

Bayesian Estimation of Input-Output Tables for Russia

Oleg Lugovoy^{‡§} Andrey Polbin[§] Vladimir Potashnikov[§]

Abstract

The paper gives an overview of a research started in December 2011 on stochastic estimation of input-output data.

Contents

INTRODUCTION	2
CONCEPTUAL FRAMEWORK.....	2
UPDATING IOT WITH BAYESIAN METHODS	3
<i>Monte Carlo Experiment</i>	5
BAYESIAN DISAGGREGATION OF IO TABLES.....	9
<i>Monte-Carlo experiment</i>	9
<i>Measurement errors in observed data</i>	10
COMPUTER IMPLEMENTATION.....	10
EXPERIMENTAL UPDATING AND DISAGGREGATION OF IO TABLES FOR RUSSIA.....	12
UPDATING 2003 IO TABLE FOR RUSSIA (OKONH) UP TO 2010	12
DISAGGREGATION OF 15 TO 69 INDUSTRIES (OKVED) FOR 2005	13
UPDATING OF 15 TO 69 INDUSTRIES (OKVED) UP TO 2012	14
CGE APPLICATION OF STOCHASTIC SAM	14
CONCLUSIONS.....	14
REFERENCES.....	15
APPENDIX	16

[‡] Environmental Defense Fund, USA

[§] The Russian Presidential Academy of National Economy and Public Administration, Moscow, Russia

Introduction

Russian statistical system is under transition for almost two decades from Soviet type Material Product System to SNA. The main transitional break in methodology took place in 2004-2005 when Russian statistical agency “Rosstat” started reporting based on the new definition of economic sectors consistent with NACEs, and stopped reporting using definition of activities inherited from the Soviet statistical system. This methodological break splits all industry level statistics into two periods with little consistency between each other. As a result, Rosstat stopped updating IOT in 2003, based on the only benchmark survey conducted in 1995. The next survey is scheduled for 2012 with expected publication of results in 2015 or later. Official backward estimation is not expected. Therefore Russian statistics will miss IOT at least from 2004 to 2010. Also quality of officially updated IOT from 1996 to 2003 based on 1995 benchmark is questionable.

The paper pursues several goals. First, to apply Bayesian statistics for updating IO tables for 1996-2004 period, i.e. within “old” definition of industries. Second, we apply Monte Carlo Markov Chains (MCMC) methods to disaggregate available in NACE classification SUTs (2004 to 2006, 15 products by 15 activities) into larger 69 by 69 format. **Third**, since the 15x15 SUTs are published by Rosstat as preliminary estimates, they are not fully consistent with other available national accounts data, such as output and value added by industries. To take into account the data uncertainty, we introduce a measurement error for the aggregated io-coefficients. As result, we estimate posterior distribution of input-output coefficients for aggregated and disaggregated matrices, which are consistent with yearly national accounts information. **Forth**, we update the estimated 15x15 and 69x69 matrices for 2007-2011 period using proposed sampling methods. Also we discuss potential ways to improve the estimates imposing additional constrains and links to national macro statistics. **Finally**, we demonstrate an application of the estimated stochastic SUTs on a simple CGE model, analyzing a response to exogenous parameter shock, and sensitivity of the results to the estimated uncertainty in input-output coefficients.

Conceptual framework

In this section we discuss a methodology for updating IOT using Bayesian framework and Monte Carlo Markov Chains method as alternative to RAS and maximum entropy methods.

Updating IOT with Bayesian methods

The basic problem of updating an IO matrix or more generally a SAM can be formulating as follows: find an unknown IO matrix with known sums of rows and columns and a known IO matrix for some previous year. Mathematically speaking, we need to find a matrix A with following restrictions:

$$\begin{aligned} Y &= AX, \\ \sum_i a_{i,j} &= \bar{a}_j, \quad a_{i,j} \geq 0 \end{aligned} \tag{1}$$

where Y, X are known vectors and $\bar{a}_j$ are known sums of columns. And there is a known matrix A^0 from previous year. In classical framework the solution of this problem is usually reduced to finding such matrix A , which minimize some distance function from known matrix A^0 under system of restrictions (1).

In this paper we propose to follow Bayesian methods in estimation of IO tables. Bayesian approach provides a natural and flexible way to incorporate any kind and amount of information either as a prior distribution or observable data. Bayesian methods also provide full density profile on estimated parameters with covariates.

In Bayesian econometrics it is assumed that a researcher has some prior beliefs about estimated parameter vector θ before observing the data, which could be summarized by prior density function $p(\theta)$. When new data comes the researcher update the beliefs about parameters according Bayes theorem:

$$p(\theta|Y) = \frac{L(Y|\theta)p(\theta)}{\int L(Y|\theta)p(\theta)d\theta} \propto L(Y|\theta)p(\theta) \tag{2}$$

where $p(\theta|Y)$ is the posterior density and $L(Y|\theta)$ is the likelihood.

Bayesian inference is easy since the posterior density contain all the information one may need. The researcher could be interested in point estimate, credible set and correlation of parameters and construct it from posterior distribution. In Bayesian framework point parameter estimate is chosen to minimize expected loss function with expectation taken with respect to the posterior distribution. The most common loss function used for Bayesian estimation is the mean square error and the corresponding point parameter estimate is simply the mean of the posterior distribution.

Despite the attractiveness of this method, in the past, Bayesian inference was not so popular due to numerical integration needed in equation (2). In some cases when the prior on θ is conjugate with posterior on θ the posterior density can be obtained analytically. But in more general setup we know posterior density up to normalizing constant. Recently developed computer-intensive sampling methods such as Monte Carlo Markov Chain (MCMC) methods

have revolutionized the application of Bayesian approach. MCMC methods are iterative sampling methods that allow sampling from posterior distribution $p(\theta|Y)$.

Heckelei *et al.* (2008) shortly discuss IOT update with Bayesian method and give an example on artificial data. In this paper authors present a Bayesian alternative to the cross-entropy method for deriving solutions to econometric models represented by undetermined system of equation. In the context of balancing an IO matrix they formulate posterior distribution in the following way:

$$p(z|data) \propto I_{\Psi}(z)p(z) \quad (3)$$

$$z = \text{vec}(A) \quad (4)$$

Equation (4) means vectorization of matrix A . In equation (3) $p(z)$ is some prior distribution, $p(z|data)$ is the posterior distribution and $I_{\Psi}(z)$ is the indicator function that assigns weights of 1 if z satisfies the constraints (1) and 0 otherwise. Authors interpret the indicator function as the likelihood function. As estimate of z Heckelei *et al.* (2008) consider mode of posterior distribution which could be found with some optimization routine. And they illustrate proposed method balancing small 4x4 matrix with independent normal prior taking A^0 as prior mean.

But proposed by Heckelei *et al.* (2008) method actually reduced to minimization yet another distance function from known matrix A^0 . In this paper we concentrate on finding full density profile of posterior distribution with MCMC techniques and applying it to real data.

For convenience we consider equality and inequality constraints of the system of restriction (1) separately. Inequality constrains could be simply introduced in prior distribution by assigning 0 value of density in inadmissible domain. For example one could specify independent truncated normal distribution between 0 and 1 for each parameter of the matrix A . On the other hand if we have certain beliefs about some parameters we could introduce it as additional linear equality constraints. For example it is convenient to assign 0 values for elements of unknown matrix A if corresponding elements in the matrix A^0 are zeros.

At the next step let us consider linear equality constraints and rewrite it in the following form:

$$Bz = T \quad (5)$$

where B is the known matrix, T is the known vector and $z = \text{vec}(A)$ is the unknown vector of estimated parameters. System (5) represents undetermined linear system of equations. And from linear algebra it is known that any solution of linear system (5) could be written in the form:

$$z = \tilde{z} + F^{(1)}\xi^{(1)} \quad (6)$$

where $\tilde{z}$ is the particular solution of the system (5) and $F^{(1)}$ is the fundamental matrix of solutions of homogeneous system $Bz=0$. And any vector $\xi^{(1)}$ solves system (5). The particular solution and the fundamental matrix could be obtained by Gaussian elimination algorithm.

Columns of the fundamental matrix $F^{(1)}=[f_1^{(1)},...,f_k^{(1)}]$ represent basis of the Euclidean subspace. At the next step we could find the basis of the orthogonal complement of this subspace $F^{(2)}=[f_1^{(2)},...,f_{n-k}^{(2)}]$. Let us consider linear transformation of the original space:

$$\begin{bmatrix} \xi^{(1)} \\ \xi^{(2)} \end{bmatrix} = \begin{bmatrix} F^{(1)} & F^{(2)} \end{bmatrix}^{-1} (z - \tilde{z}) \quad (7)$$

In the new system of coordinates prior density has the following form:

$$p_{\xi}(\xi) = \det \begin{bmatrix} F^{(1)} & F^{(2)} \end{bmatrix} p_z(\tilde{z} + F^{(1)}\xi^{(1)} + F^{(2)}\xi^{(2)}) \quad (8)$$

If we specify posterior distribution in the form (3) than posterior distribution will be the conditional distribution of random vector $\xi^{(1)}$ given the zero value of the random vector $\xi^{(2)}$:

$$p_{\xi}(\xi | data) = p_{\xi^{(1)}|\xi^{(2)}}(\xi^{(1)} | \xi^{(2)} = 0) \quad (9)$$

If prior distribution is multivariate normal distribution, posterior distribution of vector $\xi^{(1)}$ is also multivariate normal and we could compute posterior mean and covariance matrix analytically. But it doesn't guarantee nonnegative values of estimated matrix A . In general setup we use truncated prior distribution and know posterior density up to normalizing constant. To conduct inference about parameters we approximate posterior distribution (9) applying MCMC sampling methods. After generating the sample of vectors $\xi^{(1)}$ we could move to initial space using formula (6) and obtain the sample of vectors z , which represents elements of unknown matrix A .

To obtain sample from posterior distribution for examples in this paper we perform the Metropolis sampling algorithm, which is a special case of a broader class of Metropolis-Hasting algorithms, and apply a standard single-site updating scheme. As a proposal density for generating candidate parameter values we use normal distribution for each parameter of vector $\xi^{(1)}$. Standard deviations of the proposal density are iteratively selected during adaptive phase to guarantee acceptance rate for each parameter to be between 35 and 45 percent.

Monte Carlo Experiment

To illustrate the proposed Bayesian method for updating IO matrices in this section we perform Monte Carlo experiments and compare results with the RAS and the cross-entropy methods. In Bayesian framework we assume that there exist several additional known matrices

$A^{-1}, \dots, A^{-T}$ from all previous years and they could provide additional information for the estimation purpose. The main hypothesis of our set up is that incorporating additional information in Bayesian framework about variation of IO coefficients in time could improve estimate of unknown coefficients. This information could shed light on relative stability of IO coefficients.

To perform Bayesian method we need to specify some prior distribution for parameters and we assume independent truncated normal distributions for each IO coefficient and use coefficients of last known IO table as prior mean. To specify standard deviations in prior distribution we estimate standard deviation for each coefficient from all available matrices $A^0, A^{-1}, \dots, A^{-T}$ with a following formula:

$$\sigma_{i,j} = \sqrt{\frac{1}{T} \sum_{t=-T}^0 \left(a_{i,j}^t - \frac{1}{T+1} \sum_{t=-T}^0 a_{i,j}^t \right)^2} \quad (10)$$

So in Monte Carlo simulations we would assume that we don't know the true data generating process and apply the same procedure for all data sets. And for robustness of results we experiment with different stochastic processes for IO coefficients.

We perform 10 000 Monte Carlo experiments. Monte Carlo are carried out by randomly generating (data generating process would be described later) of six 4×4 matrices $A^{-4}, \dots, A^0, A$. Than we generate randomly vector X and compute vector Y from equation (1). In the next step we treat IO matrix A as unknown and estimate it with known vectors X, Y and matrices $A^{-4}, \dots, A^0$. Bayesian method is performed by assuming independent normal distribution for each parameter as prior distribution with A^0 as prior mean and estimated standard deviation for each coefficient from matrices $A^{-4}, \dots, A^0$ as prior standard deviation.

To compute posterior distribution of coefficients we apply Markov chain Monte Carlo (MCMC) method with one chain and sampled length of 50 000 simulation. In each experiment we also estimate matrix A with the RAS and the cross-entropy methods using A^0 as prior information about unknown coefficients. To compare relative performance of methods we need to specify some measure of closeness between true coefficients of matrix A and its estimated values. In contrast to the RAS and the cross-entropy method, Bayesian approach provides full probability profile of estimated parameters and as point estimates of coefficients we choose the mean of posterior distribution. To compare results for each of the methods in each Monte Carlo experiment we compute following statistics:

1. Root mean square error:

$$RMSE = \sqrt{1/16 \sum_{i=1}^4 \sum_{j=1}^4 (a_{ij} - \hat{a}_{ij})^2} \quad (11)$$

2. Mean absolute error:

$$MAE = 1/16 \sum_{i=1}^4 \sum_{j=1}^4 |a_{ij} - \hat{a}_{ij}| \quad (12)$$

3. Mean absolute percentage error:

$$MAPE = 1/16 \sum_{i=1}^4 \sum_{j=1}^4 \left| \frac{a_{ij} - \hat{a}_{ij}}{a_{ij}} \right| \quad (13)$$

And now we describe data generating process. For robustness of results we experiment with stationary and nonstationary processes for IO matrix coefficients. The main assumption in data generating is that there are different variances of error terms for coefficients stochastic processes. Coefficients are generated only for the first three rows of the matrix A^t and the last element is computed as $a_{4,j}^t = 1 - \sum_{i=1}^3 a_{i,j}^t$. And if one of the coefficients falls out the boundary restrictions we treat the current experiment as unsuccessful and through it out. We experiment with following data generating procedures:

1. *Independent IO coefficients.* We assume that each coefficient $a_{i,j}^t$ has normal distribution $a_{i,j}^t \square N(m_{i,j}, \sigma_{i,j}^2)$. At the first step of generating data in one Monte Carlo experiment we generate mean parameters $m_{i,j}$ from uniform distribution $U[0,0.5]$. Than we compute standard deviation as $\sigma_{i,j} = k_{i,j} m_{i,j}$, where $k_{i,j}$ are random variables from uniform distribution $U[0,0.05]$. So we assume that standard deviations of simulated coefficients are not greater than 5% of its value. And at the last step we generate $a_{i,j}^t$ from the distribution $N(m_{i,j}, \sigma_{i,j}^2)$.

2. *AR(1) process for IO coefficients.* We assume that each coefficient $a_{i,j}^t$ are from the following stationary process:

$$a_{i,j}^t = (1 - \rho) m_{i,j} + \rho a_{i,j}^{t-1} + \varepsilon_{i,j}^t, \quad \varepsilon_{i,j}^t \square N(0, \sigma_{i,j}^2) \quad (14)$$

For generating parameters $m_{i,j}$ and $\sigma_{i,j}$ we apply the same procedure as in the previous point. For simplicity we fix parameter $\rho = 0.7$ and the initial conditions for the $a_{i,j}^{-4}$ equal to the unconditional mean $m_{i,j}$. And at the last step we generate $\varepsilon_{i,j}^t$ from the distribution $N(0, \sigma_{i,j}^2)$.

3. *Random walk process for IO coefficients.* We assume that each coefficient $a_{i,j}^t$ are from the random walk process:

$$a_{i,j}^t = a_{i,j}^{t-1} + \varepsilon_{i,j}^t, \quad \varepsilon_{i,j}^t \square N(0, \sigma_{i,j}^2) \quad (15)$$

At the first step we generate the initial conditions for the $a_{i,j}^{-4}$ using the same procedure for generating $m_{i,j}$ in previous points and analogically generate values for standard deviations $\sigma_{i,j}$. And finally generate realizations $\varepsilon_{i,j}^t$ from the distribution $N(0, \sigma_{i,j}^2)$.

Table 1 summarizes relative performance of the Bayesian method in comparison with the RAS and the cross-entropy methods. Results indicate that the Bayesian approach outperforms the competing methods according to the introduced measures of fit. In the case of the independent process Bayesian method wins the RAS and the cross-entropy methods in 70-80 percent of simulations and in the case of random walk wins in 60-70 percent of simulations.

Table 1. Relative performance of Bayesian method.

	Independent process		AR(1) process		Random walk	
	Entropy	RAS	Entropy	RAS	Entropy	RAS
RMSE	72.2%	73.2%	67.3%	67.8%	62.0%	63.0%
MAE	76.0%	77.8%	71.2%	73.1%	66.1%	67.7%
MAPE	76.2%	78.1%	72.4%	74.1%	67.1%	69.0%

Figures 1-3 in appendix demonstrate scatter plots of the RMSE, MAE and MAPE statistics for the Monte Carlo experiments for each data generating process. On the Y-axis is the corresponding statistic obtained using the Bayesian approach. On the X-axis is the statistic obtained using one of the competing methods. All graphs demonstrate that most of the points lie below the solid line at 45 degrees. The Monte Carlo experiments show that Bayesian method with additional information about variation of IO coefficients could potentially improve estimate of unknown IO matrix.

Bayesian disaggregation of IO tables

Задача дезагрегации таблиц во многом аналогична продлению. Математически необходимо найти матрицу A размерностью 69×69 , которая удовлетворяет условию (1) и условию:

$$Kr^T * A * Kr = A_{agr} \quad (16)$$

где Kr – матрица перехода к известной агрегированной матрице A_{agr} .

$$Kr = \begin{cases} kr_{i,j} = 1, & i \in \text{sector } j \\ kr_{i,j} = 0, & i \notin \text{sector } j \end{cases} \quad (16^*)$$

Эта задача может быть представлена как система линейных уравнений, аналогичной системе (5).

В качестве дополнительной предпосылки используется условие, что промежуточный спрос на продукцию сектора не превышает конечный выпуск отрасли.

$$\sum_j a_{i,j} * x_j \leq x_i \quad (17)$$

Условие может быть сформулировано аналогично системе (5) в виде:

$$D * z \leq X \quad (18)$$

где X конечный выпуск.

ГРАФИЧЕСКАЯ ИНТЕРПРЕТАЦИЯ ЧЕМ-ТО ОТЛИЧАЕТСЯ ОТ ПРЕДЫДУЩЕЙ?

Monte-Carlo experiment

Для демонстрации работы метода проведем дезагрегацию сгенерированной случайным образом симметричной таблицы из 15 отраслей ОКВЭД до 69 отраслей. Выбранное число отраслей соответствует реальной задаче дезагрегации, рассматриваемой в следующем разделе. Предполагается, что в наличии имеется информация о совокупном выпуске и промежуточном потреблении всех 69 отраслей. Однако отраслевой импорт и конечный спрос неизвестен. Поэтому мы не обладаем достаточной информацией для определения промежуточного спроса.

Таким образом, задача состоит в семплировании 4761 ячейки «большой» матрицы посредством семплирования вектора констант, так чтобы удовлетворить условиям (5) и (17). В качестве начальных точек выбираются угловые точки, полученные с помощью симплекс метода.

Было проведено $4e6$ итераций, сгенерировано три марковских цепи с общим числом наблюдений 1950. Детальные результаты эксперимента приведены в Приложении. Обобщение результатов приводится в Table 4.

Table 2. Характеристики доверительного интервала (90%).

	$z^{true} < CI$	$z^{true} \in CI$	$CI < z^{true}$
Standart CI	3.02%	95.88%	1.09%
Specific CI	0.34%	97.98%	1.68%

По результатам эксперимента 95.8% истинных значений ячеек попадают в 10% доверительный интервал. Более качественная оценка получается, если использовать специфический доверительный интервал. По данным интервалом подразумевается односвязный интервал с минимальной длиной, который в данном случае, с вероятностью 90% содержит истинное значение. Более подробно см. Table 2. Графическое распределение не сошедшихся ячеек представлено на Figure 2 и Figure 2.

В соответствии со статистикой сходимости цепей, число не сошедшихся ячеек составляет 94% для отдельной Марковской цепи. При этом для 100% ячеек существует сошедшаяся цепь. Максимальный, по модулю коэффициент автокорреляции составляет 0.32. Результаты могут быть, очевидно, улучшены более длительным семплированием либо большим количеством цепей. Однако даже существующие результаты демонстрируют довольно хороший результат.

Measurement errors in observed data

0.5-2 pages

Computer implementation

Представленная выше методика является вычислительно интенсивной. Более того, точность полученных результатов напрямую зависит от количества и длины полученных цепей. Для получения качественных результатов за приемлемый промежуток времени необходимо снизить время необходимое для одной итерации.

Для начала заметим, что перемножение столбца констант на фундаментальную матрицу в уравнении (6) с точки зрения компьютерного времени довольно трудоемкая задача.

Сократить время вычислений можно, если учесть что на каждом шаге меняется только одна константа, и поэтому новые значения столбца z при изменении i —той константы можно вычислить по формуле:

$$z_{new} = z_{last} + \xi_i * F[, i] \quad (19)$$

.Матрица F сильно разрежена, в частности 4761 – N строк составляет единичную матрицу. Где N число линейно независимых ограничений типа равенств в уравнении (5). Для данной задачи N меньше 294.

Путем перенумерации переменных F можно представить в виде:

$$F = \begin{bmatrix} E \\ rF \end{bmatrix} \quad (20)$$

Где E единичная матрица размерность 4761 – N , а rF остаточная матрица размерностью $N \times (4761 - N)$. Тогда для расчета новые значения столбца z при изменении i —той константы можно вычислить по формуле:

$$\begin{cases} z_{new}^j = \delta \xi_i + z_{last}^j, i = j \\ z_{new}^j = z_{last}^j, j \leq (4761 - N) \cap i \neq j \\ z_{new}^j = \delta \xi_i * rF[j - (4761 - N), i] + z_{last}^j, (4761 - N) < j \leq 69 * 69 \end{cases} \quad (21)$$

При этом теперь необходимо проверять удовлетворяют ли ячейки граничным значениям только для $1 + N$ ячеек вместо 4761.

После изменения алгоритма вычисления новых значений z в этом и предыдущем пункте на каждом шаге необходимо производить приблизительно $1+N$ операций сложения, умножения и сравнения, вместо $(4761-N)^2$ (отметим, что это более чем $2e7$) операций сложения, умножения и 4761 операций сравнения.

Для проверки условия (18) можно избежать умножения матрицы D на столбец z при выполнении каждой итерации, если учесть что при изменении i —й константы изменение промежуточного спроса равно:

$$\Delta ID = \Delta(D * z) = \Delta(D * (z_{new} - z_{last})) = \Delta(D * \xi_i * F[, i]) = \delta \xi_i * D * F[, i] \quad (22)$$

При этом значения $D * F[, i]$ можно вычислить заранее, тем самым экономя время необходимое для выполнения одной итерации.

В результате на каждом шаге необходимо производить только 69 операций сложений, умножений и сравнений вместо 69^3 операций сложений, умножений и 69 операций сравнений.

В численных экспериментах проведенных ранее standard deviations of the proposal density are iteratively selected during adaptive phase to guarantee acceptance rate for each parameter to be between 35 and 45 percent. Применительно к данной задаче такой способ может требовать слишком много времени для пробных запусков.

В качестве standard deviations of the proposal density можно выбрать максимальный диапазон изменения константы, если такой существует, перемноженный на какой-то заранее выбранный коэффициент от 0 до 1. Для этого отметим, что при использовании фундаментальной матрицы, представленной в виде (20), все коэффициенты

соответствующие строкам единичной матрицы зависят только от одного столбца. Следовательно, константы ξ_i может изменяться только в диапазоне:

$$z_i^{\min} - z_{part}^i \leq \xi_i \leq z_i^{\max} - z_{part}^i \quad (23)$$

следовательно

$$range(\xi_i) \leq z_i^{\max} - z_i^{\min} \quad (24)$$

Граничные значения ячеек можно найти, используя симплекс метод для минимизации/максимизации соответствующих ячеек при условии (5) и (17).

Метод, описанный выше, хорошо приспособлен к параллельному вычислению. Для того чтобы ускорить МСМС расчет был реализован на многопроцессорном графическом ядре, при помощи технологии CUDA. Для того чтобы показать преимущество модифицированного алгоритма приведено сравнение времени необходимого чтобы обновить все (4761-N) констант, реализованных с помощью различных алгоритмов и программных обеспечений.

Таблица 1. Сравнение времени необходимого для одного прогона всех констант.

№	Реализация	Алгоритм	Время	Операций сложений и умножений	Операций сравнений
1	R	МСМС	более 20 мин	$(4761-N)^2 + 69^3 > 2e7$	4830
2	R	МСМС модифицированный	~ 0.1 с	$70+N < 364$	$70+N < 364$
3	CUDA	МСМС модифицированный	~ 7e-3 с	$70+N < 364$	$70+N < 364$

Experimental updating and disaggregation of IO tables for Russia

Updating 2003 IO table for Russia (OKONH) up to 2010

In this part of the work we illustrate an application of the proposed method to real data. There are available official Russian publications of IO accounts for the period from 1995 to 2003. These accounts used the All-Union Classifier of Economy Branches (OKONH). At different years Rosstat published IO accounts for different number of industries and the longest period of the consistent symmetric IO tables with 22 industries is the period from 1998 to 2003.

In our empirical implication we use symmetric IO tables at basic prices from 1998 to 2002 for the estimation of IO matrix coefficients of 2003. We assume that we know only vectors of total outputs and intermediate demand. Value added we also treat as unknown and estimate corresponding coefficients of IO matrix. We apply the same procedure for estimation as in the Monte Carlo experiments: we assume independent truncated normal distributions for each IO coefficient and use coefficients of 2003 IO table as prior mean. To specify standard deviations in prior distribution we estimate standard deviation for each coefficient on the period from 1998 to 2002. To compute posterior distribution of coefficients we apply Markov chain Monte Carlo (MCMC) method with one chain and sampled length of 300 000 simulation.

Figure 4 in appendix shows scatter plot of the posterior mean of the estimated coefficients in comparison to the true values. All points are concentrated around the solid line at 45 degrees and estimates are close enough to the true values. We also compare the performance of the Bayesian method to the RAS and the cross-entropy methods as in the Monte Carlo experiments. Table 2 shows closeness statistics of the three methods.

Table 3. Results of updating 2003 IO table.

	RMSE	MAE	MAPE	RMSPE
Bayes	0.0074	0.0029	0.1844	0.4502
RAS	0.0067	0.0026	0.1728	0.4604
Entropy	0.0065	0.0026	0.1797	0.4552

where RMSPE is root mean square percentage error:

$$RMSPE = \sqrt{1 / (m * n) \sum_{i=1}^m \sum_{j=1}^n \left(\frac{a_{ij} - \hat{a}_{ij}}{a_{ij}} \right)^2} \quad (16)$$

The main idea of computing the additional closeness statistic is that estimated standard deviations are approximately proportional to the coefficients values. And all other things being equal, Bayesian method should outperform the other methods according to this statistic because of the corresponding specification of the prior distribution. Nevertheless Bayesian estimate demonstrates the poorest results according the other measures of fit. But this result is not surprising because only the 1995 IO accounts were constructed on the basis of the detailed survey method. The other IO accounts based on nonsurvey methods.

Disaggregation of 15 to 69 industries (OKVED) for 2005

Задача дезагрегации проводилась для таблицы использования товаров и услуг в ценах покупателей в 2006 г. Эксперимент является аналогичным проведенному на

смоделированных данных (см. табл 2, 4). Результаты эксперимента представлены в таблице Table 5. Стоит отметить, что качество оценки сильно снизилось, по сравнению с оценкой известной матрицы, а необходимый thin увеличился. Так при thin 5000 максимальный коэффициент автокорреляции равен 0.996, а при thin 5e4 происходит снижение коэффициента автокорреляции до 0.9689.

Сравнение плотности распределения коэффициентов автокорреляции представлено на рисунке Figure 3. При этом только 10% показателей автокорреляции для thin 5000 превышают 0.43, и 5% показатель 0.15 при thin 5e4. Если учесть что в соответствии со статистикой сходимости цепей, число не сошедшихся ячеек составляет 87% для отдельной цепи. При этом для 99.6% ячеек существует сошедшаяся цепь. Следует ожидать, что при дальнейшем, значительном увеличении числа итераций и thin качество оценки улучшится.

Updating of 15 to 69 industries (OKVED) up to 2012

CGE application of stochastic SAM

Conclusions

References

Appendix

Table 4. Параметры и результаты оценки известной матрицы.

Name	Value
iteration	4e6
thin	2000
burn	1E5
success Geweke, %	94%
max ACF	.032
RMSE	1.58e-3
MAE	6.13e-3
MAPE	3.49e0

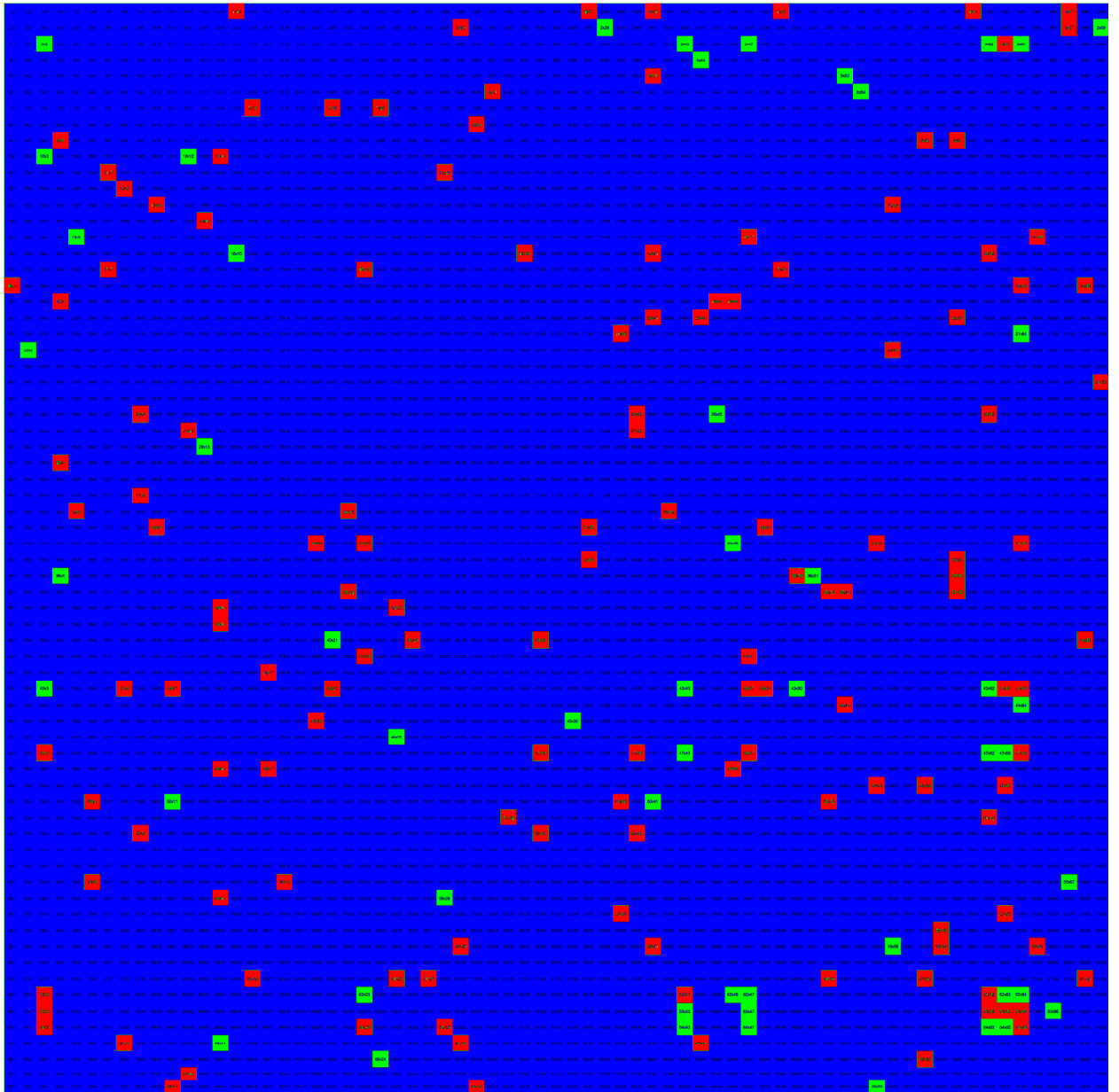


Figure 1. Распределение ячеек со стандартным доверительным интервалом, синий — истинное значение в доверительном интервале, зеленый — истинное значение правее доверительного интервала, красный — истинное значение левее доверительного интервала.

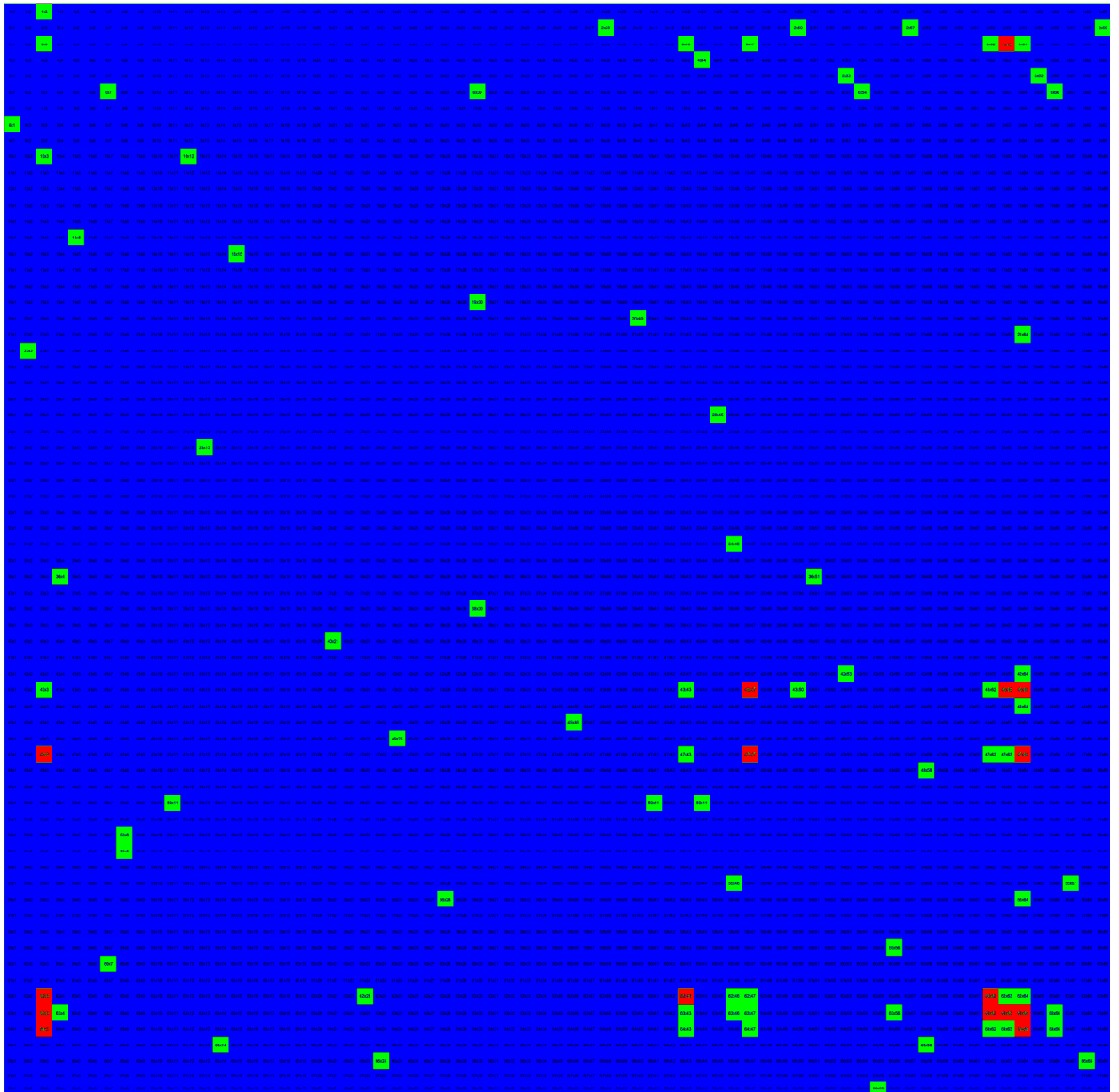


Figure 2. Распределение ячеек с измененным доверительным интервалом (подобранный так чтобы минимизировать ширину доверительного интервала), синий — истинное значение в доверительном интервале, зеленый — истинное значение правее доверительного интервала, красный — истинное значение левее доверительного интервала.

Table 5. Параметры и результаты оценки известной матрицы.

Name	Value
iteration	4e6
thin	5000
burn	1e5
success Geweke, %	87.8%
max ACF	0.996

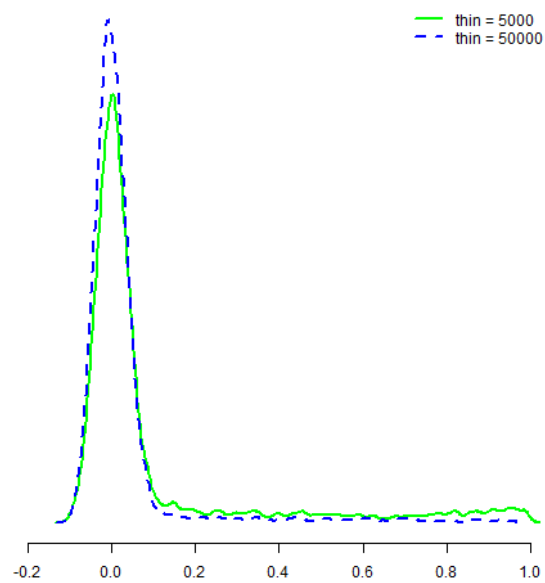


Figure 3. Плотность распределения коэффициентов автокорреляции.

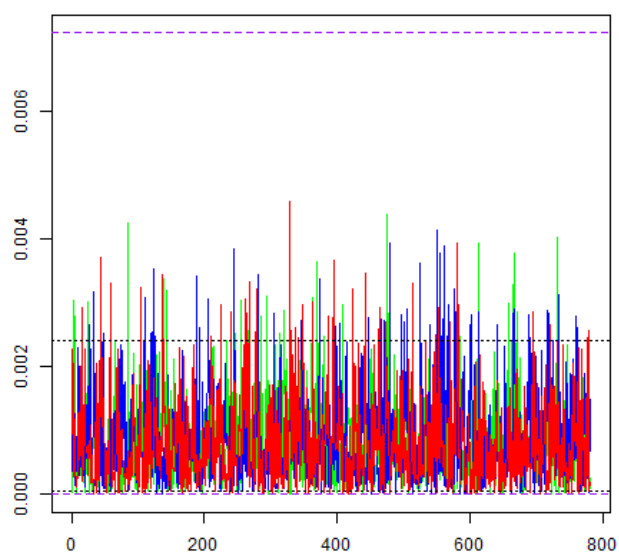


Figure 4. Пример сошедшихся цепей для ячейки 1x7.

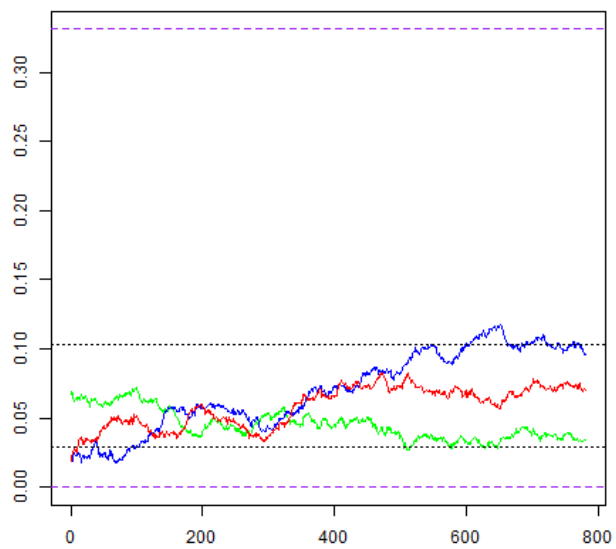


Figure 5. Пример не сошедшихся цепей для ячейки 1x9.