



A Fast and Accurate Approximation to the Distributions of Quadratic Forms of Gaussian Variables

Hong Zhang, Judong Shen & Zheyang Wu

To cite this article: Hong Zhang, Judong Shen & Zheyang Wu (2022) A Fast and Accurate Approximation to the Distributions of Quadratic Forms of Gaussian Variables, Journal of Computational and Graphical Statistics, 31:1, 304-311, DOI: [10.1080/10618600.2021.2000423](https://doi.org/10.1080/10618600.2021.2000423)

To link to this article: <https://doi.org/10.1080/10618600.2021.2000423>



Published online: 01 Jan 2022.



Submit your article to this journal [↗](#)



Article views: 123



View related articles [↗](#)



View Crossmark data [↗](#)

SHORT TECHNICAL NOTE



A Fast and Accurate Approximation to the Distributions of Quadratic Forms of Gaussian Variables

Hong Zhang^a , Judong Shen^a, and Zheyang Wu^b

^aBiostatistics and Research Decision Sciences, MRL, Merck & Co., Inc., Kenilworth, NJ; ^bDepartment of Mathematical Sciences, Worcester Polytechnic Institute, Worcester, MA

ABSTRACT

In computational and applied statistics, it is of great interest to get fast and accurate calculation for the distributions of the quadratic forms of Gaussian random variables. This article presents a novel approximation strategy that contains two developments. First, we propose a fast numerical procedure in computing the moments of the quadratic forms. Second, we establish a general moment-matching framework for distribution approximation, which covers existing approximation methods for the distributions of the quadratic forms of Gaussian variables. Under this framework, a novel moment-ratio method (MR) is proposed to match the ratio of skewness and kurtosis based on the gamma distribution. Our extensive simulations show that (i) MR is almost as accurate as the exact distribution calculation and is much faster; (ii) comparing with existing approximation methods, MR significantly improves the accuracy of approximating far right tail probabilities. The proposed method has wide applications. For example, it is a better choice than existing methods for facilitating hypothesis testing in big data analysis, where fast and accurate calculation of very small p -values are desired. An R package *Qapprox* that implements related methods is available on CRAN.

ARTICLE HISTORY

Received September 2020
Revised July 2021

KEYWORDS

Chi-square; Gamma distribution; Gaussian quadratic forms; Kurtosis; Mixture of Type I error; Moment-matching method

1. Introduction

Quadratic forms of Gaussian variables appear in many statistical applications such as hypothesis testing and statistical power calculation. For instance, recent developments in statistical genetics often adopt the quadratic forms as test statistics to detect associations between genetic variants and complex phenotypes. Well-known examples include, to name a few, the kernel machine based score statistic for genetic pathway analysis (Liu, Lin, and Ghosh 2007), the sequence kernel association test statistic (SKAT) (Wu et al. 2011) and the sum of powered score test statistic (SPU) (Pan et al. 2014) for rare-variant analysis. These new applications of analyzing big biological data require accurate calculation of small p -values. This requirement poses challenges to the calculation of distribution at the far-right tail. For example, in a typical gene-based whole-genome association study of about 20,000 human genes, the significance threshold under Bonferroni correction is at the level of $\alpha = 0.05/20000 = 2.5 \times 10^{-6}$.

In this article, we study computationally efficient and accurate algorithms to calculate the distributions of the quadratic forms of Gaussian variables. Specifically, for a vector of d Gaussian variables, $X \sim \mathcal{N}(\mu, \Sigma)$ with mean vector μ of real numbers and a positive-definite covariance matrix Σ , its quadratic form is

$$Q = X'AX, \quad (1)$$

where A is a $d \times d$ positive semi-definite matrix. The statistic Q is often used in the hypothesis testing of the mean μ . An exceedingly large Q is the evidence against the hypothesis that $\mu = 0$. We are interested in calculating the right tail probability of Q , or, equivalently the p -value

$$P(Q > q), \quad (2)$$

where $q \geq 0$ is an observed value of Q .

First, note that the distribution of Q can be calculated in an exact way. Consider a “decorrelation” of X by the inverse of the Cholesky of Σ , $U = (\Sigma^{1/2})^{-1}X \sim \mathcal{N}((\Sigma^{1/2})^{-1}\mu, I)$. The existence of $(\Sigma^{1/2})^{-1}$ is guaranteed by the positive definiteness of Σ . We can write $Q = U'MU$, where $M = (\Sigma^{1/2})'A\Sigma^{1/2}$. Denote Λ the diagonal matrix of eigenvalues of M and P the orthonormal matrix such that $M = P\Lambda P'$, we have

$$Q \stackrel{d}{=} (Z + \tilde{\mu})'\Lambda(Z + \tilde{\mu}), \quad (3)$$

where Z is a vector of d independent standard normal variables and $\tilde{\mu} = P'(\Sigma^{1/2})^{-1}\mu$. The formula (3) indicates that the distribution of Q is the same as a weighted sum of independent, potentially noncentral, chi-squared random variables. Thus, the cumulative distribution function (CDF) of Q can be obtained by inverting its characteristic function numerically (Imhof 1961; Davies 1980). We call this calculation approach the exact method since it is in theory accurate except for truncation

errors that can be bounded in numerical procedures. However, the exact approach is time-consuming for two reasons: (i) it requires eigendecomposition with a computation cost in the order of $O(d^3)$, which is heavy when d is large in big data analysis; (ii) the algorithm's speed is sensitive to the number of positive eigenvalues and the point at which the CDF is to be evaluated (Davies 1980).

To speed up the computation, various approximation methods have been proposed based on the moments of Q . In particular, the Satterthwaite–Welch method (SW) (Welch 1938; Satterthwaite 1946) matches the first two moments of Q with a gamma variable. The Hall–Buckley–Eagleson approximation (HBE) (Hall 1983; Buckley and Eagleson 1988) matches the skewness of Q with a chi-squared variable while adjusting for the mean and the variance. Wood's F approximation (Wood) (Wood 1989) matches the first three moments with a three-parameter F variable. Liu–Tang–Zhang method (LTZ) (Liu, Tang, and Zhang 2009) tries to match both skewness and kurtosis of Q with a non-central chi-squared variable while adjusting for the mean and the variance at the same time. A more comprehensive literature review can be found in (Bodenham and Adams 2016). In the applications to hypothesis testing of big data, these methods lack desired accuracy for controlling the Type I error rate at small α levels (see Figure 3). Furthermore, to obtain the moments of Q , these methods typically use eigenvalues of $A\Sigma$ or the trace of $(A\Sigma)^k$, $k = 1, 2, \dots$, which can be computationally intensive when d is large.

To address these two issues, we first propose a numerical procedure that calculates the moments of Q faster in Section 2.1. Then, in Section 2.2, we describe a general moments-matching framework, which allows flexibly incorporating the information of high moments to improve accuracy without increasing computational difficulties. Within this framework, we propose a novel approximation method in Section 2.3, that significantly improves the accuracy in computing the distribution of Q , especially for the far right tail probability. The computation time and Type I error comparison results from extensive simulation study are presented in Section 3. We conclude this article with final remarks in Section 4.

2. Methods

2.1. Computing the Moments of Q

To approximate the distribution of Q , the first step is to calculate its moments. This step takes the dominate computational time in the approximation methods. We propose a strategy to speed up this process, which we did not see in the relevant Q distribution approximation literature yet but can be applied into these methods to improve computational efficiency.

For a general quadratic form Q , the k th cumulant, c_k , is given by (Johnson, Kotz, and Balakrishnan 1995),

$$c_k = 2^{k-1}(k-1)! \left(\text{tr} \left(\Lambda^k \right) + k \tilde{\mu}' \Lambda^k \tilde{\mu} \right). \quad (4)$$

Subsequently, the mean μ_Q , variance σ_Q^2 , skewness γ_Q and kurtosis κ_Q can be calculated through c_k , $k = 1, 2, 3, 4$:

$$\mu_Q = c_1, \quad \sigma_Q^2 = c_2, \quad \gamma_Q = \frac{c_3}{c_2^{3/2}}, \quad \kappa_Q = \frac{c_4}{c_2^2} + 3, \quad (5)$$

The formula in (4) requires calculating the eigenvalues, which has a computation cost in the order of $O(n^3)$ in practice. To improve the computation, Liu, Tang, and Zhang (2009) proposed a “trace” approach by rewriting (4) to be

$$c_k = 2^{k-1}(k-1)! \left(\text{tr} \left((A\Sigma)^k \right) + k \mu' (A\Sigma)^{k-1} A \mu \right). \quad (6)$$

Although formula (6) does not need eigenvalues explicitly, the naive implementation of this approach, assuming $A\Sigma$ is known in advance, will require $k-1$ matrix multiplication. This approach could still be computationally expensive especially when k is large. For example, it requires at least three matrix multiplications to get c_1 to c_4 , which is not much faster than the eigenvalue approach (formula (4)) as evidenced by Figure 1.

We can improve the computational efficiency by not explicitly computing every matrix multiplication (see Lemma 2.1). The improvement is significant when the Gaussian variables are centered, that is, $\mu \equiv 0$.

Lemma 2.1. Let B be a $d \times d$ matrix. B^0 is defined as a $d \times d$ identity matrix. If B^{k-1} and B^k , $k \geq 1$, are known, then we can compute $\text{tr}(B^{2k-1})$ and $\text{tr}(B^{2k})$ in $O(d^2)$ time.

Proof. Through matrix multiplication, we have $\text{tr}(B^{2k-1}) = \sum_{i=1}^d \sum_{j=1}^d (B^{k-1})_{ij} (B^k)_{ji}$ and $\text{tr}(B^{2k}) = \sum_{i=1}^d \sum_{j=1}^d (B^k)_{ij} (B^k)_{ji}$. $\square$

Lemma 2.1 says that we can compute the trace of B^k by roughly $k/2$ (instead of $k-1$) matrix multiplications. In particular, it only takes a single matrix multiplication to get the first four moments. Thus, by using Lemma 2.1, the computation cost of getting the first four moments is about one-third of the naive trace approach. If only mean and variance are needed, as in the SW method, the computational efficiency is further improved since no matrix multiplication is needed by Lemma 2.1. As will be shown later, SW is not very accurate for very small p -values. However, it is still adequate when the significance level is not too stringent, for example, $\alpha = 0.05$. Because of its fast speed, it might be useful as a screening method to speed up the overall data analysis. For example, in the genome-wide association studies most genes are not trait associated and thus their p -values are expected to be relatively large. The SW can be used to remove these genes before applying more accurate but computationally more intensive methods to calculate small p -values.

2.2. A General Framework for Moment-Matching Method

Here, we describe a general moment-matching framework that covers the literature approximation methods in Introduction. Note that the distribution of a quadratic form can be entirely characterized by its moments according to the Cramer's condition (Lin 2017). Ideally, engaging more moments and more parameters could allow more flexible distribution model and thus provide more accuracy in general. However, the resulting matching equations would also be more difficult to solve, and sometimes a solution does not exist. To ease this tradeoff, our general moment-matching framework can engage more moments without adding too much computational difficulty.

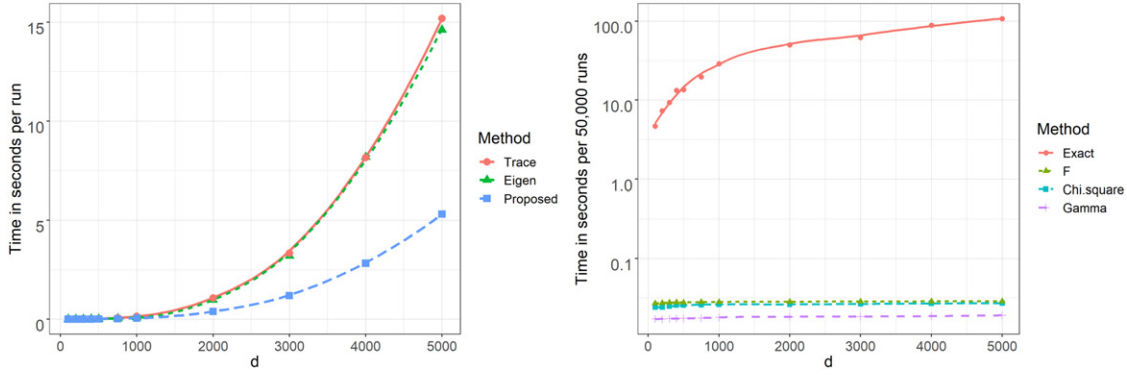


Figure 1. Computation time comparison. Left: eigenvalues/moments computation comparison. Trace: naive trace method (formula (5)). Eigen: eigenvalue decomposition (formula (4)). Proposed: the proposed moment computation method Lemma 2.1. Right: distribution computation comparison per 50,000 runs (average of 10 repetitions). Exact: Davies method. F: F distribution. Chi.square: Chi-squared distribution. Gamma: Gamma distribution. d is the dimension of the correlation matrix.

Specifically, let Y be a random variable with known CDF $F_Y(y|\theta)$, where $\theta = (\theta_1, \dots, \theta_m)$ denotes m distribution parameters. In this framework, we match the moments of Q with the moments of a linear transformation of Y :

$$T = aY + b. \quad (7)$$

Note that the standardized moments of T are always equal to those of Y , which are functions of θ but not of a and b . That is, for $k \geq 1$,

$$\tilde{\mu}_{k,Y}(\theta) := E\left(\frac{Y - \mu_Y}{\sigma_Y}\right)^k = E\left(\frac{T - \mu_T}{\sigma_T}\right)^k =: \tilde{\mu}_{k,T}(\theta).$$

The advantage of matching Q with T , instead of directly matching Q with Y , is that it has a potential to gain two “free” parameters a and b corresponding to the mean and variance. In other words, the means and variances of Q and T are automatically matched as long as θ is available. Therefore, θ could be determined by matching higher ($k \geq 3$) standardized moments (e.g., skewness $\gamma = \tilde{\mu}_3$ and kurtosis $\kappa = \tilde{\mu}_4$), which provides extra flexibility of distribution especially at the tail.

Specifically, by matching the mean and variance of Q and T , we have

$$a^* = \frac{\sigma_Q}{\sigma_Y(\theta)}, \quad b^* = \mu_Q - \frac{\sigma_Q}{\sigma_Y(\theta)} \mu_Y(\theta). \quad (8)$$

Now, θ could be determined by solving proper equations regarding the standardized higher ($k \geq 3$) moments of Q and Y . We denote such equations in the following generic way:

$$g_j(\tilde{\mu}_{k,Y}(\theta), \tilde{\mu}_{k,Q}(\theta); k = 3, 4, \dots) = 0, \quad j = 1, 2, \dots, m, \quad (9)$$

where the g_j functions of these standardized higher moments need to be designed properly in order to get solutions θ^* . Finally, with the determined a^* , b^* and θ^* the right-tail probability of Q is approximated by

$$\begin{aligned} P(Q > q) &\approx P(T > q) = P\left(Y > \frac{q - b^*}{a^*}\right) \\ &= 1 - F_Y\left(\frac{\sigma_Y}{\sigma_Q}(q - \mu_Q) + \mu_Y|\theta^*\right). \end{aligned} \quad (10)$$

The existing approximation methods for Q can fit into this general framework in (7) – (10). When choosing Y a gamma random variable $G(\alpha, 1)$ with shape parameter α and scale

parameter 1, the Satterthwaite–Welch method is equivalent to solving $\theta = \alpha$ by fixing $b^* = 0$ in (8) without engaging higher moment information in (9). Hall–Buckley–Eagleson method also uses the gamma variable; it is equivalent to solving α by matching the skewness in (9). Choosing Y a noncentral chi-squared variable $\chi_d^2(\delta)$ with degrees of freedom d and noncentrality δ , Liu–Tang–Zhang method is equivalent to solving $\theta = (d, \delta)$ by matching the skewness and kurtosis in (9). Letting Y follow an F distribution with degrees of freedom d_1 and d_2 , Wood’s F approximation is equivalent to solving $\theta = (d_1, d_2)$ by fixing $b^* = 0$ in (8) and matching the skewness in (9).

2.3. Strategies for Matching Higher Moments

Solving θ by matching standardized higher moments in (9) is not trivial. Depending on the distribution of Y , due to restrictions on the domains of these moments and their interdependence, the solution to equations in (9) often do not exist for many seemingly obvious choices of the g functions. For example, one design for (9) is to match skewness and kurtosis at the same time, such as described in the Liu–Tang–Zhang method (Liu, Tang, and Zhang 2009). The authors provided a nice result on the necessary and sufficient condition, that is, $\frac{3}{2}c_3^2 > c_4c_2$, for the solution of $\theta = (d, \delta)$ exists. For centered Gaussian variables, the condition becomes $(\sum_{i=1}^n \lambda_i^3)^2 > (\sum_{i=1}^n \lambda_i^2)(\sum_{i=1}^n \lambda_i^4)$. In this case, however, a general inequality given in Lemma 2.2 indicates that such condition will never be met. This is somewhat surprising because Q follows a chi-square distribution when $\Sigma = I$. However, when $\Sigma \neq I$, it is impossible to find a chi-squared distribution that matches Q ’s skewness and kurtosis by adding a noncentrality parameter. Even for noncentral Gaussian variables, the solution still often does not exist (see Duchesne and Lafaye De Micheaux 2010; Bodenham and Adams 2016 and the simulations in Section 3). When the solution does not exist, Liu–Tang–Zhang method steps back to matching only the skewness, which is exactly the Hall–Buckley–Eagleson method. We can also modify the method to matching the kurtosis instead, which could show some improved results in some scenarios. However, in either way it will lose the information of the unmatched moment.

Lemma 2.2 (Littlewood's inequality). Define $S_a = \sum_{i=1}^n \lambda_i^a$, $a > 0$, $\lambda_i \geq 0$, $i = 1, \dots, n$. Suppose $0 < a < b < c$, then $S_b^{c-a} \leq S_a^{c-b} S_c^{b-a}$. The equality holds if and only if $\lambda_1 = \lambda_2 = \dots = \lambda_n$.

Proof. The lemma follows a classic inequality that can be found in literature (cf. Hardy, Littlewood, and Polya (1934), p.28, Th. 18). Here we show it is equivalent to the Hölder's inequality.

Write $b = ka + (1-k)c$, $0 < k < 1$. Then $c - b = k(c - a)$, $b - a = (1-k)(c - a)$. The inequality is equivalent to

$$S_b^{c-a} \leq S_a^{k(c-a)} S_c^{(1-k)(c-a)} \iff \sum_{i=1}^n \lambda_i^{ka} \lambda_i^{(1-k)c} \leq \left(\sum_{i=1}^n \lambda_i^a \right)^k \left(\sum_{i=1}^n \lambda_i^c \right)^{1-k}.$$

Define $p = 1/k$, $q = 1/(1-k)$, $x_i = \lambda_i^{a/p}$, $y_i = \lambda_i^{c/q}$, the above inequality becomes

$$\sum_{i=1}^n x_i y_i \leq \left(\sum_{i=1}^n x_i^p \right)^{1/p} \left(\sum_{i=1}^n y_i^q \right)^{1/q},$$

which is the Hölder's inequality. The equality holds if and only if x_i and y_i , $i = 1, \dots, n$, are proportional which is equivalent to λ_i , $i = 1, \dots, n$, are equal to each other. $\square$

In order to overcome the difficulty of matching both skewness and kurtosis, while at the same time still keeping the information from both, we propose the following solution. First, we choose the gamma distribution model, that is, $Y \sim G(\alpha, 1)$. Here we fix the gamma scale parameter to be 1 because it is redundant with the coefficient parameter a in (7) and thus does not affect the calculation in (10). Second, we propose two types of g functions given below for Equation (9), which incorporate the information of both skewness and kurtosis. We choose gamma distribution because it is one of the extensions of chi-square distribution such that it becomes exact when the mean vector $\mu \equiv 0$ and covariance matrix $\Sigma = I$. More importantly, comparing with the noncentral chi-squared distribution, the solution of our matching equation using gamma distribution always exists. At the same time, it can be more accurate in approximating the right tail probability in (2) at relatively large threshold q .

Specifically, the first method is called the moment-ratio matching (MR) method. With one parameter α , we match the ratio between skewness and kurtosis. Define

$$g_1(\tilde{\mu}_{3,Y}, \tilde{\mu}_{3,Q}, \tilde{\mu}_{4,Y}, \tilde{\mu}_{4,Q}) = \frac{\tilde{\mu}_{3,Y}}{\tilde{\mu}_{4,Y} - 3} - \frac{\tilde{\mu}_{3,Q}}{\tilde{\mu}_{4,Q} - 3}.$$

The solution of Equation (9) always exists and has a simple closed form: $\alpha^* = \frac{9\tilde{\mu}_{3,Q}^2}{(\tilde{\mu}_{4,Q} - 3)^2}$.

The second method is called the minimized matching error (ME) method. That is, we choose α to minimize the Euclidean distance between $(\tilde{\mu}_{3,Y}, \tilde{\mu}_{4,Y})$ and $(\tilde{\mu}_{3,Q}, \tilde{\mu}_{4,Q})$. For that, the corresponding g function is the first-order derivative of the distance as a function of α ,

$$\begin{aligned} g_1(\tilde{\mu}_{3,Y}, \tilde{\mu}_{3,Q}, \tilde{\mu}_{4,Y}, \tilde{\mu}_{4,Q}) \\ = \frac{\partial}{\partial \alpha} [(\tilde{\mu}_{3,Y} - \tilde{\mu}_{3,Q})^2 + (\tilde{\mu}_{4,Y} - \tilde{\mu}_{4,Q})^2]. \end{aligned} \quad (11)$$

Accordingly, Equation (9) becomes

$$\tilde{\mu}_{3,Q} \alpha^{3/2} - 2(10 - 3\tilde{\mu}_{4,Q})\alpha - 36 = 0. \quad (12)$$

Since $\tilde{\mu}_{3,Q} > 0$, it is straightforward to show that there exists one and only one real root $\alpha^* > 0$.

3. Simulation Results

3.1. Computation Time Comparison

In this section, we compare the computation time of the proposed moment-ratio (MR) matching method with other existing methods. The comparison is two-fold. First, we examine the computation efficiency of different methods for obtaining the eigenvalues or the moment estimates of Q . Second, given the eigenvalues and moment estimates, we compare various methods for calculating the final right-tail probability based on different distributions.

We first consider the correlation matrix Σ to be a polynomial decaying matrix, that is, $(\Sigma)_{ij} = 1/|i - j|$, $1 \leq i, j \leq n$, $d = 100, 200, \dots, 5000$. The eigenvalue decomposition (formula (4)), the naive trace method (formula (5)) and the proposed moment computation method (Lemma 2.1) were performed 1000 times and the average run times were summarized in the left panel of Figure 1. Next, for each d , we simulate 50,000 quadratic form of centered Gaussian variables, $Q = X'X$, with correlation matrix Σ . The exact Davies method, the LTZ method (based on chi-squared distribution), the HBE method and the proposed MR method (gamma distribution) and Wood method (F distribution) were used to calculate the p -values assuming the eigenvalues of Σ are known. The total runtimes were recorded and summarized in the right panel of Figure 1. For the second comparison, the Exact and LTZ methods were implemented by a widely used R package *CompQuadForm*. All other methods were implemented by the authors. All computations were completed on an Intel Core-i5 8350U processor at 1.70 GHz with 16 GB of RAM, running Microsoft R Open 3.5.0.

As Figure 1 shows, the computation time of the trace approach and the eigenvalue approach is about the same but the proposed moment calculation method can save the computation time by 1/2 to 2/3. Compared to the exact calculation by Davies method, the approximation methods for calculating the final right-tail probability (based on chi-square or gamma or F distribution) could be hundreds of times faster. Moreover, unlike Davies method, the computation costs of the approximation methods almost do not increase as the dimension d increases. This is because, given the eigenvalues, the CDF of approximation distribution is readily available through method of moments while the Davies method needs the input of d eigenvalues for every evaluation. Meanwhile, we also note that the overall computation time is dominated by the calculation of the eigenvalues/moments.

3.2. Type I Error Comparison

In this section we systematically evaluate the accuracies of the proposed and existing Q distribution approximation methods in the context of hypothesis testing. Specifically, consider Q as the test statistic and q as its observed value. The right-tail probability

$P(Q > q)$ can be regarded as the p -value for testing the null hypothesis $H_0 : \mu_X = \mu$. In the simulation study, we fix $A = I$ and vary the correlation matrix Σ (since A can be considered as a factor that varies the correlation matrix Σ). For a given mean vector μ and correlation matrix Σ , 2×10^7 of random Gaussian vectors $X \sim \mathcal{N}(\mu, \Sigma)$ were generated, each led to an observed q . Consequently, each approximation method was applied to generate 2×10^7 p -values. The empirical Type I error rate for a given approximation method is defined as the proportion of rejections under H_0 , that is, the proportion of p -values less than a nominal level α . A good approximation method would have the empirical Type I error rate being close to the nominal α . If a ratio between the empirical Type I error rate and the nominal α is larger than 1, it indicates the approximation method leads to a liberal test (i.e., rejecting more than appropriate); a ratio less than 1 indicates a conservative test.

Motivated by genetic data, we examined a wide variety of covariance patterns that are potentially with block structures. Consider the $d \times d$ covariance matrix

$$\Sigma = \begin{bmatrix} \Sigma_{11} & \Sigma_{12} \\ \Sigma_{12}' & \Sigma_{22} \end{bmatrix},$$

where each of the blocks is of size $(d/2) \times (d/2)$. Equal correlation matrix E_k and polynomially decaying correlation matrix D_k were used to define Σ or its blocks:

$$E_k(\rho) : E_k(i, j) = \rho, \quad 1 \leq i \neq j \leq k \text{ and } 0 \leq \rho < 1, \quad (13)$$

$$D_k(\phi) : D_k(i, j) = 1/|i - j|^\phi, \quad 1 \leq i \neq j \leq k \text{ and } \phi > 0. \quad (14)$$

For example, we can define entire $\Sigma = E_d$ or D_d . Alternatively, we may define Σ to be block-diagonal, for example, $\Sigma_{11} = \Sigma_{22} = E_{d/2}$ or $D_{d/2}$. Table 1 summarizes a total of 12 correlation types for Σ . Multiple values of the matrix parameters $\rho = 0.9, 0.5, 0.1$ and $\phi = 0.2, 1, 3$ were considered to model the strong, moderate and weak correlations, respectively. The dimensions $d = 10, 50, 100$ and 500 were considered. We also tested three configurations of the mean vector: a central

Table 1. The correlation matrix Σ used in the simulations based on the correlation models in (13) and (14).

Type	I	II	III
Equal(ρ)	$\Sigma_{11} = E_{d/2}(\rho)$	$\Sigma_{11} = \Sigma_{22} = E_{d/2}(\rho)$	$\Sigma = E_d(\rho)$
Poly(ϕ)	$\Sigma_{11} = D_{d/2}(\phi)$	$\Sigma_{11} = \Sigma_{22} = D_{d/2}(\phi)$	$\Sigma = D_d(\phi)$
Inv-Equal(ρ) [*]	$\Sigma_{11} = E_{d/2}^{-1}(\rho)$	$\Sigma_{11} = \Sigma_{22} = E_{d/2}^{-1}(\rho)$	$\Sigma = E_d^{-1}(\rho)$
Inv-Poly(ϕ) [*]	$\Sigma_{11} = D_{d/2}^{-1}(\phi)$	$\Sigma_{11} = \Sigma_{22} = D_{d/2}^{-1}(\phi)$	$\Sigma = D_d^{-1}(\phi)$

^{*} Σ is standardized to become a correlation matrix.

NOTE: Three types: I (upper left correlations with $\Sigma_{22} = I, \Sigma_{12} = 0$); II (block-wise correlations); III (the whole correlation matrix follows the models).

Gaussian case $\mu = \mu_1 \equiv 0$, and two non-central Gaussian cases $\mu = \mu_2 \equiv 1$ and $\mu = \mu_3 = (\underbrace{1, \dots, 1}_{d/2}, \underbrace{0, \dots, 0}_{d/2})$. Our simulation

settings are pretty extensive; altogether there are 432 different combinations of the mean μ and correlation matrix Σ for each given d .

The approximation methods to be compared are (i) MR: the proposed MR matching method; (ii) ME: the proposed minimized matching error method; (iii) SW: Satterthwaite–Welch method; (iv) HBE: Hall–Buckley–Eagleson method; (v) Wood: Wood’s F approximation; (vi) LTZ: Liu–Tang–Zhang method; (vii) LTZ4: modified Liu–Tang–Zhang method (i.e., when skewness and kurtosis could not be matched simultaneously, we chose to match kurtosis instead of skewness); and (viii) Exact: the Davies method that inverts the characteristic function. We used R package *CompQuadForm* for the LTZ method and the Exact method. Other methods were implemented by the authors.

The boxplots in Figure 2 show the comparison of all eight methods at relatively larger nominal α levels, $\alpha = 0.05$ and 0.01 . Each box consists of multiple empirical Type I errors under different combinations of the correlation matrix Σ and mean vector μ . It is clear that there exist many scenarios in which the Wood method is ultra-conservative while the SW method could inflate the Type I errors at $\alpha = 0.01$. The other approximation methods seem to perform well as their empirical Type I error

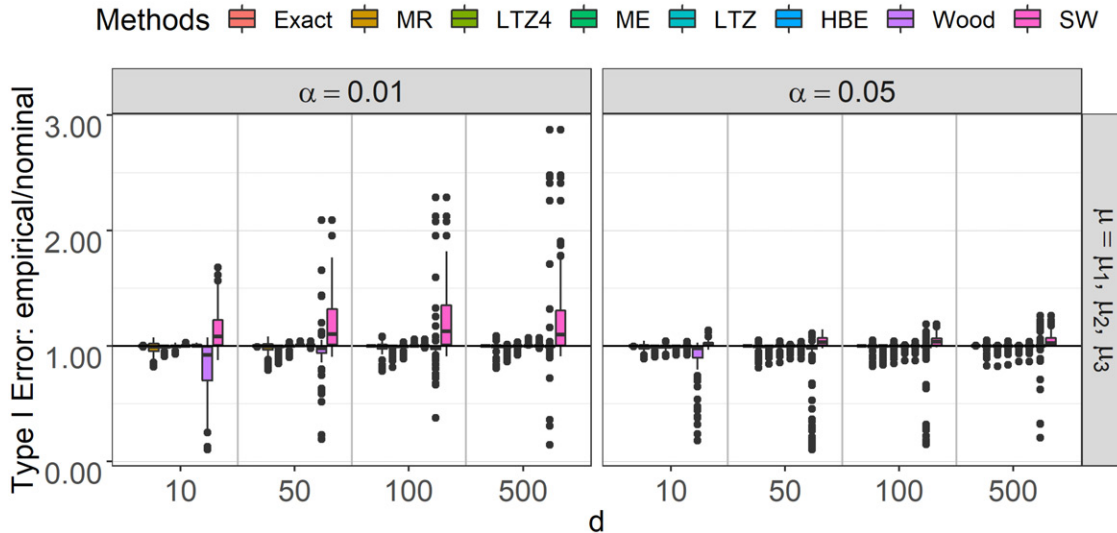


Figure 2. Boxplots of the ratios between empirical Type I error rates and the nominal levels, $\alpha = 0.01$ and 0.05 . Exact: the Davies method. MR: moment-ratio matching method. LTZ4: modified Liu–Tang–Zhang method. ME: minimized matching error method. LTZ: Liu–Tang–Zhang method. HBE: Hall–Buckley–Eagleson method. Wood: Wood’s F approximation. SW: Satterthwaite–Welch method.

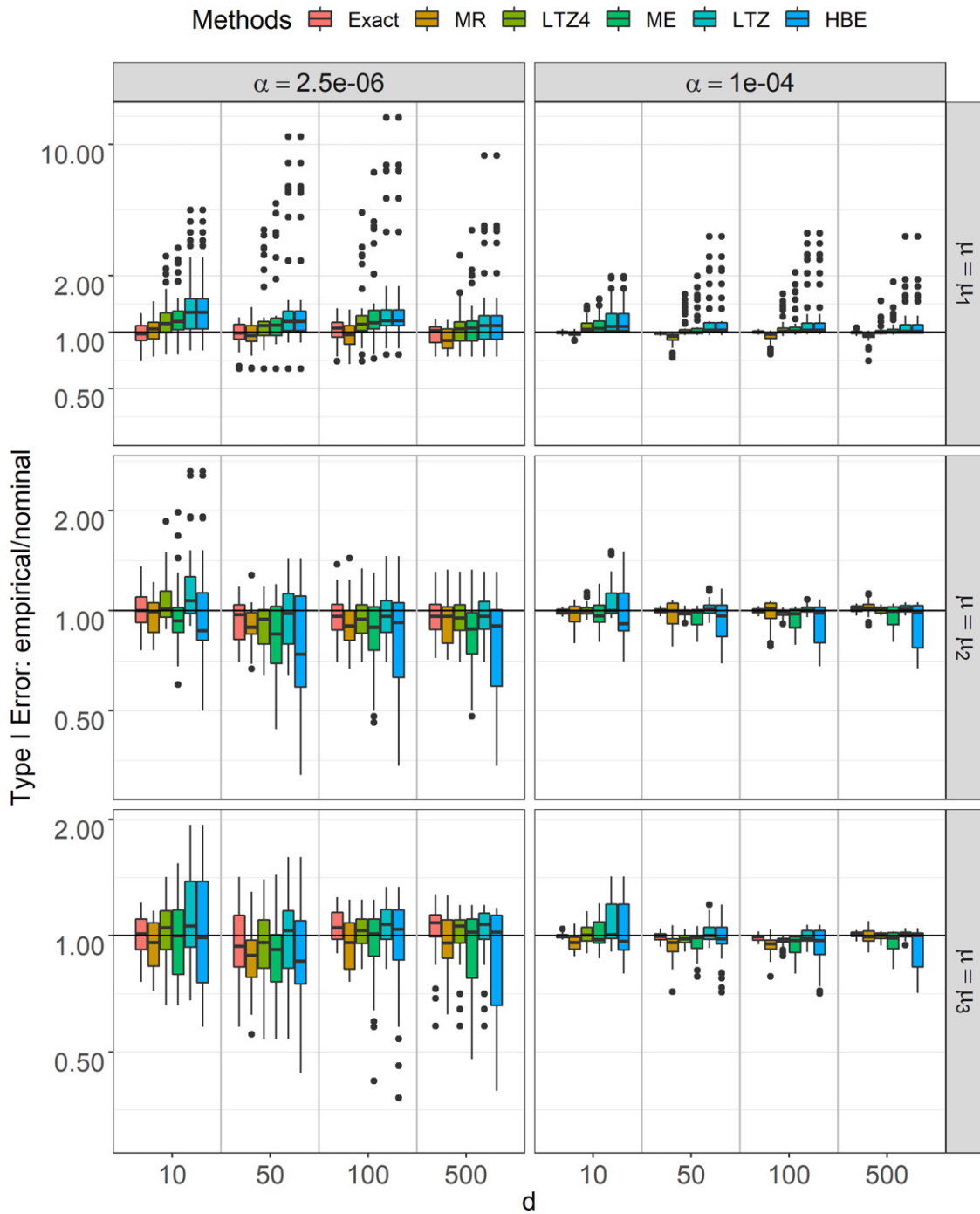


Figure 3. Boxplots of the ratios between empirical Type I error rates and the nominal levels, $\alpha = 2.5 \times 10^{-6}$ and 1×10^{-4} , stratified by the mean vector $\mu_1 \equiv 0$, $\mu_2 \equiv 1$ and $\mu_3 = (1, \dots, 1, 0, \dots, 0)$. Exact: the Davies method. MR: moment-ratio matching method. LTZ4: modified Liu-Tang-Zhang method. ME: minimized matching error method. LTZ: Liu-Tang-Zhang method. HBE: Hall-Buckley-Eagleson method.

rates are within $(0.8, 1.1)$ of the nominal levels. Unsurprisingly, the Exact method is the most accurate method with the least amount of variation around the nominal levels.

Next, we consider smaller nominal levels, $\alpha = 1 \times 10^{-4}$ and 2.5×10^{-6} . The comparison results are summarized in Figure 3 as boxplots stratified by the mean vector μ . Each box represents empirical Type I errors under different correlation matrices. The SW and Wood methods were dropped because their empirical Type I errors were further away from α . When $\mu = \mu_1 \equiv 0$ (the central Gaussian case), as expected, the exact method was accurate at both levels. The proposed moment

ratio matching method (MR) also controlled the Type I error rates very well. It only varied slightly more than the exact method when $\alpha = 2.5 \times 10^{-6}$ and is slightly conservative when $\alpha = 10^{-4}$. The minimized matching error method (ME) was slightly more inflated compared to MR. The Liu-Tang-Zhang (LTZ) method performed the same as the Hall-Buckley-Eagleson (HBE) method because it cannot match both skewness and kurtosis when $\mu \equiv 0$, as proven by Lemma 2.2. Therefore, in this case it reduced to matching the skewness only. Both methods seemed to inflate the Type I errors in multiple occasions across the dimension d . The modified Liu-Tang-Zhang (LTZ4)

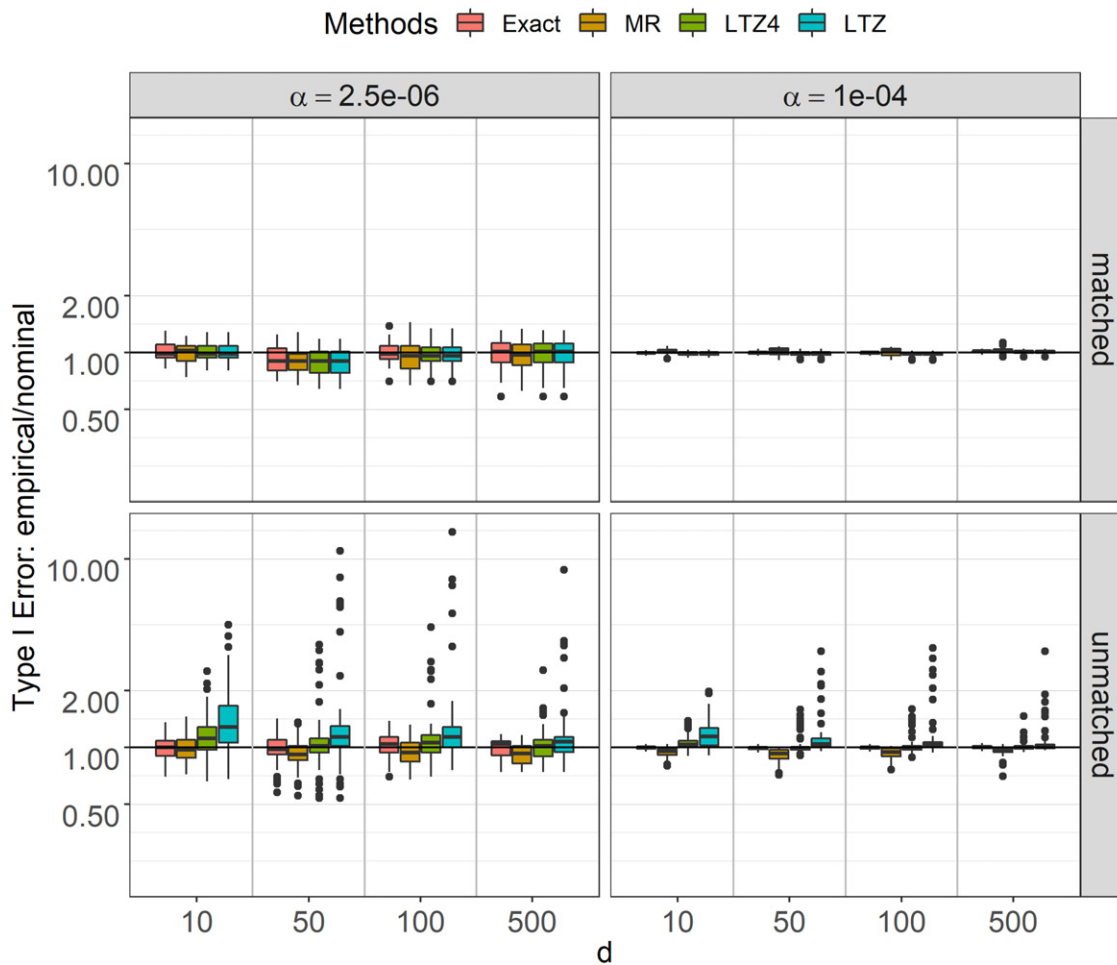


Figure 4. Boxplots of the ratios between empirical Type I error rates and the nominal levels, $\alpha = 2.5 \times 10^{-6}$ and 1×10^{-4} , stratified by whether skewness and kurtosis were matched in the LTZ method. Exact: the Davies method. MR: moment-ratio matching method. LTZ4: modified Liu-Tang-Zhang method. LTZ: Liu-Tang-Zhang method.

method seemed to perform better than the original version but still with considerable amount of inflations. When $\mu = \mu_2$ or μ_3 (the noncentral Gaussian case), the inflation of LTZ's and LTZ4's Type I errors still existed but were greatly reduced because now in various scenarios both skewness and kurtosis could be matched. The proposed MR method still controlled Type I errors well with small variation around α .

To further demonstrate the impact of unmatched moments on the Type I errors, we present the boxplots stratified by whether the skewness and kurtosis can be matched ("unmatched" vs. "matched") in the LTZ method. Figure 4 shows that when these two moments were matched, the performance of LTZ method is on par with the MR method and the exact method. However, if they are not matched, choosing either skewness or kurtosis would inflate the Type I errors while the MR method would not.

4. Conclusion

We propose an MR matching method that improves the accuracy of approximating the right-tail probability of quadratic forms of Gaussian variables. The proposed method controls the Type I error rates almost as well as the exact method while improving the computational efficiency by 1/2 to 2/3. This

method is expected to have broad applications in hypothesis testing, such as genetic association studies where Gaussian quadratic forms are frequently used as statistics to combine signals from multiple sources.

Beyond the Gaussian quadratic forms, we believe the idea of using higher-order moments without exactly matching them, as described in Section 2.2 and 2.3, is applicable to more general distribution approximation problems where matching both skewness and kurtosis is difficult or even impossible. Our results demonstrate that a wise choice of matching equation that incorporates both skewness and kurtosis, such as the MR matching, can obtain accurate approximation like both moments are matched.

Disclaimer

Hong Zhang and Judong Shen are employees of Merck Sharp & Dohme Corp., a subsidiary of Merck & Co., Inc., Kenilworth, NJ, USA and shareholders in Merck & Co., Inc., Kenilworth, NJ, USA.

ORCID

Hong Zhang  <http://orcid.org/0000-0002-8869-8671>

References

- Bodenham, D. A., and Adams, N. M. (2016), "A Comparison of Efficient Approximations for a Weighted Sum of Chi-Squared Random Variables," *Statistics and Computing*, 26, 917–928. [305,306]
- Buckley, M. J., and Eagleson, G. (1988), "An Approximation to the Distribution of Quadratic Forms in Normal Random Variables," *Australian Journal of Statistics*, 30, 150–159. [305]
- Davies, R. B. (1980), "Algorithm as 155: The Distribution of a Linear Combination of χ^2 Random Variables," *Journal of the Royal Statistical Society, Series C*, 29, 323–333. [304,305]
- Duchesne, P., and Lafaye De Micheaux, P. (2010), "Computing the Distribution of Quadratic Forms: Further Comparisons Between the Liu–Tang–Zhang Approximation and Exact Methods," *Computational Statistics & Data Analysis*, 54, 858–862. [306]
- Hall, P. (1983), "Chi Squared Approximations to the Distribution of a Sum of Independent Random Variables," *The Annals of Probability*, 1028–1036. [305]
- Hardy, G., Littlewood, J., and Polya, G. (1934), *Inequalities*, Cambridge: University Press. [307]
- Imhof, J.-P. (1961), "Computing the Distribution of Quadratic Forms in Normal Variables," *Biometrika*, 48, 419–426. [304]
- Johnson, N. L., Kotz, S., and Balakrishnan, N. (1995), *Continuous Univariate Distributions*, Wiley. [305]
- Lin, G. D. (2017), "Recent Developments on the Moment Problem," *Journal of Statistical Distributions and Applications*, 4, 1–17. [305]
- Liu, D., Lin, X., and Ghosh, D. (2007), "Semiparametric Regression of Multidimensional Genetic Pathway Data: Least-Squares Kernel Machines and Linear Mixed Models," *Biometrics*, 63, 1079–1088. [304]
- Liu, H., Tang, Y., and Zhang, H. H. (2009), "A New Chi-Square Approximation to the Distribution of Non-Negative Definite Quadratic Forms in Non-Central Normal Variables," *Computational Statistics & Data Analysis*, 53, 853–856. [305,306]
- Pan, W., Kim, J., Zhang, Y., Shen, X., and Wei, P. (2014), "A Powerful and Adaptive Association Test for Rare Variants," *Genetics*, 197, 1081–1095. [304]
- Satterthwaite, F. E. (1946), "An Approximate Distribution of Estimates of Variance Components," *Biometrics Bulletin*, 2, 110–114. [305]
- Welch, B. L. (1938), "The Significance of the Difference Between Two Means When the Population Variances Are Unequal," *Biometrika*, 29, 350–362. [305]
- Wood, A. T. (1989), "An f Approximation to the Distribution of a Linear Combination of Chi-Squared Variables," *Communications in Statistics-Simulation and Computation*, 18, 1439–1456. [305]
- Wu, M. C., Lee, S., Cai, T., Li, Y., Boehnke, M., and Lin, X. (2011), "Rare-Variant Association Testing for Sequencing Data With the Sequence Kernel Association Test," *The American Journal of Human Genetics*, 89, 82–93. [304]