



The bootstrap and Bayesian bootstrap method in assessing bioequivalence

Wan Jianping^{a,*}, Zhang Kongsheng^{a,c}, Chen Hui^b

^a Department of Mathematics, Huazhong University of Science and Technology, Wuhan 430074, PR China

^b Institute of Clinical Pharmacology, Tongji Medical College, Huazhong University of Science and Technology, Wuhan 430030, PR China

^c College of Statistics and Mathematics of Anhui Finance and Economics University, Bengbu 233030, PR China

ARTICLE INFO

Article history:

Accepted 26 August 2008

Communicated by Prof. Ji-Huan He

ABSTRACT

Parametric method for assessing individual bioequivalence (IBE) may concentrate on the hypothesis that the PK responses are normal. Nonparametric method for evaluating IBE would be bootstrap method. In 2001, the United States Food and Drug Administration (FDA) proposed a draft guidance. The purpose of this article is to evaluate the IBE between test drug and reference drug by bootstrap and Bayesian bootstrap method. We study the power of bootstrap test procedures and the parametric test procedures in FDA (2001). We find that the Bayesian bootstrap method is the most excellent.

© 2008 Elsevier Ltd. All rights reserved.

1. Introduction

The aim of bioequivalence (BE) studies is to assess the equivalence of two pharmaceutical drug of the same active drug substance [6]. BE generally have three types including average bioequivalence (ABE), population bioequivalence (PBE) and individual bioequivalence (IBE). For assessing ABE, these measures may be considered such as area under the curve (AUC) and peak concentration ($C_{\max}$) [2]. ABE focuses only on the difference of average measure between test drug (T) and reference drug (R), the interest measure may be area under curve and peak concentration. But ABE ignores the variability of the measure for T and R . PBE emphasizes the total variability of the measure in the population. IBE takes into account the within-subject variability and subject by formulation interaction for T and R . The mixed-effects model usually be used to evaluate BE.

FDA [1] studied IBE using the original method named bootstrap percentile method. Jun shao et al. [3] improved the assessing procedure of FDA [1] and applied some special bootstrap method to enhance power of the test procedure for IBE. Pigeot [4] continued to investigate IBE by bootstrap percentile method. By now we do not find these paper with respect to the Bayesian bootstrap method in assessing bioequivalence. Bayesian bootstrap is superior to bootstrap method in generating the random sample and simulated power.

The rest of this article is organized as follows. In Section 2, we introduce the bootstrap method and Bayesian bootstrap method. In Section 3, we provide a description of the statistical model and criteria for IBE in Appendix G of FDA's Guidance [2]. In Section 4, the power of different method for test procedures are simulated. Some conclusion is given in Section 5.

2. Bootstrap and Bayesian bootstrap

In 1979, Efron [5] proposed a new method named bootstrap to simulate confidence upper bound for interest parameter such as mean and variance. Now there are many different styles about the bootstrap. In this article we only concentrate the bootstrap percentile, hybrid bootstrap and Bayesian bootstrap method.

Let $X = (X_1, X_2, \dots, X_n)$ be a random sample and $x = (x_1, x_2, \dots, x_n)$ an realization of X . One bootstrap sample from $(x_1, x_2, \dots, x_n)$ is a random sample from $x_1, x_2, \dots, x_n$ with replacement. θ be an interest parameter and $\hat{\theta}$ an estimator. For

* Corresponding author.

E-mail address: hust_jp_w@yahoo.com.cn (W. Jianping).

example θ is mean and $\hat{\theta} = \bar{x}$, we can calculate the bootstrap estimator of θ , i.e. $\hat{\theta}^* = \bar{x}^*$. Repeating this step for B times, we obtain a series of bootstrap estimator of θ , i.e. $\hat{\theta}^{*b}$, $b = 1, 2, \dots, B$, sort them with order $\hat{\theta}^*(1), \hat{\theta}^*(2), \dots, \hat{\theta}^*(B)$ so that $\hat{\theta}^*(i)$ is the i th of $\hat{\theta}^{*b}$ and $\hat{\theta}^*(i) < \hat{\theta}^*(j)$ for $i < j$. Then an one-sided $100(1 - \alpha)$ percent interval for θ is approximated by $\hat{\theta}^*(B - B\alpha)$ (for convenience we choose B such that $B - B\alpha$ an integer). The hybrid bootstrap is to approximate the distribution of $\hat{\theta} - \theta$ by $\hat{\theta}^* - \hat{\theta}$. On the basis of bootstrap percentile method, the approximated confidence interval is $2\hat{\theta} - \hat{\theta}^*(B\alpha)$.

Rubin [7] discussed the Bayesian bootstrap method to construct confidence interval. The main idea is as follows: Drawing $(n - 1)$ random variables $u_1, u_2, \dots, u_{n-1}$, sorting them and calculating the gaps $g_i = u_{(i)} - u_{(i-1)}$, where $u_{(0)} = 0$ and $u_{(n)} = 1$. Then Bayesian bootstrap sample is $X^* = (g_1 X_1, g_2 X_2, \dots, g_n X_n)$. The confidence interval for B is similar to original bootstrap.

3. Statistical model and criterion for IBE

To assess IBE the s -sequence and four-period experiment usually be considered. FDA[2] studied the mixed-effect model

$$Y_{ijkl} = \mu_k + \gamma_{ikl} + \delta_{ijk} + \varepsilon_{ijkl} \quad (3.1)$$

where $i = 1, 2, \dots, s$ indicates sequence, $j = 1, 2, \dots, n_i$ indicates subject within sequence i , $k = R, T$ indicates treatment, $l = 1$ and 2 indicates replicate on treatment k for subjects within sequence i . Y_{ijkl} is the response of replicate 1 on treatment k for subject j in sequence i , γ_{ikl} represents the fixed effect of replicate 1 on treatment k in sequence i , δ_{ijk} is the random subject effect for subject j in sequence i on treatment k , and ε_{ijkl} is the random error for subject j within sequence i on replicate 1 of treatment k . The linearized criteria are as follows in FDA [2]

1. reference-scaled ($\sigma_{WR}^2 \geq \sigma_{W0}^2$):

$$\eta_1 = (\mu_T - \mu_R)^2 + \sigma_D^2 + \sigma_{WT}^2 - \sigma_{WR}^2 - \theta_1 \cdot \sigma_{WR}^2 \quad (3.2)$$

2. constant-scaled ($\sigma_{WR}^2 < \sigma_{W0}^2$):

$$\eta_2 = (\mu_T - \mu_R)^2 + \sigma_D^2 + \sigma_{WT}^2 - \sigma_{WR}^2 - \theta_1 \cdot \sigma_{W0}^2 \quad (3.3)$$

where μ_T and μ_R indicate population average responses of the log-transformed measure for the T and R formulation, respectively. σ_D^2 indicates subject-by-formulation interaction variance component, σ_{WT}^2 and σ_{WR}^2 represent the within-subject variance of the T formulation and R formulation respectively. σ_{W0}^2 represents specified constant within-subject variance and θ_1 BE limit. We consider the testing hypothesis.

$$H_0 : \eta \geq 0 \quad \text{versus} \quad H_1 : \eta < 0 \quad (3.4)$$

$$\eta = \eta_1 \text{ if } \sigma_{WT}^2 \geq \sigma_{W0}^2 \text{ and } \eta = \eta_2 \text{ if } \sigma_{WT}^2 < \sigma_{W0}^2.$$

Some statistics are defined as:

$$I_{ij} = Y_{ijT} + Y_{ijR}, \quad T_{ij} = Y_{ijT_1} - Y_{ijT_2}, \quad R_{ij} = Y_{ijR_1} - Y_{ijR_2}, \quad i = 1, 2, \dots, s, \quad j = 1, 2, \dots, n_i,$$

$$Y_{ijT} = \frac{1}{2}(Y_{ijT_1} + Y_{ijT_2}), \quad Y_{ijR} = \frac{1}{2}(Y_{ijR_1} + Y_{ijR_2}), \quad \hat{\mu}_k = \frac{1}{s} \sum_{i=1}^s \bar{Y}_{ik}, \quad k = R, T$$

$$\bar{Y}_{ik} = \frac{1}{n_i} \sum_{j=1}^{n_i} \frac{1}{2} \sum_{l=1}^2 Y_{ijkl}, \quad \hat{\Delta} = \hat{\mu}_T - \hat{\mu}_R,$$

$$M_I = \hat{\sigma}_I^2 = \frac{1}{n_I} \sum_{i=1}^s \sum_{j=1}^{n_i} (I_{ij} - \bar{I}_i)^2, \quad M_T = \hat{\sigma}_{WT}^2 = \frac{1}{2n_T} \sum_{i=1}^s \sum_{j=1}^{n_i} (T_{ij} - \bar{T}_i)^2,$$

$$M_R = \hat{\sigma}_{WR}^2 = \frac{1}{2n_R} \sum_{i=1}^s \sum_{j=1}^{n_i} (R_{ij} - \bar{R}_i)^2, \quad n_I = n_T = n_R = \sum_{i=1}^s n_i - s.$$

$$\bar{I}_i = \frac{1}{n_i} \sum_{j=1}^{n_i} I_{ij}, \quad \bar{T}_i = \frac{1}{n_i} \sum_{j=1}^{n_i} T_{ij}, \quad \bar{R}_i = \frac{1}{n_i} \sum_{j=1}^{n_i} R_{ij}$$

Then the above linearized criteria are estimated by

3. reference-scaled ($M_R \geq \sigma_{W0}^2$):

$$\tilde{\eta}_1 = \hat{\Delta}^2 + M_I + 0.5M_T - (1.5 + \theta_1)M_R \quad (3.5)$$

4. constant-scaled ($M_R < \sigma_{W0}^2$):

$$\tilde{\eta}_2 = \hat{\Delta}^2 + M_I + 0.5M_T - 1.5M_R - \theta_1 \sigma_{W0}^2 \quad (3.6)$$

To evaluate IBE, compute the 95% upper bounds of both reference-scaled and constant scaled linearized criteria. If the upper bound of either criterion is negative or zero, we can draw a conclusion that the IBE is equivalent for T and R . To calculate the upper bound there are parametric method such as FDA [2] and nonparametric method such as FDA [1] and Shao [3]. On the basis of the mixed-model in FDA[2] we study IBE using bootstrap and Bayesian bootstrap method.

Download English Version:

<https://daneshyari.com/en/article/1889060>

Download Persian Version:

<https://daneshyari.com/article/1889060>

[Daneshyari.com](https://daneshyari.com)