x}) = \sum_{i=1}^b \alpha_i^* \varphi_i(\mathbf{x})$

■ Residual: $r(\mathbf{x}) \perp g(\mathbf{x})$

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



Bias/Variance Decomposition

$$\mathbb{E}_{\epsilon} G(p) = C + B(p) + V(p)$$

$\mathbb{E}_{\epsilon}$: Expectation over noise

$$G(p) = \|\hat{f}_p - f\|^2$$

■ **Model error (constant):**

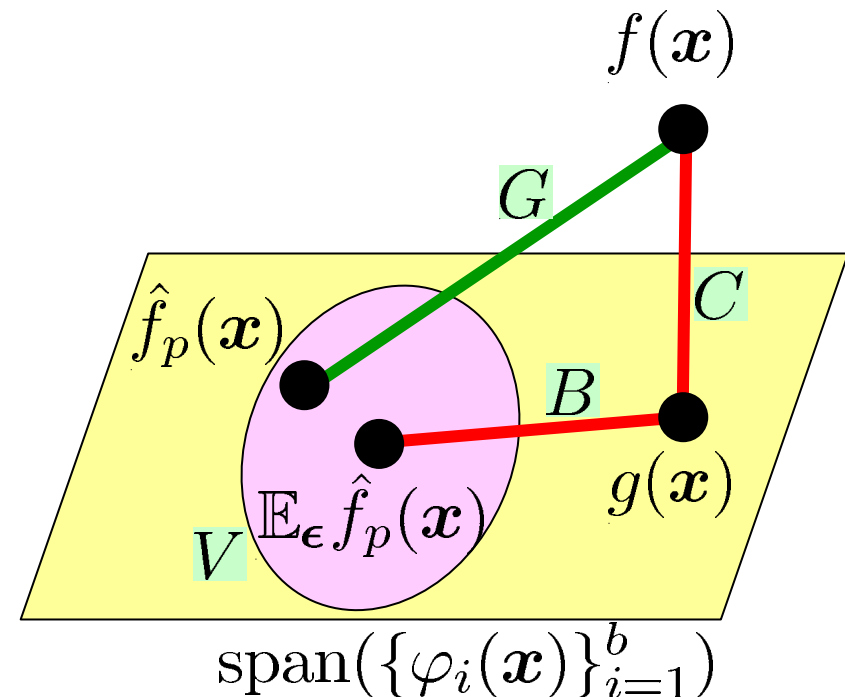
$$C = \|r\|^2$$

■ **Bias:**

$$B(p) = \|\mathbb{E}_{\epsilon} \hat{f}_p - g\|^2$$

■ **Variance:**

$$V(p) = \mathbb{E}_{\epsilon} \|\mathbb{E}_{\epsilon} \hat{f}_p - \hat{f}_p\|^2$$



Ordinary Least Squares

■ Ordinary least squares (OLS):

$$\min_{\alpha} \left[\sum_{i=1}^n \left(\hat{f}(\mathbf{x}_i) - y_i \right)^2 \right] \quad \hat{f}(\mathbf{x}) = \sum_{i=1}^b \alpha_i \varphi_i(\mathbf{x})$$

■ Solution:

$$\hat{\alpha}_O = \mathbf{L}_O \mathbf{y}$$

$$\mathbf{L}_O = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top$$

$$\alpha = (\alpha_1, \alpha_2, \dots, \alpha_p)^\top$$

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

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

Active Learning with OLS

$$\mathbb{E}_{\epsilon} G_O(p) = C + B_O(p) + V_O(p)$$

- Want to find $p(\mathbf{x})$ s.t. $\operatorname{argmin}_p \mathbb{E}_{\epsilon}[G_O(p)]$
- However, bias is hard to predict
- Variance:

$$V_O(p) = \sigma^2 \operatorname{tr}(\mathbf{U} \mathbf{L}_O \mathbf{L}_O^{\top})$$

$$U_{i,j} = \langle \varphi_i, \varphi_j \rangle_{L_2(q)}$$

$$\mathbf{L}_O = (\mathbf{X}^{\top} \mathbf{X})^{-1} \mathbf{X}^{\top}$$

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

- Variance-only active learning (OLS-Based):

$$\min_{p(\mathbf{x})} J_O(p) \quad J_O(p) = \operatorname{tr}(\mathbf{U} \mathbf{L}_O \mathbf{L}_O^{\top})$$

(e.g., Fedorov, 1972; Cohn, Ghahramani & Jordan, 1996; Fukumizu, 2000)

Can We Ignore Bias?

$$\delta = ||r||$$

- If model is correct, bias is zero.

$$\delta = 0 \quad \rightarrow \quad B_O(p) = 0$$

- If model is misspecified, bias is not zero even asymptotically

$$\delta \neq 0 \quad \rightarrow \quad \lim_{n \rightarrow \infty} B_O \neq 0$$

- Solutions:

- Take bias into account
- Use learning method with smaller bias

Importance-Weighted LS

(Shimodaira, 2000)

■ Importance-weighted LS (WLS):

$$\min_{\alpha} \left[\sum_{i=1}^n \frac{q(\mathbf{x}_i)}{p(\mathbf{x}_i)} \left(\hat{f}(\mathbf{x}_i) - y_i \right)^2 \right] \quad \hat{f}(\mathbf{x}) = \sum_{i=1}^b \alpha_i \varphi_i(\mathbf{x})$$

■ Solution:

$$\hat{\alpha}_W = L_W \mathbf{y}$$

$$L_W = (\mathbf{X}^\top \mathbf{D} \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{D}$$

$$\text{cf. } L_O = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top$$

$$\boldsymbol{\alpha} = (\alpha_1, \alpha_2, \dots, \alpha_p)^\top$$

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

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

$$\mathbf{D} = \text{diag} \left(\frac{q(\mathbf{x}_i)}{p(\mathbf{x}_i)} \right)$$

Bias

$$\delta = \|r\|$$

Model	OLS	WLS
Correct $\delta = 0$	$B_O = 0$	$B_W = 0$
Misspecified $\delta \neq 0$	$B_O \not\rightarrow 0$	$B_W \rightarrow 0$

■ Bias of WLS asymptotically vanishes!

Proposed Active Learning Method¹⁴ using WLS (ALICE)

■ Variance:

$$V_W(p) = \sigma^2 \text{tr}(\mathbf{U} \mathbf{L}_W \mathbf{L}_W^\top)$$

$$\mathbf{U}_{i,j} = \langle \varphi_i, \varphi_j \rangle_{L_2(q)}$$

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

$$\mathbf{L}_W = (\mathbf{X}^\top \mathbf{D} \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{D}$$

$$\mathbf{D} = \text{diag} \left(\frac{q(\mathbf{x}_i)}{p(\mathbf{x}_i)} \right)$$

■ Variance-only active learning with WLS: ALICE

$$\min_{p(\mathbf{x})} J(p) \quad J(p) = \text{tr}(\mathbf{U} \mathbf{L}_W \mathbf{L}_W^\top)$$

Active Learning using Importance-weighted least-square
based on Conditional Expectation of generalization error

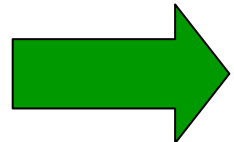
Comparison (1): OLS

- Condition for being able to ignore bias when model is misspecified:

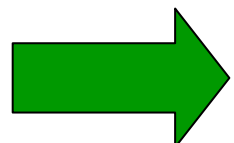
- OLS-based: $\delta = o(n^{-1/2})$

- ALICE: $\delta = o(1)$

$$\delta = \|r\|$$

 ALICE is more general

- When model is correct, bias is zero for both OLS and WLS. However, OLS has smaller variance than WLS (cf. BLUE)

 OLS-based is better

Comparison (2-1): Wiens

- **ALICE:** Minimize conditional variance given training input points

$$\min_{p(\mathbf{x})} J(p) \quad J(p) = \text{tr}(\mathbf{U} \mathbf{L}_W \mathbf{L}_W^\top)$$

- **Wiens(2000):** Minimize full expectation of variance

$$\min_{p(\mathbf{x})} J_W(p) \quad J_W(p) = \text{tr}(\mathbf{U}^{-1} \mathbf{T})$$

$$T_{i,j} = \int \varphi_i(\mathbf{x}) \varphi_j(\mathbf{x}) \frac{q(\mathbf{x})^2}{p(\mathbf{x})} d\mathbf{x}$$

- **Note:** $J_W(p)$ does not depend on $\{\mathbf{x}_i\}_{i=1}^n$

Comparison (2-2): Wiens

$$\delta = \|r\|$$

- Accuracy of generalization error prediction:
When $\delta = o(n^{-\frac{1}{4}})$, if terms of $o(n^{-3})$ are ignored

$$\mathbb{E}_{\epsilon}(\sigma^2 J_W - G_W)^2 \geq \mathbb{E}_{\epsilon}(\sigma^2 J - G_W)^2$$

ALICE is more accurate predictor
of **single-trial** generalization error G_W !

Comparison (2-3): Wiens

- **Wiens(2000):** Can analytical give optimal training input density

$$p_W^*(\mathbf{x}) = \frac{\hat{h}(\mathbf{x})}{\int \hat{h}(\mathbf{x}) d\mathbf{x}}$$

$$\hat{h}(\mathbf{x}) = q(\mathbf{x}) \sqrt{\sum_{i,j=1}^b U_{i,j}^{-1} \varphi_i(\mathbf{x}) \varphi_j(\mathbf{x})}$$

- **ALICE:** Naïvely chooses best one from several candidate densities

Comparison (3-1): K&S

- **Kanamori & Shimodaira (2003):** Minimizes generalization error without ignoring bias by choosing training input points in 2 stages
- **1st stage:** Predict generalization error using randomly chosen training input points
- **2nd stage:** Design training input density so that predicted generalization error is minimized

Comparison (3-2): K&S

- Kanamori & Shimodaira (2003): Needs randomly gathered samples in 1st stage
- Learner can't design all training input points
- Condition for being able to ignore bias when model is misspecified:
 - K&S: $\delta = O(1)$
 - ALICE: $\delta = o(1)$
- K&S is more general, when model is totally misspecified. But learning is meaningless in this case.

$$\delta = \|r\|$$

Simulations (Toy)

$$\delta = 0, 0.03, 0.3$$

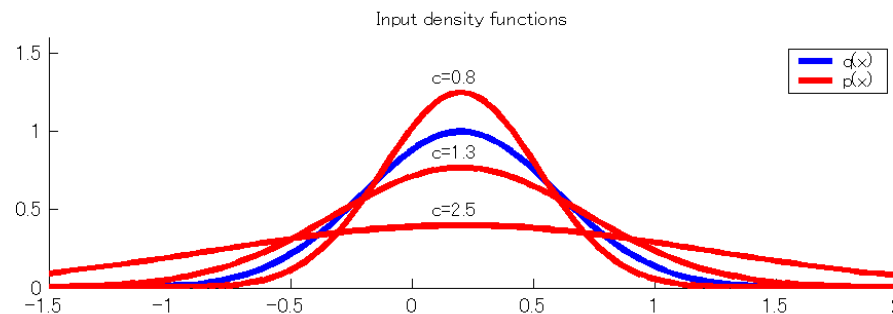
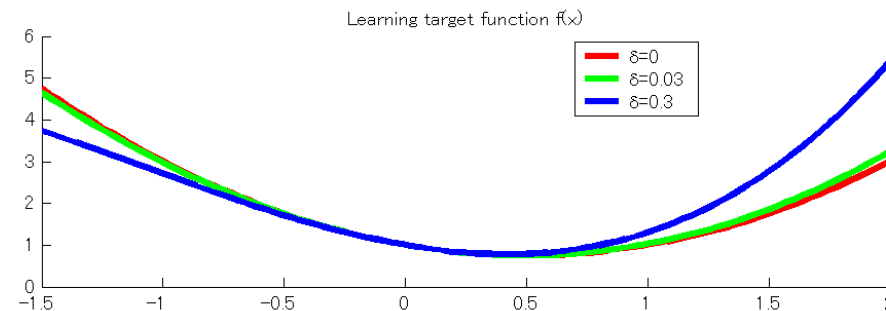
■ Learning target: $f(x) = 1 - x + x^2 + \delta x^3$

■ Model: $\hat{f}(x) = \alpha_1 + \alpha_2 x + \alpha_3 x^2$

■ Test input density: $\mathcal{N}(0.2, (0.4)^2)$

■ Training input density: $\mathcal{N}(0.2, (0.4c)^2)$

$$c = 0.8, 0.9, 1.0, \dots, 2.5$$



Obtained Generalization Error²²

Mean $\pm$ Std (1000 trials)

T-test (95%)

	$\delta = 0$	$\delta = 0.03$	$\delta = 0.3$
ALICE	2.07 $\pm$ 1.90	2.09 $\pm$ 1.90	4.28 $\pm$ 2.02
Wiens	2.42 $\pm$ 2.14	2.43 $\pm$ 2.15	4.58 $\pm$ 2.28
Wiens*	2.37 $\pm$ 2.06	2.40 $\pm$ 2.06	4.61 $\pm$ 2.21
K&S	2.96 $\pm$ 2.95	2.98 $\pm$ 2.97	5.47 $\pm$ 3.42
OLS-based	1.45 $\pm$ 1.82	2.56 $\pm$ 2.24	113 $\pm$ 63.7
Passive	3.10 $\pm$ 2.61	3.13 $\pm$ 2.61	5.75 $\pm$ 3.09

- When model is correct, OLS-based works well. But when model is misspecified, it is very poor.
- ALICE works well in all cases.

Simulations (DELVE)

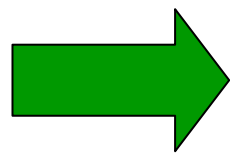
- Test input density is unknown:
Approximated by Gaussian

$$\mathcal{N}(\hat{\boldsymbol{\mu}}_{MLE}, \hat{\gamma}_{MLE}^2)$$

- Training input density is chosen from

$$\mathcal{N}(\hat{\boldsymbol{\mu}}_{MLE}, (c\hat{\gamma}_{MLE})^2) \quad c = 0.7, 0.75, 0.8, \dots, 2.4$$

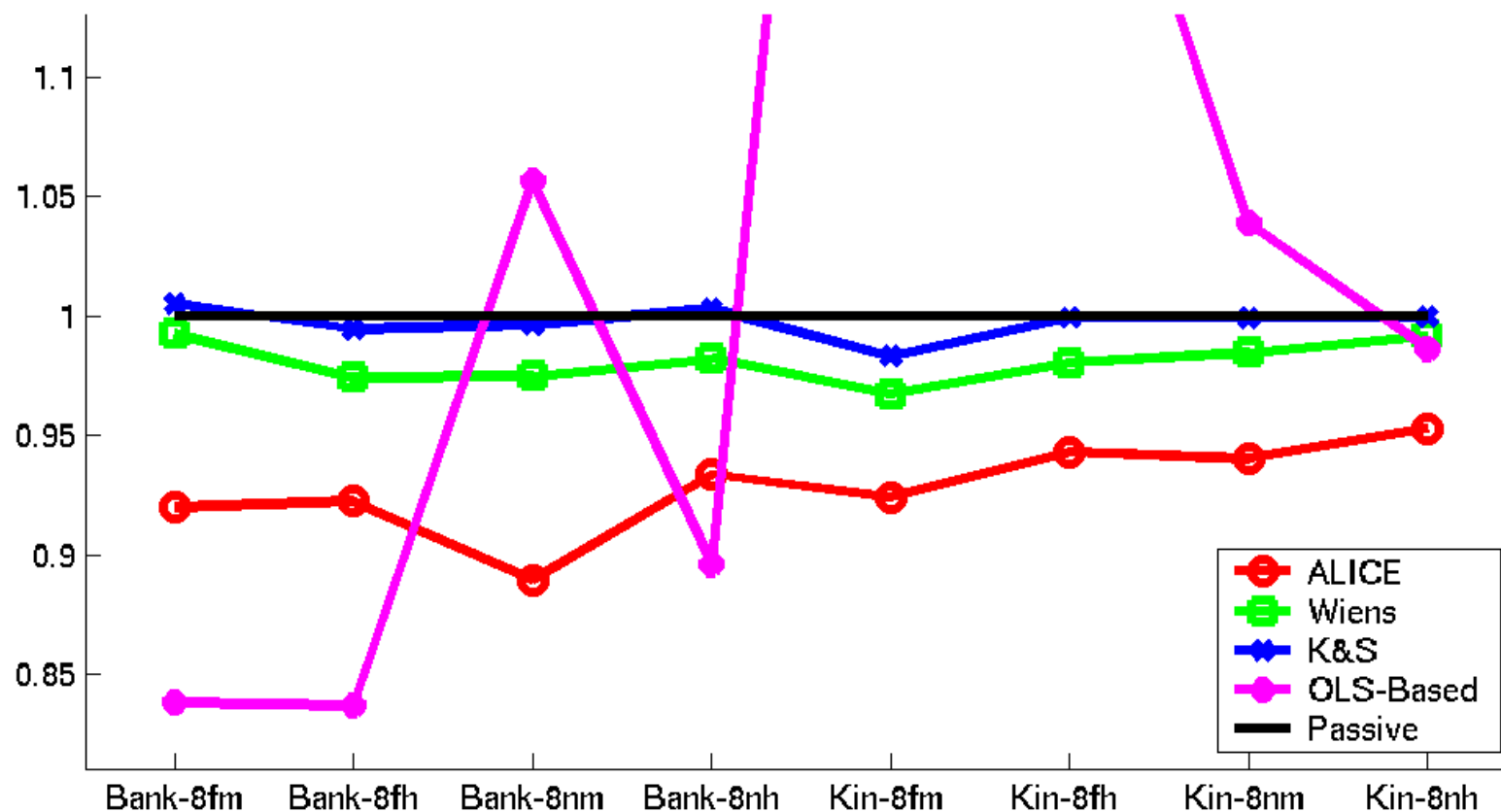
- Can not gather samples at arbitrary locations
since only 8192 samples are available



Choose samples closest to
desired locations

Obtained Test Errors

Mean over 100 trials (normalized by passive)



- OLS-based is sometimes good, but unstable.
- ALICE performs well and stable.