



Electronic Journal of Applied Statistical Analysis
EJASA, Electron. J. App. Stat. Anal.

<http://siba-ese.unisalento.it/index.php/ejasa/index>

e-ISSN: 2070-5948

DOI: 10.1285/i20705948v18n2p299

**Bayes Analysis of mRS Followup Data in the
Presence of Multiple Categories**

By Pandey et al.

15 October 2025

This work is copyrighted by Università del Salento, and is licensed under a Creative Commons Attribuzione - Non commerciale - Non opere derivate 3.0 Italia License.

For more information see:

<http://creativecommons.org/licenses/by-nc-nd/3.0/it/>

Bayes Analysis of mRS Followup Data in the Presence of Multiple Categories

Pranjal Kumar Pandey^{*a}, Priya Dev^b, Shambhavi Singh^a, Akanksha Gupta^a, Abhishek Pathak^b, and Satyanshu Kumar Upadhyay^a

^a*Department of Statistics, Banaras Hindu University, Varanasi, India*

^b*Department of Neurology, Banaras Hindu University, Varanasi, India*

15 October 2025

Situations are often encountered, especially in the medical sciences, where observing each stage of an event is necessary and overlooking it might be risky for the well-being of an individual. Keeping the same very viewpoint, this article presents the analysis of real Modified Rankin score data with multiple responses from a Bayesian perspective using a polytomous logistic regression model. The study involves utilizing the Markov chain Monte Carlo technique for acquiring samples from the resulting posterior distribution. Finally, to check the scope of the model simplification, several covariates are tested against zero and then a comparison between the full model and the simplified model is proposed based on the deviance information criterion.

keywords: Odds ratio, Logistic regression, Markov Chain Monte Carlo, Modified Rankin score, Posterior distribution.

1 Introduction

The logistic regression model has a vast number of applications where the occurrence of an incident can be coded by means of a dichotomous variable. Based on a given dataset of independent variables, logistic regression determines the outcome of an event taking place, such as having a disease or not, on the basis of the odds ratio (Gupta and Upadhyay (2019)). Relying only on the dichotomous responses may, however, reduce the quality of inferences desired. This is because, in real life scenarios, we sometimes have different levels or stages of an event happening, which has to be monitored cautiously.

^{*}Corresponding authors: pranjal2802@bhu.ac.in, akankshagupta@bhu.ac.in

Usually, such situations can often be seen in the field of medical sciences and healthcare, where assessing every change in the stages of disease or recovery is very essential. The risk of ignoring the differences can be severe as it plays a major role in the subject's mortal behaviour. Therefore, multiple response datasets for such types of events can be seen where various responses can be considered to depict the different levels of a particular incident.

Generally, there are three types of multiple responses classified on the basis of three different scales, namely, the nominal scale, the ordinal scale and the interval scale. Nominal scale is the lowest scale of polytomous responses in which no specific ordering between the categories is found and can be interchanged with one another. Say, for example, grouping on the basis of the person's eye colour. The ordinal scale, on the other hand, considers specific ordering of the categories and is not exchangeable. Examples include degrees of burn, stages of heart failure, size of malignant lump, etc. Similarly, the interval scale is one where both the ordering of the categories and numerical labels attached to them are observed. A few well known responses on the interval scale include pulse rates, blood pressure levels, age distribution, etc. where the midpoint of an interval is selected as the corresponding score.

The polytomous logistic regression model, sometimes referred to as the multinomial multiple logistic regression, is a multivariate extension of the logistic model for handling data having more than two categorizations. Not only does it have the opportunity to extend the odds ratio to encompass responses beyond binary choices, but it also utilizes the sample size across all outcome categories. This approach strengthens the estimation of parameters, thereby providing more power than the traditional binary logistic regression, which solely considers the sample size of two outcome categories. The natural extension of the generalized linear model (GLM) for polytomous data is analogous to the logistic regression model with the usual 'logit' link function. The only difference that lies here is the distribution of a random component, which is now a multivariate analogue of the Bernoulli distribution, similar to multinomial density with a total count up to unity. This model has already been considered by several authors, such as Andrich (1978), Castilla (2024), Engel (1988), Rasch (1961), etc. in a classical paradigm, whereas one may refer to Draper and Smith (1998), Gelman and Hill (2006) and Fisher and McEvoy (2022) for the Bayesian developments of the model on real life problems.

The modified Rankin Scale (mRS), a widely used scale that measures the extent of disability or reliance on everyday tasks of individuals who have experienced a stroke or other neurological impairments. Dr. John Rankin originally introduced this scale in 1957, having 5 levels taking values from 1 to 5. It was later altered by Van Swieten et al. (1988) in which the lower range was adjusted from 1 to 0. Between the years 2005 and 2008, a significant modification occurred to include the numerical value of 6, indicating patients who had passed away. In contrast to the Rankin's original scale, the modernized and final variant of the scale incorporates grade 0 to signify no disability and grade 6 for indicating death. The explanations of other grades are as follows: 1 to specify no significant disability but having symptoms of it, 2 bespeak of slight, 3 to connote moderate, 4 to imply moderately severe and 5 corresponds to severe disability. Those individuals

having an mRS score of grade 0, 1, 2, and 3 do not require assistance to perform daily life routine, but from grade 4 to 6 the requirement for assistance necessitates as per increase in severity. Despite persistent critiques concerning its subjective character which can potentially bias the results, mRS remains an integral part of hospital systems for estimating rehabilitation necessities and monitoring outpatient developments. In the recent past, extensive work has been done to design multiple tools that can offer a more systematic approach towards evaluating the mRS accurately. One may see, for instance, the techniques suggested by Bruno et al. (2010), Patel et al. (2012) and Nobels-Janssen et al. (2024) in this context.

The remaining segments of the paper are outlined as below. Section 2 explains the nature of data with deliberation of the Bayesian version of the polytomous logistic regression model implemented in the study. Section 3 discusses briefly an important model comparison tool. This section is given for completeness only. Section 4 deals with numerical illustration based on a real life dataset, while the conclusion of this study has been provided in Section 5.

2 Data Structure and Bayesian Model Formulation

Let us consider that an ordinary case-control dataset is given on n subjects. Further, suppose that the Subject i is likely to have a specific value of the categorical variable M , say, $c_i, i = 1, 2, \dots, n$, where each $c_i, i = 1, 2, \dots, n$, can take one and only one value of m categories denoted by $j = 0, 1, 2, \dots, (m - 1)$. Obviously, the categorical variable M can be considered to have a total of m possible values given as $M = j; j = 0, 1, 2, \dots, (m - 1)$. $M = 0$ portrays that subject belongs to the category 0 which can be latter considered as the baseline category, $M = 1$ conveys that the subject belongs to the category 1 and so on with $M = (m - 1)$ conveying that the subject falls in the category $(m - 1)$. We next assume that each of these j categories is likely to be affected by several other explanatory variables or covariates denoted by $\underline{F} : F_1, F_2, \dots, F_k$. The explanatory variables can assume any value, such as binary, discrete or continuous. Table 1 shows a simple illustration of such a dataset with multiple categories. It may be noted that one can, of course, consider a more generalized version where each subject can have more than one category, but the assumption considered here considerably simplifies our modelling formulation.

Table 1: A simple illustration of data with polytomous responses

Subject identity	Categorical variable M	Explanatory variables			
		F_1	F_2	$\dots$	F_k
Subject 1	c_1	F_{11}	F_{12}	$\dots$	F_{1k}
Subject 2	c_2	F_{21}	F_{22}	$\dots$	F_{2k}
Subject 3	c_3	F_{31}	F_{32}	$\dots$	F_{3k}
.	.	.	.	$\dots$	.
.	.	.	.	$\dots$	.
.	.	.	.	$\dots$	.
Subject n	c_n	F_{n1}	F_{n2}	$\dots$	F_{nk}

To proceed further, the entire setup discussed above can be rewritten, assuming that each category $M = j$ is linked with a random vector $\underline{Y}$ comprising $(m - 1)$ components considering values of either zero or one based on the category of the subject. We follow Tuerlinckx and Wang (2004) and define y_j , $j = 1, 2, \dots, (m - 1)$, as the j^{th} element of $\underline{Y}$ with the expression

$$y_j = \begin{cases} 1 & \text{if } M = j; \quad j = 1, 2, \dots, (m - 1), \\ 0 & \text{otherwise.} \end{cases}$$

Thus, this conversion allows the generation of a random vector that presents an indicator-based representation of multiple categories within the variable M as shown in Table 2.

Table 2: Representation of an indicator version of polytomous responses

M	Random vector $\underline{Y}$			
	y_1	y_2	$\dots$	y_{m-1}
0	0	0	$\dots$	0
1	1	0	$\dots$	0
2	0	1	$\dots$	0
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
$m - 1$	0	0	$\dots$	1

Utilising Table 2, the probability mass function (pmf) for the vector $\underline{Y}$ can be written as

$$P(\underline{Y}|p_1, p_2, \dots, p_{m-1}) = p_1^{y_1} p_2^{y_2} \dots p_{m-1}^{y_{m-1}} (1 - p_1 - p_2 - \dots - p_{m-1})^{1-y_1-y_2-\dots-y_{m-1}}, \quad (1)$$

which is the pmf of a multivariate Bernoulli distribution with probability p_j indicating that the subject belongs to category j . Also, $M = 0$ indicates a baseline category with

the probability $P(M = 0) = 1 - p_1 - p_2 - \cdots - p_{m-1}$. Following this, let us define the link function in order to establish a connection between the response variable and the predictors associated with it. Thus, using the baseline category, the logit link function can be defined for any category $j = 1, 2, \dots, (m - 1)$ as

$$\log \frac{p_j}{p_0} = \beta_{j0} + \sum_{l=1}^k \beta_{jl} F_l,$$

where β_{j0} is the intercept and β_{jl} is the regression coefficient associated with the corresponding explanatory variable F_l ; $l = 1, 2, \dots, k$. The probability p_j can then be obtained as

$$p_j = \frac{\exp(-\beta_{j0} - \sum_{l=1}^k \beta_{jl} F_l)}{1 + \sum_{j=1}^{m-1} \exp(-\beta_{j0} - \sum_{l=1}^k \beta_{jl} F_l)}. \quad (2)$$

Also, the probability that the subject belongs to the baseline category p_0 is

$$p_0 = \frac{1}{1 + \sum_{j=1}^{m-1} \exp(-\beta_{j0} - \sum_{l=1}^k \beta_{jl} F_l)}. \quad (3)$$

Thus, using p_j 's and p_0 from (2) and (3), respectively, the likelihood function for the sample values y_{ij} obtained from $i = 1, 2, \dots, n$ subjects corresponding to $j = 1, 2, \dots, m - 1$ categories can be written as

$$L(\boldsymbol{\beta}; \underline{Y}, \underline{F}) = \prod_{i=1}^n \left\{ \left(\frac{1}{1 + \sum_{j=1}^{m-1} \exp(-\beta_{j0} - \sum_{l=1}^k \beta_{jl} F_l)} \right)^{1 - \sum_{j=1}^{m-1} y_{ij}} \times \prod_{j=1}^{m-1} \left(\frac{\exp(-\beta_{j0} - \sum_{l=1}^k \beta_{jl} F_l)}{1 + \sum_{j=1}^{m-1} \exp(-\beta_{j0} - \sum_{l=1}^k \beta_{jl} F_l)} \right)^{y_{ij}} \right\}, \quad (4)$$

where $\boldsymbol{\beta}$ is used to denote all the regression coefficients corresponding to different categories and the explanatory variables $\underline{F}$. Obviously, the number of parameters involved with the model (that is, $(m - 1)(k + 1)$) is too large to provide a routine solution, especially using the classical paradigm and, therefore, the Bayes paradigm appears to be a viable option.

2.1 Prior Specification and Posterior Distribution

To begin with the Bayesian model formulation, let us first assign the priors to each parameter of the model. Due to non-availability of appropriate information regarding the intercepts and the regression coefficients, it is appropriate to consider $N(0, \sigma^2)$ prior for each of these parameters (see also Madigan et al. (2005), Bayarri and Berger (2023)). Moreover, by considering the hyperparameter σ^2 large enough, one can proclaim that the considered priors are almost weak and the inferences are, in general, data driven.

Madigan et al. (2005) also suggested the use of independent normal priors with more or less similar formulation. Hence, the considered priors are of the forms

$$g(\beta_{jl}) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(\frac{-\beta_{jl}^2}{2\sigma^2}\right); \quad j = 1, 2, \dots, m-1, l = 0, 1, \dots, k. \quad (5)$$

Now, using the Bayes theorem, one can combine (4) and (5) to get the joint posterior distribution up to proportionality that can be specified as

$$\begin{aligned} h(\boldsymbol{\beta}|\underline{Y}, \underline{F}, \sigma^2) &\propto \prod_{i=1}^n \left\{ \left(\frac{1}{1 + \sum_{j=1}^{m-1} \exp(-\beta_{j0} - \sum_{l=1}^k \beta_{jl} F_l)} \right)^{1 - \sum_{j=1}^{m-1} y_{ij}} \right. \\ &\quad \times \left. \prod_{j=1}^{m-1} \left(\frac{\exp(-\beta_{j0} - \sum_{l=1}^k \beta_{jl} F_l)}{1 + \sum_{j=1}^{m-1} \exp(-\beta_{j0} - \sum_{l=1}^k \beta_{jl} F_l)} \right)^{y_{ij}} \right\} \\ &\quad \times \prod_{j=1}^{m-1} \prod_{l=0}^k \frac{1}{\sqrt{2\pi}\sigma} \exp\left(\frac{-\beta_{jl}^2}{2\sigma^2}\right). \end{aligned} \quad (6)$$

Solving the aforementioned posterior distribution (6) proves quite arduous analytically and, therefore, Markov chain Monte Carlo (MCMC) simulation appears to be an important alternative to drawing sample based posterior inferences. From a range of available possibilities, we can employ the Metropolis algorithm to generate samples from the posterior distribution (6). Before commenting on the implementation of the Metropolis algorithm, let us briefly review the algorithm in a general setup. In scenarios where direct sampling becomes challenging due to the presence of intractable integrals, the Metropolis algorithm offers a straightforward approach through efficient sampling from known proposal distribution, say, $q(\theta'|\theta)$ where θ is the current posterior output and θ' is the next proposal obtained from $q(\theta'|\theta)$. However, the acceptance of θ' is based on the probability given by

$$\alpha = \left(\frac{h(\theta')}{h(\theta)}, 1 \right),$$

where $h(\theta)$ is used to define the intended posterior. If the generated value θ' is accepted, the chain is allowed to move to θ' to proceed for the next iteration, otherwise the previous value θ itself is retained, and the chain proceeds for the next iteration treating θ as its previous value (see Upadhyay et al. (2001)). This process may be continued to get a single long run of the chain and the convergence monitoring may be done on the basis of ergodic averages. Once the convergence is obtained, one can pick up equally distant observations to form independent samples from the intended posterior. The gaps are chosen to make serial correlation negligibly small in order to get independent samples. It may be noted that although the paper advocates for a single long run of the chain, one can use several other alternatives to get independent samples from the intended posterior (see also Upadhyay et al. (2001), Keil et al. (2023)).

The present paper considers a multivariate normal kernel as the candidate generating density. The mean vector of the normal density was taken to be the maximum likelihood

(ML) estimates of the parameters, whereas the variance-covariance matrix was obtained through a Hessian-based approximation evaluated on the ML estimates. A scaling constant ϕ_s was also considered as a multiplier of the normal standard deviation in order to maintain a rationally high probability of acceptance (see, for example, Upadhyay et al. (2001)). The value of the scaling constant ϕ_s is generally recommended in the range 0.5 to 1.0 (see, for example, Mishra and Upadhyay (2019)).

3 Model Comparison

The polytomous responses affected by several factors represented in Table 1 have suggested analysing a polytomous logistic regression model for the dataset with $(m-1)(k+1)$ regression coefficients, although one cannot deny the possibility of having some of the regression coefficients being insignificant in the sense that they are close to zero. If one finds some of the regression coefficients close to zero, one can undoubtedly ignore these regression coefficients and come up with a simplified model. The question arises about whether one should really consider the initially proposed model or whether the simplified model can be better answered if one entertains the comparison of the two models using some well-known Bayesian tool of model comparison. Some of the most frequently used model comparison tools are the Akaike information criterion (AIC), Bayesian information criterion (BIC) and Deviance information criterion (DIC), etc. This paper considers DIC as the criterion for model comparison if the scope of simplified models appears to be a possibility (Zhang and Yang (2023)). The DIC proves to be quite beneficial where the posterior samples are obtained through MCMC simulation. The criterion is based on deviance, which can be defined as

$$D(\theta) = -2 \times \log L(\theta),$$

where θ is the unknown parameter of the model and $L(\theta)$ is the corresponding likelihood function. Following Spiegelhalter et al. (2002), one can write

$$p_D = \overline{D(\theta)} - D(\bar{\theta}),$$

where the term p_D is known as an effective number of parameters, $\overline{D(\theta)}$ is the posterior mean of deviance and $D(\bar{\theta})$ is a point estimate of deviance with $\bar{\theta}$ being the posterior mean of θ . As such, the DIC can be defined as

$$\text{DIC} = D(\bar{\theta}) + 2p_D. \quad (7)$$

The DIC so defined is calculated for each model under consideration and the model that provides the least value of DIC can finally be recommended as an appropriate model for the data in hand.

4 Numerical Illustration

To illustrate the model formulation given in Section 2, a real-life dataset of mRS scores was used from 280 subjects collected by the team of one of the co-authors at Sir Sunderlal Hospital, Banaras Hindu University. This dataset was previously analysed by Pandey

et al. (2024) from a Bayesian perspective where the authors considered the impact of covariates, presuming only the binary outcomes of the mRS score. The present illustration generalizes the assumption given in Pandey et al. (2024) by taking the outcome variable as a categorical variable with multiple categories with the intent of assessing the influence of associated explanatory variables on each category of the outcome variable individually.

Table 3: Description of the associated explanatory variables

Explanatory variable	Interpretation
Diabetes	Individual is diabetic or not
Dyslipidemia	Individual is suffering from dyslipidemia or not
Dyslipidemia-DAA	Individual is diagnosed with dyslipidemia at the time of admission
FamilyDM	Family history of diabetes
Ryletube	Ryle tube required for the individual or not
Location	Location where the incidence of stroke happened
Midline shift	Information about the midline shift occurred in the brain
Ratio	Ratio of left to right hemisphere of the brain
RBC	Red blood cell level in the blood

In the original dataset, the variable mRS score contains seven categorical responses from levels 0 to 6 where, as already mentioned, the category 0 corresponds to no disability condition and, therefore, it is considered as the baseline category. These mRS scores are further affected by nine explanatory variables. An interpretation of these variables is provided in Table 3 for clarification. It is important to mention that the complete dataset is not shown in the paper because of its confidentiality and space restrictions, although interested readers may approach one of the co-authors from the Department of Neurology for the purpose of academic use only.

Before proceeding further, let us standardize the predictors since there might be disparities between their scales. As such, the technique adopted by Gelman (2008) is used in which the binary variables are moved to have an average of zero with a deviation equal to unity, while the rest of the predictors are scaled to have a zero mean with standard deviation 0.5. Now the coefficients associated with these standardized predictors are assigned the independent prior distribution $N(0, 30)$ as mentioned previously. The joint posterior distribution for all the regression coefficients was obtained accordingly.

To obtain the posterior samples, the joint posterior distribution was subjected to the Metropolis algorithm with the settings described earlier in Subsection 2.1 and convergence based on ergodic averages was observed as the single run of the chain progressed. The value of the scaling constant ϕ_s was found to be 0.5. The plots of the ergodic averages for some of the important regression coefficients are shown in Appendix. These plots not only confirm the convergence of the running Metropolis chain but also confirm that the convergence is obtained within somewhat less than 5K iterations. A similar finding was noticed for other variates too, although not shown in the paper. After as-

sessing the convergence based on ergodic averages, 1000 samples with a constant gap of 10 units were obtained. The motive behind acquiring samples at the equidistant gap was to reduce the presence of serial correlation among the generating variates. Table 4 shows the posterior based inferences in the form of posterior means, standard deviations and the highest posterior density intervals with 0.95 coverage probability (0.95 HPDI) of the parameters drawn from these finally generated posterior samples of size 1000. In the table, the coefficients are presented as β_{jl} where $j = 1, 2, \dots, 6$ and $l = 0, 1, \dots, 9$. In the second subscript, the first value of l is reserved for the intercept, whereas the remaining values, $1, 2, \dots, 9$, are used to denote Diabetes, Dyslipidemia, Dyslipidemia-DAA, FamilyDM, Ryletube, Location, Midline shift, Ratio, RBC, respectively. Thus, β_{j0} shows the various intercepts and the rest of the coefficients explain the impact of different covariates on the mRS score for each category: $j, j = 1, 2, \dots, 6$. Say, for example, β_{11} is the regression coefficient associated with the explanatory variable diabetes for the mRS score category 1. A similar interpretation can be given to all other coefficients.

Now coming on to the results given in Table 4, it is worth noting that the positive values of the estimated posterior mean indicate that the corresponding parameters are contributing to increasing the mRS score, a fact that can be considered against the health of an individual. The negative values of the estimated posterior mean, on the other hand, can be considered to have a negative impact on the mRS score, but this fact cannot be proclaimed in totality unless specific data are made available to focus on the negative impact of the corresponding estimated regression coefficients.

Looking further at the results reported in Table 4, it is evident that in comparison to the baseline category, the occurrence of midline shift in the brain, dyslipidemia (already occurred or diagnosed at the time of admission), family history of diabetes and the requirement for ryle tube are contributing significantly in increasing the mRS score of the individuals from 0 to 1, which means that these factors are causing little deterioration in the brain of healthy individuals. Moreover, while studying the increase in mRS score from 0 to 2, the variables showing major effects were found to be diabetes, location of stroke in the brain, left to right hemisphere ratio of the brain, and family history of diabetes. The roles of other variables such as midline shift and the requirement for ryle tube are comparatively less effective but cannot be overlooked (see Table 4). Next, observing the increase in mRS score from 0 to 3, dyslipidemia diagnosed at the time of admission, ryle tube, midline shift, left to right hemisphere ratio, and family history of diabetes appear to be important factors causing the increase in the mRS score. One can also notice that some of the predictor variables are acting differently (that is, either having significant or negligible impact) up to the category score 3. The reason might be attributed to the non-availability of data corresponding to some predictors to obtain precise estimates for all the categories. Also, the degree of disability is not too severe to reach any precise conclusion. Moving towards the inferences for higher individuals' mRS scores from the baseline category, we can see that almost all the predictor variables play a significant role in its increase, indicating a rapid decrease in the individuals' brain health.

Table 4: Estimated posterior characteristics of different regression parameters.

Coefficients	Posterior mean	Posterior SD	0.95 HPDI	Coefficients	Posterior mean	Posterior SD	0.95 HPDI
β_{10}	3.24	1.18	(-0.40,6.95)	β_{40}	5.16	1.48	(0.65,9.71)
β_{11}	-0.28	0.06	(-0.50,0.14)	β_{41}	0.95	0.24	(0.12,1.85)
β_{12}	1.45	0.41	(0.20,2.71)	β_{42}	1.32	0.38	(0.15,2.51)
β_{13}	0.54	0.19	(-0.06,1.21)	β_{43}	1.46	0.31	(0.45,2.55)
β_{14}	0.64	0.18	(0.10,1.26)	β_{44}	1.63	0.38	(0.41,2.82)
β_{15}	0.87	0.23	(0.15,1.60)	β_{45}	1.23	0.31	(0.13,2.40)
β_{16}	0.06	0.006	(-0.13,0.08)	β_{46}	0.66	0.21	(-0.08,1.53)
β_{17}	1.47	0.39	(0.20,2.76)	β_{47}	0.52	0.15	(0.05,1.12)
β_{18}	-0.31	0.09	(-0.65,-0.09)	β_{48}	1.66	0.40	(0.28,2.97)
β_{19}	0.06	0.005	(-0.10,0.21)	β_{49}	0.06	0.005	(-0.06,0.18)
β_{20}	3.51	1.11	(0.12,7.12)	β_{50}	5.34	1.37	(1.12,9.72)
β_{21}	1.18	0.26	(0.30,2.05)	β_{51}	1.71	0.37	(0.51,2.98)
β_{22}	-0.03	0.006	(-0.13,0.09)	β_{52}	1.46	0.35	(0.17,2.71)
β_{23}	-0.03	0.005	(-0.12,0.10)	β_{53}	1.45	0.33	(0.28,2.56)
β_{24}	0.86	0.22	(0.15,1.55)	β_{54}	1.63	0.43	(0.25,3.10)
β_{25}	0.43	0.16	(-0.05,1.10)	β_{55}	1.56	0.41	(0.20,2.95)
β_{26}	0.78	0.21	(0.11,1.57)	β_{56}	1.93	0.48	(0.32,3.53)
β_{27}	0.66	0.21	(0.03,1.39)	β_{57}	1.47	0.37	(0.22,2.75)
β_{28}	0.81	0.23	(0.11,1.53)	β_{58}	1.66	0.32	(0.41,2.77)
β_{29}	-0.03	0.005	(-0.11,0.10)	β_{59}	0.56	0.21	(-0.11,1.19)
β_{30}	4.10	1.13	(0.54,7.83)	β_{60}	5.20	1.24	(0.97,9.15)
β_{31}	0.01	0.006	(-0.09,0.10)	β_{61}	0.70	0.21	(0.04,1.41)
β_{32}	-0.53	0.17	(-1.14,0.18)	β_{62}	1.03	0.29	(0.11,2.15)
β_{33}	1.39	0.34	(0.21,2.52)	β_{63}	0.93	0.26	(0.10,1.97)
β_{34}	0.86	0.22	(0.14,1.77)	β_{64}	0.89	0.23	(0.17,1.86)
β_{35}	1.36	0.35	(0.18,2.67)	β_{65}	0.56	0.19	(-0.19,1.31)
β_{36}	-0.06	0.005	(-0.16,0.08)	β_{66}	0.36	0.14	(-0.16,0.93)
β_{37}	0.95	0.31	(-0.21,2.43)	β_{67}	0.72	0.23	(-0.12,1.61)
β_{38}	0.95	0.28	(0.05,1.99)	β_{68}	1.23	0.31	(0.20,2.31)
β_{39}	-0.28	0.11	(-0.78,0.15)	β_{69}	0.37	0.13	(-0.13,0.85)

Initially, it was hypothesised that some of the variables such as Diabetes and FamilyDM might be closely associated, leading to an investigation of potential correlation between them. Similarly, the variables Dyslipidemia and Dyslipidemia-DAA might have some positive correlation. As such, we worked on the correlation coefficients between different covariates and, in most cases, the values were found to be quite small, leading to almost independence of these covariates, or more appropriately, very weak associations between the different covariates. Truly speaking, a real conclusion is difficult to draw due to the non-availability of appropriate data that can aptly convey a message on the independence or dependence of these covariates. Say, for instance, considering the variables Diabetes and FamilyDM, it was observed that the data size was either too small or it was independently provided for the two sets of variables.

Besides, it can also be observed that some of the covariates have a negligible effect on the mRS score of a certain category. This fact can be concluded by observing the value of the associated regression coefficients close to zero with a narrow 0.95 HPDI encompassing zero in between as well. These coefficients are β_{16} , β_{22} , β_{23} , β_{29} , β_{31} , β_{36} and β_{49} . If one assumes that these coefficients are equal to zero, the model becomes notably simpler and then drawing inferences from the simplified model is expected to be

easier than the originally considered model.

Further assessment to determine whether the simplified model is really beneficial, a comparison of the reduced model, that is, the model excluding β_{16} , β_{22} , β_{23} , β_{29} , β_{31} , β_{36} , β_{49} and the full model (which contains all the covariates) is proposed using DIC. The corresponding values of DIC for the two models are shown in Table 5. It can be seen that the reduced model provides a smaller value of DIC in comparison to the corresponding value for the full model. Thus, one can confidently recommend the reduced model for developing the desired inferences.

Table 5: DIC values for the full model and the reduced model

Model	DIC
Full model	3490.61
Reduced model	3426.14

5 Conclusion

It is common to observe multiple explanatory variables in medical studies. Some of these explanatory variables have little to no effect on the outcome variable. Using a polytomous logistic regression model, the present paper successfully delves into an examination of mRS scores ranging between 0 – 6 in neurological patients based on a real dataset. A thorough Bayes analysis is presented by employing the Metropolis algorithm, demonstrating that the procedure is both standard and adept at offering nearly all the desired inferential aspects. The findings suggest that when it comes to individuals with high mRS scores, ranging between 4 and 6, almost every explanatory variable is playing a significant role in raising the mRS score and, thereby, increasing the severity of disability. However, with low mRS scores where the degree of disability is not that severe, some of the explanatory variables such as location in the brain where the stroke occurs, pre-existing dyslipidemia, diabetes and red blood cell level, etc. can be considered to have minimal impact on the mRS score. The paper suggests that if one removes the associated regression coefficients and considers a simplified model, the simplified model stands better than the full model observed on the basis of DIC. As such, the simplified model can be considered appropriate for further inferential developments. A word of remark: our study never suggests giving up these explanatory variables, but rather proposes to give slightly less attention to these explanatory variables if the mRS scores are low, that is, less than or equal to three.

Acknowledgement

The authors express their thankfulness to the Department of Neurology, Institute of Medical Sciences, Banaras Hindu University, for providing the dataset used in the paper

that finally resulted in this collaborative study.

The authors also acknowledge their gratitude to the editor and the anonymous reviewers for their valuable and constructive suggestions that helped in improving the earlier version of the manuscript.

Disclosure Statement

The authors report that there are no potential conflicts of interest to declare.

References

- Andrich, D. (1978). Application of a psychometric rating model to ordered categories which are scored with successive integers. *Applied Psychological Measurement*, 2(4):581–594.
- Bayarri, M. J. and Berger, J. O. (2023). The interplay of Bayesian and frequentist analysis. *Statistical Science*, 38(1):1–30.
- Bruno, A., Shah, N., Lin, C., Close, B., Hess, D. C., Davis, K., Baute, V., Switzer, J. A., Waller, J. L., and Nichols, F. T. (2010). Improving Modified Rankin Scale assessment with a simplified questionnaire. *Stroke*, 41(5):1048–1050.
- Castilla, E. (2024). A new robust approach for the polytomous logistic regression model based on rényi's pseudodistances. *Biometrics*, 80(4):ujae125.
- Draper, N. R. and Smith, H. (1998). *Applied Regression Analysis*. John Wiley & Sons.
- Engel, J. (1988). Polytomous logistic regression. *Statistica Neerlandica*, 42(4):233–252.
- Fisher, T. and McEvoy, D. (2022). Bayesian multinomial logistic regression. *Journal of Statistical Modeling*, 45(3):123–145.
- Gelman, A. (2008). Scaling regression inputs by dividing by two standard deviations. *Statistics in Medicine*, 27(15):2865–2873.
- Gelman, A. and Hill, J. (2006). *Data Analysis Using Regression and Multi-level/Hierarchical Models*. Cambridge University Press.
- Gupta, A. and Upadhyay, S. (2019). On the use of a logistic regression model in the gene-environment problem: A Bayesian approach. *American Journal of Mathematical and Management Sciences*, 38(4):363–372.
- Keil, A. P., Edwards, J. K., Naimi, A. I., and Cole, S. R. (2023). The Metropolis algorithm: A useful tool for epidemiologists. *arXiv preprint arXiv:2306.16188*.
- Madigan, D., Genkin, A., Lewis, D. D., and Fradkin, D. (2005). Bayesian multinomial logistic regression for author identification. In *AIP Conference Proceedings*, volume 803, pages 509–516. American Institute of Physics.
- Mishra, R. K. and Upadhyay, S. K. (2019). Parametric Bayes analyses to study the age-specific fertility patterns. *American Journal of Mathematical and Management Sciences*, 38(2):151–173.

- Nobels-Janssen, E., Postma, E., Abma, I., et al. (2024). Validity of the Modified Rankin Scale in patients with Aneurysmal Subarachnoid Hemorrhage: A randomized study. *BMC Neurology*, 24:23.
- Pandey, P. K., Dev, P., Gupta, A., Pathak, A., Shukla, V., and Upadhyay, S. (2024). Analysis of wide Modified Rankin Score dataset using Markov Chain Monte Carlo simulation. *International Journal of Statistics in Medical Research*, 13:13–18.
- Patel, N., Rao, V. A., Heilman-Espinoza, E. R., Lai, R., Quesada, R. A., and Flint, A. C. (2012). Simple and reliable determination of the Modified Rankin Scale score in neurosurgical and neurological patients: The mRS-9Q. *Neurosurgery*, 71(5):971–975.
- Rasch, G. (1961). On general laws and the meaning of measurement in psychology. In *Proceedings of the Fourth Berkeley Symposium on Mathematical Statistics and Probability*, volume 4, pages 321–333.
- Spiegelhalter, D. J., Best, N. G., Carlin, B. P., and Van Der Linde, A. (2002). Bayesian measures of model complexity and fit. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 64(4):583–639.
- Tuerlinckx, F. and Wang, W. (2004). Models for Polytomous data. In De Boeck, P. and Wilson, M., editors, *Explanatory Item Response Models: A Generalized Linear and Nonlinear Approach*, pages 75–109. Springer-Verlag, New York, NY.
- Upadhyay, S., Vasishta, N., and Smith, A. (2001). Bayes inference in life testing and reliability via Markov Chain Monte Carlo simulation. *Sankhyā: The Indian Journal of Statistics, Series A (1961-2002)*, 63(1):15–40.
- Van Swieten, J., Koudstaal, P., Visser, M., Schouten, H., and Van Gijn, J. (1988). Interobserver agreement for the assessment of handicap in stroke patients. *Stroke*, 19(5):604–607.
- Zhang, Y. and Yang, Y. (2023). Information criteria for model selection. *WIREs Computational Statistics*.

Appendix

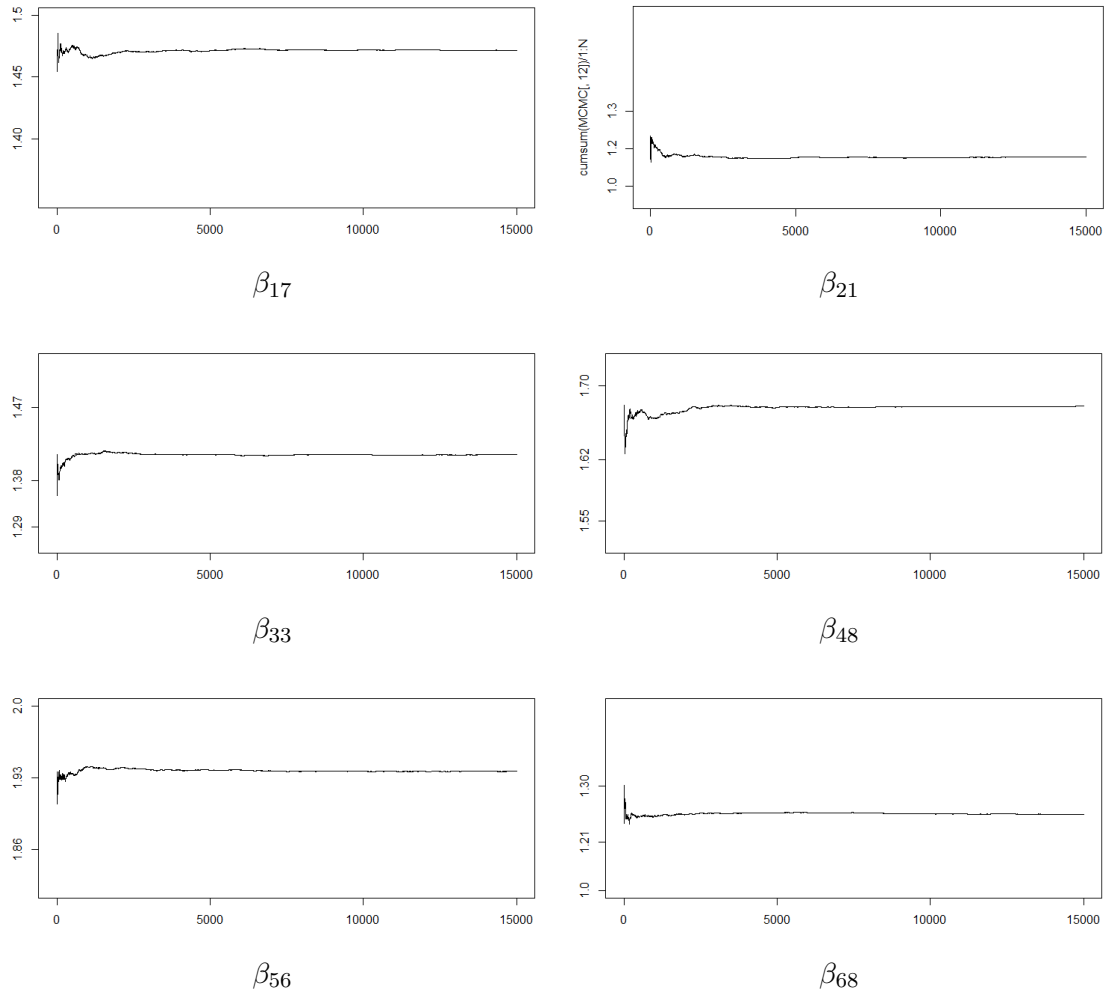


Figure 1: Plots of the ergodic averages for some of the regression coefficients.