

COUNT DATA REGRESSION MODELS: NUMBER OF
CAESAREAN SECTION DELIVERIES USING INTEGRATED
NESTED LAPLACE APPROXIMATION IN THE STATE OF
ANDHRA PRADESH, INDIA.

¹SRINU SETTI*, ²B.MUNISWAMY

Abstract: *Count Data denotes the frequency of an event transpiring within a designated time interval. For instance, examine the prevalence of caesarean sections that women experience during their lifetime. Nearly every academic discipline, such as management, economics, medicine, and industrial organizations, utilizes count data. The count data is extensively utilized across various fields including marketing, public health, and biomedical science. This study aims to estimate the relevant parameters for the Number of C-Section Deliveries (NCSD) among women in Andhra Pradesh (AP), India, who are between the ages of 15 and 49. The secondary dataset from NFHS-5 is used for the analysis. This study uses Integrated Nested Laplace Approximation (INLA) to fit the NCSD model. The PRM and NBRM are used to find the best fit. Using the information criterion, the DIC and WAIC of NBRM are 7050.75 and 7050.74, respectively, which is less than the DIC and WAIC of PRM, which are 8092.54 and 8102.35, respectively. Hence, it is concluded that NBRM is the best fit of NCSD and that Breech Presentation, Heart Disease, High Blood Pressure, Prolonged Labour, Child is Twin, Age of the respondent, and Education are significant determinants of the NCSD. Therefore, the government policy makers need to consider these variables while making the health care policies for women aged 15 to 49 years old, who are of childbearing Age.*

Keywords: Number of Caesarean Section Deliveries, PRM, NBRM, INLA, DIC, WAIC

1. Introduction

Count data regression methods are utilized when the dependent variable takes on non-negative integer values. Long (1997) and Cameron & Trivedi (1996) provide a thorough overview of count regression models. Count data models are commonly utilized in empirical research. Count was utilized in specific recent investigations. These represent the models. Yang (2007) examines the factors affecting potential admission into a sector through a Poisson distribution count model. Hellström and Nordström (2008) analyze household decisions about the number of nights allocated to monthly leisure outings through count data modelling. Nelson and Young (2008) employ Poisson and negative binomial count regressions to analyze the influence of several factors on alcohol advertising in magazines (Nan-Ting Chou & David

Steenhard, 2009). This research uses INLA to analyze fertility count data. INLA is selected due to its scalability for high-dimensional fertility data, avoiding Markov chain Monte Carlo's (MCMC) convergence issues. Several texts exclusively focused on count regression, specifically PRM and NBRM; include works by Cameron and Trivedi (2013) and Hilbe (2014). The conventional negative binomial regression model, referred to as NB2, is founded on the Poisson-gamma mixture distribution. This formulation is favoured for its capacity to model Poisson heterogeneity through a gamma distribution (NCSS, LLC, 2021). To ascertain the parameters of the variables when the dependent variable is count data and the independent variables are categorical data. The secondary data, National Family Health Survey (NFHS-5), conducted between 2019 and 2021, which is from the Demography and Health Surveys (DHS) during the phase VII, is used for the research. Estimating the mean or mode of the NCSD for the women aged between 15 and 49 in Andhra Pradesh, India is the primary goal (Hilbe, 2011). The mean or mode of the resultant posterior distribution of a parameter is referred to as the parameter of interest (Hilbe, 2011).

2. Materials & Methods

2.1. Variables considered for the study

The variables are as follows:

NCSD	=	"Number of caesarean section deliveries"
BP	=	"Breech presentation"
HD	=	"Currently has heart disease"
HBP	=	"High blood pressure"
PL	=	"Prolonged labour"
CT	=	"Child is twin"
Age	=	"Current Age"
EL	=	"Education level"
TR	=	"Type of place of residence"

Where

y = "Number of caesarean section deliveries"

x_1 = "Breech presentation"

x_2 = "Currently has heart disease"

x_3 = "High blood pressure"

x_4 = "Prolonged labour"

x_5 = "Child is twin"

x_6 = "Current Age"

x_7 = "Educational level"

Hence, we have the mathematical model as:

$$\begin{aligned}
 & \text{Number of caesarean section deliveries} \\
 &= \beta_0 + \beta_1(\text{Breech presentation}) \\
 &+ \beta_2(\text{Currently has heart disease}) \\
 &+ \beta_3(\text{High blood pressure}) + \beta_4(\text{Prolonged labour}) \\
 &+ \beta_5(\text{Child is twin}) + \beta_6(\text{Current Age}) \\
 &+ \beta_7(\text{Educational level})
 \end{aligned} \tag{1}$$

The above model can be written in Statistical terms as:

$$\begin{aligned}
 NCSD = & \beta_0 + \beta_1(BP) + \beta_2(HD) + \beta_3(HBP) + \beta_4(PL) + \beta_5(CT) \\
 & + \beta_6(Age) + \beta_7(EL) + U
 \end{aligned} \tag{2}$$

$$y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \beta_3 x_3 + \beta_4 x_4 + \beta_5 x_5 + \beta_6 x_6 + \beta_7 x_7 + U \tag{3}$$

Where y is the predictand variable, $x_1, x_2, x_3, x_4, x_5, x_6, x_7$ are predictor variables, $\beta_0, \beta_1, \beta_2, \beta_3, \beta_4, \beta_5, \beta_6, \beta_7$ are parameters and U is disturbance term

Then the structure of a count model will be

$$\log(\mu) = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \beta_3 x_3 + \beta_4 x_4 + \beta_5 x_5 + \beta_6 x_6 + \beta_7 x_7 \tag{4}$$

In GLM theory, the link function linearizes the relationship between the linear predictor, $x'\beta$, and the fitted value, μ or $\hat{y}$ or estimated mean, $E(y)$. Consequently, μ is delineated in relation to the inverse relationship.

$$\mu = e^{\beta_0 + \beta_1 x_1 + \beta_2 x_2 + \beta_3 x_3 + \beta_4 x_4 + \beta_5 x_5 + \beta_6 x_6 + \beta_7 x_7} \tag{5}$$

$$\log(\mu_i) = \sum_{i=1}^n \beta_0 + \beta_1 X_{1i} + \beta_2 X_{2i} + \dots + \beta_j X_{ji} \quad (6)$$

2.2. Multicollinearity detection by Variance Inflation Factor (VIF)

Montgomery, D.C. et al. (2003), state that the presence of more than two regressor variables results in the effects of multicollinearity. The estimates of $\hat{\beta}_j$ will be large if multicollinearity exists between regressors x_i and x_j . Based on VIF values, one can say that there is no multicollinearity if the VIF value is 1. Moderate multicollinearity exists if the VIF value lies in the interval 1-5. High multicollinearity exists if the VIF value is greater than 5. Serious multicollinearity exists if the VIF value is greater than 10.

2.3. Heterogeneity Test

This involves comparing a model with a random effect for group membership to a model without it, using model comparison metrics like the WAIC or the Deviance DIC provided by INLA. It includes creating an INLA model with the outcome variable regressed on covariates but without a random effect for the group variable, allowing it to be a null model. Then, create a new INLA model identical to the null model, but include a random effect for the group variable. Extract the DIC or WAIC values from both models. A larger difference in DIC/WAIC between the null and alternative models suggests significant evidence of heterogeneity. If the null model DIC/WAIC value is less than the alternative model DIC/WAIC value, then it is evident that heterogeneity exists.

2.4. Poisson Regression Model (PRM)

The probability density function of Poisson random variables is defined by

$$f(y_i, \mu_i) = \frac{e^{-\mu_i} \mu_i^{y_i}}{y_i!}, y_i = 0, 1, 2, \dots; \mu_i > 0. \quad (7)$$

And its log likelihood function is

$$\ell(\mu, y) = \sum_{i=1}^n [y_i \log \mu_i - \mu_i - \log y_i!] \quad (8)$$

A key characteristic of the Poisson distribution and PRM is equi-dispersion, indicating the same mean and variance of the distribution.

$$E(y_i) = Var(y_i) = \mu_i \quad (9)$$

2.5. Over-dispersion Test

B. Muniswamy & Aragaw Eshetie Aguade (2019) explain that the efficiency of parameter estimates is still rather high when the typical Poisson model is used on over-dispersion data, but their standard errors are understated. As a result, test significance levels and coverage probabilities of confidence intervals are no longer reliable and may produce wildly deceptive results (Heinzl & Mittlböck, 2003). To put it another way, excessive dispersion will result in an underestimation of standard errors, which will lead to incorrect analysis conclusions (Ismail & Jemain, 2007). Over-dispersion must therefore be managed. Using an NBRM for over-dispersion is an additional strategy (Zhao et al., 2009; Zhu & Zhang, 2006).

Hilbe, J.M. (2011) writes that the classic basic count response model is the PRM. The PRM's central distribution assumption is equi-dispersion of PRM. Real data rarely matches this assumption. When the response variance exceeds the mean, over-dispersion in PRM takes place. When the variation is smaller than the mean, under-dispersion is present. In actuality, these hardly ever happen.

The primary causes of over-dispersion are a positive association between response and excessive variability in response counts or probabilities. Over-dispersion arises when the distributional assumptions of the data are violated, particularly when observations are clustered, contravening the assumption of probability independence. Excessive variability may cause the standard errors of the estimations to be underestimated or diminished. A variable may appear to be a significant predictor when it is not.

A model is considered over-dispersed if the ratio of the Pearson or Chi-square statistic to the degrees of freedom exceeds 1.0. The fraction of either is referred to as the dispersion. If greater than 1.25, then an adjustment may be necessary. If greater than 1.5, then it is classified as over-dispersed.

2.5.1. Pearson's Chi-square Test

Pearson's Chi-square test is calculated by dividing the Pearson (or χ^2) statistic by the degrees of freedom. The fraction of either is referred to as the dispersion. Equation (10) serves as a criterion for assessing over-dispersion in a PRM. If the quotient exceeds 1, it indicates over-dispersion.

$$Dispersion = \frac{Chi - square}{degree\ of\ freedom} \quad (10)$$

Where Chi-square is

$$\chi^2 = \sum_{i=1}^n \frac{(y_i - \hat{\mu}_i)^2}{\hat{\mu}_i} \quad (11)$$

Then create a diagnostic plot in R. The plot is designed to visualize the relationship between the fitted values and the squared Pearson residuals.

2.6. Negative Binomial Regression Model (NBRM)

If the principal properties of the Poisson distribution and the equi-dispersion of the PRM are violated, then the mean and variance of the distribution will not be equivalent. Over-dispersion occurs when the variance exceeds the mean of the response variable. Under-dispersion occurs when the variance is less than the mean of the response variable. The Binomial regression model is more suitable for under-dispersed data. The NBRM is more suitable for over-dispersed data.

Then the NBRM is

$$p\left(y_i; \frac{1}{\alpha}, \mu_i\right) = \frac{\Gamma\left(y_i + \frac{1}{\alpha}\right)}{\Gamma\left(\frac{1}{\alpha}\right) \Gamma(y_i + 1)} \left(\frac{1}{1 + \alpha \mu_i}\right)^{\frac{1}{\alpha}} \left(\frac{\alpha \mu_i}{1 + \alpha \mu_i}\right)^{y_i} \quad (12)$$

Where α is the parameter that indicates the degree of over-dispersion, and $y_i = 0, 1, 2, 3, \dots$

The log likelihood function of NBRM is

$$\begin{aligned} \ell\left(y_i; \frac{1}{\alpha}, \mu_i\right) = \sum_{i=1}^n & \left[y_i \log\left(\frac{\alpha \mu_i}{1 + \alpha \mu_i}\right) - \left(\frac{1}{\alpha}\right) \log(1 + \alpha \mu_i) \right. \\ & \left. + \log \Gamma\left(y_i + \left(\frac{1}{\alpha}\right)\right) - \log \Gamma(y_i + 1) - \log \Gamma\left(\frac{1}{\alpha}\right) \right] \end{aligned} \quad (13)$$

Hence, the mean and variance of NBRM are

$$E(y_i) = \mu_i \quad (14)$$

and

$$Var(y_i) = \mu_i(1 + \alpha \mu_i) \quad (15)$$

2.7. Model selection

When comparing multiple models to choose the most suitable one for the data, DIC and WAIC can be employed as alternatives to the likelihood ratio test. Similar to the Akaike Information Criterion (AIC), the DIC is the predominant method for determining the model that best fits the data by comparing two or more models using INLA. It seeks to equilibrate the goodness of fit with the model's complexity. It is analogous to the coefficient of multiple determinations (R^2); however, it imposes a penalty based on the number of parameters included in the model (i.e., the model's complexity). In contrast to R^2 , an effective model is characterised by the lowest DIC value (Dejen Tesfaw Molla, 2013).

2.7.1. Marginal Log-likelihood (MLIK)

The marginal likelihood $p(y)$ is a helpful metric when comparing models. Bayes factors, for instance, are defined as ratios of the marginal likelihoods of two competing models; this likelihood is also used in the DIC computation.

$$\tilde{p}(y) = \int \frac{p(\eta, \theta | y)}{\tilde{p}(\eta | \theta, y)} |_{\eta=\eta^*(\theta)} d\theta \quad (16)$$

In actuality, it is $\tilde{p}(\theta | y)$'s normalizing constant. This approximation permits the deviation from Gaussian since $\tilde{p}(\theta | y)$ is handled non-parametrically. However, if the posterior marginal of θ is multimodal, this approach might not work. This applies broadly to the INLA technique and is not unique to the evaluation of the marginal likelihood. Thankfully, unimodal posterior distributions are usually produced by the latent Gaussian models (H. Rue, Martino, and Chopin 2009).

2.7.2. Deviance Information Criterion (DIC)

The most commonly utilized information criterion in the frequentist statistical framework is AIC. The DIC proposed by Spiegelhalter et al. (2002) represents a major advancement in model selection within the Bayesian literature during the last two decades. DIC can be seen as a Bayesian alternative to AIC. DIC can be seen as a Bayesian alternative to AIC. Analogous to AIC, it evaluates the predictive accuracy of hypothetically duplicated data against observed data, balancing a metric of model adequacy with a metric of complexity. DIC, however, takes into account earlier information, unlike AIC. DIC possesses several beneficial attributes and is constructed with the posterior mean of the log-likelihood or the deviation (Yong Li et al., 2022).

$$DIC = -2(d - P_{dic}) \quad (17)$$

Where

$$d = \frac{1}{S} \sum_{s=1}^S \log[P(y|\theta_s)] \quad (18)$$

$$P_{dic} = \max [\log[P(y|\theta)]] - d \quad (19)$$

Where $\max[\]$ is the maximum function over all posterior samples, and S is the number of posterior samples (Nathan J. Evans, 2019).

2.7.3. Watanabe – Akaike Information Criterion (WAIC)

As stated by Watanabe (2010), Andrew Gelman et al. (2014), and Gaya & Ketz (2024), the predicted log predictive densities for each data point are used to compute WAIC. WAIC is adequate to determine the best model among all candidate models if the posterior distribution of θ is derived directly from the probability of y_i without the use of latent processes and if the observation process is free of sampling bias (Gaya & Ketz, 2024).

WAIC is defined as:

$$WAIC = -2(lpd - P_{waic}) \quad (20)$$

Where

$$lpd = \sum_{i=1}^n \log \left[\frac{1}{S} \sum_{s=1}^S p(y|\theta_s) \right] \quad (21)$$

$$P_{waic} = \sum_{i=1}^n var [\log [p(y|\theta)]] \quad (22)$$

Where S is the number of posterior samples, n is the number of data points, $var[\]$ is the variance function over the posterior samples, lpd is the log predictive density of the data, and P_{waic} is the effective number of parameters (Gaya & Ketz, 2024).

3. Test Results & Applications

3.1. Multicollinearity

Certain variables can be multicollinear when constructing a linear regression model. When multiple independent variables correlate with one another, this is

known statistically as multicollinearity. As a result of this multicollinearity, statistical judgments become less reliable. In a regression model analysis, multicollinearity is the absence of unique information about the regression model due to a high correlation between two or more independent predictor variables. Therefore, these variables need to be eliminated when creating a multiple regression model.

The variance inflation factor (VIF) measures the extent of correlation among independent variables in a regression model, serving to detect multicollinearity within the model. A VIF value below 1 indicates the absence of a connection. A moderate correlation is shown when the VIF value ranges from 1 to 5. A VIF value exceeding 5 indicates a strong association.

Table 1: Summary of VIF values

Covariates	GVIF	DF	$GVIF^{1/(2*DF)}$
BP	6.1722	2	1.5762
HD	1.0026	1	1.0013
HBP	12.9088	2	1.8955
PL	20.2855	2	2.1223
CT	1.0358	3	1.0059
Age	1.1436	1	1.0697
EL	1.1440	3	1.0227

Table 1 gives the summary of VIF values. The $GVIF^{\frac{1}{2*DF}}$ value of BP is 1.5762, HD is 1.0013, HBP is 1.8955, PL is 2.1223, CT is 1.0059, Age is 1.0697, and EL is 1.0227. Since all the $GVIF^{\frac{1}{2*DF}}$ values are between 1 and 5, it is concluded that moderate multicollinearity is present. Where (GVIF) is the Generalized Variance Inflation Factor and (DF) is the Degrees of Freedom. Hence, it is concluded that all the covariates are to be included in the model, for they are all important variables in the study.

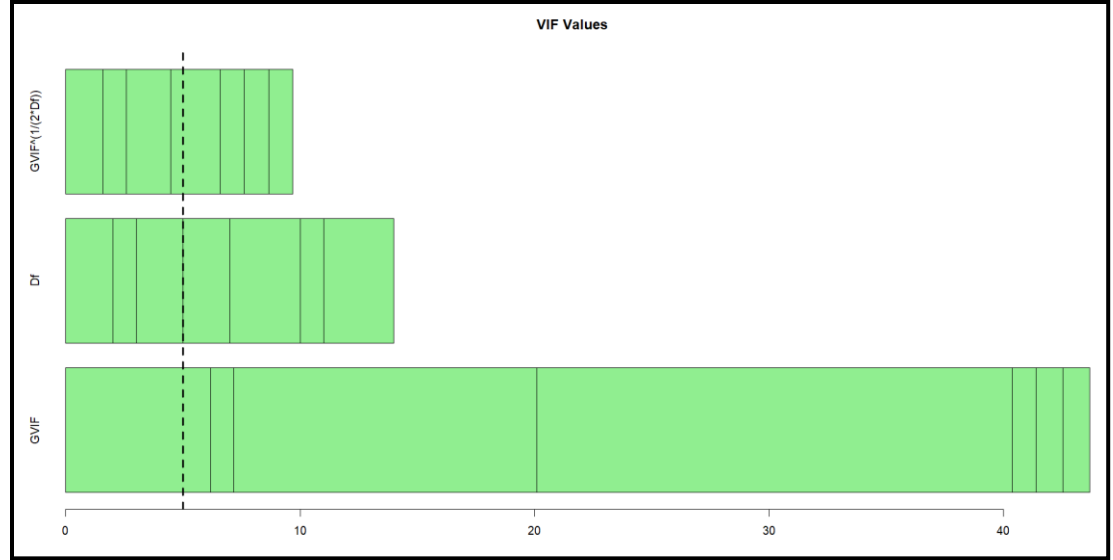


Figure 1: VIF plot

Figure 1 showcases the VIF values. The plot explains that the VIF value is between 1 and 5, which is a sign of moderately multicollinearity present. That says all the regressors BP, HD, HBP, PL, CT, Age, and EL are to be included for the study.

3.2. Heterogeneity

The DIC values of PRM without random effect, type of residence (TR), and PRM with random effect, TR, are computed and compared for the heterogeneity.

Table 2: Heterogeneity Test

PRM	DIC
without random effect, TR	8092.69
with random effect, TR	8194.00

Table 2 gives the DIC values of PRM without and with random effect; TR. The PRM with random effect TR, DIC value 8194.00, is greater than the PRM without random effect TR, DIC value 8092.69, which indicates that there is no significant evidence of heterogeneity.

3.3. Application - PRM using INLA

Modelling the NCSD using INLA, fitted in PRM. The mean, standard deviation, 0.025 quantile, 0.5 quantile, 0.975 quantile, mode, and Kullback-Leibler Divergence (KLD) are estimated.

Table 3: Values of the parameters of interest, and KLD - PRM

Fixed effect:	Mean	Standard deviation	0.025 quantile	0.5 quantile	0.975 quantile	Mode	KLD
(Intercept)	-1.112	0.161	-1.428	-1.112	-0.796	-1.112	0.00
BPYes	1.082	0.070	0.945	1.082	1.219	1.082	0.00
BPDon't know	0.656	0.115	0.431	0.656	0.881	0.656	0.00
HDYes	-0.920	0.501	-1.901	-0.920	0.062	-0.920	0.00
HDDon't know	0.000	31.623	-61.980	0.000	61.980	0.000	0.00
HBPYes	0.640	0.061	0.521	0.640	0.759	0.640	0.00
HBPDon't know	0.036	0.146	-0.249	0.036	0.321	0.036	0.00
PLYes	-0.487	0.063	-0.611	-0.487	-0.364	-0.487	0.00
PLDon't know	-0.348	0.187	-0.714	-0.348	0.018	-0.348	0.00
CT1st of multiple	0.204	0.213	-0.214	0.204	0.621	0.204	0.00
CT2nd of multiple	0.095	0.210	-0.316	0.095	0.506	0.095	0.00
CT3rd of multiple	-105.924	13.394	-132.175	-105.924	-79.673	-105.924	0.00
CT4th of multiple	0.000	31.623	-61.980	0.000	61.980	0.000	0.00
CT5th of multiple	0.000	31.623	-61.980	0.000	61.980	0.000	0.00
Age	0.026	0.005	0.017	0.026	0.035	0.026	0.00
ELPrimary	0.391	0.101	0.193	0.391	0.588	0.391	0.00
ELSecondary	0.031	0.076	-0.119	0.031	0.181	0.031	0.00
ELHigher	0.449	0.081	0.290	0.449	0.609	0.449	0.00

Table 3 displays estimations of PRM. The mean or mode of the covariates BP Yes, BP Don't know and HBP Yes are more than 0.500 where as the covariates HD Don't know, CT4th of multiple and CT5th of multiple are 0.000. A posterior distribution is correctly approximated by a Gaussian distribution when the KLD score is zero.

3.4. Pearson's Chi-square Test

Pearson's Chi-square test performs the test for dispersion. The null hypothesis is that there is no over-dispersion in PRM. The result is as follows:

Table 4: Chi-square Test

Data	χ^2	df	Dispersion
PRM	3169.961	2818	1.1249

Table 4 shows the Chi-square test for dispersion. The dispersion value, 1.1249, is greater than 1, which indicates that the PRM may not be appropriate due to over-dispersion. Reject the null hypothesis and accept the alternative hypothesis. It is concluded that the PRM is over-dispersed.

Over-dispersion in a PRM is verified by using Pearson's chi-square statistic and creating a diagnostic plot.

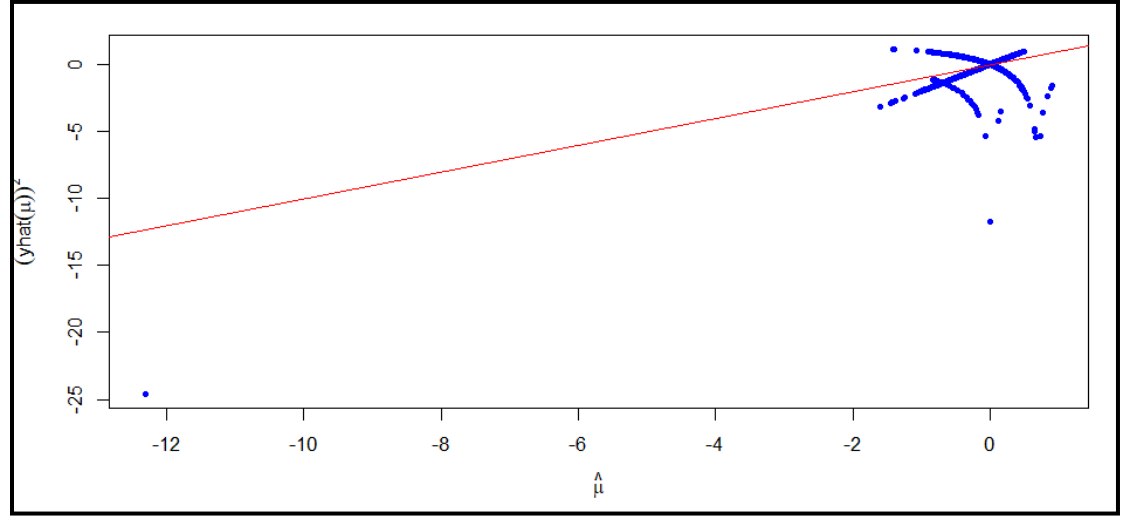


Figure 2: Diagnostic plot for over-dispersion

Figure 2 is a diagnostic plot for over-dispersion in Poisson model data. Some points above the reference red line suggest over-dispersion, which confirms NBRM adoption.

3.5. Application - NBRM using INLA

Modelling the NCSD using INLA, fitted in NBRM. The mean, standard deviation, 0.025 quantile, 0.5 quantile, 0.975 quantile, mode, and KLD are estimated.

Table 6: Values of the parameters of interest, KLD and hyperparameters - NBRM

		Standard	0.025	0.5	0.975		
Fixed effect:	Mean	deviation	quantile	quantile	quantile	Mode	KLD
(Intercept)	-1.081	0.173	-1.420	-1.081	-0.742	-1.081	0.00
BPYes	-0.316	0.075	-0.463	-0.316	-0.170	-0.316	0.00
BPDon't know	0.513	0.129	0.260	0.513	0.765	0.513	0.00
HDYes	-0.977	0.520	-1.997	-0.977	0.043	-0.977	0.00
HDDon't know	0.000	31.623	-62.009	0.000	62.009	0.000	0.00
HBPYes	0.032	0.066	-0.097	0.032	0.161	0.032	0.00
HBPDon't know	0.017	0.158	-0.293	0.017	0.327	0.017	0.00

PLYes	-0.067	0.068	-0.201	-0.067	0.067	-0.067	0.00
PLDon't know	-0.750	0.205	-1.151	-0.750	-0.347	-0.750	0.00
CT1st of multiple	0.201	0.231	-0.251	0.201	0.654	0.201	0.00
CT2nd of multiple	0.082	0.227	-0.364	0.082	0.529	0.082	0.00
CT3rd of multiple	-31.384	13.481	-57.833	-31.379	-4.962	-31.379	0.00
CT4th of multiple	0.000	31.623	-62.009	0.000	62.009	0.000	0.00
CT5th of multiple	0.000	31.623	-62.009	0.000	62.009	0.000	0.00
Age	0.022	0.005	0.012	0.022	0.032	0.022	0.00
ELPrimary	0.107	0.106	-0.101	0.107	0.315	0.107	0.00
ELSecondary	0.491	0.081	0.333	0.491	0.649	0.491	0.00
ELHigher	0.701	0.087	0.531	0.701	0.871	0.701	0.00
Model hyperparameters:	6.89	3.66	3.34	5.84	16.17	4.81	

Table 6 presents the estimation of the NBRM parameters. The mean or mode of the covariates BP, Don't know, and EL Higher is more than 0.500, whereas the covariates HD, Don't know, CT4th of multiple, and CT5th of multiple are 0.000.

Table 7: Comparison of models

Model	Model selection criteria		
	MLIK	DIC	WAIC
PRM	-3597.74	8092.54	8102.35
NBRM	-3593.64	7050.75	7050.74

Table 7 illustrates that the marginal log-likelihood of NBRM, -3593.64, exceeds that of PRM, -3597.74. The DIC for NBRM is 7050.75, which is lower than the PRM value of 8092.54. Additionally, the WAIC for NBRM is 7050.74, which is also lower than the PRM value of 8102.35. Therefore, NBRM aligns more effectively with the NCSD model. The NBRM exhibits lower DIC and WAIC values than PRM. Consequently, this substantiates that NBRM is suitable and superior to PRM.

4. Discussion

4.1. Findings

This study succinctly outlines the INLA algorithm for estimating the marginal posterior mean or mode of parameters and hyperparameters in Bayesian spatial and spatio-temporal models. The parameters of interest, mean or mode of

PRM, are as follows. Mean or mode of the intercept is -1.112, BP Yes & BP Don't know is 1.082 & 0.656, HD Yes & HD Don't know is -0.92 & 0.000, HBP Yes & HBP Don't know is 0.640 & 0.036, PL Yes & PL Don't know is -0.487 & -0.348, CT1st multiple, CT2nd multiple, CT3rd multiple, CT4th multiple & CT5th multiple is 0.204, 0.095, -105.924, 0.000 & 0.000, Age is 0.026, EL Primary, EL Secondary & EL Higher" is 0.391, 0.031, 0.449.

The parameters of interest, mean or mode of NBRM, are as follows. Mean or mode of the intercept is -1.081, BP Yes & BP Don't know is -0.316 & 0.513, HD Yes & HD Don't know is -0.977 & 0.000, HBP Yes & HBP Don't know is 0.032 & 0.017, PL Yes & PL Don't know is -0.067 & -0.750, CT1st multiple, CT2nd multiple, CT3rd multiple, CT4th multiple & CT5th multiple is 0.201, 0.082, -31.384, 0.000 & 0.000, Age is 0.022, EL Primary, EL Secondary & EL Higher is 0.107, 0.491 & 0.701 respectively. The NCSD dataset relevant to PRM and NBRM is employed to demonstrate the estimated solution using INLA. The algorithm INLA provides substantial computing advantages compared to various techniques for tackling issues related to random and fixed effects within designated regions and timeframes in spatial-temporal analysis. Using the INLA algorithm this work calculates the fixed effects of additive models. NBRM fits the best for the NCSD, which is over-dispersed. The alternative models can be used to estimate NCSD. The prevalence of c-section deliveries in AP from 2019 to 2021 is examined by Bayesian spatial-temporal modelling using the INLA framework. The NCSD model can be compared with other regression models.

4.2. Conclusion

The study is conducted using the INLA package in R. The NBRM; DIC, 7050.75 and WAIC, 7050.74 demonstrate a superior match in modeling the NCSD compared to the PRM; DIC 8092.54 and WAIC 8102.35, as indicated by DIC and WAIC values. The INLA offers an effective approach for modeling in PRM and NBRM. The PRM can be compared with the count data regression models that evaluate over-dispersion, which is recommended for further research. This study aimed to utilize regression models dealing count data to examine the NCSD, employing empirical data from NFHS-5. The factors of interest are assessed, focusing on the NCSD childbearing women in Andhra Pradesh, India. NBRM is recognized as the most appropriate model, indicating that BP, HD, HBP, PL, CT, Age, and EL are significant factors influencing the NCSD. Therefore, government policymakers need to consider these variables while making health care policies for women aged 15 to 49 years old who are childbearing.

Acknowledgement

The first author is grateful to the Ministry of Tribal Affairs, Government of India, for the financial support through NFST Fellowship (Award No-202021-NFST-AND-02662).

Declaration of conflicting interests

The authors declared no conflicts of interest

References

1. Aragaw Eshetie Aguade and B. Muniswamy.: "Proposed Score Test for Over-dispersion Parameter in the Multilevel Negative Binomial Regression Model", Journal of Emerging Technologies and Innovative Research, Volume 5, Issue 12, (2018) 709-720.
2. B. Muniswamy and Aragaw Eshetie Aguade.: "Proposed Score Test of Homogeneity of Groups in the Multilevel Poisson Regression Model for Clustered Discrete Data", International Journal of Scientific Research and Reviews, 8(1), (2019) 1036-1052.
3. Cameron, A. C. and Trivedi, P. K.: Count Data Models for Financial Data, Handbook of Statistics, Vol. 14, Statistical Methods in Finance, Amsterdam, North-Holland, (1996) 363-392.
4. Cameron, A. C. and Trivedi, P. K.: Regression Analysis of Count Data. Cambridge University Press, 2013.
5. Evans, N. J.: Assessing the practical differences between model selection methods in inferences about choice response time tasks, Psychonomic Bulletin and Review, 26 (2019) 1070-1098
6. Gaya, H. E., & Ketz, A. C.: Comparison of WAIC and posterior predictive approaches for N-mixture models. Scientific Reports, 14(1) (2024).
7. Gelman, Andrew, Jessica Hwang, and Aki Vehtari.: "Understanding Predictive Information Criteria for Bayesian Models." Statistics and Computing 24 (6): (2014) 997-1016.
8. Heinzl, H., Mittlböck, M.: R-squared measures for the inverse Gaussian regression model. Comput. Statist., 17, (2002) 525-544.
9. Hellström, Jörgen and Jonas Nordström.: A Count Data Model with Endogenous Household Specific Censoring: the Number of Nights to Stay, Empirical Economics, Vol. 35, No. 1, (2008) 179-193.
10. Hilbe, J. M.: Negative Binomial Regression. Cambridge University Press, 2011.

11. Hilbe, J., Arizona State University, & Jet Propulsion Laboratory, California Institute of Technology. Modeling count data, Cambridge University Press, 2014.
12. Ismail, N., and Jemain, A.: Handling over dispersion with negative binomial and generalized Poisson regression models, Causality Society forum, (2007) 103–158.
13. Li, Y., Wang, N., Yu, J., & Zeng, T.: Deviance Information Criterion for model selection: Theoretical justification and Applications (2022b) 1–108.
14. Long, J. Scott. Regression Models for Categorical and Limited Dependent Variables: Advanced Quantitative Techniques in the Social Sciences Series 7, SAGE Publications, Thousand Oaks, CA. 1997.
15. Molla, D. T.: The Power of Statistical Tests for Over-dispersed and Zero-Inflated Count Data in Some Discrete Regression Models. Andhra University, Visakhapatnam [Ph.D. Thesis]. (2013).
16. Montgomery. D.C., Peck. E.A., and Geofferey Vining, G.: Introduction to Linear Regression Analysis, John Wiley and sons, 2003.
17. Nan-Ting Chou and David Steenhard.: A Flexible Count Data Regression Model Using SAS PROC NLMIXED, Statistics and Data Analysis, SAS Global Forum 2009 paper 250–2009, (2009) 1-14.
18. NCSS, LLC.: Negative Binomial Regression, In NCSS Statistical Software, Chapter 326, (2021) 1–36.
19. Nelson, Jon P. and Douglas Young.: Effects of youth, price, and audience size on alcohol advertising in magazines, Health Economics, 17(4), (2008) 551–556.
20. Rue, H., S. Martino, and N. Chopin.: “Approximate Bayesian Inference for Latent Gaussian Models Using Integrated Nested Laplace Approximations (with Discussion).” Journal of the Royal Statistical Society: Series B (Methodological) 71 (2): (2009) 319–92.
21. Spiegelhalter, D. J., Best, N. G., Carlin, B. P., and Van Der Linde, A.: Bayesian measures of model complexity and fit. Journal of the Royal Statistical Society: Series B (Statistical Methodology), 64(4): (2002) 583-639.
22. Watanabe, S.: Asymptotic equivalence of Bayes cross-validation and widely applicable information criterion in singular learning theory. J. Mach. Learn. Res., 11: (2010) 3571–3594.
23. Yang, Chih-Hai.: What factors inspire the high entry flow in Taiwan's manufacturing industries - A count entry model approach, Applied Economics. Vol. 39, (2007) 1817-1831.

24. Zhao, Y., James, H., Cheryl, L.: Score tests for over dispersion in zero-inflated Poisson mixed models, Epidemiology and Biostatistics, Arnold School of Public Health, University of South Carolina, SC 29208, USA, (2009).
25. Zhu, H., and Zhang, H.: Generalized score test of homogeneity for mixed effects models, The Annals of Statistics, 34, (2006) 1545 – 1569.

SRINU SETTI: RESEARCH SCHOLAR, DEPARTMENT OF STATISTICS,
ANDHRA UNIVERSITY, VISAKHAPATNAM, ANDHRA PRADESH-530003,
INDIA

B.MUNISWAMY: PROFESSOR, DEPARTMENT OF STATISTICS, ANDHRA
UNIVERSITY, VISAKHAPATNAM, ANDHRA PRADESH-530003, INDIA

**Corresponding author Email : srinujesusj@gmail.com*