

DL Book 5-2

capacity, overfitting, underfitting

訓練集合の汎化能力の正負と、テスト集合にどう影響を及ぼすか $\rightarrow$ statistical learning theory.

訓練データとテストデータは

data generating process と

呼び出すデータ集合の確率分布から生成される

iid assumption

各データ集合の事例が互いに独立である (independent)

訓練とテストの集合が同一の分布に従う

これを仮定するのが一般論 (identically distributed)

標本平均の中心極限定理がどの程度当てはまるか?

1. Train loss だけだと

2. Train loss と test loss だけだと

過少適合 = 未学習 (underfitting)

過剰適合 = 過学習 (overfitting)

train loss ↓ + test loss ↑

model の capacity と control する

→ 多様性 (diversity) 増加 → 適合能力

1つ control 方法は hypothesis space

のサイズ

The model specifies which family of functions
the learning algorithm can choose
from when varying the parameters

in order to reduce a
training objective

関代2:12.02

(truly objective to what it is
 param 1: 変数と関係

what it is
 変数

The model specifies which family of function
 the learning algorithm can choose from

which どの関数 which 10:00

e.g. learning algo 10:00

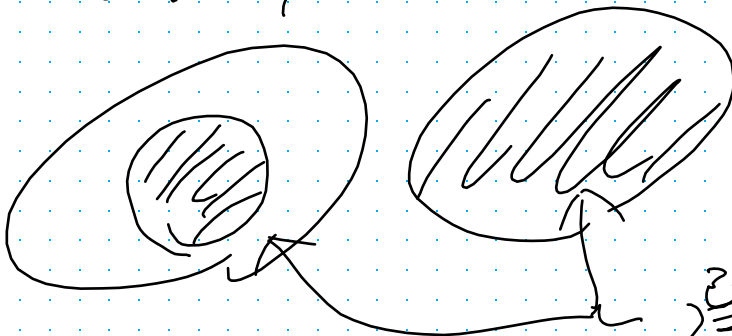
↳ early stopping 10:00

$H_{w, \text{early}}$

$H_{w, \text{early}}$

10:00

$H_{w, \text{early}} = \text{set of}$



family

選択可能な family of function

モデルの表現能力 (representation capacity)

関数集合から最適な関数
を見出すのが難しい

Early Stopping

とここで制限される

↓
学習に過剰損失を大幅に減少させる
関数を見出す

学習アルゴリズムの effective capacity が

モデルの表現力 (representational capacity) に近い
= 表現容量 可能性がある

Occam's razor

統計的学習 (SLT: Statistical Learning Theory) は

representation capacity を定量化する必要がある

いくつかあり e.g. VC次元 (Vapnik-Chervonenkis Dim.)

representation capacity 不定量化

↓

汎化 gap は ϵ 以下の容量が増えたと ϵ 以下に定量化

train sample が増えたと ϵ 以下

汎化 gap は減少する量に ϵ 以下

上に 有界 である

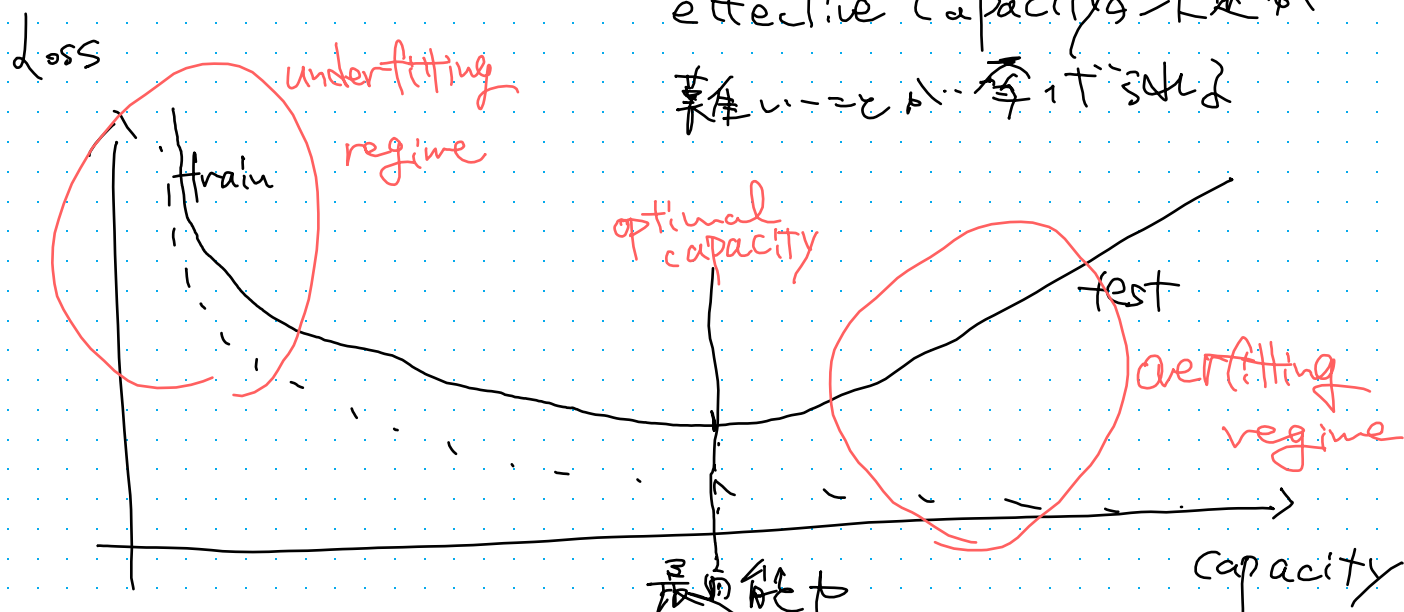
この Bound は ML の性能評価の根拠

↑

Deep の場合は使えない

!! Bound は strict ではない DL algo による

effective capacity の深さから
難い ϵ 以下に ϵ 以下に ϵ 以下



任意の high capacity という最も極端な状態に達するには

non-parametric モデルの概念を導入する

実際には実装による、理論的な抽象化にのみ着目。

→ 複雑さを訓練集合のサイズに関する関数にする

実用的な non-parametric モデルを設計することを目指す。

eg. nearest neighbor regression

$i = \underset{j}{\operatorname{argmin}} \|X_{i,:} - x\|_2^2$ のときの $\hat{y} = y_i$ である

最近傍に近づくための y_i である

すべての $X_{i,:}$ に対する y_i の平均を取りに換えて
アベリッジをとる

X と y を併行

テスト点 x を分類する場合、モデルは訓練集合の中の最近の
要素を調べる

$X_{i,:}$ の意味 Python の slice と同じ

$$A = \begin{pmatrix} 0 & 1 & 2 & 3 \\ 4 & 5 & 6 & 7 \\ 8 & 9 & 10 & 11 \end{pmatrix}$$

$$A_{1,:} = (4, 5, 6, 7)$$

$$e_j = \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 1 \\ \vdots \\ 0 \end{pmatrix}^T \leftarrow j \text{ 番目}$$

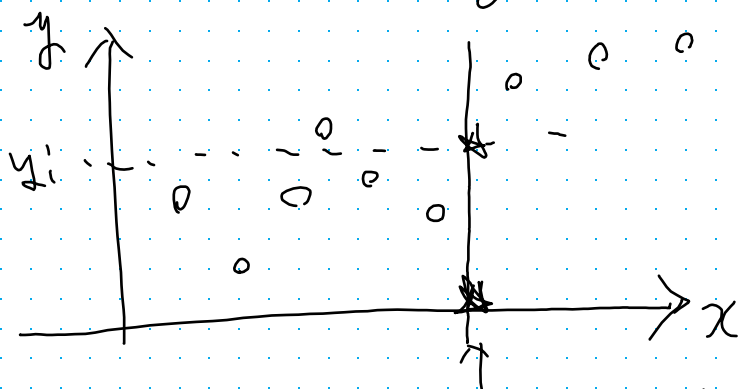
$$e_j^T W = W_{j,:}$$

j 番目の W の row

$$W = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$$

$$e_1^T W = (0 \ 1) \begin{pmatrix} a & b \\ c & d \end{pmatrix} = (0 \times a + 1 \times c, 0 \times b + 1 \times d) \\ = (c \ d)$$

$$i = \underset{j}{\operatorname{argmin}} \|X_{i,:} - x\|_2^2 \quad \text{or} \quad \hat{y} = y_i \text{ とする}$$



1つ x_i は y_i とする

1つの train sample の近隣の y の平均を $\hat{y}$ と推定する

理想的なモデルは、
データ生成する真の確率分布を完全に知っている oracle
 $p(x, y)$ がある Oracle が予測した際に 生じる誤差を
"ベイズ" 誤差 (Bayes Error) と呼ぶ

train sample $\uparrow$ 汎化が悪くなる
サンプル数は 多くする
100% 学習は できても Bayes Error
に 漸近する

訓練集合の "ベイズ" 誤差を小さく
train loss を 増やす (過適合が 難しくなる)
test loss は 訓練データと 違う場合の 仮説が 正解のため

最適な容量でのテスト誤差は Bayes Error に 漸近

訓練誤差は memorizing できる 場合がある
"ベイズ" 誤差を 下図のとおり

train sample の 容量 n が大きい

optimal capacity が存在する
後に示す。

no free lunch theorem

どの分類アルゴリズムも

(データ生成分布 P について平均すると)

過剰に観測している点で分類するときの

Error Rate は同じ

→ 逆に換えて

他の ML アルゴリズムも 全般的に P について

最適なアルゴリズムは存在しない

データ生成する全分布

を平均した場合には成り立ち

e.g.

Model	<u>MNIST</u>	<u>WILDS</u>	<u>PAcS</u>
A	90%	40%	10%
B	80%	30%	60%
C	50%	40%	80%

(こういう感じで、考え方は「テストセット全体に対して」
 良い性能とあるモデルがモデルは多く、

test error は 全て同じになる
 〃

from Dist (全部の平均)

現実世界では 確率分布の種類に依存を設けることは、

その分布に好いよ、性能を発揮するモデルを比較できる

正則化 1-71-ラング theorem は特定のタスクに特化したように
機能するML アルゴリズムを設計しなきゃいけないことと
意味して子。

e.g.
weight decay を含めて 重みの優先度を調整する
↓

一般化 (regularizer) : 10+L2コスト
optimization

この2つが大事。

解きたいタスクに適した正則化の形を選択する
必要がある

1. (a) Explain the notations
of underfitting, overfitting
Bias, variance, model capacity

Underfitting

when model has low capacity,
(both of representation and effective
error of prediction of training sample
and unseen samples
is still high. in this case.

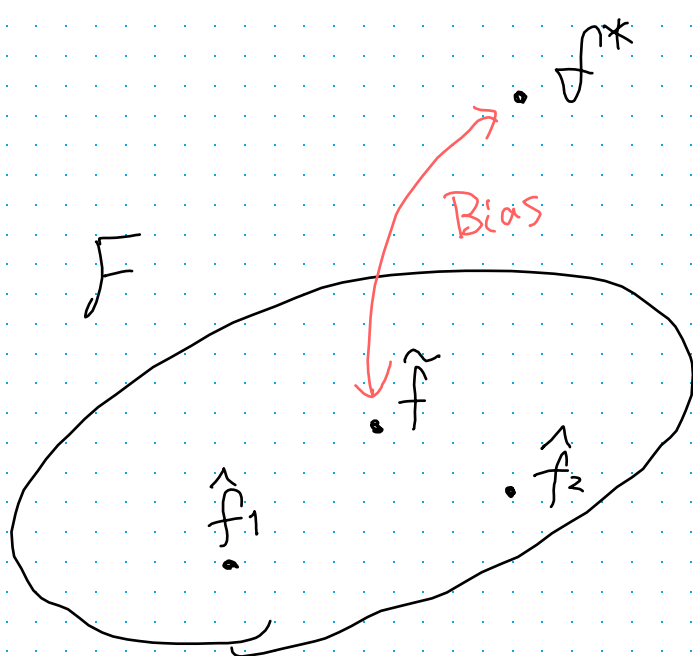
model is underfitting.

Overfitting A statistical model fits to training sample itself.

(when model capacity is enough to memorize training sample.)

It leads to degrade test error.

Bias variance



$$f^* = \arg \min_f R(f)$$

$$\hat{f} = \arg \min_{f \in F} R(f)$$

$$\hat{f}_1 = \arg \min_{f \in F} \hat{R}(f, \mathcal{Q}_1)$$

$$R(f) = \mathbb{E}_{x, y \sim \mathcal{Q}} p(f(x), y)$$

$$\hat{R}(f) = \frac{1}{m} \sum_{i=1}^m p(f(x_i), y_i)$$

$$\text{Bias} = \mathbb{E}[\hat{f}(x)] - f(x)$$

$$\text{Variance} = \mathbb{E}[\hat{f}(x)^2] - \mathbb{E}[\hat{f}(x)]^2$$

Bias-Variance Decomposition

誤差の期待値を bias-variance-wise
の3つの項に分解.

$$\mathbb{E}[(y - \hat{f}(x))^2]$$

$$= (\text{Bias})^2 + \text{Var} + \sigma^2$$

$$\hat{f}(x) = \arg\min_{f \in \mathcal{F}} R$$

$$f(x) = \arg\min_f R$$

PRML (上) p(46)

最適予測

$$h(x) = \mathbb{E}[t|x] = \int t p(t|x) dt.$$

こゝで t は ラベル のようなものと見ていい!

期待二乗損失は

$$\mathbb{E}[L] = \int \int \{y(x) - h(x)\}^2 p(x) dx + \int \int \{h(x) - t\}^2 p(x, t) dx dt$$

この二項は $y(x)$ と t が独立で $t \sim \mathcal{D}_1$ に含まれる
本質的にノイズのみ依存する

この二項を最小化する

関数 $y(x)$ を求める必要がある

分布 $p(x, y)$ に関する多数のデータ集合 D に対し

学習アルゴリズムを設計

$\Rightarrow$ 予測func $f(x; D)$ を設計する

と仮定する

このとき異なるデータ集合は異なるfuncを学習するため、

二乗損失も異なる

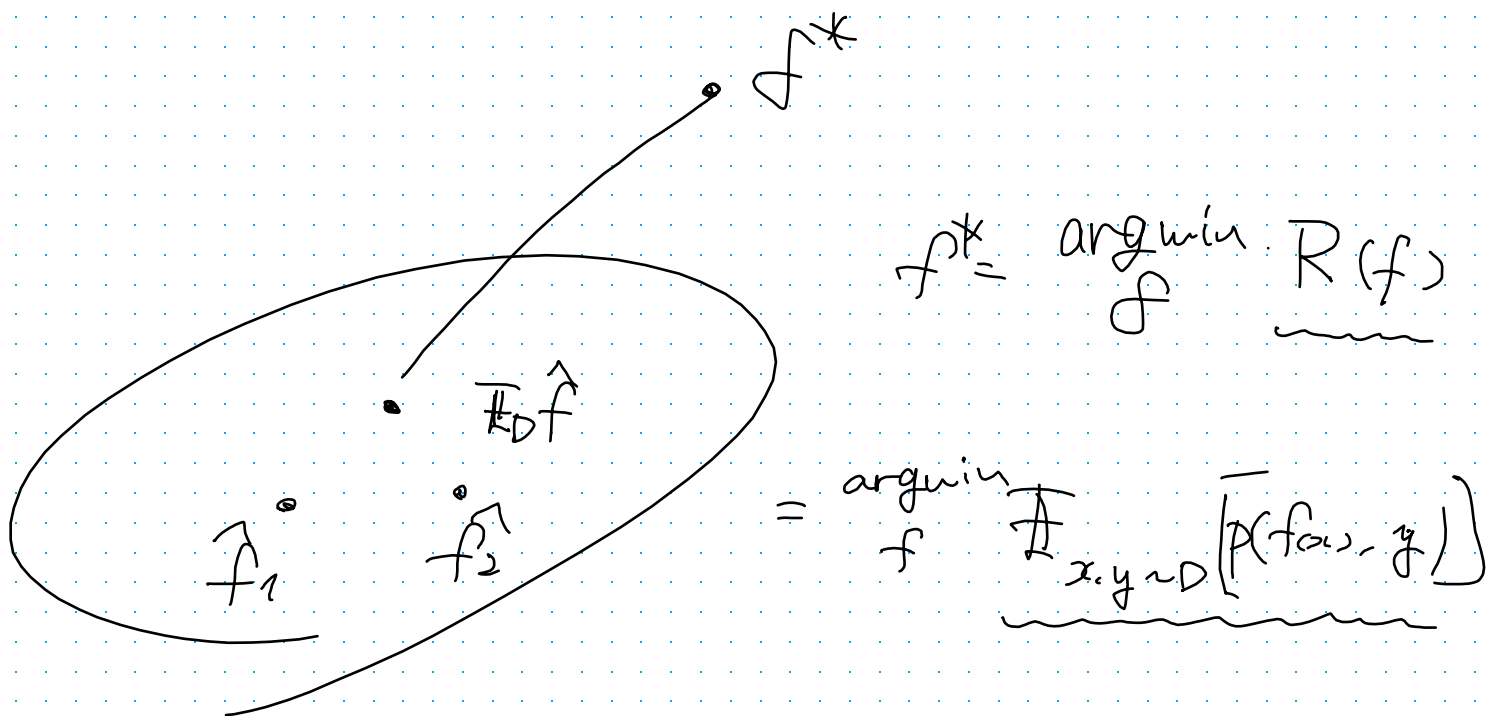
学習アルゴリズムの性能は

“データ集合のとり方に関する平均の意味での平均誤差”

$$\left\{ f(x; D) - h(x) \right\}^2 = \left\{ y(x; D) - \mathbb{E}_D[y(x; D)] + \mathbb{E}_D[y(x; D)] - h(x) \right\}^2$$

これに対し D の期待値を平均すると

$$\mathbb{E}_D \left[\left\{ y(x; D) - h(x) \right\}^2 \right] = \left\{ \mathbb{E}_D[y(x; D)] - h(x) \right\}^2 + \mathbb{E}_D \left[\left\{ y(x; D) - \mathbb{E}_D[y(x; D)] \right\}^2 \right]$$



$$\tilde{f} = \underset{f \in F}{\operatorname{argmin}} R(f)$$

$$= \underset{f \in F}{\operatorname{argmin}} \mathbb{E}_{x, y \sim D} [\bar{P}(f(x), y)]$$

$$\hat{f}_1 = \underset{f \in F}{\operatorname{argmin}} \hat{R}(f, D_1)$$

$$= \underset{f \in F}{\operatorname{argmin}} \mathbb{E}_{x, y \sim D_1} [\bar{P}(f(x), y)]$$

$$= \underset{f \in F}{\operatorname{argmin}} \frac{1}{M} \sum_{m=1}^M P(f(x), y)$$

$$\mathbb{E}_D \left[\left\{ y(x; D) - h(x) \right\}^2 \right] = \left\{ \mathbb{E}_D \left[y(x; D) \right] - h(x) \right\}^2 \leftarrow \text{Bias}^2$$

$$+ \mathbb{E}_D \left[\left\{ y(x; D) - \mathbb{E}_D \left[y(x; D) \right] \right\}^2 \right] \leftarrow \text{var}$$

Because

$$\left\{ y(x; D) - h(x) \right\}^2 = \left\{ y(x; D) - \mathbb{E}_D \left[y(x; D) \right] \right. \\ \left. + \mathbb{E}_D \left[y(x; D) \right] - h(x) \right\}^2$$

$$= \left\{ y(x; D) - \mathbb{E}_D \left[y(x; D) \right] \right\}^2 + \left\{ \mathbb{E}_D \left[y(x; D) \right] - h(x) \right\}^2 \\ - 2 \left\{ y(x; D) - \mathbb{E}_D \left[y(x; D) \right] \right\} \left\{ \mathbb{E}_D \left[y(x; D) \right] - h(x) \right\}$$

take expectation over $D \rightarrow 0$

$$\mathbb{E}_D \left[\left\{ y(x; D) - h(x) \right\}^2 \right] = \left\{ \text{Bias} \right\}^2 + \left\{ \text{Variance} \right\}$$

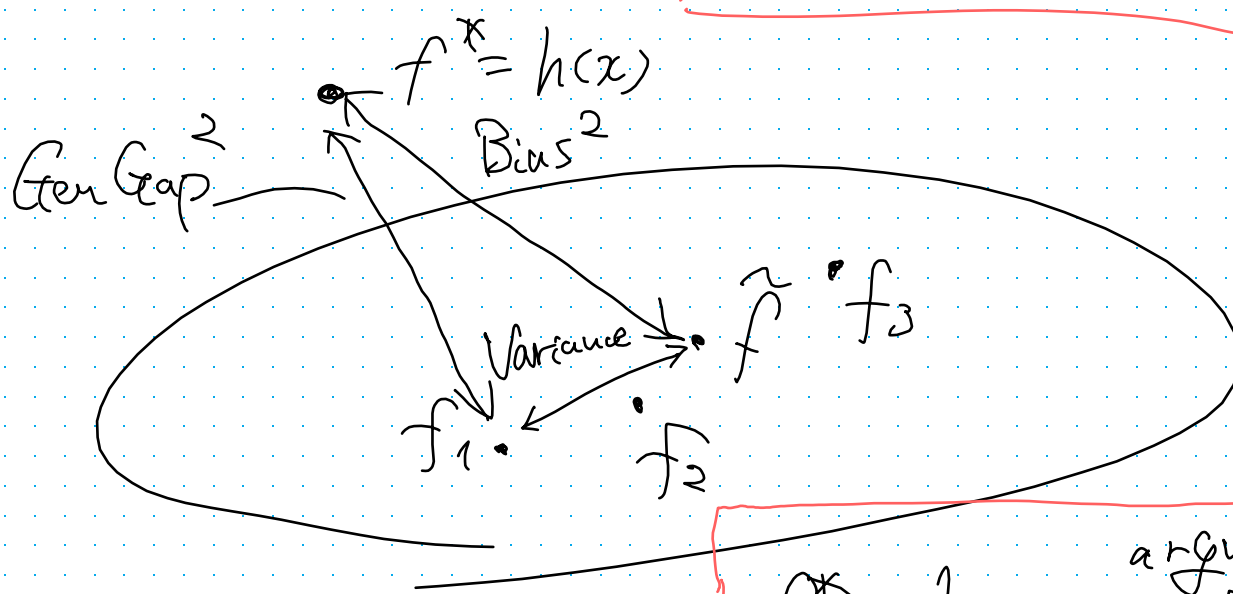
$$\mathbb{E}_D \left[\left(\int y(x; D) - h(x) \right)^2 \right] = \int \left(\mathbb{E}_D \left[\int y(x; D) \right] - h(x) \right)^2 \leftarrow \text{bias}^2$$

$$+ \mathbb{E}_D \left[\left(\int y(x; D) - \mathbb{E}_D \left[\int y(x; D) \right] \right)^2 \right] \leftarrow \text{var}$$

$$f_1 = y(x; D_1) = \left(\underset{f \in F}{\text{argmin}} y(x; D_1) \right)$$

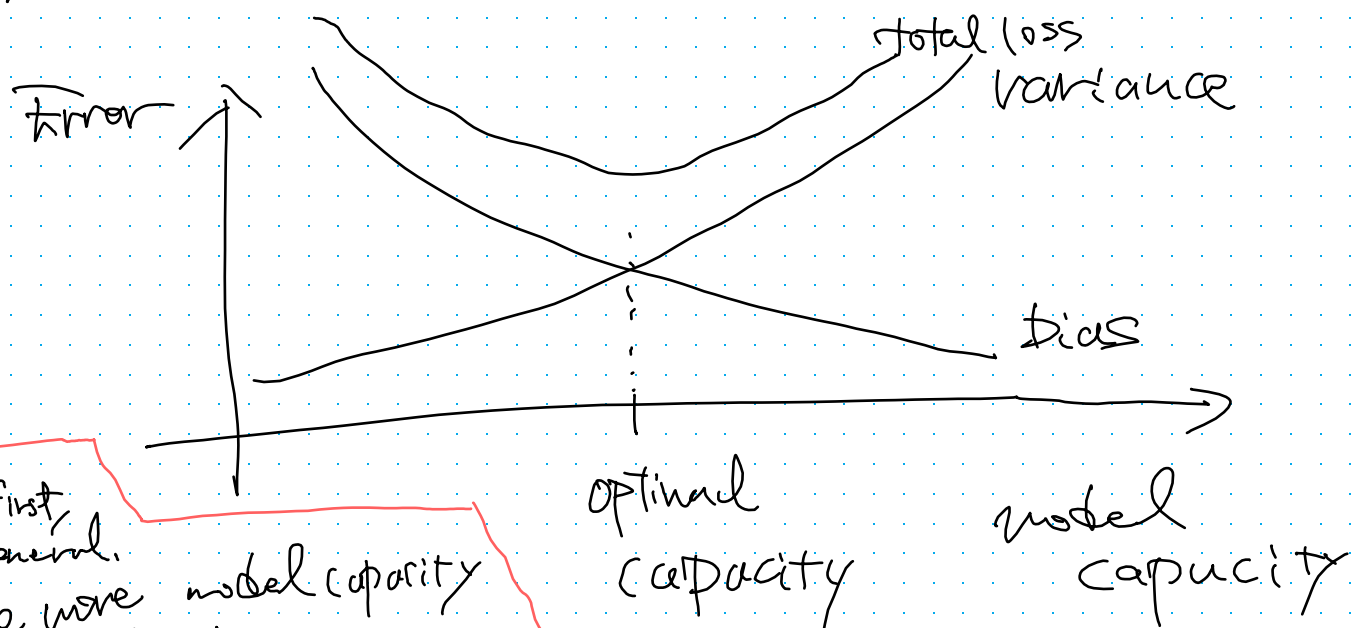
$$f_N = y(x; D_N)$$

$$\mathbb{E}_D \left[\int y(x; D) \right] = \tilde{f} = \frac{1}{N} \sum_{i=1}^N f_i$$



$$f^* = h(x) = \underset{f}{\text{argmin}} y(x; D)$$

1 (b) How do these related to one another



at first,
in general.

The more
increased,

model capacity

the more

the bias decreased

if model capacity is lower than

optimal capacity, model is not

fit to train samples, it is

called under-fitting regime.

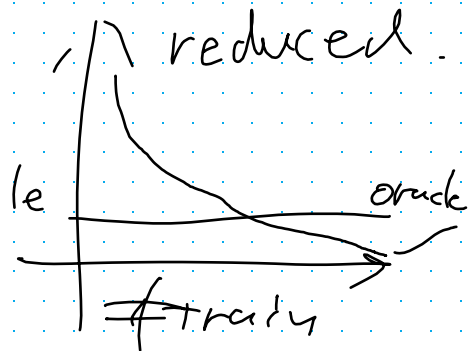
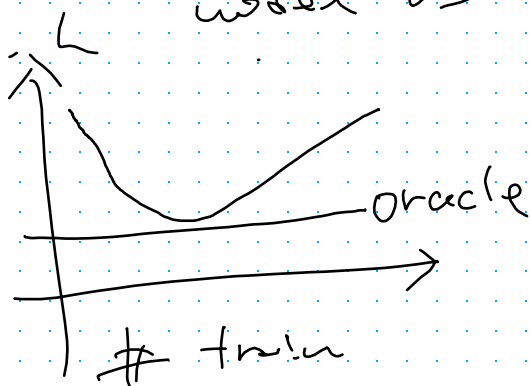
when optimal capacity is capacity when min
total loss is achieved.

The more model capacity increase,
the more bias increase,
it leads to overfitting.

1(c) How does the size of training set
affect the bias-variance trade-off

For nonparametric model, the more # sample increase
the better generalization Error
and reduce this Trade-off

For parametric model,
the more # training sample increase,
the more variance reduce and if we use
model as asymptotically unbiased, the bias
reduced.



Optimal capacity Capacity

DL Book 3.4

Bias

$\overline{\theta}$ は θ に対する
期待値

$$\text{bias}(\hat{\theta}_m) = \underbrace{\overline{\theta}}_{\substack{\uparrow \\ \text{推定値の期待値}}} - \underbrace{\theta^*}_{\substack{\uparrow \\ \text{真の値}}}$$

推定量 $\hat{\theta}_m$ は $\text{bias}(\hat{\theta}_m) = 0$ である場合には不偏
(unbiased)

と等しい $\overline{\theta} = \theta^*$ である意味する

$$\lim_{m \rightarrow \infty} \text{bias}(\hat{\theta}_m) = 0 \text{ ならば}$$

asymptotically unbiased

s^2 : 分散 (sample variance)

$$\hat{\sigma}_m^2 = \frac{1}{m} \sum_{i=1}^m (x^{(i)} - \bar{x}_m)^2$$

これは平均からの分散

unbiased sample variance

$$\tilde{\sigma}_m^2 = \frac{1}{m-1} \sum_{i=1}^m (x^{(i)} - \bar{\mu}_m)^2$$

$$\mathbb{E}[\tilde{\sigma}_m^2] = \sigma^2$$

Variance (一定量の分散は有限集合内の
事例数 m の増加と共に減らされる)

$$MSE = \mathbb{E}[(\hat{\theta}_m - \theta^*)^2] = \text{Bias}(\hat{\theta}_m)^2 + \text{Var}(\hat{\theta}_m)$$

$$P(|\hat{\theta}_m - \theta^*| > \epsilon) \rightarrow 0 \quad \text{consistency}$$

$$P(\lim_{n \rightarrow \infty} x^n = m) = 1 \text{ のとき収束}$$

almost sure convergence

consistency によらず、一定量の bias は
無限事例の増えと共に減らされる

$$m \uparrow \quad \text{bias} \downarrow \quad \text{as} \quad \lim_{m \rightarrow \infty} \hat{\theta}_m = \theta^*$$

$$\left[\begin{array}{l} \mathbb{E}[\hat{\theta}_m] = \theta^* \Leftrightarrow \hat{\theta}_m \text{ is unbiased (unbiased)} \\ \lim_{m \rightarrow \infty} \mathbb{E}[\hat{\theta}_m] = \theta^* \Leftrightarrow \hat{\theta}_m \text{ is asymptotically unbiased (asymptotically unbiased)} \end{array} \right]$$

分散はサンプル数 m の関数として減少する

2(a) Diff of Hyperparam & param.

→ which is trained by using training data.

Hyperparam affect dynamics of training

as a training configuration and used for model selection with validation data.

2(b) Two ML Model
clearly distinguishing the two categories.

We assume x is input
 y is prediction.

① MLP

$$\hat{y} = f(x; \theta)$$

this θ is model parameter

$$L = \text{loss}(f, y)$$

$$\theta_{t+1} \leftarrow \theta_t - \eta \frac{\partial}{\partial \theta} L$$

η is called learning rate

this is also hyperparameter

②

(Linear Regression)

$$\hat{y} = W^T x + b$$

$$W \in \mathbb{R}^{n \times k} \quad x \in \mathbb{R}^{k \times 1}, \quad b \in \mathbb{R}^n$$

$$L = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 + \lambda \sum_{i=1}^n W_i^2$$

this λ is regularization term as hyperparam.

$$\left(\begin{matrix} \text{ } & \leftarrow & \text{ } \\ \text{ } & \text{ } & \text{ } \\ \text{ } & \text{ } & \text{ } \end{matrix} \right) \times \begin{bmatrix} 1 \\ \text{ } \\ \text{ } \end{bmatrix} = \begin{bmatrix} 1 \\ \text{ } \\ \text{ } \end{bmatrix}$$

is model param

2(c) How can you choose hyperparam.

if statistical model is regular
we can use information criterion
such as AIC (Akaike's Information
Criterion)

if statistical model is singular
such as deep neural networks.

choose hyperparam configuration
where best validation
loss is recorded.

At first, we assume to use hold-out
strategy, that divides train data into
train, val, and test data.

Thanks to Joann's

3. Give two example for each of these problem.

o news text

identify groups of articles that are similar.
e.g. similar topic.

unsupervised
clustering
problem

(a)

k-means as ML model
nonparametric

(b)

hierarchical clustering method
by using measure of
single linkage method or
ward method.

Thanks to Ioannis!

- o Predict the likelihood of a patient having a given disease given the presence / absence of symptom and results of medical test

Supervised classification problem

- (a) Using logistic regression.

classification problem "Linear Regression and Confusion matrix"

- (b) Using NN model w/ softmax in the last layer.

height

Thanks to Joannis.

- Predict whether a flower will have an ~~eight~~ height higher than 20cm 30 days after it has been sown given characteristics/measurements of soil and environment. It will grow in.

種をまく (sow) の PP

生育する soil (土壌) と環境の特徴 + 測定値
を元に、種をまいてから 30日後に
花の高さが 20cm 以上になるかを予測したい。

Supervised

classification

(height が高い
かどうか?)

problem

(20cm 以上かそれ以下か?)

① decision tree

② SVM

どうにかしたい
問題が成立しない。

Q4 (a) 最大推定

モデルを固定したとき、事例集合から出現する
probability θ max とする param θ を求める

$$\hat{\theta}_M = \arg \max_{\theta} P(X; \theta)$$

$$= \arg \max_{\theta} \prod_{i=1}^n p(x_i; \theta)$$

$$= \arg \max_{\theta} \log p(x_1; \theta) p(x_2; \theta) \dots p(x_n; \theta)$$

$$= \arg \max_{\theta} \sum_{i=1}^n \log p(x_i; \theta)$$

at first, let θ as model param
and x as input

MLE is a method to choose model θ where the likelihood based on model θ is maximized.

$$\hat{\theta}_{MLE} = \underset{\theta}{\operatorname{argmax}} L(x|\theta)$$

(4) Explain the concept of identically and independently distributed (i.i.d) random variables

1. identically means distributed from same distribution

e.g. $x_1 \sim p(x) \quad x_2 \sim p(x)$

2. independently distributed means

$$P(X_1 = x_1, X_2 = x_2) = P(X_1 = x_1) P(X_2 = x_2)$$

独立変数の同分布列

random variable $X_1, \dots, X_n$ 同分布列

任意の実数 x に対し

$$P(X_1 = x) = P(X_2 = x) = \dots = P(X_n = x)$$

が成り立つこと

4.(c) Let $\theta > 0$ $x_1, \dots, x_n \in \mathbb{R}$ are

drawn i.i.d.

from the dist w/ density function

$$f_\theta(x) = \frac{1}{2\theta} \exp\left(-\frac{|x|}{\theta}\right)$$

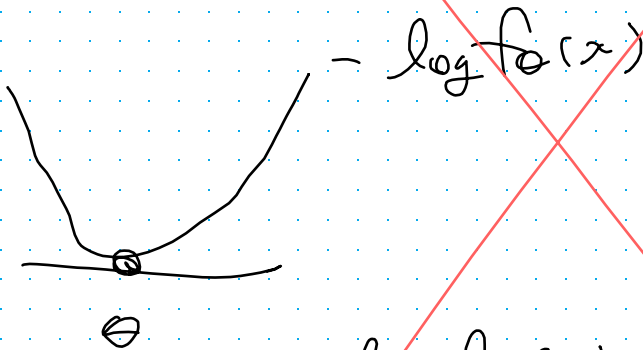
consider negative log-likelihood

$$-\log L(\theta) = -\log(2\theta) + \frac{|x|}{\theta}$$

$$\hat{\theta}_{ML} = \underset{\theta}{\operatorname{argmax}} f_{\theta}(x)$$

$$= \underset{\theta}{\operatorname{argmin}} -\log f_{\theta}(x)$$

$$= \underset{\theta}{\operatorname{argmin}} \left\{ \frac{|x|}{\theta} - \log 2\theta \right\}$$



$$\frac{\partial -\log f_{\theta}(x)}{\partial \theta} = -(x|\theta^{-2} - \log 2 - (\log \theta)')$$

$$\Leftrightarrow \frac{\partial}{\partial \theta} \text{NLL} = -\frac{|x|}{\theta^2} - \frac{1}{\theta} - \log 2$$

$$\frac{\partial}{\partial \theta} \text{NLL} = 0 \quad \Leftrightarrow$$

$$f_{\theta}(x) = \frac{1}{2\theta} \exp\left(-\frac{|x|}{\theta}\right)$$
$$= f(x; \theta)$$

We first obtain the likelihood by multiplying the probability density function for each X_i :

$$L(\theta) = \prod_{i=1}^n f_{\theta}(x_i) = \prod_{i=1}^n \frac{1}{2\theta} \exp\left(-\frac{|x_i|}{\theta}\right)$$
$$= \left(\frac{1}{2\theta}\right)^n \exp\left(-\frac{\sum_{i=1}^n |x_i|}{\theta}\right)$$

Instead of directly maximizing the likelihood, we instead maximize the log likelihood

$$\log L(\theta) = -n \log(2\theta) - \frac{\sum_{i=1}^n |x_i|}{\theta}$$

$$\log L(\theta) = -n \log(2\theta) - \frac{\sum_{i=1}^n |x_i|}{\theta}$$

To maximize this function,
we take a derivative wrt θ

$$\frac{d}{d\theta} L(\theta) = \frac{d}{d\theta} (\log 2 + \log \theta) \times -n$$

$$= \frac{d}{d\theta} \left(- \frac{\sum_{i=1}^n |x_i|}{\theta} \right)$$

$$= -\frac{n}{\theta} + \frac{\sum_{i=1}^n |x_i|}{\theta^2}$$

we set this derivative equal to zero
then solve for θ

$$\theta = \frac{1}{n} \sum_{i=1}^n |x_i|$$

(d) $x_1, \dots, x_n \in \mathbb{R}$ would have been identically distributed but not independent would that have changed your answer to the previous question? why?

Ans

Yes, it changed.

Cuz, we cannot use following equation

$$P(x_1, x_2, x_3, \dots, x_n) = \prod_{i=1}^n P(x_i)$$