

对数几率回归

张腾

2026 年 4 月 29 日

对数几率回归 (logistic regression, LR) 虽然名字中包含“回归”二字，但其本质上是一种广泛使用的线性分类模型。它通过非线性变换将线性模型的输出映射到 $[0, 1]$ 使结果具有概率意义。

1 二分类

对于二分类问题，设标记集合 $\mathcal{Y} = \{1, -1\}$ 。由贝叶斯定理，样本 $\mathbf{x}$ 属于正类的条件概率为：

$$\mathbb{P}(y = 1|\mathbf{x}) = \frac{\mathbb{P}(\mathbf{x}, y = 1)}{\mathbb{P}(\mathbf{x})} = \frac{\mathbb{P}(\mathbf{x}, y = 1)}{\mathbb{P}(\mathbf{x}, y = 1) + \mathbb{P}(\mathbf{x}, y = -1)} = 1 / \left(1 + \frac{\mathbb{P}(\mathbf{x}, y = -1)}{\mathbb{P}(\mathbf{x}, y = 1)} \right) \quad (1)$$

令

$$a = \ln \frac{\mathbb{P}(\mathbf{x}, y = 1)}{\mathbb{P}(\mathbf{x}, y = -1)} = \ln \frac{\mathbb{P}(\mathbf{x}, y = 1)}{1 - \mathbb{P}(\mathbf{x}, y = 1)}$$

为正负类联合概率比值的对数，即对数几率 (log-odds)，则式(1)可写为：

$$\mathbb{P}(y = 1|\mathbf{x}) = \frac{1}{1 + \exp(-a)} \triangleq \sigma(a)$$

其中 $\sigma: \mathbb{R} \mapsto (0, 1)$ 为对数几率函数 (logistic function)，其将对数几率映射为正类条件概率。LR 的核心假设是：对数几率 a 与输入特征 $\mathbf{x}$ 呈线性关系，即 $a = \mathbf{w}^\top \mathbf{x} + b$ 。因此，模型预测的概率为：

$$\begin{aligned} \mathbb{P}(y = 1|\mathbf{x}) &= \sigma(\mathbf{w}^\top \mathbf{x} + b) = \frac{1}{1 + \exp(-(\mathbf{w}^\top \mathbf{x} + b))} \\ \mathbb{P}(y = -1|\mathbf{x}) &= \frac{\exp(-(\mathbf{w}^\top \mathbf{x} + b))}{1 + \exp(-(\mathbf{w}^\top \mathbf{x} + b))} = \frac{1}{1 + \exp(\mathbf{w}^\top \mathbf{x} + b)} = \sigma(-(\mathbf{w}^\top \mathbf{x} + b)) \end{aligned}$$

给定数据集 $\mathcal{D} = \{(\mathbf{x}_i, y_i)\}_{i \in [m]} \subseteq \mathbb{R}^n \times \{1, -1\}$ ，我们采用极大似然估计法求解参数 $\mathbf{w}$ 和 b ：

$$\begin{aligned} \operatorname{argmax}_{\mathbf{w}, b} \ln \mathbb{P}(\mathcal{D}|\mathbf{w}, b) &= \operatorname{argmax}_{\mathbf{w}, b} \ln \prod_{i \in [m]} \mathbb{P}(\mathbf{x}_i, y_i|\mathbf{w}, b) && (\text{IID 假设}) \\ &= \operatorname{argmax}_{\mathbf{w}, b} \sum_{i \in [m]} \ln \mathbb{P}(\mathbf{x}_i, y_i|\mathbf{w}, b) && (\ln(a \cdot b) = \ln a + \ln b) \\ &= \operatorname{argmax}_{\mathbf{w}, b} \sum_{i \in [m]} (\ln \mathbb{P}(y_i|\mathbf{x}_i, \mathbf{w}, b) + \ln \mathbb{P}(\mathbf{x}_i|\mathbf{w}, b)) && (\mathbb{P}(\mathbf{x}_i, y_i) = \mathbb{P}(y_i|\mathbf{x}_i)\mathbb{P}(\mathbf{x}_i)) \\ &= \operatorname{argmax}_{\mathbf{w}, b} \sum_{i \in [m]} \ln \mathbb{P}(y_i|\mathbf{x}_i, \mathbf{w}, b) && (\mathbf{x}_i \text{ 不依赖于 } \mathbf{w}, b, \text{ 故 } \mathbb{P}(\mathbf{x}_i|\mathbf{w}, b) = \mathbb{P}(\mathbf{x}_i)) \\ &= \operatorname{argmax}_{\mathbf{w}, b} \sum_{i \in [m]} \ln \frac{1}{1 + \exp(-y_i(\mathbf{w}^\top \mathbf{x}_i + b))} \\ &= \operatorname{argmin}_{\mathbf{w}, b} \sum_{i \in [m]} \ln(1 + \exp(-y_i(\mathbf{w}^\top \mathbf{x}_i + b))) \end{aligned}$$

上述目标函数也可从交叉熵的角度来推导：类别标记的独热编码可表示为 $[(1+y)/2, (1-y)/2]$ ，模型预测分布为 $[\sigma(\mathbf{w}^\top \mathbf{x} + b), \sigma(-(\mathbf{w}^\top \mathbf{x} + b))]$ ，其经验交叉熵损失 ℓ_{CE} 化简后与上述极大似然形式一致：

$$\begin{aligned}\ell_{\text{CE}} &= - \sum_{i \in [m]} \left(\frac{1+y_i}{2} \ln \sigma(\mathbf{w}^\top \mathbf{x}_i + b) + \frac{1-y_i}{2} \ln \sigma(-(\mathbf{w}^\top \mathbf{x}_i + b)) \right) \\ &= \sum_{i \in [m]} \left(\frac{1+y_i}{2} \ln(1 + \exp(-(\mathbf{w}^\top \mathbf{x}_i + b))) + \frac{1-y_i}{2} \ln(1 + \exp(\mathbf{w}^\top \mathbf{x}_i + b)) \right) \\ &= \sum_{i \in [m]} (\mathbb{I}(y_i = 1) \ln(1 + \exp(-(\mathbf{w}^\top \mathbf{x}_i + b))) + \mathbb{I}(y_i = -1) \ln(1 + \exp(\mathbf{w}^\top \mathbf{x}_i + b))) \\ &= \sum_{i \in [m]} \ln(1 + \exp(-y_i(\mathbf{w}^\top \mathbf{x}_i + b)))\end{aligned}$$

2 多分类

对于多分类问题，设标记集合 $\mathcal{Y} = [c]$ 。LR 引入 c 个线性判别函数 $\mathbf{w}_j^\top \mathbf{x} + b_j$ ，其中 $j \in [c]$ ，并通过 softmax 函数将输出转化为概率分布：

$$\mathbb{P}(y = k | \mathbf{x}) = \sigma_k(\mathbf{x}) = \frac{\exp(\mathbf{w}_k^\top \mathbf{x} + b_k)}{\sum_{j \in [c]} \exp(\mathbf{w}_j^\top \mathbf{x} + b_j)} = \frac{\exp((\mathbf{w}_k - \mathbf{w}_c)^\top \mathbf{x} + b_k - b_c)}{\sum_{j \in [c-1]} \exp((\mathbf{w}_j - \mathbf{w}_c)^\top \mathbf{x} + b_j - b_c) + 1} \quad (2)$$

其中最后一个等号是因为 softmax 函数的参数存在冗余性，分子、分母同除以某个 $\exp(\mathbf{w}_j^\top \mathbf{x} + b_j)$ 概率值不变，故通常令 $\mathbf{w}_c = \mathbf{0}$ 、 $b_c = 0$ ，即 $\exp(\mathbf{w}_c^\top \mathbf{x} + b_c) = 1$ 以消除冗余。当 $c = 2$ 时，式(2)即退化为二分类的对数几率函数：

$$\mathbb{P}(y = 1 | \mathbf{x}) = \frac{\exp(\mathbf{w}_1^\top \mathbf{x} + b_1)}{\exp(\mathbf{w}_1^\top \mathbf{x} + b_1) + 1} = \frac{1}{1 + \exp(-(\mathbf{w}_1^\top \mathbf{x} + b_1))}$$

给定数据集 $\mathcal{D} = \{(\mathbf{x}_i, y_i)\}_{i \in [m]} \subseteq \mathbb{R}^n \times [c]$ ，极大似然估计的目标是最大化对数似然：

$$\begin{aligned}\arg\max_{\mathbf{w}_{j \in [c]}, b_{j \in [c]}} \ln \mathbb{P}(\mathcal{D} | \mathbf{w}_{j \in [c]}, b_{j \in [c]}) &= \arg\max_{\mathbf{w}_{j \in [c]}, b_{j \in [c]}} \ln \prod_{i \in [m]} \mathbb{P}(\mathbf{x}_i, y_i | \mathbf{w}_{j \in [c]}, b_{j \in [c]}) \\ &= \arg\max_{\mathbf{w}_{j \in [c]}, b_{j \in [c]}} \sum_{i \in [m]} \ln \mathbb{P}(y_i | \mathbf{x}_i, \mathbf{w}_{j \in [c]}, b_{j \in [c]}) \\ &= \arg\max_{\mathbf{w}_{j \in [c]}, b_{j \in [c]}} \sum_{i \in [m]} \ln \frac{\exp(\mathbf{w}_{y_i}^\top \mathbf{x}_i + b_{y_i})}{\sum_{j \in [c]} \exp(\mathbf{w}_j^\top \mathbf{x}_i + b_j)} \\ &= \arg\max_{\mathbf{w}_{j \in [c]}, b_{j \in [c]}} \sum_{i \in [m]} \ln \sigma_{y_i}(\mathbf{x}_i)\end{aligned}$$

为了求解该优化问题，我们计算其梯度：

$$\begin{aligned}\frac{\partial \sigma_{y_i}(\mathbf{x}_i)}{\partial \mathbf{w}_k} &= \frac{\partial}{\partial \mathbf{w}_k} \left(\frac{\exp(\mathbf{w}_{y_i}^\top \mathbf{x}_i + b_{y_i})}{\sum_{j \in [c]} \exp(\mathbf{w}_j^\top \mathbf{x}_i + b_j)} \right) \\ &= \frac{\mathbb{I}(k = y_i) \exp(\mathbf{w}_{y_i}^\top \mathbf{x}_i + b_{y_i}) \mathbf{x}_i \sum_{j \in [c]} \exp(\mathbf{w}_j^\top \mathbf{x}_i + b_j) - \exp(\mathbf{w}_{y_i}^\top \mathbf{x}_i + b_{y_i}) \exp(\mathbf{w}_k^\top \mathbf{x}_i + b_k) \mathbf{x}_i}{(\sum_{j \in [c]} \exp(\mathbf{w}_j^\top \mathbf{x}_i + b_j))^2} \\ &= \mathbb{I}(k = y_i) \sigma_{y_i}(\mathbf{x}_i) \mathbf{x}_i - \sigma_{y_i}(\mathbf{x}_i) \sigma_k(\mathbf{x}_i) \mathbf{x}_i \\ &= \sigma_{y_i}(\mathbf{x}_i) (\mathbb{I}(k = y_i) - \sigma_k(\mathbf{x}_i)) \mathbf{x}_i\end{aligned}$$

于是对数似然关于 $\mathbf{w}_k$ 的梯度为

$$\frac{\partial \sum_{i \in [m]} \ln \sigma_{y_i}(\mathbf{x}_i)}{\partial \mathbf{w}_k} = \sum_{i \in [m]} \frac{1}{\sigma_{y_i}(\mathbf{x}_i)} \frac{\partial \sigma_{y_i}(\mathbf{x}_i)}{\partial \mathbf{w}_k} = \sum_{i \in [m]} (\mathbb{I}(k = y_i) - \sigma_k(\mathbf{x}_i)) \mathbf{x}_i$$

同理可得对数似然关于 b_k 的梯度为

$$\frac{\partial \sum_{i \in [m]} \ln \sigma_{y_i}(\mathbf{x}_i)}{\partial b_k} = \sum_{i \in [m]} \frac{1}{\sigma_{y_i}(\mathbf{x}_i)} \frac{\partial \sigma_{y_i}(\mathbf{x}_i)}{\partial b_k} = \sum_{i \in [m]} (\mathbb{I}(k = y_i) - \sigma_k(\mathbf{x}_i))$$

令梯度为零，我们得到最优解应满足的平衡方程组：

$$\sum_{i \in [m]} \mathbb{I}(k = y_i) \mathbf{x}_i = \sum_{i \in [m]} \sigma_k(\mathbf{x}_i) \mathbf{x}_i, \quad \sum_{i \in [m]} \mathbb{I}(k = y_i) = \sum_{i \in [m]} \sigma_k(\mathbf{x}_i), \quad \forall k \in [c] \quad (3)$$

式(3)表明：在最优解处，第 k 类样本的特征之和（及数量）等于所有样本以 softmax 变换给出的概率为权重的加权和。这与马尔可夫过程中的细致平稳 (detailed balance) 条件形式相似。

3 最大熵视角

最大熵 (maximum entropy) 原理是概率建模的一个重要准则：在满足已知约束的前提下，应选择熵最大的分布，即对未知信息不做任何额外假设（视为等概率）。例如，对于一个未知骰子，在没有任何信息的情况下，只能假设掷出各数的概率均为 $1/6$ ；若已知掷出 1 的概率为 $1/4$ ，则掷出其它数的概率应均分剩余的 $3/4$ 。

定理 3.1. 对数几率回归是满足平衡方程组(3)约束下的最大熵模型。

证明. 设 $\pi_k(\mathbf{x}_i)$ 为模型将样本 $\mathbf{x}_i$ 预测为第 k 类的概率，最大化所有样本的预测熵之和：

$$\begin{aligned} \max_{\pi_1, \dots, \pi_c} \quad & \sum_{i \in [m]} \left(- \sum_{k \in [c]} \pi_k(\mathbf{x}_i) \ln \pi_k(\mathbf{x}_i) \right) \\ \text{s.t.} \quad & \pi_k(\mathbf{x}_i) \geq 0, \quad \forall i \in [m], \quad \forall k \in [c] && (\text{概率非负}) \\ & \sum_{k \in [c]} \pi_k(\mathbf{x}_i) = 1, \quad \forall i \in [m] && (\text{概率归一化}) \\ & \sum_{i \in [m]} \mathbb{I}(k = y_i) \mathbf{x}_i = \sum_{i \in [m]} \pi_k(\mathbf{x}_i) \mathbf{x}_i, \quad \forall k \in [c] && (\text{特征矩匹配}) \\ & \sum_{i \in [m]} \mathbb{I}(k = y_i) = \sum_{i \in [m]} \pi_k(\mathbf{x}_i), \quad \forall k \in [c] && (\text{样本数匹配}) \end{aligned}$$

引入拉格朗日乘子： α_i 对应概率归一化、 β_k 对应特征矩匹配、 γ_k 对应样本数匹配，构建拉格朗日函数（暂忽略概率非负约束，因其由指数形式自动满足）：

$$\begin{aligned} L(\pi_k, \alpha_i, \beta_k, \gamma_k) = & - \sum_{i \in [m]} \sum_{k \in [c]} \pi_k(\mathbf{x}_i) \ln \pi_k(\mathbf{x}_i) - \sum_{i \in [m]} \alpha_i \left(\sum_{k \in [c]} \pi_k(\mathbf{x}_i) - 1 \right) \\ & - \sum_{k \in [c]} \beta_k^\top \left(\sum_{i \in [m]} (\mathbb{I}(k = y_i) - \pi_k(\mathbf{x}_i)) \mathbf{x}_i \right) - \sum_{k \in [c]} \gamma_k \left(\sum_{i \in [m]} (\mathbb{I}(k = y_i) - \pi_k(\mathbf{x}_i)) \right) \end{aligned}$$

对 π_k 求偏导并令其为零：

$$\frac{\partial L}{\partial \pi_k} = -\ln \pi_k(\mathbf{x}_i) - 1 - \alpha_i + \beta_k^\top \mathbf{x}_i + \gamma_k = 0$$

解得：

$$\pi_k(\mathbf{x}_i) = \exp(-1 - \alpha_i + \beta_k^\top \mathbf{x}_i + \gamma_k)$$

利用归一化约束 $\sum_{k \in [c]} \pi_k(\mathbf{x}_i) = 1$ 消去 α_i :

$$\exp(1 + \alpha_i) = \sum_{k \in [c]} \exp(\beta_k^\top \mathbf{x}_i + \gamma_k) \implies \pi_k(\mathbf{x}_i) = \frac{\exp(\beta_k^\top \mathbf{x}_i + \gamma_k)}{\sum_{j \in [c]} \exp(\beta_j^\top \mathbf{x}_i + \gamma_j)}$$

这正是 softmax 函数的形式。

