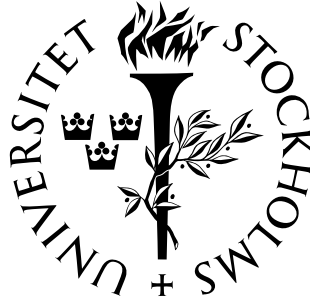




Mathematical Statistics
Stockholm University

**Finite Sample Size Bounds on the Variance
Estimator in Non-Gaussian General Linear
Models**

Mathias Lindholm and Felix Wahl

Research Report 2018:2

ISSN 1650-0377

Postal address:

Mathematical Statistics
Dept. of Mathematics
Stockholm University
SE-106 91 Stockholm
Sweden

Internet:

<http://www.math.su.se>



Finite Sample Size Bounds on the Variance Estimator in Non-Gaussian General Linear Models

Mathias Lindholm and Felix Wahl

January 2018

Abstract

We consider bounds on the variance of the standard unbiased variance estimator in a general non-Gaussian linear model for finite sample sizes. In particular we obtain bounds that are sharp in the sense that both the lower and upper bound will converge to the same asymptotic limit when scaled with the sample size. Further, these bounds are independent of covariate information. Due to this we may also obtain unconditional variance bounds for the situation with random covariates. The above results rely on a result in Atiullah (1962) which is stated without proof. We provide a proof of this result, both for easy reference, but also since the derivation of the variance bounds rely on an observation in the construction of this proof.

Keywords: General linear models, non-Gaussian error terms, moments of variance estimators, finite sample size bounds, random covariates, unconditional bounds

1 Introduction

In the present note we will consider the variance of the standard unbiased variance estimator in a general linear model (GLM) with non-Gaussian error terms. For this class of models the covariance of the generalized least squares estimator $\hat{\beta}$ is well-known, but an explicit expression for the variance of $\hat{\sigma}^2$ is not as widely known. In [2] an expression for this variance is provided without proof for a mixed GLM with non-Gaussian homoscedastic error terms. This result allows us to obtain finite sample size bounds on the variance of $\hat{\sigma}^2$ which are independent of the covariates in the non-Gaussian GLM. Moreover, the bounds which we obtain are sharp in the sense that

both the lower and upper bound, when scaled by the sample size n , converges to the asymptotic variance of $\sqrt{n}\hat{\sigma}^2$ as $n \rightarrow \infty$. Further, we note that these results remain valid if the covariates are random, but conditioned upon. Consequently, since the variance bounds are independent of covariate information it is possible to obtain *unconditional* variance bounds for the situation when the covariates themselves are random, given suitable regularity conditions. This is interesting since the unconditional variance of the standard unbiased variance estimator is in general not explicitly computable due to the randomness of the covariates.

The problem formulation above of course relies on the theory of random quadratic forms. For more on this topic, see e.g. [4, 5, 8] and the references therein. One can also note that the special case of a sample variance in the non-Gaussian setting, i.e. an intercept only GLM, was treated already in e.g. [3, Eq. (27.4.2)]. Other similar results are obtained in the theory of minimum variance component estimation, see e.g. [6, 7] and the proof of [8, Thm. 3.4].

One direct application of the bounds derived in the present note is that they provide a simple way of assessing consistency of $\hat{\sigma}^2$.

Another motivating example is the situation when we want to calculate conditional moments of a non-linear function of the parameter estimators. This is typically computationally infeasible, but the results in the present note provide exact expressions for second order Taylor approximations of conditional moments. Analogous approximations of unconditional moments will not result in exact expressions, but we may obtain sharp finite sample size bounds.

2 The General Linear Model

Let $\mathbf{Y}$ be a random $n \times 1$ vector and let $\mathbf{X}$ be a random $n \times p$ matrix of almost surely full column rank, $\text{rank}(\mathbf{X}) = p$, with $n > p$. Further, let $\mathbf{\Sigma}$ be a symmetric almost surely strictly positive definite $n \times n$ matrix. This ensures that we may define $\mathbf{\Sigma}^{1/2}$ in the standard way using orthogonalization. The class of GLMs which will be studied in the present note are of the form

$$\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \sigma\mathbf{\Sigma}^{1/2}\mathbf{e}, \quad (1)$$

where $\boldsymbol{\beta}$ is a $p \times 1$ vector, $\sigma > 0$ is a scalar and $\mathbf{e}$ is some random $n \times 1$ vector whose elements are independent. Moreover, we assume that $\mathbf{e}$ has, *conditional* on $\mathbf{X}$ and $\mathbf{\Sigma}$, mean $\mathbf{0}$ and covariance $\mathbf{I}_n$, together with common central fourth moments μ_4 . Here $\mathbf{I}_n$ denotes the $n \times n$ identity matrix. The

standard generalized least squares estimator of β , conditional on $\mathbf{X}$ and Σ , is given by:

$$\hat{\beta} = (\mathbf{X}'\Sigma^{-1}\mathbf{X})^{-1}\mathbf{X}'\Sigma^{-1}\mathbf{Y},$$

see for instance [8, Sec. 3.10] for this and more on the general linear model. Moreover, an unbiased estimator of σ^2 (conditional on $\mathbf{X}$ and Σ), and the estimator we will focus on in the present note, is given by:

$$\hat{\sigma}^2 = \frac{1}{n-p} (\mathbf{Y} - \mathbf{X}\hat{\beta})' \Sigma^{-1} (\mathbf{Y} - \mathbf{X}\hat{\beta}). \quad (2)$$

The results below will of course remain valid, with the obvious changes, if we consider an estimator normalized with some other, non-degenerate, function of n and p , for instance simply n . It is important to note that $\mathbf{X}$ is random, but we assume that it is possible to observe perfectly. That is, we are not dealing with an errors-in-variables model, which would lead to problems such as biased estimators.

As stated in the introduction, the result of [2] concerns general linear mixed models with non-Gaussian homoscedastic errors. In order to see how this result may be used in the above setting, let $\widetilde{\mathbf{Y}} := \Sigma^{-1/2}\mathbf{Y}$ and $\widetilde{\mathbf{X}} := \Sigma^{-1/2}\mathbf{X}$, which gives us that we can restate (1) as the linear model

$$\widetilde{\mathbf{Y}} = \widetilde{\mathbf{X}}\beta + \sigma\mathbf{e}, \quad (3)$$

and may hence rephrase the estimator $\hat{\sigma}^2$ in terms of (3). Further, define

$$\mathbf{U} := \Sigma^{-1/2}\mathbf{X}(\mathbf{X}'\Sigma^{-1}\mathbf{X})^{-1}\mathbf{X}'\Sigma^{-1/2},$$

which corresponds to the projection matrix associated with the linear model (3), and we note for future reference that $\text{tr}(\mathbf{U}) = \text{rank}(\mathbf{X}) = p$. If, in addition, we introduce

$$\mathbf{K} := \mathbf{I}_n - \mathbf{U}$$

we obtain the following result, which is a special case of a result stated without proof in [2]:

Proposition 1. *(Atiqullah (1962) for the special case (1), see [2]) Let*

- (i) Σ be a random symmetric almost surely positive definite $n \times n$ matrix and $\mathbf{X}$ be a random $n \times p$ matrix of almost surely full column rank,

(ii) the error term components defining the random $n \times 1$ vector $\mathbf{e}$ be independent with, conditional on $\mathbf{X}$ and $\mathbf{\Sigma}$, mean 0, variance 1 and common fourth central moments μ_4 .

It then holds that

$$\text{Var}(\hat{\sigma}^2 | \mathbf{X}, \mathbf{\Sigma}) = \frac{\sigma^4}{n-p} \left(2 + \frac{\mu_4 - 3}{n-p} \sum_{i=1}^n \mathbf{K}_{ii}^2 \right). \quad (4)$$

The proof of Proposition 1 relies on that (2) may be rewritten according to

$$\hat{\sigma}^2 = \frac{\sigma^2}{n-p} \mathbf{e}' \mathbf{K} \mathbf{e}, \quad (5)$$

which together with an application of [8, Thm. 1.6] yields the desired result. Further, the derivation of the variance bounds stated below in Corollary 1 and 2, rely on that $\mathbf{U}$ and $\mathbf{K}$ are idempotent, facts used in the derivation of (5).

Proof. (Proof of Proposition 1) In order to see how this result is obtained, one can note that if we let $\hat{\tilde{\mathbf{e}}} := \tilde{\mathbf{Y}} - \tilde{\mathbf{X}}\hat{\boldsymbol{\beta}}$, it follows that

$$\hat{\sigma}^2 = \frac{1}{n-p} \hat{\tilde{\mathbf{e}}}' \hat{\tilde{\mathbf{e}}},$$

which may be expressed in terms of $\mathbf{e}$ by noting that

$$\begin{aligned} \hat{\tilde{\mathbf{e}}} &= (\tilde{\mathbf{X}}\boldsymbol{\beta} + \sigma\mathbf{e}) - \mathbf{U}\tilde{\mathbf{Y}} \\ &= \tilde{\mathbf{X}}\boldsymbol{\beta} + \sigma\mathbf{e} - \mathbf{U}(\tilde{\mathbf{X}}\boldsymbol{\beta} + \sigma\mathbf{e}) \\ &= \sigma(\mathbf{I}_n - \mathbf{U})\mathbf{e} = \sigma\mathbf{K}\mathbf{e}. \end{aligned}$$

Hence, we have

$$\hat{\sigma}^2 = \frac{\sigma^2}{n-p} \mathbf{e}' \mathbf{K}' \mathbf{K} \mathbf{e}.$$

Further, since $\mathbf{U}$ is a projection matrix, it follows that $\mathbf{U}$ is idempotent, and moreover, $\mathbf{U}$ is, by definition, symmetric. Thus, $\mathbf{K}$ will inherit these properties as well, since $(\mathbf{I}_n - \mathbf{U})(\mathbf{I}_n - \mathbf{U}) = \mathbf{I}_n - 2\mathbf{U} + \mathbf{U}\mathbf{U}$ and, trivially, $\mathbf{I}_n - \mathbf{U}$ is symmetric if $\mathbf{U}$ is symmetric. Due to this we arrive at

$$\hat{\sigma}^2 = \frac{\sigma^2}{n-p} \mathbf{e}' \mathbf{K} \mathbf{e},$$

which is exactly (5). Continuing, recall from above that $\text{tr}(\mathbf{U}) = p$, which immediately yields $\text{tr}(\mathbf{K}) = \text{tr}(\mathbf{I}_n - \mathbf{U}) = \text{tr}(\mathbf{I}_n) - \text{tr}(\mathbf{U}) = n - p$. Finally, an application of [8, Thm. 1.6], which is also stated without proof in [1], gives us

$$\begin{aligned}\text{Var}(\hat{\sigma}^2 | \mathbf{X}, \boldsymbol{\Sigma}) &= \frac{\sigma^4}{(n-p)^2} \text{Var}(\mathbf{e}' \mathbf{K} \mathbf{e} | \mathbf{X}, \boldsymbol{\Sigma}) \\ &= \frac{\sigma^4}{n-p} \left(2 + \frac{\mu_4 - 3}{n-p} \sum_{i=1}^n \mathbf{K}_{ii}^2 \right).\end{aligned}$$

which corresponds to the result of Proposition 1. $\square$

Remark 1. *Note that for the application of [8, Thm. 1.6] in the proof of Proposition 1 we need not assume that the error terms have common (conditional) third moments. This can be seen by inspecting the proof of [8, Thm. 1.6] and noting that any parts involving the third moments disappear since the first moment is 0.*

Further, note that it is possible to generalize Proposition 1 in the obvious way when the error terms have finite, but not necessarily common, fourth moments. This however is not very illuminating and is hence omitted.

3 Results: bounds on the variance of $\hat{\sigma}^2$

Based on Proposition 1 it is natural to approach $\mathbf{K}$ in order to obtain bounds on the variance of $\hat{\sigma}^2$. If we exploit the properties of $\mathbf{K}$ directly we arrive at the following naive bounds:

Corollary 1. *If the assumptions of Proposition 1 hold: Then*

$$\text{Var}(\hat{\sigma}^2 | \mathbf{X}, \boldsymbol{\Sigma}), \text{Var}(\hat{\sigma}^2) \in \left[\frac{2\sigma^4}{n-p}, \frac{\sigma^4(\mu_4 - 1)}{n-p} \right].$$

Proof. (Proof of Corollary 1) Since $\mathbf{K}$ is idempotent and symmetric it follows that

$$\mathbf{K}_{ii} = \mathbf{K}_{ii}^2 + \sum_{j \neq i} \mathbf{K}_{ij}^2,$$

and in turn that

$$0 \leq \sum_{i=1}^n \mathbf{K}_{ii}^2 \leq \sum_{i=1}^n \mathbf{K}_{ii} = n - p.$$

This together with (4) directly yields

$$\frac{2\sigma^4}{n-p} \leq \text{Var}(\hat{\sigma}^2|\mathbf{X}, \mathbf{\Sigma}) \leq \frac{\sigma^4(\mu_4 - 1)}{n-p}.$$

It now trivially follows that the inequality holds for $\text{Var}(\hat{\sigma}^2)$ as well. $\square$

Note that the bounds in Corollary 1 are not sufficient in order to show that the variance of $\sqrt{n}\hat{\sigma}^2$ will converge to a limit point. However, by using that $\mathbf{K} = \mathbf{I}_n - \mathbf{U}$ and exploiting the properties of $\mathbf{U}$ the lower variance bound may be tightened:

Corollary 2. *If the assumptions of Proposition 1 hold: Then*

$$\text{Var}(\hat{\sigma}^2|\mathbf{X}, \mathbf{\Sigma}), \text{Var}(\hat{\sigma}^2) \in [\nu_n - \kappa_n, \nu_n]$$

with $\nu_n := \sigma^4 \frac{\mu_4 - 1}{n-p}$ and $\kappa_n := \sigma^4 \frac{(\mu_4 - 3)p}{(n-p)^2}$.

Proof. (Proof of Corollary 2) Since $\mathbf{K} = \mathbf{I}_n - \mathbf{U}$ we can rewrite (4) as follows

$$\text{Var}(\hat{\sigma}^2|\mathbf{X}, \mathbf{\Sigma}) = \frac{\sigma^4}{n-p} \left(2 + \frac{\mu_4 - 3}{n-p} \sum_{i=1}^n (1 - U_{ii})^2 \right).$$

Expanding the square and noting that $\sum_{i=1}^n U_{ii} = p$ yields

$$\text{Var}(\hat{\sigma}^2|\mathbf{X}, \mathbf{\Sigma}) = \frac{\sigma^4}{n-p} \left(\mu_4 - 1 - \frac{(\mu_4 - 3)p}{n-p} + \frac{\mu_4 - 3}{n-p} \sum_{i=1}^n U_{ii}^2 \right).$$

Now, as before for $\mathbf{K}$, since $\mathbf{U}$ is idempotent and symmetric we know that

$$0 \leq \sum_{i=1}^n U_{ii}^2 \leq \sum_{i=1}^n U_{ii} = p, \quad (6)$$

and therefore

$$\frac{\sigma^4(\mu_4 - 1)}{n-p} - \frac{\sigma^4(\mu_4 - 3)p}{(n-p)^2} \leq \text{Var}(\hat{\sigma}^2|\mathbf{X}, \mathbf{\Sigma}) \leq \frac{\sigma^4(\mu_4 - 1)}{n-p}.$$

As before the inequality trivially follows for $\text{Var}(\hat{\sigma}^2)$. $\square$

We can also state a finite sample upper bound on the difference between the conditional and unconditional variances together with convergence of these using the bounds in Corollary 2.

Corollary 3. *If the assumptions of Proposition 1 hold: Then*

$$|\text{Var}(\hat{\sigma}^2|\mathbf{X}, \mathbf{\Sigma}) - \text{Var}(\hat{\sigma}^2)| \leq \kappa_n, \text{ for all } n,$$

and

$$n \text{Var}(\hat{\sigma}^2) \rightarrow \nu, \quad n \text{Var}(\hat{\sigma}^2|\mathbf{X}, \mathbf{\Sigma}) \rightarrow \nu, \quad \text{uniformly as } n \rightarrow \infty$$

where $\nu = \sigma^4(\mu_4 - 1)$.

Proof. (Proof of Corollary 3) The first part follows trivially from Corollary 2, but we also get that

$$n \text{Var}(\hat{\sigma}^2|\mathbf{X}, \mathbf{\Sigma}) \in [n\nu_n - n\kappa_n, n\nu_n]$$

and, since $\lim_n n\nu_n = \nu$ and $\lim_n n\kappa_n = 0$, it follows that

$$\lim_{n \rightarrow \infty} n \text{Var}(\hat{\sigma}^2|\mathbf{X}, \mathbf{\Sigma}) = \nu,$$

due to the assumptions on $\mathbf{X}$ and $\mathbf{\Sigma}$. That is, $n \text{Var}(\hat{\sigma}^2|\mathbf{X}, \mathbf{\Sigma}) \rightarrow \nu$ uniformly. By the same argument it follows that $n \text{Var}(\hat{\sigma}^2) \rightarrow \nu$ uniformly in n . $\square$

Remark 2. *One example of a more restricted sub-class of the GLM from (1) is when we explicitly include an intercept, i.e. $\widetilde{\mathbf{X}}$ contains a column of ones. In [8, Eq. (10.12)] it is shown that in this situation it holds that*

$$1/n \leq \mathbf{U}_{ii} \leq 1$$

which gives us that the lower bound in (6) can be tightened according to

$$\frac{1}{n} \leq \sum_{i=1}^n \mathbf{U}_{ii}^2.$$

Hence, in this situation we can tighten the lower bound in Corollary 2 to $\nu_n - \kappa_n + \frac{\mu_4 - 3}{(n-p)n}$.

Another natural sub-class of models contained in the GLM from (1) is to assume that $\mathbf{\Sigma}$ is diagonal. In this situation it holds that $\text{diag}(\mathbf{U}) = \text{diag}(\mathbf{P})$ where

$$\mathbf{P} := \mathbf{X}(\mathbf{X}'\mathbf{\Sigma}^{-1}\mathbf{X})^{-1}\mathbf{X}'\mathbf{\Sigma}^{-1}$$

is the projection matrix of the general linear model in (1). Thus, assuming diagonal $\mathbf{\Sigma}$ does not make any parts of the variance formula, or its bounds, any more explicit.

Acknowledgements

The authors acknowledge beneficial discussion with Rolf Sundberg.

References

- [1] Atiqullah, M. (1962). The estimation of residual variance in quadratically balanced least-squares problems and the robustness of the F-test. *Biometrika*, 49(1-2), pp. 83–91.
- [2] Atiqullah, M. (1962). On the effect of non-normality on the estimation of components of variance. *Journal of the Royal Statistical Society. Series B (Methodological)*, 140–147.
- [3] Cramér, H. (1946). *Mathematical Methods of Statistics*. Princeton: Princeton university press.
- [4] Eaton, M.L. (1983). *Multivariate Statistics: A Vector Space Approach*. Wiley, New York.
- [5] Mathai, A.M., Provost, S.B., & Hayakawa, T. (2012). *Bilinear Forms and Zonal Polynomials* (Vol. 102). Springer Science & Business Media.
- [6] Rao, C.R. (1970) Estimation of Heteroscedastic Variances in Linear Models. *Journal of the American Statistical Association*, 65(329), pp. 161–172.
- [7] Rao, C.R. (1971) Minimum Variance Quadratic Unbiased Estimation of Variance Components. *Journal of Multivariate Analysis*, 1, pp. 445–456.
- [8] Seber, G., & Lee, A. (2003). *Linear Regression Analysis* (2nd ed.). John Wiley & Sons, Inc.