

Gaussian Quadrature with Correlation and Broader Sampling

Paul V. Preckel¹

Monika Verma

Thomas Hertel

William Martin

January 3, 2010

Introduction

Gaussian quadrature (GQ) has proven to be a useful approach to performing systematic sensitivity analysis on computable general equilibrium models. By systematic sensitivity analysis, we refer to the case where the analyst makes explicit assumptions regarding the probability distributions of the exogenous model inputs that are not known with certainty, and estimates statistics regarding the probability distributions of the model results (e.g. means, variances, confidence bounds, etc.). This activity is greatly facilitated by the incorporation of the procedure into the GEMPACK modeling system (Arndt 1996; Arndt and Pearson 1998).

Users have noted two shortcomings with the system. First, the initial implementation allowed for only two cases of stochastic relationships between the model inputs – either perfect correlation or stochastic independence. Second, it has been noted that the sample points for the model inputs at which the model is evaluated do not range very far from the mean – plus or minus only one standard deviation with the Liu (1997) quadrature and plus or minus the square root of two times the standard deviation for the quadrature due to Stroud (1957). The purpose of this note is to clarify and demonstrate methods for correcting these shortcomings. (Another shortcoming relates to the perceived incompatibility of using the automated SSA system simultaneously for both “shocks” and “parameters.” This point is addressed in a working paper by Horridge and Pearson [2010], and because it appears they have found a way to implement this feature in the existing system, it is not addressed here.)

Correlation between Model Inputs

There is really not much problem with allowing for correlation between parameters in a Gaussian quadrature (GQ). As indicated in Preckel and DeVuyst (1992), adding correlation between parameters is straightforward and achieved by a linear transformation of an uncorrelated quadrature. That paper recommended using an Eigen system decomposition of the covariance matrix to implement the transformation. Others such as Artavia et al. (2009) have proposed using the Cholesky factorization ($R^t R$ where R is an upper triangular matrix).

While algebraically these appear to achieve the same end, they are not geometrically equivalent. This is illustrated in Figure 1. The independent quadrature (black stars) is a two-point, three-moment Cartesian product quadrature with equally weighted points. It is transformed via the Cholesky decomposition to get the red stars for a highly correlated approximate distribution. In

¹ This paper was prepared for presentation at the 14th Annual GTAP Conference, June 16-18, 2011, Venice, Italy. The author is grateful for useful comments from M. Horridge.

contrast, using the Eigen system results in the blue points. It is interesting that for a polynomial of degree less than or equal to two, the blue and red points will result in the same expectation. However, for higher order polynomials or other nonlinear functions the results will be different. Which approach is better? It is not obvious, and here we provide simulation evidence that suggests that they are similarly effective.

Broader Sampling

The idea with Gaussian quadrature-based systematic sensitivity analysis (SSA) is to assume an explicit joint distribution for the model inputs (SSA inputs), evaluate the model whose sensitivity is being investigated at a limited number of vectors and form the weighted sum of the associated results (the SSA outputs) to get an approximation to the mathematical expectation of the model results. (One may also use a related weighted sum to approximate the expectation of some function of the model results such as their variance.)

The motivation underlying Gaussian quadrature for the selection of the points and weights is that they should be chosen so that the mathematical expectations of all polynomials up to a certain degree are exactly equal to the true expectations – that is, there is no approximation error for polynomials of limited degree. Thus, a degree o quadrature is a set of vectors x^i and weights p^i that satisfy the following system of equations:

$$(1) \quad \sum_{i=1}^n p^i \prod_{j=1}^m [x_j^i]^{k_j} = E \left\{ \prod_{j=1}^m [x_j]^{k_j} \right\} \quad \text{for } k_j = 0, 1, \dots, m \text{ such that } \sum_{j=1}^m k_j \leq o.$$

The construction of analytic formulas for Gaussian quadratures in arbitrary dimensions is a challenging task. Analytic formulas are known for arbitrary dimension for degree $o = 2$. In addition, quadratures that include some moments of degree three for symmetric distributions and arbitrary dimensions are also known. (These are addressed in Stroud [1957] and Liu [1997]. Note that while both Stroud and Liu claim that their formulas are degree $o = 3$, they are not. While all polynomials of degree less than or equal to two are exact and expected third central moments for the individual variables are equal to zero, as will always be the case with a symmetric distribution, third order polynomials containing a product of three different variables or the product of one variable and the square of another variable may not be exact. While this does not negate their usefulness, it is not technically accurate to call these degree three quadratures.)

For convenience, we focus on the case where the distribution is symmetric, with mean equal to the zero vector and covariance matrix equal to the identity. The Stroud quadrature implemented in the GEMPACK software for automated SSA samples has all vectors on the surface of a hypersphere centered on the origin for which no vector component exceeds $\pm \sqrt{2}$. Thus, no component of any vector for which the model is evaluated is greater than $\sqrt{2}$ in absolute magnitude. Similarly for the Liu quadrature (also implemented in GEMPACK for automated SSA), no component of any vector for which the model is evaluated is greater than 1.0 in absolute magnitude.

Our goal is to develop a new quadrature that samples more broadly but which satisfies the following system of moment equations (the same system satisfied by the Stroud and Liu quadratures):

$$(2) \quad \begin{aligned} \sum_{i=1}^n p^i &= 1, \\ \sum_{i=1}^n p^i x_j^i &= 0 \text{ for all } j, \\ \sum_{i=1}^n p^i x_j^i x_k^i &= \delta_{jk} \text{ for all } j \text{ and } k, \text{ and} \\ \sum_{i=1}^n p^i [x_j^i]^3 &= 0, \end{aligned}$$

where $\delta_{jk} = 1$ if $j = k$ and $\delta_{jk} = 0$ otherwise. Denote a given quadrature satisfying this system by $[p^i, x^i]_{i=1}^n$. This quadrature could be derived by the Stroud or Liu formula, or any other formula that satisfies the system above. To construct a new quadrature that samples more broadly, we conceptually (a) make two copies of the quadrature above, (b) stretch one to achieve the desired broadening of the sample, (c) shrink the other and allocate probability across the two copies so as to make it possible to maintain the variances. To operationalize these ideas, it is useful to introduce three scalars: an expansion factor α , a contraction factor β and a probability allocation factor q . The final quadrature will be the following:

$$(3) \quad \left\{ [qp^i, \alpha x^i]_{i=1}^n, [(1-q)p^i, \beta x^i]_{i=1}^n \right\}.$$

This quadrature has twice as many vectors ($2n$) as the original quadrature and can be chosen to expand the sampling as much as desired by controlling the parameter α , where $\alpha > 1 > \beta$. Once α is chosen, the values for β and q follow from the system of equations above. A small technical point is that there is an extra degree of freedom that must be nailed down in order to solve for β and q . Here this degree of freedom has been removed by choosing to specify the kurtosis for *one* of the variables. Denoting this kurtosis by κ , the parameters of the desired quadrature can be found by solving the following system:

$$(4) \quad \begin{aligned} \alpha^2 q + \beta^2 (1-q) &= 1 \\ \alpha^4 q + \beta^4 (1-q) &= \kappa. \end{aligned}$$

Focusing on only positive values for α and β , the solution is:

$$(5) \quad \begin{aligned} q &= \frac{1 - \kappa}{\alpha^4 - 2\alpha^2 + \kappa} \text{ and} \\ \beta &= \sqrt{\frac{\alpha^2 - \kappa}{\alpha^2 - 1}}. \end{aligned}$$

It is straightforward to verify that these choices satisfy the systems (4) and (2).

Experiments

We now present two experiments designed to examine the effect of the theory derived above. The first of experiments focuses on the linear transformations, while the second demonstrates the impact of broader sampling.

Experiment 1 – Design

For the first experiment, three CGE models are used to provide evidence of the importance of the choice of transformation used for incorporating imperfect covariance between SSA input variables. The first case, due to Shoven and Whalley (1984), is a two household, two factor, two consumption good CGE model focused on tax policy with CES specifications of all production and preference relationships. The SSA input variables are the elasticities of substitution for the two households between the two consumption goods and the elasticities of substitution for the two factors in the two sectors for a total of four SSA input variables. The assumed distributions for the SSA input variables (for this problem as well as to the other two) are normal with mean equal to the original values in the deterministic model (i.e. with no distribution assumed for the SSA parameters). The covariance structure is discussed below. The SSA output variables are the prices of the factors and the consumption goods. Note that this is not a policy simulation. Rather, we simply are estimating statistics of these SSA output variables whose variability is driven by the distributional assumptions regarding the SSA input variables.

The second case is from Hansen's thesis (1968) and is published in Scarf and Hansen (1973). This is a four consumer, three factor, 14 good model with an activity analysis representation of productive possibilities. The SSA input variables are the elasticities of substitution between goods in the preference relationships for each of the four households, and the level of output for the domestic agriculture sector. To provide a special challenge for the methods, the activity levels for four productive sectors are chosen – import activities 2, 3, 5 and 7 – to be the SSA output variables. We focus on these particular activities because they exhibit some “action” when the SSA input variables are varied. In particular, these activities have the potential to have zero levels for some draws of the input variables, thus illustrating the performance of the methods when the SSA output variables are non-differentiable functions of the SSA input variables. As with the Shoven and Whalley case, this is not a policy simulation. We are estimating statistics (mean and variance) of these SSA output variables whose variability is driven by the distributional assumptions regarding the SSA input variables.

The third case is a highly aggregated version of the GTAP model that includes five regions, eight commodities and five factors. The SSA input variables are sector productivity indices (five regions times eight commodities for a total of 40 SSA input variables). Assumed distributions for the input variables are joint normal with mean equal to zero. The SSA output variables in this case are the percent changes in the factor prices for land in Africa by sector (land is not mobile across sectors in this model). This is a policy simulation. The percentage changes in factor prices are jointly driven by the productivity indices and by a full import tariff elimination policy.

One goal of this research is to assess the performance of the GQ approach with the two alternative methods for incorporating correlation between the SSA input variables. This is of importance for both the Monte Carlo and the GQ approaches. It is important for the former because the vast majority of methods for pseudo random number generation are designed to generate univariate random numbers. Thus to generate correlated random numbers, one first generates vectors of independent univariate numbers and applies a linear transformation to obtain correlations. This is especially easy to do with normal random numbers where a linear transformation of a vector of independent normal random variables is distributed according to a joint normal distribution, with parameters that are easily derived from the parameters of the transformation. For example, if x is a vector of independent identically distributed normal random variables with mean zero and variance one, then the vector of random variables $y = b + Ax$ where A is a square matrix is distributed as joint normal with mean b and covariance matrix AA^t , where the t indicates a transpose. If a given covariance structure, Σ , is desired, then y will have the desired covariance structure for any matrix A such that $\Sigma = AA^t$. Note that the choice of A is not unique. Two approaches that are common in the literature are to choose A as a lower triangular matrix R^t such that $R^t R = \Sigma$. Another option is to choose $A = QD^{1/2}$ where Q is an orthonormal (i.e. its transpose is its inverse) matrix whose columns are the eigenvectors of Σ , and $D^{1/2}$ is a diagonal matrix whose diagonal elements are the square roots of the eigenvalues of Σ (in order corresponding to the ordering of the eigenvectors in Q). Of course, there is an infinity of other choices, but methods are well developed for calculating these two types of decompositions of any positive definite matrix. For the purposes below, we refer to the linear transformation based on the Cholesky decomposition as the $R^t R$ approach and the transformation based on the Eigen system as the QDQ^t approach.

Because we want to avoid dependence of our results on a particular covariance structure, we will draw a series randomly generated covariance structures and apply our comparison of methods for each. We then report average results where the averages are taken over the draws of alternative covariance structures. Now consider the comparisons for each covariance structure. The quantities of interest are the expected values and variances of the SSA outputs where their variability is driven by the randomness of the SSA inputs.

We take as our basis of comparison the Monte Carlo approach. It is well established that the Monte Carlo method under fairly mild conditions produces estimates of an expectation that can be made arbitrarily good *provided that the sample size, or number of draws, is large enough* (see e.g. Haber [1970]). The italicized phrase ending the previous sentence provides our motivation for considering Gaussian quadrature. The trouble is that the sample size needed for reasonably accurate estimates is typically quite large and therefore not practical for systematic sensitivity analysis with models for which an individual solution may require significant computing time.

For the exercise presented here, we selected three models that can be solved quickly enough that using Monte Carlo methods to assess the expected values and variances of the SSA outputs is practical. To do this, we generate a (fairly) large sample of random vectors from the SSA input variable distributions, solve the model for each vector and average each SSA output variable to obtain an estimate of its mean, and use the typical sample based formula to estimate its variance. Unfortunately, it is not clear what linear transformation should be used to

incorporate correlation between the SSA input variables for the Monte Carlo. As a check, we use both the R^tR and QDQ^t approaches and compare the results to see if they are very different from each other.

Our ultimate goal is to assess the effectiveness of Gaussian quadrature for the correlated case. To generate the Gaussian quadrature estimates, we use a procedure that parallels the Monte Carlo approach. Namely, we first generate multi-variate Gaussian quadratures (GQs) that have zero mean, variances equal to one, and covariances that are all zero. We then apply a linear transformation to each of the points in these quadratures to obtain a multi-variate quadrature with the desired mean vector and covariance matrix. (Because mathematical expectation is a linear operator, the algebra for the impact of the linear transformation on the moments of the Gaussian quadrature estimates is identical to the Monte Carlo case.) As with Monte Carlo, there is ambiguity in the selection of the matrix in the linear transformation. To get evidence on the efficacy of using the R^tR or QDQ^t approaches for generating this linear transformation, the estimates are calculated twice – once for the GQ that is transformed using the R^tR approach, and once for the GQ transformed using the QDQ^t approach. However, a basis of comparison is needed. While the best available estimates of the means and variances of the SSA outputs are the Monte Carlo results, we are unsure if the results based on R^tR or on QDQ^t are superior. Hence, we compare the GQ results to each of these Monte Carlo results in turn.

One final detail should be explained before turning to the description of the experiment – the random generation of the covariance matrices. These are generated in a manner that facilitates the computations – namely, by first generating the elements of the Eigen system, Q and D , then by calculating the covariance matrix $\Sigma = QDQ^t$, and finally by calculating the Cholesky factors for Σ . The matrix Q is generated as a product of $n - 1$ plane rotations (one for each of the off-diagonal pairs of variables), where the angle of the rotation is randomly generated from a uniform distribution on $[0, 2\pi]$, and the i -th diagonal element of D is taken to be equal to ten percent of the mean value of the i -th SSA input variable. (Because the means for the GTAP model SSA input parameters are zero, some other scheme was necessary. The diagonal elements of D for the GTAP experiment are uniformly generated random numbers on the interval $[50, 100]$.) The latter selection is somewhat arbitrary, but serves to ensure that the likelihood of obtaining a parameter value outside the range of values that ensures regularity (e.g. a negative substitution elasticity) is highly unlikely. (Operationally, in the rare instances where negative elasticities are generated, they are simply reset to a small value. This happens very rarely.)

So the procedure for our experiment is as follows:

- Generate a covariance matrix for the SSA input variables.
 - Generate a Monte Carlo sample of vectors of independent normal random variables.
 - Transform the Monte Carlo sample via the R^tR approach to reflect the mean vector and the current generated covariance matrix.
 - Solve the model for each sample point.

- Average the results to get Monte Carlo estimates of the SSA output means, and use the usual sample variance formula to get Monte Carlo estimates of the SSA output variances.
 - Transform the Monte Carlo sample via the QDQ^t approach to reflect the mean vector and the current generated covariance matrix.
 - Solve the model for each sample point.
 - Average the results to get Monte Carlo estimates of the SSA output means, and use the usual sample variance formula to get Monte Carlo estimates of the SSA output variances.
- Generate a multivariate Gaussian quadrature (GQ) for independent random variables using the Stroud method (see Arndt [1996]).
- Transform the GQ sample via the R^tR approach to reflect the mean vector and the current generated covariance matrix.
 - Solve the model for each sample point.
 - Weight the GQ results by the associated probabilities and sum to get GQ estimates of the SSA output means and variances.
 - Transform the GQ sample via the QDQ^t approach to reflect the mean vector and the current generated covariance matrix.
 - Solve the model for each sample point.
 - Weight the GQ results by the associated probabilities and sum to get GQ estimates of the SSA output means and variances.

This procedure is applied for each randomly generated covariance matrix, and the average absolute percentage errors are calculated for the means and variances of each of the SSA output variables, where the average is across the alternative covariance matrices.

Experiment 1 – Results

Simulation results are presented in Table 1. All Monte Carlo estimates are based on 1,000 model solutions, and each of the GQ estimates is based on $2n$ solutions, where n is the number of SSA input variables (n is four in the Shoven and Whalley case, five in the Hansen case, and forty in the GTAP case). The values reported in the table are averages across 100 draws of the covariance matrix for the first two problems and averages across 10 draws of the covariance matrix for the GTAP problem. The latter reflects the fact that each instance of the GTAP problem takes an order of magnitude longer to solve than the smaller problems and that due to the large number of SSA input variables, the amount of computation for the GQ estimates is also substantially larger (although still much smaller than the associated Monte Carlo estimates). Future work will expand the sample sizes for the GTAP model, likely using supercomputer resources.

The first column in Table 1 displays the average relative difference between the Monte Carlo results based on the QDQ^t transformation divided by Monte Carlo results based on the R^tR

transformation. Consider the first block of these numbers, which are for the Shoven and Whalley problem. The top half of these (labeled Mean) are the percentage errors in the estimates of the means of the factor prices (Labor and Capital) and consumption goods prices (Manuf. and Non-manuf.) relative to the R^tR based Monte Carlo estimates. That is, the R^tR based Monte Carlo is taken to be the basis for comparison for this column. These are quite small considering the small size of the Monte Carlo estimates (only 1,000 draws), with the largest amounting to less than 0.3 percent. The bottom half of this block corresponds to the estimates of the variances of the model results. These are quite a bit larger, with the largest being on the order of 9 percent squared, or about 3 percent in standard deviation terms. This is not surprising given that the variance is a more nonlinear function of the SSA input parameters than the mean.

The second column of Table 1 is similar to the first except that the method evaluated is the Monte Carlo based on the R^tR transformation and rather than taking the Monte Carlo results based on the R^tR transformation as the basis for comparison, it uses the Monte Carlo results based on the QDQ^t transformation as the standard. These are quite close to the values in column 1, leading us to feel some confidence that the choice of transformation is not so important when the estimation method is Monte Carlo.

The third column gives the average percentage errors in the GQ estimates of the statistics for the SSA output variables where the GQ is transformed based on the QDQ^t transformation, and the Monte Carlo results based on the R^tR transformation are taken as the basis for comparison. These are of similar magnitude to either of the first two columns, doing slightly better for most of the SSA output variable statistics, with the exceptions of the variances for prices of the consumer goods. The fourth column gives the average percentage errors in the GQ estimates of the statistics for the SSA output variables where the GQ is transformed based on the R^tR transformation, and the Monte Carlo results based on the R^tR transformation are taken as the basis for comparison. The differences between columns three and four are not large.

Columns five and six of Table 1 are analogous to columns three and four, but they report the average percentage errors arising from the GQ that is transformed based on the R^tR transformation. The results in columns five and six are also quite similar to each other, indicating that the choice of standard of comparison (between our two alternative Monte Carlo approaches) is not critical. Now compare columns three and five. This comparison reflects the difference between the GQ transformed using the QDQ^t approach versus the GQ transformed using the R^tR approach (using the Monte Carlo results based on an R^tR -based transformation as the basis for comparison). It is interesting that the percentage errors are uniformly lower for the GQ using the QDQ^t -based transformation, although the differences are quite small. Columns four and six also reflect the difference in choice of the transformation approach for the GQ, but with a different basis of comparison – the Monte Carlo results using a QDQ^t -based transformation. Some of the average percentage errors in column four are smaller than the ones in column six, and some are larger.

Based on this evidence from the Shoven and Whalley case, we draw a couple of conclusions. First, the choice of transformation method between our two candidate alternatives (R^tR and QDQ^t) does not have a large impact on the estimated statistics. Second, the choice of

transformation method for the GQs does not have a large impact on the estimates of the SSA output variable statistics.

Now consider the second block of numbers, which correspond to the case due to Hansen. These tell a similar story to the one for the Shoven and Whalley case. That is, there is little difference between the SSA output percentage error statistics for the Monte Carlo results based on R^tR or QDQ^t transformations (i.e. columns one and two are of very similar magnitude, matching to at least two significant digits). Similarly, the percentage error estimates for the GQ estimates are of similar magnitude regardless of what transformation is used for the basis of comparison. As for the Shoven and Whalley case, the percentage errors for the means are quite small and the percentage errors for the variances are quite a bit larger.

The results based on the GTAP model are a bit different (see the third and final block of numbers in Table 1). The average percentage differences between the Monte Carlo results based on the two transformations is much larger – on the order of 5-20 percent for the mean estimates and 4-7 percent for the variance estimates. (The mean estimates for Other Ag. and Services are exceptions to this rule. This is due to the combination of low mean and high sensitivity of the results to the SSA input variables for these two SSA output variables. In particular, this means that the base for the percentage error calculations is small – in the case of the mean for Services the value obtained from the Monte Carlo with the RtR transformation is -0.023 percent.) The variances are of much more substantial magnitude, and the percentage error estimates for them are in line with those for the other two models.

An interesting pattern emerges from these three cases. The percentage errors for the estimates of the means and variances of the SSA outputs are of about the same magnitude for all cases. This could be caused by inaccuracy of the Monte Carlo estimates. For purposes of argument, assume that the GQ estimates are exact for a particular statistic and that the Monte Carlo estimate is inaccurate by 1 percent. This means that because the basis of comparison is the Monte Carlo, each of the GQ estimates will be deemed to be inaccurate by 1 percent. The theory for error analysis indicates that the error in a Monte Carlo estimate should decline (in a probabilistic sense) at a rate of one over the square root of d , where d is the number of draws. Thus, if the number of draws is increased by a factor of 100, then the number of correct digits in the Monte Carlo estimate should increase by 1 (see Haber [1970]). Thus, the conjecture that inaccuracy of the Monte Carlo results could be causing the numbers in each row in Table 1 to be so similar can be tested by simply increasing the number of draws.

Experiment 2 – Design

The second experiment uses the same set of models and the same choices of input and output SSA variables. Three levels of sample expansion (α in the theory section) are employed for illustrative purposes: 1.5, 2.0 and 3.0. The Stroud quadrature is employed for this experiment. An expansion of 1.5 means that the most distant points will be just over two standard deviations from the mean. In the univariate normal case, values this extreme occur with probability less than four percent of the time. An expansion of 2 means that the most distant points will be over 2.8 standard deviations from the mean, which will occur in the univariate normal case on the order of about one half of one percent of the time. An expansion of 3 means that the most distant

points will be over 4 standard deviations from the mean, which will occur in the univariate normal case on the order of one hundredth of one percent of the time. So, while the expansion factors may not seem large, they represent substantial deviations in the normal case. As mentioned in the theory section, it is necessary to specify the kurtosis for one of the variables in order to fully determine the expanded quadrature. Here we have arbitrarily chosen to set the kurtosis for the first variable to 3.0 – the kurtosis value for the standard normal distribution.

As in the first experiment, random correlations between the SSA input variables are incorporated. These correlations are reflected in the quadrature by applying the R^tR -based transformation to the quadrature vectors, and results are compared to the Monte Carlo results that employ the R^tR -based transformation to incorporate correlation. These are the same Monte Carlo results used for comparison in the first experiment.

Experiment 2 – Results

Table 2 displays the results of the simulation exercise. The figures in the table are average absolute percentage errors in the mean SSA output variables and variances of SSA output variables, where the Monte Carlo results based on the R^tR approach to incorporating correlations are treated as the correct values. As in Table 1, the first block of numbers corresponds to the model due to Shoven and Whalley. Generally, the error estimates are slightly larger with the sample expansion – about 1-2 percent for the means and 1-33 percent for the variances. (Results for the unexpanded sample are found in column five of Table 1.) Exceptions for this rule are the percentage errors for the mean factor price for Non-manufactures which are 1-2 percent smaller than with no expansion, and the variance for Manufactures for the highest level of expansion, which is about 5 percent smaller than with no expansion. As the expansion increases, the errors generally grow for the mean estimates, but generally fall for the variance estimates. The variance of Non-manufactures is an exception to this rule.

The second block of numbers corresponds to the model due to Hansen. With the expansion factor of 1.5, there is almost no impact on the errors with the exception of the estimates of the variances for Import activities 3 and 5, which increase by about 10 percent and 4 percent, respectively. Increasing the expansion factor to 2.0 has little impact on the error estimates for the means or variances, with the exception of the variance for Import activity 5, which increases to nearly 10 percent. The further increase to an expansion factor of 3.0 has generally negative impacts with increases in the errors greater than 1 percent for the mean estimates for Import activities 5 and 7, and for the variance estimates for all of the Import activities except 2. Indeed, the average percentage error estimate increases very substantially for the variance estimate for Import activity 5, which increases to over 48 percent.

The third block of numbers corresponds to the GTAP model and yields generally larger impacts. With a few exceptions, the percentage errors for the means are larger for each of the expansions relative to the unexpanded quadrature. The exceptions are for average percent error estimates of the mean for Services for expansions of 2.0 and 3.0, for the variance for Rice and Other Grains for the 3.0 expansion, for the variance for Other Ag. at the 1.5 and 3.0 expansions, and for the variance for manufacturing at the 1.5 expansion. These GTAP results are a bit

perplexing, and a significant effort to redesign the generation scheme for the covariance matrix for the SSA input variables is likely warranted.

Conclusions

This technical report documents several experiments designed to explore the performance of Gaussian quadrature (GQ) approaches to systematic sensitivity analysis (SSA). In particular, a first experiment was designed to test the importance of the choice of linear transformation for incorporating covariance between SSA input variables on the accuracy of the estimates of statistics for SSA output variables. A second experiment was designed to demonstrate and test a procedure for broadening the SSA sampling procedure while increasing the number of model evaluations needed to perform the SSA by only a factor of two.

These experiments were implemented for three models that are well known in the literature: a model of tax policy due to Shoven and Whalley (1984), a model that includes an activity analysis representation of production due to Hansen (1968), and the GTAP model due to Hertel (1997). The results for all three models support the view that the particular choice of linear transformation for incorporating covariance between SSA input variables is not crucial. Further increasing the sample size of the Monte Carlo estimates that are used as the basis of comparison may add to our confidence in this regard. Results also suggest that the payoff to expansion of the range of the SSA sampling does not have a dramatic, or even consistent, impact on the accuracy of the results. The latter may seem surprising at first, but given that the results (SSA output variables) of those models are quite smooth as functions of the SSA input variables, this could have been expected.

The accomplishments of this project are manifold. First, evidence has been produced that suggests that the choice of linear transformation between the two obvious contenders – the Cholesky factorization and the Eigen system decomposition – is not terribly important from a performance perspective. Given that the Eigen system is much more difficult to compute than the Cholesky factorization, the preference for the Cholesky approach is clear. Second, an approach to broader sampling that (a) satisfies the same moment conditions as the Stroud and Liu quadratures, (b) allows the user to choose how much to expand the range of sampling, and (c) can be applied for any number of SSA input variables is derived and demonstrated. The demonstration indicated only modest changes in the performance of SSA for the first two models. However for models that exhibit different behaviors (e.g. tariff regime changes) for more extreme perturbations of the SSA input variables, the differences in performance may be more substantial. Third, a GAMS utility program was created that implements the calculations needed to incorporate covariance and to broaden the sampling. (The program is found in the appendix of this report and can be obtained electronically by e-mailing a request to preckel@purdue.edu.) This utility should serve as a blueprint for implementing these SSA features in GEMPACK at some point in the future. Fourth, methods were developed for formally assessing the performance of SSA relative to Monte Carlo methods that are applicable not only to models implemented in GAMS, but also to models implemented in GEMPACK.

Potential Future Work

Several potentially productive avenues for future work are clear. These are briefly described below in order of increasing additional work needed to complete.

First, there is a paper that should be written that focuses on the issue of what method to choose for incorporating covariance between SSA input variables into sampling distributions via linear transformations. This will require some additional work to improve the experimental design, but beyond that, the main improvement will be to increase Monte Carlo sample sizes to increase confidence that the results are robust. In particular, the Shoven and Whalley and Hansen experiments could be refocused on a policy assessment. This could be the published tax exercise for the Shoven and Whalley model, and perhaps a tariff or subsidy exercise for the Hansen model. In addition, greater effort should be expended to specify SSA input variable variances that are better in line with what has been used in past SSA exercises. This paper will present results similar to those in Table 1. The software necessary to execute the computations for this paper has been developed in this project, and should take only minor tweaking to incorporate these modifications described here. However, the computational effort will be substantial, especially for the GTAP model results. It may be best to do the initial version of this paper as a GTAP Conference paper, with the ultimate goal of submission to a peer reviewed journal.

Second, a paper should be written on the approach to broadening the range of sampling for SSA that explains how the modified GQ is constructed and demonstrates the application of the approach to some test cases to demonstrate the impact on performance of SSA. This paper will present results similar to those in Table 2 (plus column 5 in Table 1). It may be useful to include a model among the test cases that incorporates regime changes and thus will be likely to exhibit greater differences in the SSA results depending upon the sampling range. In the case of regime changes, this could also provide some evidence of the impact of failing to have differentiability of the SSA outputs as functions of the SSA inputs that are typically associated with regime changes. The modifications to the experimental design suggested for the first paper should also be incorporated in this paper. Because the Monte Carlo results are the most computationally demanding part of the experiment, the current GAMS programs incorporate the GQ calculations for both the correlated and broader sampling cases. Thus, the marginal computational effort to produce the second paper should be minor. As with the first paper, it may be best to do the initial version of this paper as a GTAP Conference paper, with the ultimate goal of submission to a peer reviewed journal.

Third, a more serious application that incorporates empirically motivated correlations between SSA inputs may serve to motivate modelers to seriously consider employing this potential feature of SSA. Perhaps cases illustrating the impact of SSA with and without correlations (but with similar variances for the SSA inputs) would help make the point. Having good motivation for the correlations will be essential.

Fourth (although this could possibly be combined with the third paper), a more serious application that uses broader sampling for a model that has regime changes could be examined. Perhaps cases illustrating the impact of using SSA with and without broader sampling (but the

same means and covariance structure for the SSA inputs) would be useful for establishing the need for broadening the sample in some cases.

References

- Arndt, C. 1996. "An Introduction to Systematic Sensitivity Analysis via Gaussian Quadrature," GTAP Technical Paper No. 2, July.
- Arndt, C. and K.R. Pearson 1998. "How to Carry Out Systematic Sensitivity Analysis via Gaussian Quadrature with GEMPACK," GTAP Technical Paper No. 3, April.
- Artavia, M., H. Grethe, T. Möller and G. Zimmermann, 2009. "Correlated Order Three Gaussian Quadratures in Stochastic Simulation Modelling," Presented at the Twelfth Annual Conference on Global Economic Analysis, Santiago, Chile, June 10-12.
- Haber, S. 1970. "Numerical Evaluation of Multiple Integrals," SIAM Review, 12(4):481-526.
- Hansen, T. 1968. "On the Approximation of a Competitive Equilibrium," Ph.D. Thesis, Yale University.
- Hertel, T.W., ed. 1997. Global Trade Analysis: Modeling and Applications, Cambridge: Cambridge University Press.
- Horridge J.M. and K.R. Pearson 2010. "Systematic Sensitivity Analysis with Respect to Correlated Variations in Parameters and Shocks," Working Paper, Center of Policy Studies, Monash University, Victoria, Australia, January.
- Liu, S. 1997. "Gaussian Quadrature and Applications," Ph.D. Dissertation, Department of Agricultural Economics, Purdue University, June.
- Preckel, P.V., and E. DeVuyst, 1992. "Efficient Handling of Probability Information for Decision Analysis under Risk," American Journal of Agricultural Economics, 74(3):655-662.
- Scarf, H. and T. Hansen, 1973. "The Computation of Economic Equilibria," New Haven: Yale University Press.
- Shoven, J, and Whalley, J, 1984. "Applied General Equilibrium Models of Taxation and International Trade: An Introduction and Survey," Journal of Economic Literature 22:1007-1051.
- Stroud, A.H. 1957. "Remarks on the Disposition of Points in Numerical Integration Formulas," Mathematical Tables and Aids to Computation, 11:257-261.

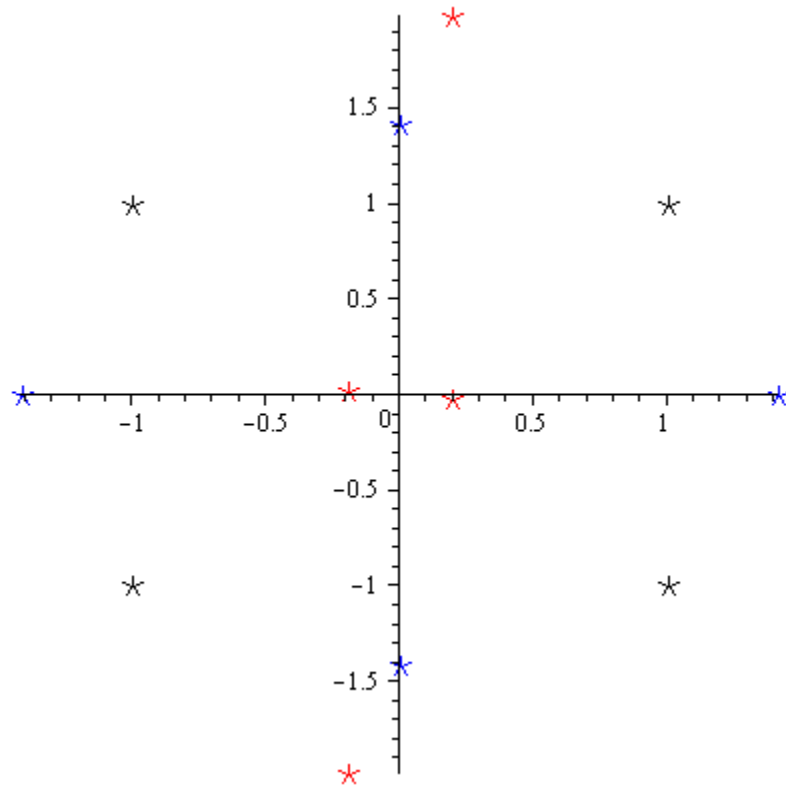


Figure 1. An Independent Order-Two Symmetric Quadrature (Black Stars) and Two Algebraically Equivalent Correlated Order-Two Quadratures Based on Cholesky Transformation (Red Stars) and Eigen system Transformation (Blue Stars)

Table 1. Average Monte Carlo and GQ Absolute Percent Error Estimates of Price Distribution Parameters

Method		Monte Carlo QDQ ^t	Monte Carlo R ^t R	GQ QDQ ^t	GQ QDQ ^t	GQ R ^t R	GQ R ^t R
Relative to Monte Carlo Based on:		RtR	QDQt	RtR	QDQt	RtR	QDQt
Shoven and Whalley							
Mean	Labor	0.1742	0.1741	0.1206	0.1157	0.1215	0.1166
	Capital	0.2963	0.2965	0.2053	0.1970	0.2070	0.1986
	Manuf.	0.0697	0.0696	0.0489	0.0456	0.0496	0.0446
	Nonmanuf.	0.1289	0.1290	0.0931	0.0853	0.0930	0.0829
Variance	Labor	5.9284	5.9799	4.6558	5.2792	4.8628	4.8220
	Capital	5.9284	5.9799	4.6558	5.2792	4.8628	4.8220
	Manuf.	8.9376	9.1553	13.3685	13.1081	12.8495	12.1856
	Nonmanuf.	7.5166	7.6369	8.4936	8.9160	8.6164	8.4628
Hansen							
Mean	Imp 2	0.0597	0.0597	0.0414	0.0420	0.0414	0.0420
	Imp 3	0.0715	0.0715	0.0497	0.0517	0.0497	0.0517
	Imp 5	0.0012	0.0012	0.0009	0.0008	0.0009	0.0008
	Imp 7	0.0583	0.0584	0.0413	0.0420	0.0413	0.0421
Variance	Imp 2	4.5774	4.5925	4.0417	3.8079	4.0474	3.8096
	Imp 3	4.5949	4.5829	3.8524	3.5395	3.9012	3.6103
	Imp 5	4.9115	4.8854	5.2403	5.0235	4.6650	4.5799
	Imp 7	4.6260	4.6230	3.8805	3.4272	3.8655	3.4888
GTAP							
Mean	Rice	5.9598	6.3338	4.2840	3.7017	4.1738	4.4370
	Wheat	3.0659	3.1371	3.1337	2.9955	3.0660	3.0470
	Oth. Grains	6.6560	7.0182	5.8745	5.0401	5.6294	4.5715
	Other Ag.	33.1664	73.7542	29.2256	68.1431	27.1316	57.5977
	Extraction	18.0168	15.7275	22.7017	17.2371	15.2701	11.4120
	Manuf.	15.4388	16.6505	8.7632	14.5773	8.5875	11.9114
	Services	265.4627	266.4477	232.8846	353.4413	216.2967	239.7209
	Cons.Gds.	18.8702	21.3717	15.5039	21.1546	11.8331	18.2367
Variance	Rice	6.0124	6.1308	5.6806	4.8581	4.0825	5.2797
	Wheat	5.9521	6.2408	6.0655	5.9636	5.8892	4.3263
	Oth. Grains	5.2165	5.3551	3.4814	3.7231	5.4385	4.1646
	Other Ag.	5.1349	5.1887	4.5688	3.5342	7.9535	6.7364
	Extraction	4.6144	4.7399	3.3522	3.6082	5.0509	4.4744
	Manuf.	6.3713	6.3087	4.9155	4.1150	5.4729	5.9048
	Services	5.9986	5.9893	5.3112	3.3176	7.6696	8.1298
	Cons. Gds.	6.5570	6.4387	5.7710	5.8744	7.6232	6.6174

Table 2. Average GQ Absolute Percent Error Estimates of Price Distribution Parameters with Broader Sampling

		Sampling Expansion Factor		
		1.5	2.0	3.0
Shoven and Whalley				
Mean	Labor	0.1236	0.1240	0.1238
	Capital	0.2104	0.2110	0.2107
	Manuf.	0.0500	0.0504	0.0507
	Nonmanuf.	0.0915	0.0919	0.0923
Variance	Labor	5.7324	5.4106	4.9232
	Capital	5.7324	5.4106	4.9232
	Manuf.	15.6219	14.6791	12.2232
	Nonmanuf.	10.8700	11.4903	11.6601
Hansen				
Mean	Imp 2	0.0414	0.0414	0.0412
	Imp 3	0.0498	0.0498	0.0497
	Imp 5	0.0009	0.0009	0.0009
	Imp 7	0.0412	0.0413	0.0418
Variance	Imp 2	4.0593	4.0601	4.0658
	Imp 3	4.2909	4.3082	4.3349
	Imp 5	4.8336	5.1268	6.9284
	Imp 7	3.8672	3.8552	4.1696
GTAP				
Mean	Rice	5.8476	4.7670	4.7484
	Wheat	3.7620	3.9633	3.2728
	Oth. Grains	9.9953	5.9734	6.5200
	Other Ag.	54.2889	25.1965	31.4372
	Extraction	20.8218	17.9915	19.6705
	Manuf.	26.0533	9.9087	11.4139
	Services	1464.6826	139.9958	184.4650
	Cons. Goods	65.0538	13.2274	18.5636
Variance	Rice	6.7170	6.3274	4.0326
	Wheat	4.5782	7.7064	8.1167
	Oth. Grains	6.7601	8.8299	4.7882
	Other Ag.	4.3892	8.4500	4.5162
	Extraction	5.1491	9.8872	7.9730
	Manuf.	5.3147	7.3339	5.5802
	Services	8.3984	14.8438	9.7022
	Cons. Goods	14.0273	19.4205	10.7050

Appendix – A GAMS Utility for Calculation of Correlated and Expanded Quadratures

```

$ontext
GQ Utility Program Version 1.0
Direct questions/problems to Paul Preckel preckel@purdue.edu
December 28, 2010
$offtext
*Boilerplate
option decimals=8 ;
sets
  v          Variables (elements should be 1*n)      / 1*3 /
  p          Points (elements should be 1*2n)        / 1*6 /
  p2         Points for expanded GQ (should be 4*n) / 1*12 / ;
abort$(card(p) ne 2*card(v) or card(p2) ne 2*card(p))
  'p must have 2*(elements of v) and p2 must have 2*(elements of p).';
alias (v,vv,vvv) ;
*User Inputs
parameter mu(v)      Mean vector for random variables
  / 1  7,  2  8,  3  9 / ;
table vcv(v,v)       Covariance matrix for random variables
  1      2      3
1      10     4     5
2       4     8     3
3       5     3     9 ;
scalars
  expand      Expansion factor for expanded quadrature      / 2 /
  kurt        Kurtosis for first variable in expanded quadrature / 3 / ;
*End of User Inputs
parameter
  r(v,v)      Cholesky factors of covariance matrix / 1.1 0 /
  muchk(v)    Verification for the mean
  vcvchk(v,v) Verification for the covariance matrix
  skwchk(v)   Verification for the skewness ;
loop(v,
  abort$(vcv(v,v) - sum(vv$(ord(vv) lt ord(v)),sqr(r(vv,v))))**0.5 le 0)
  'The covariance matrix appears to not be positive definite.' ;
  r(v,v) = (vcv(v,v)
    - sum(vv$(ord(vv) lt ord(v)),sqr(r(vv,v))))**0.5 ;
  loop(vv$(ord(vv) gt ord(v)),
    r(v,vv) = (vcv(v,vv)
      -sum(vvv$(ord(vvv) lt ord(v)),r(vvv,v)*r(vvv,vv)))/r(v,v) ;
  ) ;
) ;
display 'Cholesky factors of the covariance matrix:',r ;
vcvchk(v,vv) = sum(vvv,r(vvv,v)*r(vvv,vv)) ;
display 'These two matrices should be equal (vcvchk = RtR).',vcv,vcvchk ;
parameter
  ggpt(p,v)    Stroud-based GQ points
  prb(p)       Probabilities for Stroud GQ points ;
* Generate Stroud points
ggpt(p,v) = 2**0.5*
  (cos(ord(v)*ord(p)*pi/card(v))$(2*trunc(ord(v)/2) ne ord(v))
  + sin((ord(v)-1)*ord(p)*pi/card(v))$(2*trunc(ord(v)/2) eq ord(v))) ;
* Treat final component of points when v is odd.
ggpt(p,v)$(ord(v) eq card(v) and 2*trunc(ord(v)/2) ne ord(v))
  = power(-1,ord(p)) ;
prb(p) = 1/card(p) ;
display 'GQ points and weights before incorporating covariance',ggpt,prb ;
ggpt(p,v) = sum(vv,r(vv,v)*ggpt(p,vv)) + mu(v) ;
display 'GQ points and weights incorporating correlation',ggpt,prb ;
* Verification that the mean, covariance and skew are as advertised for the GQ.

```

```

muchk(v)      = sum(p,prb(p)*ggpt(p,v)) ;
vcvchk(v,vv) = sum(p,prb(p)*(ggpt(p,v)-muchk(v))*(ggpt(p,vv)-muchk(vv))) ;
skwchk(v)     = sum(p,prb(p)*power((ggpt(p,v)-muchk(v))/(vcvchk(v,v)**0.5),3)) ;
display 'These should be equal',mu,muchk ;
display 'These should be equal',vcv,vcvchk ;
display 'These should be zero',skwchk ;
* Now expand to broaden sampling
parameter
  ggpt2(p2,v) Stroud-based GQ points with expansion
  prb2(p2)    Probabilities for Stroud-based GQ points with expansion
  alpha(v)    Expansion factor by variable
  beta(v)     Contraction factor by variable
  kurtchk(v)  Kurtosis check ;
scalar
  gqk         Original kurtosis in the GQ before expansion
  q           Share of probability to expanded quadrature
  contract    Amount by which the contracted quadrature should be shrunk ;
* Generate two copies of Stroud points (noting that they cycle).
ggpt2(p2,v) = 2**0.5*
  (cos(ord(v)*ord(p2)*pi/card(v))$(2*trunc(ord(v)/2) ne ord(v))
  + sin((ord(v)-1)*ord(p2)*pi/card(v))$(2*trunc(ord(v)/2) eq ord(v))) ;
* Deal with final component if v is odd.
ggpt2(p2,v)$(ord(v) eq card(v) and 2*trunc(ord(v)/2) ne ord(v))
  = power(-1,ord(p2)) ;
* Generate the unexpanded probabilities. (note these sum to 2.)
prb2(p2) = 1/card(p) ;
* Now calculate the distribution expansion/contraction factors and probability
* allocation.
gqk      = sum(p2$(ord(p2) le card(p2)/2),
  prb2(p2)*power(sum(v$(ord(v) eq 1),ggpt2(p2,v)),4)) ;
q        = (kurt - gqk)/(expand**4*gqk - 2*expand**2*gqk + kurt) ;
contract = ((expand**2*gqk - kurt)/(gqk*(expand**2 - 1)))*0.5 ;
alpha(v) = expand ;
beta(v)  = contract ;
* Now do the expansion/contraction.
ggpt2(p2,v) = (alpha(v)$(ord(p2) le card(p2)/2)
  + beta(v)$(ord(p2) gt card(p2)/2))*ggpt2(p2,v) ;
prb2(p2) = q*prb2(p2)$(ord(p2) le card(p2)/2)
  + (1-q)*prb2(p2)$(ord(p2) gt card(p2)/2) ;
display 'Expanded GQ before incorporating correlation',ggpt2,prb2 ;
ggpt2(p2,v) = sum(vv,r(vv,v)*ggpt2(p2,vv)) + mu(v) ;
display 'Expanded GQ after incorporating correlation',ggpt2,prb2 ;
* Verification that the mean, covariance and skew are as advertised for the GQ.
muchk(v)      = sum(p2,prb2(p2)*ggpt2(p2,v)) ;
vcvchk(v,vv) = sum(p2,prb2(p2)*(ggpt2(p2,v)-muchk(v))*(ggpt2(p2,vv)-muchk(vv))) ;
skwchk(v)     = sum(p2,prb2(p2)*power((ggpt2(p2,v)-muchk(v))/(vcvchk(v,v)**0.5),3)) ;
kurtchk(v)    = sum(p2,prb2(p2)*power((ggpt2(p2,v)-muchk(v))/(vcvchk(v,v)**0.5),4)) ;
display 'The following should be equal',mu,muchk ;
display 'The following should be equal',vcv,vcvchk ;
display 'The following should be zero',skwchk ;
display 'The scalar kurt and first component of kurtchk should be equal',
  kurt,kurtchk ;

```