



机器学习（第6讲）

主讲：张磊

E-mail: leizhang@cqu.edu.cn
Lab Website: <http://www.leizhang.tk>





第六章：概率建模



脑认知



猫



第六章：概率建模

□产生式思考

对于任意一个样本，我们定义模型如下：

$$t_n = \mathbf{w}^T \mathbf{x}_n + \varepsilon_n$$

其中， ε_n 是一个随机变量。所谓产生式考虑，是指第n次的观测响应，是通过 $\mathbf{w}^T \mathbf{x}_n$ 生成。同时，伴随着一个随机变量，可正可负。

ε_n 是一个随机变量，存在某种概率分布。

我们假设， $\varepsilon_n \sim \mathcal{N}(\mu, \sigma^2)$ 高斯分布（正态分布）

$$p(\varepsilon_n) = \mathcal{N}(\mu, \sigma^2)$$

该假设已在前面的课程给出解释（中心极限定理、易于分析）。



第六章：概率建模

□产生式思考

对于任意一个样本，我们定义模型如下：

$$t_n = \mathbf{w}^T \mathbf{x}_n + \varepsilon_n$$

现在，模型由两部分组成：

1. 数据生成的决定性部分， $\mathbf{w}^T \mathbf{x}_n$ ，表示一种趋势或倾向；
2. 随机组成部分 ε_n ，可以理解成噪声。

注：这里的随机变量 ε_n （或噪声）并非限定为高斯分布，也并非限定为可加性噪声。但为了便于解释和理解，选择可加性高斯分布噪声，可使我们获得最优参数 $\hat{\mathbf{w}}$ 的精确表达式。



第六章：概率建模

□产生式思考

将模型重新定义如下：

$$t_n = \mathbf{w}^T \mathbf{x}_n + \varepsilon_n, \text{ 其中, } \varepsilon_n \sim \mathcal{N}(\mu, \sigma^2)$$

这里令 $\mu=0$ ，即 $\varepsilon_n \sim \mathcal{N}(0, \sigma^2)$

由于 $\mathbf{w}^T \mathbf{x}_n$ 是一个常量，因此加上一个随机变量 ε_n 后， t_n 也是一个随机变量（这是赋予随机变量的原因）。

而且， $t_n \sim \mathcal{N}(\mathbf{w}^T \mathbf{x}_n, \sigma^2)$ （为什么？高斯随机变量加上一个常数，依然是高斯，仅均值发生变化），从而概率密度函数可以写成：

$$p(t_n | \mathbf{x}_n, \mathbf{w}, \sigma^2) = \mathcal{N}(\mathbf{w}^T \mathbf{x}_n, \sigma^2) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(t_n - \mathbf{w}^T \mathbf{x}_n)^2}{2\sigma^2}}$$

可以看出，预测响应 t_n 的密度依赖于特定的 $\mathbf{x}_n, \mathbf{w}$ (均值)和 σ^2 (方差)的值。



第六章：概率建模

□问题

我们看下面这个概率密度函数

$$p(t_n|\mathbf{x}_n, \mathbf{w}, \sigma^2) = \mathcal{N}(\mathbf{w}^T \mathbf{x}_n, \sigma^2)$$

对于给定的 $\mathbf{w}, \sigma^2$ ，我们可以计算出在观测 $\mathbf{x}_n$ 下，使得 $p(t_n|\mathbf{x}_n, \mathbf{w}, \sigma^2)$ 最大的 t_n ，即预测值。

对于已知的高斯分布来说，只要均值和方差确定，即可绘制出连续或离散的概率密度。

例：设第10个苹果的直径为0.5，即 $\mathbf{x}_{10}=[0.5, 1]^T$ ；设 $\mathbf{w}=[0.6, -0.1]^T$, $\sigma^2 = 0.1$ ；另：该苹果实际重量 $t_{10}=2$ 。

分析：针对该苹果，可以绘制其重量的概率密度函数曲线

第六章：概率建模

□问题

✓ 如果 $\mathbf{w}, \sigma^2$ 已经最优

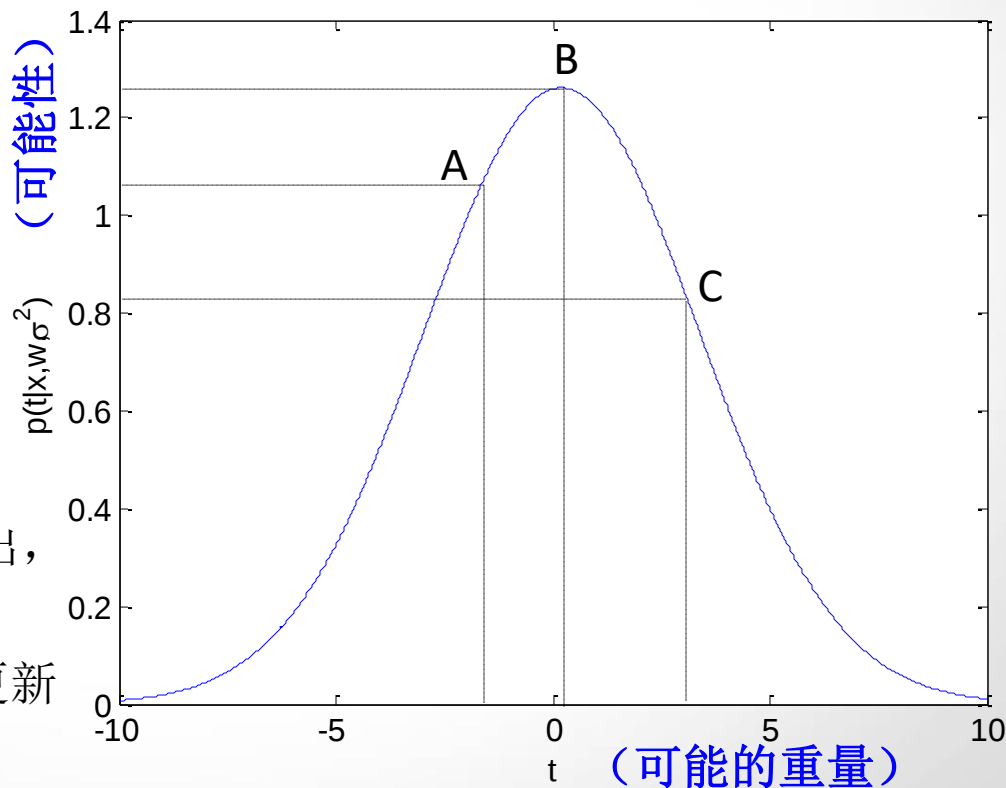
通过右边的概率密度函数可以得出：

1. B点的可能性最大，即为苹果重量的最佳估计值 $\mathbf{w}^T \mathbf{x}_{10}$
2. C点的可能性最小。

✓ 如果 $\mathbf{w}, \sigma^2$ 还不是最优

实际重量 $t_{10}=2$ ，根据图可以看出，最佳的可能值是在B和C之间。

但是数据不可能变，因此只能更新模型参数。



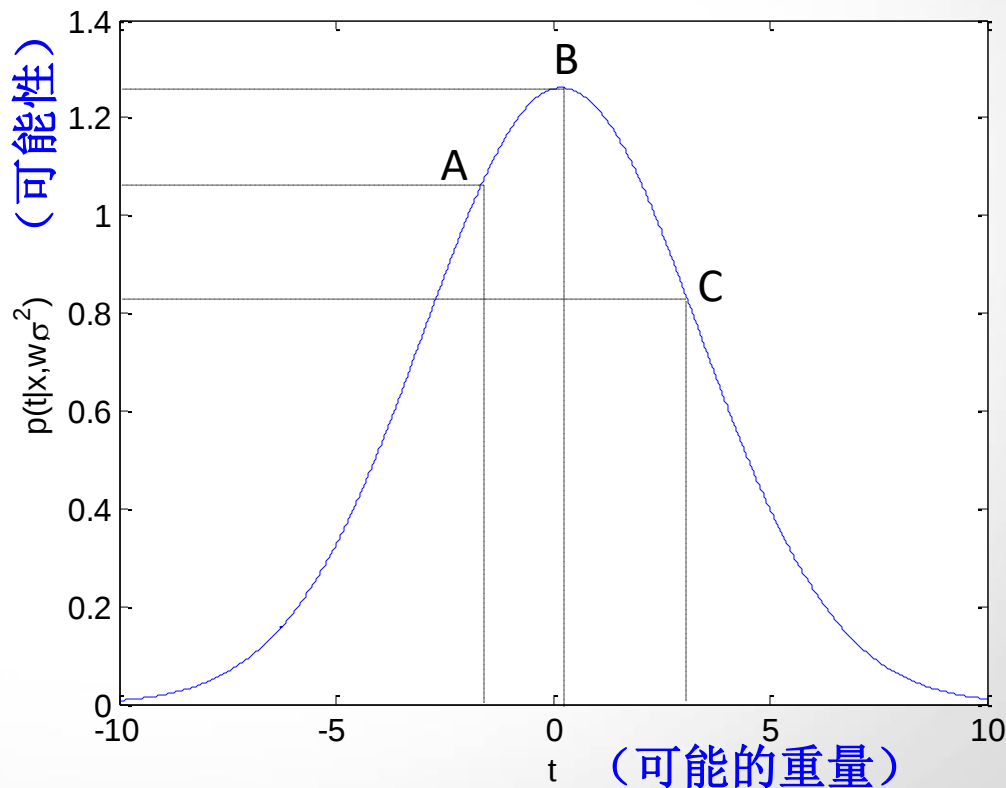
第六章：概率建模

□ 定义

这种通过寻找参数，来最大化似然值的方式，称为基于似然函数最大化的线性建模——概率建模，是机器学习中的一种重要方法。

□ 数据集的似然值

对机器学习，我们感兴趣的是整个数据集的所有样本的似然值，而并非单个样本的似然值。





第六章：概率建模

□ 数据集的似然值

- 对于单个样本的概率密度表达式为

$$p(t_n|\mathbf{x}_n, \mathbf{w}, \sigma^2) = \mathcal{N}(\mathbf{w}^T \mathbf{x}_n, \sigma^2)$$

- 如果有N个数据样本点，那么我们感兴趣的是它们的联合条件概率密度

$$p(t_1, t_2, \dots, t_N | \mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N, \mathbf{w}, \sigma^2)$$

- 根据向量和矩阵表示法，可以写成紧缩的表达形式：

$$p(\mathbf{t}|\mathbf{X}, \mathbf{w}, \sigma^2)$$

- 假设数据的噪声是独立的，即噪声的联合概率密度可表达为

$$p(\boldsymbol{\varepsilon}) = p(\varepsilon_1, \varepsilon_2, \dots, \varepsilon_N) = \prod_{n=1}^N p(\varepsilon_n)$$

这一独立性假设，可以使联合密度进行分解为更易操作的对象，使问题求解更简单



第六章：概率建模

□ 数据集的似然值

➤ 基于噪声的独立性假设，对N个数据样本点，它们的联合条件概率密度

$$\begin{aligned} p(\mathbf{t}|\mathbf{X}, \mathbf{w}, \sigma^2) &= p(t_1, t_2, \dots, t_N | \mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N, \mathbf{w}, \sigma^2) = \prod_{n=1}^N p(t_n | \mathbf{x}_n, \mathbf{w}, \sigma^2) \\ &= \prod_{n=1}^N \mathcal{N}(\mathbf{w}^T \mathbf{x}_n, \sigma^2) \end{aligned}$$

注：这里不是说 $\mathbf{t}$ 是独立的，否则不存在对数据的建模。而是条件独立，可以根据模型 $t_n = \mathbf{w}^T \mathbf{x}_n + \varepsilon_n$ 理解。当 $\mathbf{w}^T \mathbf{x}_n$ 给定时， t_n 独立，即条件独立。
或者说，只有当模型确定了， t_n 才是一个独立的随机变量。

另： $\prod_{n=1}^N$ 表示连乘符号。



第六章：概率建模

□ 最大似然建模

- 最大似然建模，就是为了寻找 $\mathbf{w}$ 和 σ^2 ，使得似然值最大化。也就是最大化下列似然函数。

$$\begin{aligned}\max_{\mathbf{w}, \sigma^2} p(\mathbf{t}|\mathbf{X}, \mathbf{w}, \sigma^2) &= p(t_1, t_2, \dots, t_N | \mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N, \mathbf{w}, \sigma^2) \\ &\propto \max_{\mathbf{w}, \sigma^2} \prod_{n=1}^N p(t_n | \mathbf{x}_n, \mathbf{w}, \sigma^2) = \prod_{n=1}^N \mathcal{N}(\mathbf{w}^T \mathbf{x}_n, \sigma^2)\end{aligned}$$

换句话说，当给定 $\mathbf{w}$ 和 σ^2 时，可以获得数据集的似然值 $L = p(\mathbf{t}|\mathbf{X}, \mathbf{w}, \sigma^2)$ ，但是，该似然值并非最大，因此，我们要改变 $\mathbf{w}$ 和 σ^2 使得 L 最大。

当然，为什么不改变 $\mathbf{X}$ 呢？数据！！

因为数据集是固定的，不同的参数值会产生不同似然值，建模的目的是要求出最佳的参数值。



第六章：概率建模

□ 最大似然建模

➤ 最大似然建模求解：即求解 $\mathbf{w}$ 和 σ^2 ，使得似然值最大化。

$$\max_{\mathbf{w}, \sigma^2} \prod_{n=1}^N \mathcal{N}(\mathbf{w}^T \mathbf{x}_n, \sigma^2)$$

很显然，要求解上面的模型，可以最大化似然值的自然对数，用 $\log$ 表示，即，

$$\max_{\mathbf{w}, \sigma^2} \prod_{n=1}^N \mathcal{N}(\mathbf{w}^T \mathbf{x}_n, \sigma^2) \propto \max_{\mathbf{w}, \sigma^2} \log \prod_{n=1}^N \mathcal{N}(\mathbf{w}^T \mathbf{x}_n, \sigma^2)$$

➤ 似然函数的自然对数可以简化

$$\begin{aligned} \log \prod_{n=1}^N \mathcal{N}(\mathbf{w}^T \mathbf{x}_n, \sigma^2) &= \log \mathcal{N}(\mathbf{w}^T \mathbf{x}_1, \sigma^2) + \cdots + \log \mathcal{N}(\mathbf{w}^T \mathbf{x}_N, \sigma^2) \\ &= \sum_{n=1}^N \log \mathcal{N}(\mathbf{w}^T \mathbf{x}_n, \sigma^2) \end{aligned}$$



第六章：概率建模

□ 最大似然建模

➤ 似然函数的自然对数可以简化

$$\begin{aligned}\log L &= \log \prod_{n=1}^N \mathcal{N}(\mathbf{w}^T \mathbf{x}_n, \sigma^2) = \sum_{n=1}^N \log \mathcal{N}(\mathbf{w}^T \mathbf{x}_n, \sigma^2) \\ &= \sum_{n=1}^N \log \left(\frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(t_n - \mathbf{w}^T \mathbf{x}_n)^2}{2\sigma^2}} \right) \\ &= \sum_{n=1}^N \left(-\frac{1}{2} \log(2\pi) - \log \sigma - \frac{1}{2\sigma^2} (t_n - \mathbf{w}^T \mathbf{x}_n)^2 \right) \\ &= -\frac{N}{2} \log(2\pi) - N \log \sigma - \frac{1}{2\sigma^2} \sum_{n=1}^N (t_n - \mathbf{w}^T \mathbf{x}_n)^2\end{aligned}$$



第六章：概率建模

□ 最大似然建模

似然函数的自然对数最大化模型可以简化

$$\max_{\mathbf{w}, \sigma^2} \log L = -\frac{N}{2} - \log(2\pi) - N \log \sigma - \frac{1}{2\sigma^2} \sum_{n=1}^N (t_n - \mathbf{w}^T \mathbf{x}_n)^2$$

➤ 利用前面学习的内容，对 $\mathbf{w}$ 求偏导，

$$\begin{aligned} \frac{\partial \log L}{\partial \mathbf{w}} &= \frac{1}{\sigma^2} \sum_{n=1}^N \mathbf{x}_n (t_n - \mathbf{x}_n^T \mathbf{w}) = \frac{1}{\sigma^2} \sum_{n=1}^N \mathbf{x}_n t_n - \frac{1}{\sigma^2} \sum_{n=1}^N \mathbf{x}_n \mathbf{x}_n^T \mathbf{w} \\ &= \frac{1}{\sigma^2} (\mathbf{X}^T \mathbf{t} - \mathbf{X}^T \mathbf{X} \mathbf{w}) \end{aligned}$$

令 $\frac{\partial \log L}{\partial \mathbf{w}} = \mathbf{0}$ ，可以得出 $\hat{\mathbf{w}} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{t}$ ($\mathbf{w}$ 的最大似然解) 似曾相识! ?
将噪声假设为高斯分布，最小化平方损失等同于最大似然解。



第六章：概率建模

□ 最大似然建模

似然函数的自然对数最大化模型可以简化

$$\max_{\mathbf{w}, \sigma^2} \log L = -\frac{N}{2} - \log(2\pi) - N \log \sigma - \frac{1}{2\sigma^2} \sum_{n=1}^N (t_n - \mathbf{w}^T \mathbf{x}_n)^2$$

➤ 对 σ 求偏导,

$$\frac{\partial \log L}{\partial \sigma} = -\frac{N}{\sigma} + \frac{1}{\sigma^3} \sum_{n=1}^N (t_n - \mathbf{w}^T \mathbf{x}_n)^2$$

令 $\frac{\partial \log L}{\partial \sigma} = 0$, 可以得出 $\widehat{\sigma^2} = \frac{1}{N} \sum_{n=1}^N (t_n - \mathbf{w}^T \mathbf{x}_n)^2 = \frac{1}{N} (\mathbf{t} - \mathbf{X}\mathbf{w})^T (\mathbf{t} - \mathbf{X}\mathbf{w})$

最终, $\widehat{\sigma^2} = \frac{1}{N} (\mathbf{t}^T \mathbf{t} - \mathbf{t}^T \mathbf{X}\mathbf{w})$



第六章：概率建模

课外任务：

根据前面的内容，请自行推导最大似然建模方法的整个过程。

第六章：概率建模

□ Logistic回归

回顾线性回归算法中的二类分类问题

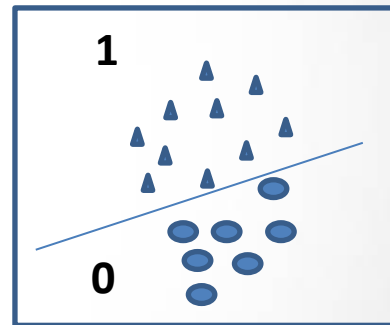
$$h(\mathbf{x}) = \mathbf{w}_0 + \mathbf{w}_1\mathbf{x}_1 + \cdots + \mathbf{w}_d\mathbf{x}_d$$

分类边界为

$$h(\mathbf{x}) = 0$$

□ 当 $h(\mathbf{x}) > 0$, $\mathbf{x}$ 的属性为 $Y=1$

□ 当 $h(\mathbf{x}) < 0$, $\mathbf{x}$ 的属性为 $Y=0$



也就是说，对于一次观测 $\mathbf{x}$ ，其仅仅给出“是”或者“不是”的猜测（**简单、粗暴**），给出的预测是一个“肯定发生”或“肯定不发生”的事件。

如果提供一个关于可能性大小的结果岂不是更好？？换句话说，“可能性”是关于概率的一个概念，如果对于一个观测，能够计算出该观测的事件发生的概率（即 $Y=1$ 发生的概率），通常比直接给出“是”或者“不是”更加有用。



第六章：概率建模

□ Logistic回归

几个特点：

- **Logistic**回归模型作为一种概率模型，可用于预测某事件发生的概率；
- **Logistic**回归模型不要求因变量在正态分布的前提假设下完成。
- 该模型在疾病诊断、预测中应用较广，主要用于分析不同变量因素对疾病的影响。在医学领域，一些随机事件只具有两种**互斥**结果的离散性随机事件。这类只具有两种互斥结果的离散型随机事件的规律性进行描述的一种概率分布，叫**二项分布**。
- 二项分布概率公式：如果进行n次伯努利试验，取得成功次数(事件发生)为k(k=0,1,...,n)的概率，表示为

$$P(\xi = k) = C_n^k p^k (1 - p)^{n-k}$$

其中，事件发生的概率为p, 事件不发生的概率为1-p。 $\xi \sim B(n, p)$ ，当n=1时，二项分布也可称为伯努利分布或0-1分布（1次伯努利试验）



第六章：概率建模

□ Logistic回归

条件概率：

对于观测输入变量 $\mathbf{x}$, 其响应 Y 的条件概率分布表达为

$$P(Y|X)$$

这个概率表示模型的准确性。

设想，如果对于“今天是否下雪”的预测结果是51%，但是并没有下雪。那么这个预测也要比预测为下雪的可能性为99%要好！

设 $P(Y = 1|X = x) = p(x; \mathbf{w})$, 其中 $\mathbf{w}$ 为参数，那么有

$$P(Y = 0|X = x) = 1 - p(x; \mathbf{w})$$

Y 的取值为0
或1



第六章：概率建模

□ Logistic回归

条件概率：

对于一个观测 x_i ,其响应 $Y = y_i$ 的条件概率可以写成

$$P(Y = y_i | X = x_i) = p(x_i; \mathbf{w})^{y_i} (1 - p(x_i; \mathbf{w}))^{1-y_i}$$

由于 N 次观测的独立性，联合条件概率分布(似然)可以写为

$$\prod_{i=1}^N P(Y = y_i | X = x_i) = \prod_{i=1}^N p(x_i; \mathbf{w})^{y_i} (1 - p(x_i; \mathbf{w}))^{1-y_i}$$

N 个观测的似然函数

Y 的取值为0
或1



第六章：概率建模

□ Logistic回归—与线性回归的关系

主要思想：logistic回归模型的输出是某次观测事件发生的概率，通过观测样本 $\mathbf{x}$ 的特征与最佳估计参数相乘、求和后，然后利用Logistic函数（sigmoid）进行非线性计算（概率），然后与阈值进行比较，得出该观测 $\mathbf{x}$ 的输出。

相乘、求和：

$$h(\mathbf{x}) = \mathbf{w}_0 + \mathbf{w}_1\mathbf{x}_1 + \cdots + \mathbf{w}_d\mathbf{x}_d$$

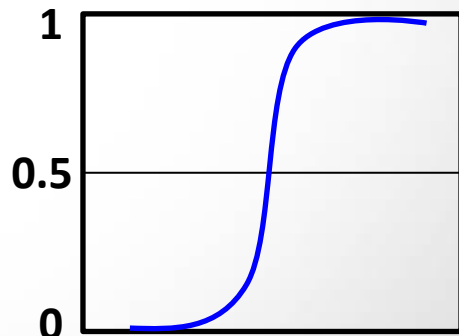
条件概率 $P(y = 1|\mathbf{x})$ 的计算表达式为

$$P(y = 1|\mathbf{x}) = p(\mathbf{x}) = \frac{1}{1 + e^{-h(\mathbf{x})}}$$

条件概率 $P(y = 0|\mathbf{x})$ 的计算表达式为

$$P(y = 0|\mathbf{x}) = 1 - P(y = 1|\mathbf{x}) = 1 - p(\mathbf{x}) = \frac{1}{1 + e^{h(\mathbf{x})}}$$

Sigmoid函数



定义域 $(-\infty, \infty)$ ；值域为 $(0, 1)$

第六章：概率建模

□ Logistic回归—与线性回归的关系

$P(y = 1|x) = \frac{1}{1+e^{-h(x)}}$ 为事件发生的概率;

$P(y = 0|x) = \frac{1}{1+e^{h(x)}}$ 为事件不发生的概率;

那么, 事件发生与事件不发生的概率比值为

$$\frac{P(y = 1|x)}{P(y = 0|x)} = \frac{p(x)}{1 - p(x)} = \frac{1 + e^{h(x)}}{1 + e^{-h(x)}} = e^{h(x)}$$

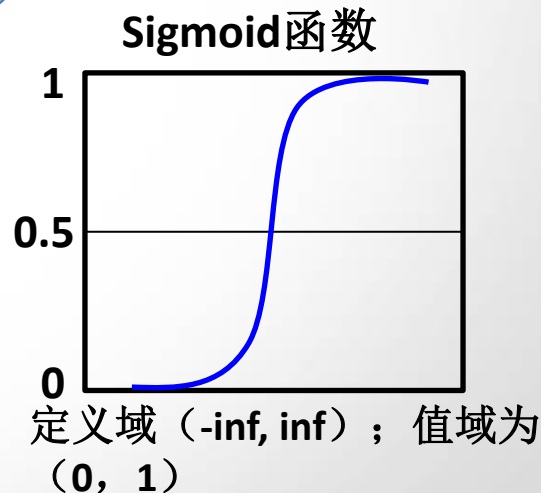
两边同时取对数,

$$\log \frac{p(x)}{1 - p(x)} = h(x) = w_0 + w_1x_1 + \dots + w_dx_d$$

得到线性函数 (即**Logistic 回归模型**)。

利用上式, 可以求出 $p(x) = \frac{e^{h(x)}}{1+e^{h(x)}} = \frac{1}{1+e^{-h(x)}}$ 介于 (0,1)。

采用Logit变换, 建立与属性变量之间的关系。
当 $0 < p < 1$ 时, logit变换可取任意值



第六章：概率建模

□ Logistic回归—与线性回归的关系

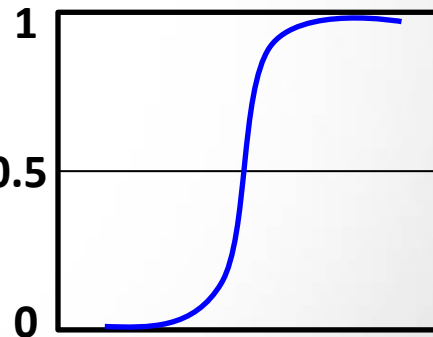
□ $p(x) = \frac{e^{h(x)}}{1+e^{h(x)}} = \frac{1}{1+e^{-h(x)}}$ 介于 $(0,1)$ 。

根据观测样本 x 的概率表示可以看出，当 $p(x) > 0.5$, $Y=1$;
当 $p(x) < 0.5$, $Y=0$ 。

□ 这也同时意味着， $h(x) > 0$ 时， $Y=1$ ； $h(x) < 0$ 时， $Y=0$ 。
即，将两类分开的判决边界为 $h(x) = 0$ 。

□ 但Logistic回归不同于线性分类器，不是为了找分类边界，
而是找出一个概率值，对“可能性”进行建模。

Sigmoid函数



定义域 $(-\infty, \infty)$;
值域为 $(0, 1)$

在Logistic回归模型中，若要计算 $p(x;w)$ ，关键是要求出参数 w 。



第六章：概率建模

□ Logistic回归

最大似然函数：

对于N次观测，似然函数可以写为

$$\prod_{i=1}^N P(Y = y_i | X = x_i) = \prod_{i=1}^N p(x_i; \mathbf{w})^{y_i} (1 - p(x_i; \mathbf{w}))^{1-y_i}$$

注：N次观测的独立性。 $\mathbf{w}$ 是模型参数。

我们的目标是能够求出使这一似然函数取**最大值**时的参数估计 $\mathbf{w}$ 。

$$\hat{\mathbf{w}} = \max_{\mathbf{w}} \prod_{i=1}^N P(Y = y_i | X = x_i) = \prod_{i=1}^N p(x_i; \mathbf{w})^{y_i} (1 - p(x_i; \mathbf{w}))^{1-y_i}$$

接下来，如何解这个模型是关键。



第六章：概率建模

□ Logistic回归

最大似然函数：

对似然函数取自然对数，可得出对数似然函数 $L(\mathbf{w})$ ，

$$\begin{aligned} L(\mathbf{w}) &= \sum_{i=1}^N [y_i \log p(x_i; \mathbf{w}) + (1 - y_i) \log(1 - p(x_i; \mathbf{w}))] \\ &= \sum_{i=1}^N \left[\log(1 - p(x_i; \mathbf{w})) + y_i \log \frac{p(x_i; \mathbf{w})}{1 - p(x_i; \mathbf{w})} \right] \\ &= \sum_{i=1}^N \left[\log \left(1 - \frac{1}{1 + e^{-h(x)}} \right) + y_i h(x) \right] \\ &= \sum_{i=1}^N \left[-\log(1 + e^{h(x)}) + y_i h(x) \right] \\ &= \sum_{i=1}^N \left[-\log(1 + e^{w_0 + \mathbf{w}^T \mathbf{x}_i}) \right] + \sum_{i=1}^N y_i (w_0 + \mathbf{w}^T \mathbf{x}_i) \end{aligned}$$

为了估计能使 $L(\mathbf{w})$ 取得最大值的参数 $\mathbf{w}=[w_0, w_1, \dots, w_d]$ ，对函数求导。



第六章：概率建模

□ Logistic回归

最大似然函数：

化简后的对数似然函数 $L(\mathbf{w})$,

$$L(\mathbf{w}) = \sum_{i=1}^N \left[-\log \left(1 + e^{w_0 + \mathbf{w}^T \mathbf{x}_i} \right) \right] + \sum_{i=1}^N y_i (w_0 + \mathbf{w}^T \mathbf{x}_i)$$

为了估计能使 $L(\mathbf{w})$ 取得最大值的参数 $\mathbf{w}=[w_0, w_1, \dots, w_d]$, 分别求函数对 $w_0, w_1, \dots, w_d$ 的导数, 显然存在 $d+1$ 个方程。

$$\begin{aligned} \frac{\partial L(\mathbf{w})}{\partial w_j} &= - \sum_{i=1}^N \frac{1}{1 + e^{w_0 + \mathbf{w}^T \mathbf{x}_i}} \cdot e^{w_0 + \mathbf{w}^T \mathbf{x}_i} \cdot x_{ij} + \sum_{i=1}^N y_i x_{ij} \\ &= \sum_{i=1}^N (y_i - p(x_i; \mathbf{w})) x_{ij} \end{aligned}$$

注：此时导数表达式是一个超越方程，不存在闭合解（解析解、精确解）。怎么办？

超越方程：非多项式函数，如指数函数、对数函数、三角函数等



第六章：概率建模

□ Logistic回归

Newton-Raphson（牛顿-拉斐森）优化算法：

我们考虑一个单变量函数的情况，即 $f(\beta)$ ，求该函数的最优解 β^* 。我们假设函数 $f(\cdot)$ 是平滑的，那么意味着该函数在 β^* 处的导数为0，二阶导数大于0（极值理论）。

那么根据泰勒展开， $f(\beta)$ 的表达式可以写成

$$f(\beta) \approx f(\beta^*) + \frac{1}{2}(\beta - \beta^*)^2 \frac{d^2 f}{d\beta^2} \Big|_{\beta=\beta^*}$$

既然 β^* 是最小值点，那么 $f(\beta) > f(\beta^*)$ ，因此， $\frac{d^2 f}{d\beta^2} \Big|_{\beta=\beta^*} > 0$ 。

而且，根据上式， $\frac{df}{d\beta} \Big|_{\beta=\beta^*} = 0$

换句话说，在最优解附近， $f(\beta)$ 是一个近二次的函数表达式。



第六章：概率建模

□ Logistic回归

Newton-Raphson（牛顿-拉斐森）优化算法：

既然在最优解附近， $f(\beta)$ 是一个近二次的函数表达式，而我们的目的是求出这个最优解，因此，可以将 $f(\beta)$ 进行二次泰勒展开（为什么不更高次？？极值定理）

首先，猜测一个初始解 $\beta^{(0)}$ ，在该点，二次泰勒展开：

$$f(\beta) \approx f(\beta^{(0)}) + (\beta - \beta^{(0)}) \frac{df}{d\beta} \Big|_{\beta=\beta^{(0)}} + \frac{1}{2} (\beta - \beta^{(0)})^2 \frac{d^2f}{d\beta^2} \Big|_{\beta=\beta^{(0)}}$$

为了方便，上式可以重新写成，

$$f(\beta) \approx f(\beta^{(0)}) + (\beta - \beta^{(0)}) f'(\beta^{(0)}) + \frac{1}{2} (\beta - \beta^{(0)})^2 f''(\beta^{(0)})$$

此时， $f(\beta)$ 是关于 β 的二次函数，可导，因此便可以利用导数=0求解。



第六章：概率建模

□ Logistic回归

Newton-Raphson（牛顿-拉斐森）优化算法：

$$f(\beta) \approx f(\beta^{(0)}) + (\beta - \beta^{(0)})f'(\beta^{(0)}) + \frac{1}{2}(\beta - \beta^{(0)})^2 f''(\beta^{(0)})$$

若要求函数 $f(\beta)$ 的最小值，我们令

$$\frac{df(\beta)}{d\beta} = f'(\beta^{(0)}) + f''(\beta^{(0)})(\beta - \beta^{(0)}) = 0$$

从而可以计算出

$$\beta = \beta^{(0)} - \frac{f'(\beta^{(0)})}{f''(\beta^{(0)})}$$

令此时的最优解为 $\beta = \beta^{(1)}$ （并非最优）

$$\beta^{(1)} = \beta^{(0)} - \frac{f'(\beta^{(0)})}{f''(\beta^{(0)})}$$



第六章：概率建模

□ Logistic回归

Newton-Raphson（牛顿-拉斐森）优化算法：

当计算出新的次优解 $\beta^{(1)}$ ，函数 $f(\beta)$ 的泰勒展开为

$$f(\beta) \approx f(\beta^{(1)}) + (\beta - \beta^{(1)})f'(\beta^{(1)}) + \frac{1}{2}(\beta - \beta^{(1)})^2 f''(\beta^{(1)})$$

若要求函数 $f(\beta)$ 的最小值，我们令

$$\frac{df(\beta)}{d\beta} = f'(\beta^{(1)}) + f''(\beta^{(1)})(\beta - \beta^{(1)}) = 0$$

从而可以计算出

$$\beta = \beta^{(1)} - \frac{f'(\beta^{(1)})}{f''(\beta^{(1)})}$$

令此时的最优解为 $\beta = \beta^{(2)}$ （并非最优）

$$\beta^{(2)} = \beta^{(1)} - \frac{f'(\beta^{(1)})}{f''(\beta^{(1)})}$$



第六章：概率建模

□ Logistic回归

Newton-Raphson（牛顿-拉斐森）优化算法：

为了找到最优解，通常迭代多次，设为T次。

当计算出新的次优解 $\beta^{(t)}$ ，函数 $f(\beta)$ 的泰勒展开为

$$f(\beta) \approx f(\beta^{(t)}) + (\beta - \beta^{(t)})f'(\beta^{(t)}) + \frac{1}{2}(\beta - \beta^{(t)})^2 f''(\beta^{(t)})$$

若要求函数 $f(\beta)$ 的最小值，我们令

$$\frac{df(\beta)}{d\beta} = f'(\beta^{(t)}) + f''(\beta^{(t)})(\beta - \beta^{(t)}) = 0$$

从而可以计算出

$$\beta = \beta^{(t)} - \frac{f'(\beta^{(t)})}{f''(\beta^{(t)})}$$

令此时的最优解为 $\beta = \beta^{(t+1)}$ （并非最优）

解更新的迭代公式

$$\beta^{(t+1)} = \beta^{(t)} - \frac{f'(\beta^{(t)})}{f''(\beta^{(t)})}$$



第六章：概率建模

□ Logistic回归

Newton-Raphson（牛顿-拉斐森）优化算法（多个变量的情况——多维）：

一维时，解更新的迭代公式

$$\beta^{(t+1)} = \beta^{(t)} - \frac{f'(\beta^{(t)})}{f''(\beta^{(t)})}$$

对于**多维**时，模型由 $f(\beta) = \beta_0 + \beta_1 x$ 转化为 $f(\beta) = \beta_0 + \beta_1 x_1 + \cdots + \beta_d x_d$

那么如何计算多维的函数关于变量的一阶导数和二阶导数？？



第六章：概率建模

□ Logistic回归

Newton-Raphson（牛顿-拉斐森）优化算法（多个变量的情况——多维）：

对于**多维**时，模型由 $f(\beta) = \beta_0 + \beta_1 x$ 转化为 $f(\beta) = \beta_0 + \beta_1 x_1 + \cdots + \beta_d x_d$
回顾第三章（数学基础）中，关于对矩阵或向量的导数。

$$f'(\beta^{(t)}) = \nabla f(\beta^{(t)}) = \left[\frac{\partial f}{\partial \beta_1}, \frac{\partial f}{\partial \beta_2}, \cdots, \frac{\partial f}{\partial \beta_d} \right]^T$$

$$f''(\beta^{(t)}) = \nabla^2 f(\beta^{(t)}) = \begin{bmatrix} \frac{\partial^2 f}{\partial \beta_1^2} & \cdots & \frac{\partial^2 f}{\partial \beta_1 \partial \beta_d} \\ \vdots & \ddots & \vdots \\ \frac{\partial^2 f}{\partial \beta_d \partial \beta_1} & \cdots & \frac{\partial^2 f}{\partial \beta_d^2} \end{bmatrix} \Big|_{\beta=\beta^{(t)}} = H \text{ (Hessian矩阵)}$$



第六章：概率建模

□ Logistic回归

Newton-Raphson（牛顿-拉斐森）优化算法（多个变量的情况——多维）：

对于多维时，函数 $f(\beta)$ 的解更新优化方法为：

$$\beta^{(t+1)} = \beta^{(t)} - \frac{f'(\beta^{(t)})}{f''(\beta^{(t)})}$$



$$\beta^{(t+1)} = \beta^{(t)} - H^{-1}(\beta^{(t)}) \nabla f(\beta^{(t)})$$

迭代方法完全相同，只是在写法上，变成了矩阵向量的方式。



第六章：概率建模

□ 知识回顾（多元函数）

多元函数 $f(x_1, x_2, \dots, x_n)$ 在 $X^{(0)}$ 处的泰勒展开：

$$f(X) = f(X^{(0)}) + \nabla f(X^{(0)})^T \Delta X + \frac{1}{2} \Delta X^T G(X^{(0)}) \Delta X + \dots$$

其中，

$$(1) \nabla f(X^{(0)}) = \left[\frac{\partial f}{\partial x_1}, \frac{\partial f}{\partial x_2}, \dots, \frac{\partial f}{\partial x_n} \right] \bigg|_{X^{(0)}}^T$$

一阶梯度向量

$$(2) G(X^{(0)}) = \begin{bmatrix} \frac{\partial^2 f}{\partial x_1^2} & \frac{\partial^2 f}{\partial x_1 \partial x_2} & \cdots & \frac{\partial^2 f}{\partial x_1 \partial x_n} \\ \frac{\partial^2 f}{\partial x_2 \partial x_1} & \frac{\partial^2 f}{\partial x_2^2} & \cdots & \frac{\partial^2 f}{\partial x_2 \partial x_n} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial^2 f}{\partial x_n \partial x_1} & \frac{\partial^2 f}{\partial x_n \partial x_2} & \cdots & \frac{\partial^2 f}{\partial x_n^2} \end{bmatrix} \bigg|_{X^{(0)}}$$

海塞矩阵(Hessian Matrix, 二阶梯度)



第六章：概率建模

□ Logistic回归

最大似然函数：

化简后的对数似然函数 $L(\mathbf{w})$,

$$L(\mathbf{w}) = \sum_{i=1}^N \left[-\log \left(1 + e^{w_0 + \mathbf{w}^T \mathbf{x}_i} \right) \right] + \sum_{i=1}^N y_i (w_0 + \mathbf{w}^T \mathbf{x}_i)$$

为了估计能使 $L(\mathbf{w})$ 取得最大值的参数 $\mathbf{w}=[w_0, w_1, \dots, w_d]$, 分别求函数对 $w_0, w_1, \dots, w_d$ 的导数, 显然存在 $d+1$ 个方程。

$$\begin{aligned} \frac{\partial L(\mathbf{w})}{\partial w_j} &= - \sum_{i=1}^N \frac{1}{1 + e^{w_0 + \mathbf{w}^T \mathbf{x}_i}} \cdot e^{w_0 + \mathbf{w}^T \mathbf{x}_i} \cdot x_{ij} + \sum_{i=1}^N y_i x_{ij} \\ &= \sum_{i=1}^N (y_i - p(x_i; \mathbf{w})) x_{ij} \end{aligned}$$

此时导数表达式是一个超越方程, 不存在闭合解。采用牛顿-拉斐森方法求解。



第六章：概率建模

□ Logistic回归

最大似然函数：

Hessian矩阵H的计算：

$$p(x) = \frac{e^{h(x)}}{1 + e^{h(x)}}$$

$$\frac{\partial^2 L(w)}{\partial w_j^2} = - \sum_{i=1}^N x_{ij}^2 \cdot p(x_i; \mathbf{w}) \cdot (1 - p(x_i; \mathbf{w}))$$
$$\frac{\partial^2 L(w)}{\partial w_j \partial w_l} = - \sum_{i=1}^N x_{ij} x_{il} \cdot p(x_i; \mathbf{w}) \cdot (1 - p(x_i; \mathbf{w}))$$

$$w^{(t+1)} = w^{(t)} - H^{-1}(w^{(t)}) \nabla f(w^{(t)})$$

迭代公式



第六章：概率建模

□ 牛顿法存在的问题

- ✓ 经典牛顿法虽然具有二次收敛性，但是要求初始点需要尽量靠近极小点，否则有可能不收敛。
- ✓ 计算过程中需要计算目标函数的二阶偏导数，难度较大。
- ✓ 目标函数的Hessian矩阵无法保持正定，会导致算法产生的方向不能保证是 f 在 x_k 处的下降方向，从而令牛顿法失效；
- ✓ 如果Hessian矩阵奇异，牛顿方向可能根本是不存在的。

□ 补充：正定矩阵的定义、性质、判定方法

定义：对于 n 阶方阵 A ，如果对任何非零向量 z ，都有 $z^T A z > 0$ ，就称 A 为正定矩阵。

性质：正定矩阵一定是非奇异的、可逆的矩阵。

判定：对称阵 A 为正定的充分必要条件为 A 的特征值均大于0。



第六章：概率建模

□ 本章小结

- **Logistic**回归在本质上是线性回归。
- **Logistic**回归解决的是因变量的离散型问题，即分类问题（二类问题）。
- **Logistic**回归适合于当因变量属于二项分布这类问题。
- **Logistic**回归模型的输出是介于0-1之间的数，是一个概率值，基于**sigmoid**函数。
- 该模型的求解采用最大似然估计和牛顿法。
- 建模过程不需要对属性变量进行高斯分布的假设，但依然需要条件独立假设。



第六章：概率建模

课外任务：

根据前面的内容，请自行推导**Logistic**回归建模方法的整个过程。