



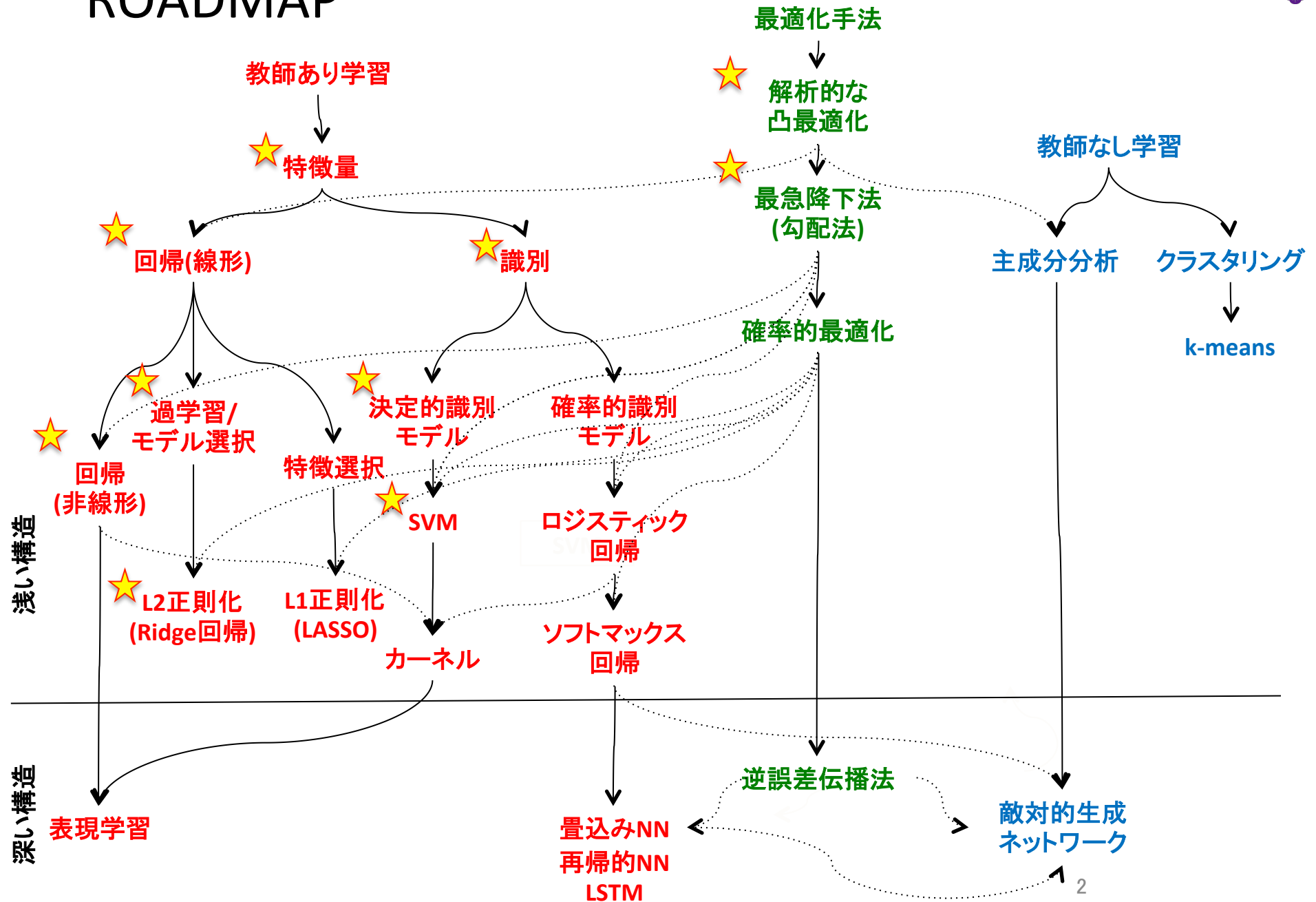
機械学習(6)

カーネル/確率的識別モデル1

情報科学類 佐久間 淳

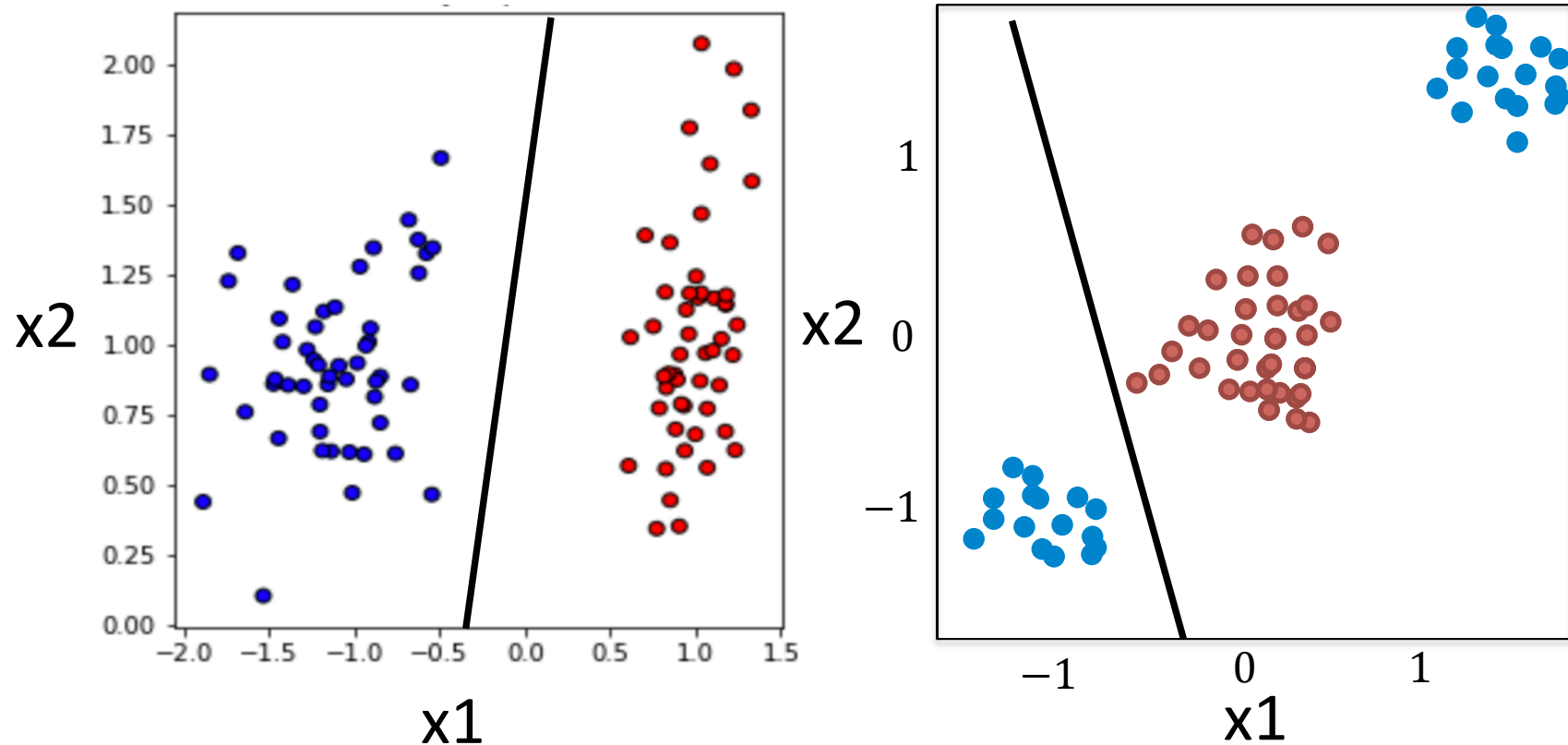


ROADMAP

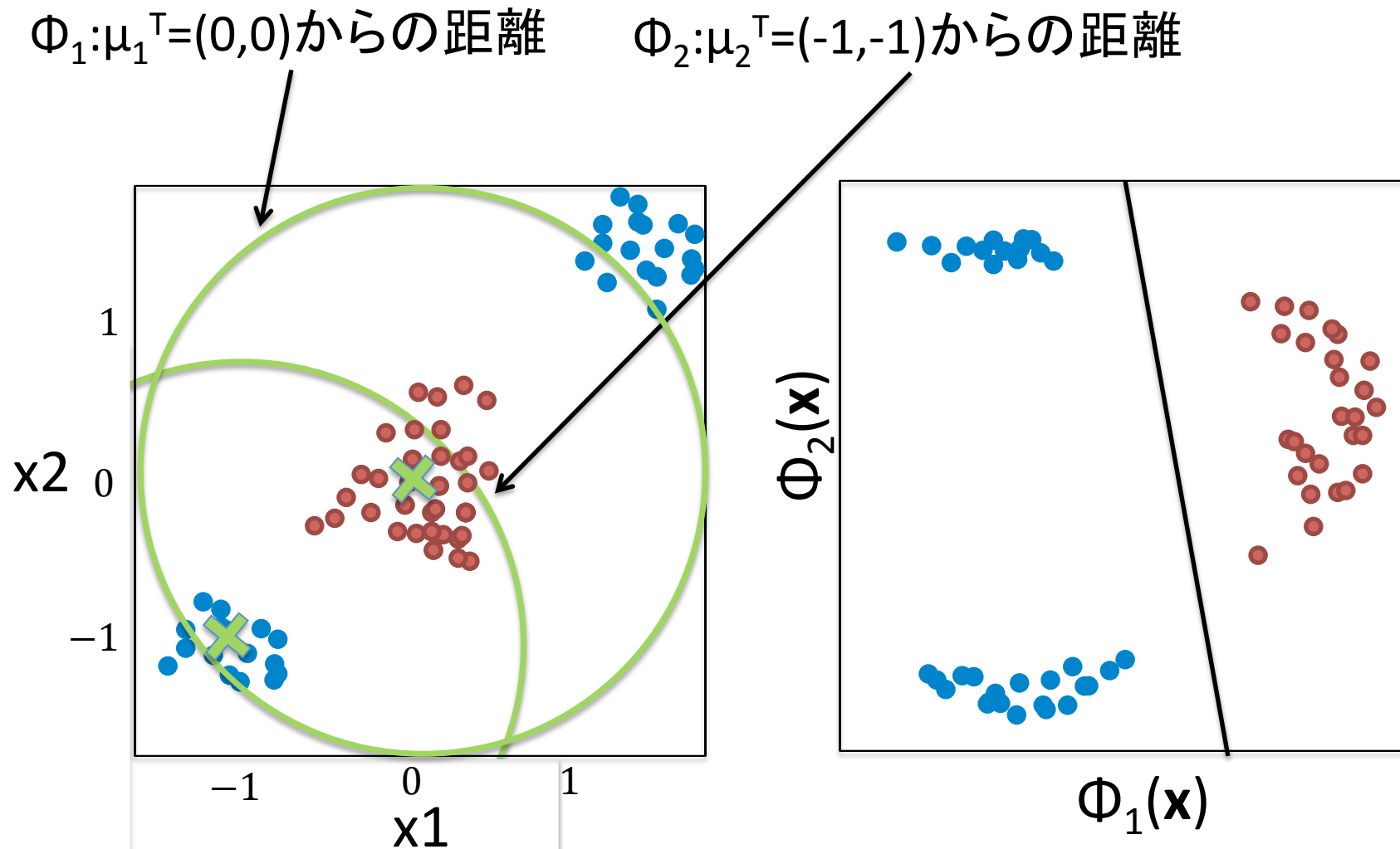




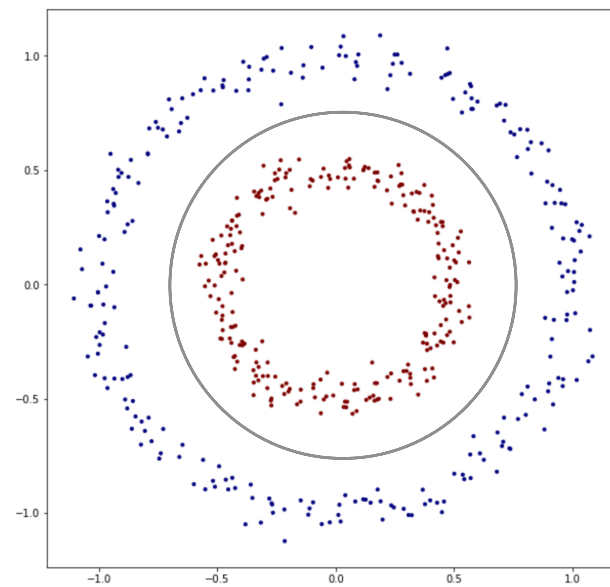
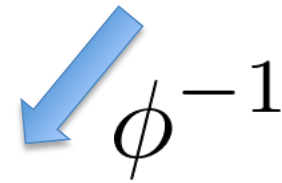
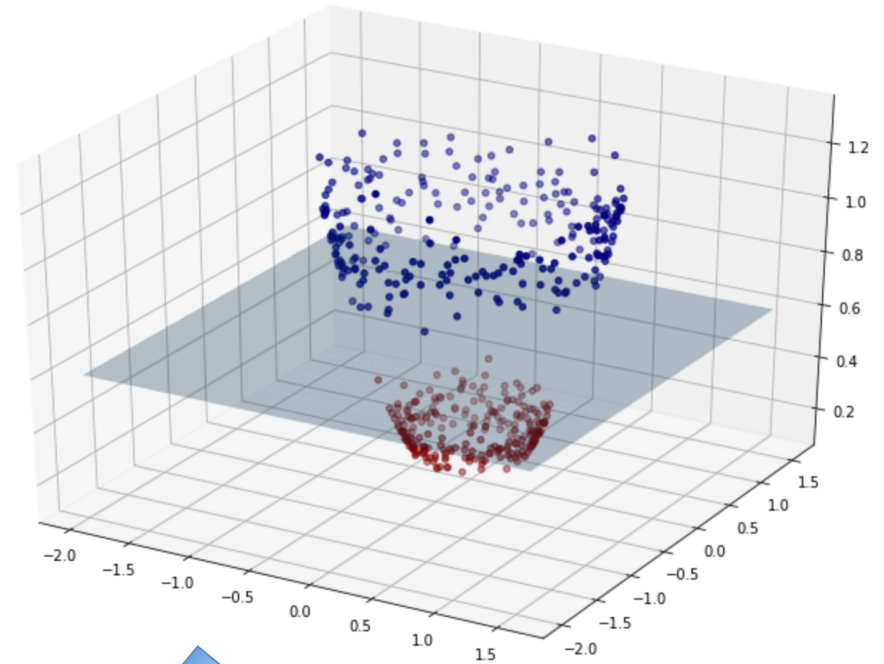
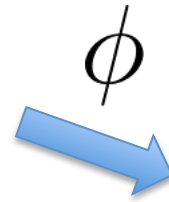
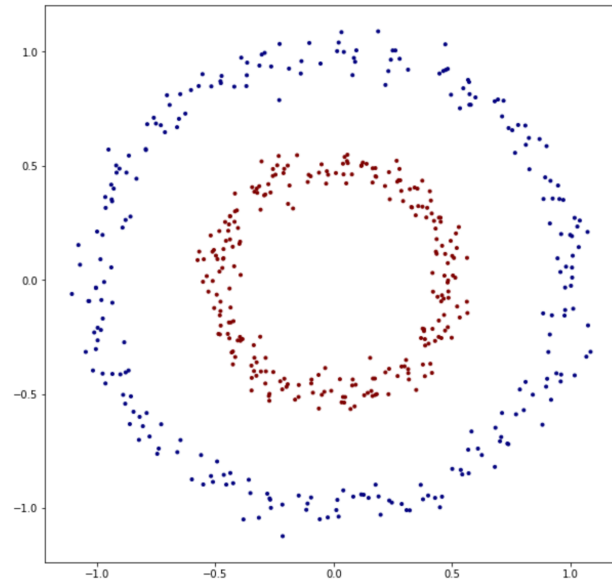
線形分類器からカーネルへ



- 青と赤を線形モデル(超平面)で分離したい
- 線形モデルでは限界がある



- 青と赤を超平面(線形モデル)で分離したい
 - 左はどんな直線でも線形分離できない
 - 非線形な変換を導入すると線形分離可能

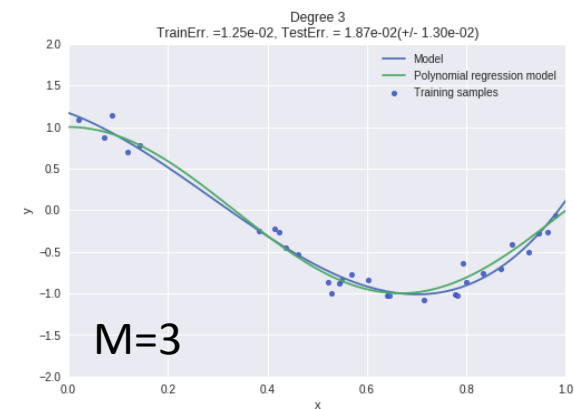
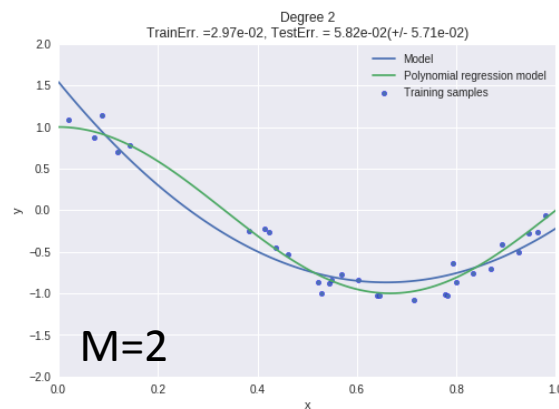
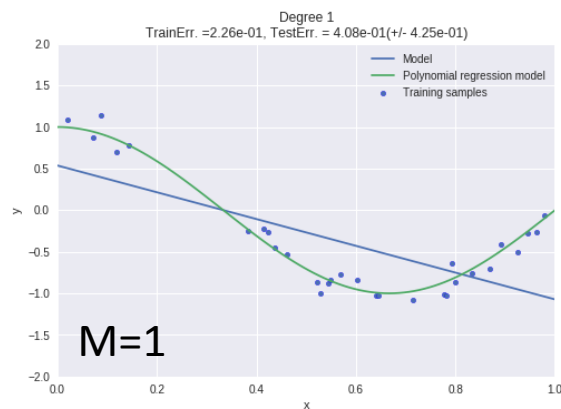




再掲

1次元多項式特徴量による回帰

- 特徴量関数 $\phi : \mathbb{R} \rightarrow \mathbb{R}^{M+1}$
 - 多項式特徴量 $\phi(x) = (x^0, x^1, x^2, \dots, x^M)$
- 多項式特徴量による線形回帰モデル
$$t = w^T \phi(x)$$





交互作用項を持つ多項式特徴量

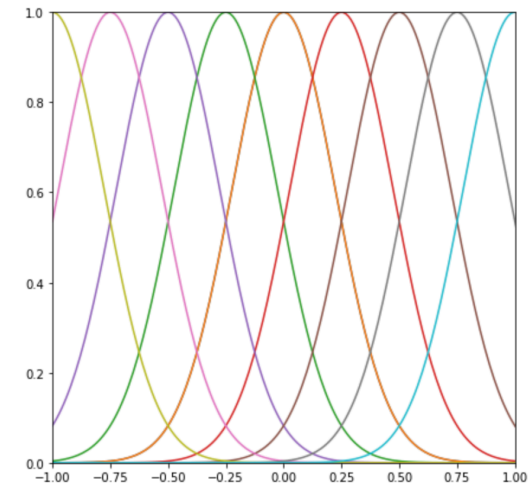
- 2次元2次多項式特徴量
 $\phi(x)^T = (1, x_1, x_2, x_1^2, x_2^2, x_1 x_2)$
- 2次元3次多項式特徴量
 $\phi(x)^T = (1, x_1, x_2, x_1^2, x_2^2, x_1 x_2, x_1^3, x_2^3, x_1^2 x_2, x_1 x_2^2)$
- 3次元2次多項式特徴量
 $\phi(x)^T = (1, x_1, x_2, x_3, x_1^2, x_2^2, x_3^2, x_1 x_2, x_2 x_3, x_1 x_3)$
- 特徴量同士の積(交互作用項)が含まれることに注意



その他の特徴量関数

- ガウスカーネル特徴量

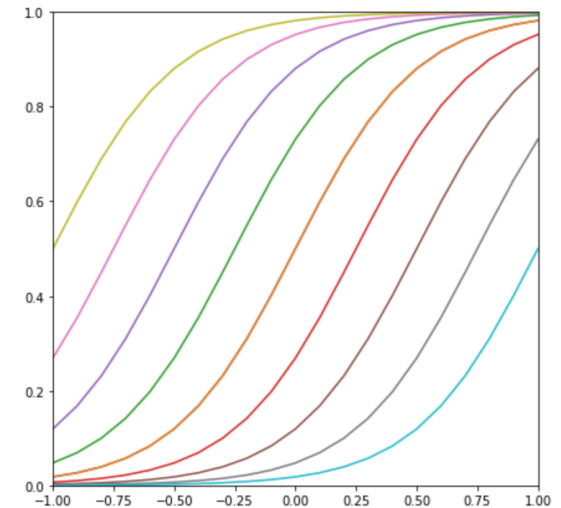
$$\phi_j(x) = \exp \left\{ -\frac{(x - \mu_j)^2}{2s^2} \right\}$$



- シグモイド関数

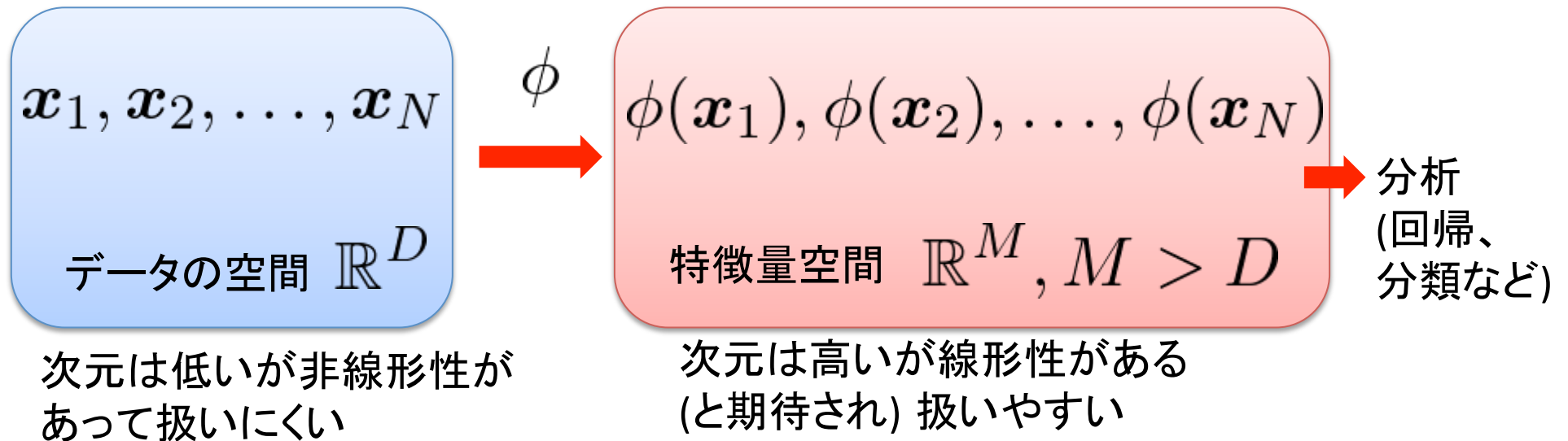
$$\phi_j(x) = \sigma \left(\frac{x - \mu_j}{s} \right),$$

$$\sigma(a) = \frac{1}{1 + \exp(-a)}$$





特徴量の考え方





特徴量次元の高さを実感する

- 交互作用項を持つ多項式特徴量
- 2次元ベクトルの2次までの組み合わせ
 $\phi(x)^T = (1, x_1, x_2, x_1^2, x_2^2, x_1x_2)$
- 2次元ベクトルの3次までの組み合わせ
 $\phi(x)^T = (1, x_1, x_2, x_1^2, x_2^2, x_1x_2, x_1^3, x_2^3, x_1^2x_2, x_1x_2^2)$
- D次元ベクトルM次多項式だと...?
 - 特徴量次元数が組合せ爆発
- そのような膨大な次元数は扱えない？



多項式カーネル

- 特徴量ベクトル (2次元3次多項式)

$$\phi(\mathbf{x})^T = (x_1^3, x_2^3, 1, \sqrt{3}x_1^2x_2, \sqrt{3}x_1x_2^2, \sqrt{3}x_2^2, \sqrt{3}x_2, \sqrt{3}x_1^2, \sqrt{3}x_1, \sqrt{6}x_1x_2),$$

$$\phi(\mathbf{x}')^T = (x_1'^3, x_2'^3, 1, \sqrt{3}x_1'^2x_2', \sqrt{3}x_1'x_2'^2, \sqrt{3}x_2'^2, \sqrt{3}x_2', \sqrt{3}x_1'^2, \sqrt{3}x_1', \sqrt{6}x_1'x_2').$$

- 多項式カーネル

- 3次多項式カーネル $k(\mathbf{x}, \mathbf{x}') = (\mathbf{x}^T \mathbf{x}' + 1)^3$

- M次多項式カーネル $k(\mathbf{x}, \mathbf{x}') = (\mathbf{x}^T \mathbf{x}' + 1)^M$

- カーネルトリック

$$k(\mathbf{x}, \mathbf{x}') = \phi(\mathbf{x})^T \phi(\mathbf{x}')$$

6.1 多項式カーネル

交互作用項を含む多項式特徴量と多項式カーネルについて以下の問題に答えなさい。

1. 2次元ベクトルの3次多項式特徴量を考える。また3次多項式カーネルを $k(\mathbf{x}, \mathbf{x}') = (\mathbf{x}^T \mathbf{x}' + 1)^3$ と定義する。二次多項式特徴量同士の内積は、カーネルを通じて $k(\mathbf{x}, \mathbf{x}') = \phi(\mathbf{x})^T \phi(\mathbf{x}')$ と計算できることを示せ。
2. 多項式次元数の増加につれて交互作用項の数は指数的に増加するため、計算量的な意味で扱いが厄介であるが、予測においては重要な役割を果たすことがある。例えば、spam メール判定ルールとして、 x_1 :実行可能ファイルが添付ファイルに含まれる (false=0/true=1), x_2 :同一メールが K 以上の異なるドメインのメールアドレスに送信されている (false=0/true=1)、という二つの特徴量を用いて、両者が真である場合のみスパムメールと判定する ($x_1 \wedge x_2 = 1$) によってスパム判定することを考えよう。このような分類は、交互作用項 $x_1 x_2$ を使って自然に表現できる $x_1 x_2 \geq 1$ 。この例にならって、交互作用項が分類の予測に役立つ特徴量と分類器の例をあげなさい。



カーネルトリック (ガウシアンカーネル) と 無限次元特徴量

- ガウシアンカーネル $k(\mathbf{x}, \mathbf{x}') = \exp\left(-\frac{\|\mathbf{x} - \mathbf{x}'\|^2}{\sigma^2}\right)$
 $k(\mathbf{x}, \mathbf{x}') = \exp\left(-\frac{\mathbf{x}^T \mathbf{x}}{\sigma^2}\right) \exp\left(\frac{2\mathbf{x}^T \mathbf{x}'}{\sigma^2}\right) \exp\left(-\frac{\mathbf{x}'^T \mathbf{x}'}{\sigma^2}\right)$

テイラー展開 $\exp(x) = \sum_{n=0}^{\infty} \frac{1}{n!} x^n$ より

$$\begin{aligned} \exp\left(\frac{2\mathbf{x}^T \mathbf{x}'}{\sigma^2}\right) &= \sum_{k=0}^{\infty} \frac{1}{k!} \left(\frac{2\mathbf{x}^T \mathbf{x}'}{\sigma^2}\right)^k = \sum_{k=0}^{\infty} \frac{1}{\sqrt{k!}} \left(\frac{\sqrt{2}\mathbf{x}^T}{\sigma}\right)^k \cdot \frac{1}{\sqrt{k!}} \left(\frac{\sqrt{2}\mathbf{x}'}{\sigma}\right)^k \\ &= \sum_{k=0}^{\infty} \varphi_k(\mathbf{x})^T \varphi_k(\mathbf{x}') = \psi(\mathbf{x})^T \psi(\mathbf{x}') \end{aligned}$$

ここで、 $\varphi_k(\mathbf{x}) = \frac{1}{\sqrt{k!}} \left(\frac{\sqrt{2}\mathbf{x}}{\sigma}\right)^k$, $\psi(\mathbf{x})^T = (\varphi_0(\mathbf{x})^T, \varphi_1(\mathbf{x})^T, \dots)$
としている。

$$\therefore k(\mathbf{x}, \mathbf{x}') = \exp\left(-\frac{\mathbf{x}^T \mathbf{x}}{\sigma^2}\right) \psi(\mathbf{x})^T \psi(\mathbf{x}') \exp\left(-\frac{\mathbf{x}'^T \mathbf{x}'}{\sigma^2}\right) = \phi(\mathbf{x})^T \phi(\mathbf{x}')$$

無限次元の特徴量も、内積に限って言えば、有限の時間で計算できる



カーネルをどう活用するか？

- カーネルを使えば、内積だけなら、非常に高い次元(あるいは無限次元)特徴量も(有限の)効率的な時間で扱える
 - 過学習が心配？→正則化すればよい
- もし望む計算が特徴量 $\phi(x)$ を經由せず、カーネル $k(x, x')$ だけで達成できるなら、特徴量次元数がどんなに高くても困らない!
- そんなことができるか...?



SVMカーネル版

- 特徴量ベクトルによるSVM



目的関数 $\min_{\mathbf{w}} \|\mathbf{w}\|_2^2 + C \sum_{i=1}^N \max\{0, 1 - t_i \mathbf{w}^T \phi(\mathbf{x}_i)\}$

分類器 $t_i = \text{sgn}[\mathbf{w}^T \phi(\mathbf{x}_i)]$

- カーネルで表現したSVM

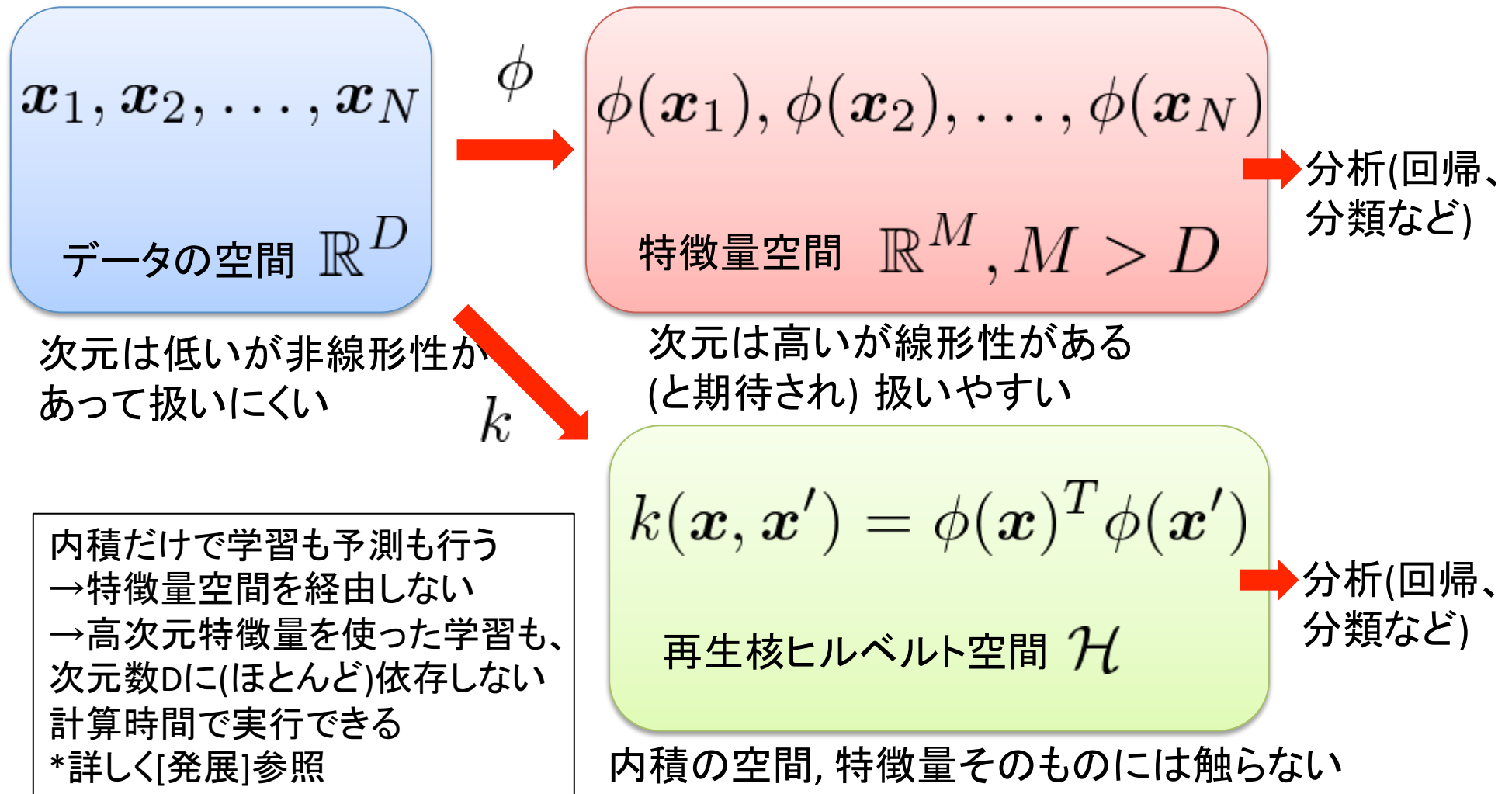
特徴量にアクセスせず、内積
(カーネル)だけで目的関数、
分類器を表現

目的関数 $\max_{\alpha_1, \dots, \alpha_N} \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i,j \in [N]} \alpha_i \alpha_j t_i t_j k(\mathbf{x}_i, \mathbf{x}_j) \text{ sub. to } 0 \leq \alpha_i \leq C, \sum_{i=1}^N y_i \alpha_i = 0$

分類器 $t_i = \text{sgn} \left[\sum_{i: \alpha_i \neq 0} \alpha_i k(\mathbf{x}, \mathbf{x}_i) \right]$



カーネルの考え方



演習 様々なカーネル



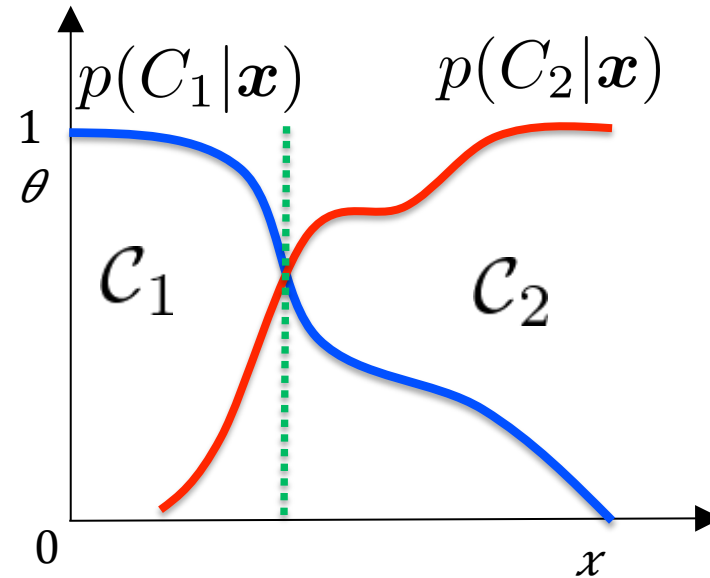
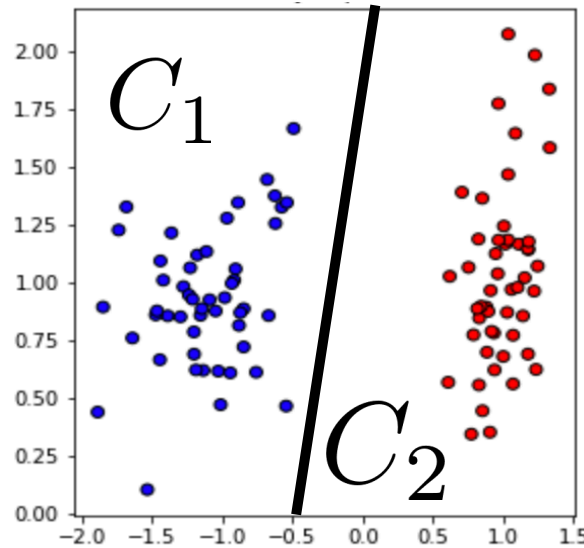
6.2 カーネルと SVM

svm_with_kernels.ipynb を実行して以下の問いに答えなさい。

1. "We first create the data" では、プログラムは3種類のデータを生成します。線形モデルでの分類性能が、最も良いと考えられるデータセットはどれか。最も悪いと考えられるデータセットはどれか。
2. "Linear kernel" では、特徴量やカーネルを用いない線形モデルの SVM を学習する。ここで、パラメータ C は、SVM の定義で現れたパラメータ C である。この C を大きくすると (あるいは小さくすると) どのような挙動が期待できるか？実際にパラメータを変化させたとき、何が起こったか考察しなさい。
3. "Gaussian kernel" ではガウシアンカーネル (円形基底カーネル, RBF カーネル) における3つのデータの分類性能と決定境界を表示する。パラメータ s は、スライドにおけるガウシアンカーネルのパラメータ σ^2 に相当する。 s を変化させた時に、分類器の挙動はどのように変化するか 考察しなさい。なお、このプログラムではパラメータ C は $[1, 30000]$ の区間で、1000 刻みで最適な値 (最もテスト誤差を小さくする値) が選ばれている。



2つのアプローチ



- 決定的識別モデル
 - 識別面 $f(x)$ を学習 $f(x) > 0 \Rightarrow x$ のラベルは C_1
- 確率的識別モデル
 - 条件付き確率 $p(C_1|x)$ を学習
 - $p(C_1|x) > 0.5$ ならば C_1 のラベルは



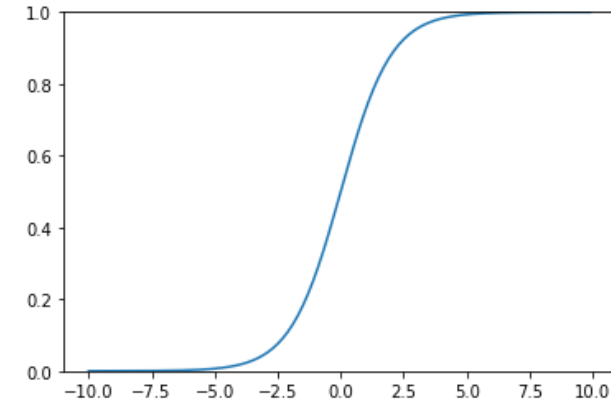
確率的識別モデルの考え方

- 決定的識別モデル 注意: ラベルは(-1,1)
 - 最小マージンを最大にする識別平面(超平面)
 - SVM=ヒンジ損失+L2正則化
- 確率的識別モデル 注意: ラベルは(0,1)
 - ラベル{0,1}を直接予測する代わりに, 分類結果=1である確率 $P(y=1|x) = f(x)$ を予測
- 懸念: 連続値の予測なので回帰が使えるが...
 - 回帰 $f(x) = w^T x$ $[-\infty, \infty]$ の値を返す $0 \leq f(x) \leq 1$
 - $f(x)$ は確率なので $[0, 1]$ に入っている必要がある



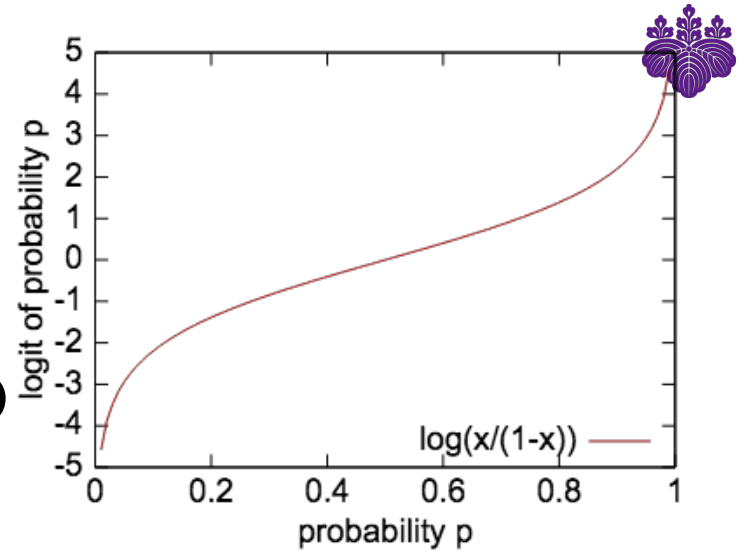
ロジスティックシグモイド関数

- 関数形 $\sigma(a) = \frac{1}{1 + \exp(-a)}$
 - 全実数を[0,1]に押し込む
 - [-3,3]付近は感度が高い(線形関数に近い)
 - 絶対値が大きな値に対しては感度が低い
 - 神経細胞の挙動を模倣している側面がある
- 性質
 1. $\sigma(-a) = 1 - \sigma(a)$ $y=1/2$ に対して対称
 2. $\frac{d}{da}\sigma(a) = \sigma(a)(1 - \sigma(a))$ 符号はどうなっている？



ロジスティック回帰

- x のラベルが $t=1$ である確率を x の式でモデル化したい
- 確率の線形回帰によるモデル化
 - $P(t = 1|x) = f(x) = w^T x$
 - 確率の値が無限大に発散しうる
- 確率 p のロジット: $\text{logit}(p) = \log \frac{p}{1-p}$ 確率を実数値 $\mathbb{R}$ に写像
 $\text{logit} : [0, 1] \mapsto [-\infty, \infty]$
- ロジスティック回帰:
 - $p_x = \text{Pr}(t = 1|x)$ 書く(簡単のため)
 - 確率の代わりにロジットを線形回帰でモデリング
$$\text{logit}(p_x) = w^T x \iff p_x = \sigma(w^T x)$$





演習 ロジスティックシグモイド関数

6.3 ロジスティックシグモイド関数

ロジスティックシグモイド関数 $\sigma(a) = \frac{1}{1+\exp(-a)}$ について以下の問いに答えよ。

1. $\sigma(-a) = 1 - \sigma(a)$ を示せ.
2. $\frac{d}{da}\sigma(a) = \sigma(a)(1 - \sigma(a))$ を示せ.
3. $a \in \mathbb{R}$ のとき, $\sigma(a) \in [0, 1]$ を示せ.
4. $\log \left(\frac{P(t=1|\mathbf{x})}{1-P(t=1|\mathbf{x})} \right) = \mathbf{w}^T \mathbf{x}$ は $P(t = 1|\mathbf{x}) = \sigma(\mathbf{w}^T \mathbf{x})$ と等価であることを示せ.



6.4 ロジスティック回帰による予測

LR_prediction.ipynb を実行して以下の問いに答えなさい。プログラムは 10 点の二次元のラベル付きサンプル (ラベルは 0 or 1) を生成します。生成された特徴とラベルは以下の通り

	1	2	3	4	5	6	7	8	9	10
特徴 x_1	1.5	-0.5	1.0	1.5	0.5	1.5	-0.5	1.0	0.0	0.0
特徴 x_2	-0.5	-1.0	-2.5	-1.0	0.0	-2.0	-0.5	-1.0	-1.0	0.5
ラベル t	1	0	0	1	1	1	0	1	0	0

このデータを超平面 $\mathbf{w}^T \mathbf{x} = 0$ に基づくロジスティック回帰 $y = \sigma(\mathbf{w}^T \mathbf{x})$ で分類することを考えます。ここでは三つの超平面 $\mathbf{w}_1^T = (6, 3, -2)$, $\mathbf{w}_2^T = (4.6, 1, -2.2)$, $\mathbf{w}_3^T = (1, -1, -2)$ を考えます。エクセルやプログラムなどを用いて、以下の問いに答えなさい。この課題については、結果は manaba からダウンロードしたエクセルシートに結果を記入して提出すること。

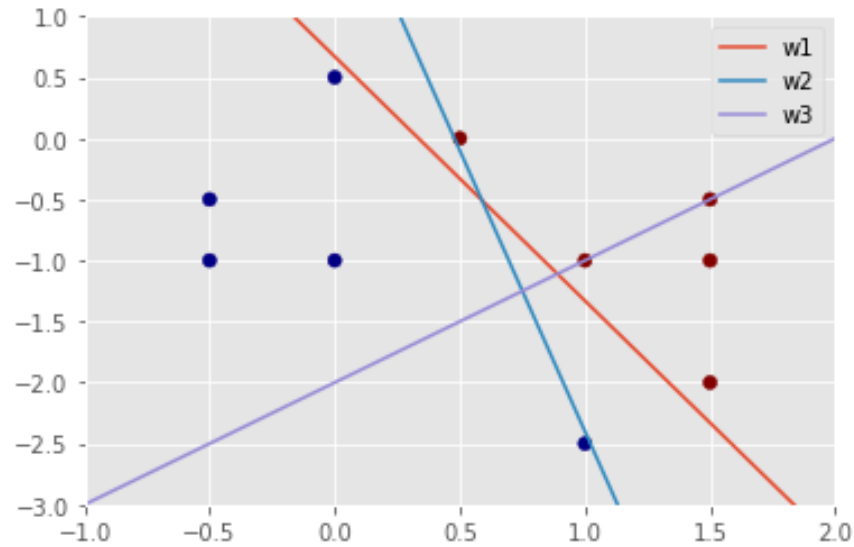
(参考) プログラムは三つのモデルに基づくロジスティック回帰で予測された $P(t=1|\mathbf{x})$ の等高線を表示しています。回答の参考にしてください。

1. $\mathbf{w}_i^T \mathbf{x}_j$ ($i = 1, 2, 3, j = 1, \dots, 10$) を求めよ
2. $\sigma(\mathbf{w}_i^T \mathbf{x}_j)$ ($i = 1, 2, 3, j = 1, \dots, 10$) を求めよ
3. モデル $\sigma(\mathbf{w}_i^T \mathbf{x})$ による $\mathbf{x}_j$ の分類結果 (確率ではなく分類結果) を求めよ
4. $\mathbf{w}_1, \mathbf{w}_2, \mathbf{w}_3$ で実現されるロジスティック回帰モデルの、観測値 $(x_1, \dots, x_{10})$ に対する正解率を求めよ

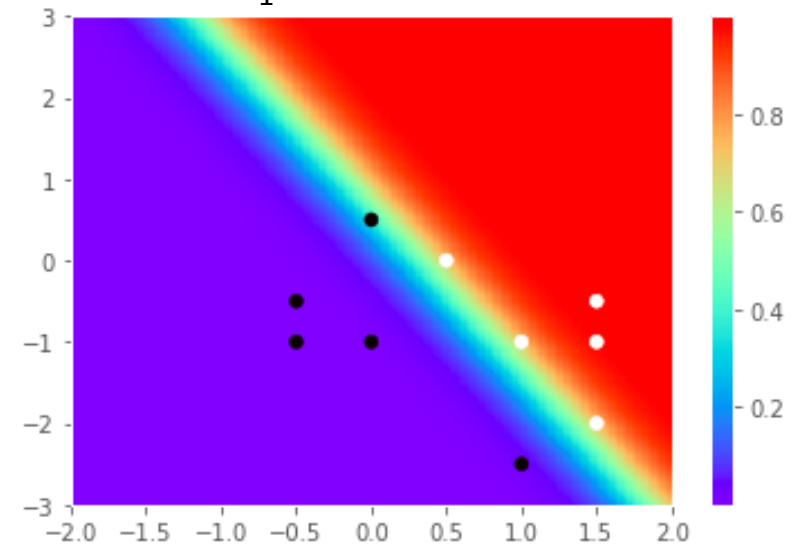


演習6.4

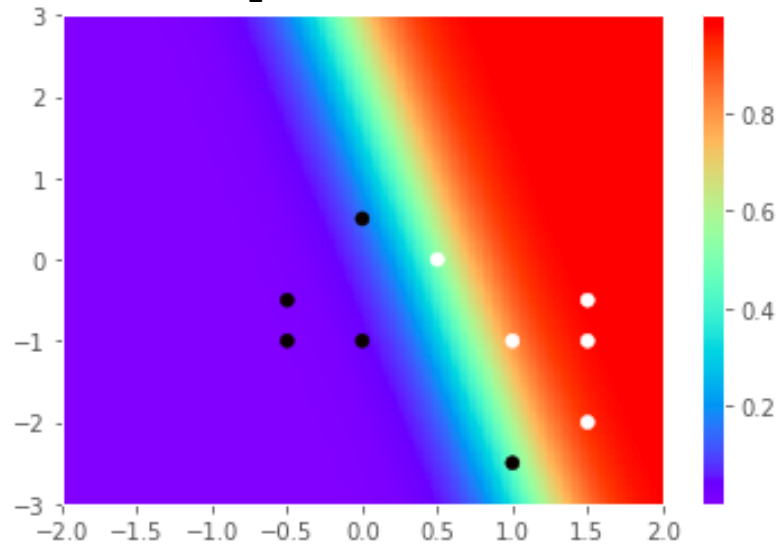
サンプルの分布



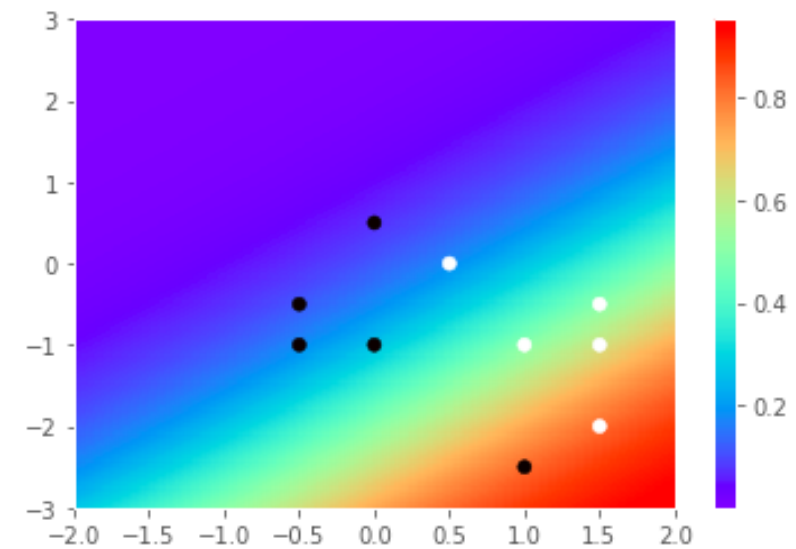
モデル w_1 におけるロジット等高線



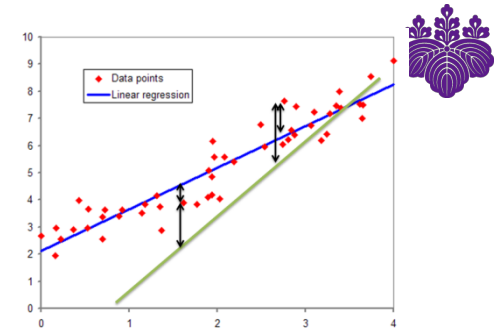
モデル w_2 におけるロジット等高線



モデル w_3 におけるロジット等高線



ロジスティック回帰の損失



- 回帰の場合は損失として二乗誤差を考えた

正解ラベル

$$t_1 = 2.1, t_2 = 4.6, \dots, t_N = 32.7$$

予測

$$w^T x_1 = 2.33, w^T x_2 = 4.51, \dots, \\ w^T x_N = 33.3$$

$$\text{回帰 } t = f(x) = w^T x$$

$$\text{損失 } E(w) = \sum_{i=1}^N (t_i - w^T x_i)^2$$

- ロジスティック回帰は損失として何を考えるか?

正解ラベル

$$t_1 = 0, t_2 = 1, t_3 = 1, \dots, t_N = 0$$

予測

$$\sigma(w^T x_1) = 0.12, \sigma(w^T x_2) = 0.86, \dots, \\ \sigma(w^T x_N) = 0.56$$

$$\text{分類 } P(C_1|x) = \sigma(w^T x)$$

損失 ???



確率分布(ベルヌーイ分布)

確率 μ で表($x=1$)が出て, 確率 $1-\mu$ で裏($x=0$)が出るコインを考える

確率変数 パラメータ

$$B(x; \mu) = \mu^x (1 - \mu)^{1-x} \quad x \in \{0, 1\}$$

- 分布 B はパラメータ μ でパラメトライズされている
- 真のパラメータは神のみぞ知る(コインに固有の値)
- 人間が観測できるのは、有限回のコインの出目
- コインの出目から、真のパラメータ μ を知るには？



尤度と最尤推定

- 尤度(likelihood)
 - なんらかの前提条件に従って観測値が出現する場合に、逆に観測値から前提条件が「何々であった」と推測する尤もらしさ(もっともらしさ)を表す数値
 - E.g. 前提: 確率0.3で表が出るコイン、観測値: その出目
- 尤度関数 $L(\theta) = \prod_{i=1}^N p(x_i; \theta)$ ベルヌーイ分布の場合 $L(\mu) = \prod_{i=1}^N \mu^{x_i} (1 - \mu)^{1-x_i}$
 - 観測データ $(x_1, \dots, x_N)$ の前提条件(パラメータ) θ でパラメトライズされた確率分布 に対する尤度
- 最尤推定 $\max_{\theta} L(\theta)$
 - 観測データに対して、最大の尤度関数値
- 最尤推定量 $\theta^* = \arg \max_{\theta} L(\theta)$
 - 尤度を最大にするパラメータ



最尤推定

- 尤度関数を最大にする
パラメータ

→ **最尤推定量**

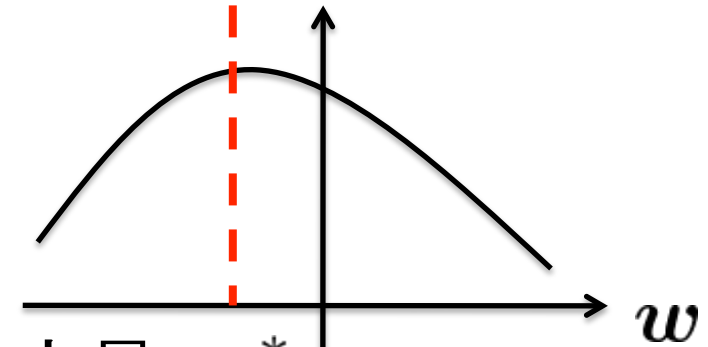
最も尤もらしい推定量

- $f(x)$ が凸関数のとき
 $\log f(x)$
も凸関数

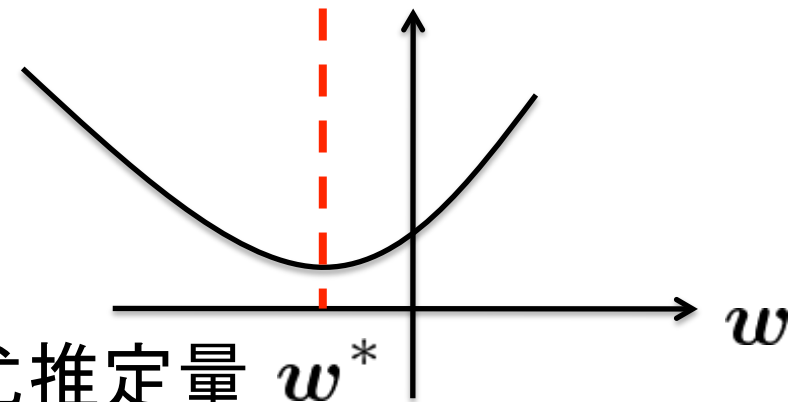
- 尤度最大化よりも**負の対数尤度最小化**のほうが

ベルヌーイ分布の尤度 $L(\mu) = \prod_{i=1}^N \mu^{x_i} (1 - \mu)^{1-x_i}$ 計算に便利なのが多い

ベルヌーイ分布の対数尤度 $-\log L(\mu) = -\sum_{i=1}^N x_i \log \mu + (1 - x_i) \log(1 - \mu)$



最尤推定量 w^*



最尤推定量 w^*



演習 ベルヌーイ分布の尤度

6.5 ベルヌーイ分布の尤度

A 君がコインを 7 回投げたところ、その観測値として (表, 裏, 表, 裏, 裏, 裏, 裏) を得た。以下の問いに答えなさい。

1. このコインの出目を $\mu = 0.25$ および $\mu = 0.1$ のベルヌーイ分布としたときの尤度をそれぞれ求めよ。ただしベルヌーイ分布では表に対応する確率変数の値を 1, 裏に対応する確率変数の値を 0 とする。
2. $\mu = 0.25$ のベルヌーイ分布に従うコインと $\mu = 0.1$ のベルヌーイ分布に従うコイン、この観測値においてはどちらがより尤もらしいか？
3. この観測値に対するパラメータ μ におけるベルヌーイ分布の対数尤度を μ の式で表し、最尤推定量を求めよ



ロジスティック回帰の損失の考え方

- 例えば以下のような状況を考える:

	$(x_1, t_1=1)$	$(x_2, t_2=0)$	$(x_3, t_3=0)$	$(x_4, t_4=1)$
モデル1による予測確率 $\sigma(w_1^T x_i)$	0.9	0.1	0.1	0.9
モデル2による予測確率 $\sigma(w_2^T x_i)$	0.8	0.2	0.2	0.8
実現値(訓練データのラベル) t_i	1	0	0	1

- w_1 と w_2 , どちらがよりよいモデルと言えるか?
- モデルを w でパラメトライズされたベルヌーイ分布、訓練データ中のラベルを観測値とする
- 尤度を最大にする w を尤も「良い」予測を与えるモデルと考える



ロジスティック回帰の損失

- 損失関数(尤度最大化の枠組みで)
 - ロジスティック回帰が出力する予測 $\hat{y}_n = \sigma(\mathbf{w}^T \mathbf{x}_n)$
 - 実現値 t_n
 - ロジスティック回帰の予測をベルヌーイ分布とみなした時の実現値(訓練データ)の尤度

$$L(\mathbf{w}) = \prod_{n=1}^N \hat{y}_n^{t_n} (1 - \hat{y}_n)^{(1-t_n)}$$

$L(\mathbf{w})$ はロジスティック回帰 $\hat{y}_n = \sigma(\mathbf{w}^T \mathbf{x}_n)$ が
訓練データに対してどの程度尤もらしいかを評価



ロジスティック回帰の尤度関数

$$L(\boldsymbol{w}) = \prod_{n=1}^N \underbrace{\hat{t}_n^{t_n}}_{\text{予測が1に近い値}} \underbrace{(1 - \hat{t}_n)^{(1-t_n)}}_{\text{予測が0に近い値}}$$

	$\hat{t}_n^{t_n}$	$(1 - \hat{t}_n)^{(1-t_n)}$	$\hat{t}_n^{t_n}(1 - \hat{t}_n)^{(1-t_n)}$
$t_n = 1, \hat{t}_n \simeq 1$	(1に近い値) ¹	(0に近い値) ⁰	1に近い値
$t_n = 1, \hat{t}_n \simeq 0$	(0に近い値) ¹	(1に近い値) ⁰	0に近い値
$t_n = 0, \hat{t}_n \simeq 0$	(0に近い値) ⁰	(1に近い値) ¹	1に近い値
$t_n = 0, \hat{t}_n \simeq 1$	(1に近い値) ⁰	(0に近い値) ¹	0に近い値

予測が当たってるほうが尤度は確かに大きくなる

交差エントロピー誤差関数の最小化



- 尤度関数最大化の代わりに負の対数尤度関数(交差エントロピー損失と呼ばれる)を最小化

$$E(\mathbf{w}) = -\ln L(\mathbf{w}) = -\sum_{n=1}^N \{t_n \ln \hat{t}_n + (1 - t_n) \ln(1 - \hat{t}_n)\}$$

- 凸関数(証明は略)

- 微分できれば最小化できる
- でもシグモイド関数が挟まっているので解析的な解法(微分をゼロにする $\mathbf{w}$ を求める)ことはできない
- どうやって最小化する?

$$\hat{t}_n = \sigma(\mathbf{w}^T \mathbf{x}_n)$$

演習 ロジスティック回帰と対数尤度



6.6 ロジスティック回帰と対数尤度

LR_prediction.ipynb を実行して以下の問いに答えなさい。用いるデータやモデルは 6.4 と同じである。

1. 訓練サンプル $(x_1, \dots, x_{10})$ について交差エントロピー損失 $E(w_1), E(w_2), E(w_3)$ を求めなさい
2. 最も尤度の高いモデルはどれか
3. w_1, w_2, w_3 どれが分類モデルとして適切か, 理由とともに答えよ



まとめ

- 確率的識別モデルでは、データがクラス1に属する確率をシグモイドロジスティック関数 $P(C_1|x)=\sigma(w^Tx)$ を用いてモデル化する
- 誤差最小化の代わりに対数尤度最大化を行う



ROADMAP

