

How to reduce bias in the estimates of count data regression?

Ashwini Joshi, Cytel Inc., Pune, India

Sumit Singh, Cytel Inc., Pune, India

ABSTRACT

Maximum likelihood estimation is a method commonly used for parameter estimation. The main disadvantage of this method is that it can be heavily biased for small samples. Firth (Biometrika, 1993) suggested method for reduction in bias through a penalization of the likelihood. This bias reduction method is used frequently. LogXact®, SAS® and STATA® provided this method for binary regression.

In clinical trials, regression of count data such as lesions in multiple sclerosis is performed quite often. The Maximum Likelihood Estimate for count data also suffers from small sample bias and an infinite estimate in the case of sparse data. In this paper we propose the use of bias reduced MLE for count data regression. We study important aspects of these estimates including shorter confidence intervals in the case of sparsity. We study these using examples of Poisson and Negative Binomial regression.

We will demonstrate this method using a SAS Proc.

INTRODUCTION

Non-linear regression has always been computationally challenging area of study. When the relation in the response and predictors is not linear, estimates are generally found by maximizing the log likelihood. In Non linear regression one has to deal with two important issues – maximization of the log likelihood and bias in the estimation. Maximization problem is solved by numerical iterative methods like Newton-Raphson or Quasi Newton Method. In the case of small sample study, convergence issues arise while maximizing the log likelihood and small sample bias gets introduced.

Most of the bias reduction methods in practice are corrective in nature. That means first the biased estimate is computed and then modifications are done to remove the highest order term in bias.

Firth [1993] suggested a new approach of modifying the score equation or likelihood function to obtain less biased estimates. Firth used Jeffrey's prior to penalize the likelihood function of Generalized linear models. Heinze et al. (2002, 2010) have applied this method to unconditional and conditional binary logistic regression. Heinze et al. have also shown other beneficial properties of the bias- reduction method suggested by Firth for binary logistic regression. I. Kosmidis (2007) has taken an overview of bias-reduction methods for all the exponential family models.

This method is used quite often. Firth's paper is cited six times as high in 2013 as in 2008 (Geroldinger, 2014). Many statistical software implemented this idea for logistic regression.

BIAS REDUCTION METHOD

In the regression analysis, regression model is fitted to the underlying data. Using the regression model one can compute the point and interval estimate of the parameter. The difference in the expected value of the estimator and true value of the parameter is called as bias.

Suppose Y is the response or dependent term in the regression analysis and $X = (X_1, X_2, \dots, X_p)$ are p model parameters or predictors. And $\hat{\beta} = (\hat{\beta}_1, \hat{\beta}_2, \hat{\beta}_3, \dots, \hat{\beta}_p)^T$ are the estimated values of the coefficients.

The bias in the estimated value of the parameter $\hat{\beta}$ can be expanded as

$$\text{Bias} = E(\hat{\beta}) - \beta = \frac{b_1(\beta)}{n} + \frac{b_2(\beta)}{n^2} + \dots$$

where n is interpreted as the measurement of unit of information such as the number of observations or Fisher's Information. β is the true but unknown value of the parameter and each $b_i(\beta)$ is first order function of β .

As discussed earlier, existing bias reduction methods are corrective in nature. Cox & Snell (1968, §3) derived the expression for First order bias term for a very general family of models. Based upon the work by Cox & Snell (1968), Anderson & Richardson (1979) and Schaefer (1983) calculated the first-order bias terms for logistic regressions. Cordeiro & McCullagh (1991) further extended these results for the class of generalized linear models. Jackknife method is another popular method for bias reduction.

But both these methods may not work for small sample or sparse data. These methods can be applied only when Maximum likelihood estimate exists for the complete data or for sub-sample in the case of Jackknife. But in the case of small sample or sparse data, Maximum likelihood estimate may not exist; in such situations these existing methods can not be used for bias reduction.

Firth [1993] suggested a completely different approach of bias reduction for exponential family models. Firth pointed out that bias in the maximum likelihood estimate arises due to combination of

- (i) Unbiasedness of the score function, $E(U(\beta)) = 0$ at the true value of beta, and
- (ii) Curvature of the score function indicated by $\frac{\partial^2 U(\beta)}{\partial \beta^2} \neq 0$

Firth illustrated that introducing a small bias or penalty in the score function or likelihood function will help in reducing the bias in the estimate. He further showed that for generalized linear model with canonical link, this penalty for likelihood function would be Jeffrey's Prior.

This method is already studied and implemented in many Software (viz. LogXact®, R package – BRGLM, SAS® and STATA®) for binary regression. As the method is applicable to all exponential family models (Firth (1993), I. Kosmidis (2007)), we have implemented it for count response models.

In this paper we will demonstrate the use of this method for Poisson and Negative binomial regression models. With examples in clinical trial domain we will assess other properties of the bias reduction method like shorter confidence interval, existence in the case of quasi separation in the data or sparse data.

PROFILE LIKELIHOOD BASED CONFIDENCE INTERVAL

In Maximum likelihood estimation, it is assumed that MLE follows Normal distribution. Since standard errors of GLM are based on asymptotic variance, they may not be a good estimator of standard deviation for small sample data and hence Wald CIs may not perform well.

When this is the case, inference resulting from profile based likelihood confidence intervals is much more reliable.

This more robust construction of confidence intervals is derived from the asymptotic χ^2 distribution of the generalized likelihood ratio test. In case of the highly biased estimate use of profile likelihood based confidence interval is highly recommended.

If $l_{\max}$ is the maximum likelihood then the endpoints of profile likelihood based confidence interval for the j^{th} element of parameter array are obtained by solving for values of β_j that satisfy $l_j^*(\beta_j) = l_0$

where

$l_j^*(\beta_j)$ is the maximum likelihood obtained for the set of all parameter value arrays with the j^{th} element fixed at β_j and

$l_0 = l_{\max} - 0.5\chi_{(1,1-\alpha)}^2$, $\chi_{(1,1-\alpha)}^2$: the $100*(1-\alpha)$ percentile of the chi-square distribution with one degree of freedom.

The confidence limits can be found iteratively by approximating the log likelihood function by a quadratic function in a neighborhood of β_j .

In the next section, we will apply Firth's method to Poisson regression model. Profile likelihood based confidence interval is also computed for Poisson Estimates.

POISSON AS EXPONENTIAL FAMILY MODEL

Generalized linear models consist of three components: a random component, a systematic component, and a link.

The random component of a generalized linear model specifies the distribution of the responses $y_1, y_2, \dots, y_n$. Specifically, exponential family model specifies that the responses are independent and have natural exponential distribution

$$f(y_j; \psi_j) = a(y_j) \exp(y_j \psi_j - b(\psi_j))$$

where ψ_j is termed the natural parameter and $b(\psi_j)$ is a function of ψ_j such that

$$\mu_j = \frac{db(\psi_j)}{d\psi_j} \text{ is the mean of response } y_j.$$

The systematic component is the linear predictor $\eta_j = \sum_i x_{ij} \beta_i$

The link g is a monotonic function that relates the response mean to the covariate values through $g(\mu_j) = \eta_j$
Poisson's PMF is given by

$$P(y_i | \lambda_i) = \frac{e^{-\lambda_i} \lambda_i^{y_i}}{y_i!}$$

The log likelihood function is given by

$$l(y | \lambda) = \sum_{i=1}^n (y_i \log(\lambda_i) - \lambda_i - \log(y_i!))$$

Using the Link function; $\log(\lambda_i) = x_i \beta$ log likelihood function can be written as

$$l(\beta | y, \lambda) = \sum_{i=1}^n [y_i (\log(G_i) + x_i \beta) - G_i \exp(x_i \beta) - \log(y_i!)]$$

Here $\beta = (\beta_1, \beta_2, \beta_3, \dots, \beta_p)^T$ are unknown parameters and G_i is the offset term.

The score function is obtained by taking the First derivative of log likelihood function w. r. t. β .

The first derivative will be a (p × 1) vector whose rth term is as follows

$$\sum_{i=1}^n (y_i x_{ir} - G_i x_{ir} \exp(x_i \beta))$$

The maximum likelihood estimate is obtained by solving 'p' score equations given by

$$\sum_{i=1}^n (y_i x_{ir} - G_i x_{ir} \exp(x_i \beta)) = 0 \text{ for } r = 1, 2, 3, \dots, p$$

For Firth's penalty term, compute Fisher's information matrix as

$$I(\beta) = - \frac{\partial^2 (l(\beta))}{\partial \beta^2}$$

Second derivative of log likelihood function will be a (p × p) matrix whose (r, s)th term will be

$$\sum_{i=1}^n (-x_{ir} x_{is} G_i \exp(x_i \beta))$$

Firth corrected log likelihood is

$$FL(\beta) = L(\beta) + \frac{1}{2} \log(|I(\beta)|)$$

To find the bias reduced MLEs (BRMLE), we have to solve following 'p' score equations:

$$\frac{\partial (FL(\beta))}{\partial \beta_k} = 0 \text{ for } k = 1, 2, 3, \dots, p$$

It can be written as

$$\frac{\partial (FL(\beta))}{\partial \beta_k} = \frac{\partial (L(\beta))}{\partial \beta_k} + \frac{1}{2} \text{trace} \left(I(\beta)^{-1} \left(\frac{\partial (I(\beta))}{\partial \beta_k} \right) \right) = 0 \text{ for } k = 1, 2, 3, \dots, p$$

$\frac{\partial(I(\beta))}{\partial\beta_k}$ is the derivative of Information matrix w.r.t. β_k for $k = 1, 2, 3, \dots, p$.

(r,s)th element of $\frac{\partial(I(\beta))}{\partial\beta_k}$ is given by

$$-\frac{\partial^3(L(\beta))}{\partial\beta_r\partial\beta_s\partial\beta_k} = \sum_{i=1}^n (x_{ir}x_{is}x_{ik}G_i \exp(x_i\beta))$$

Using all these formulation, we have developed a SAS Proc to compute the bias reduced MLE. We will demonstrate its use in the next example.

Example1:

In Clinical trials, count response is observed in situations like multiple sclerosis where number of relapses is one of the important outcomes.

A small sample is considered from a data on multiple sclerosis. Our interest is to check if number of relapses depends on baseline EDSS (Expanded Disability Status Scale) score noted at the start of the trial. Baseline EDSS score is converted to binary variable by following transformation

Baseline EDSS Score <=3.5 then EDSS = 1
Else EDSS = 2

Poisson regression model is fitted to the data using the model

Count ~ EDSS + offset(time)

Running following command in Proc LogXact (Software developed by Cytel Inc.) we can get bias reduced estimate. We have also computed profile likelihood based confidence interval to see improvement due to penalization.

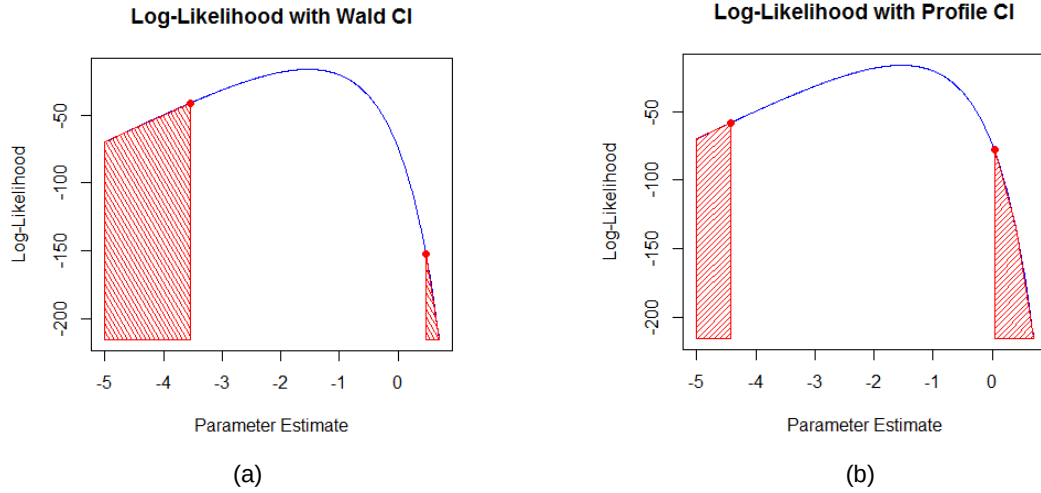
```
DATA MultSclerosys;
INPUT subID EDSS Age SexCode RaceCode Arm count Time;
DATA LINES;
1      1      19      2      1      1      0      758
2      1      54      1      2      2      0      57
3      1      24      2      1      1      1      749.
.
.

30     1      27      1      1      0      1      364
31     2      53      2      1      1      0      751
32     2      46      1      1      0      1      449
RUN;

PROC LOGXACT DATA= MultSclerosys;
MODEL count/Time=EDSS / LINK=POISSON;
ES /AS PMLE EDSS;
RUN;
```

Table1: Comparison of results obtained from traditional ML estimator and bias reduced ML estimator

Parameter	MLE	SE (MLE)	WALD CI_Lower	WALD CI_Upper	Profile CI_Lower	Profile CI_Upper
Constant	-5.2358	1.1002	-7.3922	-3.0793	-7.0468	-2.2763
EDSS	-1.5366	1.026	-3.5475	0.4743	-4.4246	0.0354
Parameter	BRMLE	SE (BRMLE)	WALD CI_Lower	WALD CI_Upper	Profile CI_Lower	Profile CI_Upper
Constant	-5.5893	0.9337	-7.4193	-3.7593	-7.2078	-3.2941
EDSS	-1.1571	0.8473	-2.8178	0.5036	-3.3561	0.2065



Plot1: Log-likelihood function along with Wald CI limits and profile likelihood CI limit.

From the table one can easily see that the bias reduced MLE has lesser standard error as compared to ML estimator. Graphically we can also see how Profile likelihood based CI works for asymmetry in the likelihood. Plot 1(a) shows the left and right tail area beyond Wald CI limits by the shaded portion and Plot 1(b) shows the tail areas beyond Profile likelihood based CI limits.

In the next section we will apply Firth's method to Stratified data. Heinze (2010) has explored this method for stratified logistic regression. It is also available in various softwares including LogXact. Here we will use it for stratified Poisson regression.

Poisson Regression for Stratified data

In Stratified data, population is divided into homogeneous subpopulations or subgroups. Regression model is applied to each subpopulation with same list of covariates.

Conditional likelihood function is used. Conditioning is done for nullifying the effect of stratification.

Suppose γ are stratum specific constants and β values are the parameters to be estimated. γ is the nuisance parameter hence it should be eliminated. If t_β and t_γ are sufficient statistics for β and γ respectively. Note that t_β is a $S \times 1$ vector whose elements are the sum of the responses in each of the S strata. The conditional likelihood function, given t_β , is then only a function of β and is given by

$$L(\beta | t_\beta) = \frac{C(t_\beta, t_\gamma) e^{t_\beta' \beta}}{\sum_{u_\beta \in \Omega} C(u_\beta, t_\gamma) e^{u_\beta' \beta}}$$

where

$C(u_\beta, t_\gamma)$ is computed for the set consisting of all permutations of Y leading to the same $X.Y$ value.

and Ω is the reference set consisting of all possible values of $X.Y$ such that stratum specific sum of Y is same as the observed one.

$$\Omega = \{X.Y \mid E.Y = t_\gamma; Y_i = \{0, 1, 2, \dots\} \text{ for } i = 1, 2, 3, \dots, n\}$$

Then the conditional log-likelihood function is given by

$$l(\beta | \gamma) = \log(L(\beta | \gamma)) = \log(C(t_\beta, t_\gamma)) + t_\beta' \beta - \log\left(\sum_{u_\beta \in \Omega} C(u_\beta, t_\gamma) e^{u_\beta' \beta}\right)$$

PhUSE 2015

Firth's method is applied to this conditional log likelihood function and various properties are studied.

The main advantage of bias reduced ML estimate is that it exists in the case of quasi separation where MLE does not exist. Before going to the example we will briefly explain the concept of separation in the data.

Separation or Quasi-separation:

Problem of Separation (Perfect Fit) occurs when

- Response value divides set of covariate or combinations of covariates
- A covariate value predicts value of the response

X1	X2	Y
-5	-1	0
-4	-2	0
-1	1	0
3	2	1
7	5	1
8	1	1

In this data, notice that

For $X1 < 0$, $Y = 0$: Separation

For $X2 < 0$, $Y = 0$

And $X2 > 1$, $Y = 1$

For $X2 = 1$ there is an overlap
: Quasi-separation

In the next example we will show how Bias reduced MLE works better than MLE in the case of quasi-separation.

Example2:

This example is from Heinze and Puhr (2010) regarding an animal experiment conducted by Bergmeister et al. (2008) to investigate heparin-crosslinked and non-heparinized, xenogeneic vascular substitutes in a rat model. Seventy six rats were divided into strata followed up for one day, three days, seven days, ten days, one month, three months, and six months, and each stratum consisted of equal numbers of rats of both treatment groups. The event of interest was aneurysm formation. In total, only four aneurysms were observed, only among non-heparinized implants and in only two strata.

Since there were only four events, conditional poisson regression was preferred as it nullifies the cluster effect from the likelihood. Using Aneur as Response variable, strata as Stratum and NoHep in the model terms and run following SAS code

In this example quasi-separation is observed in the data. Hence MLE does not exist or Standard errors are very large. In this situation Bias reduced MLE method works better.

```
data Animal;
input Hep Strata Patent Aneur NoHep;
cards;
1 1 1 0 0
1 1 1 0 0
.
.
.
0 182.6 1 0 1
0 182.6 1 0 1
;
run;
PROC LOGXACT DATA=Animal;
ST strata;
MODEL Aneur= NoHep / LINK=POISSON;
ES/AS NoHep;
RUN;

PROC LOGXACT DATA=Animal;
ST strata;
MODEL Aneur= NoHep / LINK=POISSON;
ES/AS PMLE NoHep;
RUN;
```

PhUSE 2015

Running above code MLE and BRMLE are obtained. Comparison of results obtained from traditional ML estimator and bias reduced ML estimator is given in the following table

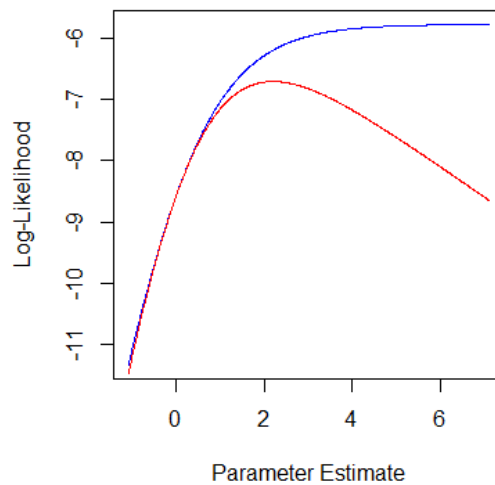
Table2: Comparison of results obtained from traditional ML estimator and Penalized ML estimator

Parameter	MLE	SE (MLE)	WALD CI_Lower	WALD CI_Upper	Profile CI_Lower	Profile CI_Upper
NoHep	28.9996	611802.1173	-1199081.1166	1199139.1158	0.4839	INF

Parameter	BRMLE	SE (BRMLE)	WALD CI_Lower	WALD CI_Upper	Profile CI_Lower	Profile CI_Upper
NoHep	2.1970	1.6665	-1.0693	5.4632	-0.0397	7.0839

Following graph shows log likelihood versus beta parameter plot for Animal data.

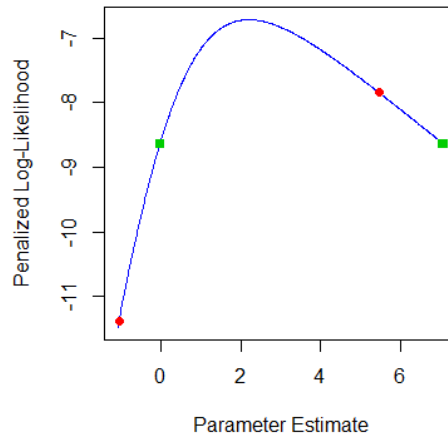
Log-Likelihood and penalized likelihood



Plot2: Log-likelihood and penalized log likelihood functions for Animal data.

Following plot shows wald and Profile likelihood based confidence interval for bias reduced MLE. Here as the likelihood function is not symmetric about its MLE Profile likelihood based CI are more suited than wald CI.

Penalized likelihood with Wald and profile



Plot3: Penalized Log-likelihood function along with Wald CI limits (denoted by red points) and profile likelihood CI limit (denoted by green square).

NEGATIVE BINOMIAL REGRESSION

When responses are counts, the natural distribution for modeling the data is the Poisson distribution. But if the data is over dispersed or suppose one desires to observe a fixed number of successful outcomes (or failures), but the overall number of trials necessary to achieve this is unknown. In these cases, the natural distribution for modeling the data is the Negative Binomial distribution.

We have developed Firth's method for two models of Negative binomial:

1. NB-C with canonical link of Negative Binomial PMF
2. NB-2 with log link

Example 3:

Consider the subset of data from the US national Medicare inpatient hospital database (Medpar) for the 1991 Medicare files for the state of Arizona. The study investigates specific cardiovascular conditions for patients and the data set consists of 1495 observations on ten different variables.

The data set, which is a part of the MedparNegativeBinomial.sas distributed with the package, contains the following variables:

los: the length of the hospital stay for each patient in number of days
hmo: binary variable indicating if the patient belongs to a health maintenance organization (1) or private pay (0)
white: binary variable indicating if the patient identifies themselves as Caucasian (1) versus non-white (0)
died: binary variable indicating if the patient died
age80: binary variable indicating if the patient age is 80 or over
type: a three-level factor predictor related to the type of admission. (1) = elective, or referred, (2)=urgent, and (3)=emergency
type1: binary variable indicating if the patient had an elective admission
type2: binary variable indicating if the patient had an urgent admission
type3: binary variable indicating if the patient had an emergency admission
provnum: provider ID (integer)

It is of interest to see if the data collected indicates if the length of hospital stay is related to the certain patient attributes. Define a Negative Binomial model that regresses the number of days of hospitalization on the type of admission and conduct asymptotic inference on the type regression coefficient. We can define such a negative binomial regression model and estimate its parameters. This model is expressed as

$\text{los} \sim \text{type}$

In this model, los is the Response variable and type is a model covariate. The Asymptotic regression coefficients are estimated by the ES/AS statement as shown in the SAS code given below:

```
Data Medpar;
Input Sr los hmo white died age80 type type1 type2 type3 provnum$;
datalines;
1 4 0 1 0 0 1 1 0 0 30001
2 9 1 1 0 0 1 1 0 0 30001
. . . . .
. . . . .
1495 32 0 1 0 0 2 0 1 0 32003
run;
PROC LOGXACT DATA=Medpar ;
MODEL los=type / LINK=NEGBIN2;
ES/AS type;
CLASS _LO type;
RUN;
```

Table3: results obtained from ML estimator

Parameter	MLE	SE (MLE)	WALD CI Lower	WALD CI Upper	Profile CI Lower	Profile CI Upper
Constant	2.1785	0.0222	2.1349	2.2221	2.1348	2.2220
type_2	0.2374	0.0502	0.1390	0.3358	0.1396	0.3364
type_3	0.7249	0.0757	0.5765	0.8732	0.5793	0.8762

Again fit the previous model (Negative Binomial Model: NB2, Ancillary Parameter Value:

PhUSE 2015

Estimate), this time allowing for bias reduced estimates of the parameters. This can be done using the PMLE option in the PROC syntax.

```
Data Medpar;
Input Sr los hmo white died age80 type type1 type2 type3 provnum$;
datalines;
1 4 0 1 0 0 1 1 0 0 30001
2 9 1 1 0 0 1 1 0 0 30001
. . . . .
. . . . .
1495 32 0 1 0 0 2 0 1 0 32003
run;
PROC LOGXACT DATA=Medpar ;
MODEL los=type / LINK=NEGBIN2;
ES/AS PMLE type;
CLASS_LO type;
RUN;
```

Notice that, as the data is large, MLE does not have small sample bias hence MLE and BRMLE estimates are close. Next table gives comparison of MLE and BRMLE estimates obtained from NB-2 model.

Table4: Comparison of results obtained from traditional ML estimator and bias reduced ML estimator

Parameter	MLE	SE (MLE)	WALD CI Lower	WALD CI Upper	Profile CI Lower	Profile CI Upper
Constant	2.1785	0.0222	2.1349	2.2221	2.1348	2.2220
type_2	0.2374	0.0502	0.1390	0.3358	0.1396	0.3364
type_3	0.7249	0.0757	0.5765	0.8732	0.5793	0.8762
Parameter	BRMLE	SE (BRMLE)	WALD CI Lower	WALD CI Upper	Profile CI Lower	Profile CI Upper
Constant	2.1781	0.0222	2.1345	2.2217	2.1347	2.2219
type_2	0.2368	0.0502	0.1384	0.3352	0.1391	0.3359
type_3	0.7235	0.0756	0.5752	0.8717	0.5776	0.8741

CONCLUSION

Overall, bias reduced estimates show shorter confidence intervals for count data. It also exists in some of the situations where MLE fails like small or sparse data or separation in the data. But if data is large enough to follow asymptotic assumptions, this new method does not show any improvements over MLE.

REFERENCES

- Anderson, J., Richardson, S. (1979). Logistic discrimination and bias correction in maximum likelihood estimation. *Technometrics* 21, 71–78.
- Cordeiro, G., McCullagh, P. (1991). Bias correction in generalized linear models. *Journal of the Royal Statistical Society, Series B: Methodological* 53, 629–643.
- Cox, D., Snell, E. (1968). A general definition of residuals (with discussion). *Journal of the Royal Statistical Society, Series B: Methodological* 30, 248–275.
- Firth D (1993). “Bias reduction of maximum likelihood estimates”. *Biometrika* (1993), 80, 1, pp. 27-38
- Hardin J and Hilbe J (2001). *Generalized Linear Model and Extensions*. Stata Press, College Station, Texas.
- Heinze G. (2006) A comparative investigation of methods for logistic regression with separated or nearly separated data. *Statistics in Medicine*, 25, 4216-4226.
- Heinze G., Ploner M. (2004), A SAS macro, S-PLUS library, and R package to perform logistic regression without convergence problem.
- Heinze G., Puh R. (2010). “Bias-reduced and separation-proof conditional logistic regression with small or sparse data sets”. *Statistics in Medicine*
- Heinze G and Schemper M (2002). A solution to the problem of separation in regression. *Statistics in Medicine*, 21, 2409-2419.
- Hilbe J (2011), “*Negative Binomial Regression*”, 2nd ed. 2011, University Press, Cambridge.
- Kosmidis I (2007). “*Bias Reduction in Exponential Family Nonlinear Models*”.

PhUSE 2015

Kosmidis I (2010). "A generic algorithm for reducing bias in parametric estimation". Electronic Journal of Statistics. Vol. 4 (2010) 1097–1112
LogXact 11 User Manual.
Schaefer, R.(1983). Bias correction in maximum likelihood logistic regression. Statistics in Medicine 2, 71–78.
Venzon D., Moolgavkar S.(1988) , A method for computing Profile-Likelihood- Based intervals, Applied Statistics, 37;1:87:94

ACKNOWLEDGMENTS

The authors would like to thank Dr. Sharayu Paranjape and Dr. Pralay Senchaudhuri for reviewing and providing their valuable feedback. The authors would also like to thank all the colleagues at Cytel who gave their valuable inputs as and when required.

CONTACT INFORMATION

Your comments and questions are valued and encouraged. Contact authors at:

Ashwini Joshi

Email: ashwini.joshi@cytel.com

Sumit Singh

Email: sumit.singh@cytel.com

Cytel Statistical Software & Services Pvt. Ltd. (Subsidiary of Cytel Inc.)
6th Floor, Lohia-Jain IT Park, A Wing, Chandani Chowk, Paud Road,
Kothrud, Pune 411 038, India

Brand and product names are trademarks of their respective companies.