

The background features two wireframe hands, one in the upper right and one in the lower left, reaching towards the center. The background is a solid blue color with faint, concentric circular patterns and small star-like sparkles.

第4讲 线性模型

周文晖

杭州电子科技大学



线性模型基本形式

线性方程, 基本形式, ...



线性回归

参数估计, 最小二乘, 对数几率回归, ...



线性判别分析

基本概念, 广义线性判别, 几何性质 ...



经典线性判别准则

Fisher判别, 感知器准则、梯度法 ...



非线性判别函数

分段线性, ...



线性模型基本形式

线性方程, 基本形式, ...



线性回归

参数估计, 最小二乘, 对数几率回归, ...



线性判别分析

基本概念, 广义线性判别, 几何性质 ...



经典线性判别准则

Fisher判别, 感知器准则、梯度法 ...



非线性判别函数

分段线性, ...

统计决策方法的局限

$$P(\omega_i | X) = \frac{P(\omega_i)P(X | \omega_i)}{P(X)}$$

统计决策方法思路：通过贝叶斯公式将后验估计问题转为类先验概率和类条件概率计算；
将诊断推理转为因果推理。

局限：

- 必须精确知道类先验概率和类条件概率
- 类条件概率估计
 - 估计类条件概率密度函数：参数估计和非参数估计
 - 直接估计类条件概率，利用条件独立性，采用朴素贝叶斯方法
- 挑战：类条件概率密度函数难以估计；类条件概率计算（尤其是联合概率计算）困难。

线性模型基本形式——数值表示

线性模型的一般形式

$$f(\mathbf{x}) = w_1x_1 + w_2x_2 + \dots + w_dx_d + b$$

$\mathbf{x} = (x_1; x_2; \dots; x_d)$ 是由 d 个属性（特征）描述的一个示例（样本）。

一个样本表示为一个 d 维特征空间上的点，其中 x_i 是 $\mathbf{x}$ 在第 i 个属性（第 i 维）上的取值。

线性模型的向量形式

$$f(\mathbf{x}) = \mathbf{w}^T \mathbf{x} + b \quad \text{其中 } \mathbf{w} = (w_1; w_2; \dots; w_d)$$

模型的参数由权重向量 $\mathbf{w}$ 与偏置项 b 构成，通过训练得到。

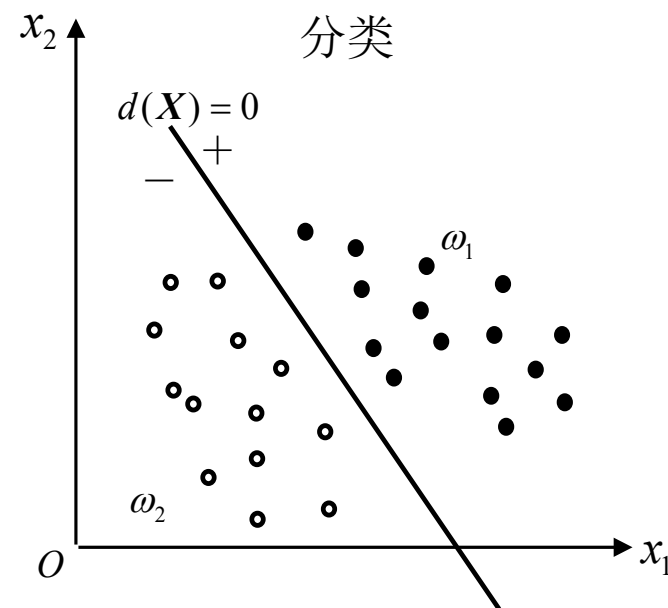
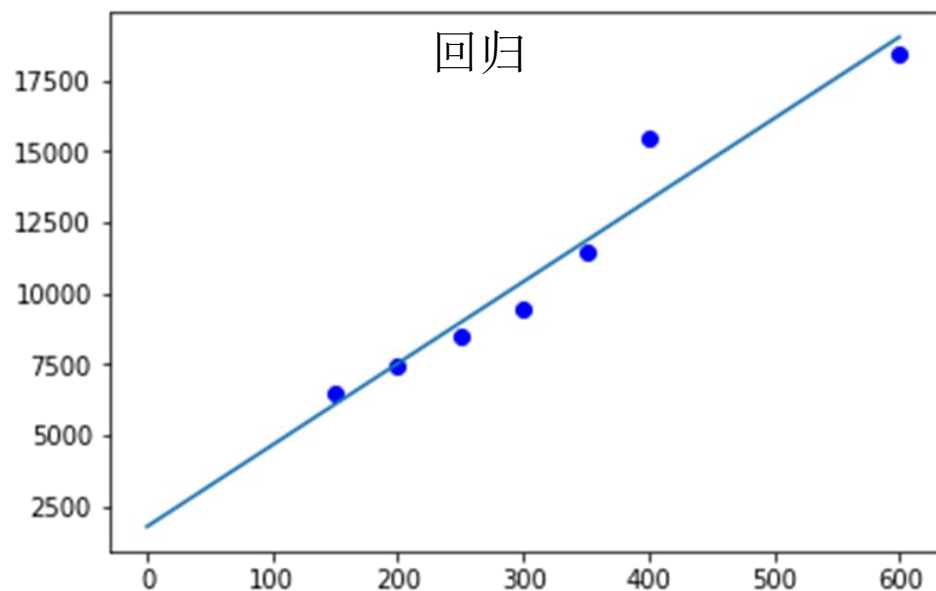
线性模型基本形式——几何表示

线性模型的几何意义： d 维特征空间中的一个超平面 $f(\mathbf{x}) = w_1x_1 + w_2x_2 + \dots + w_dx_d + b$

一维特征空间：线性模型表示为一条直线；

二维特征空间：线性模型表示为一个平面

线性模型可以用来解决模式识别的两个基本问题：



线性模型基本形式——增广表示

线性模型的增广向量形式

将权重向量和特征向量进行增广合并，以简化表达，即 $[\mathbf{w}, b] \rightarrow \mathbf{w}$, $[\mathbf{x}, 1] \rightarrow \mathbf{x}$

$$f(\mathbf{x}) = w_1x_1 + w_2x_2 + \dots + w_nx_n + b \cdot 1 = [w_1, w_2, \dots, w_n, b] \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \\ 1 \end{bmatrix} = \mathbf{w}^T \mathbf{x}$$

式中： $\mathbf{x} = [x_1, x_2, \dots, x_n, 1]^T$ 为增广模式向量。

$\mathbf{w} = [w_1, w_2, \dots, w_n, w_{n+1}]^T$ 为增广权向量。

线性模型优点

优点：◇ 形式简单、易于建模

◇ 可解释性

◇ 非线性模型的基础

- 许多非线性模型可在线性模型基础上引入层级结构或高维映射而得到。

一个例子：

- 综合考虑色泽、根蒂和敲声来判断西瓜好不好
- 其中根蒂的系数最大，表明根蒂最要紧；而敲声的系数比色泽大，说明敲声比色泽更重要

$$f_{\text{好瓜}}(\mathbf{x}) = 0.2 \cdot x_{\text{色泽}} + 0.5 \cdot x_{\text{根蒂}} + 0.3 \cdot x_{\text{敲声}} + 1$$



线性模型基本形式

线性方程, 基本形式, ...



线性回归

参数估计, 最小二乘, 对数几率回归, ...



线性判别分析

基本概念, 广义线性判别, 几何性质 ...



经典线性判别准则

Fisher判别, 感知器准则、梯度法 ...



非线性判别函数

分段线性, ...

线性回归

给定数据集 $D = \{(\mathbf{x}_1, y_1), (\mathbf{x}_2, y_2), \dots, (\mathbf{x}_m, y_m)\}$

其中 $\mathbf{x}_i = (x_{i1}; x_{i2}; \dots; x_{id}) \quad y_i \in \mathbb{R}$

线性回归（linear regression）问题：

◇ 学得一个线性模型以尽可能准确地预测实值输出标记。

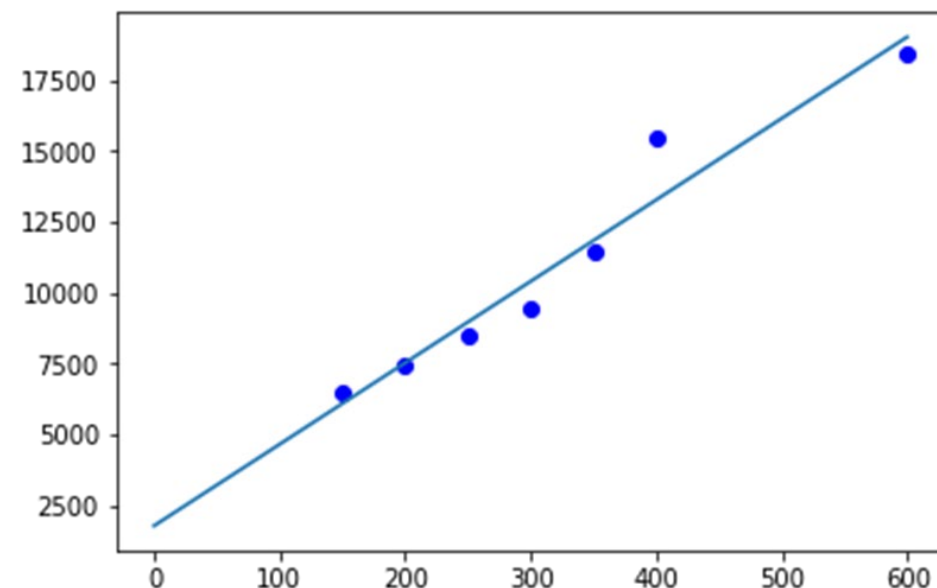
例，设定一元线性回归的线性预测函数为

$$f(\mathbf{x}) = w \cdot \mathbf{x} + b$$

通过学习到 w 和 b ，使得 $f(x_i) \simeq y_i$

寻找 w 和 b 的最佳值为线性回归模型的参数估计过程。

面积	价格
150	6450
200	7450
250	8450
300	9450
350	11450
400	15450
600	18450



线性回归——最小二乘法

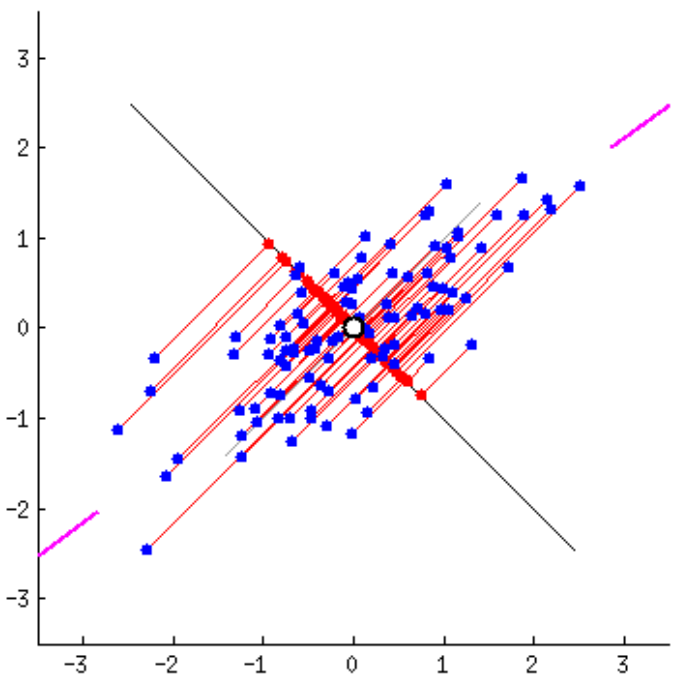
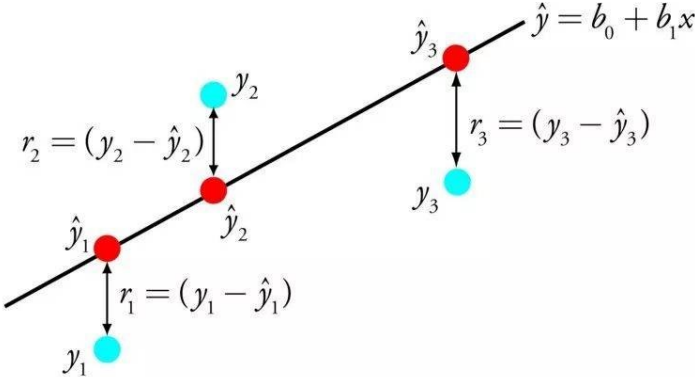
为确定 w 和 b ，关键是使得 $f(\mathbf{x})$ 与 $\mathbf{y}$ 之间的差距最小。

均方误差是回归任务中最常用的性能度量，线性回归问题转化为寻找使得均方误差最小的 w 和 b ：

$$\begin{aligned}(w^*, b^*) &= \arg \min_{(w, b)} \sum_{i=1}^m (f(x_i) - y_i)^2 \\ &= \arg \min_{(w, b)} \sum_{i=1}^m (y_i - wx_i - b)^2\end{aligned}$$

采用均方误差（MSE）的意义：

- (1) 均方误差的几何意义：欧几里得距离，或简称欧氏距离（Euclidean distance）
- (2) 基于均方误差最小化的模型参数求解过程称为“最小二乘法”（Least square method）



线性回归——最小二乘法参数估计

求解 w 和 b 的最佳值，使得均方误差 $E_{(w,b)} = \sum_{i=1}^m (y_i - wx_i - b)^2$ 最小的过程，称为线性回归模型的最小二乘“参数估计”。

对于一元线性回归问题，其求解思路：(1) 将 $E_{(w,b)}$ 分别对 w 和 b 求偏导；

(2) 令偏导为零，得到最优解的闭式（closed-form）解。

对 w 求偏导：

$$\begin{aligned}\frac{\partial E_{(w,b)}}{\partial w} &= \frac{\partial}{\partial w} \left[\sum_{i=1}^m (y_i - wx_i - b)^2 \right] = \sum_{i=1}^m \frac{\partial}{\partial w} [(y_i - wx_i - b)^2] = \sum_{i=1}^m [2 \cdot (y_i - wx_i - b) \cdot (-x_i)] \\ &= \sum_{i=1}^m [2 \cdot (wx_i^2 - y_i x_i + bx_i)] = 2 \cdot \left(w \sum_{i=1}^m x_i^2 - \sum_{i=1}^m y_i x_i + b \sum_{i=1}^m x_i \right) = 2 \left(w \sum_{i=1}^m x_i^2 - \sum_{i=1}^m (y_i - b) x_i \right)\end{aligned}$$

线性回归——最小二乘法参数估计

对 b 求偏导：

$$\begin{aligned}\frac{\partial E(w,b)}{\partial b} &= \frac{\partial}{\partial b} \left[\sum_{i=1}^m (y_i - wx_i - b)^2 \right] = \sum_{i=1}^m \frac{\partial}{\partial b} [(y_i - wx_i - b)^2] = \sum_{i=1}^m [2 \cdot (y_i - wx_i - b) \cdot (-1)] \\ &= \sum_{i=1}^m [2 \cdot (b - y_i + wx_i)] = 2 \cdot \left[\sum_{i=1}^m b - \sum_{i=1}^m y_i + \sum_{i=1}^m wx_i \right] = 2 \left(mb - \sum_{i=1}^m (y_i - wx_i) \right)\end{aligned}$$

令 $\frac{\partial E(w,b)}{\partial b}$ 为零，求解 b 的最优解

$$\frac{\partial E(w,b)}{\partial b} = 2 \left(mb - \sum_{i=1}^m (y_i - wx_i) \right) = 0 \quad \Rightarrow \quad b = \frac{1}{m} \sum_{i=1}^m (y_i - wx_i)$$

令 $\frac{1}{m} \sum_{i=1}^m y_i = \bar{y}$, $\frac{1}{m} \sum_{i=1}^m x_i = \bar{x}$, 则 $b = \bar{y} - w\bar{x}$.

线性回归——最小二乘法参数估计

令 $\frac{\partial E(w,b)}{\partial w}$ 为零，求解 w 的最优解

$$\frac{\partial E(w,b)}{\partial w} = 2 \left(w \sum_{i=1}^m x_i^2 - \sum_{i=1}^m (y_i - b) x_i \right) = 0 \quad \Rightarrow \quad w \sum_{i=1}^m x_i^2 = \sum_{i=1}^m y_i x_i - \sum_{i=1}^m b x_i \quad \text{将 } b = \bar{y} - w\bar{x} \text{ 代入左式,}$$

$$\Rightarrow w \sum_{i=1}^m x_i^2 = \sum_{i=1}^m y_i x_i - \sum_{i=1}^m (\bar{y} - w\bar{x}) x_i \quad \Rightarrow \quad w \sum_{i=1}^m x_i^2 = \sum_{i=1}^m y_i x_i - \bar{y} \sum_{i=1}^m x_i + w\bar{x} \sum_{i=1}^m x_i$$

$$\Rightarrow w \left(\sum_{i=1}^m x_i^2 - \bar{x} \sum_{i=1}^m x_i \right) = \sum_{i=1}^m y_i x_i - \bar{y} \sum_{i=1}^m x_i \quad \Rightarrow \quad w = \frac{\sum_{i=1}^m y_i x_i - \bar{y} \sum_{i=1}^m x_i}{\sum_{i=1}^m x_i^2 - \bar{x} \sum_{i=1}^m x_i}$$

$$\begin{aligned} \bar{y} \sum_{i=1}^m x_i &= \frac{1}{m} \sum_{i=1}^m y_i \sum_{i=1}^m x_i = \bar{x} \sum_{i=1}^m y_i \\ \bar{x} \sum_{i=1}^m x_i &= \frac{1}{m} \sum_{i=1}^m x_i \sum_{i=1}^m x_i = \frac{1}{m} (\sum_{i=1}^m x_i)^2 \end{aligned} \quad \Rightarrow \quad w = \frac{\sum_{i=1}^m y_i (x_i - \bar{x})}{\sum_{i=1}^m x_i^2 - \frac{1}{m} (\sum_{i=1}^m x_i)^2}$$

多元线性回归

对于更一般情况, $f(\mathbf{x}_i) = \mathbf{w}^T \mathbf{x}_i + b$, 使得 $f(\mathbf{x}_i) \simeq y_i$ 。

类似于一元线性回归中的最小二乘法, 多元线性回归问题转为

$$(\mathbf{w}^*, b^*) = \arg \min_{(\mathbf{w}, b)} \sum_{i=1}^m (f(\mathbf{x}_i) - y_i)^2 = \arg \min_{(\mathbf{w}, b)} \sum_{i=1}^m (y_i - (\mathbf{w}^T \mathbf{x}_i + b))^2$$

为便于讨论, 采用增广向量形式 $\hat{\mathbf{w}} = (\mathbf{w}; b) = (w_1; \dots; w_d; b) \in \mathbb{R}^{(d+1) \times 1}$, $\hat{\mathbf{x}}_i = (x_1; \dots; x_d; 1) \in \mathbb{R}^{(d+1) \times 1}$, 有

$$\hat{\mathbf{w}}^* = \arg \min_{\hat{\mathbf{w}}} \sum_{i=1}^m (y_i - \hat{\mathbf{x}}_i^T \hat{\mathbf{w}})^2$$

根据向量内积的定义可知, 上式可以写成如下向量内积的形式

$$\hat{\mathbf{w}}^* = \arg \min_{\hat{\mathbf{w}}} \begin{bmatrix} y_1 - \hat{\mathbf{x}}_1^T \hat{\mathbf{w}} & \dots & y_m - \hat{\mathbf{x}}_m^T \hat{\mathbf{w}} \end{bmatrix} \begin{bmatrix} y_1 - \hat{\mathbf{x}}_1^T \hat{\mathbf{w}} \\ \vdots \\ y_m - \hat{\mathbf{x}}_m^T \hat{\mathbf{w}} \end{bmatrix} = \arg \min_{\hat{\mathbf{w}}} (\mathbf{y} - \mathbf{X} \hat{\mathbf{w}})^T (\mathbf{y} - \mathbf{X} \hat{\mathbf{w}})$$

$$\begin{aligned} \begin{bmatrix} y_1 - \hat{\mathbf{x}}_1^T \hat{\mathbf{w}} \\ \vdots \\ y_m - \hat{\mathbf{x}}_m^T \hat{\mathbf{w}} \end{bmatrix} &= \begin{bmatrix} y_1 \\ \vdots \\ y_m \end{bmatrix} - \begin{bmatrix} \hat{\mathbf{x}}_1^T \hat{\mathbf{w}} \\ \vdots \\ \hat{\mathbf{x}}_m^T \hat{\mathbf{w}} \end{bmatrix} \\ &= \mathbf{y} - \begin{bmatrix} \hat{\mathbf{x}}_1^T \\ \vdots \\ \hat{\mathbf{x}}_m^T \end{bmatrix} \cdot \hat{\mathbf{w}} \\ &= \mathbf{y} - \mathbf{X} \hat{\mathbf{w}} \end{aligned}$$

多元线性回归——最小二乘法

令 $E_{\hat{\mathbf{w}}} = (\mathbf{y} - \mathbf{X}\hat{\mathbf{w}})^T (\mathbf{y} - \mathbf{X}\hat{\mathbf{w}})$ ，对 $\hat{\mathbf{w}}$ 求导，并令导数为零，得到 $\hat{\mathbf{w}}$ 最优解的闭式解。

将 $E_{\hat{\mathbf{w}}}$ 展开，有 $E_{\hat{\mathbf{w}}} = \mathbf{y}^T \mathbf{y} - \mathbf{y}^T \mathbf{X} \hat{\mathbf{w}} - \hat{\mathbf{w}}^T \mathbf{X}^T \mathbf{y} + \hat{\mathbf{w}}^T \mathbf{X}^T \mathbf{X} \hat{\mathbf{w}}$

对 $\hat{\mathbf{w}}$ 求导 $\frac{\partial E_{\hat{\mathbf{w}}}}{\partial \hat{\mathbf{w}}} = \frac{\partial \mathbf{y}^T \mathbf{y}}{\partial \hat{\mathbf{w}}} - \frac{\partial \mathbf{y}^T \mathbf{X} \hat{\mathbf{w}}}{\partial \hat{\mathbf{w}}} - \frac{\partial \hat{\mathbf{w}}^T \mathbf{X}^T \mathbf{y}}{\partial \hat{\mathbf{w}}} + \frac{\partial \hat{\mathbf{w}}^T \mathbf{X}^T \mathbf{X} \hat{\mathbf{w}}}{\partial \hat{\mathbf{w}}}$

由矩阵微分公式 $\frac{\partial \mathbf{a}^T \mathbf{x}}{\partial \mathbf{x}} = \frac{\partial \mathbf{x}^T \mathbf{a}}{\partial \mathbf{x}} = \mathbf{a}$ ， $\frac{\partial \mathbf{x}^T \mathbf{A} \mathbf{x}}{\partial \mathbf{x}} = (\mathbf{A} + \mathbf{A}^T) \mathbf{x}$ ，可得

$$\frac{\partial E_{\hat{\mathbf{w}}}}{\partial \hat{\mathbf{w}}} = 0 - \mathbf{X}^T \mathbf{y} - \mathbf{X}^T \mathbf{y} + (\mathbf{X}^T \mathbf{X} + \mathbf{X}^T \mathbf{X}) \hat{\mathbf{w}} = 2\mathbf{X}^T (\mathbf{X} \hat{\mathbf{w}} - \mathbf{y})$$

令上式为零可得 $\hat{\mathbf{w}}$ 最优解的闭式解

多元线性回归——最小二乘法

求解 $\hat{\mathbf{w}}$ 最优解的闭式解: $\frac{\partial E_{\hat{\mathbf{w}}}}{\partial \hat{\mathbf{w}}} = 2\mathbf{X}^T (\mathbf{X}\hat{\mathbf{w}} - \mathbf{y}) = 0$

(1) 当 $\mathbf{X}^T\mathbf{X}$ 是满秩矩阵或正定矩阵（即可逆），则有

$$\hat{\mathbf{w}}^* = (\mathbf{X}^T\mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

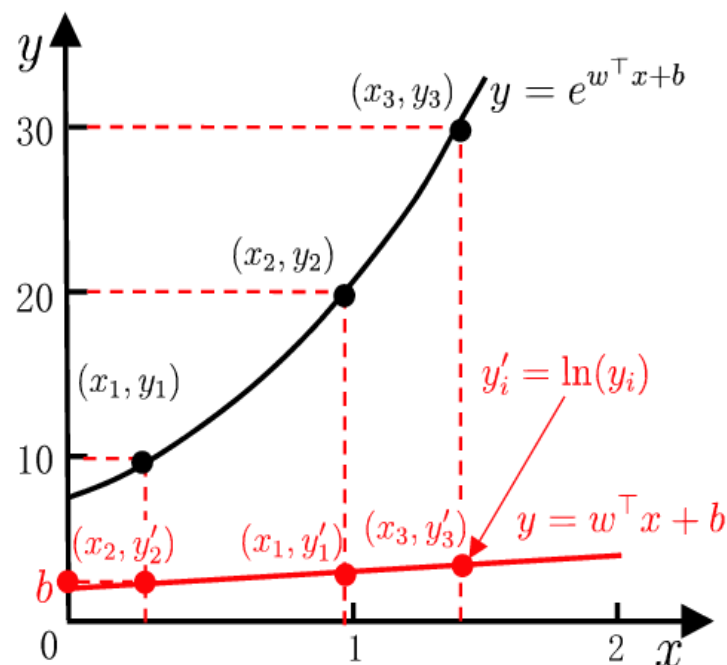
其中 $(\mathbf{X}^T\mathbf{X})^{-1}$ 是 $\mathbf{X}^T\mathbf{X}$ 的逆矩阵，线性回归模型为 $f(\hat{\mathbf{x}}_i) = \hat{\mathbf{x}}_i^T (\mathbf{X}^T\mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$

(2) 若 $\mathbf{X}^T\mathbf{X}$ 不是满秩矩阵，则可能存在多个解，则根据归纳偏好选择合适的解。

线性回归→广义线性回归

从线性回归推广到广义线性回归。

如**对数线性回归模型**（log-Linear Regression）：将输出标记的对数作为线性模型逼近的目标。



$$y = e^{w^T x + b}$$
$$\updownarrow$$
$$\ln y = w^T x + b$$

更一般形式，**广义线性回归模型**（GLM, Generalized Linear Model）

$$y = g^{-1}(w^T x + b)$$

$g(\cdot)$ 称为联系函数（link function），为单调可微函数

广义线性回归意义：

- 1) 非线性模型→线性模型
- 2) 输入空间 $\mathbf{X}$ 到输出空间 $\mathbf{y}$ 的非线性映射。

对数几率回归（逻辑回归、罗杰斯特回归）

几率（odds）：即胜算的意思，来自于统计学领域。

以中国 vs. 巴西足球赛为例，中国队赢面是1，输面是99，则中国队赢的几率是1/99，巴西队赢的几率是99。两队输赢程度是一样的，但 $1/99=0.01010\dots$ 和 99 这两个数的尺度不同，难以做出直观的判断，而**对数几率 $\log(\text{odds})$** 可解决这个问题。

	中国队赢	巴西队赢
几率	1/99	99
对数几率	-4.60	4.60

对 odds 加了 log 后，中国队赢和巴西队赢这两种情况的 $\log(\text{odds})$ 的绝对值都是 4.6，一眼可看出两者的输赢程度相同，并通过正负号可判断赢面多还是赢面少。

当 $\log(\text{odds})$ 为 0 时，赢面和输面一样多。

对数几率回归 (log odds, logit)

在线性回归模型中, 假定样本 $\mathbf{x}$ 和真实 y 呈线性关系 $y = \mathbf{w}^T \mathbf{x} + b$;

在二分类问题中, 考虑 y 的几率形式, 并假定样本 $\mathbf{x}$ 和 y 的对数几率呈线性关系, 即**对数几率回归**。

在对数几率回归中, 模型输出 y 定义为样本为正例的概率, 则对数几率为:

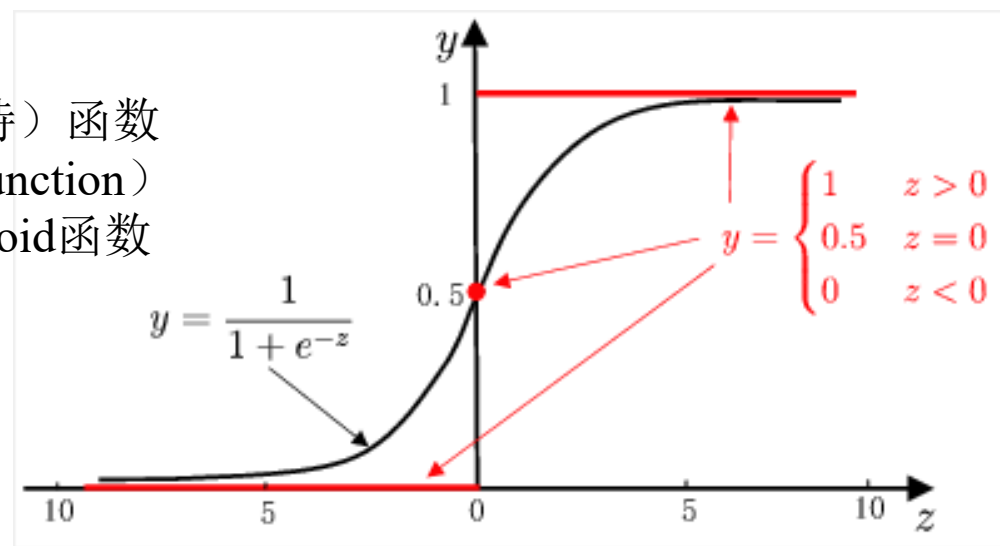
$$\ln \frac{y}{1-y} \quad 1-y \text{ 为样本为负例的概率}$$

因此, 有 $\ln \frac{y}{1-y} = \mathbf{w}^T \mathbf{x} + b$

➡ $\frac{y}{1-y} = e^{\mathbf{w}^T \mathbf{x} + b}$

➡ $y = \frac{e^{\mathbf{w}^T \mathbf{x} + b}}{1 + e^{\mathbf{w}^T \mathbf{x} + b}} \xrightarrow{z = \mathbf{w}^T \mathbf{x} + b} \frac{e^z}{1 + e^z} = \frac{1}{1 + e^{-z}}$

逻辑 (斯特) 函数
(logistic function)
或称 Sigmoid 函数



对数几率回归——对数几率函数

对数几率函数（Sigmoid函数）优点：

(1) 可将 $(-\infty, +\infty)$ 结果，映射到 $(0, 1)$ 间，作为概率；

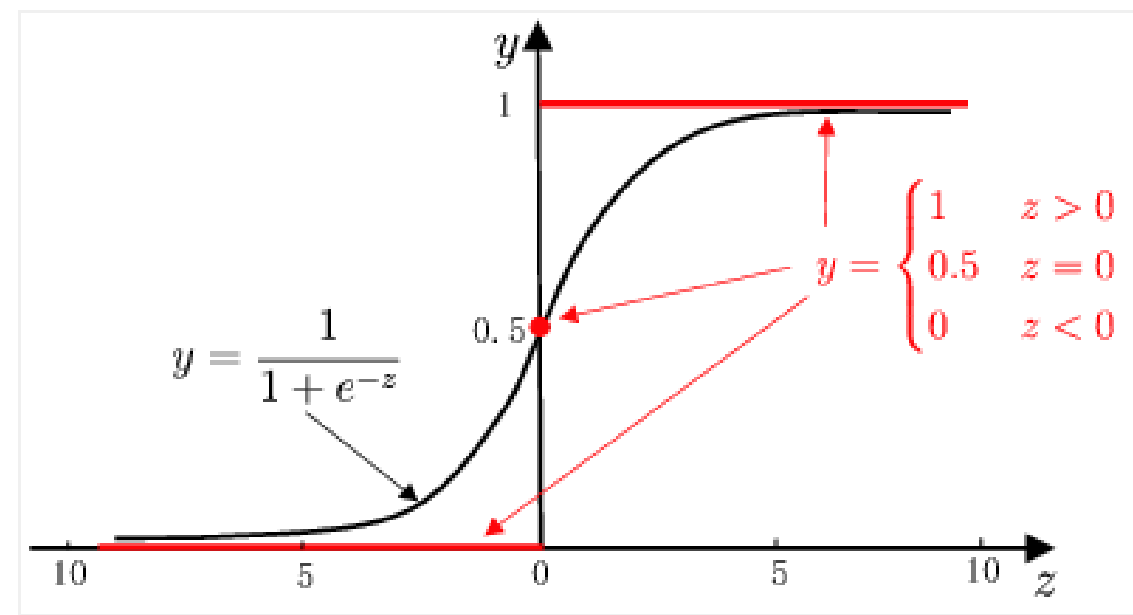
(2) 当 $z < 0$ 时， $\text{Sigmoid}(z) < 0.5$

当 $z > 0$ 时， $\text{Sigmoid}(z) > 0.5$

因此可将 0.5 作为分类的决策边界

(3) 单调可微、任意阶可导

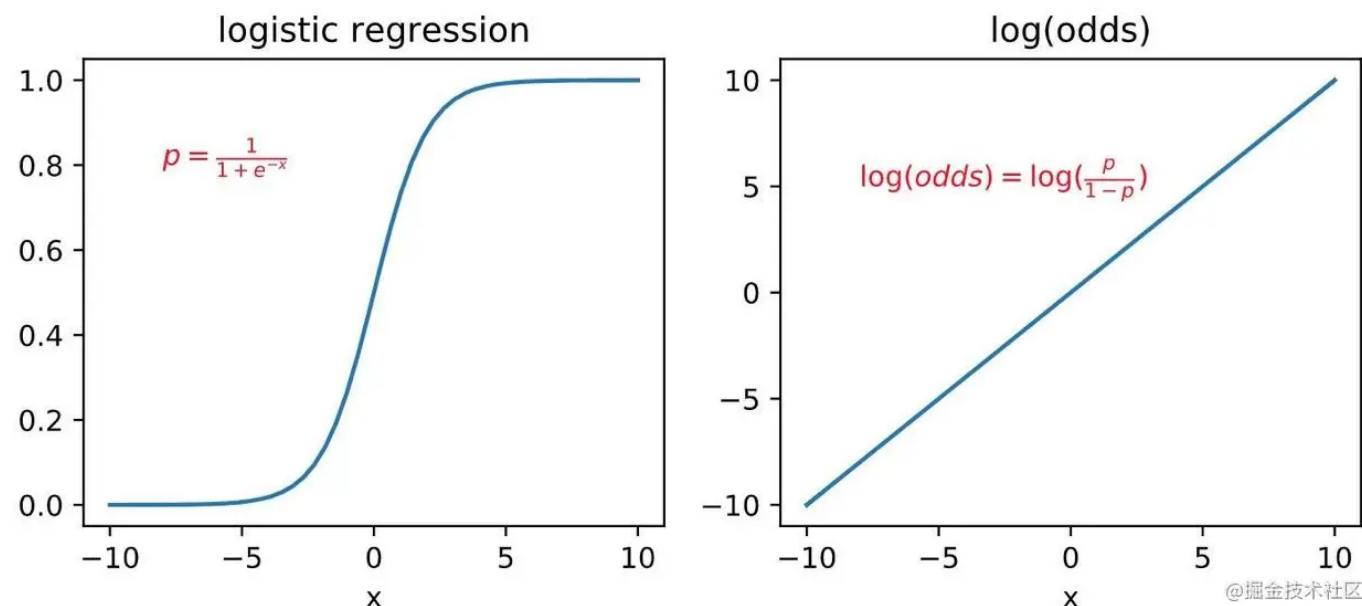
$$y' = \left(\frac{1}{1 + e^{-z}} \right)' = \frac{-e^{-z}}{(1 + e^{-z})^2} = \frac{1}{1 + e^{-z}} \left(1 - \frac{1}{1 + e^{-z}} \right) = y(1 - y)$$



对数几率回归——特点

对数几率回归的特点：

- 1、虽然名字里有“回归”，但实际上是最为经典的分类算法之一；
- 2、直接对分类的可能性建模，无需事先假设数据分布，避免了假设分布不准确带来的问题；
- 3、分类的软判决，可得到的是“类别”的近似概率预测；
- 4、对数几率函数是任意阶可导的凸函数，可直接应用现有数值优化算法求取 w 和 b 的最优解。



Sigmoid函数与logit 函数转换关系；

逻辑（斯特）函数（logistic function）又称为Sigmoid函数；

$\log(odds)$ 又称为 logit 函数；

对数几率回归——最大似然法

求解 $y = \frac{1}{1 + e^{-(\mathbf{w}^T \mathbf{x} + b)}}$ 中 $\mathbf{w}$ 和 b 的最优解，即求解 $\ln \frac{y}{1-y} = \mathbf{w}^T \mathbf{x} + b$ 中 $\mathbf{w}$ 和 b 的最优解

将 y 视为类后验概率 $p(y = 1 | \mathbf{x})$ ，则对数几率为：
$$\ln \frac{p(y = 1 | \mathbf{x})}{p(y = 0 | \mathbf{x})} = \mathbf{w}^T \mathbf{x} + b$$

显然有：
$$p(y = 1 | \mathbf{x}) = \frac{e^{\mathbf{w}^T \mathbf{x} + b}}{1 + e^{\mathbf{w}^T \mathbf{x} + b}} = p_1(\mathbf{x}) \quad p(y = 0 | \mathbf{x}) = \frac{1}{1 + e^{\mathbf{w}^T \mathbf{x} + b}} = p_0(\mathbf{x})$$

为便于讨论，令 $\boldsymbol{\beta} = (\mathbf{w}; b)$ ， $\hat{\mathbf{x}} = (\mathbf{x}; 1)$ ，则 $\mathbf{w}^T \mathbf{x} + b$ 可简写为 $\boldsymbol{\beta}^T \hat{\mathbf{x}}$

采用对数形式：
$$\ln p(y_i = 1 | \mathbf{x}) = \boldsymbol{\beta}^T \hat{\mathbf{x}} - \ln(1 + e^{\boldsymbol{\beta}^T \hat{\mathbf{x}}}) \quad \ln p(y_i = 0 | \mathbf{x}) = -\ln(1 + e^{\boldsymbol{\beta}^T \hat{\mathbf{x}}})$$

考虑到 y 取值只是 0 或 1，则统一的对数形式：
$$\ln p(y_i | \mathbf{x}) = y_i \boldsymbol{\beta}^T \hat{\mathbf{x}} - \ln(1 + e^{\boldsymbol{\beta}^T \hat{\mathbf{x}}})$$

对数几率回归——最大似然法

似然函数 (likelihood function)

$$l(\boldsymbol{\beta}) = p(\mathbf{y} | \mathbf{x}; \boldsymbol{\beta}) = p(y_1 | x_1, y_2 | x_2, \dots, y_m | x_m; \boldsymbol{\beta}) = \prod_{i=1}^m p(y_i | x_i; \boldsymbol{\beta}) \quad \text{当满足i.i.d}$$

对数似然函数

$$\ell(\boldsymbol{\beta}) = \ln \prod_{i=1}^m p(y_i | x_i; \boldsymbol{\beta}) = \sum_{i=1}^m \ln p(y_i | x_i; \boldsymbol{\beta}) = \sum_{i=1}^m y_i \boldsymbol{\beta}^T \hat{\mathbf{x}}_i - \ln(1 + e^{\boldsymbol{\beta}^T \hat{\mathbf{x}}_i})$$

$\boldsymbol{\beta}$ 的最优解应满足似然函数值最大（对数似然函数值也为最大），即有最大似然估计：

最大似然估计：

如果在参数 $\boldsymbol{\beta} = \hat{\boldsymbol{\beta}}$ 下， $\ell(\boldsymbol{\beta})$ 最大，则 $\hat{\boldsymbol{\beta}}$ 应是最佳的参数值，称作最大似然估计量。

为了便于求解，将最大化似然函数转化为最小化问题 $J(\boldsymbol{\beta}) = \frac{1}{m} \sum_{i=1}^m -y_i \boldsymbol{\beta}^T \hat{\mathbf{x}}_i + \ln(1 + e^{\boldsymbol{\beta}^T \hat{\mathbf{x}}_i})$

加负号将最大化转为最小化问题。

对数几率回归——最大似然法求解

求解 β 的最优解，即 $\hat{\beta} = \arg \min_{\beta} J(\beta)$

可以采用经典数值优化算法，如梯度下降法、牛顿法等求解 β 的最优解，

牛顿法求解

牛顿法第 $t+1$ 轮迭代解的更新公式 $\beta^{t+1} = \beta^t - \left(\frac{\partial^2 J(\beta)}{\partial \beta \cdot \partial \beta^T} \right)^{-1} \frac{\partial J(\beta)}{\partial \beta}$

其中关于 β 的一阶、二阶导数分别为：

$$\frac{\partial J(\beta)}{\partial \beta} = -\frac{1}{m} \sum_{i=1}^m \hat{x}_i (y_i - p_1(\hat{x}_i; \beta))$$

$$\frac{\partial^2 J(\beta)}{\partial \beta \cdot \partial \beta^T} = \frac{1}{m} \sum_{i=1}^m \hat{x}_i \hat{x}_i^T p_1(\hat{x}_i; \beta) (1 - p_1(\hat{x}_i; \beta))$$

$$\text{其中 } p_1(\hat{x}_i; \beta) = \frac{e^{\beta^T \hat{x}}}{1 + e^{\beta^T \hat{x}}}$$

对数几率回归——最大似然法求解

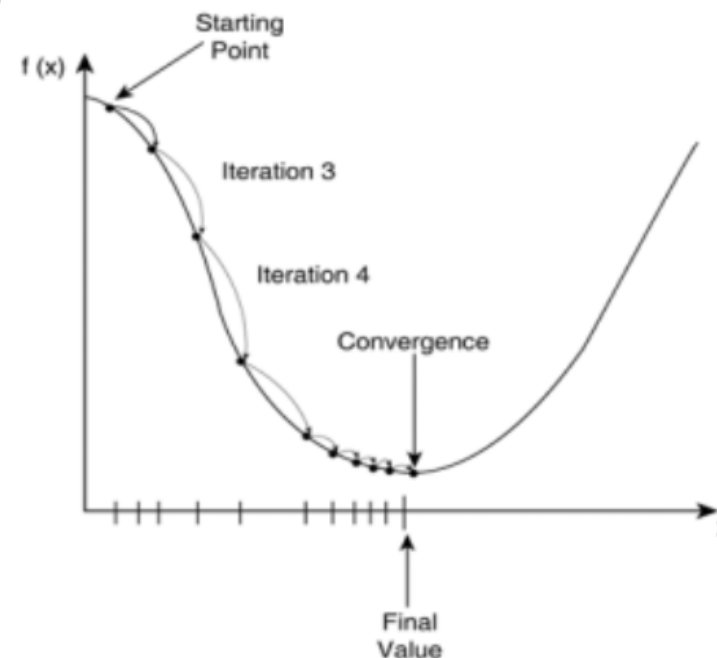
$$\hat{\beta} = \arg \min_{\beta} J(\beta)$$

梯度下降法求解

$$\frac{\partial J(\beta)}{\partial \beta} = -\frac{1}{m} \sum_{i=1}^m \hat{x}_i (y_i - p_1(\hat{x}_i; \beta))$$

将 $\frac{\partial J(\beta)}{\partial \beta}$ 带入参数 β 的更新公式: $\beta^* = \beta - \eta \frac{\partial J}{\partial \beta}$, 得到最终的参数 β 更新公式

$$\beta^* = \beta + \frac{\eta}{m} \sum_{i=1}^m \hat{x}_i (y_i - p_1(\hat{x}_i; \beta)) \quad \eta \text{ 为学习率}$$





- ✓ 线性模型基本形式
线性方程, 基本形式, ...
- ✓ 线性回归
参数估计, 最小二乘, 对数几率回归, ...
- ✓ **线性判别分析**
基本概念, 广义线性判别, 几何性质 ...
- ✓ 经典线性判别准则
Fisher判别, 感知器准则、梯度法 ...
- ✓ 非线性判别函数
分段线性, ...

线性判别分析——基本概念

判别函数 (discriminant function)

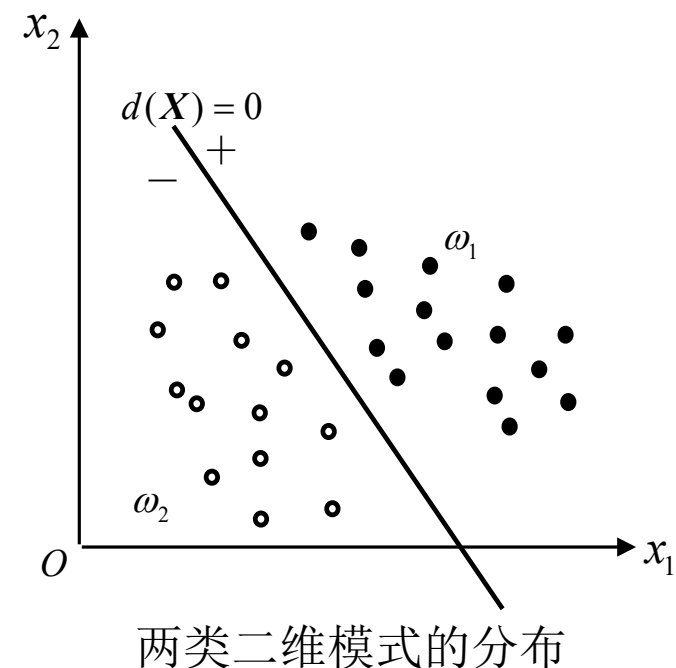
定义：直接用来对模式进行分类的准则函数。

若分属于 ω_1 , ω_2 的两类模式可用一方程 $d(\mathbf{X}) = 0$ 来划分，那么称 $d(\mathbf{X})$ 为判别函数，或称判决函数、决策函数。

例：一个二维的两类判别问题，模式分布如图示，这些分属于 ω_1 , ω_2 两类的模式可用一直线方程 $d(\mathbf{X})=0$ 来划分。

$$d(\mathbf{X}) = w_1x_1 + w_2x_2 + w_3 = 0$$

式中： x_1, x_2 为坐标变量, w_1, w_2, w_3 为方程参数。



判别函数 (discriminant function)

将某一未知模式 $\mathbf{X}$ 代入: $d(\mathbf{X}) = w_1x_1 + w_2x_2 + w_3$

若 $d(\mathbf{X}) > 0$, 则 $\mathbf{X} \in \omega_1$ 类;

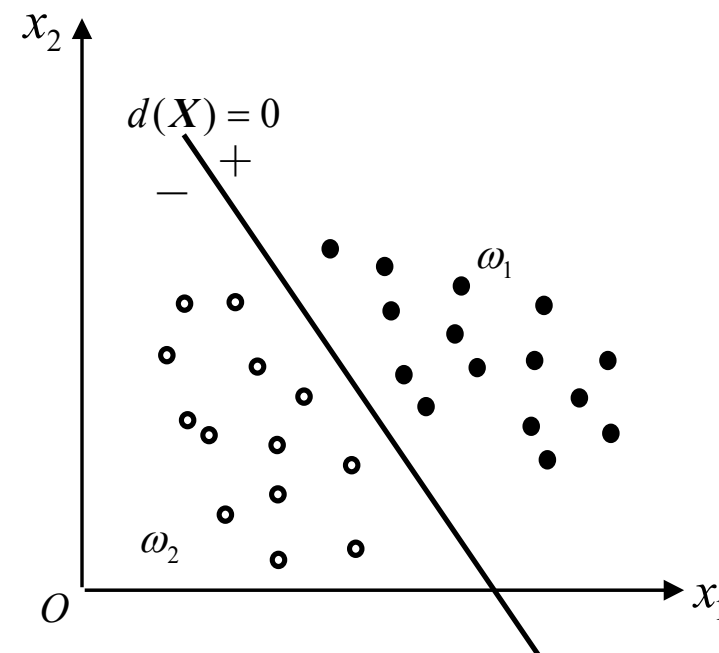
若 $d(\mathbf{X}) < 0$, 则 $\mathbf{X} \in \omega_2$ 类;

若 $d(\mathbf{X}) = 0$, 则 $\mathbf{X} \in \omega_1$ 或 $\mathbf{X} \in \omega_2$ 或拒绝

方程 $d(\mathbf{X})$ 定义了一个决策面,

维数=3时: 判别面为一个平面。

维数>3时: 判别面为一个超平面。

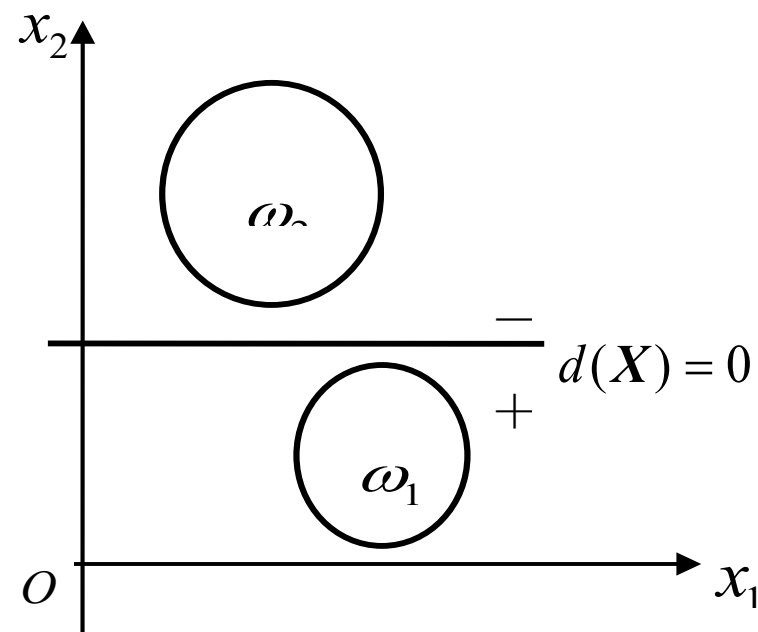


判别函数 (discriminant function)

判别函数正负值的确定

判别界面的正负侧，是在训练判别函数的权值时确定的。

$d(X)$ 表示的是一种分类的标准，它可以是1、2、3维的，也可以是更高维的。



判别函数正负的确定

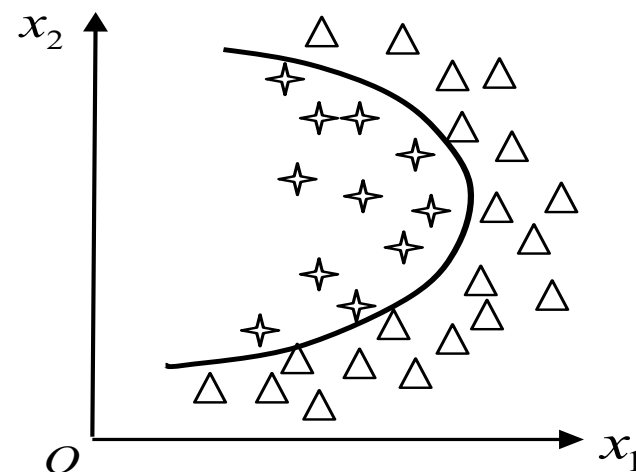
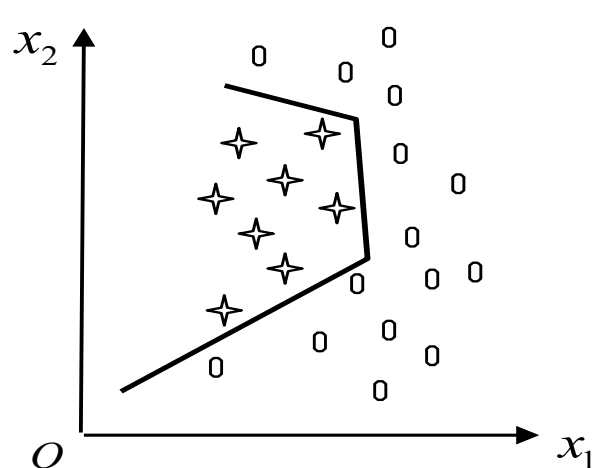
判别函数 (discriminant function)

确定判别函数的两个因素

- 1) 判决函数 $d(\mathbf{X})$ 的几何性质。它可以是线性的或非线性的函数，维数在特征提取时已经确定。
- 2) 判决函数 $d(\mathbf{X})$ 的系数，由所给的模式样本确定。

如：已知三维线性分类——判决函数的性质就确定了判决函数的形式： $d(\mathbf{X}) = w_1x_1 + w_2x_2 + w_3x_3 + w_4$

例：非线性判决函数



线性判别函数

线性判别函数的一般形式

将二维模式推广到 n 维，线性判别函数的一般形式为： $d(\mathbf{X}) = w_1x_1 + w_2x_2 + \dots + w_nx_n + w_{n+1} = \mathbf{W}_0^T \mathbf{X} + w_{n+1}$

式中： $\mathbf{X} = [x_1, x_2, \dots, x_n]^T$ ， $\mathbf{W}_0 = [w_1, w_2, \dots, w_n]^T$ ：权向量，即参数向量。

增广向量的形式：

$$d(\mathbf{X}) = w_1x_1 + w_2x_2 + \dots + w_nx_n + w_{n+1} \cdot 1 = [w_1, w_2, \dots, w_n, w_{n+1}] \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \\ 1 \end{bmatrix} = \mathbf{W}^T \mathbf{X}$$

式中： $\mathbf{X} = [x_1, x_2, \dots, x_n, 1]^T$ 为增广模式向量。

$\mathbf{W} = [w_1, w_2, \dots, w_n, w_{n+1}]^T$ 为增广权向量，

线性判别函数

线性判别函数的性质

1. 两类情况 $d(\mathbf{X}) = \mathbf{W}^T \mathbf{X} \begin{cases} > 0, & \text{若 } \mathbf{X} \in \omega_1 \\ < 0, & \text{若 } \mathbf{X} \in \omega_2 \end{cases}$ $d(\mathbf{X}) = 0$: 不可判别情况, 可以 $\mathbf{X} \in \omega_1$ 或 $\mathbf{X} \in \omega_2$ 或拒绝

2. 多类情况

对 M 个线性可分模式类, $\omega_1, \omega_2, \dots, \omega_M$, 有三种划分方式:

$\omega_i / \overline{\omega_i}$ 两分法 ω_i / ω_j 两分法 ω_i / ω_j 两分法特例

线性判别函数

(1)多类情况1: $\omega_i/\bar{\omega}_i$ 两分法

用线性判别函数将属于 ω_i 类的模式与其余不属于 ω_i 类的模式分开。

$$d_i(\mathbf{X}) = \mathbf{W}_i^T \mathbf{X} \begin{cases} > 0, & \text{若 } \mathbf{X} \in \omega_i \\ < 0, & \text{若 } \mathbf{X} \in \bar{\omega}_i \end{cases} \quad i = 1, \dots, M$$

每个类别用一个判别面分割, 因此 M 类需要 M 个线性判别函数。

识别分类时:

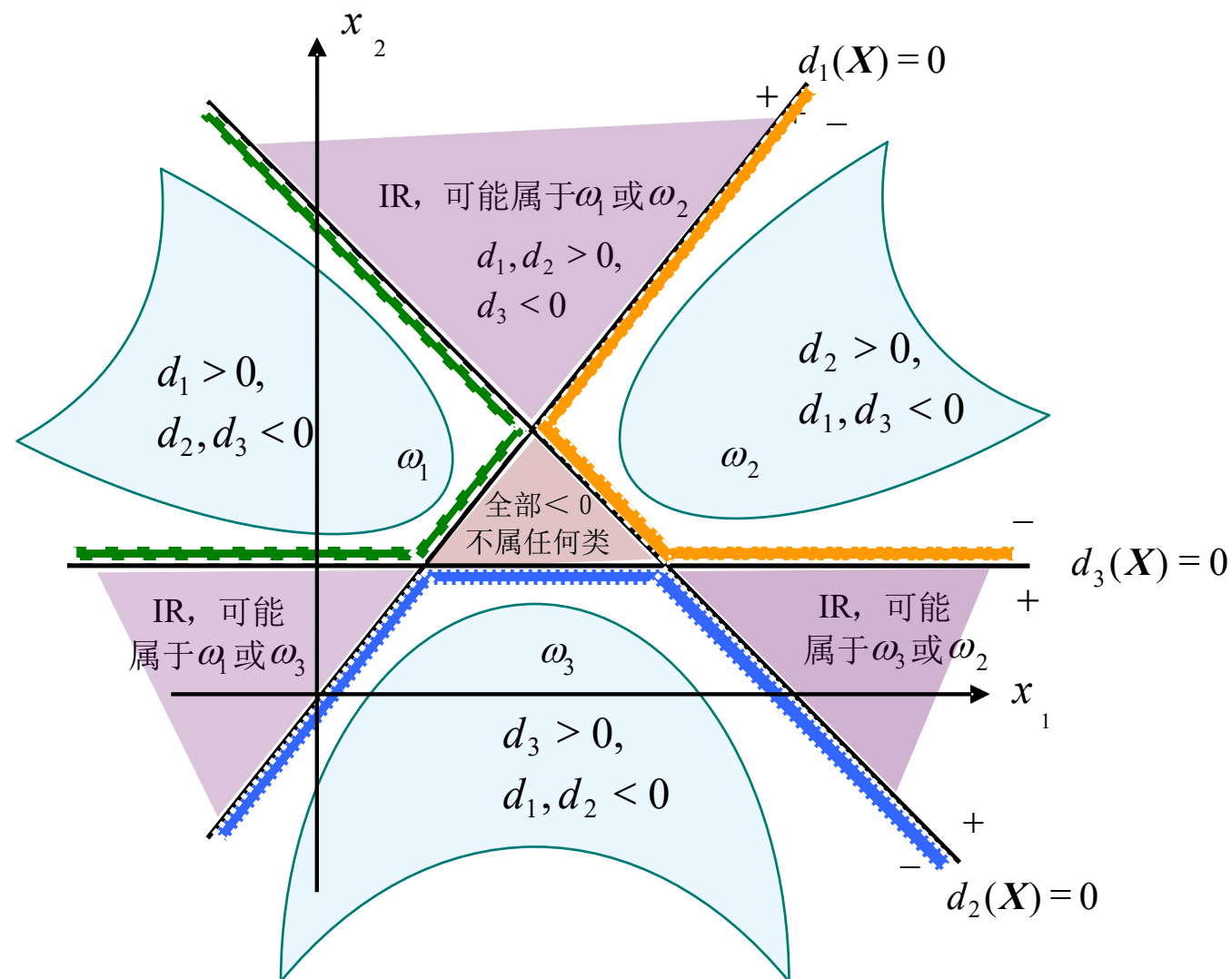
- 将某个待分类模式 $\mathbf{X}$ 分别代入 M 个类的 $d(\mathbf{X})$ 中, 若只有 $d_i(\mathbf{X}) > 0$, 其他 $d(\mathbf{X})$ 均 < 0 , 则判为 ω_i 类。

线性判别函数

(1)多类情况1: $\omega_i/\bar{\omega}_i$ 两分法

此法将 M 个多类问题分成 M 个两类问题，识别每一类均需 M 个判别函数。识别出所有的 M 类仍是这 M 个函数。

对某一模式区， $d_i(\mathbf{X}) > 0$ 的条件超过一个，或全部的 $d_i(\mathbf{X}) < 0$ ，分类失效。相当于不确定区 (indefinite region, IR)。



线性判别函数

(1)多类情况1: $\omega_i/\bar{\omega}_i$ 两分法

例 设有一个三类问题, 其判别式为:

$$d_1(X) = -x_1 + x_2 + 1$$

$$d_2(X) = x_1 + x_2 - 4$$

$$d_3(X) = -x_2 + 1$$

现有一模式, $X = [-7, 5]^T$, 试判定应属于哪类?
并画出三类模式的分布区域。

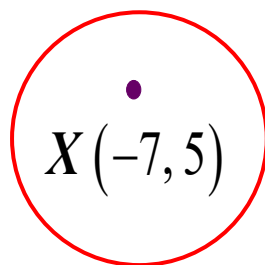
解: 将 $X = [-7, 5]^T$ 代入上三式, 有:

$$d_1(X) = 7 + 5 + 1 = 13 > 0$$

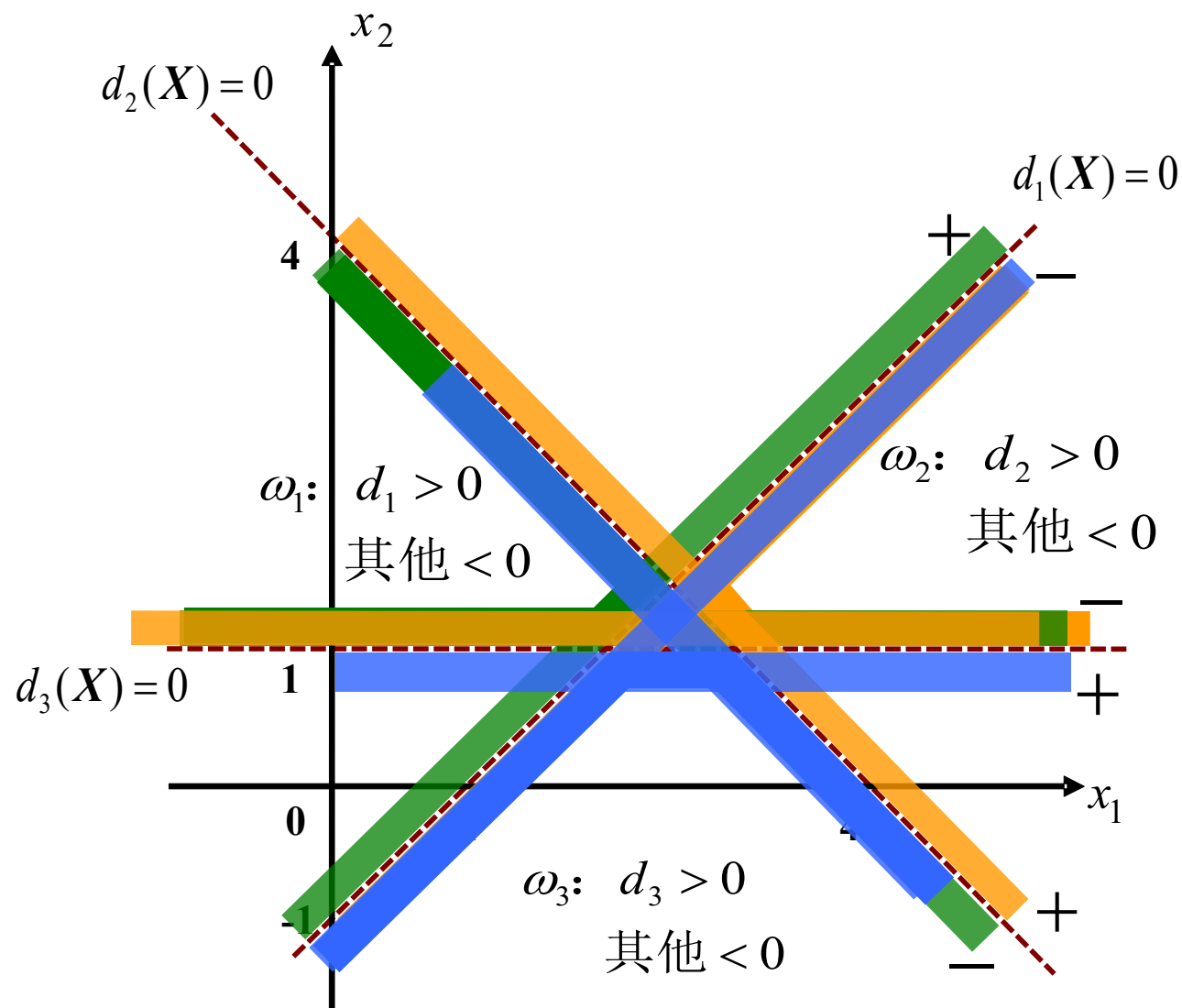
$$d_2(X) = -7 + 5 - 4 = -6 < 0$$

$$d_3(X) = -5 + 1 = -4 < 0$$

$$\because d_1(X) > 0, \quad d_2(X), d_3(X) < 0 \quad \therefore X \in \omega_1$$



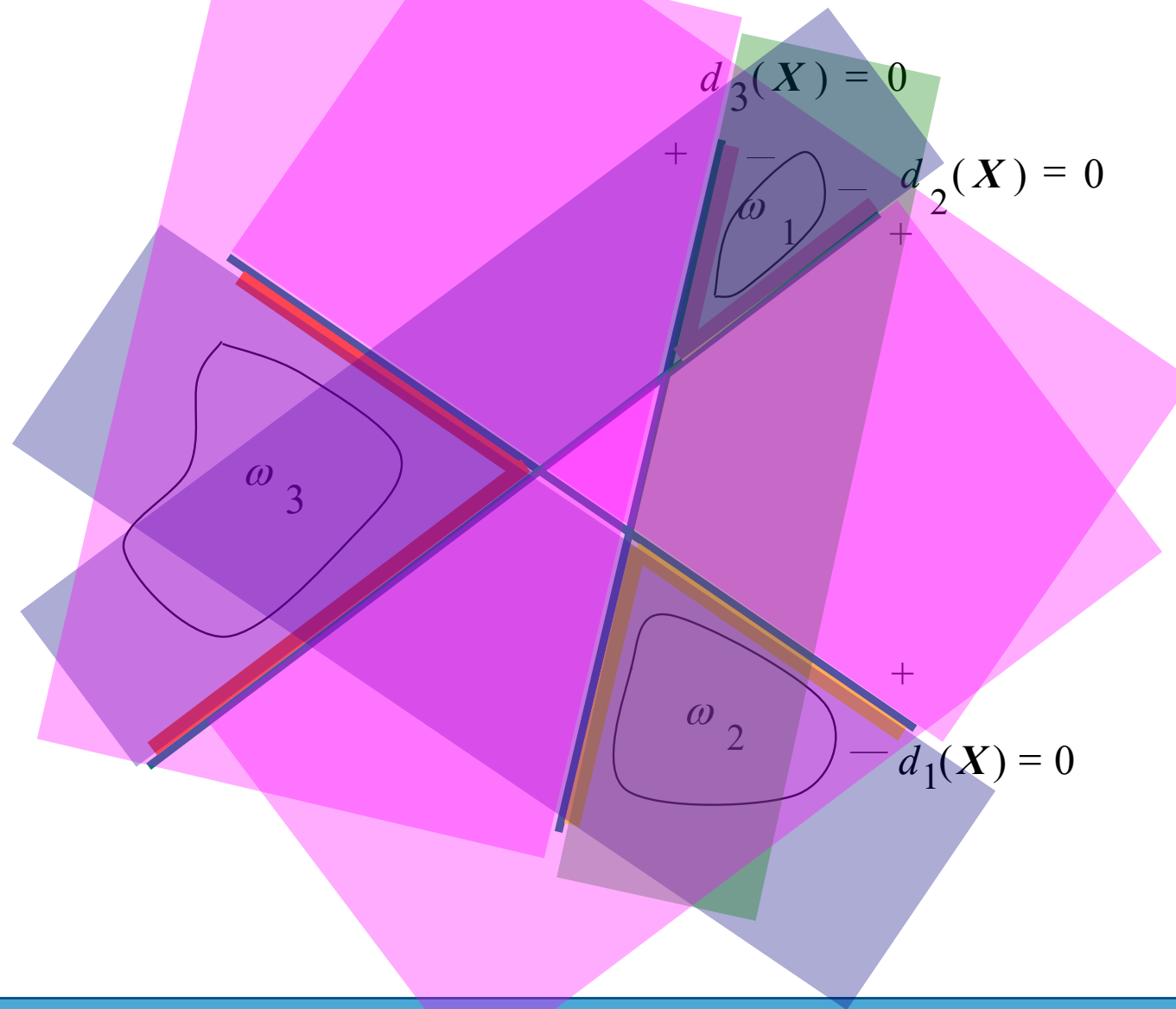
三个判别界面分别为: $d_1(X)=0$ $d_2(X)=0$ $d_3(X)=0$



线性判别函数

(1)多类情况1: $\omega_i/\bar{\omega}_i$ 两分法

例 已知 $d_i(\mathbf{X})$ 的位置和正负侧,
分析三类模式的分布区域



线性判别函数

(2) 多类情况2: ω_i/ω_j 两分法

一个判别界面只能分开两个类别，不能把其余所有的类别都分开。判决函数为：

$$d_{ij}(\mathbf{X}) = \mathbf{W}_{ij}^T \mathbf{X} \quad \text{且} \quad d_{ji} = -d_{ij}$$

判别函数性质: $d_{ij}(\mathbf{X}) > 0, \forall j \neq i; j = 1, 2, \dots, M$, 则 $\mathbf{X} \in \omega_i$

识别分类时:

M 类需要 $M(M-1)/2$ 个线性判别函数。

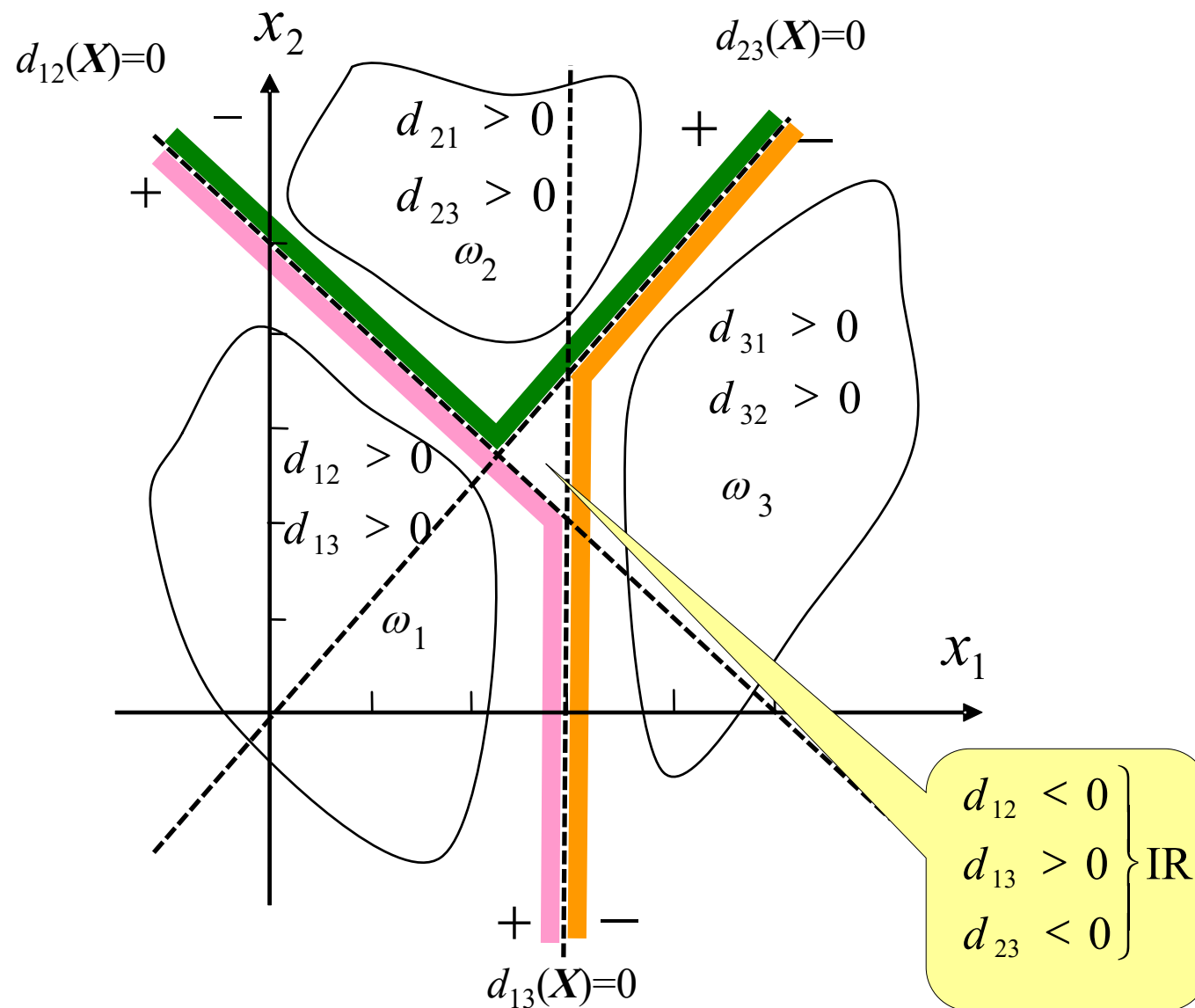
- 在 M 类模式中，与 i 有关的 $M-1$ 个判决函数全为正时， $\mathbf{X} \in \omega_i$ 。其中若有一个为负，则为IR区。

如：三类问题，如果 $d_{12}(\mathbf{X}) > 0, d_{13}(\mathbf{X}) > 0$; 则 $\mathbf{X} \in \omega_1$ 类，而 $d_{23}(\mathbf{X})$ 在判别 ω_1 类模式时不起作用。

线性判别函数

(2) 多类情况2: ω_i/ω_j 两分法

如: 三类问题



线性判别函数

(2) 多类情况2: ω_i/ω_j 两分法

例 一个三类问题, 三个判决函数为:

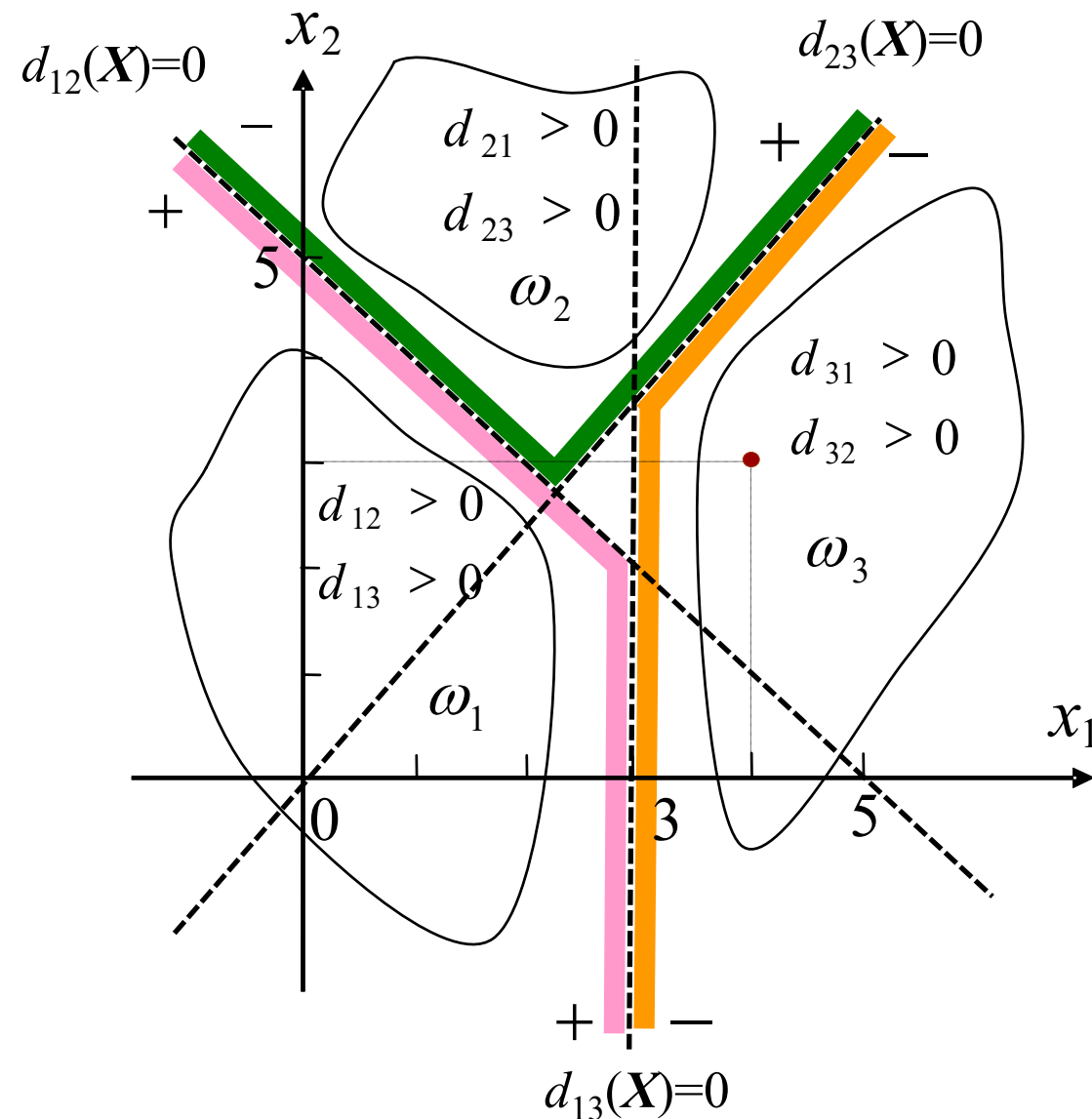
$$d_{12}(\mathbf{X}) = -x_1 - x_2 + 5 \quad d_{13}(\mathbf{X}) = -x_1 + 3 \quad d_{23}(\mathbf{X}) = -x_1 + x_2$$

问模式 $\mathbf{X} = [4, 3]^T$ 属于哪类?

解: 计算得 $d_{12}(\mathbf{X}) = -2$, $d_{13}(\mathbf{X}) = -1$, $d_{23}(\mathbf{X}) = -1$

可写成: $d_{21}(\mathbf{X}) = 2$, $d_{31}(\mathbf{X}) = 1$, $d_{32}(\mathbf{X}) = 1$

$$\left. \begin{array}{l} d_{31}(\mathbf{X}) > 0 \\ d_{32}(\mathbf{X}) > 0 \end{array} \right\} \Rightarrow \mathbf{X} = [4, 3]^T \in \omega_3 \quad \text{与 } d_{12}(\mathbf{X}) \text{ 值无关。}$$



线性判别函数

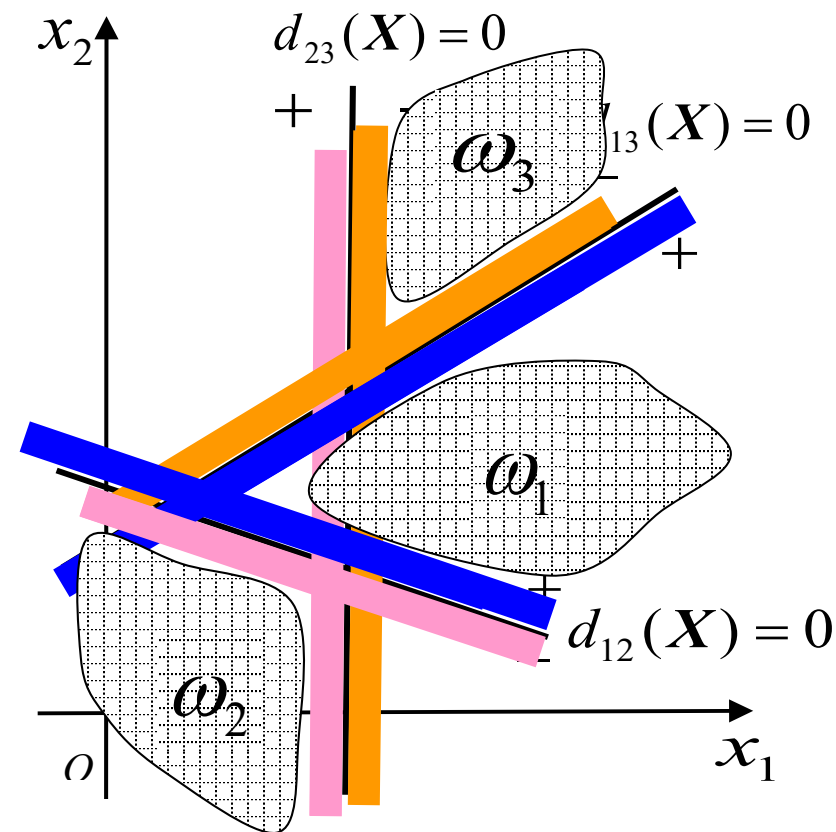
(2) 多类情况2: ω_i/ω_j 两分法

分类时：每分离出一类，需要与 i 有关的 $M-1$ 个判决函数；要分开 M 类模式，共需 $M(M-1)/2$ 个判决函数。对三类问题需要 $3(3-1)/2=3$ 个判决函数。

即：每次从 M 类中取出两类的组合：

$$C_M^2 = \frac{M(M-1)}{2!}$$

例 已知 $d_{ij}(\mathbf{X})$ 的位置和正负侧，分析三类模式的分布区域。



线性判别函数

(3) 多类情况3: ω_i/ω_j 两分法特例

当 ω_i/ω_j 两分法中的判别函数 $d_{ij}(\mathbf{X})$, 可以分解为 $d_{ij}(\mathbf{X}) = d_i(\mathbf{X}) - d_j(\mathbf{X})$ 时, 那么 $d_i(\mathbf{X}) > d_j(\mathbf{X})$ 就相当于多类情况2中的 $d_{ij}(\mathbf{X}) > 0$ 。

因此对判别函数 $d_i(\mathbf{X}) = \mathbf{W}_i^T \mathbf{X}$, $i = 1, \dots, M$ 的 M 类情况, 判别函数性质为:

$$d_i(\mathbf{X}) > d_j(\mathbf{X}), \quad \forall j \neq i; \quad j = 1, 2, \dots, M, \quad \text{则 } \mathbf{X} \in \omega_i$$

或:

$$d_i(\mathbf{X}) = \max \{d_k(\mathbf{X}), \quad k = 1, \dots, M\}, \quad \text{则 } \mathbf{X} \in \omega_i$$

识别分类时:

判别界面需要做差值。对 ω_i 类, 应满足 $d_i >$ 其他所有 d

线性判别函数

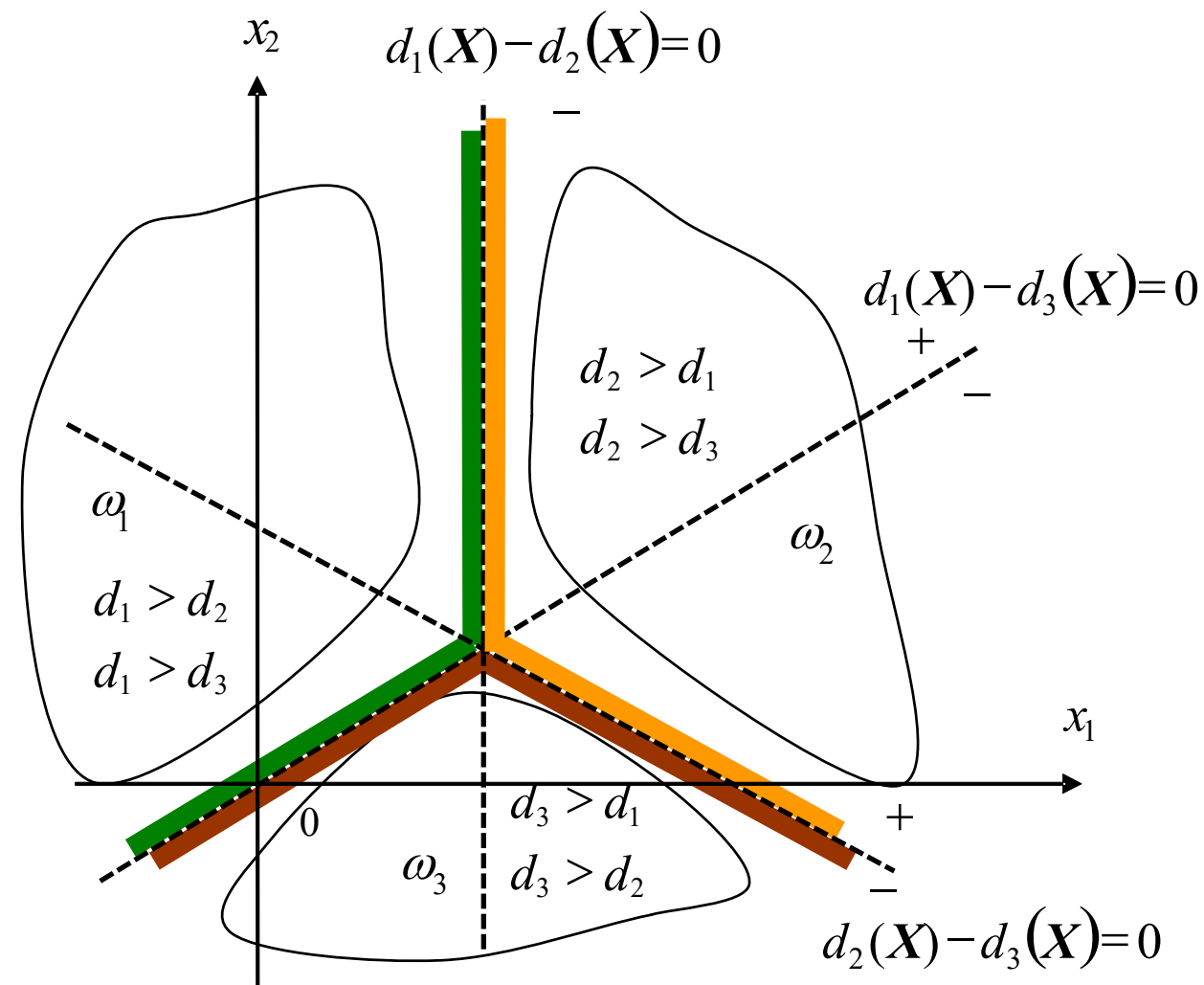
(3) 多类情况3: ω_i/ω_j 两分法特例

特点:

a) 是第二种情况的特例。由于 $d_{ij}(\mathbf{X}) = d_i(\mathbf{X}) - d_j(\mathbf{X})$, 若在第三种情况下可分, 则在第二种情况下也可分, 但反过来不一定。

b) 除边界区外, 没有不确定区域。如右图三条判别线相交于一点, 该点满足 $d_1 = d_2 = d_3$ 。

c) 把 M 类情况分成了 $(M-1)$ 个两类问题。并且 ω_i 类的判别界面全部与 $\omega_j (\forall j \neq i)$ 类的判别界面相邻。



线性判别函数

(3) 多类情况3: ω_i/ω_j 两分法特例

例 一个三类模式 ($M=3$) 分类器, 其判决函数为:

$$d_1(\mathbf{X}) = -x_1 + x_2$$

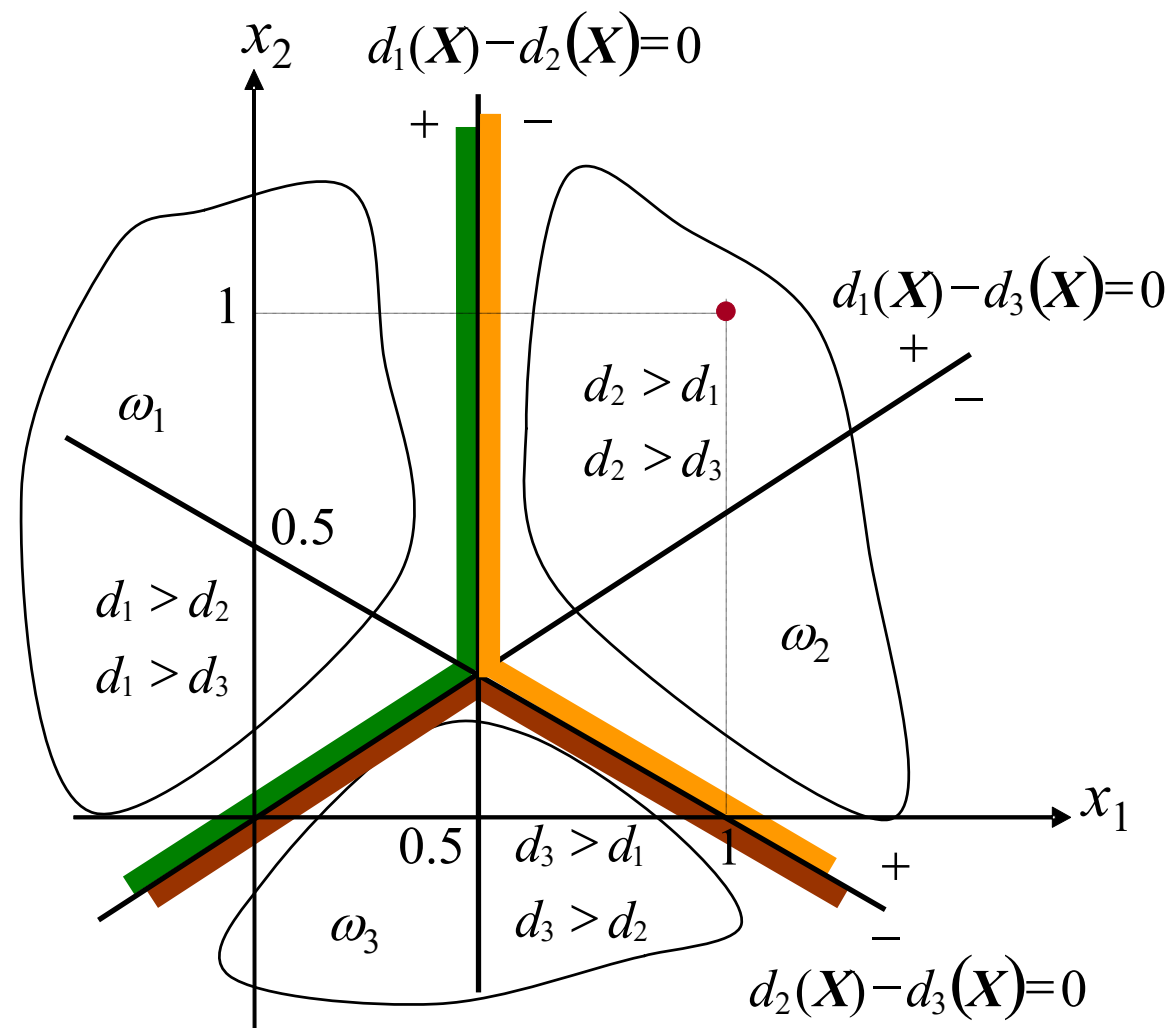
$$d_2(\mathbf{X}) = x_1 + x_2 - 1$$

$$d_3(\mathbf{X}) = -x_2$$

试判断 $\mathbf{X}_0 = [1, 1]^T$ 属于哪一类, 且分别给出三类的判决界面。

解:

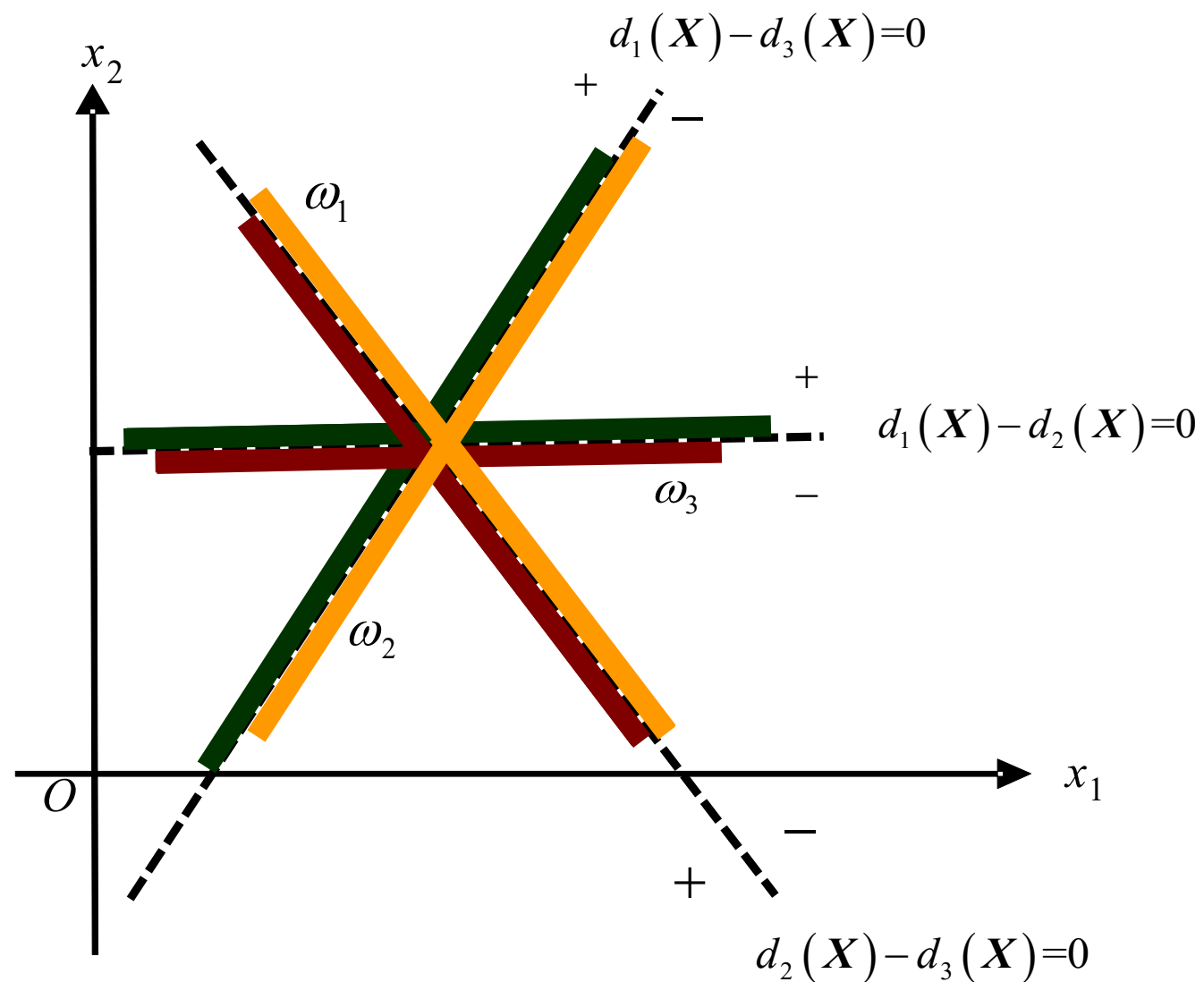
$$\left. \begin{array}{l} d_1(\mathbf{X}_0) = -1 + 1 = 0 \\ d_2(\mathbf{X}_0) = 1 + 1 - 1 = 1 \\ d_3(\mathbf{X}_0) = -1 \end{array} \right\} \Rightarrow \left. \begin{array}{l} d_2(\mathbf{X}_0) > d_1(\mathbf{X}_0) \\ d_2(\mathbf{X}_0) > d_3(\mathbf{X}_0) \end{array} \right\} \Rightarrow \mathbf{X}_0 \in \omega_2$$



线性判别函数

(3) 多类情况3: ω_i/ω_j 两分法特例

例 已知判决界面的位置和正负侧，分析三类模式的分布区域。



线性判别函数

小结 (1) 明确概念：线性可分。

一旦线性判别函数的系数 $\mathbf{W}_k$ 被确定以后，这些函数就可以作为模式分类的基础。

(2) $\omega_i/\overline{\omega}_i$ 与 ω_i/ω_j 分法的比较：

对于 M 类模式的分类， $\omega_i/\overline{\omega}_i$ 两分法共需要 M 个判别函数，但 ω_i/ω_j 两分法需要 $M(M-1)/2$ 个。

当时 $M>3$ 时，后者需要更多个判别式（缺点），但对模式的线性可分的可能性要更大一些（优点）。

原因：

一种类别模式的分布要比 $M-1$ 类模式的分布更为聚集， ω_i/ω_j 分法受到的限制条件少，故线性可分的可能性大。

广义线性判别函数

目的：对非线性边界：通过某映射，把模式空间 X 变成 X^* ，以便将 X 空间中非线性可分的模式集，变成在 X^* 空间中线性可分的模式集。

非线性多项式函数

非线性判别函数的形式之一是非线性多项式函数。

设一训练用模式集， $\{X\}$ 在模式空间 X 中线性不可分，非线性判别函数形式如下：

$$d(X) = w_1 f_1(X) + w_2 f_2(X) + \dots + w_k f_k(X) + w_{k+1} = \sum_{i=1}^{k+1} w_i f_i(X)$$

式中 $\{f_i(X), i=1, 2, \dots, k\}$ 是模式 X 的单值实函数， $f_{k+1}(X)=1$ 。

$f_i(X)$ 取什么形式及 $d(X)$ 取多少项，取决于非线性边界的复杂程度。

广义线性判别函数

广义形式的模式向量定义为: $\mathbf{X}^* = [x_1^*, x_2^*, \dots, x_k^*, 1]^T = [f_1(\mathbf{X}), f_2(\mathbf{X}), \dots, f_k(\mathbf{X}), 1]^T$

$$d(\mathbf{X}) = \mathbf{W}^T \mathbf{X}^* = d(\mathbf{X}^*), \quad \mathbf{W} = [w_1, w_2, \dots, w_k, w_{k+1}]^T$$

上式是线性的, 讨论线性判别函数并不会失去一般性的意义。

例: 有一个三次判别函数: $z=g(x)=x^3+2x^2+3x+4$ 。试建立一映射 $x \rightarrow \mathbf{y}$, 使得 z 化为 $\mathbf{y}$ 的线性判别函数。

解: 映射 $\mathbf{X} \rightarrow \mathbf{Y}$ 如下:

$$\mathbf{y} = \begin{bmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \end{bmatrix} = \begin{bmatrix} 1 \\ x \\ x^2 \\ x^3 \end{bmatrix}, \quad \mathbf{w} = \begin{bmatrix} a_1 \\ a_2 \\ a_3 \\ a_4 \end{bmatrix} = \begin{bmatrix} 4 \\ 3 \\ 2 \\ 1 \end{bmatrix}$$

$$z = g(x) = h(\mathbf{y}) = \mathbf{w}^T \mathbf{y} = \sum_{i=1}^4 w_i y_i$$

广义线性判别函数

例 假设 X 为二维模式向量, $f_i(\mathbf{X})$ 选用二次多项式函数, 原判别函数为

$$d(\mathbf{X}) = w_{11}x_1^2 + w_{12}x_1x_2 + w_{22}x_2^2 + w_1x_1 + w_2x_2 + w_3$$

广义线性判别函数:

$$\begin{aligned} \text{定义: } x_1^* &= f_1(\mathbf{X}) = x_1^2 & x_2^* &= f_2(\mathbf{X}) = x_1x_2 \\ x_3^* &= f_3(\mathbf{X}) = x_2^2 & x_4^* &= f_4(\mathbf{X}) = x_1 & x_5^* &= f_5(\mathbf{X}) = x_2 \end{aligned}$$

$$\text{即: } \mathbf{X}^* = [x_1^2, x_1x_2, x_2^2, x_1, x_2, 1]^T \quad \mathbf{w} = [w_{11}, w_{12}, w_{22}, w_1, w_2, w_3]^T$$

$$d(\mathbf{X}) \text{线性化为: } d(\mathbf{X}^*) = \mathbf{w}^T \mathbf{X}^*$$

广义线性判别函数

例： 设在三维空间中一个类别分类问题拟采用二次曲面。如欲采用**广义线性方程**求解，试问其**广义样本向量**与**广义权向量**的表达式，其维数是多少？

答：设二次曲面为：

二次曲面

$$ax_1^2 + bx_2^2 + cx_3^2 + dx_1x_2 + ex_1x_3 + fx_2x_3 + gx_1 + hx_2 + lx_3 + m = 0$$

广义权向量

$$\mathbf{w} = (a, b, c, d, e, f, g, h, l, m)^T$$

广义样本向量

$$\mathbf{y} = (x_1^2, x_2^2, x_3^2, x_1x_2, x_1x_3, x_2x_3, x_1, x_2, x_3, 1)^T$$

维数为10

广义线性判别函数

$$z = g(\mathbf{x}) = h(\mathbf{y}) = \mathbf{w}^T \mathbf{y}$$

广义线性判别函数

存在的问题：

- 1) 非线性变换可能非常复杂。
- 2) 维数大大增加： 维数灾难。

随着小样本学习理论和支持向量机的迅速发展，广义线性判别函数的“维数灾难”问题在一定程度上找到了解决的办法。

线性判别函数的几何性质

模式空间与超平面

1. 概念

模式空间：以 n 维模式向量 $\mathbf{X}$ 的 n 个分量为坐标变量的欧氏空间。

模式向量的表示：点、有向线段。

线性分类：用 $d(\mathbf{X})$ 进行分类，相当于用超平面 $d(\mathbf{X})=0$ 把模式空间分成不同的决策区域。

2. 讨论

设判别函数： $d(\mathbf{X}) = \mathbf{W}_0^T \mathbf{X} + w_{n+1}$

式中 $\mathbf{W}_0 = [w_1, w_2, \dots, w_n]^T$, $\mathbf{X} = [x_1, x_2, \dots, x_n]^T$ 。

超平面： $d(\mathbf{X}) = \mathbf{W}_0^T \mathbf{X} + w_{n+1} = 0$

线性判别函数的几何性质

(1) 模式向量 $\mathbf{X}_1$ 和 $\mathbf{X}_2$ 在超平面上

$$\mathbf{W}_0^T \mathbf{X}_1 + w_{n+1} = \mathbf{W}_0^T \mathbf{X}_2 + w_{n+1}$$

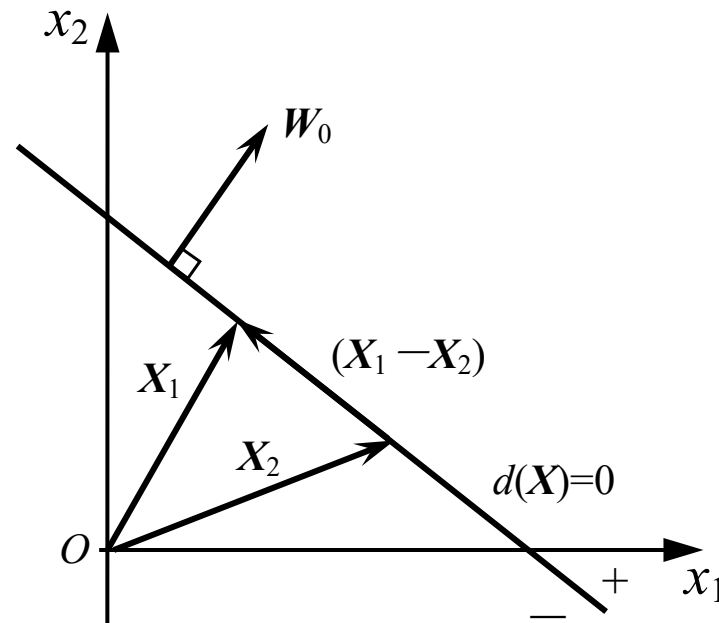
$$\mathbf{W}_0^T (\mathbf{X}_1 - \mathbf{X}_2) = 0$$

—— $\mathbf{W}_0$ 是超平面的法向量，方向由超平面的负侧指向正侧。

设超平面的单位法线向量为 $\mathbf{U}$:

$$\mathbf{U} = \frac{\mathbf{W}_0}{\|\mathbf{W}_0\|}$$

$$\|\mathbf{W}_0\| = \sqrt{w_1^2 + w_2^2 + \cdots + w_n^2}$$



线性判别函数的几何性质

(2) X 不在超平面上, 将 X 向超平面投影得向量 X_p , 构造向量 R :

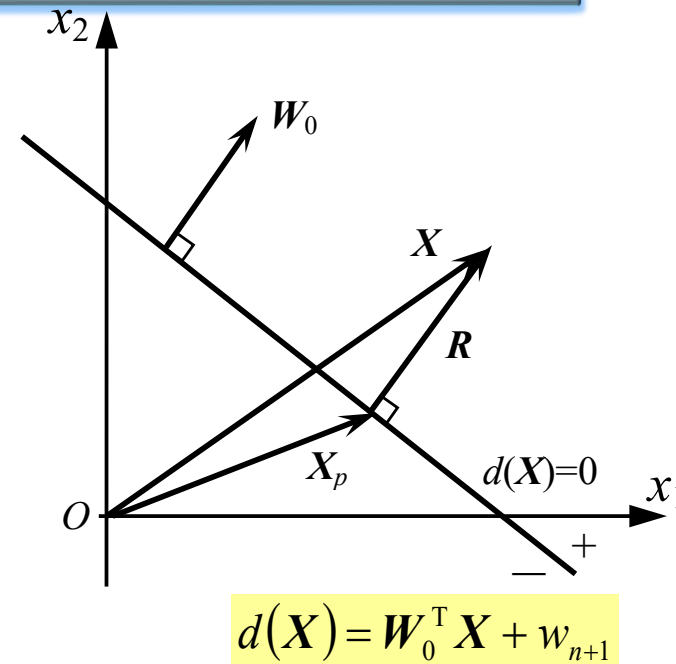
$$R = r \cdot U = r \frac{W_0}{\|W_0\|}$$

r : X 到超平面的垂直距离。有: $X = X_p + R = X_p + r \frac{W_0}{\|W_0\|}$

$$d(X) = W_0^T \left(X_p + r \frac{W_0}{\|W_0\|} \right) + w_{n+1} = W_0^T X_p + w_{n+1} + r \frac{W_0^T W_0}{\|W_0\|} \quad \text{由于 } X_p \text{ 在超平面上, 即 } d(X_p)=0;$$

$$d(X) = r \frac{W_0^T W_0}{\|W_0\|} = r \|W_0\| \quad \text{—— 判别函数 } d(X) \text{ 正比于点 } X \text{ 到超平面的代数距离。}$$

$$X \text{ 到超平面的距离: } r = \frac{d(X)}{\|W_0\|} \quad \text{—— 点 } X \text{ 到超平面的代数距离 (带正负号) 正比于 } d(X) \text{ 函数值。}$$



线性判别函数的几何性质

(3) X 在原点 $d(X) = \mathbf{W}_0^T \mathbf{X} + w_{n+1} = w_{n+1}$

得
$$r_0 = \frac{d(X)}{\|\mathbf{W}_0\|} = \frac{w_{n+1}}{\|\mathbf{W}_0\|}$$

——超平面的位置由阈值权 w_{n+1} 决定:

$w_{n+1} > 0$ 时, 原点在超平面的正侧;

$w_{n+1} < 0$ 时, 原点在超平面负侧;

$w_{n+1} = 0$ 时, 超平面通过原点。



线性模型基本形式

线性方程, 基本形式, ...



线性回归

参数估计, 最小二乘, 对数几率回归, ...



线性判别分析

基本概念, 广义线性判别, 几何性质 ...



经典线性判别准则

Fisher判别, 感知器准则、梯度法 ...



非线性判别函数

分段线性, ...

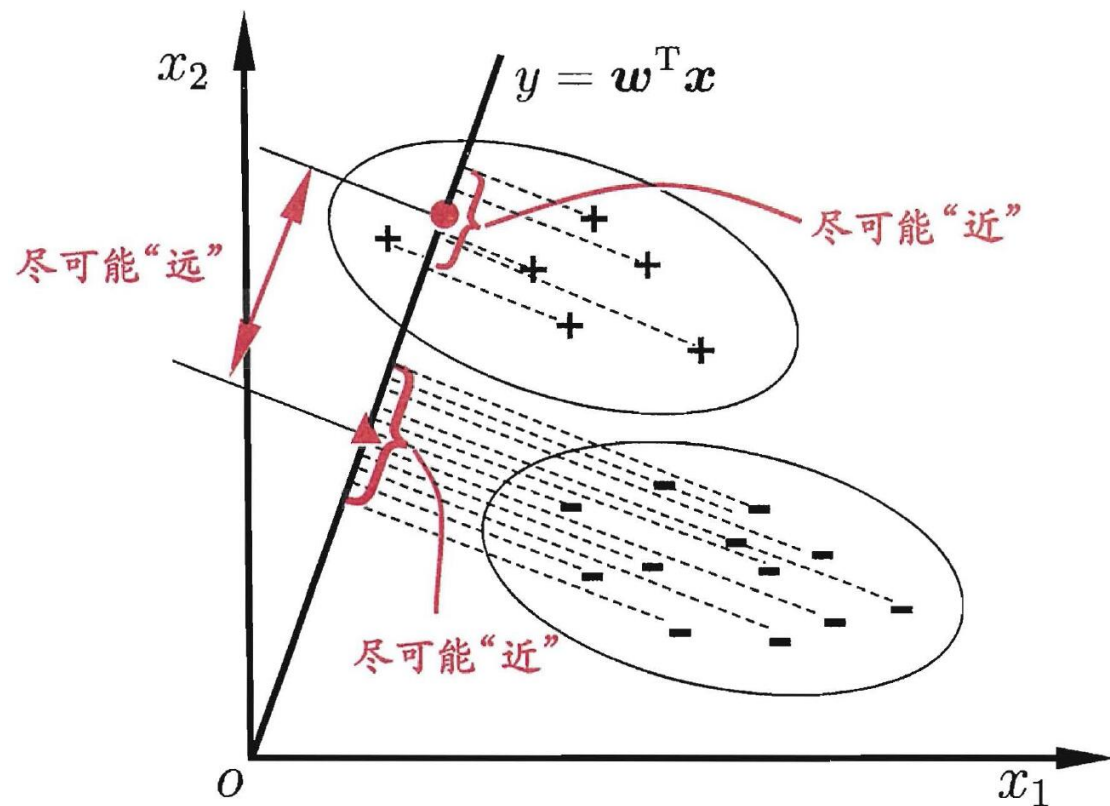
Fisher线性判别 (Fisher Linear Discrimination, FLD)

Fisher准则是英国统计学家R. A. Fisher在1936年首先提出的一种线性判别方法。

Fisher准则的基本原理： 投影（降维）

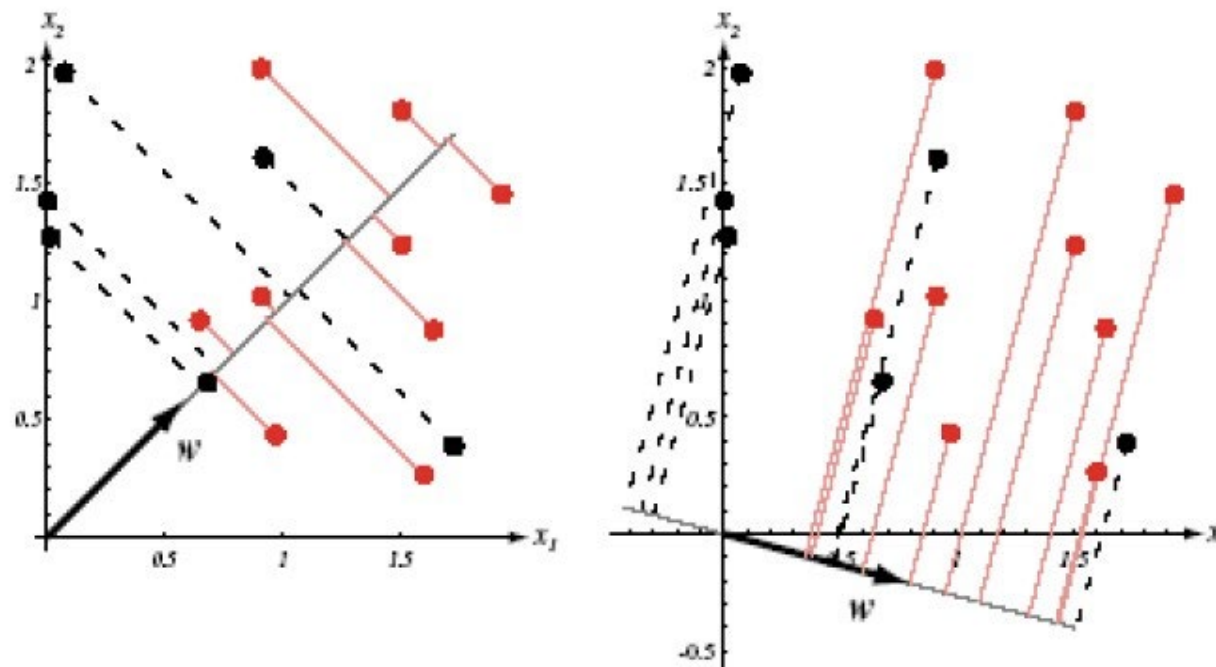
找到一个**最合适的投影轴**，将多维数据投影到该投影轴上，使两类样本在该投影轴上投影之间的距离尽可能远，而每一类样本的投影尽可能紧凑，从而使分类效果为最佳。

用投影后数据的统计性质 (均值和离散度的函数) 作为判别优劣的标准。



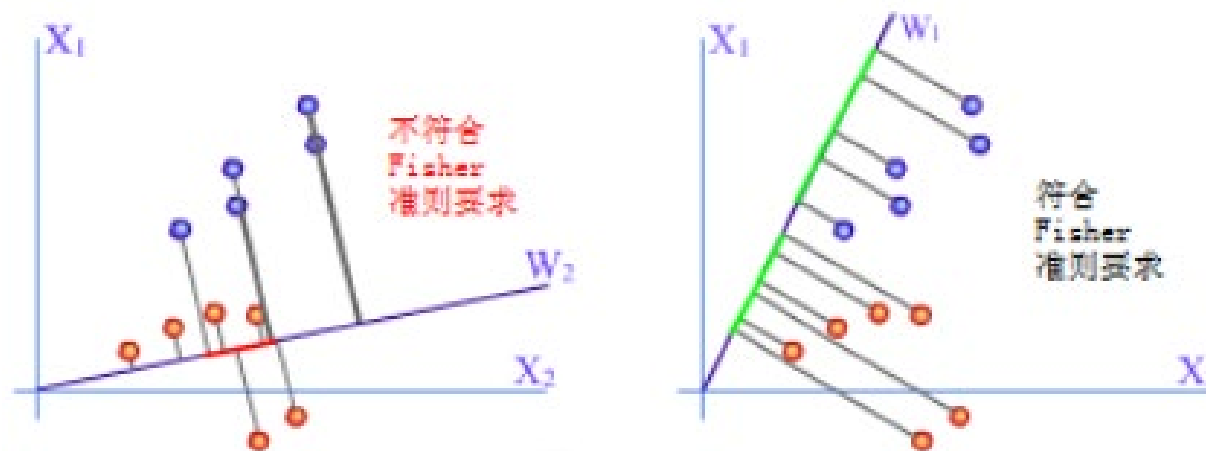
Fisher线性判别

把同一组样本点向两个不同的方向作投影，显然右图更易分开。



问题：

如何根据实际情况找到一条最好的、最易于分类的投影线？



Fisher线性判别 (Fisher Linear Discrimination, FLD)

d 维到一维的数学变换

样本集: N 个 d 维样本 $\mathcal{X} = \{x_1, \dots, x_N\}$

ω_1 类包含样本 $\mathcal{X}_1 = \{x_1^1, \dots, x_{N_1}^1\}$, ω_2 类包含样本 $\mathcal{X}_2 = \{x_1^2, \dots, x_{N_2}^2\}$

在 Y 空间的一维投影 $y_i = w^T x_i$, $i = 1, \dots, N$

在 $\mathcal{X}$ 空间, 类均值向量 $m_i = \frac{1}{N_i} \sum_{x_j \in \mathcal{X}_i} x_j$, $i = 1, 2$

类内离散度矩阵 within-class scatter $S_i = \sum_{x_j \in \mathcal{X}_i} (x_j - m_i)(x_j - m_i)^T$, $i = 1, 2$

总类内离散度矩阵 $S_w = S_1 + S_2$

样本类间离散度矩阵 $S_b = (m_1 - m_2)(m_1 - m_2)^T$

离散度矩阵在形式上与协方差矩阵很相似, 但协方差矩阵是一种总体期望值, 而离散矩阵只是表示有限个样本在空间分布的离散程度。

Fisher线性判别 (Fisher Linear Discrimination, FLD)

d 维到一维的数学变换

一维投影空间 Y : 各类样本均值 $\tilde{m}_i = \frac{1}{N_i} \sum_{y \in \psi_i} y, \quad i = 1, 2$

样本类内离散度 $\tilde{S}_i = \sum_{y \in \psi_i} (y - \tilde{m}_i)^2, \quad i = 1, 2$

总类内离散度 $\tilde{S}_w = \tilde{S}_1 + \tilde{S}_2$

样本类间离散度 $\tilde{S}_b = (\tilde{m}_1 - \tilde{m}_2)^2$

以上定义描述 d 维空间样本点到一向量投影后的分散情况。样本离散度的定义与随机变量方差相类似。

Fisher线性判别 (Fisher Linear Discrimination, FLD)

d 维到一维投影后, 样本 $\mathbf{x}$ 与其投影 y 的统计量之间的关系:

各类样本均值 $\tilde{m}_i = \frac{1}{N_i} \sum_{y \in \mathcal{Y}_i} y = \frac{1}{N_i} \sum_{y \in K_i} \mathbf{w}^T \mathbf{x} = \mathbf{w}^T \mathbf{m}_i, \quad i = 1, 2$

样本类内离散度 $\tilde{S}_i = \sum_{y \in \mathcal{Y}_i} (y - \tilde{m}_i)^2 = \sum_{x \in K_i} (\mathbf{w}^T \mathbf{x} - \mathbf{w}^T \mathbf{m}_i)^2 = \mathbf{w}^T \left[\sum_{x \in K_i} (\mathbf{x} - \mathbf{m}_i)(\mathbf{x} - \mathbf{m}_i)^T \right] \mathbf{w} = \mathbf{w}^T S_i \mathbf{w}$

总类内离散度 $\tilde{S}_w = \tilde{S}_1 + \tilde{S}_2 = \mathbf{w}^T (S_1 + S_2) \mathbf{w} = \mathbf{w}^T S_w \mathbf{w}$

样本类间离散度 $\tilde{S}_b = (\tilde{m}_1 - \tilde{m}_2)^2 = (\mathbf{w}^T \mathbf{m}_1 - \mathbf{w}^T \mathbf{m}_2)^2 = \mathbf{w}^T (\mathbf{m}_1 - \mathbf{m}_2)(\mathbf{m}_1 - \mathbf{m}_2)^T \mathbf{w} = \mathbf{w}^T S_b \mathbf{w}$

Fisher线性判别 (Fisher Linear Discrimination, FLD)

Fisher判别准则： d 维到一维投影后，在一维Y空间里各个样本尽可能做到：

1、分得开

两类均值差 $(\tilde{m}_1 - \tilde{m}_2)^2$ 越大越好

2、各类样本内部尽量密集

类内离散度 $\tilde{S}_1^2$ 和 $\tilde{S}_2^2$ 越小越好

Fisher准则函数(Fisher's Criterion):
$$J_F(\mathbf{w}) = \frac{(\tilde{m}_1 - \tilde{m}_2)^2}{\tilde{S}_1 + \tilde{S}_2} = \frac{\tilde{S}_b}{\tilde{S}_w} = \frac{\mathbf{w}^T \mathbf{S}_b \mathbf{w}}{\mathbf{w}^T \mathbf{S}_w \mathbf{w}}$$

Fisher最佳投影方向的求解 $\mathbf{w}^* = \underset{\mathbf{w}}{\operatorname{argmax}} J_F(\mathbf{w})$

采用拉格朗日乘子算法求解，得： $\mathbf{w}^* = \mathbf{S}_w^{-1}(\mathbf{m}_1 - \mathbf{m}_2)$

Fisher线性判别 (Fisher Linear Discrimination, FLD)

实现步骤: 1) 计算各类样品均值向量 $\mathbf{m}_i$, $\mathbf{m}_i$ 是各个类的均值, N_i 是类 ω_i 的样品个数

$$\mathbf{m}_i = \frac{1}{N_i} \sum_{x_j \in X_i} x_j, \quad i = 1, 2$$

2) 计算样品类内离散度矩阵 $\mathbf{S}_i$ 和总类内离散度矩阵 $\mathbf{S}_w$

$$\mathbf{S}_i = \sum_{x_j \in X_i} (x_j - \mathbf{m}_i)(x_j - \mathbf{m}_i)^T, \quad i = 1, 2$$

$$\mathbf{S}_w = \mathbf{S}_1 + \mathbf{S}_2$$

3) 计算样品类间离散度矩阵 $\mathbf{S}_b$

$$\mathbf{S}_b = (\mathbf{m}_1 - \mathbf{m}_2)(\mathbf{m}_1 - \mathbf{m}_2)^T$$

4) 求最佳投影向量 $\mathbf{w}^*$

$$\mathbf{w}^* = \mathbf{S}_w^{-1}(\mathbf{m}_1 - \mathbf{m}_2)$$

Fisher线性判别 (Fisher Linear Discrimination, FLD)

实现步骤: 5) 对训练集内所有样品进行投影 $y = (\mathbf{w}^*)^T \mathbf{x}$

6) 计算在一维投影空间上的分割阈值 y_0

在一维投影空间 Y 上的各样本均值 $\tilde{m}_i = \frac{1}{N_i} \sum_{y \in \psi_i} y \quad i=1,2$

类内离散度 $\tilde{S}_i^2$ 和总类内离散度 $\tilde{S}_w$ 为: $\tilde{S}_i = \sum_{y \in \psi_i} (y - \tilde{m}_i)^2$, 且 $\tilde{S}_w = \tilde{S}_1 + \tilde{S}_2$

分割阈值 y_0 的选取可以有不同方案:

$$y_0 = \frac{N_1 \tilde{m}_1 + N_2 \tilde{m}_2}{N_1 + N_2} \quad \text{或} \quad y_0 = \frac{\tilde{m}_1 + \tilde{m}_2}{2} + \frac{\ln(p(\omega_1)/p(\omega_2))}{N_1 + N_2 - 2}$$

7) 对于给定的 $\mathbf{X}$, 计算它在投影向量 $\mathbf{w}^*$ 上的投影点 Y

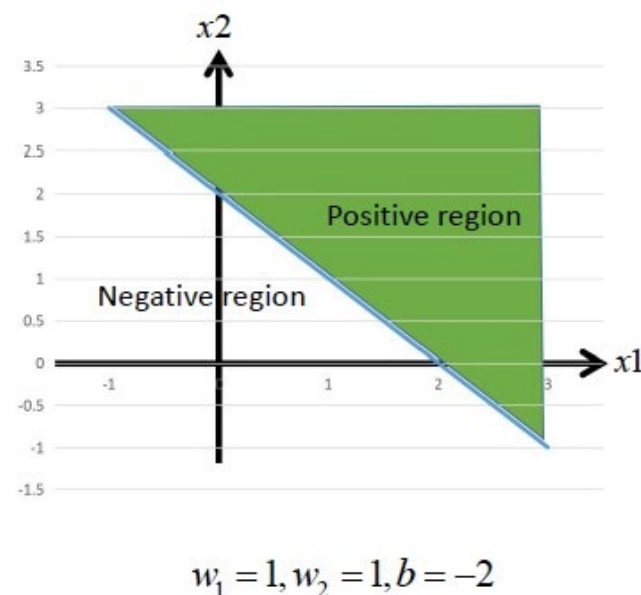
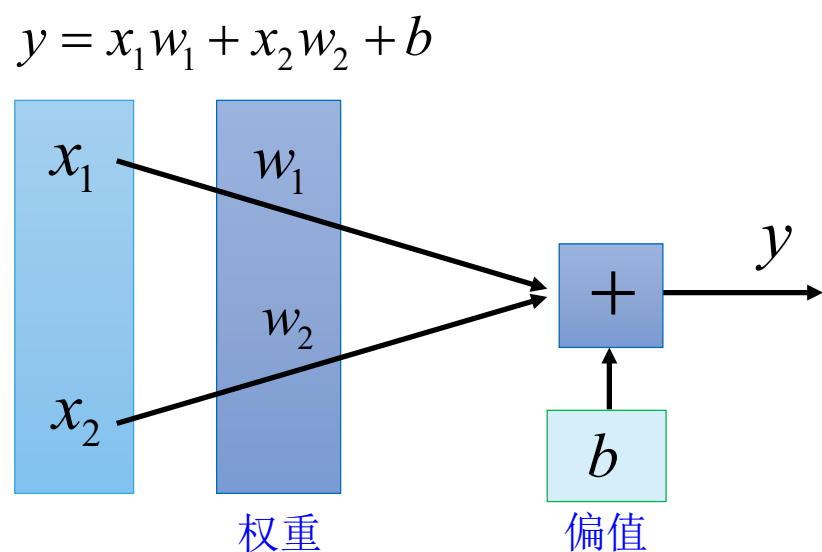
根据决策规则分类
$$\begin{cases} Y > y_0 \Rightarrow \mathbf{X} \in \omega_1 \\ Y < y_0 \Rightarrow \mathbf{X} \in \omega_2 \end{cases}$$

感知器准则 (Perceptron)

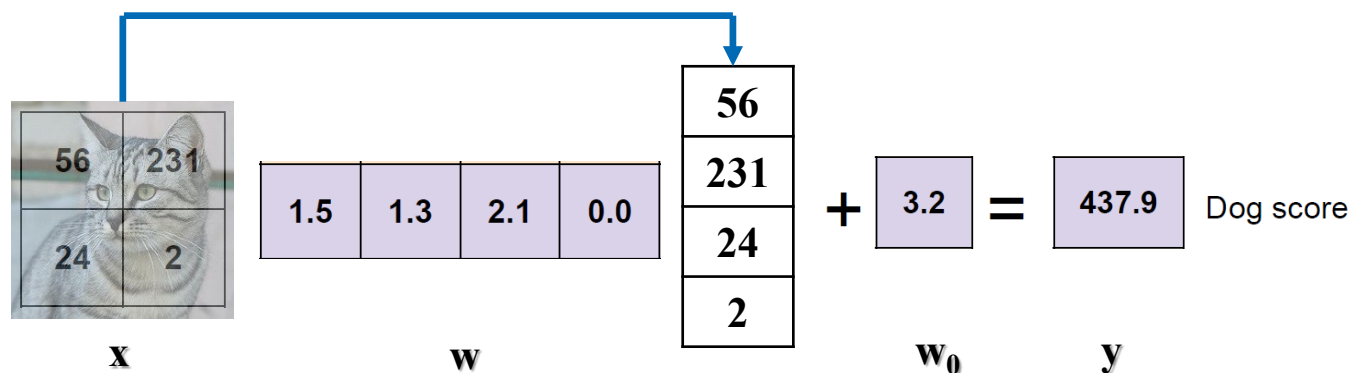
感知器准则是五十年代由Rosenblatt提出的一种自学习判别函数生成方法；

Rosenblatt将其用于脑模型感知器(Perceptron)，因此被称为感知准则函数；

特点：随意确定的判别函数初始值，在对样本分类训练过程中逐步修正直至最终确定。

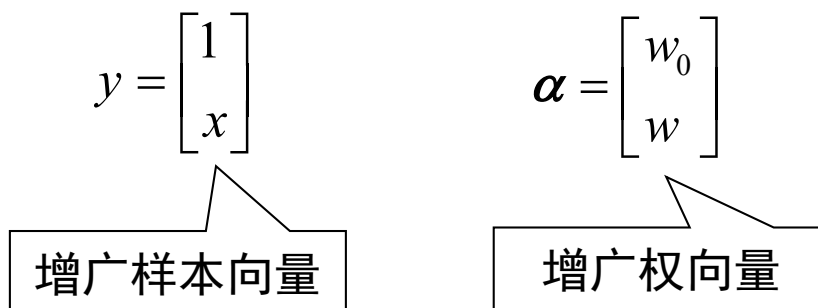


感知器准则 (Perceptron)



基本概念

1、线性判别函数的齐次简化 $g(x) = w^T x + w_0 = \alpha^T y$



只要求出权向量，分类器即设计完成。

问题为：如何通过各种算法，利用已知类别的模式样本训练出权向量 w 。

感知器准则 (Perceptron)

基本概念

2、线性可分性：训练样本集中的两类样本在特征空间可以用一个线性分界面正确无误地分开。在线性可分条件下，对合适的（广义）权向量 $\mathbf{a}$ 应有：

如果 $\mathbf{y}_i \in \omega_1$, 则 $\mathbf{a}^T \mathbf{y}_i > 0$

如果 $\mathbf{y}_i \in \omega_2$, 则 $\mathbf{a}^T \mathbf{y}_i < 0$

3、样本的规范化

$$\mathbf{y}'_n = \begin{cases} \mathbf{y}_i & \text{如果 } \mathbf{y}_i \in \omega_1 \\ -\mathbf{y}_j & \text{如果 } \mathbf{y}_j \in \omega_2 \end{cases} \Rightarrow \mathbf{a}^T \mathbf{y}'_n > 0 \quad i = 1, \dots, N$$

无需考虑样本原来的类别标志，只要找一个对全部样本 $\mathbf{y}'_n$ 都满足 $\mathbf{a}^T \mathbf{y}'_n > 0$ 的权向量 $\mathbf{a}$ 就可以。上述过程称为样本的规范化。 $\mathbf{y}'_n$ 称为规范化增广样本向量。

感知器准则 (Perceptron)

基本概念

4、解向量和解区：

解向量：在线性可分情况下，满足 $\alpha^T \mathbf{y}_n > 0$ 的权向量，记作 α^* 。

解区：权值空间中所有解向量组成的区域。

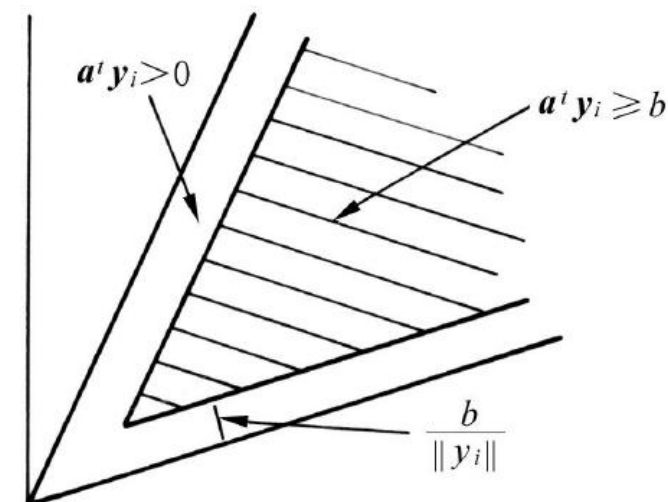
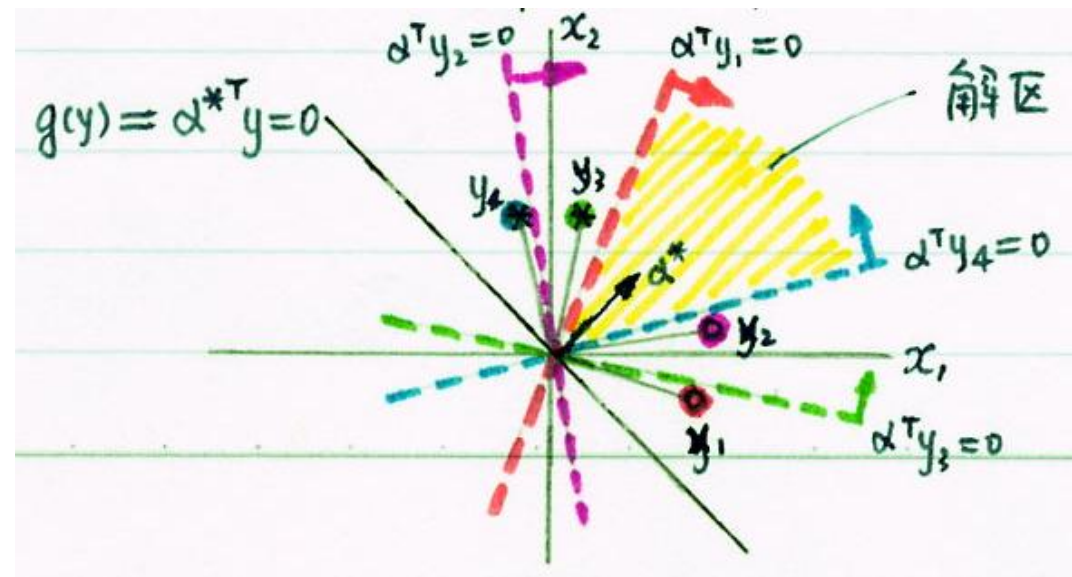
对每个样本 $\mathbf{y}_i$ ， $\alpha^T \mathbf{y}_i = 0$ 为权空间中的一个超平面 $\hat{H}_i$ ，解区只可能在 $\hat{H}_i$ 的正侧。所有样本对应的超平面 $\hat{H}_i$ 的正侧的交集就是解区。

余量：解区中所有权向量都是解，但越靠近解区中间的解向量，所得的分类面离各类样本越远（最近距离越远），对将来的样本错分的可能性越小。

引入余量 $b > 0$ ，要求解向量满足 $\alpha^T \mathbf{y}_i \geq b$ ， $i=1, \dots, N$

若要求 b 最大则为“最优分类面”目的：

1. 解更可靠，（推广性更强）
2. 防止算法收敛到解区边界



感知器准则 (Perceptron)

基本概念

5、训练与学习

训练：用已知类别的模式样本指导机器对分类规则进行反复修改，最终使分类结果与已知类别信息相同的过程。

学习：从分类器的角度讲 $\left\{ \begin{array}{l} \text{非监督学习} \\ \text{监督学习} \end{array} \right. \longleftrightarrow \text{训练}$

6、确定性分类器

处理确定可分情况的分类器。通过几何方法将特征空间分解为对应不同类的子空间，又称为几何分类器。

7、感知器

感知器作为人工神经网络的基石，由于无法实现非线性分类而被冷落。但其“赏罚概念（reward-punishment concept）”得到广泛应用。

感知器准则 (Perceptron)

感知器算法 (perception approach)

问题提出:

两类线性可分的模式类: ω_1, ω_2 , 设 $d(X) = W^T X$ 其中, $W = [w_1, w_2, \dots, w_n, w_{n+1}]^T$, $X = [x_1, x_2, \dots, x_n, 1]^T$

应具有性质 $d(X) = W^T X \begin{cases} > 0, & \text{若 } X \in \omega_1 \\ < 0, & \text{若 } X \in \omega_2 \end{cases}$

对样本进行规范化处理, 即 ω_2 类样本全部乘以(-1), 则有:

$$d(X) = W^T X > 0$$

感知器算法通过对已知类别的训练样本集的学习, 寻找一个满足上式的权向量。

感知器准则 (Perceptron)

感知器算法 (perception approach)

步骤： (1) 选择 N 个分属于 ω_1 和 ω_2 类的模式样本构成训练样本集 $\{\mathbf{X}_1, \dots, \mathbf{X}_N\}$ 构成增广向量形式，并进行规范化处理。任取权向量初始值 $\mathbf{W}(1)$ ，开始迭代。迭代次数 $k=1$ 。

(2) 用全部训练样本进行一轮迭代，计算 $\mathbf{W}^T(k) \mathbf{X}_i$ 的值，并修正权向量。

分两种情况，更新权向量的值：

① 若 $\mathbf{W}^T(k) \mathbf{X}_i \leq 0$ ，分类器对第 i 个模式做了错误分类，权向量校正为： $\mathbf{W}(k+1) = \mathbf{W}(k) + \eta \mathbf{X}_i$

② 若 $\mathbf{W}^T(k) \mathbf{X}_i > 0$ ，分类正确，权向量不变： $\mathbf{W}(k+1) = \mathbf{W}(k)$ η ：正的校正增量，学习率。

统一写为：

$$\mathbf{W}(k+1) = \begin{cases} \mathbf{W}(k) & \text{若 } \mathbf{W}^T(k) \mathbf{X}_i > 0 \\ \mathbf{W}(k) + \eta \mathbf{X}_i & \text{若 } \mathbf{W}^T(k) \mathbf{X}_i \leq 0 \end{cases}$$

(3) 分析分类结果：只要有一个错误分类，回到 (2)，直至对所有样本正确分类。

感知器准则 (Perceptron)

感知器算法 (perception approach)

感知器算法是一种赏罚过程

分类正确时，对权向量“赏”——这里用“不罚”，即权向量不变；

分类错误时，对权向量“罚”——对其修改，向正确的方向转换。

感知器算法的收敛性

收敛性：经过算法的有限次迭代运算后，求出了一个使所有样本都能正确分类的 W ，则称算法是收敛的。

可以证明感知器算法是收敛的。收敛条件：模式类别线性可分。

理论证明，只要模式类是线性可分的，无论 a 的初值是什么，经过有限次叠代，都可收敛。

感知器准则 (Perceptron)

例 已知两类训练样本 $\omega_1: X_1=[0,0]^T, X_2=[0,1]^T$; $\omega_2: X_3=[1,0]^T, X_4=[1,1]^T$

用感知器算法求出将模式分为两类的权向量解和判别函数。

解: 所有样本写成增广向量形式; 进行规范化处理, 属于 ω_2 的样本乘以 (-1) 。

$$X_1=[0,0,1]^T \quad X_2=[0,1,1]^T \quad X_3=[-1,0,-1]^T \quad X_4=[-1,-1,-1]^T$$

任取 $W(1)=\mathbf{0}$, 取 $\eta=1$, 迭代过程为:

第一轮:

$$W^T(1)X_1=[0,0,0] \cdot \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix} = 0, \leq 0, \text{故 } W(2)=W(1)+X_1=[0,0,1]^T$$

$$W^T(3)X_3=[0,0,1] \cdot \begin{bmatrix} -1 \\ 0 \\ -1 \end{bmatrix} = -1, \leq 0, \text{故 } W(4)=W(3)+X_3=[-1,0,0]^T$$

$$W^T(2)X_2=[0,0,1] \cdot \begin{bmatrix} 0 \\ 1 \\ 1 \end{bmatrix} = 1, > 0, \text{故 } W(3)=W(2)=[0,0,1]^T$$

$$W^T(4)X_4=[-1,0,0] \cdot \begin{bmatrix} -1 \\ -1 \\ -1 \end{bmatrix} = 1, > 0, \text{故 } W(5)=W(4)=[-1,0,0]^T$$

有两个 $W^T(k)X_i \leq 0$ 的情况 (错判), 进行第二轮迭代。

感知器准则 (Perceptron)

$$X_1 = [0, 0, 1]^T \quad X_2 = [0, 1, 1]^T \quad X_3 = [-1, 0, -1]^T \quad X_4 = [-1, -1, -1]^T \quad W(5) = [-1, 0, 0]^T$$

第二轮:

$$W^T(5)X_1 = 0 \leq 0, \text{ 故 } W(6) = W(5) + X_1 = [-1, 0, 1]^T \quad W^T(6)X_2 = 1 > 0, \text{ 故 } W(7) = W(6) = [-1, 0, 1]^T$$
$$W^T(7)X_3 = 0 \leq 0, \text{ 故 } W(8) = W(7) + X_3 = [-2, 0, 0]^T \quad W^T(8)X_4 = 2 > 0, \text{ 故 } W(9) = W(8) = [-2, 0, 0]^T$$

第三轮:

$$W^T(9)X_1 = 0 \leq 0, \text{ 故 } W(10) = W(9) + X_1 = [-2, 0, 1]^T \quad W^T(10)X_2 = 1 > 0, \text{ 故 } W(11) = W(10)$$
$$W^T(11)X_3 = 1 > 0, \text{ 故 } W(12) = W(11) \quad W^T(12)X_4 = 1 > 0, \text{ 故 } W(13) = W(12)$$

第四轮:

$$W^T(13)X_1 = 1 > 0, \text{ 故 } W(14) = W(13) \quad W^T(14)X_2 = 1 > 0, \text{ 故 } W(15) = W(14)$$
$$W^T(15)X_3 = 1 > 0, \text{ 故 } W(16) = W(15) \quad W^T(16)X_4 = 1 > 0, \text{ 故 } W(17) = W(16)$$

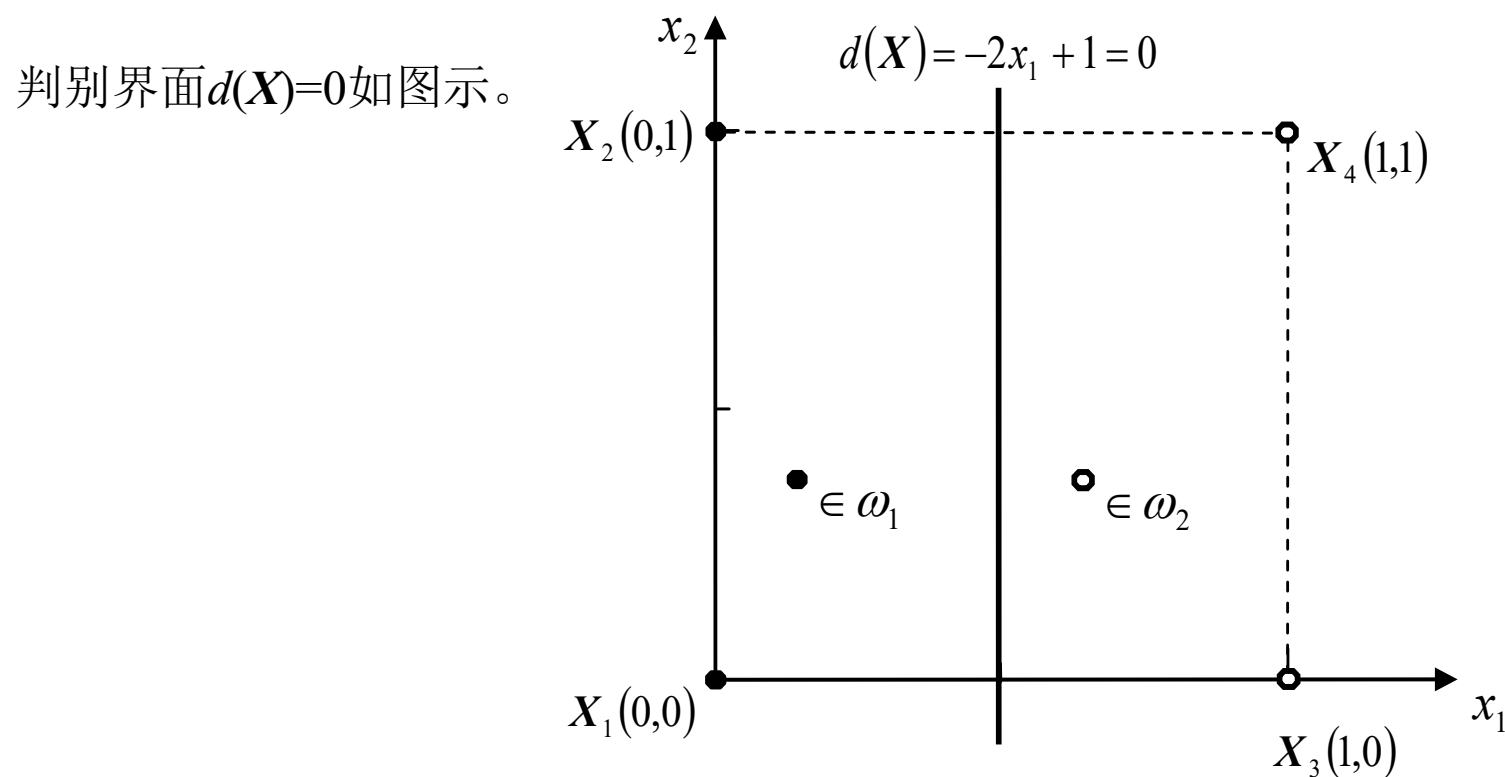
该轮迭代的分类结果全部正确, 故解向量相应的判别函数为:

$$d(X) = -2x_1 + 1$$

感知器准则 (Perceptron)

例 已知两类训练样本 $\omega_1: X_1=[0,0]^T, X_2=[0,1]^T$; $\omega_2: X_3=[1,0]^T, X_4=[1,1]^T$

用感知器算法求出将模式分为两类的权向量解和判别函数。

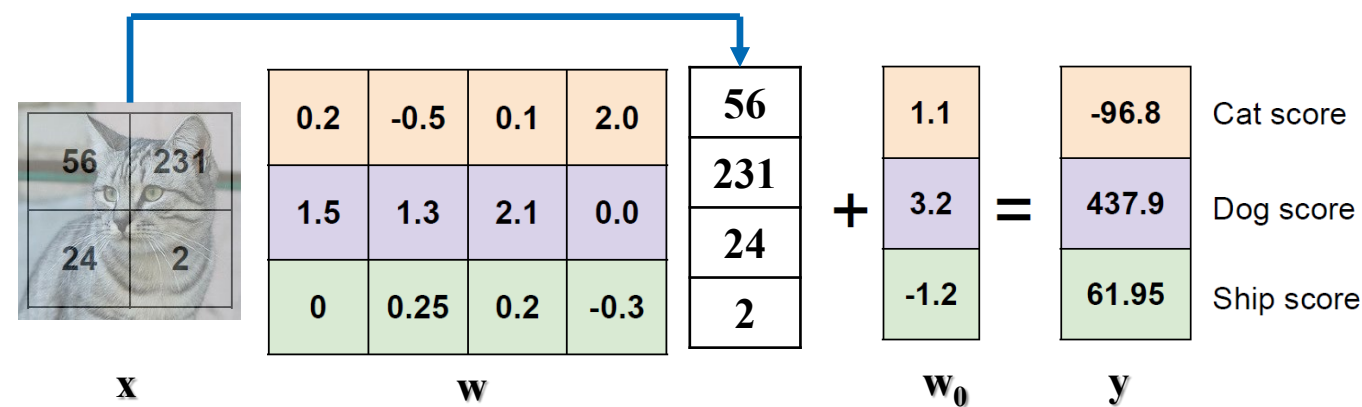


当 η 、 $W(1)$ 取其他值时，结果可能不一样，所以感知器算法的解不是单值的。

感知器准则 (Perceptron)

感知器算法用于多类情况（采用 ω_i/ω_j 两分法特例的多类分类方式）

对于 M 类模式应存在 M 个判决函数： $\{ d_i, i=1,...,M \}$ 。若 $X \in \omega_i$, 则 $d_i(X) > d_j(X), \forall j \neq i, j=1,...,M$



算法主要内容:

设有 M 种模式类别: $\omega_1, \omega_2, \dots, \omega_M$, 其权向量初值为: $W_j(1), j=1, \dots, M$

训练样本为增广向量形式, 但不需要规范化处理。

感知器准则 (Perceptron)

感知器算法用于多类情况

第 k 次迭代时，一个属于 ω_i 类的模式样本 $\mathbf{X}$ 被送入分类器，计算所有判别函数

$$d_j(k) = \mathbf{W}_j^T(k) \mathbf{X}, \quad j = 1, \dots, M$$

分二种情况修改权向量：

① 若 $d_i(k) > d_j(k)$, $\forall j \neq i; j = 1, 2, \dots, M$ 则权向量不变； $\mathbf{W}_j(k+1) = \mathbf{W}_j(k)$, $j = 1, 2, \dots, M$

② 若第 l 个权向量使 $d_i(k) \leq d_l(k)$ ，则相应的权向量作调整，即：

$$\begin{cases} \mathbf{W}_i(k+1) = \mathbf{W}_i(k) + \eta \mathbf{X} \\ \mathbf{W}_l(k+1) = \mathbf{W}_l(k) - \eta \mathbf{X} \\ \mathbf{W}_j(k+1) = \mathbf{W}_j(k), j \neq i, l \end{cases}, \quad \eta \text{ 为正的校正增量, 学习率}$$

可以证明：对于采用 ω_i/ω_j 两分法特例的多类分类方式，经过有限次迭代后算法收敛。

感知器准则 (Perceptron)

感知器算法用于多类情况

例 设有三个线性可分的模式类，三类的训练样本分别为 $\omega_1: \mathbf{X}_1 = [0,0]^T$; $\omega_2: \mathbf{X}_2 = [1,1]^T$; $\omega_3: \mathbf{X}_3 = [-1,1]^T$

现采用 ω_i/ω_j 两分法特例的多类分类方式，试用感知器算法求出判别函数。

解：增广向量形式： $\mathbf{X}_1 = [0,0,1]^T$, $\mathbf{X}_2 = [1,1,1]^T$, $\mathbf{X}_3 = [-1,1,1]^T$ 注意，这里任一类的样本都不能乘以 (-1) 。

任取初始权向量 $\mathbf{W}_1(1) = \mathbf{W}_2(1) = \mathbf{W}_3(1) = [0,0,0]^T$; $\eta = 1$

第一次迭代： $d_1(1) = \mathbf{W}_1^T(1)\mathbf{X}_1 = 0$ $d_2(1) = \mathbf{W}_2^T(1)\mathbf{X}_1 = 0$ $d_3(1) = \mathbf{W}_3^T(1)\mathbf{X}_1 = 0$

$\mathbf{X}_1 \in \omega_1$, 但 $d_1(1) > d_2(1)$ 且 $d_1(1) > d_3(1)$ 不成立，因此三个权向量都需要修改：

$$\mathbf{W}_1(2) = \mathbf{W}_1(1) + \mathbf{X}_1 = [0,0,1]^T$$

$$\mathbf{W}_2(2) = \mathbf{W}_2(1) - \mathbf{X}_1 = [0,0,-1]^T$$

$$\mathbf{W}_3(2) = \mathbf{W}_3(1) - \mathbf{X}_1 = [0,0,-1]^T$$

感知器准则 (Perceptron)

感知器算法用于多类情况

例 设有三个线性可分的模式类，三类的训练样本分别为 $\omega_1: \mathbf{X}_1 = [0,0]^T$; $\omega_2: \mathbf{X}_2 = [1,1]^T$; $\omega_3: \mathbf{X}_3 = [-1,1]^T$
现采用 ω_i/ω_j 两分法特例的多类分类方式，试用感知器算法求出判别函数。

$$\text{第二次迭代: } d_1(2) = \mathbf{W}_1^T(2)\mathbf{X}_2 = 1 \quad d_2(2) = \mathbf{W}_2^T(2)\mathbf{X}_2 = -1 \quad d_3(2) = \mathbf{W}_3^T(2)\mathbf{X}_2 = -1$$

$\mathbf{X}_2 \in \omega_2$, 但 $d_2(2) > d_1(2)$ 且 $d_2(2) > d_3(2)$ 不成立, 修改权向量:

$$\mathbf{W}_1(3) = \mathbf{W}_1(2) - \mathbf{X}_2 = [-1, -1, 0]^T$$

$$\mathbf{W}_2(3) = \mathbf{W}_2(2) + \mathbf{X}_2 = [1, 1, 0]^T$$

$$\mathbf{W}_3(3) = \mathbf{W}_3(2) - \mathbf{X}_2 = [-1, -1, -2]^T$$

第三次迭代: $\mathbf{X}_3 \in \omega_3$, 但 $d_3(3) > d_1(3)$ 且 $d_3(3) > d_2(3)$ 不成立, 修改为权向量。

以上进行的一轮迭代运算中, 三个样本都未正确分类, 进行下一轮迭代。

感知器准则 (Perceptron)

感知器算法用于多类情况

例 设有三个线性可分的模式类，三类的训练样本分别为 $\omega_1: \mathbf{X}_1 = [0,0]^T$; $\omega_2: \mathbf{X}_2 = [1,1]^T$; $\omega_3: \mathbf{X}_3 = [-1,1]^T$
现采用 ω_i/ω_j 两分法特例的多类分类方式，试用感知器算法求出判别函数。

第四次迭代:

在第五、六、七迭代中，对所有三个样本都已正确分类。权向量的解：

$$\mathbf{W}_1 = \mathbf{W}_1(7) = \mathbf{W}_1(6) = \mathbf{W}_1(5) = [0, -2, 0]^T$$

$$\mathbf{W}_2 = \mathbf{W}_2(7) = \mathbf{W}_2(6) = \mathbf{W}_2(5) = [2, 0, -2]^T$$

$$\mathbf{W}_3 = \mathbf{W}_3(7) = \mathbf{W}_3(6) = \mathbf{W}_3(5) = [-2, 0, -2]^T$$

判别函数: $d_1(\mathbf{X}) = -2x_2$ $d_2(\mathbf{X}) = 2x_1 - 2$ $d_3(\mathbf{X}) = -2x_1 - 2$

梯度法

梯度下降——数学基础知识

导数 Derivation

设函数 $y = f(x)$ 在 $U(x_0, \delta)$ 内有定义,

如果极限 $\lim_{\Delta x \rightarrow 0} \frac{\Delta y}{\Delta x} = \lim_{\Delta x \rightarrow 0} \frac{f(x_0 + \Delta x) - f(x_0)}{\Delta x}$ 存在,

则称这个极限值为 $y = f(x)$ 在点 x_0 处的导数。

记为 $y'|_{x=x_0}$

导数的意义:

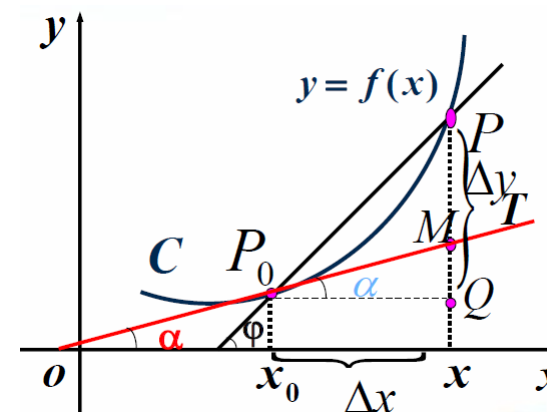
1、变化率、变化速度

反映变量 y 在点 x_0 处的变化率。

或是 y 随自变量 x 变化而变化的快慢程度。

如: 速度 $v = \Delta d / \Delta t$, 加速度 $a = \Delta v / \Delta t$

2、导数的几何意义——曲线上某一点处的切线斜率



$$\text{切线斜率 } \tan \alpha = \lim_{\varphi \rightarrow \alpha} \tan \varphi = \lim_{\Delta x \rightarrow 0} \frac{\Delta y}{\Delta x} = f'(x_0)$$

梯度法

梯度下降——数学基础知识

偏导数 Partial Derivation

二元偏导数定义：设二元函数 $z = f(x, y)$ 在 $P_0(x_0, y_0)$ 的某邻域内有定义，

如果极限 $\lim_{\Delta x \rightarrow 0} \frac{\Delta y}{\Delta x} = \lim_{\Delta x \rightarrow 0} \frac{f(x_0 + \Delta x, y_0) - f(x_0, y_0)}{\Delta x}$ 存在，

则称此极限值为函数 f 在 P_0 处对于 x 的偏导数，记为 $\left. \frac{\partial f}{\partial x} \right|_{P_0}$ 。

同理有对于 y 的偏导数

$$\left. \frac{\partial f}{\partial y} \right|_{P_0} = \lim_{\Delta y \rightarrow 0} \frac{f(x_0, y_0 + \Delta y) - f(x_0, y_0)}{\Delta y}$$

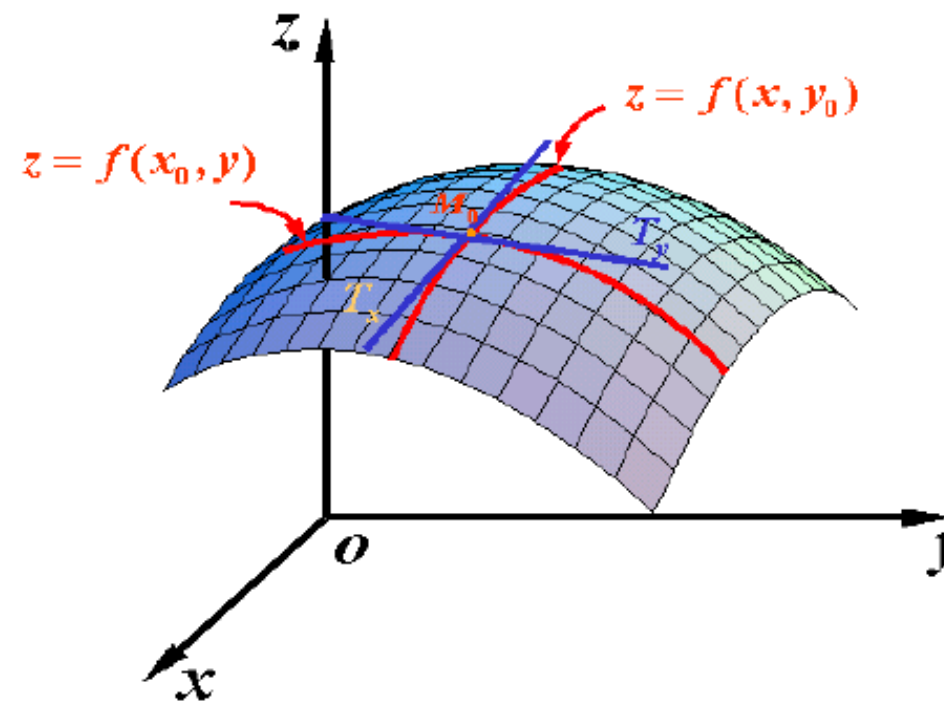
梯度法

梯度下降——数学基础知识

偏导数的几何意义

偏导数 $f'_x(x_0, y_0)$ 就是曲线 $z = f(x, y_0)$ 在 x_0 处的 x 方向切线斜率

偏导数 $f'_y(x_0, y_0)$ 就是曲线 $z = f(x_0, y)$ 在 y_0 处的 y 方向切线斜率



梯度法

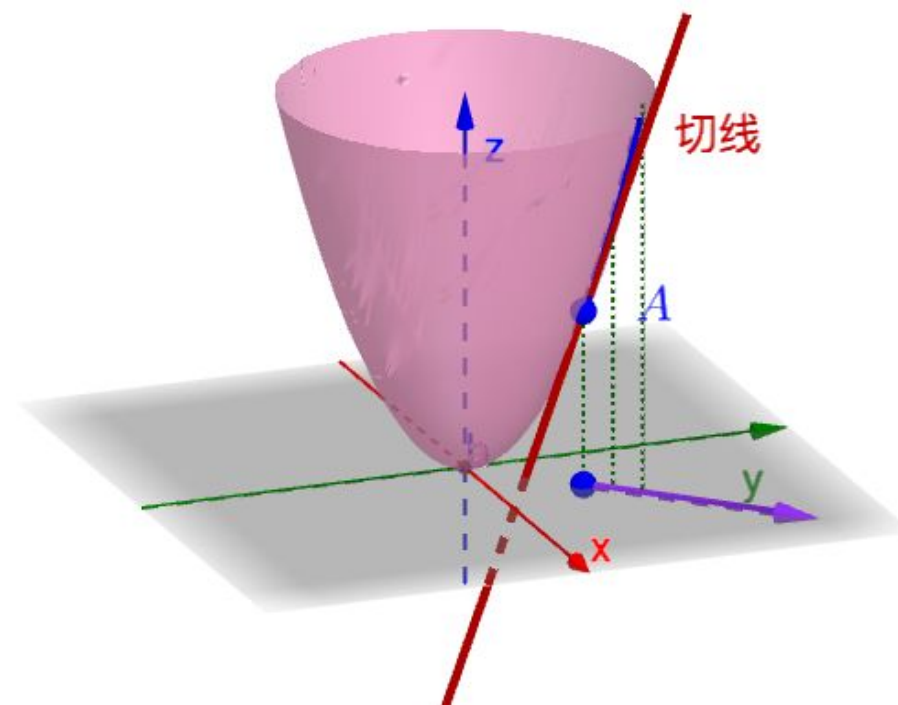
梯度下降——数学基础知识

方向导数 Directional Derivative

设 f 是定义于 R^n 中某区域 D 上的函数，点 $P_0 \in D$ ， l 为一给定的非零向量， P 为一动点，向量 P_0P 与 l 的方向始终一致，如果极限

$$\lim_{\|P_0P\| \rightarrow 0} \frac{f(P) - f(P_0)}{\|P_0P\|}$$

存在，则称此极限为函数 f 在 P_0 处沿 l 方向的方向导数，记为 $\left. \frac{\partial f}{\partial l} \right|_{P_0}$ 。



梯度法

梯度下降——数学基础知识

梯度 Gradient

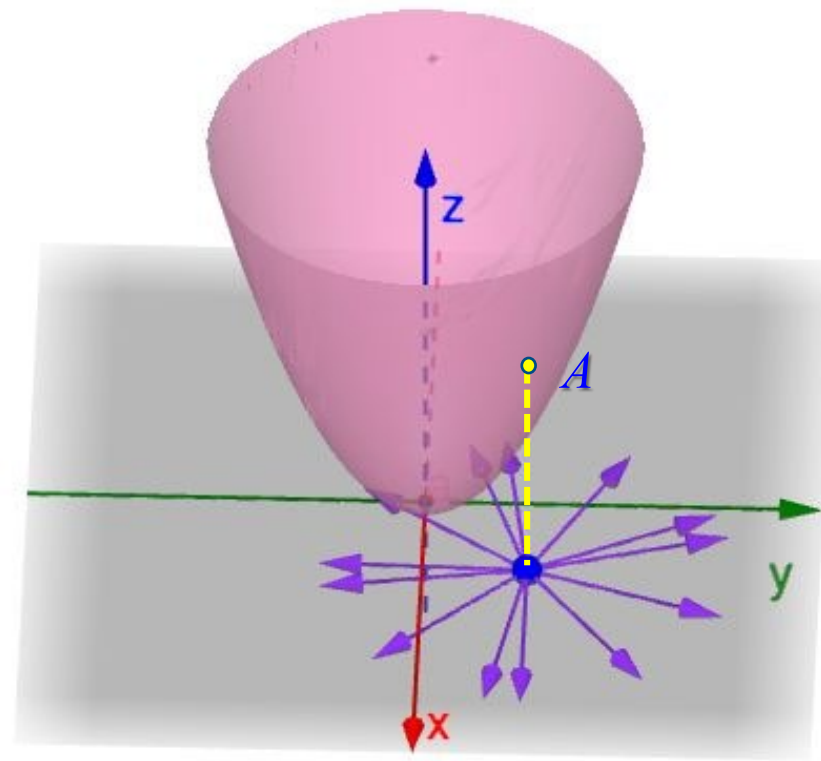
显然A点不止一个方向导数，而是360度都有方向导数。

梯度的几何意义：是一个矢量，其方向上的方向导数最大，其大小正好是此最大方向导数。

二元函数梯度定义：设二元函数 $z = f(x, y)$ 在 $P_0(x_0, y_0)$ 点存在偏导数，该点处梯度为： $\nabla G = \left(\frac{\partial f}{\partial x} \Big|_{P_0}, \frac{\partial f}{\partial y} \Big|_{P_0} \right)$

梯度方向为： $\theta = \tan^{-1} \left(\frac{\partial f}{\partial y} / \frac{\partial f}{\partial x} \right)$

梯度幅值为： $\|\nabla f\| = \sqrt{\left(\frac{\partial f}{\partial x} \right)^2 + \left(\frac{\partial f}{\partial y} \right)^2}$



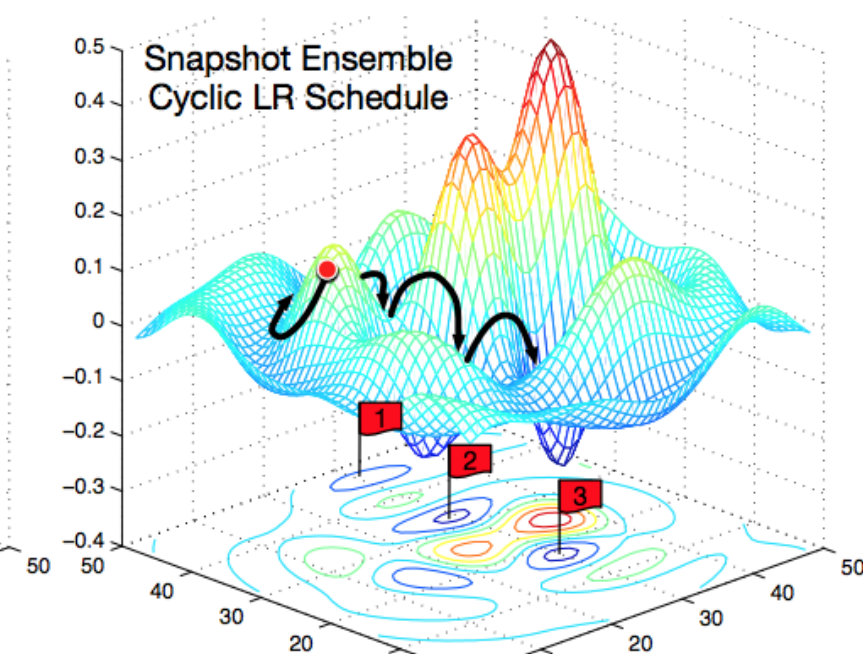
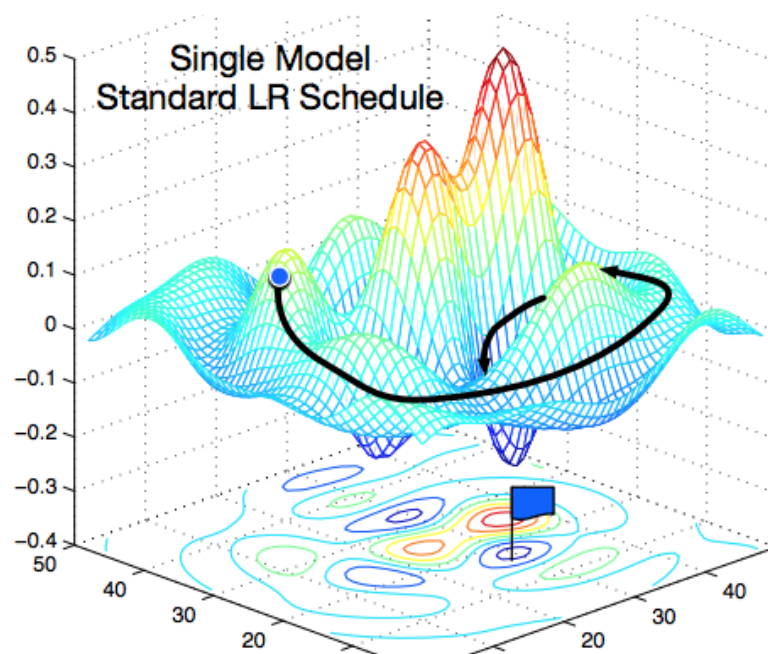
梯度法

梯度下降——数学基础知识

梯度 Gradient

设 f 是 R_n 中区域 D 上的数量场, 如果 f 在 $P_0 \in D$ 处可微, 称向量 $\left(\frac{\partial f}{\partial x_1}, \frac{\partial f}{\partial x_2}, \dots, \frac{\partial f}{\partial x_n} \right) \bigg|_{P_0}$ 为 f 在 P_0 处的梯度。

梯度下降



梯度（下降）法

1、梯度概念小结

设函数 $f(\mathbf{Y})$ 是向量 $\mathbf{Y} = [y_1, y_2, \dots, y_n]^T$ 的函数，则 $f(\mathbf{Y})$ 的梯度定义为：

$$\nabla f(\mathbf{Y}) = \frac{d}{d\mathbf{Y}} f(\mathbf{Y}) = \left[\frac{\partial f}{\partial y_1}, \frac{\partial f}{\partial y_2}, \dots, \frac{\partial f}{\partial y_n} \right]^T$$

梯度向量的最重要性质之一：指出函数 f 在其自变量增加时，增长最快的方向。

即：梯度的**方向**是函数 $f(\mathbf{Y})$ 在 $\mathbf{Y}$ 点增长最快的方向，

梯度的**模**是 $f(\mathbf{Y})$ 在增长最快的方向上的增长率 (增长率最大值)。

梯度（下降）法

梯度（下降）法基本原理

2、梯度算法 设两个线性可分的模式类 ω_1 和 ω_2 ，共 N 个样本， ω_2 类样本乘(-1)。

将两类样本分开的判决函数 $d(\mathbf{X})$ 应满足： $d(\mathbf{X}_i) = \mathbf{W}^T \mathbf{X}_i > 0 \quad i = 1, 2, \dots, N$ —— N 个不等式

梯度算法的目的是求一个满足上述条件的权向量，即联立不等式求解 $\mathbf{W}$ 的问题，转换成求准则函数极小值的问题。

用负梯度向量的值对权向量 $\mathbf{W}$ 进行修正，实现使准则函数（损失函数）达到极小值的目的。

准则函数（损失函数）的选取原则：具有唯一的最小值，并且这个最小值发生在 $\mathbf{W}^T \mathbf{X}_i > 0$ 时。

基本思路：

定义一个对错误分类敏感的准则函数 $J(\mathbf{W}, \mathbf{X})$ ，在 J 的梯度方向上对权向量进行修改。一般关系

表示成从 $\mathbf{W}(k)$ 导出 $\mathbf{W}(k+1)$ ： $\mathbf{W}(k+1) = \mathbf{W}(k) + \eta(-\nabla J) = \mathbf{W}(k) - \eta \nabla J$

具体为：

$$\mathbf{W}(k+1) = \mathbf{W}(k) - \eta \left[\frac{\partial J(\mathbf{W}, \mathbf{X})}{\partial \mathbf{W}} \right]_{\mathbf{W}=\mathbf{W}(k)}$$

梯度法

梯度法步骤:

(1) 将样本写成规范化增广向量形式, 选择准则函数, 设置初始权向量 $\mathbf{W}(1)$, 括号内为迭代次数 $k=1$ 。

(2) 依次输入训练样本 $\mathbf{X}$ 。设第 k 次迭代时输入样本为 $\mathbf{X}_i$, 此时已有权向量 $\mathbf{W}(k)$, 求 $\nabla J(k)$:

$$\nabla J(k) = \left. \frac{\partial J(\mathbf{W}, \mathbf{X}_i)}{\partial \mathbf{W}} \right|_{\mathbf{W}=\mathbf{W}(k)}$$

权向量修正为: $\mathbf{W}(k+1) = \mathbf{W}(k) - \eta \nabla J(k)$

迭代次数 k 加1, 输入下一个训练样本, 计算新的权向量, 直至对全部训练样本完成一轮迭代。

(3) 在下一轮迭代中, 如果有样本使 $\nabla J \neq 0$, 则回到(2)再次迭代。直至 $\mathbf{W}$ 不再变化, 算法收敛。

随着权向量 $\mathbf{W}$ 向理想值接近, 准则函数关于 $\mathbf{W}$ 的导数 $\nabla J(k)$ 越来越趋近于零, 这意味着准则函数 J 越来越接近最小值。当 最终 $\nabla J = 0$ 时, J 达到最小值, 此时 $\mathbf{W}$ 不再改变, 算法收敛。

梯度法

梯度算法是求解权向量的一般解法，算法的具体计算形式取决于准则函数 $J(W, X)$ 的选择， $J(W, X)$ 的形式不同，得到的具体算法不同。

感知器算法是梯度法的特例。

固定增量法

准则函数： $J(W, X) = \frac{1}{2} (|W^T X| - W^T X)$ 该函数有唯一最小值“0”，且发生在 $W^T X > 0$ 的时候，求 $W(k)$ 的递推公式：

$$\text{设 } X = [x_1, x_2, \dots, x_n, 1]^T, \quad W = [w_1, w_2, \dots, w_n, w_{n+1}]^T$$

$$\text{步骤1: 求 } J \text{ 的梯度 } \nabla J = \frac{\partial J(W, X)}{\partial W} = ?$$

步骤2: 求 $W(k+1)$

梯度法

固定增量法

$$J(W, X) = \frac{1}{2} (|W^T X| - W^T X)$$

步骤1: 求 J 的梯度 $\nabla J = \frac{\partial J(W, X)}{\partial W} = ?$ 方法: 函数对向量求导=函数对向量的分量求导, 即 $\frac{\partial f}{\partial \mathbf{X}} = \left[\frac{\partial f}{\partial x_1}, \dots, \frac{\partial f}{\partial x_n} \right]^T$

① 首先求 $W^T X$ 部分:

$$\frac{\partial (W^T X)}{\partial W} = \frac{\partial}{\partial W} \left(\sum_{i=1}^n w_i x_i + w_{n+1} \right) = \left[\frac{\partial}{\partial w_1} \left(\sum_{i=1}^n w_i x_i + w_{n+1} \right), \dots, \frac{\partial}{\partial w_k} \left(\sum_{i=1}^n w_i x_i + w_{n+1} \right), \dots, \frac{\partial}{\partial w_{n+1}} \left(\sum_{i=1}^n w_i x_i + w_{n+1} \right) \right]^T = [x_1, \dots, x_k, \dots, x_n, 1]^T = X$$

② 由①的结论 $\frac{\partial (W^T X)}{\partial W} = X$ 有:

$$\begin{cases} W^T X > 0 \text{ 时}, \frac{\partial (|W^T X|)}{\partial W} = \frac{\partial (W^T X)}{\partial W} = X \\ W^T X \leq 0 \text{ 时}, \frac{\partial (|W^T X|)}{\partial W} = \frac{\partial (-W^T X)}{\partial W} = -X \end{cases}$$

③ $\nabla J = \frac{\partial J(W, X)}{\partial W}$ 有:

$$\therefore \nabla J = \frac{\partial J(W, X)}{\partial W} = \begin{cases} 0 & \text{当 } W^T X > 0 \text{ 时} \\ -X & \text{当 } W^T X \leq 0 \text{ 时} \end{cases}$$

梯度法

固定增量法

步骤2: 求 $W(k+1)$

$$W(k+1) = W(k) - \eta \nabla J = W(k) - \eta \left[\frac{\partial J(W, X)}{\partial W} \right]_{W=W(k)} = W(k) + \begin{cases} 0, & \text{若 } W^T(k)X > 0 \\ \eta X, & \text{若 } W^T(k)X \leq 0 \end{cases}$$

上式即为固定增量算法，与感知器算法形式相同。由此可以看出，感知器算法是梯度法的特例。

即：梯度法是将感知器算法中联立不等式求解 W 的问题，转换为求函数 J 极小值的问题，将原来有多个解的情况，变成求最优解的情况。

只要模式类是线性可分的，算法就会给出解。

最小平方误差算法 (Least Mean Square Error, LMSE; 亦称Ho-Kashyap算法)

感知器算法、梯度算法等类似方法，**只有当模式类可分离时才收敛**，在不可分的情况下，算法会来回摆动，始终不收敛。当一次次迭代而又不见收敛时，造成不收敛现象的原因分不清，有两种可能：

- a) 迭代过程本身收敛缓慢
- b) 模式本身不可分

LMSE算法特点：

对可分模式收敛。

对于类别不可分的情况也能指出来。

最小平方误差算法

分类器的不等式方程

两类分类问题的解相当于求一组线性不等式的解。如果给出分属于 ω_1 和 ω_2 ，两个模式类的训练样本集（规范化增广样本向量） $\{\mathbf{X}_i, i=1,2,\dots,N\}$ ，应满足： $\mathbf{W}^T \mathbf{X}_i > 0$

上式展开写为：

$$\begin{array}{ll} \omega_1 \text{类} & \begin{cases} w_1 x_{11} + w_2 x_{12} + \dots + w_n x_{1n} + w_{n+1} > 0 & \text{对 } \mathbf{X}_1 \\ w_1 x_{21} + w_2 x_{22} + \dots + w_n x_{2n} + w_{n+1} > 0 & \text{对 } \mathbf{X}_2 \\ \dots & \dots \end{cases} \\ \omega_2 \text{类} & \begin{cases} -w_1 x_{N1} - w_2 x_{N2} - \dots - w_n x_{Nn} - w_{n+1} > 0 & \text{对 } \mathbf{X}_N \end{cases} \end{array}$$

写成矩阵形式为：

$$\begin{bmatrix} x_{11} & x_{12} & \dots & x_{1n} & 1 \\ x_{21} & x_{22} & \dots & x_{2n} & 1 \\ \dots & \dots & \dots & \dots & \dots \\ -x_{N1} & -x_{N2} & \dots & -x_{Nn} & -1 \end{bmatrix}_{N \times (n+1)} \begin{bmatrix} w_1 \\ w_2 \\ \dots \\ w_n \\ w_{n+1} \end{bmatrix}_{(n+1) \times 1} > \begin{bmatrix} 0 \\ 0 \\ \dots \\ 0 \\ 0 \end{bmatrix}_{N \times 1} = \mathbf{0}$$

改写为 $\Rightarrow \mathbf{XW} > \mathbf{0}$

通过解不等式组 $\mathbf{XW} > \mathbf{0}$ ，求出 $\mathbf{W}$

最小平方误差算法

LMSE算法原理

LMSE算法把对满足 $\mathbf{XW} > \mathbf{0}$ 的求解，改为满足 $\mathbf{XW} = \mathbf{B}$ 的求解。

式中： $\mathbf{B} = [b_1, b_2, \dots, b_i, \dots, b_N]^T$ 为各分量均为正值的矢量。 $\therefore$ 两式等价。

说明：

① 当方程组中当行数 $\gg$ 列数时，通常无精确解，称为矛盾方程组（或超定方程组），一般求近似解。

在模式识别中，通常训练样本数 N 通常总是大于模式的维数 n ，

因此方程的个数(行数) $\gg$ 模式向量的维数(列数) 是矛盾方程组，只能求近似解 $\mathbf{W}^*$ ，即

$$\|\mathbf{XW}^* - \mathbf{B}\| = \text{极小}$$

② LMSE算法的出发点：选择一个准则函数，使得当 J 达到最小值时，

$\mathbf{XW} - \mathbf{B}$ 可得到近似解（最小二乘近似解）。

最小平方误差算法

LMSE算法原理

准则函数定义为: $J(W, X, B) = \frac{1}{2} \|XW - B\|^2$

“最小二乘”:

—— 最小: 使方程组两边误差最小, 也即使 J 最小。

—— 二乘: 次数为2, 乘了两次

} 最小平方 (误差算法)

③ LMSE算法的思路: 对 $XW > 0$ 求解



转化为

对 $XW = B$ 求解



转化为

通过求准则函数极小找 W 、 B

最小平方误差算法

LMSE算法原理

考察向量($XW-B$) 有:

$$XW - B = \begin{bmatrix} x_{11} & \dots & x_{1n} & 1 \\ x_{21} & \dots & x_{2n} & 1 \\ \dots & \dots & \dots & 1 \\ x_{i1} & \dots & x_{in} & 1 \\ \dots & \dots & \dots & 1 \\ -x_{N1} & \dots & -x_{Nn} & 1 \end{bmatrix}_{N \times (n+1)} \begin{bmatrix} w_1 \\ w_2 \\ \dots \\ w_n \\ w_{n+1} \end{bmatrix}_{(n+1) \times 1} - \begin{bmatrix} b_1 \\ \dots \\ b_i \\ \dots \\ b_N \end{bmatrix}_{N \times 1} = \begin{bmatrix} x_{11}w_1 + \dots + x_{1n}w_n + w_{n+1} - b_1 \\ \dots \\ x_{i1}w_1 + \dots + x_{in}w_n + w_{n+1} - b_i \\ \dots \\ -x_{N1}w_1 - \dots - x_{Nn}w_n - w_{n+1} - b_N \end{bmatrix}_{N \times 1} = \begin{bmatrix} W^T X_1 - b_1 \\ \dots \\ W^T X_i - b_i \\ \dots \\ W^T X_N - b_N \end{bmatrix}$$

$$\|XW - B\|^2 = \left(\sqrt{\text{向量各分量的平方和}} \right)^2 = \text{向量各分量的平方和}$$

$$\|XW - B\|^2 = (W^T X_1 - b_1)^2 + \dots + (W^T X_N - b_N)^2 = \sum_{i=1}^N (W^T X_i - b_i)^2$$

最小平方误差算法

LMSE算法原理

准则函数:
$$J(\mathbf{W}, \mathbf{X}, \mathbf{B}) = \frac{1}{2} \|\mathbf{XW} - \mathbf{B}\|^2 = \frac{1}{2} \sum_{i=1}^N (\mathbf{W}^T \mathbf{X}_i - b_i)^2$$

$\mathbf{XW}=\mathbf{B}$ 的近似解也称“最优近似解”——使方程组两边所有误差之和最小（即最优）的解。

可以看出:

- ① 当函数 J 达到最小值, 等式 $\mathbf{XW}=\mathbf{B}$ 有最优解。即又将问题转化为求准则函数极小值的问题。
- ② 因为 J 有两个变量 $\mathbf{W}$ 和 $\mathbf{B}$, 有更多的自由度供选择求解, 故可望改善算法的收敛速率。

最小平方误差算法

LMSE递推公式

与问题相关的两个梯度：
$$\frac{\partial J}{\partial \mathbf{W}} = \mathbf{X}^T (\mathbf{X}\mathbf{W} - \mathbf{B}) \quad \frac{\partial J}{\partial \mathbf{B}} = -\frac{1}{2} [(\mathbf{X}\mathbf{W} - \mathbf{B}) + |\mathbf{X}\mathbf{W} - \mathbf{B}|]$$

求递推公式：

(1) 求 $\mathbf{W}$ 的递推关系

使 J 对 $\mathbf{W}$ 求最小，令 $\frac{\partial J}{\partial \mathbf{W}} = 0$ ，得：
$$\mathbf{X}^T (\mathbf{X}\mathbf{W} - \mathbf{B}) = 0 \Rightarrow \mathbf{X}^T \mathbf{X}\mathbf{W} = \mathbf{X}^T \mathbf{B} \Rightarrow \mathbf{W} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{B} = \mathbf{X}^\# \mathbf{B}$$

式中： $\mathbf{X}^\# = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T$ 称为 $\mathbf{X}$ 的伪逆，

$\mathbf{X}$ 为 $N \times (n+1)$ 长方阵， $\mathbf{X}^\#$ 为 $(n+1) \times N$ 长方阵。

可知：只要求出 $\mathbf{B}$ ，就可求出 $\mathbf{W}$ 。

最小平方误差算法

LMSE递推公式

(2) 求 $\mathbf{B}(k+1)$ 的迭代式

利用梯度算法公式 $\mathbf{W}(k+1) = \mathbf{W}(k) - \eta \left[\frac{\partial J(\mathbf{W}, \mathbf{X})}{\partial \mathbf{W}} \right]_{\mathbf{W}=\mathbf{W}(k)}$ 有:

$$\mathbf{B}(k+1) = \mathbf{B}(k) - \eta' \left[\frac{\partial J}{\partial \mathbf{B}} \right]_{\mathbf{B}=\mathbf{B}(k)}$$

代入 $\frac{\partial J}{\partial \mathbf{B}} = -\frac{1}{2}[(\mathbf{XW} - \mathbf{B}) + |\mathbf{XW} - \mathbf{B}|]$, 有 $\mathbf{B}(k+1) = \mathbf{B}(k) + \frac{\eta'}{2}[(\mathbf{XW}(k) - \mathbf{B}(k)) + |\mathbf{XW}(k) - \mathbf{B}(k)|]$

令 $\eta = \eta'/2$, 定义 $\mathbf{e}(k) = \mathbf{XW}(k) - \mathbf{B}(k)$, 有:

$$\mathbf{B}(k+1) = \mathbf{B}(k) + \eta [\mathbf{e}(k) + |\mathbf{e}(k)|]$$

最小平方误差算法

LMSE递推公式

(3) 求 $W(k+1)$ 的迭代式

$$W = X^{\#} B \quad B(k+1) = B(k) + \eta [e(k) + |e(k)|]$$

$$\begin{aligned} W(k+1) &= X^{\#} B(k+1) = X^{\#} \{B(k) + \eta [e(k) + |e(k)|]\} \\ &= X^{\#} B(k) + X^{\#} \eta e(k) + X^{\#} \eta |e(k)| \\ &= W(k) + \eta X^{\#} |e(k)| \end{aligned}$$

$=0$

$$\begin{aligned} X^{\#} e(k) &= (X^T X)^{-1} X^T [XW(k) - B(k)] \\ &= (X^T X)^{-1} X^T [XX^{\#} B(k) - B(k)] \\ &= 0 \end{aligned}$$

最小平方误差算法

LMSE递推公式总结 设初值 $\mathbf{B}(1)$ ，各分量均为正值，括号中数字代表迭代次数。

$$\mathbf{W}(1) = \mathbf{X}^\# \mathbf{B}(1)$$

.....

$$\mathbf{e}(k) = \mathbf{X}\mathbf{W}(k) - \mathbf{B}(k)$$

$$\mathbf{W}(k+1) = \mathbf{W}(k) + \eta \mathbf{X}^\# |\mathbf{e}(k)|$$

$$\mathbf{B}(k+1) = \mathbf{B}(k) + \eta [\mathbf{e}(k) + |\mathbf{e}(k)|]$$

$\mathbf{W}(k+1)$ 、 $\mathbf{B}(k+1)$ 互相独立，先后次序无关。

或另一算法：先算 $\mathbf{B}(k+1)$ ，再算 $\mathbf{W}(k+1)$ 。

$$\mathbf{W}(1) = \mathbf{X}^\# \mathbf{B}(1)$$

.....

$$\mathbf{e}(k) = \mathbf{X}\mathbf{W}(k) - \mathbf{B}(k)$$

$$\mathbf{B}(k+1) = \mathbf{B}(k) + \eta [\mathbf{e}(k) + |\mathbf{e}(k)|]$$

$$\mathbf{W}(k+1) = \mathbf{X}^\# \mathbf{B}(k+1)$$

求出 $\mathbf{B}$ ， $\mathbf{W}$ 后，再迭代出下一个 $\mathbf{e}$ ，

从而计算出新的 $\mathbf{B}$ ， $\mathbf{W}$ 。

最小平方误差算法

3) 模式类别可分性判别

可以证明：当模式类线性可分，且学习率 η 满足 $0 < \eta \leq 1$ 时，该算法收敛，可求得解 W 。

理论上不能证明该算法到底需要迭代多少步才能达到收敛，通常在每次迭代计算后检查一下 $XW(k)$ 和误差向量 $e(k)$ ，从而可以判断是否收敛。

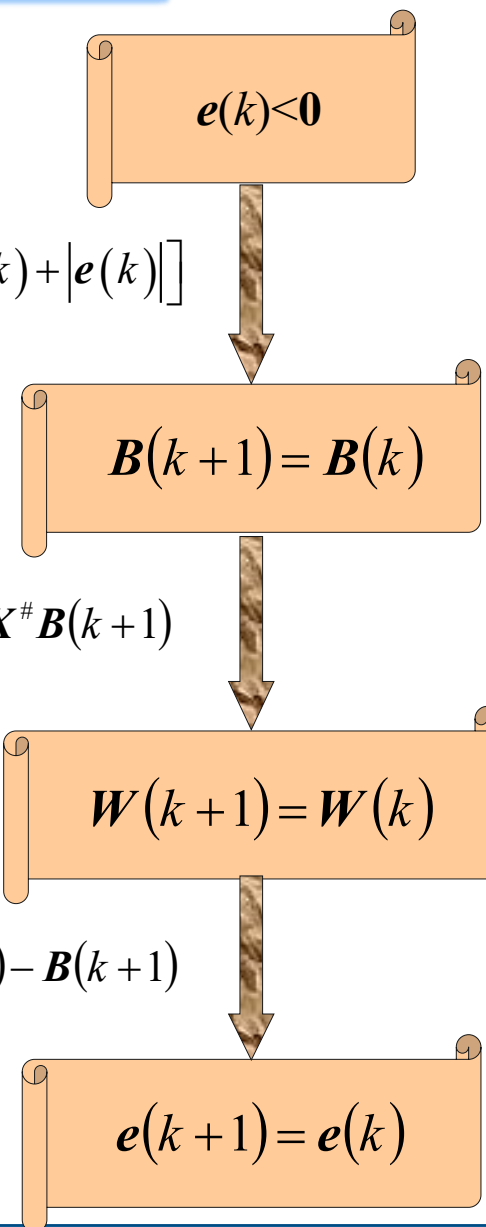
分以下几种情况：

- ① 如果 $e(k)=0$ ，表明 $XW(k)=B(k) > 0$ ，有解。
- ② 如果 $e(k)>0$ ，表明 $XW(k)>B(k) > 0$ ，隐含有解。
继续迭代，可使 $e(k) \rightarrow 0$ 。
- ③ 如果 $e(k)<0$ （所有分量为负数或零，但不全为零），
停止迭代，无解。此时若继续迭代，数据不再发生变化。

$$B(k+1) = B(k) + \eta [e(k) + |e(k)|]$$

$$W(k+1) = X^{\#} B(k+1)$$

$$e(k+1) = XW(k+1) - B(k+1)$$



最小平方误差算法

综上所述：

- 1、只有当 $\mathbf{e}(k)$ 中有大于零的分量时，才需要继续迭代，一旦 $\mathbf{e}(k)$ 的全部分量只有0和负数，则立即停止。事实上，往往早在 $\mathbf{e}(k)$ 全部分量都达到非正值以前，就能看出其中有些分量向正值变化得极慢，可及早采取对策。
- 2、通过反证法可以证明：在线性可分情况下，算法进行过程中不会出现 $\mathbf{e}(k)$ 的分量全为负的情况；若出现 $\mathbf{e}(k)$ 的分量全为负，则说明模式类线性不可分。

最小平方误差算法

4) LMSE算法描述

(1) 根据 N 个分属于两类的样本，写出规范化增广样本矩阵 $\mathbf{X}$ 。

(2) 求 $\mathbf{X}$ 的伪逆矩阵 $\mathbf{X}^\# = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top$ 。

(3) 设置初值 η 和 $\mathbf{B}(1)$ ， η 为正的校正增量， $\mathbf{B}(1)$ 的各分量大于零，迭代次数 $k=1$ 。开始迭代：

计算 $\mathbf{W}(1) = \mathbf{X}^\# \mathbf{B}(1)$

.....

(4) 计算 $\mathbf{e}(k) = \mathbf{X}\mathbf{W}(k) - \mathbf{B}(k)$ ，进行可分性判别。

如果 $\mathbf{e}(k)=\mathbf{0}$ ，线性可分，解为 $\mathbf{W}(k)$ ，算法结束。

如果 $\mathbf{e}(k)>\mathbf{0}$ ，线性可分，若进入(5)可使 $\mathbf{e}(k) \rightarrow \mathbf{0}$ ，得最优解。

如果 $\mathbf{e}(k)<\mathbf{0}$ ，线性不可分，停止迭代，无解，算法结束。

否则，说明 $\mathbf{e}(k)$ 的各分量值有正有负，进入(5)。

最小平方误差算法

4) LMSE算法描述

- (5) 计算 $\mathbf{W}(k+1)$ 和 $\mathbf{B}(k+1)$ 。
- 方法1: 分别计算 $\mathbf{W}(k+1) = \mathbf{W}(k) + c\mathbf{X}^\#|\mathbf{e}(k)|$ $\mathbf{B}(k+1) = \mathbf{B}(k) + c[\mathbf{e}(k) + |\mathbf{e}(k)|]$
- 方法2: 先计算 $\mathbf{B}(k+1) = \mathbf{B}(k) + c[\mathbf{e}(k) + |\mathbf{e}(k)|]$ 再计算 $\mathbf{W}(k+1) = \mathbf{X}^\# \mathbf{B}(k+1)$

迭代次数 k 加1, 返回(4)。

3. 算法特点

- (1) 算法尽管略为复杂一些, 但提供了线性可分的测试特征。
- (2) 同时利用 N 个训练样本, 同时修改 $\mathbf{W}$ 和 $\mathbf{B}$, 故收敛速度快。
- (3) 计算矩阵 $(\mathbf{X}^\top \mathbf{X})^{-1}$ 复杂, 但可用迭代算法计算。

最小平方误差算法

例1 已知两类模式训练样本: $\omega_1 : [0,0]^T, [0,1]^T$, $\omega_2 : [1,0]^T, [1,1]^T$;

试用LMSE算法求解权向量。

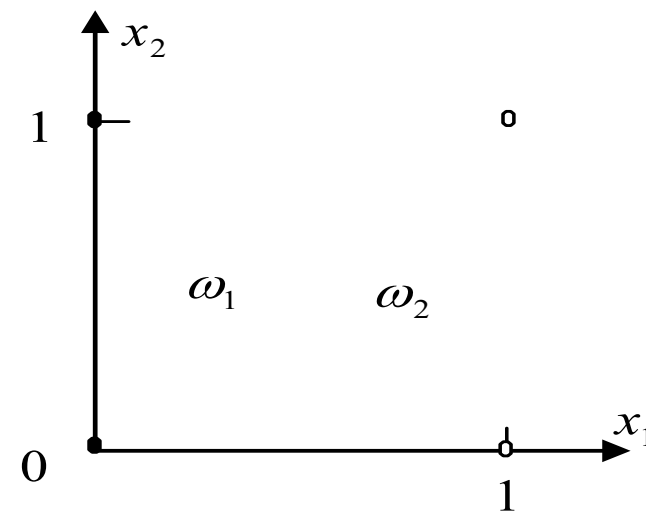
解: (1) 写出规范化增广样本矩阵: $\mathbf{X} = \begin{bmatrix} 0 & 0 & 1 \\ 0 & 1 & 1 \\ -1 & 0 & -1 \\ -1 & -1 & -1 \end{bmatrix}$

(2) 求伪逆矩阵 $\mathbf{X}^\# = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T$

$$\mathbf{X}^T \mathbf{X} = \begin{bmatrix} 0 & 0 & -1 & -1 \\ 0 & 1 & 0 & -1 \\ 1 & 1 & -1 & -1 \end{bmatrix} \begin{bmatrix} 0 & 0 & 1 \\ 0 & 1 & 1 \\ -1 & 0 & -1 \\ -1 & -1 & -1 \end{bmatrix} = \begin{bmatrix} 2 & 1 & 2 \\ 1 & 2 & 2 \\ 2 & 2 & 4 \end{bmatrix}$$

$$(\mathbf{X}^T \mathbf{X})^{-1} = \frac{1}{|\mathbf{X}^T \mathbf{X}|} \begin{bmatrix} 4 & 0 & -2 \\ 0 & 4 & -2 \\ -2 & -2 & 3 \end{bmatrix} = \frac{1}{4} \begin{bmatrix} 4 & 0 & -2 \\ 0 & 4 & -2 \\ -2 & -2 & 3 \end{bmatrix}$$

$$\mathbf{X}^\# = \frac{1}{4} \begin{bmatrix} 4 & 0 & -2 \\ 0 & 4 & -2 \\ -2 & -2 & 3 \end{bmatrix} \begin{bmatrix} 0 & 0 & -1 & -1 \\ 0 & 1 & 0 & -1 \\ 1 & 1 & -1 & -1 \end{bmatrix} = \frac{1}{4} \begin{bmatrix} -2 & -2 & -2 & -2 \\ -2 & 2 & 2 & -2 \\ 3 & 1 & -1 & 1 \end{bmatrix}$$



最小平方误差算法

例1 已知两类模式训练样本： $\omega_1 : [0,0]^T, [0,1]^T$, $\omega_2 : [1,0]^T, [1,1]^T$;

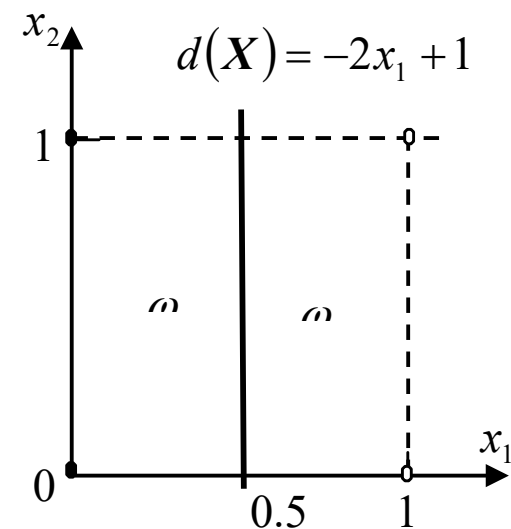
试用LMSE算法求解权向量。

解： (3) 取 $B(1) = [1,1,1,1]^T$ 和 $\eta=1$ 开始迭代：

$$W(1) = X^{\#} B(1) = \frac{1}{4} \begin{bmatrix} -2 & -2 & -2 & -2 \\ -2 & 2 & 2 & -2 \\ 3 & 1 & -1 & 1 \end{bmatrix} \begin{bmatrix} 1 \\ 1 \\ 1 \\ 1 \end{bmatrix} = \frac{1}{4} \begin{bmatrix} -8 \\ 0 \\ 4 \end{bmatrix} = \begin{bmatrix} -2 \\ 0 \\ 1 \end{bmatrix}$$

$$e(1) = XW(1) - B(1) = \begin{bmatrix} 0 & 0 & 1 \\ 0 & 1 & 1 \\ -1 & 0 & -1 \\ -1 & -1 & -1 \end{bmatrix} \begin{bmatrix} -2 \\ 0 \\ 1 \end{bmatrix} - \begin{bmatrix} 1 \\ 1 \\ 1 \\ 1 \end{bmatrix} = \begin{bmatrix} 1 \\ 1 \\ 1 \\ 1 \end{bmatrix} - \begin{bmatrix} 1 \\ 1 \\ 1 \\ 1 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \end{bmatrix}$$

解为 $W(1)$, 判断函数为： $d(X) = -2x_1 + 1$



最小平方误差算法

例 已知模式训练样本: $\omega_1: [0,0]^T, [1,1]^T$, $\omega_2: [0,1]^T, [1,0]^T$;

用LMSE算法求解权向量。

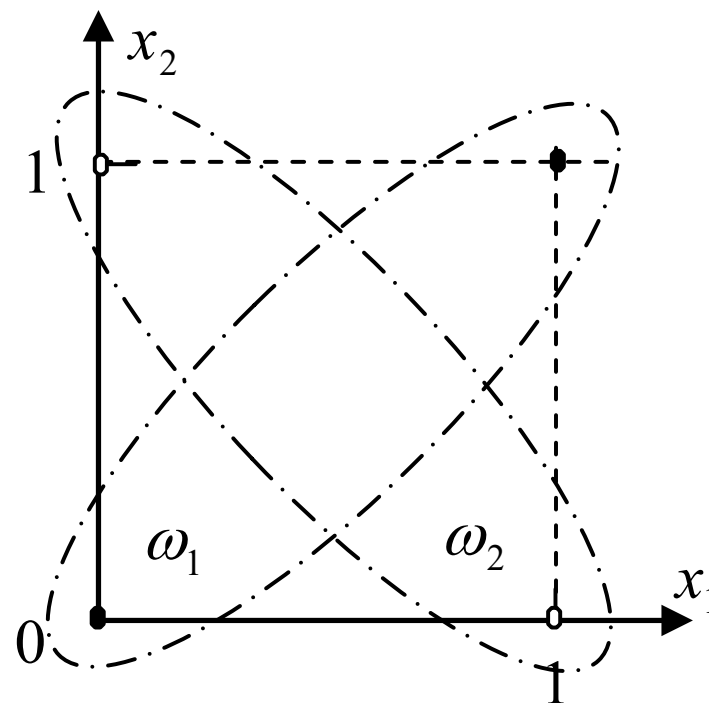
解: (1) 规范化增广样本矩阵:
$$\mathbf{X} = \begin{bmatrix} 0 & 0 & 1 \\ 1 & 1 & 1 \\ 0 & -1 & -1 \\ -1 & 0 & -1 \end{bmatrix}$$

(2) 求 $\mathbf{X}^\# = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T$:
$$\mathbf{X}^\# = \frac{1}{4} \begin{bmatrix} -2 & 2 & 2 & -2 \\ -2 & 2 & -2 & 2 \\ 3 & -1 & -1 & -1 \end{bmatrix}$$

(3) 取 $\mathbf{B}(1) = [1,1,1,1]^T$ 和 $\eta=1$, 迭代: $\mathbf{W}(1) = \mathbf{X}^\# \mathbf{B}(1) = [0 \ 0 \ 0]^T$

$$\mathbf{e}(1) = \mathbf{XW}(1) - \mathbf{B}(1) = [0 \ 0 \ 0 \ 0]^T - [1 \ 1 \ 1 \ 1]^T = [-1 \ -1 \ -1 \ -1]^T$$

$\mathbf{e}(1)$ 全部分量为负, 无解, 停止迭代。为线性不可分模式。





- ✓ 线性模型基本形式
线性方程, 基本形式, ...
- ✓ 线性回归
参数估计, 最小二乘, 对数几率回归, ...
- ✓ 线性判别分析
基本概念, 广义线性判别, 几何性质 ...
- ✓ 经典线性判别准则
Fisher判别, 感知器准则、梯度法 ...
- ✓ **非线性判别函数**
分段线性, ...

非线性判别函数

线性判别函数的特点：形式简单，容易学习；
用于线性可分的模式类。

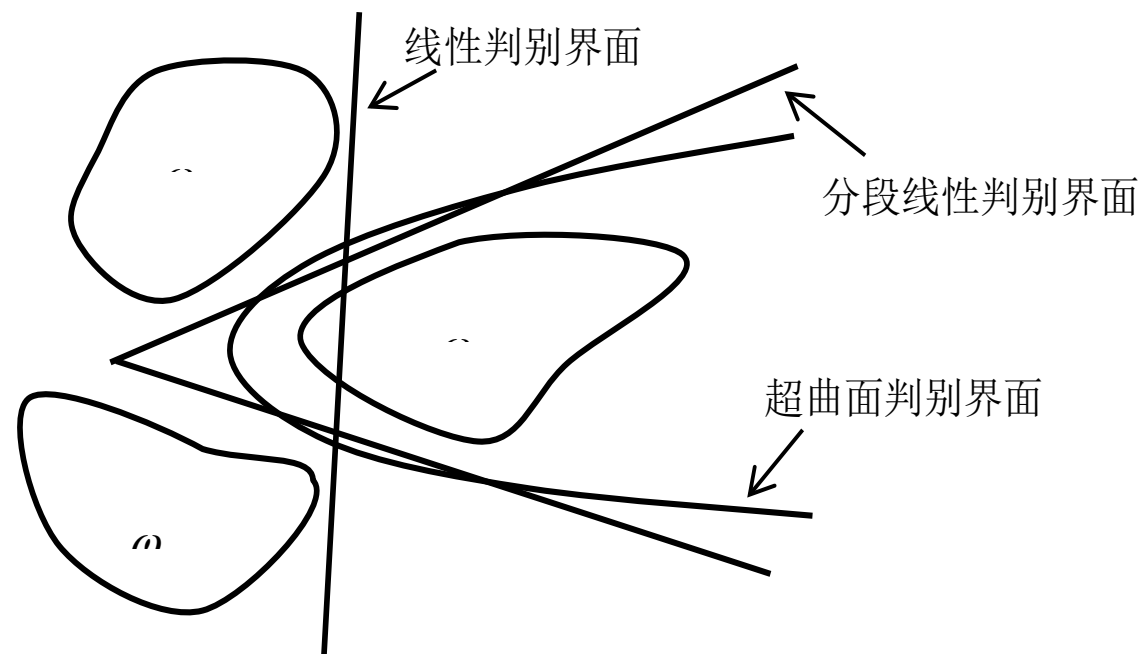
非线性判别函数：用于线性不可分情况。分段线性、超曲面。

分段线性判别函数

基本组成为超平面。

特点

- * 相对简单；
- * 能逼近各种形状的超曲面。



非线性判别函数

1. 一般分段线性判别函数

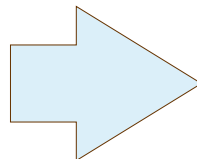
设有 M 类模式，将 ω_i 类 ($i=1,2,\dots,M$) 划分为 l_i 个子类: $\omega_i: \{\omega_i^1, \omega_i^2, \dots, \omega_i^{l_i}\} \quad i=1,2,\dots,M$

其中第 n 个子类的判别函数: $d_i^n(\mathbf{X}) = (\mathbf{W}_i^n)^T \mathbf{X} \quad n=1,2,\dots,l_i; \quad i=1,2,\dots,M$

ω_i 类的判别函数定义为: $d_i(\mathbf{X}) = \max\{d_i^n(\mathbf{X}), \quad n=1,2,\dots,l_i\}$

M 类的判决规则: 若 $d_j(\mathbf{X}) = \max\{d_i(\mathbf{X}), \quad i=1,2,\dots,M\}$, 则 $\mathbf{X} \in \omega_j$

用各类判别函数进行分类判决实际是
用各类选出的子类判别函数进行判决



判别面由各子类的
判别函数决定

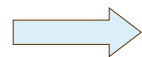
非线性判别函数

1. 一般分段线性判别函数

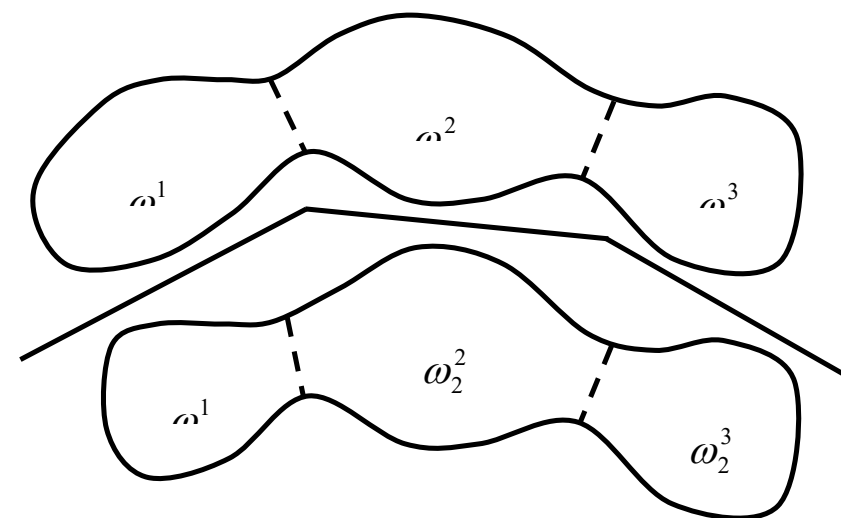
若 ω_i 类的第 n 个子类和 ω_j 类的第 m 个子类相邻，判别界面方程为：

$$d_i^n(\mathbf{X}) = d_j^m(\mathbf{X})$$

子类之间的判别界面组成各类之间的判别界面



类间判别界面分段线性



非线性判别函数

2. 基于距离的分段线性判别函数

1) 最小距离分类器

设 ω_1 类均值向量: $\mathbf{M}_1 = \frac{1}{N_1} \sum_{i=1}^{N_1} \mathbf{X}_i$ ω_2 类均值向量: $\mathbf{M}_2 = \frac{1}{N_2} \sum_{i=1}^{N_2} \mathbf{X}_i$ N_1, N_2 : 两类样本数。

任一模式 $\mathbf{X}$ 到 $\mathbf{M}_1$ 和 $\mathbf{M}_2$ 的欧氏距离平方:

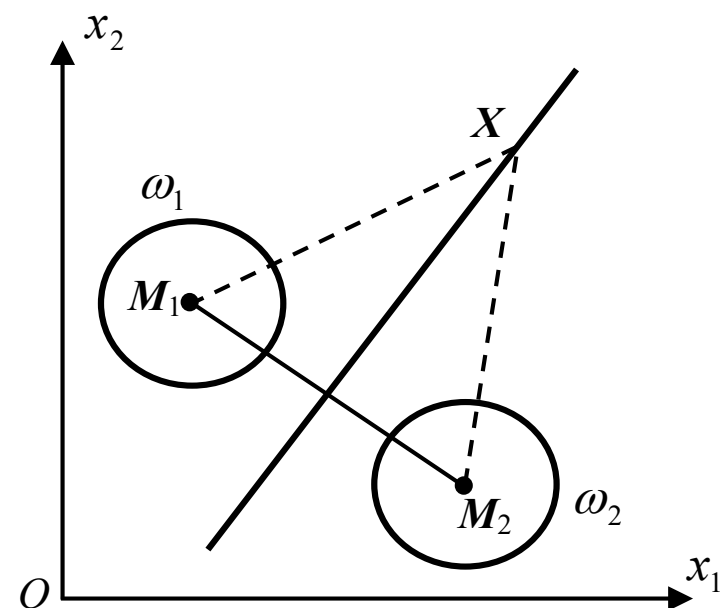
$$d_1(\mathbf{X}) = \|\mathbf{X} - \mathbf{M}_1\|^2 \qquad d_2(\mathbf{X}) = \|\mathbf{X} - \mathbf{M}_2\|^2$$

判决规则: 若 $d_1(\mathbf{X}) < d_2(\mathbf{X})$, 则 $\mathbf{X} \in \omega_1$

 若 $d_2(\mathbf{X}) < d_1(\mathbf{X})$, 则 $\mathbf{X} \in \omega_2$

判别界面方程:

$$\|\mathbf{X} - \mathbf{M}_1\|^2 = \|\mathbf{X} - \mathbf{M}_2\|^2$$



—— 最小距离分类器

非线性判别函数

2. 分段线性距离分类器

设： M 类模式，其中 ω_i 类划分为 l_i 个子类，第 n 个子类的均值向量为 $\mathbf{M}_i^n$ 。每个子类的判别函数：

$$d_i^n(\mathbf{X}) = \|\mathbf{X} - \mathbf{M}_i^n\|^2, \quad n = 1, 2, \dots, l_i; \quad i = 1, 2, \dots, M$$

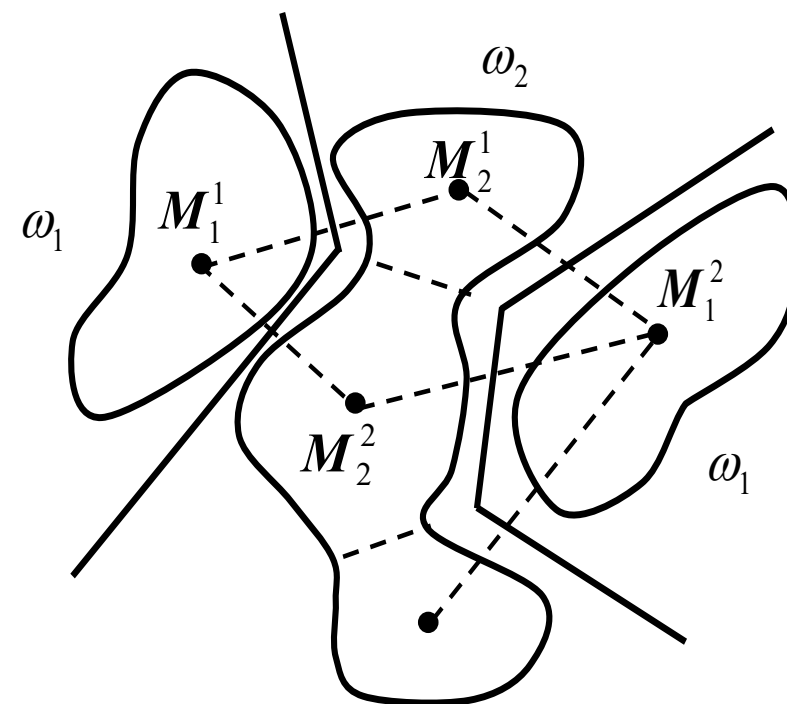
每类的判别函数：

$$d_i(\mathbf{X}) = \min \{d_i^n(\mathbf{X}), \quad n = 1, 2, \dots, l_i\}, \quad i = 1, 2, \dots, M$$

判决规则：

$$\text{若 } d_j(\mathbf{X}) = \min \{d_i(\mathbf{X}), \quad i = 1, 2, \dots, M\}$$

$$\text{则 } \mathbf{X} \in \omega_j$$



类别不平衡问题

类别不平衡 (class imbalance)

- 不同类别训练样例数相差很大情况 (正类为小类)

类别平衡正例预测 $\frac{y}{1-y} > 1$  $\frac{y}{1-y} > \frac{m^+}{m^-}$ 正负类比例

类别不平衡学习的一个基本策略: 再缩放 (rescaling)

- 欠采样 (undersampling)
 - 去除一些反例使正反例数目接近 (EasyEnsemble [Liu et al.,2009])
- 过采样 (oversampling)
 - 增加一些正例使正反例数目接近 (SMOTE [Chawla et al.2002])
- 阈值移动 (threshold-moving)

Any Questions?

