

Lecture 9 | 28.04.2025

Marginal models for non-normal response

GLM extensions for the longitudinal data

❑ **Marginal models**

- ❑ primary interest is given to the conditional mean structure
- ❑ separate models for the mean and the covariance structure

❑ **Random effects models**

- ❑ one equation used to account for both—the mean and the covariance
- ❑ mostly used when some subject specific inference is of the main interest

❑ **Transition models**

- ❑ primary interest again with respect to the mean structure
- ❑ the correlation structure due to historical observations within the subject

All three categories of the regression models for correlated (repeated) observations above lead to the same model (with the same interpretation) for the Gaussian type of the data but for the discrete data different models typically produce different interpretations (due to the non-linearity involved in the models)

Marginal models

- ❑ For simplicity, let $\mathbf{Y}_i = (Y_{i1}, \dots, Y_{in_i})^\top$ denote a vector of $n_i \in \mathbb{N}$ correlated binary responses for some individual $i \in \{1, \dots, N\}$
- ❑ The idea is to model $P[\mathbf{Y}_i = \mathbf{y} | \mathbb{X}_i]$, for $\mathbf{y} \in \{0, 1\}^{\otimes n_i}$ by using the marginals of the joint distribution (conditionally on $\mathbb{X}_i$)
- ❑ The saturated model has $2^{n_i} - 1$ parameters and different “marginals”
 - ❑ First order marginals $\mu_j = P[Y_{ij} = 1]$, for $j = 1, \dots, n_i$
 - ❑ Second order marginals $\mu_{jk} = P[Y_{ij} = 1, Y_{ik} = 1]$, for $j \neq k$
 - ❑ Third order marginals $\mu_{jkl} = P[Y_{ij} = 1, Y_{ik} = 1, Y_{il} = 1]$, for $j \neq k \neq l$
 - ❑ ...
 - ❑ The n_i^{th} order marginal $\mu_{1, \dots, n_i} = P[\mathbf{Y}_i = \mathbf{1}]$, where $\mathbf{1} = (1, \dots, 1)^\top \in \mathbb{R}^{n_i}$
- ❑ Which marginals should be used in the model and how should they explain the overall joint probability $P[\mathbf{Y}_i = \mathbf{y} | \mathbb{X}_i]$ (always conditionally on $\mathbb{X}_i$)
 - ❑ full log-linear model
 - ❑ log-linear model for the first and higher order marginals (GEE formulation)
 - ❑ Bahadur model for the first order marginals and correlations
 - ❑ marginal model for μ_j and marginal odds ratios
 - ❑ ...

Full log-linear model

- the joint probability $P[\mathbf{Y}_i = \mathbf{y}]$ for $\mathbf{y} = (y_1, \dots, y_{n_i})^\top \in \{0, 1\} \times \dots \times \{0, 1\}$ can be expressed as $P[\mathbf{Y}_i = \mathbf{y}] = P[Y_{i1} = y_1 \wedge Y_{i2} = y_2 \wedge \dots \wedge Y_{in_i} = y_{n_i}]$
- there are many options how to define a saturated model (with 2^{n_i} total parameters but only $2^{n_i} - 1$ free parameters)
- one commonly used model (proposed by Bishop et al. 1975) is a log-linear model which can be formulated as

$$P[\mathbf{Y}_i = \mathbf{y}] = c(\theta_i) \exp \left\{ \sum_{j=1}^n \theta_{ij}^{(1)} y_j + \sum_{j_1 < j_2} \theta_{ij_1 j_2}^{(2)} y_{j_1} y_{j_2} + \dots + \theta_{ij_1 \dots j_{n_i}}^{(n_i)} y_1 \dots y_{n_i} \right\}$$

for the vector $\theta_i = (\theta_{i1}^{(1)}, \dots, \theta_{in_i}^{(1)}, \theta_{i12}^{(2)}, \dots, \theta_{i1\dots n_i}^{(n)})^\top \in \mathbb{R}^{2^{n_i}-1}$ which is the canonical vector of the unknown model parameters (where $\theta_{ij}^{(1)}$ are conditional log odds and $\theta_{i\dots}^{(k)}$ for $k \geq 2$ are conditional log odds ratios)

- however, the association between Y_j and Y_k depends on all other values of Y_l for $l \neq j, k$, respectively

$$\log \left[\frac{P[Y_{ij} = 1 | Y_{ik} = y_k, Y_{il} = 0 \forall l \neq j, k]}{P[Y_{ij} = 0 | Y_{ik} = y_k, Y_{il} = 0 \forall l \neq j, k]} \right] = \theta_{ij}^{(1)} + \theta_{ijk}^{(2)} y_k$$

- it would be more interesting to model $P[Y_j = 1 | \mathbf{X}] = E[Y_j | \mathbf{X}] = \mu_j$
(respectively $P[Y_{ij} = 1 | \mathbf{X}_i] = E[Y_{ij} | \mathbf{X}_i] = \mu_{ij}$, for a given subject $i \in \{1, \dots, N\}$)

Marginal models towards GEE

□ Mean structure

The marginal (conditional) expectation of the response depends (non-linearly) on a linear combination of the explanatory variables (i.e., linear predictor)

$$h(\mu_{ij}) = \mathbf{X}_{ij}^{\top} \boldsymbol{\beta}, \quad \text{for } \mu_{ij} = E[Y_{ij} | \mathbf{X}_{ij}] \text{ and } \boldsymbol{\beta} \in \mathbb{R}^p$$

for a known, strictly monotone, and twice continuously differentiable function h

□ Variance structure

The marginal (conditional) variance of the response depends on the marginal mean (and, optionally, some other overdispersion parameter $\phi > 0$) as

$$\text{Var}(Y_{ij} | \mathbf{X}_{ij}) = v(\mu_{ij})\phi, \quad \text{for } \phi > 0$$

for a known positive and continuously differentiable function v

□ Covariance structure

The correlation between two observations Y_{ij} and Y_{ik} (within the same subject $i \in \{1, \dots, N\}$) is assumed to be modeled as

$$\text{Cor}(Y_{ij}, Y_{ik} | \mathbf{X}_{ij}, \mathbf{X}_{ik}) = \rho(\mu_{ij}, \mu_{ik}, \boldsymbol{\alpha}), \quad \text{for } \boldsymbol{\alpha} \in \mathbb{R}^q$$

for a known covariance function ρ

Key pivots of the marginal models

- Instead of specifying the whole distribution (i.e., the exponential family of distributions) which is required, for instance, for the likelihood based estimation in GLM, only a specification of the first two moments (and their mutual relationship) is provided (quasi-likelihood and GEE instead)
- For $\mathbf{Y}_i = (Y_{i1}, \dots, Y_{in_i})^\top$ and $\mathbb{X}_i = (\mathbf{X}_{i1}, \dots, \mathbf{X}_{in_i})^\top$ we separately specify the mean $E[\mathbf{Y}_i | \mathbb{X}_i] = \mathbb{X}_i \boldsymbol{\beta}$ and the covariance $\text{Var}[\mathbf{Y}_i | \mathbb{X}_i] = \mathbb{V}_i(\mathbb{X}_i, \boldsymbol{\beta}, \boldsymbol{\phi}, \boldsymbol{\alpha})$
- As the distribution is not provided the likelihood based on the data can not be constructed. Therefore, in a sense, this is not a parametric model (where the parameters specify the whole distribution) but rather a semi-parametric one (the parameters only specify the first two moments)
- In a log-linear model (with the first ($\boldsymbol{\beta}$) and higher order ($\boldsymbol{\alpha}$) marginals) the score function for $\boldsymbol{\beta} \in \mathbb{R}^p$ leads to a GEE formulation with the equations

$$\frac{\partial \boldsymbol{\mu}}{\partial \boldsymbol{\beta}^\top} [\text{Var}(\mathbf{Y})]^{-1} (\mathbf{Y} - \boldsymbol{\mu}) = \mathbf{0}$$

- Therefore, in order to use the quasi-likelihood estimation approach appropriately, one has to correctly specify both, the mean and the variance-covariance function (structure)

Overview: Maximum likelihood for GLM

- for some generic random vector $(Y, \mathbf{X}^\top)^\top \sim F_{(Y, \mathbf{X})}$ we assume that $f(y|\mathbf{X}) = \exp\{\phi^{-1}(y\theta - \psi(\theta)) + c(y, \phi)\}$ (i.e., the exponential family)
- the linear predictor $\theta = \mathbf{X}^\top \beta$ is associated with the conditional mean $E[Y|\mathbf{X}] = \mu$ via the link function g , such that $g(\mu) = \theta \equiv \theta(\beta) = \mathbf{X}^\top \beta$
- as far as $f(y|\mathbf{X})$ is a probability density function, it holds that $\int f(y|\mathbf{X}) dy = 1$ (by integrating with respect to an appropriate measure)
- thus, the first and the second partial derivatives of $\int f(y|\mathbf{X}) dy = 1$ with respect to θ (scores) yield the following two equations

$$\frac{\partial}{\partial \theta} : \int [y - \psi'(\theta)] f(y|\mathbf{X}) dy = 0$$

and

$$\frac{\partial^2}{\partial \theta^2} : \int [\psi^{-1}(y - \psi'(\theta))^2 - \psi''(\theta)] f(y|\mathbf{X}) dy = 0$$

which gives $\mu = E[Y|\mathbf{X}] = \psi'(\theta)$ and $\text{Var}[Y|\mathbf{X}] = \phi \psi''[(\psi')^{-1}(\mu)]$

Score equations under MLE

- for the random sample $\mathcal{D}_S = \{(Y_i, \mathbf{X}_i); i = 1, \dots, N\}$ we have $f(y|\mathbf{X}_i) = \exp\{\phi^{-1}(y\theta_i - \psi(\theta_i)) + c(y, \phi)\}$ where $\theta_i = \mathbf{X}_i^\top \beta \equiv \theta_i(\beta)$
- as $\theta_i = \mathbf{X}_i^\top \beta$ and we aim to estimate the unknown parameter vector $\beta \in \mathbb{R}^p$, we need partial derivatives with respect to $\beta \in \mathbb{R}^p$

□ Log-likelihood

$$\ell(\beta, \phi, \mathcal{D}_S) = \frac{1}{\phi} \sum_{i=1}^N [Y_i \theta_i - \psi(\theta_i)] + \sum_{i=1}^N c(Y_i, \phi)$$

□ First order derivatives wrt. β

$$\frac{\partial \ell(\beta, \phi, \mathcal{D}_S)}{\partial \beta} = \frac{1}{\phi} \sum_{i=1}^N \frac{\partial \theta_i}{\partial \beta} [Y_i - \psi'(\theta_i)]$$

□ Score equations for β

$$S(\beta) = \sum_{i=1}^N \frac{\partial \theta_i}{\partial \beta} [Y_i - \psi'(\theta_i)] = \mathbf{0}$$

Score equations under MLE – continuation

- since $\mu_i = \psi'(\theta_i)$ and $v_i = v(\mu_i) = \psi''(\theta_i)$ the score equations become

$$S(\beta) = \sum_{i=1}^N \frac{\partial \mu_i}{\partial \beta} v_i^{-1} [Y_i - \psi'(\theta_i)] = \mathbf{0}$$

and, thus, the estimation of $\beta \in \mathbb{R}^p$ depends on the exponential distribution only through the mean μ_i and the variance function $v_i = v(\mu_i)$

- the solution to the equations above is obtained numerically (e.g., Newton-Raphson, iterative re-weighted LS, Fisher scoring)
- the inference about $\beta \in \mathbb{R}^p$ is based on the classical maximum likelihood theory (i.e., asymptotic Wald tests, likelihood ratio tests, score tests)
- the estimation of the over-dispersion parameter based on residuals

$$\hat{\phi} = \frac{1}{N - p} \sum_{i=1}^N \frac{[Y_i - \hat{\mu}_i]^2}{v_i(\hat{\mu}_i)}$$

General Estimating Equations (GEE)

- for independent observations $Y_1, \dots, Y_N$ (within the GLM framework) the corresponding score equations for estimating $\beta \in \mathbb{R}^p$ are

$$S(\beta) = \sum_{i=1}^N \frac{\partial \mu_i}{\partial \beta} v_i^{-1} [Y_i - \mu_i] = \mathbf{0}$$

- for longitudinal observations $\mathbf{Y}_1, \dots, \mathbf{Y}_n$ (within the GEE framework) the score equations for $\beta \in \mathbb{R}^p$ can be seen as a multivariate extension

$$S(\beta) = \sum_{i=1}^N \sum_{j=1}^{n_i} \frac{\partial \mu_{ij}}{\partial \beta} v_{ij}^{-1} [Y_{ij} - \mu_{ij}] = \mathbf{0}$$

- which can be also expressed in a more common (as a sum of independent subjects) matrix notation

$$S(\beta) = \sum_{i=1}^N \mathbb{D}_i^\top [\mathbb{V}_i(\alpha)]^{-1} [\mathbf{Y}_i - \boldsymbol{\mu}_i] = \mathbf{0}$$

where $\mathbb{D}_i = \left(\frac{\partial \mu_{ij}}{\partial \beta_k} \right)_{j,k=1}^{n_i, p}$ and $\mathbb{V}_i(\alpha) \equiv \mathbb{V}_i(\mathbb{X}_i, \beta, \phi, \alpha)$

Correlation structure within $\mathbf{Y}_i$

- Note, that the variance matrix $\text{Var}[\mathbf{Y}_i|\mathbf{X}_i]$ is relatively complex and specific structural decomposition is typically used to model the variance-covariance structure more carefully

$$\text{Var}[\mathbf{Y}_i|\mathbf{X}_i] = \mathbb{V}_i(\mathbf{X}_i, \beta, \phi, \alpha) = \phi \mathbb{A}_i^{1/2}(\beta) \mathbb{R}_i(\alpha) \mathbb{A}_i^{1/2}(\beta)$$

where the matrix $\mathbb{A}_i^{1/2}(\beta)$ models the variance of the repeated observations for the given subject $i \in \{1, \dots, N\}$

$$\mathbb{A}_i^{1/2}(\beta) = \begin{pmatrix} \sqrt{v_{i1}(\mu_{i1})} & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & \sqrt{v_{in_i}(\mu_{in_i})} \end{pmatrix}$$

and the correlation of the repeated observations is modeled by $\mathbb{R}_i(\alpha)$

- Recall, that the covariances and variances follow from the mean structure... for modeling purposes we specify the mean structure and the correlations (i.e., the working correlation matrix)

Statistical properties and inference

- the GEE estimates of $\beta \in \mathbb{R}^p$ are consistent even if the working correlation matrix is incorrect

$$\sqrt{N}(\hat{\beta}_N - \beta) \underset{as.}{\sim} N_p(\mathbf{0}, \mathbb{I}_0^{-1} \mathbb{I}_1 \mathbb{I}_0^{-1})$$

where $\mathbb{I}_0$ is the limit matrix of $\sum_{i=1}^N \mathbb{D}_i^\top [\mathbb{V}_i(\beta)]^{-1} \mathbb{D}_i$ and, analogously also $\mathbb{I}_1$ is the limit matrix of $\sum_{i=1}^N \mathbb{D}_i^\top [\mathbb{V}_i(\beta)]^{-1} \text{Var}(\mathbf{Y}_i) [\mathbb{V}_i(\beta)]^{-1} \mathbb{D}_i$

- note that if the variance matrix $\text{Var}(\mathbf{Y}_i)$ is correctly specified, the asymptotic variance only reduces to $\mathbb{I}_0^{-1}$ (i.e., the likelihood variance)
- the estimate for $\mathbb{I}_1$ can be obtained by replacing $\text{Var}(\mathbf{Y}_i)$ by $(\mathbf{Y}_i - \hat{\mu}_i)(\mathbf{Y}_i - \hat{\mu}_i)^\top$ which usually leads to a good estimate of $\mathbb{I}_1$ even if it is a bad estimate for $\text{Var}(\mathbf{Y}_i)$

Alternatives and extensions

- ❑ **Prentice's two sets of GEE**

(two sets of GEE equations—one to obtain the estimates for β and the other one to get the estimates for α)

- ❑ **GEE based on linearization**

(linearization in a form of $Y_{ij} = \mu_{ij} + \varepsilon_{ij}$, where $\varepsilon_{ij} = \mu_{ij}$ with probability $1 - \mu_{ij}$ and $\varepsilon_{ij} = 1 - \mu_{ij}$ with probability μ_{ij})

- ❑ **GEE2 up to GGE k generalizations**

(extended marginal mean structure considering the first and second (possibly up to k^{th} order) marginals and pairwise associations)

- ❑ **Alternating Logistic Regression (ALR)**

(parameters β and α are estimated in two separate alternating regression formulations that are iterated until convergence)

Alternating logistic regression

- when the response variable only takes two possible values $Y_{ij} \in \{0, 1\}$ (i.e., logistic regression) the mean and the correlation structure can be estimated using two alternating regression models
- first order marginals are used to model the conditional mean structure using the equation $\text{logit}(E[Y_{ij}|\mathbf{X}_{ij}]) = \mathbf{X}_{ij}^\top \boldsymbol{\beta}$ and the marginal odds ratios are used to model the associations

$$\text{logit } P[Y_{ij} = 1 | Y_{ik} = y_{ik}] = \alpha_{ijk} y_{ik} + \ln \left(\frac{P[Y_{ij} = 1, Y_{ik} = 1] P[Y_{ij} = 0, Y_{ik} = 0]}{P[Y_{ij} = 0, Y_{ik} = 1] P[Y_{ij} = 1, Y_{ik} = 0]} \right)$$

where $\alpha_{ijk} y_{ik} \in \mathbb{R}$ is modeled as a predictor variable and some unknown parameter and the odds ratio is an offset parameter (intercept)

- the alternating logistic regression is (almost) as efficient as GEE2 and (almost) as computationally easy as GGE

Summary

- ❑ Marginal models for correlated observations are specifically suitable for population interpretation and population based inference
- ❑ Different strategies can be used to build the model using the marginals of the joint distribution $P[\mathbf{Y}_i = \mathbf{y} | \mathbb{X}_i]$
- ❑ Different models imply different interpretation of the estimated parameters and also different limitations for a practical utilization
- ❑ For a subject specific interpretation another models need to be used – for instance, models with random effects