

## 1 GLM: Logistic regression

**Logistic regression:** A model predicts if the patient is a vegan ( $y = 1$ ) or not ( $y = 0$ ) by a result of cholesterol test and gets the result  $x$  mmol/L. A binary response modeled with Bernoulli distribution  $y \sim \mathcal{B}(p)$ , where  $p := \mathbb{P}[y = 1]$  is the probability of being a vegan.

Bernoulli distribution belongs to the exponential family; in canonical form, it is:

$$y \sim f(y|\theta) = (1 - \sigma(\theta)) \cdot e^{\theta \cdot y},$$

where  $\theta = \text{logit}(p) = \log\left(\frac{p}{1-p}\right)$  is the logit function.

Fitting the parameter  $\theta$  on the historical data:

$$\theta^* = \arg \max_{\theta} \sum_{(x^*, y^*) \in (X, Y)^\ell} f(y = y^* | x = x^*, \theta)$$

Then, to make a prognosis, the probability of being a vegan is calculated:

In logistic regression, the connection between the input  $x$  and the probability  $p$  is modeled as:

$$\underbrace{\log \frac{p}{1-p}}_{\text{logit } p} = \beta_0 + \beta_1 \cdot x$$

The inverse of the logit function is the sigmoid function  $\sigma(\theta) = (1 + e^{-\theta})^{-1}$ :

$$p = \text{logit}^{-1}(\beta_0 + \beta_1 \cdot x) = \sigma(\beta_0 + \beta_1 \cdot x)$$

The parameters  $\beta_0$  and  $\beta_1$  are trained on the historical data:

$$\ell(\beta_0, \beta_1) = \log \prod_i p_i^{y_i} (1 - p_i)^{1-y_i} \rightarrow \max_{\beta_0, \beta_1}.$$

Then, to make a prognosis, the probability of a bad outcome is calculated:

$$\hat{p}(x) = \sigma(\beta_0 + \beta_1 \cdot x).$$

**NB:** We used distribution of  $y \sim \mathcal{B}(p)$  to derive the loss function:

$$\ell = \log \prod_i \mathbb{P}[y_i = 1]^{y_i} \cdot \mathbb{P}[y_i = 0]^{1-y_i},$$

but we ignore the distribution of  $y$  when we make a prediction, we are only interested in

**Exponential family:** Bernoulli distribution  $\mathcal{B}(p)$  belongs to the exponential family, the exponential family parameter  $\theta$  can be calculated from the Bernoulli parameter  $p$ , then

$$y \sim \text{Exp}(\theta).$$

## 2 Generalized Linear Models (GLM)

**Introduction to GLM:** A generalized linear model (GLM) extends ordinary linear regression by allowing for response variables that follow any exponential family distribution. The general form is:

$$Y \sim f(y \mid \boldsymbol{\theta}) = \exp[\boldsymbol{\theta} \cdot T(y) - A(\boldsymbol{\theta}) + C(y)] \quad (1)$$

**Making Predictions:** To make a prediction in GLM, we estimate the conditional expectation (canonical mean parameter):

$$\hat{y}(\mathbf{x}) := \mathbb{E}[Y \mid X = \mathbf{x}] \equiv \mu \quad (2)$$

For most cases, the sufficient statistics  $T$  is trivial, and we can obtain the needed expectation from the distribution parameters:

$$\hat{y}(\mathbf{x}) := \mathbb{E}[Y \mid \boldsymbol{\theta}] = \mathbb{E}[Y \mid \boldsymbol{\theta} = F\boldsymbol{\beta}] = \mathbb{E}[Y \mid X = \mathbf{x}] \quad (3)$$

For example, in logistic regression, the mean parameter corresponds to probability:

$$\mu := \mathbb{E}[Y \mid \boldsymbol{\theta}] = \mathbb{P}[Y = 1 \mid \boldsymbol{\theta}] = \mathbb{P}[Y = 1 \mid X = \mathbf{x}] = \hat{p} \quad (4)$$

### Mean and Link Functions:

#### 231 Mean Function

The mean function describes the expected value of the response variable  $Y$  (or sufficient statistics  $T(Y)$ ) given current parameters:

$$\boldsymbol{\mu} := \mathbb{E}[T(Y) \mid \boldsymbol{\theta}] \quad (5)$$

#### 232 Link Function

The link function connects linear parameters  $\boldsymbol{\eta} = X\boldsymbol{\beta}$  (linear predictor) to the expected value (canonical mean):

$$\boldsymbol{\mu} = \boldsymbol{\psi}(\boldsymbol{\eta}) \quad (6)$$

Its inverse calculates parameters:

$$\boldsymbol{\eta} = \boldsymbol{\psi}^{-1}(\boldsymbol{\mu}) \quad (7)$$

**GLM as Linear + Nonlinear Transforms:** GLM combines linear and nonlinear transformations:

1. Linear predictor computation:

$$\boldsymbol{\eta}(\mathbf{x}) = \boldsymbol{\beta}^\top \mathbf{x} = M\mathbf{x} \quad (8)$$

$$\boldsymbol{\eta} = X\boldsymbol{\beta} \quad (9)$$

2. Link function application:

$$\boldsymbol{\mu} : \boldsymbol{\eta} = \boldsymbol{\psi}(\boldsymbol{\mu}) \quad (10)$$

$$\boldsymbol{\psi}(\mathbb{E}[y|\mathbf{x}]) = \boldsymbol{\beta}^\top \mathbf{x} \quad (11)$$

$$\boldsymbol{\psi}(\mathbb{E}[y|X]) = X\boldsymbol{\beta} \quad (12)$$

3. Final prediction via inverse link function:

$$\hat{y}(\mathbf{x}) = \mathbb{E}[y|\mathbf{x}] = \boldsymbol{\psi}^{-1}(\boldsymbol{\beta}^\top \mathbf{x}) \quad (13)$$

**Logistic Regression as GLM:** Logistic regression is a special case of GLM using Bernoulli distribution:

$$y_i \sim \mathcal{B}(p), \quad \mathbb{P}[y_i = 1] = p \quad (14)$$

In canonical form:

$$y_i \sim f(y \mid \theta) = (1 - \sigma(\theta)) \cdot e^{\theta \cdot y} \quad (15)$$

The link function can be found from:

$$\psi = (\nabla_{\theta} A)^{-1} \quad (16)$$

Where:

$$\frac{\partial A}{\partial \theta} = \sigma(\theta) = \frac{1}{1 + e^{-\theta}} \quad (17)$$

$$\psi(\mu) = \sigma^{-1}(p) = \ln \frac{p}{1-p} = \text{logit } p \quad (18)$$

### 3 GLM: Cross-entropy and log-loss

**Model:** Logistic regression represents a special case of GLM where the binary response variable  $Y$  follows a Bernoulli distribution:

$$y_i \sim \mathcal{B}(p), \quad p := \mathbb{P}[y_i = 1] \quad (19)$$

Here,  $p$  represents the success probability in a single trial. The canonical form of the Bernoulli distribution is:

$$y_i \sim f(y|\theta) = \sigma(-\theta) \cdot e^{\theta y}, \quad \sigma(\theta) = \frac{1}{1 + e^{-\theta}} \quad (20)$$

Starting from the general GLM form:

$$Y \sim f(y|\theta) = \exp[\theta \cdot T(y) - A(\theta) + C(y)] \quad (21)$$

We can derive both cross-entropy and log-loss directly, assuming only the Bernoulli distribution of  $Y$ .

**Link Function:** The link function  $\psi$  connects the response variable's mean  $\mu = \mathbb{E}[Y]$  to the distribution's canonical parameters  $\theta$ :

$$\mu = \psi(\theta) \quad (22)$$

In GLM, we assume the canonical parameters are linear:

$$\theta_i = \mathbf{x}_i^\top \boldsymbol{\beta}, \quad \theta = \mathbf{X} \boldsymbol{\beta} \quad (23)$$

where  $\boldsymbol{\beta}$  represents the linear coefficients corresponding to features in  $\mathbf{x}$ .

For the Bernoulli distribution, the link function takes the form:

$$\psi(\mu) = \log \frac{\mu}{1 - \mu} = \text{logit } \mu \quad (24)$$

**Cross-entropy Loss:** We begin with the log-likelihood function  $l(\theta)$  for the Bernoulli-distributed response variable  $Y$ , assuming  $\theta = \mathbf{x}^\top \boldsymbol{\beta}$ :

$$\begin{aligned} l(\theta) &= \log \prod_i f(y_i|\theta) \\ &= \log \prod_i \sigma(-\theta) \cdot e^{\theta y_i} \\ &= \sum_i \{\theta \cdot y_i + \log \sigma(-\theta)\} \\ &= \sum_i \left\{ \theta \cdot y_i + \log \frac{1}{1 + e^{-(\theta)}} \right\} \\ &= \sum_i \left\{ y_i \log \frac{\mu}{1 - \mu} + \log \frac{1}{1 + \frac{\mu}{1 - \mu}} \right\} \quad (25) \\ &= \sum_i \left\{ y_i \log \frac{\mu}{1 - \mu} + \log \frac{1 - \mu}{1 - \mu + \mu} \right\} \\ &= \sum_i \{y_i \log \mu - y_i \log(1 - \mu) + \log(1 - \mu)\} \\ &= \sum_i \{y_i \log \mu + (1 - y_i) \log(1 - \mu)\} \\ &= \sum_i \{y_i \log p + (1 - y_i) \log(1 - p)\} \\ &= l(p(\boldsymbol{\beta})) \rightarrow \max_{\boldsymbol{\beta}} \end{aligned}$$

**Log-loss:** The log-loss function  $\ell(M)$  can be derived by taking the negative log-likelihood:

$$\begin{aligned} -l(\theta) &= -\sum_i \left\{ \theta \cdot y_i + \log \frac{e^{-\theta}}{1 + e^{-\theta}} \right\} \\ &= \sum_i \begin{cases} -\log e^\theta + \log \frac{e^{-\theta}}{1 + e^{-\theta}}, & \text{if } y = 1 \\ -\log \frac{e^{-\theta}}{1 + e^{-\theta}}, & \text{if } y = 0 \end{cases} \\ &= \sum_i \begin{cases} \log(1 + e^{-\theta}), & \text{if } y = 1 \\ \log(1 + e^\theta), & \text{if } y = 0 \end{cases} \quad (26) \\ &= \sum_i \log(1 + e^{\theta \cdot \text{sgn } y_i}) \\ &= \sum_i \log(1 + e^{\langle \mathbf{x}_i, \boldsymbol{\beta} \rangle \cdot \text{sgn } y_i}) \\ &= \sum_i \log(1 + e^{-M_i}) \\ &= \ell(M(\boldsymbol{\beta})) \rightarrow \min_{\boldsymbol{\beta}} \end{aligned}$$

**Making Predictions:** To make a prediction:

$$\hat{p}(\mathbf{x}) = \mu(\mathbf{x}) = \psi(\theta = \mathbf{x} \cdot \boldsymbol{\beta}) = \frac{1}{1 + e^{\mathbf{x} \cdot \boldsymbol{\beta}}} \quad (27)$$