



Revista de Economía del Rosario

ISSN: 0123-5362

luis.gutierrez@urosario.edu.co

Universidad del Rosario

Colombia

Dawkins, Christina

Extended Sensitivity Analysis for Applied General Equilibrium Models

Revista de Economía del Rosario, vol. 8, núm. 2, julio-diciembre, 2005, pp. 85-111

Universidad del Rosario

Bogotá, Colombia

Available in: <http://www.redalyc.org/articulo.oa?id=509555104001>

- How to cite
- Complete issue
- More information about this article
- Journal's homepage in redalyc.org

redalyc.org

Scientific Information System

Network of Scientific Journals from Latin America, the Caribbean, Spain and Portugal

Non-profit academic project, developed under the open access initiative

Extended Sensitivity Analysis for Applied General Equilibrium Models

Christina Dawkins*

Ministry of Finance, Government of British Columbia
Victoria, British Columbia, Canada

Recibido: febrero 2005 - Aprobado: abril 2005

Abstract. Previous sensitivity analysis procedures for applied general equilibrium models have focussed on the values of exogenously assigned elasticity parameters, while the calibrated parameters –those that are obtained from combining elasticity information with flow or stock data– have been largely ignored. Calibrated parameters are central to a model’s specification, and uncertainty surrounding their values affects the credibility of the model’s results. This paper introduces and illustrates a calibrated parameter sensitivity analysis (CPSA) which, when combined with previous elasticity sensitivity analysis procedures in an ‘extended sensitivity analysis’, allows modelers to undertake sensitivity analysis over the full set of model parameters.

Key Words: Sensitivity Analysis, Calibration.

JEL classification: C63, C68.

Resumen. Los ejercicios de sensibilidad llevados a cabo hasta el momento acerca de los parámetros utilizados por los modelos de equilibrio general aplicados se han concentrado únicamente en la valoración de las elasticidades, ignorando completamente aquellos obtenidos del mismo proceso de calibración. Dada la importancia de ambos grupos de parámetros para la credibilidad de estos modelos, este artículo presenta e ilustra un procedimiento para el análisis de los mismos.

Palabras Clave: Análisis de sensibilidad, Calibración.

Clasificación JEL: C63, C68.

* I thank Jesús Otero, Carlo Perroni, Geoff Reed, Jeremy Smith, and John Whalley for their comments on earlier drafts. The views in this paper are those of the author only and do not necessarily represent the views of the Ministry of Finance, British Columbia.

Address for correspondence: 957 Maddison St., Victoria, British Columbia, Canada V8S 4C6. Telephone: 1 250 920 9904. E-mail: Christina.Dawkins@gov.bc.ca.

1. Introduction

Previous sensitivity analysis procedures for applied general equilibrium models (Pagan and Shannon, 1985; Pagan and Shannon, 1987; Wigle, 1991; Harrison and Vinod, 1992; Harrison, Jones, Kimbell and Wigle, 1992; DeVuyst and Preckel, 1997) have focussed on the values of exogenously assigned elasticity parameters, while the calibrated parameters—those that are obtained from combining elasticity information with flow or stock data—have been largely ignored.¹ This omission stems partly from the perception that whereas a model's elasticity values are often obtained through informed guesswork and can, therefore, be very uncertain, the calibrated parameter values have a more solid empirical foundation in data. However, the considerable uncertainty surrounding the data used for calibration introduces uncertainty into the calibrated parameter values, making them also candidates for sensitivity analysis. This uncertainty arises through measurement error and is augmented by the consistency adjustments made to the data so that they meet the equilibrium conditions of the model.

The main difficulty for calibrated parameter sensitivity analysis lies in the requirement that the set of calibrated parameters be consistent with an observed “benchmark” equilibrium. Unlike the exogenously specified elasticities, the set of calibrated parameters in an applied general equilibrium model is jointly determined by the benchmark equilibrium data. A given perturbation to one calibrated parameter requires changes in other parameters to ensure that the system remains a benchmark equilibrium but no such realignment of the remaining parameters is unique and, therefore, no single change to the model results can be determined from a given perturbation. Thus, the approach of perturbing individual parameters to observe the effect on model results that has been adopted in previous elasticity-based sensitivity analysis procedures is unsuitable for sensitivity analysis with respect to a model's calibrated parameters.

This paper proposes a calibrated parameter sensitivity analysis procedure (CPSA) that circumvents the joint determination problem by conducting sensitivity analysis over sets of calibrated parameters rather than individual parameter values. Central to CPSA is that a matrix balancing algorithm provides a unique transformation from an unbalanced data matrix into a benchmark equilibrium data set and, therefore, also yields a fixed mapping from a given set of unbalanced data into a set of jointly determined calibrated parameters. Unlike the elements of the benchmark data set, the elements of the unbalanced data matrix are independent variables and can be individually perturbed.

¹ See Mansur and Whalley (1984) for a discussion of calibration.

The CPSA methodology generates a random sample of unadjusted data matrices, each comprised of perturbed elements of the original data matrix. Underlying CPSA is the assumption that the elements of the unbalanced data matrix are observations of stochastically independent random variables for which the modeler can determine *a priori* distributions. Where these random variables are discretely distributed, the support of their joint distribution forms the population of unadjusted matrices from which the modeler samples. If the random variables in the data matrix are continuously distributed, discrete approximations to their distributions are found using the Gaussian quadrature methodology, which specifies discrete approximations that match the lower order moments of the original distributions. The population from which the modeler then samples is given by the support of the ensuing approximate discrete joint distribution. Sampling from the support of a joint distribution allows the modeler to attach a probability of being the true data matrix to each unadjusted matrix in the sample.

A fixed algorithm transforms each sample matrix into a balanced benchmark equilibrium data set. The subsequent benchmark data sets map into corresponding sets of calibrated model parameters that are used to solve the model. Each set of model results is weighted by the probability attached to the unadjusted matrix used in its derivation. From the sample, modelers can then determine confidence intervals for the solution values. Thus, CPSA passes the modeler's knowledge of uncertainty in the unbalanced data, through the calibrated parameters and into a measure of robustness for the model results. In doing so, it completes the framework for reporting the model's sensitivity to its full numerical specification.

This paper is organized as follows: Section 2 elaborates on calibration in applied general equilibrium models and on the problems associated with sensitivity analysis for calibrated parameters. Section 3 presents and illustrates the CPSA methodology using a simple applied general equilibrium model. Section 4 proposes and applies an extended sensitivity analysis procedure in which CPSA is combined with elasticity sensitivity analysis. The application examines the sensitivity of personal tax incidence results in a model of Côte d'Ivoire due to Chia, Wahba, and Whalley (1992) to the parameters calibrated from the consumption expenditure matrix and to selected elasticities. Section 5 concludes with comments on the limitations of the procedure.

2. The Framework for Calibrated Parameter Sensitivity Analysis

An applied general equilibrium model can be written as a system of m simultaneous equations in which a vector of parameters, α , and a vector of exogenous variables, $\mathbf{w}$,

generate a vector of m endogenous variables, $\mathbf{Y}$.² This relationship can be expressed in terms of a mapping, $F: \mathbb{R}^m \rightarrow \mathbb{R}^m$, such that³

$$F(\boldsymbol{\alpha}, \mathbf{w}, \mathbf{Y}) = 0. \quad (1)$$

F can be considered to represent the chosen model structure, and $\boldsymbol{\alpha}$ to summarize its parameterization. To parameterize a given model, modelers must specify values for the vector $\boldsymbol{\alpha}$. Ideally, they should be able to draw on econometric estimates with well defined statistical properties to assign values to these parameters, but in practice the magnitude of the data requirements make such an approach intractable.⁴ Instead, the values for parameters, $\hat{\boldsymbol{\alpha}}$, are inferred from a set of known values for $\mathbf{Y}$ and $\mathbf{w}$, $\hat{\mathbf{Y}}$ and $\hat{\mathbf{w}}$ that solve

$$F(\hat{\boldsymbol{\alpha}}, \hat{\mathbf{w}}, \hat{\mathbf{Y}}) = 0. \quad (2)$$

If the dimensionality of $\boldsymbol{\alpha}$ is greater than m , model parameterization becomes the two stage process discussed in Mansur and Whalley (1984) and Shoven and Whalley (1992). This procedure partitions the vector of parameters $\boldsymbol{\alpha}$ into two subsets: $\boldsymbol{\alpha}_1$, a set of parameters that the modeler is free to specify exogenously, and $\boldsymbol{\alpha}_2$, the set of ‘calibrated’ parameters. If $\hat{\boldsymbol{\alpha}}_1$ is the vector of exogenously specified values for $\hat{\boldsymbol{\alpha}}_1$, then calibration yields values for $\boldsymbol{\alpha}_2$, $\hat{\boldsymbol{\alpha}}_2$, that ensure that for $\hat{\boldsymbol{\alpha}}_1$ and $\hat{\mathbf{w}}$ the model produces $\hat{\mathbf{Y}}$ as a solution. The vector of calibrated parameter values is a function of the exogenously specified parameters and the known solution:

$$\hat{\boldsymbol{\alpha}}_2 = G(\hat{\boldsymbol{\alpha}}_1, \hat{\mathbf{w}}, \hat{\mathbf{Y}}), \quad (3)$$

where G is an implicit function of F .⁵

² In a simple applied general equilibrium framework, the vector $\mathbf{Y}$ includes an income for each agent, a price for each commodity and factor, and an activity level for each production sector. Agents’ factor and commodity endowments are included in the vector $\mathbf{w}$, while policy parameters (such as tax rates), the CES elasticities of substitution, input shares and scale parameters in utility and production functions comprise $\boldsymbol{\alpha}$. Each value in $\mathbf{Y}$ is associated with an equilibrium condition: equilibrium incomes are values that satisfy budget balance constraints for agents; equilibrium prices satisfy market clearing conditions for commodities and factors; equilibrium activity levels satisfy zero profit conditions in production sectors. These equilibrium conditions also form the basis of the more sophisticated structures discussed in Shoven and Whalley (1992), including models with taxes, joint production, nested functions, intermediate demands, decreasing returns to scale production and intertemporal frameworks.

³ A general equilibrium is characterized by a set of complementary slackness conditions where, if equilibrium prices are zero, excess supply can be positive and where, if activity levels are zero, excess profits can be negative. The discussion here is restricted to the case in which prices and activity levels are strictly positive and the equilibrium conditions are satisfied with equality.

⁴ Tractability issues surrounding the econometric estimation of applied general equilibrium models are discussed in Mansur and Whalley (1984). Econometrically derived model parameterizations have been undertaken although the data requirements make such an approach rare. Examples include Clements (1980), Jorgenson, Slesnick and Wilcoxon (1991), and McKittrick (1995).

⁵ Calibration can only be undertaken if the equations in G satisfy the conditions of implicit functions, that is, if the equations of F are continuously differentiable with respect to $\mathbf{Y}$, $\mathbf{w}$, and $\boldsymbol{\alpha}$ and if at $\hat{\mathbf{Y}}$, $\hat{\mathbf{w}}$ and $\hat{\boldsymbol{\alpha}}_1$, the determinant of the Jacobian matrix given by the derivatives of F with respect to $\boldsymbol{\alpha}_2$, is non-zero.

Once values for the calibrated parameters have been found, the vector of model parameter values $\hat{\alpha}$, and the exogenous variables $\hat{w}$, can be used in (2) to solve for Y in a “replication test.” If the solution values for Y are the same as $\hat{Y}$, the calibration procedure has found parameters that are consistent.⁶

Sensitivity analysis is a means of characterising the robustness of model results to uncertainties in this model parameterization process. Typically, the choice of $\hat{\alpha}_1$ is surrounded by a high degree of uncertainty.⁷ In response to this uncertainty, sensitivity analysis procedures that vary the values of α_1 to observe the effect on model results, have been developed.⁸

Ideally, a sensitivity analysis procedure directed towards the vector of calibrated parameters, α_2 , should also examine the link between the parameter values and the model results directly.⁹ This approach, however, is infeasible. The calibrated parameters cannot be individually perturbed because they are jointly determined through the requirement that the known values from which they are derived, $\hat{w}$ and $\hat{Y}$, satisfy the equilibrium conditions of the model; a change to any single calibrated parameter would require other calibrated parameters to change in order to preserve the equilibrium system. Since many such adjustments are possible, no unique effect on model results can be observed from a specific change to a single calibrated parameter.

The joint determination of the calibrated parameters is more evident if the vectors $\hat{w}$ and $\hat{Y}$ are transformed into a square transactions matrix, termed a “benchmark equilibrium data set” (BED). In the BED, a row, representing receipts, and a column, representing outlays, are assigned to each market, production sector, and agent defined in the model. If the Harberger (1962) convention is adopted whereby units trans-

⁶ Policy analysis is then undertaken by perturbing some of the model parameters, computing a new equilibrium and comparing the subsequent vector of endogenous variables to the base case vector. The perturbation of the model parameters captures proposed policy changes such as a change in a tax rate, or the removal of a quota. The counterfactual solution is the measure of the effects of the new policy scenario. It predicts how the economy is likely to respond to the change in the policy regime, while the model’s base case or pre-change solution is the observed outcome from the economy under the existing policy regime.

⁷ The vector α_1 is comprised largely of elasticities of substitution and transformation. The values for these elasticities are obtained, where possible, from literature based econometric estimates but such estimates are scarce and dated. Modelers occasionally undertake their own estimation for these values. Typically, the number of elasticities in an applied general equilibrium model is prohibitively large, and insufficient data exists for their estimation. As a result, modelers often derive elasticity values using ‘best guesses.’

⁸ Elasticities are not the only exogenous parameters for which sensitivity analysis has been undertaken. Rutström (1991), for example, conducts sensitivity analysis over the values of the minimum requirement parameters in a linear expenditure system.

⁹ In so far as the calibrated parameters are functions of the exogenously specified parameters, previous sensitivity analyses capture some of the uncertainty in α_1 . The approach here, however, provides a framework for addressing the *full* uncertainty in the calibrated parameter values.

acted are defined as the quantity that sells for one unit of currency, the equilibrium conditions of the model are reflected in a ‘biproportionality’ condition for the BED: the BED must satisfy the condition that for a square matrix $[x_{ij}]$,

$$\sum_j x_{ij} = \sum_i x_{ij} \quad \forall i = j. \quad (4)$$

Biproportionality ensures that budget balance holds for agents (incomes equal expenditures), sectors make zero profits (sales equal production costs), and because prices are unity, markets clear (quantities demanded equal quantities supplied).

A model’s calibrated parameters are functions of ratios of elements of the BED so that, for example, shares of commodity c where $c = 1, \dots, C$ in the consumption of agent q are calculated from the ratio of x_{cq} to $\sum_{c=1}^C x_{cq}$. Perturbing one calibrated parameter is equivalent to changing a ratio or element in the BED. Changing one off-diagonal element violates the matrix biproportionality condition, and for calibration, the modeler must rebalance the perturbed matrix into a BED.¹⁰ This rebalancing process, however, is not unique.

For example, consider a model with a fixed labour endowment and several production sectors. If the modeler wishes to observe the effects on model results of changing the input share of labour in one production sector, such an input share in at least one other sector would also have to change to maintain the base-period equilibrium condition that the labour market clears. The modeler, however, has no way of determining which of the remaining production sectors should absorb this change. Because several options exist for meeting the model’s consistency requirements, and because each could lead to a different model result, the initial perturbation does not lead to a unique change in the model results. Sensitivity analysis for the input share of labour in a single production sector is, therefore, impossible. A similar argument holds for any individual calibrated parameter value.

Unlike the elements of the BED, however, the *unadjusted* data has no consistency restrictions on the values it can assume. The following section describes a calibrated parameter sensitivity analysis that circumvents the problem of joint determination by perturbing the raw data set from which the calibrated parameters are derived. The procedure allows modelers to observe the effect on model results of varying the entire vector of calibrated parameters, rather than of perturbing individual elements of that vector.

¹⁰ Changing the ratio of a diagonal element of the BED would preserve the biproportionality condition since rows and columns would be affected equally. Diagonal elements in the BED, however, denote transactions from an agent, a market, or a sector to itself. Since such transactions are devoid of behavioral significance in an applied general equilibrium model, the diagonal elements of the BED are defined to be zero.

3. Calibrated Parameter Sensitivity Analysis

Calibrated parameter sensitivity analysis is addressed by turning to the derivation of the BED from a matrix of initial, unbalanced estimates.¹¹ The derivation of the BED falls into the class of matrix balancing problems in which an initial, unbalanced matrix is transformed into a balanced matrix that satisfies a set of linear restrictions and is close to the original matrix under some metric. Many algorithms exist for undertaking such matrix balancing procedures and the modeler must choose among these to derive the BED.¹² In the sensitivity analysis developed here, the adjustment algorithm remains constant, but data used as inputs into that algorithm are perturbed. The calibrated parameter sensitivity analysis maps perturbations in the data through the fixed adjustment algorithm, the resulting BED, the configuration of calibrated parameters, and ultimately into the model results.¹³

3.1. *The CPSA Methodology*

The general approach for the calibrated parameter sensitivity analysis procedure is one of randomized sampling over alternative values of the initial raw data matrix. Randomized sampling, the approach developed for elasticity sensitivity analysis in Harrison and Vinod (1992), avoids the prohibitive computational requirements of unconditional systematic sensitivity analysis discussed in Wigle (1991). It has the additional advantage over the Pagan and Shannon (1985; 1987) sensitivity procedure of providing global rather than local analysis, which strengthens sensitivity results for non-linear models with large uncertainties in the parameter values. The CPSA procedure employs Gaussian quadrature to find a discrete population of matrices from which to sample, following the sensitivity methodology used for elasticity parameters in DeVuyst and Preckel (1997). Unlike the discrete approximation methodology employed in Harrison and Vinod, Gaussian quadrature ensures that the moments of the sampling distribution match those of the underlying distribution.¹⁴

¹¹ The derivation of the matrix of unbalanced estimates itself can be a lengthy process. Modelers typically begin with data from disparate sources of varying quality. The data within each source may also vary in its reliability. Modelers are faced with missing values, conflicting data, and with measurement classifications that are inappropriate to the model.

¹² Günlük-^aenesen and Bates (1988) summarize the general approaches.

¹³ The sensitivity of model results to alternative BEDs has been undertaken elsewhere. Roberts (1994) examines the effects the choice of benchmark year for the BED. Adams and Higgs (1990) argue for the use of a synthetic 'typical' BED rather than one derived from a particular year of record. Wiese (1995) derives two BEDs using alternative accounting assumptions for employer contributions to health insurance and traces the effects of these assumptions on model results. These all argue for particular structures of the BED rather than providing a systematic sensitivity analysis of the type proposed here.

¹⁴ See DeVuyst and Preckel (1997) for a comparison of the two methods.

The CPSA is a procedure in which the data in the initial matrix are considered observations of random variables for which the modeler can determine *a priori* distributions. Where these distributions are continuous, each is approximated by a set of discrete points and associated probabilities. Together, the distributions form a discrete approximation to the joint distribution for the set of variables comprising the data matrix. A sample of unbalanced data sets is drawn randomly from the joint distribution. Each unbalanced data set in the sample is balanced using a constant adjustment algorithm, resulting in a series of BEDs, each of which is then used to calibrate and solve the model.

Let the vector $\mathbf{X}$, with elements $x_j, j = 1, \dots, N$, be the vector of data elements required to calibrate an applied general equilibrium model, so that this vector includes all the data variables for which the modeler must specify values. Let the vector $\bar{\mathbf{A}}$ with elements $\bar{a}_j$ be the best initial estimate of $\mathbf{X}$. The CPSA is a procedure in which the x_j are viewed as random variables, and the $\bar{a}_j$ as realizations from the probability density functions of the x_j . The CPSA methodology is comprised of the following four steps.

Step 1. Specification of the a priori distributions for the x_j

The modeler specifies an *a priori* distribution for each x_j , denoted here by $\{x_j\}$, where $\{x_j\}$ is the probability density function if x_j is a continuous random variable, and the probability mass function if x_j is a discrete random variable. For the purposes of simplicity, the x_j are assumed to be independently distributed. The random variable x_j must have finite moments, and the support of $\{x_j\}$ must be consistent with the model structure. Because the $\bar{a}_j$ are assumed to be the best initial estimates for x_j in the specified distribution, $E(x_j) = \bar{a}_j$, the variance, $E(x_j - \bar{a}_j)^2$, will be informed by the reliability of the data sources as well as the prior modifications undertaken to generate the unadjusted data.

Step 2. Discrete specification of the continuous $\{x_j\}$

In step 2, a discrete approximation is found to each continuous $\{x_j\}$, where the discrete approximation is comprised of K pairs of points, $\hat{a}_j^k, k = 1, \dots, K$, and probabilities p_j^k , such that $\sum_k p_j^k = 1$. A discrete approximation is obtained using Gaussian quadrature. For each x_j , Gaussian quadrature chooses K pairs $(\hat{a}_j^k, p_j^k)$ such that

$$\sum_{k=1}^K p_j^k (\hat{a}_j^k)^l = E(x_j)^l, \quad (5)$$

where $l = 0, 1, \dots, 2K-1$ are the moments of $\{x_j\}$.

The Gaussian quadrature approximation is found as follows (see Miller and Rice, 1983; Preckel and DeVuyst, 1992). For each j , the modeler first solves the linear system

of K equations, where the m th equation, $m = 1, \dots, K$, (and dropping the j subscript) is given by

$$\sum_{l=0}^{K-1} c_l E(x^{l+m-1}) = -E(x)^{K+m-1}, \quad (6)$$

for the coefficients c_l . The solution values for the c_l are then substituted into the polynomial

$$\sum_{l=0}^K c_l E(x)^l = 0, \quad (7)$$

and its roots are found. These roots are the $\hat{a}_j^k$ points for discrete approximation. The final step is to find the probabilities for the $\hat{a}_j^k$. These are given by solving equation (5) for the values of the p_j^k .

Two, three and four point discrete approximations arising from applying Gaussian quadrature to uniform, normal and exponential distributions are given in Miller and Rice (1983), and these provide the discrete approximations employed in the examples of the CPSA that follow.¹⁵ As an example, let the raw data vector be

$$\bar{\mathbf{A}} = [1 \ 2]$$

where element $\bar{\mathbf{a}}_1$ is distributed $N(1, 0.02)$, and element $\bar{\mathbf{a}}_2$ is distributed $N(2, 0.04)$.

A three point Gaussian quadrature would approximate the distribution for $\bar{\mathbf{a}}_1$ by the three point and probability pairs

$$\begin{aligned} (\hat{a}_1^1 &= 0.755, & p_1^1 &= 0.1667) \\ (\hat{a}_1^2 &= 1.000, & p_1^2 &= 0.6666) \\ (\hat{a}_1^3 &= 1.245, & p_1^3 &= 0.1667) \end{aligned}$$

and $\bar{\mathbf{a}}_2$ by

$$\begin{aligned} (\hat{a}_2^1 &= 1.654, & p_2^1 &= 0.1667) \\ (\hat{a}_2^2 &= 2.000, & p_2^2 &= 0.6666) \\ (\hat{a}_2^3 &= 2.346, & p_2^3 &= 0.1667). \end{aligned}$$

Step 3. Construction of a joint distribution

The joint distribution for the elements of $\mathbf{X}$, denoted here by $\{\mathbf{X}\}$, is derived from probability mass function representations of the elements in $\mathbf{X}$. If the *a priori* distributions are discrete, these probability mass functions are simply the $\{x_j\}$, whereas if the *a priori* distributions are continuous, the probability mass functions are given by the Gaussian quadrature discrete approximations to the $\{x_j\}$. Let each x_j have a probability mass function representation with K point and probability pairs. The joint distribution (see Preckel and DeVuyst, 1992) is given by the N^K vector and probability pairs:

¹⁵ Other procedures for undertaking more complicated Gaussian quadratures are cited in DeVuyst and Preckel (1997).

$$\{\mathbf{X}\} = \left(\left[\hat{a}_1^{k_1}, \hat{a}_2^{k_2}, \dots, \hat{a}_N^{k_N} \right], \prod_{j=1}^N p_j^{k_j} \right) \quad \forall \quad k_1 = 1, \dots, K; k_2 = 1, \dots, K; \dots; k_N = 1, \dots, K. \quad (8)$$

Where the x_j are discretely distributed, $\{\mathbf{X}\}$ is the true joint distribution. If $\{\mathbf{X}\}$ is formed from Gaussian quadrature approximations to continuous probability density functions for the x_j , the joint distribution also preserves up to and including the $2K-1$ moments of the original, continuous joint distribution. Because this joint distribution is formed under the assumption of stochastic independence of the x_j , the covariances and higher order cross-moments are zero.¹⁶

The joint distribution for the above example would consist of the nine vector and probability pairs:

$$\begin{aligned} ([\hat{a}_1^1, \hat{a}_2^1] &= [0.755, 1.645], p_1^1 p_2^1 = 0.0278) \\ ([\hat{a}_1^1, \hat{a}_2^2] &= [0.755, 2.000], p_1^1 p_2^2 = 0.1111) \\ ([\hat{a}_1^1, \hat{a}_2^3] &= [0.755, 2.346], p_1^1 p_2^3 = 0.0278) \\ ([\hat{a}_1^2, \hat{a}_2^1] &= [1.000, 1.645], p_1^2 p_2^1 = 0.1111) \\ ([\hat{a}_1^2, \hat{a}_2^2] &= [1.000, 2.000], p_1^2 p_2^2 = 0.4444) \\ ([\hat{a}_1^2, \hat{a}_2^3] &= [1.000, 2.346], p_1^2 p_2^3 = 0.1111) \\ ([\hat{a}_1^3, \hat{a}_2^1] &= [1.245, 1.645], p_1^3 p_2^1 = 0.0278) \\ ([\hat{a}_1^3, \hat{a}_2^2] &= [1.245, 2.000], p_1^3 p_2^2 = 0.1111) \\ ([\hat{a}_1^3, \hat{a}_2^3] &= [1.245, 2.346], p_1^3 p_2^3 = 0.0278). \end{aligned}$$

Step 4. Sampling

If N^K is sufficiently small, a simple random sample is taken from the points in the support of $\{\mathbf{X}\}$ by labeling each point with an integer in the interval $[1, 2, \dots, N^K]$ and drawing a simple random sample from that interval. If N^K is large and a simple random sample cannot be easily generated, random sampling from the support of $\{\mathbf{X}\}$ is achieved using the completely randomized factorial sampling design used in Harrison and Vinod (1992): each point in the sample is generated by randomly selecting its elements from the supports of the discrete representations of the $\{x_j\}$, so that a sample data set is generated by randomly selecting from the values $[\hat{a}_j^1, \hat{a}_j^2, \dots, \hat{a}_j^K]$ for each $j = 1, \dots, N$.

In the above example, the modeler would construct a random perturbed data set by first choosing randomly from the three possible values for the first data element:

¹⁶ The assumption of stochastic independence for such data is supported in applications of the Stone-Byron adjustment algorithm in the social accounting literature, which requires an *a priori* specification of the variance-covariance matrix for a social accounting matrix. For an example, see Crossman (1988). CPSA can, in principle, be extended to the case where elements of $\mathbf{X}$ are jointly distributed. Preckel and DeVuyt (1992) give a Gaussian quadrature joint distribution for the case in which the x_j are joint normally distributed.

$\hat{a}_1^1 = 0.755$, $\hat{a}_1^2 = 1.000$, and $\hat{a}_1^3 = 1.245$, and then choosing randomly from the three possible values for the second data element: $\hat{a}_2^1 = 1.654$, $\hat{a}_2^2 = 2.000$, and $\hat{a}_2^3 = 2.346$.

Let $\mathbf{S}_t$ be a random vector from the support of $\{\mathbf{X}\}$, and let P_t be the probability mass of $\mathbf{S}_t$. The modeler applies the same adjustment algorithm to $\mathbf{S}_t$ as was applied to the original unbalanced data set, $\bar{\mathbf{A}}$, to generate a BED. The BED is then used to calibrate and solve the model. Let $\mathbf{R}_t$ denote the vector of model results arising from the unbalanced data vector $\mathbf{S}_t$.

The process is repeated T times to generate a sample of model results. The sample size must be sufficiently large that the sample moments are consistent estimators of the population moments. To ensure that all vectors in the support of $\{\mathbf{X}\}$ have the same probability of being sampled, sampling is undertaken with replacement, allowing the possibility that the same vector may be drawn more than once. Each $\mathbf{R}_t$, $t = 1, \dots, T$, is weighted by P_t to find the expectations, standard deviations, and confidence intervals for the model results.

3.2. *An Illustration of the CPSA Using a Simple Tax Model*

The Shoven and Whalley (1984) simple $2 \times 2 \times 2$ model, with two consumers (rich and poor), two factors of production (capital and labour), and two commodities (manufactured and non-manufactured goods), is used to illustrate the CPSA. Table 1 summarizes the model structure. The base case version of the model has no taxes. In the counterfactual experiment, a 50 percent tax is levied on the use of capital in the manufacturing sector, resulting in welfare changes for both consumers. These welfare changes, measured by the Hicksian equivalent variation as a proportion of base income, provide the basis for the sensitivity analysis.

Table 1. Structure of the Tax Model Used to Illustrate the CPSA

Production	• Output is produced using capital and labour combined in proportions implied by CES technology in each sector. • The elasticity of substitution in the production of manufactured goods is 2.0 and in that of non-manufactured goods, 0.5. • Share parameters for the CES function are calibrated from the BED.
Consumption	• The utility of each consumer is a CES function of manufactured and non-manufactured goods. • The rich consumer's utility function has an elasticity of substitution of 1.5 and the poor consumer's has one of 0.75. • Share parameters for the CES function are calibrated from the BED.
Endowments	• The rich consumer is endowed with capital and the poor consumer with labour.
Equilibrium Conditions	• Markets clear for all goods and factors. • Zero profits are made in each sector. • Each consumer's expenditure equals his/her income.
Counterfactual	• A 50 percent tax is levied on the use of capital in the production of manufactured goods. • The rich consumer receives 40 percent of tax revenues and the poor consumer receives 60 percent. • Welfare changes for each consumer are measured by equivalent variation as a proportion of base income: $EV^i = (U_c^i - U_b^i) / U_b^i$ where U_b^i is the utility of consumer i, $i = \{\text{rich, poor}\}$, in the base case and U_c^i is utility after the imposition of the tax.

The initial, unbalanced data set used for this model is given in Table 2. It is derived by choosing a random value from a uniform distribution in which the expected value of each data point is the value used in Shoven and Whalley (1984), and the standard deviation is ten percent of the Shoven and Whalley value.¹⁷ The adjustment algorithm used to balance the data is the commonly used constrained quadratic minimization algorithm in which each term is weighted by the unadjusted data value. Thus, if $\bar{a}_j$, $j = 1, \dots, 10$, denotes each of the ten unbalanced, non-zero data values in Table 2, this algorithm finds balanced values, $\bar{q}_j$, such that the expression $\sum_j ((\bar{q}_j - \bar{a}_j) / \bar{a}_j)^2$ is minimized subject to the constraints of the specific experiment.

¹⁷ The choice of a uniform distribution is arbitrary, but the value for the standard deviation is roughly consistent with data. A time series of annual values for value added in manufacturing for the United States was found to have a standard deviation of 10.2 percent. The time series was constructed using annual data for 1970 to 1992 taken from the International Bank for Reconstruction and Development (1993) data base. The series “value added in manufacturing” given in current USD was deflated by the ratio of current USD to constant 1985 USD GDP at factor cost to generate a constant value series.

Table 2. Unbalanced Transactions Values for the Illustrative Tax Model (in units of currency)		
1. Consumption by Households		
	Goods	
	Manufactures	Non-Manufactures
Rich	17.2	25.8
Poor	22.0	52.7
2. Factor Demands by Sector		
	Sectors	
	Manufactures	Non-Manufactures
Capital	7.1	30.4
Labour	34.0	56.6
3. Factor Endowments		
	Factors	
	Capital	Labour
Rich	48.3	0.0
Poor	0.0	59.0
Note: Values were derived as random numbers drawn from uniform distributions with means equal to the Shoven and Whalley (1984) balanced values and standard deviations equal to 10 percent of those balanced values.		
Known Totals (values used in the Shoven and Whalley (1984) model)		
Rich Household's Endowment of Capital		34.3
Poor Household's Endowment of Labour		60.0
Total Demand for Capital		34.3
Total Demand for Labour		60.0
Total Output of Manufactured Goods		34.9
Total Output of Non-Manufactured Goods		59.4

Two sets of experiments, each of which imposes different constraints on the adjustment algorithm, are performed to illustrate the CPSA procedure. The first set assumes that the modeler knows with certainty the aggregate incomes, demands and outputs in the economy, and that they are the values used by Shoven and Whalley. These ‘known’ control totals are given in the final section of Table 2. In this case, the constraints on the adjusted data are that the adjusted endowments equal the known endowments, that the sum of demands for each good equals the known value for the output of each sector, and that the sum of the input demands for each factor equals the value of the known total endowment for each factor.

The second set of experiments assumes that these totals are unknown. In this case, the adjustment constraints are simply that the data meet the equilibrium conditions of the model: markets clear, sectors make zero profits, and the households exhibit budget balance. Both sets of balanced data are given in Table 3, together with the central case welfare results. The robustness of these equivalent variations to uncertainty in the initial data is the focus of the CPSA exercise.

Table 3. Balanced Data and Central Case Model Results for the Illustrative Tax Model

Case 1: Benchmark Values for Data Balanced Using Known Totals	
	Value in units of currency
Rich Consumption of Manufactured Goods	15.1
Rich Consumption of Non-Manufactured Goods	19.2
Poor Consumption of Manufactured Goods	19.8
Poor Consumption of Non-Manufactured Goods	40.2
Input of Capital to the Manufacturing Sector	7.7
Input of Labour to the Manufacturing Sector	27.2
Input of Capital to the Non-Manufacturing Sector	26.6
Input of Labour to the Non-Manufacturing Sector	32.8
Capital Endowment of the Rich Consumer	34.3
Labour Endowment of the Poor Consumer	60.0
Hicksian Equivalent Variation from a 50 % tax on Capital in Manufacturing	
	proportion of base income
Rich Consumer's EV (proportion of base income)	-0.1223
Poor Consumer's EV (proportion of base income)	0.0610
Case 2: Benchmark Values for Data Balanced Using Equilibrium Constraints	
	value in units of currency
Rich Consumption of Manufactured Goods	16.3
Rich Consumption of Non-Manufactured Goods	26.0
Poor Consumption of Manufactured Goods	20.5
Poor Consumption of Non-Manufactured Goods	52.1
Input of Capital to the Manufacturing Sector	8.3
Input of Labour to the Manufacturing Sector	28.5
Input of Capital to the Non-Manufacturing Sector	34.0
Input of Labour to the Non-Manufacturing Sector	44.1
Capital Endowment of the Rich Consumer	42.3
Labour Endowment of the Poor Consumer	72.6
Hicksian Equivalent Variation from a 50% tax on Capital in Manufacturing	
	proportion of base income
Rich Consumer's EV (proportion of base income)	-0.1126
Poor Consumer's EV (proportion of base income)	0.0572

In the first step of the CPSA, uniform distributions are specified for the initial data elements, where the expected value of each data element is given by the central unad-justed data value in Table 2 and its standard deviation is ten percent of that value.¹⁸ The second step of CPSA uses Gaussian quadrature to find discrete approximations to

¹⁸ Although the data generating process for the economy in this example has been specified to generate the unbalanced data set in Table 2, it would be unknown for a modeler undertaking CPSA. Here, the modeler has correctly specified the shape of the distribution, and the proportional magnitudes of the variances, but has the expectation that the error of the initial estimate is zero which, of course, it is not.

these continuous distributions. In this example, they are represented by the three-point approximations given in Table 4, which preserve up to and including the fifth moments of the original distributions.

The discrete approximations to the distributions for the individual data elements are then used to characterize the joint probability density function. The support of this joint distribution is given by the combinations arising when each of the ten data elements assumes one of the three values in the support of its discrete distribution. The result is a set of 10^3 possible configurations, each of which has a probability of being true given by the product of the probabilities of its ten constituent points.

Table 4. 3 Point Discrete Gaussian Quadrature Approximations for the Assumed Continuous Data Distributions in the Illustrative Model
(values in units of currency)

	Expected Value and Bounds of the Assumed Uniform Distribution	1 st Point Pr.: 0.278	2 nd Point Pr.: 0.444	3 rd Point Pr.: 0.278
Rich Consumption of Manufactured Goods	17.2 [14.2, 20.1]	14.8	17.2	19.5
Rich Consumption of Non-Manufactured Goods	25.8 [21.3, 30.3]	22.3	25.8	29.3
Poor Consumption of Manufactured Goods	22.0 [18.2, 25.8]	19.1	22.0	25.0
Poor Consumption of Non-Manufactured Goods	52.7 [43.6, 61.9]	45.7	52.7	59.8
Input of Capital to Manufacturing Sector	7.1 [5.8, 8.3]	6.1	7.1	8.0
Input of Labour to Manufacturing Sector	34.0 [28.1, 39.9]	29.5	34.0	38.6
Input of Capital to Non- Manufacturing Sector	30.4 [25.2, 35.7]	26.3	30.4	24.5
Input of Labour to Non- Manufacturing Sector	56.6 [46.8, 66.4]	49.0	56.6	64.2
Capital Endowment of Rich Consumer	48.3 [39.9, 56.6]	41.8	48.3	54.7
Labour Endowment of Poor Consumer	59.0 [48.8, 69.2]	51.1	59.0	66.9

Note: Pr. denotes probability.

A random, unadjusted data configuration is drawn from this support. This data configuration is derived by choosing randomly from the three point discrete distribution for each data element. For example, the three points in the distribution for the rich consumer's endowment of capital are 41.8, 48.3 and 54.7. The data configuration is constructed by randomly choosing one of these three values, then randomly choosing

one of the three possible values for the poor consumer’s endowment of labour, and so on until a random value has been chosen for each of the ten data elements. The probability associated with the configuration is given by the product of the probabilities of its ten constituent data elements.

This process is repeated to generate a sample of fifty unadjusted data configurations. In the first experiment, each configuration is then adjusted into a BED using the known totals as constraints, and the model is calibrated and solved. Attached to each result is the probability that the configuration used in its derivation is true. The process is the same in the second experiment, except the data are adjusted using only the model’s equilibrium conditions as the balancing constraints.

Table 5 presents the summary statistics that characterize the output of the CPSA sensitivity procedure. It gives the means, standard deviations and 95 percent confidence intervals for the results of both experiments undertaken in this illustrative example. From Table 5, the modeler could conclude that the model results are robust to the uncertainty in the initial data values: the signs of the welfare changes are preserved, and the central case variants lie well inside the 95 percent confidence interval. Where the control totals are known, the standard deviations of the results are lower than where the model’s equilibrium conditions alone provide the underlying adjustment consistency constraints. This result is consistent with the additional information introduced into the system by known totals.

Table 5. CPSA on the Welfare Effects of Imposing a 50 Percent Tax on the Use of Capital in the Manufacturing Sector Hicksian Equivalent Variations Measured as a Proportion of Base Income

Case 1: Data Balanced Using Known Control Totals				
	Central Case ¹	Mean	Standard Deviation	95% Confidence Interval ²
EV Rich	-0.1223	-0.1219	0.0095	[-0.1644, -0.0794]
EV Poor	0.0610	0.0609	0.0050	[0.0385, 0.0833]
Case 2: Data Balanced Using Equilibrium Constraints				
	Central Case ¹	Mean	Standard Deviation	95% Confidence Interval ²
EV Rich	-0.1126	-0.1117	0.0097	[-0.1551,-0.0683]
EV Poor	0.0572	0.0572	0.0055	[0.0326,0.0818]

Note 1: The central case uses the raw data given in Table 3.2.
Note 2: Confidence intervals are derived using Chebychev’s Theorem

4. Extended Sensitivity Analysis

While CPSA allows modelers to undertake sensitivity analysis with respect to the values of the calibrated parameters, the elasticity parameters remain a highly uncertain component of the modeling process. The ‘extended sensitivity analysis’ proposed in this section combines CPSA with the existing elasticity sensitivity analysis methodology advocated in DeVuyst and Preckel (1997), so that modelers can report the sensitivity of their model results to uncertainty in all of the model’s parameters.

4.1. *Extended Sensitivity Analysis Methodology*

The extended sensitivity analysis methodology requires a simple modification to the CPSA procedure described in Section 3. Instead of specifying *a priori* distributions just for the initial data elements, distributions are specified for both the initial data and for the exogenous elasticity parameters. Thus, if N is the number of data elements, and J is the number of exogenously specified parameters, the modeler must specify $(N+J)$ probability distributions. The conditions that apply to the distributions for the data in CPSA also apply to the elasticities in extended sensitivity analysis: they must have finite moments, and the model must be solvable over their supports.

The remaining steps in extended sensitivity analysis follow those in CPSA, except they apply to both the initial data and to the elasticities. Gaussian quadrature is used to construct discrete approximations to the continuous distributions of both the data and the elasticities, and a joint distribution of the data and elasticities is created from these discrete approximations. If K is the number of points in the support of each discrete distribution, the joint probability density function contains $(N+J)^K$ points. Each point in the support of the joint distribution is comprised of an unadjusted data set and a set of elasticity values. Its probability mass is given by multiplying the product of the N probabilities of its unadjusted data values with the product of its J elasticity probabilities.

Random samples, each of which is comprised of an unbalanced data and a set of elasticity values, are then drawn from the joint distribution. As in CPSA, the data component of each random sample is constructed by sequentially choosing a random value for each data element from the K values in the support of its individual discrete distribution. Similarly, the elasticity component of the random sample is derived by sequentially selecting a random value for each elasticity from the K values in the support of its distribution.

The data in each sample are balanced by applying the same adjustment algorithm as was applied to the central case data. Together with the elasticities in the sample, the balanced data are then used to calibrate and solve the model. Means, standard deviations and confidence intervals for the true model results are calculated from the probability weighted sample model results, as in CPSA.

4.2. *An Illustration of Extended Sensitivity Analysis: The Côte d'Ivoire Model*

This extended sensitivity analysis methodology is illustrated using an existing model developed by Chia, Wahba, and Whalley (1992) for tax incidence analysis in Côte d'Ivoire. While the simple Shoven and Whalley example was chosen to illustrate CPSA in Section 3 on the basis that its small dimension offers transparency, the Côte d'Ivoire model is used to illustrate extended sensitivity analysis because it typifies the policy modeling exercises for which such sensitivity analyses are important.

The attempt here at realism is hampered by a lack of knowledge about the unadjusted data and the reliability of the elasticities actually used in the Côte d'Ivoire model. The lack of such information means that several assumptions about the data, elasticities and data adjustments are made in the illustration of extended sensitivity analysis that follows. This obstacle, however, is not unique to the Côte d'Ivoire model. In practically all cases, the information required to undertake extended sensitivity analysis is not available to anybody other than the original modeler, and usually this information has been discarded early in the modeling process. Typically, modelers have no use for unadjusted data; they report only the adjusted version of the data and the central case elasticities. Extended sensitivity analysis, therefore, also has normative implications for modelers: they must maintain a version of the unadjusted data, record their assessment of the reliability of both the unadjusted data and of the elasticities, and report their adjustment procedure in detail if extended sensitivity analysis is ever to be undertaken.

The incidence analysis in Chia, Wahba, and Whalley is undertaken for six taxes/subsidies by replacing each with an equal yield, neutral tax on consumption, and finding the associated welfare change for each of seven household types. The exercise that follows examines the sensitivity of the personal income tax incidence results to uncertainty in the consumption expenditure data and in the values of the consumption and production elasticities of substitution.

The welfare changes on which the Chia, Wahba and Whalley tax incidence results are based, derive from household utility functions that are defined over the consumption of goods and services in the model.¹⁹ The data-based component of the extended

¹⁹ The Côte d'Ivoire model identifies seven socio-economically based household types, each of which receives utility from the consumption of ten goods and services. Incomes derive mainly from capital and labour endowments, as well as interhousehold transfers. Households pay personal income tax and make social security transfers to the government, but also receive income from the government in the form of education and other transfers. The model distinguishes fifteen production sectors, each of which produces output using value added and intermediate goods. All twelve formal sectors pay production taxes and all formal sectors, except the government services sector and the gas, electricity and water sector, also receive subsidies. Eight of the formal production sectors trade internationally, and since Côte d'Ivoire is modeled as a small, open, price-taking economy, exporters

sensitivity analysis is undertaken for this consumption expenditure data. Changes in utility arise directly from changes in consumption levels, but the extent to which a change in the consumption of a particular good translates into a change in utility is determined by the share parameter of that good in the CES utility function. Through calibration, the values of the consumption expenditure data (together with the elasticity of substitution in consumption) determine the values of these share parameters.

The consumption data are assumed to have been obtained from a household survey that reports mean consumption by household type, and Chia, Wahba, and Whalley are assumed to have derived an unbalanced estimate of total consumption expenditure by each household type from scaling the survey data by the number of households in each group. Because the actual household survey data are unknown, they are approximated by the artificially constructed household survey data given in the Appendix in Table A.1. The elements of this artificial data set are randomly drawn from a normal distribution with an expected value equal to the known, adjusted value used by Chia, Wahba, and Whalley, and a standard deviation equal to the proportions of the base case noted in Table A.1.²⁰

The balanced values of the consumption expenditure data for the Côte d'Ivoire model are assumed to have been derived in a two stage process. In the first stage, aggregate values for the total final household consumption of each good consistent with the values for total production, exports, government consumption and intermediate demand would have been found. These values are assumed to be the totals used by Chia *et al.* and are given in the second section of Table A.2. Similarly, aggregate household consumption expenditure would have been specified. These values are presented in the first section of Table A.2, and are also the values used in the original model. In the second stage, an adjustment algorithm would have been applied to the initial, unbalanced consumption data under consistency conditions implied by the aggregate values from the first stage.

The unbalanced data in Table A.1 are scaled by the number of households of each type given in Table A.3, and are then adjusted using the prevalent RAS adjustment algorithm, where the consistency constraints are that i) the total consumption of each

face a perfectly elastic demand function for their output. Traditional exports and exports of primary processed goods are taxed. Imports, used in the production of intermediate goods and in household consumption, are subject to tariffs. The Ivorian price stabilization policy for coffee, cocoa and other exports is captured in the model. In 1986, the benchmark year, the fund experienced a net inflow of revenues and thus the traditional export sector pays into the stabilization fund, while the non-traditional export sector receives only a proportion of those revenues.

²⁰ Chia, Wahba, and Whalley list their primary data sources as the national accounts, the Banque de données financières (the financial database from which balance of payments data was obtained), tax data and household budget survey data, but do not state explicitly which elements of the BED derive from which source. As a result, the sensitivity analysis presented here provides an illustration of the methodology rather than insight into the specific Côte d'Ivoire model results.

good by each household type, summed across household types is equal to the aggregate final household consumption for that good from section 2 of Table A.2, and ii) the sum across goods of total consumption expenditure by household type is equal to the disposable income of each household type, net of interhousehold transfers and savings, given in section 1 of Table A.2.²¹

Together with the original elasticity values, the resulting balanced matrix is used to calibrate and solve the model to obtain incidence results for the removal of the personal income tax. These central case results are given in column (1) of Table 6. Their robustness to uncertainty in the values of the initial household consumption data in Table A.1 as well as to uncertainty in the central case values of selected production and consumption elasticity values is the focus of the subsequent sensitivity analysis.

Table 6. Extended Sensitivity Analysis Results for Personal Income Tax Incidence in a Model of Côte d’Ivoire
Hicksian Equivalent Variation as a Percentage of Benchmark Gross Income

	Central Case (1)	Expected Value (2)	Standard Deviation (3)	95% Confidence Interval (4)
Export Croppers	-0.224	-0.214	0.013	[-0.272, -0.156]
Savannah Croppers	-0.685	-0.690	0.019	[-0.775, -0.605]
Other Food Croppers	-1.600	-1.603	0.020	[-1.692, -1.514]
Government Employees	3.493	3.490	0.009	[3.450, 3.530]
Formal Households	-0.605	-0.607	0.007	[-0.638, -0.576]
Small Businesses	-1.617	-1.614	0.006	[-1.641, -1.587]
Inactive	2.666	2.661	0.015	[2.594, 2.728]

Note: Confidence intervals are derived using Chebychev’s Theorem.

The illustration of extended sensitivity analysis presented here considers the uncertainty in the calibrated parameters given by the consumption expenditure matrix together with uncertainty in the values of three sets of elasticities used in CES functions in the model; the elasticity of substitution of consumption goods in preferences,²²

²¹ Let the unbalanced data be represented in the matrix form of Table A.1, so that the element $\bar{\mathbf{a}}_{ij}$ denotes the expenditure by household j on good i , and let the total consumption of each product be given in the first section of Table A.2, and the total expenditure by household be given in the second section of Table A.2. The RAS algorithm, attributed to Bacharach (1970) is a scaling algorithm in which each row of the initial matrix is scaled by the ratio of the known row total (section 1 of Table A.2) to the actual total. The columns of the ensuing updated matrix are scaled by the ratio of the known column totals (section 2 of Table A.2) to the updated matrix column totals. This process is applied iteratively until the deviation of the updated matrix totals from the control totals is deemed to be sufficiently close to zero.

²² In the central case, these are all 1 implying Cobb-Douglas preferences for households. Sensitivity analysis with respect to this value can therefore also be interpreted as sensitivity over the choice of functional form.

the Armington elasticity of substitution between domestic and imported goods in consumption, and the elasticity of substitution between capital and labour in production.

The elements of the consumption data matrix are assumed to be uniformly distributed where the standard deviation differs by good: rice, construction, and financial services are assumed to be the most reliably reported goods with standard deviations of 10 percent of their base value; transportation and non-financial services, the least reliably reported with standard deviations of 30 percent of their base value; and the data on the remaining goods are assumed to be of intermediate reliability with standard deviations of 20 percent of their base values.

Likewise, the elasticities are also assumed to be uniformly distributed. The bounds for the distributions of the production elasticities are assumed to be the central values ± 0.35 , while those of other elasticities are assumed to be ± 40 percent of their initial values. The central case values and bounds of these elasticities are given in Table A.4.

The uniform distributions for the data and the elasticities are then approximated with three-point discrete approximations obtained from Gaussian quadrature. The support of each approximate distribution has a low, middle and high value. The low value in each approximation derived through Gaussian quadrature is given by the lower bound of the distribution plus 11.27 percent of the range and is associated with a probability of 0.28. The middle value is the lower bound plus 50 percent of the range (the central case value) with a probability of 0.44, and the high value, the upper bound minus 11.27 percent of the range, is associated with a probability of 0.28.

With 31 elasticities and 70 data elements in the consumption expenditure matrix, the support of the discrete joint probability distribution approximation has 3^{101} points. The probability associated with any one of those points is given by the product of the probabilities of its data and elasticity components. A random sample of 500 points is drawn from the joint probability distribution on the assumption that this number is sufficiently large that the sample mean and standard deviation can be used to derive confidence intervals for the model's welfare results.

The sensitivity results are reported in columns (2), (3), and (4) of Table 6. The confidence intervals in Table 6 suggest that if the specified distributions for the data and the elasticities are true, the model results are robust to uncertainty in the parameters, in the sense that the signs of the welfare effects do not change. Furthermore, at the 95 percent confidence interval, the ranking of the incidence of the Ivorian personal income tax among household groups remains the same as in the central case: government employees bear most of the burden of the tax with inactive households assuming a secondary burden. Thus, if the many assumptions made about the source and nature of uncertainty in the data for the Côte d'Ivoire model hold, the central case model results could be confidently presented as inputs into a debate on tax policy reform.

5. Conclusions

Among the criticisms leveled against applied general equilibrium models is one of empirical weakness –model parameterization relies on point observations which lack the statistical rigour of time series data. One means of addressing this criticism is for modelers to incorporate whatever information they do have about the quality of those single observations into the modeling process via sensitivity analyses. Existing sensitivity analysis methodologies, however, are restricted to the set of exogenously specified parameters. In contrast, the CPSA methodology presented here provides measures of the sensitivity of model results to uncertainty in the data used to derive the benchmark data set, and hence, to uncertainty in the values of the calibrated parameters. When the CPSA is combined with existing exogenous parameter sensitivity analysis procedures in an ‘extended sensitivity analysis’, modelers can undertake sensitivity analysis for the full set of model parameters.

Extended parameter sensitivity analysis has been described and implemented for the reconciliation of unbalanced matrices into microconsistent data sets using formal matrix adjustment algorithms. In practice, modelers make limited use of such algorithms. Much of the adjustment to the values found in primary data sources occurs in the *ad hoc* procedures used to derive consistent control totals for submatrices, which are then ‘fine-tuned’ via formal adjustment algorithms. The next challenge is to capture the sensitivity of model results to these larger adjustments. While the approach of the extended sensitivity analysis procedure is sufficiently general to address such issues, it would require parametric representations of those larger adjustments to do so. Since many of these adjustments are *ad hoc*, records of how and why the data has changed are scarce. Paradoxically, broader sensitivity analyses will require modelers to take greater notice of how they adjust their data, but will dispense with the need to describe that process in detail by summarizing the uncertainties in those adjustments via confidence intervals over the model results.

References

- Adams, P., and Higgs, P. (1990). “Calibration of Applied General Equilibrium Models from Synthetic Benchmark Equilibrium Data Sets.” *Economic Record* 66, 110-126.
- Bacharach, M. (1970). *Biproportional Matrices and Input-Output Change*. Cambridge University Press: Cambridge.
- Byron, R. P. (1978). “The Estimation of Large Social Accounting Matrices.” *Journal of the Royal Statistical Society A* 141, 359-367.
- Chia, NC., Wahba, S., Whalley, J. (1992). “A General Equilibrium-Based Social Policy Model for Côte d’Ivoire.” *Poverty and Social Policy Series Paper 2*, The World Bank.

- Clements, K. W. (1980). "A General Equilibrium Econometric Model of an Open Economy." *International Economic Review* 21, 469-488.
- Crossman, P. (1988). "Balancing the Australian National Accounts." *Economic Record* 64, 39-46.
- DeVuyst, E. A., Preckel, P. V. (1997). "Sensitivity Analysis Revisited: A Quadrature-Based Approach." *Journal of Policy Modeling* 19, 175-185.
- Günlük-Senesen, G., Bates, J. M. (1988). "Some Experiments with Methods of Adjusting Unbalanced Data Matrices." *Journal of the Royal Statistical Society A* 151, 473-480.
- Harberger, A. C. (1962). "The Incidence of the Corporation Income Tax." *Journal of Political Economy* 70, 215-240.
- Harrison, G. W., Jones, R. C., Kimbell, L. J., and Wigle, R. M. (1992). "How Robust is Applied General Equilibrium Analysis?" *Journal of Policy Modeling* 15, 99-115.
- Harrison, G. W., and Vinod, H. D. (1992). "The Sensitivity Analysis of Applied General Equilibrium Models: Completely Randomized Factorial Sample Designs." *The Review of Economics and Statistics* 74, 357-362.
- International Bank for Reconstruction and Development, The World Bank (1993). Socio-economic Time-Series Access and Retrieval System, CDROM Data Base Version 3.0.
- Jorgenson, D. (1984). Econometric Methods for Applied General Equilibrium Analysis. In Scarf, H. E., Shoven, J. B. (eds.), *Applied General Equilibrium Analysis*. University of Chicago Press: Chicago.
- Jorgenson, D. W., Slesnick, D. T., Wilcoxon, P. J. (1992). "Carbon Taxes and Economic Welfare." *Brookings Papers on Economic Activity Microeconomics*, 393-454.
- Mansur, A., Whalley, J. (1984). Numerical Specification of Applied General Equilibrium Models: Estimation, Calibration and Data. In Scarf, H. E., Shoven, J. B. (eds.), *Applied General Equilibrium Analysis*. University of Chicago Press: Chicago.
- McKittrick, R. R. (1995). The Econometric Critique of Applied General Equilibrium Modeling: The Role of Parameter Estimation. University of British Columbia, Department of Economics Discussion Paper 95/27.
- Miller, A., Rice, T. (1983). "Discrete Approximations of Probability Distributions." *Management Science* 29, 352-362.
- Pagan, A., Shannon, J. (1985). Sensitivity Analysis for Linearized Computable General Equilibrium Models. In J. Piggott and J. Whalley (eds.), *New Developments in Applied General Equilibrium Analysis*. Cambridge University Press: Cambridge.
- Pagan, A., Shannon, J. (1987). "How Reliable are ORANI Conclusions?" *Economic Record* 63, 33-45.
- Preckel, P. V., DeVuyst, E. (1992). "Efficient Handling of Probability Information for Decision Analysis under Risk." *American Journal of Agricultural Economics*, 655-662.

- Roberts, B. M. (1984). "Calibration Procedure and the Robustness of CGE Models: Simulations with a Model for Poland." *Economics of Planning* 27, 189-210.
- Rutström, E. E. (1991). The Political Economy of Protectionism in Indonesia: A Computable General Equilibrium Analysis. Economic Research Institute, The Stockholm School of Economics: Stockholm.
- Shoven, J. B., Whalley, J. (1992). Applying General Equilibrium. Cambridge University Press: Cambridge.
- Shoven, J. B., Whalley, J. (1984). "Applied General Equilibrium Models of Taxation and International Trade: An Introduction and Survey." *Journal of Economic Literature* 22, 1007-1051.
- Stone, R. (1977). The Development of Economic Data Systems. Foreword to G. Pyatt and J. Round, *Social Accounting for Development Planning*. Cambridge University Press, New York.
- Wiese, A. M. (1995). "On the Construction of the Total Accounts From the U.S. National Income and Product Accounts: How Sensitive Are Applied General Equilibrium Results to Initial Conditions." *Journal of Policy Modeling* 17, 139-162.
- Wigle, R. M. (1991). The Pagan-Shannon Approximation: Unconditional Systematic Sensitivity in Minutes. In Piggott, J., Whalley, J. (eds.), *Applied General Equilibrium*. Physica-Verlag: Heidelberg; 35-49.

Appendix A Extended Sensitivity Analysis Specifications for the Côte d'Ivoire Model

Table A.1. Artificial Household Survey consumption Expenditure Data Annual consumption Expenditure in Millions of CFA Francs								
Consumption Good	Export Croppers	Savannah Croppers	Household Type				Small Businesses	Inactive
			Other Food Croppers	Government Employees	Formal Sector			
Rice	3,260	4,054	1,970	10,000	18,014	8,879	10,635	
Other Subsistence Agricultural Products	35,669	52,039	47,922	35,616	40,742	35,082	27,816	
Traded Agricultural Products	218	0	307	7,815	8,790	4,265	8,796	
Primary Processed	35,651	36,759	41,063	56,167	105,131	58,828	51,389	
Manufactures	20,209	13,847	12,817	40,251	64,624	29,694	21,507	
Electricity, Gas, Water	1,817	2,530	1,659	5,401	5,965	1,951	1,884	
Construction	2,460	1,871	1,784	6,340	5,103	3,956	2,214	
Transport	6,171	0	9,450	34,028	30,515	12,630	34,329	
Financial Services	576	353	385	4,757	2,916	1,214	1,165	
Non-Financial Services	6,392	7,723	6,904	18,904	34,329	1,086	13,319	

Notes: Unbalanced data were derived as random numbers drawn from a normal distribution with mean equal to the balanced value in the Chia, Wahba, and Whalley model and standard deviation as the following: Rice, Construction and Financial Services, 10% of the balanced value; Other Subsistence Agricultural Products, Traded Agricultural Goods, Primary Processed Goods, Manufactures, and Electricity, Gas, Water, 20% of the balanced value; Transport and Non-Financial Services, 30% of the balanced value.

Table A.2. Control Totals Used for the Consumption Expenditure Matrix in the RAS and Stone-Byron Adjustment Algorithms
(Million CFA Francs)

1. Column Control Totals:	
Aggregate Consumption Expenditure by Household Type	
Export Croppers	296,186
Savannah Food Croppers	157,369
Other Food Croppers	185,213
Government Employees	375,647
Formal Sector Households	285,345
Small Businesses	418,530
Inactive	334,426
2. Row Control Totals:	
Aggregate Consumption Expenditure by Product	
Rice	86,484
Other Subsistence Agricultural	516,210
Traded Agricultural Products	50,164
Primary Processed	617,750
Manufactured Goods	341,565
Electricity, Gas, Water	30,864
Construction	37,600
Transport	201,072
Financial Services	17,509
Non-Financial Services	153,498

Table A.3. Number of Households by Type

Export Croppers	2,436,000
Savannah Food Croppers	1,320,000
Other Food Croppers	1,524,000
Government Employees	1,416,000
Formal Sector Households	912,000
Small Businesses	2,580,000
Inactive	1,812,000

Table A.4. Elasticities of Substitution and Bounds Used in the Systematic Sensitivity Analysis

1. Elasticity of Substitution Between Capital and Labour in Production Sector (bounds are central case value ± 0.35)			
	Central Value	Lower Bound	Upper Bound
Food	0.4	0.05	0.75
Traditional Exports	0.4	0.05	0.75
Non-Traditional Exports	0.5	0.15	0.85
Formal Sector Primary Processing	0.8	0.45	1.05
Formal Sector Manufacturing	0.8	0.45	1.05
Gas and Electricity	0.8	0.45	1.05
Transportation	0.5	0.15	0.85
Formal Sector Services	0.8	0.45	1.15
Financial Services	0.8	0.45	1.15
Informal Sector Services	0.9	0.55	1.25
Informal Sector Primary Processing	0.9	0.55	1.25
Informal Sector Manufacturing	0.9	0.55	1.25
Informal Sector Construction	0.4	0.05	0.75
Formal Sector Construction	0.4	0.05	0.75
2. Elasticity of Substitution Between Goods in Utility (bounds are central case value ± 40)			
	Central Value	Lower Bound	Upper Bound
All Households	1	0.6	1.4
3. Elasticity of Substitution Between Imports and Domestic Goods in Consumption (bounds are central case value ± 40)			
	Central Value	Lower Bound	Upper Bound
All goods	2	1.6	2.4