

Sample size calculation for comparing two groups of count data in clinical trials

Lijuan Lin, Sanofi Corporation;

ABSTRACT

Count data are common in clinical trials, in which each subject may experience the same type of event more than once. Negative binomial models have been widely used to model count data in recent years as it provides a convenient way to account for the overdispersion for the data. Several sample size calculation methods have recently been developed and applied in PASS with simplified assumptions including a common dispersion parameter and a common follow-up time. However, dispersion parameter could differ by treatments in practice, and it is often challenging to specify at study design stage. We developed a SAS macro to implement the method proposed by Wang and Fan (2019), which was derived using the method of moments by matching the first and second moments of the distribution of the count data and can be applied to count data following any distribution in addition to the negative binomial distribution. Furthermore, the method was generalized to noninferiority tests and equivalence tests for comparing the rate ratio of two groups of count data and the precision of the method was evaluated using simulations.

INTRODUCTION

Count data are common in clinical trials, in which each subject may experience the same type of event more than once. Examples include relapses in multiple sclerosis (MS), bleeding episode in Hemophilia disease, tumor recurrence in cancer trial. Recently, negative binomial models have become increasingly popular for analyzing over dispersed count data because its greater flexibility in modeling the mean-variance relationship than the highly restrictive Poisson model. Some simulation-based methods were proposed to calculate the required sample size (Aban et al. 2008; Rettiganti et al. 2012). It will be convenient to have a simple explicit formula to calculate sample size for given assumptions and desired power. An explicit formula can also help trial designers to better understand how the parameters such as effect size, duration of the treatment period, and dispersion parameter are correlated to each other and how they affect the sample size. Such understanding may help planners to choose the best design to serve the objectives of the trials during the planning stage.

Great efforts have been made in the past few years and several sample size calculation methods have been developed for negative binomial models. Zhu and Lakkis (2014) derived an explicit sample size formula for the negative binomial model, assuming all subjects have a common follow-up time and a common dispersion parameter. Tang (2015) derived another sample size formula that allows unequal follow-up times for different subjects. Later on, Tang (2018) and Zhu (2017) generalized their sample size methods to noninferiority or equivalence trials. All those methods fully utilizing the negative binomial distribution function through the likelihood function to calculate the required sample size assuming a common dispersion parameter. It is no doubt that when the collected count data were analysis by likelihood-based approaches, those sample size calculation methods might be the most appropriate. However, those methods cannot be applied if the distribution of the count data is not negative binomial, making these formulas not generalizable when assessing potentially different study designs. In addition, to estimate sample sizes for Poisson or negative binomial models, it is necessary to have a reliable estimate of the dispersion parameters, which is often challenging to specify at the design stage of a trial. Furthermore, all those sample size methods mentioned above assume both treatments shared the same dispersion parameter, which may not be the real case in practice especially when two treatment groups are expected to be different.

Wang and Fan (2019) proposed a sample size formula for comparing two groups of count data using the method of moments, which does not need to make any specific distribution assumption of the count data, but only rely on the mean (the first moment) and variance (the second moment) of distribution. It was established that the proposed formula can be used even when the study is analyzed with a likelihood-based approach and can be viewed as an upper bound of the required sample size by likelihood-based

approaches at study design stage.

In this paper, we developed SAS macro to implement the method proposed by Wang and Fan (2019) and extended the method to noninferiority tests and equivalence tests for comparing the rate ratio of two groups of count data. The precision of the extended method was evaluated through simulations. The SAS macro provide a simple and flexible way to help designers to calculate the sample sizes for count data under different trial design without any specific distribution assumption.

SAMPLE SIZE CALCULATION FORMULA

SAMPLE SIZE CALCULATION FOR TEST OF DIFFERENCE IN RATE RATIO

Let Y_{ji} be the number of events during time T_{ji} for subject i ($i=1, 2, \dots, n_j$) in group j ($j=0$ for control group and $j=1$ for active treatment group), where n_0 and n_1 are sample sizes for groups 0 and 1, respectively. Assume the counts $Y_{01}, \dots, Y_{0n_0}$ are independently distributed with means $r_0 T_{01}, \dots, r_0 T_{0n_0}$ and variances $\sigma_{01}^2, \dots, \sigma_{0n_0}^2$, respectively, and $Y_{11}, \dots, Y_{1n_1}$ are independently distributed with means $r_1 T_{11}, \dots, r_1 T_{1n_1}$ and variances $\sigma_{11}^2, \dots, \sigma_{1n_1}^2$, respectively. With such assumptions, each group is assumed to have a common standardized rate, r_j , for $j=0, 1$.

For negative binomial regression, count data can be modeled using the log-transformed regression with an offset term, which can be expressed as $\log(r_j T_{ji}) = \beta_0 + \beta_1 X_{ji} + \log(T_{ji})$. X_{ji} is the treatment group indicator for subject i , taking the value 0 if $j=0$ for control group and 1 if $j=1$ for active treatment group. It is easy to see that $r_1/r_0 = e^{\beta_1}$. The tested null hypothesis is that there is no difference in the event rate between the two treatment groups, and the alternative is that the rates are different, that is,

$$H_0: r_1/r_0 = 1, H_a: r_1/r_0 \neq 1,$$

or equivalently,

$$H_0: \beta_1 = 0, H_a: \beta_1 \neq 0.$$

Wang and Fan (2019) who provide a general sample size formula for comparing two groups of count data using method of moments which does not need any further distribution assumption. Then the null hypothesis $H_0: \beta_1 = 0$ can be tested by a two-sided Wald test with the statistic

$$Z = \frac{\sqrt{n_0} \hat{\beta}_1}{\sqrt{\hat{\sigma}_0^2 / (\hat{r}_0^2 \bar{T}_0^2) + \hat{\sigma}_1^2 / (\theta \hat{r}_1^2 \bar{T}_1^2)}}$$

where $\hat{r}_0, \hat{r}_1, \hat{\sigma}_0^2$ and $\hat{\sigma}_1^2$ are consistent estimators of r_0, r_1, σ_0^2 and σ_1^2 .

Suppose we seek a power of $1 - \beta$ to reject the null hypothesis $H_0: \beta_1 = 0$ in favor of the alternative $H_a: \beta_1 \neq 0$ at a two-sided significant level α , the sample size can be calculated as

$$n_0 \geq \frac{(Z_{\alpha/2} \sqrt{V_0} + Z_\beta \sqrt{V_1})^2}{(\log(r_1/r_0))^2}$$

$$n_1 = \theta n_0$$

where Z_β is the upper β th quantile of the standard normal distribution, r_1/r_0 is the rate ratio under H_a or the assumed true rate ratio. And V_0 and V_1 are the variance term evaluated under H_0 and H_a , respectively. More specifically,

$$V_0 = \frac{\tilde{\sigma}_0^2}{\tilde{r}_0^2 \bar{T}_0^2} + \frac{\tilde{\sigma}_1^2}{\theta \tilde{r}_1^2 \bar{T}_1^2}$$

$$V_1 = \frac{\sigma_0^2}{r_0^2 \bar{T}_0^2} + \frac{\sigma_1^2}{\theta r_1^2 \bar{T}_1^2}$$

where $\tilde{r}_0, \tilde{r}_1, \tilde{\sigma}_0^2$ and $\tilde{\sigma}_1^2$ are means and variances evaluate under H_0 , and r_0, r_1, σ_0^2 and σ_1^2 are those

evaluate under H_a or the assumed true rates. The V_0 can be determined by using either reference group rate or true rates. As previous studies have shown that using true rates may yield more robust estimated sample size than using reference group, the true rates are used in this manuscript to calculate the required sample sizes, which means $\tilde{r}_0 = r_0$, $\tilde{r}_1 = r_1$, $\tilde{\sigma}_0^2 = \sigma_0^2$ and $\tilde{\sigma}_1^2 = \sigma_1^2$.

In addition, for a given sample size n_0 , the power can be calculated as

$$1 - \beta = \Phi \left(\frac{\sqrt{n_0} |\log(r_1/r_0)| - Z_{\alpha/2} \sqrt{V_0}}{\sqrt{V_1}} \right)$$

where Φ is the cumulative distribution function for a standard normal distribution.

SAMPLE SIZE CALCULATION FOR NONINFERIORITY TESTS FOR THE RATE RATIO

In Poisson or negative binomial models, the natural metric for between treatment comparison is the rate ratio. Therefore, the noninferiority (NI) margin is usually defined for the rate ratio. Let M denote the pre-specified noninferiority margin for rate ratio of the new treatment versus the reference. Clearly if a lower rate is better, then $M > 1$; otherwise, $M < 1$. In this paper, to simplify the presentation, we assume that the event under study is undesirable, so a lower rate is better. The null and alternative hypotheses can be written as

$$H_0: r_1/r_0 \geq M, H_a: r_1/r_0 < M$$

For a one-sided significant level α , the null hypothesis can be rejected if the upper boundary of the one-sided $(1 - \alpha)$ confidence interval for the rate ratio r_1/r_0 is less than M , in which the confidence interval for the rate ratio can be calculated using the Wald method.

Suppose we seek a power $1 - \beta$ to reject the null hypothesis $H_0: r_1/r_0 \geq M$ in favor of the alternative $H_a: r_1/r_0 < M$ at a one-sided significance level α . Under the assumption that $\beta_1 = \log(\hat{r}_1/\hat{r}_0)$ asymptotically follows a normal distribution, therefore by using true rate, the sample size can be estimated as

$$n_0 \geq \frac{(Z_\alpha \sqrt{V_0} + Z_\beta \sqrt{V_1})^2}{(\log(M) - \log(r_1/r_0))^2}$$

$$n_1 = \theta n_0$$

SAMPLE SIZE CALCULATION FOR EQUIVALENCE TESTS FOR THE RATE RATIO

Sometimes the objective of a clinical trial is to demonstrate the clinical equivalence of two different treatments. In the context of count data, an equivalence trial is to demonstrate that the event rate ratio of two treatments is within prespecified boundaries. Therefore, the null and alternative hypotheses can be written as

$$H_0: r_1/r_0 \leq L \text{ or } r_1/r_0 \geq U$$

$$H_a: L < r_1/r_0 < U$$

where $L < 1$ and $U > 1$ are the prespecified lower and upper equivalence boundaries on the rate ratio. For an α level test, the null hypothesis can be rejected if two-sided $(1 - 2\alpha)$ confidence interval for the rate ratio r_1/r_0 fall completely within the interval (L, U) .

Then by using true rates the power of the above test can be calculated as

$$\Phi \left(\frac{\sqrt{n_0} |\log(r_1/r_0) + \log(L)| - Z_\alpha \sqrt{V_0}}{\sqrt{V_1}} \right) + \Phi \left(\frac{\sqrt{n_0} |\log(r_1/r_0) - \log(U)| - Z_\alpha \sqrt{V_0}}{\sqrt{V_1}} \right) - 1$$

There is no closed-form sample size formula, the sample size can be obtained by searching n_0 until the desired power is achieved. As rate ratio is the natural metric for between treatment comparison, the equivalence boundaries are usually chosen in the way that the logarithm of the boundary values will be symmetric about zero, that is, $-\log(L) = \log(U)$. In this case, the sample size can be obtained by

$$n_0 \geq \frac{(Z_{\alpha}\sqrt{V_0} + Z_{\beta/2}\sqrt{V_1})^2}{(\log(U) - \log(r_1/r_0))^2}$$

SAS MACRO TO IMPLEMENT SAMPLE SIZE CALCULATION

```
%macro Count_power_ssize(mu0, mu1,T0, T1, ssize_ratio, sig_level=0.05,
alternative="two.sided", test="Difference", higher="worse",sd0=, sd1=,
power=, Margin=, n0=, n1=);

/* check input parameters */
  %if &mu0 <= 0 or &mu1 <= 0 or &T0 <= 0 or &T1 <= 0 or &ssize_ratio <=
0 %then %do;
    %put ERROR: All numeric parameters must be greater than 0;
    %goto ExitMacro;
  %end;

/* one-side or two-side */
  %let sided = 2;
  %if &alternative. = "one.sided" %then %let sided = 1;
  %let sig_level_new = &sig_level./&sided.;

/* Calculate event rate for each groups */
  %let rr0 = %sysevalf(&mu0./&T0.);
  %let rr1 = %sysevalf(&mu1./&T1.);

/* Calculate V1*/
  %let V1 = %sysevalf(&sd0.**2 / (&rr0.**2 * &T0.**2) + &sd1.**2 /
(&ssize_ratio. * &rr1.**2 * &T1.**2));
/* Calculate V0? use true rate*/
  %let V0 = %sysevalf(&sd0.**2 / (&rr0.**2 * &T0.**2) + &sd1.**2 /
(&ssize_ratio. * &rr1.**2 * &T1.**2));

/* calculate power under different type of trial desgin*/
  %if %quote(&power.)=%str() %then %do;
    %if %upcase(&test.) = "DIFFERENCE" %then %do;
      %let power = probnorm((sqrt(&n0.) *
abs(log(&rr1./&rr0.)) - probit(1 - &sig_level_new.) * sqrt(&V0.)) /
sqrt(&V1.));
    %end;
    %if %upcase(&test.) = "NONINFERIORITY" %then %do;
      %if %upcase(&higher.) = "BETTER" %then %do;
        %let power = probnorm((sqrt(&n0.) *
(log(&rr1./&rr0.)-log(&Margin.)) - probit(1 - &sig_level_new.) * sqrt(&V0.))
/ sqrt(&V1.));
      %end;
      %if %upcase(&higher.) = "WORSE" %then %do;
        %let power = probnorm((sqrt(&n0.) *
(log(&Margin.) - log(&rr1./&rr0.)) - probit(1 - &sig_level_new.) *
sqrt(&V0.)) / sqrt(&V1.));
      %end;
    %end;
    %if %upcase(&test.) = "EQUIVALENCE" %then %do;
```

```

                                %let power = probnorm((sqrt(&n0.) *
abs(log(&rr1./&rr0.) + log(&Margin.)) - probit(1 - &sig_level_new.) *
sqrt(&V0.)) / sqrt(&V1.)) +
                                probnorm((sqrt(&n0.) *
abs(log(&rr1./&rr0.) - log(&Margin.)) - probit(1 - &sig_level_new.) *
sqrt(&V0.)) / sqrt(&V1.))-1;
                                %end;
                                %end;

/* calculate sample size under different type of trial desgin*/
%else %do;
    %if %upcase(&test.) = "DIFFERENCE" %then %do;
        %let n0 = ((probit(1 - &sig_level_new.) * sqrt(&V0.)
+ probit(&power.) * sqrt(&V1.))/log(&rr1./&rr0.))**2;
    %end;
    %if %upcase(&test.) = "NONINFERIORITY" %then %do;
        %let n0 = ((probit(1 - &sig_level_new.) * sqrt(&V0.)
+ probit(&power.) * sqrt(&V1.))/(log(&Margin.) - log(&rr1./&rr0.))**2;
    %end;
    %if %upcase(&test.) = "EQUIVALENCE" %then %do;
        %let n0 = ((probit(1 - &sig_level_new.) * sqrt(&V0.)
+ probit((1+&power.)/2)*sqrt(&V1.))/(log(&Margin.)-log(&rr1./&rr0.))**2;
    %end;
    %end;
    %let n1 = &n0.*&ssize_ratio.;
    data rlt;
        mu0 = &mu0.;
        mu1 = &mu1.;
        sd0 = &sd0.;
        sd1 = &sd1.;
        duration0 = &T0.;
        duration1 = &T1.;
        EventRate0 = &rr0.;
        EventRate1 = &rr1.;
        SSizeRatio = &ssize_ratio.;
        sig_level = &sig_level.;
        alternative = &alternative.;
        Test = &test.;
        if upcase(&test.) = "NONINFERIORITY" then Higher = &higher.;
        Power = &power.;
        N0 = ceil(&n0.);
        N1 =ceil(&n1.);
        Ntotal = N0 +N1;
    run;
    proc print data=rlt;
    run;

    %ExitMacro:
%mend Count_power_ssize;

/* example 1: rate ratio difference test */
%Count_power_ssize(mu0=4, mu1=3, sd0=5, sd1=5, T0=1, T1=1, ssize_ratio=1,
sig_level=0.05, test="difference", alternative="two.sided", power=0.8);

/* example 2: rate ratio Noninferiority test */

```

```
%Count_power_ssize(mu0=4, mu1=4, sd0=5, sd1=5, T0=1, T1=1, ssize_ratio=1,
sig_level=0.025, test="noninferiority", alternative="one.sided", power=0.8,
higher="worse", Margin=1.4);
```

```
/* example 3: rate ratio Equivalence test */
%Count_power_ssize(mu0=4, mu1=4, sd0=3, sd1=5, T0=1, T1=1, ssize_ratio=1,
sig_level=0.025, test="EQUIVALENCE", alternative="one.sided", power=0.8,
Margin=0.7);
```

SIMULATIONS TO ASSESS THE ACCURACY OF SAMPLE SIZE CALCULATION FORMULA

The accuracy of moment-based sample size calculation formula for comparing the difference of rate ratio for two group of count data has been demonstrated in different scenario through number of simulations in previously study (Wang and Fan, 2019). The simulation in this paper was focus on assess the accuracy of the extended method for non-inferiority and equal trials. Various values for noninferiority margins M or equivalence margins (L, U) and true rate ratio r_1/r_0 , were simulated. For each parameter setting, sample sizes were calculated with a target power of 80% for 0.025 level noninferiority or equivalence tests. The number of events for each patient was generated using random number generators for negative binomial distributions with the average exposure time and the mean event rate for the treatment group to which the patient belongs. When assuming a common dispersion across treatment, the generated number of events were analyzed using negative binomial regression with the logarithm as the link function, treatment group as the independent variable, and the logarithm of the exposure time as the offset term. When the dispersion is differed by treatments, then negative binomial regression is performed separately for each treatment group and the test statistic is calculated based on the analytic formula. For each scenario with the calculated sample size, 10,000 datasets were generated and analyzed, and the percentage of the significant instances was reported as the empirical power. An appropriate sample size calculation method should result in an empirical power for the calculated sample size is close to the nominal power (80%).

Table 1 Required sample size and corresponding simulated powers for negative binomial distribution.

Control Group			Treatment Group			Rate Ratio	Margin	Estimated sample size	PASS sample size	Simulated power, %
Mean	SD	Duration	Mean	SD	Duration					
Noninferiority tests										
Different noninferiority margin assuming a common dispersion across treatments (k = 1.3125)										
4	5	1	4	5	1	1	1.5	300	300	80.12
4	5	1	4	5	1	1	1.4	434	434	79.90
4	5	1	4	5	1	1	1.3	714	714	80.29
4	5	1	4	5	1	1	1.2	1476	1476	80.00
Different noninferiority margin assuming differ dispersion										
4	3	1	4	5	1	1	1.5	204	-	80.68
4	3	1	4	5	1	1	1.4	296	-	81.01
4	3	1	4	5	1	1	1.3	486	-	80.38
4	3	1	4	5	1	1	1.2	1004	-	80.63
Different rate ratio										
4	3	1	3	5	1	0.75	1.4	136	-	81.44
4	3	1	3.5	5	1	0.875	1.4	186	-	81.47
4	3	1	4	5	1	1	1.4	296	-	81.01
4	3	1	4.5	5	1	1.125	1.4	590	-	80.82
Equivalence test										
Different equivalence margin assuming a common dispersion across treatments (k = 1.3125)										
4	5	1	4	5	1	1	0.7	518	518	80.16
4	5	1	4	5	1	1	0.75	794	794	80.06
4	5	1	4	5	1	1	0.8	1320	1320	79.76
4	5	1	4	5	1	1	0.85	2488	2488	79.89
Different equivalence margin assuming differ dispersion										
4	3	1	4	5	1	1	0.7	352	-	79.50
4	3	1	4	5	1	1	0.75	540	-	79.68
4	3	1	4	5	1	1	0.8	898	-	80.12

4	3	1	4	5	1	1	0.85	1692	-	79.67
Different rate ratio										
4	3	1	3.5	5	1	0.875	0.7	822	-	79.02
4	3	1	4	5	1	1	0.7	352	-	79.50
4	3	1	4.5	5	1	1.125	0.7	496	-	80.50
4	3	1	5	5	1	1.25	0.7	1376	-	80.10

Margin: For noninferiority test, the margin indicated NI margin. For equivalence test, margin indicates the lower boundary (L) of equivalence interval, and $-\log(L) = \log(U)$.

PASS only provide sample size calculation based on common dispersion.

Give the assumption that different treatment group have a common dispersion parameter of the negative binomial distribution, the estimated sample size based on method of moment are consistent with the sample size obtained by PASS. When the dispersion is different across treatment groups, PASS is unable to provide sample size calculations as it does not allow to input differ dispersion parameter. As it shown in Table 1, the extended moment-based sample size calculation formula could provide reasonably accurate sample sizes for the intended power (80% in simulation studies). In general, the sample sizes calculated by extended moment-based method closely achieved the targeted 80% power.

CONCLUSION

The paper provides a SAS macro to implement the moment-based sample size calculation for count data, which has been proposed by Wang and Fan and extended to noninferiority and equivalence trials in this paper. The simulation results under different scenario confirmed the appropriateness of using the moment-based formula to estimate required sample size. Unlike the likelihood-based method, the moment-based method does not need to use the specific distribution assumption of the count data but only need the knowledge of the first moment (mean) and the second moment (variance) of the distribution, which is simpler and more flexible for counts data.

REFERENCES

- Aban, I. B., Cutter, G. R., & Mavinga, N. (2008). Inferences and Power Analysis Concerning Two Negative Binomial Distributions with An Application to MRI Lesion Counts Data. *Computational statistics & data analysis*, 53(3), 820–833. <https://doi.org/10.1016/j.csda.2008.07.034>
- Rettiganti, M., & Nagaraja, H. N. (2012). Power analyses for negative binomial models with application to multiple sclerosis clinical trials. *Journal of biopharmaceutical statistics*, 22(2), 237–259. <https://doi.org/10.1080/10543406.2010.528105>
- Zhu, H., & Lakkis, H. (2014). Sample size calculation for comparing two negative binomial rates. *Statistics in medicine*, 33(3), 376–387. <https://doi.org/10.1002/sim.5947>.
- Tang Y. (2015). Sample Size Estimation for Negative Binomial Regression Comparing Rates of Recurrent Events with Unequal Follow-Up Time. *Journal of biopharmaceutical statistics*, 25(5), 1100–1113. <https://doi.org/10.1080/10543406.2014.971167>.
- Zhu, H. (2017). Sample size calculation for comparing two Poisson or negative binomial rates in noninferiority or equivalence trials. *Statistics in Biopharmaceutical Research* 9:107–115. <https://doi.org/10.1080/19466315.2016.1225594>.
- Tang Y. (2018). Sample size for comparing negative binomial rates in noninferiority and equivalence trials with unequal follow-up times. *Journal of biopharmaceutical statistics*, 28(3), 475–491. <https://doi.org/10.1080/10543406.2017.1333998>.
- Wang, L., & Fan, C. (2019). Sample size calculations for comparing two groups of count data. *Journal of biopharmaceutical statistics*, 29(1), 115–127. <https://doi.org/10.1080/10543406.2018.1489409>

CONTACT INFORMATION

Your comments and questions are valued and encouraged. Contact the author at:

Lijuan Lin (Grace Lin)
lijuanlin94@163.com

Any brand and product names are trademarks of their respective companies.