

**Consolidating trade flows and market balances globally
using a Highest Posteriori Density estimator**

- Peter Witzke, EuroCARE Bonn, peter.witzke@eurocare-bonn.de
- Wolfgang Britz, University of Bonn

Abstract

Global trade models require consolidated data with identical bilateral trade flows matrices, with exports and imports matching trade flows and closed market balances. A new methodology was developed and tested to consolidate trade matrices and market balances for more than 200 countries, around 140 products and years 1990-2004. A Bayesian Highest Posteriori Density (HPD) estimator was used in a two step procedure: first trade flow notifications of importers and exporters are rendered consistent and complete, if necessary using back- or forecasting. This step includes a so-called “revealed data quality indicator” which ranks the quality of the trade notifications per country and product. Second the aggregated totals from the estimated trade flows are used as constraints in an estimation giving complete and consistent time series of market balances. This step expresses the precision of information regarding the various elements of these balances in terms of a priori standard deviations entering the HPD estimator.

1 Motivation and overview

In an explorative study for the FAO we investigated the question whether mechanical procedures could play a greater role in the compilation of statistical databases offered by the organization. More specifically a feasible methodology was developed and tested over the last year to consolidate trade matrices and market balances for more than 200 countries, around 140 products and years 1990-2004. While the overall goal was to make best use of the available raw data, the procedure also had to reflect restrictions in hard- and software and potential applicability within FAO. Considerable simplification was achieved, of course at the cost of losing some information, by the decision to consolidate processing of primary products and output of secondary products (product trees) already prior to the estimation with a few important exceptions: Sugar and oils from soybeans, rapeseed, mustard, sunflower, cottonseed, and sesame are represented explicitly and linked to processing of corresponding raw products. Furthermore the current procedure is limited to quantity information. The highly desirable consistent integration of given value information with the quantity data, for example to calculate export unit values perfectly matching to export quantities, has been postponed for the time being.

The final solution is a two step procedure:

1. In the first step, *trade matrices are balanced and completed* using for each trade flow exclusively import and export notifications from the two reporting countries. Completion relies on simple regressions on time or partner notifications and the ‘Hodrick Prescott filter’ for smoothing and interpolating the estimates. The two completed bilateral notifications are merged with weights depending on a ‘revealed

data quality indicator'. After aggregation to trade totals this step gives input data for step two which is not altered anymore.

2. In the second step, the extended version of market balances used in FAO, that is *supply utilization accounts (SUAs)*, are consolidated, i.e. gaps in time series closed and eventual imbalances removed. Technically, the estimation process is broken down to individual countries, but all products and all years are processed simultaneously. This step explicitly minimizes the deviations of final estimates from a priori information. Because this a priori information includes both parts considered firm (such as trade totals from step 1) as well as imprecise information (such as 'reasonable' stock-to-use ratios) the estimation procedure has to be quite flexible.

An earlier attempt to integrate trade matrices and SUAs simultaneously has been abandoned for practical reasons: A fully simultaneous estimation over years, countries and products would have been impossible and an approximate iterative solution turned out very time consuming. The simultaneous solution also would not reflect data availability over the year (first trade, then other data) and the consideration that some division of labour may be useful within FAO.

Similar data consolidation tasks are often tackled using the generalized cross entropy approach (GCE). This technique allows to set ranges for missing data values and provides means of differentiating the reliability of various sources in the exercise (e.g. ROBINSON, CATTANBO, and EL-SAID 2000; BRITZ and WIECK 2002, ROBILLIARD and ROBINSON 2003). In this study we used a Bayesian Highest Posterior Density (HDP) estimator instead. In certain respects this gives a more transparent objective function than CGE and it also has computational advantages as will be explained below.

2 Estimating trade matrices

The methodology to consolidate and complete trade matrices has to face two different problem settings. The standard setting arises if both importer and exporter have notified a certain non zero trade flow for a specific product in a given year but the two notifications do not match. This case requires a reasonable compromise between the conflicting pieces of information.

The second setting is characterized by one or both of the trading partners reporting a zero trade flow for the item and year in question. Under certain conditions it may be useful here to estimate the missing data prior to tackling the standard consolidation problem. This completion relies on simple regressions on time or partner notifications and the ‘Hodrick Prescott filter’ (Hodrick, Prescott 1997) for smoothing and interpolating the estimates. It is useful to address completion before consolidation.

2.1 Identifying and completing gaps for trade flows

Many trade flows emanating from a given country, for a certain product and year are reported zero while others are reported non zero. Only in the case that a country does not report *any* trade of a particular product in a given year, notifications are considered a case of missing, but potentially nonzero data notifications which require completion. The first, very simple, but always possible option for completion is to use simple regressions on time. The second option, regressions on the trading partner notifications will usually yield better results, but may not be possible due to missing degrees of freedom. As mentioned above, a trade flow reported as zero is considered an ordinary observation if the country reports any positive trade flow for the same year and product. For both regressions we calculated R^2 to choose which should be used in the following.

Rather than using the regression results directly it was decided to smooth them using the Hodrick-Prescott-Filter (HP),

$$\min_{\hat{y}_t} obj = wgt \cdot \sum_{1 \leq t \leq T} \left[(y_{t+1}^{HP} - y_t^{HP}) - (y_t^{HP} - y_{t-1}^{HP}) \right]^2 + \sum_t (y_t^{HP} - y_t)^2 \quad (1)$$

where y_t^{HP} are the filtered estimates, y_t are the Hodrick-Prescott-Filter input values and wgt (chosen as unity) is a weighting factor attached to the first component of the objective obj . Where available the HP input values are notifications from the given country, but before the first or after the last year with any notifications one of the two regression options will provide the input values. In case of missing values within a time series no input values will be provided. In these cases the HP Filter provides a convenient interpolation between the next known points.

2.2 Consolidating trade flow information

There may be many reasons for discrepancies between import and export notifications: erroneous or missing data, differences in product classification and notification periods, re-exporting, notified exports having not yet arrived at import destinations or being lost during transport etc. Nonetheless for modelling purposes it is usually imperative to merge this information and for this purpose it is natural to think of some weighted average. Given that we are interested in reliable and accurate results the weights should decline with the likelihood of measurement errors of the information, be it notified or estimated as discussed above.

In line with the SUA estimation described later, we express the uncertainty or lack of accuracy of initial (a priori) information in terms of a standard deviation and specify the weights as reciprocals of these. The starting point for these standard deviations will be the

residual standard deviation $\tau_{r,s,p}^{rep}$ of the trends for trade flow notifications from region r to s for product p and reporter rep ($= ex, im$). When assessing the import and export notifications the lower weight would thus be attached to the more erratic of the two. The reader is reminded that zeros are counted as observations if any trade of that product was notified for the given year.

This standard deviation from the trends is modulated according to a relative transformed revealed data quality indicator for notifications. This indicator is based on the sum of absolute deviations to notifications of trading partners relative to total notified trade (QI), for example for the export notifications of a country r for a certain product p :

$$QI_{r,p}^{ex} = \frac{\sum_{t,s} abs(y_{r,s,p,t}^{ex} - y_{r,s,p,t}^{im})}{\sum_{t,s} y_{r,s,p,t}^{ex}} \quad (2)$$

where $y_{r,s,p,t}^{rep}$ ($rep = ex, im$) is the trade flow from country r to country s reported either by the exporter ($rep = ex$) or by the importer ($rep = im$).

The idea of this indicator is the following: We expect that countries with high reporting quality are close in their notifications to all other high quality reporting countries, thereby receiving a rather small indicator value. Poor quality reporting is characterised by unsystematic errors producing deviations from most other countries in reporting and high values of QI . For further use the quality indicator is transformed and mapped into a $[0,1]$ range such that the 10 most reliable reporters are characterised by a transformed quality indicator $TQI = 1$ and ‘bad’ reporters by $TQI = 0.05$ (minimum to ensure that all countries carry some weight).

The transformed indicator TQI is used to modulate the standard deviation associated with a trade flow notification from region r to s for product p and reporter rep ($= ex, im$) $\sigma_{r,s,p}^{rep}$

around the standard error of a simple trend $\tau_{r,s,p}^{rep}$ for the reported trade flow in question. Increasing the standard deviation for the ‘bad’ reporter as in Equation (3) (for the case of an exporter) reflects the belief that accuracy is low in case of a low TQI (whereas a scaling factor of 0.01 turned out useful in numerical terms).

$$\sigma_{r,s,p}^{ex} = 0.01\tau_{r,s,p}^{ex} \left/ \frac{TQI_{r,p}^{ex}}{0.5(TQI_{r,p}^{ex} + TQI_{s,p}^{im})} \right. \quad (3)$$

If one of the trade flows to be merged is an estimate (interpolation or regression) which is considered inherently inferior to official notifications we are increasing the standard deviation to express a lower confidence in this value. In the case of gaps filled with HP filter interpolations the standard error has been increased to 20 times of the modulated trend standard error from Equation (3) with an additional factor capturing the distance $dist_{r,s,p,t}$ to the next notified year:

$$\sigma_{r,s,p,t}^{ex} = 0.2 \cdot (1 + dist_{r,s,p,t}) \cdot \tau_{r,s,p}^{ex} \left/ \frac{TQI_{r,p}^{ex}}{0.5(TQI_{r,p}^{ex} + TQI_{s,p}^{im})} \right. \quad (4)$$

Because the gap is at least one observation σ_{rspt}^{rep} will be increased at least by $20 \cdot (1+1) = 40$ compared to the standard case with given notifications, all else equal. If e.g. the importer is reporting in the example above, it will dominate the interpolated information from the exporter. (Note that the time index is omitted above in the exposition of the standard case above.)

In case of missing notifications at the beginning or end of a series we use (HP filtered) regression forecasts or backcasts. Because this is again only second choice information the standard error for these estimates has been set to (100% of) the forecast error of the regression used in the particular case, for example:

$$\sigma_{rspt}^{ex} = \left(1 + \frac{1}{Nobs} + \frac{(t - \bar{t})^2}{\sum_{t \in notified} (t - \bar{t})^2} \right) \cdot \tau_{rsp}^{ex} \Bigg/ \frac{TQI_{r,p}^{ex}}{0.5(TQI_{r,p}^{ex} + TQI_{s,p}^{im})} \quad (5)$$

where $\bar{t}$ is the mean year among notified years t and $Nobs$ is the total number of observations (notified years). Note that this is again for the case of an exporter. To preserve comparability with earlier equations only the case for the trend forecast is shown, albeit typically the partner regression will perform better.

Merged trade flows are calculated as weighted averages of two notifications – either given, interpolated or back- or forecasted - with reciprocal standard deviations as defined above serving as weights:

$$\hat{y}_{r,s,p,t} = \left(\frac{y_{r,s,p,t}^{ex}}{\sigma_{r,s,p,t}^{ex}} + \frac{y_{r,s,p,t}^{im}}{\sigma_{r,s,p,t}^{im}} \right) \Bigg/ \left(\frac{1}{\sigma_{r,s,p,t}^{ex}} + \frac{1}{\sigma_{r,s,p,t}^{im}} \right) \quad (6)$$

where $y_{r,s,p,t}^{rep}$ are notified values or HP filtered estimates of reporter rep ($= ex, im$) for a trade from region r to region s in terms of product p , year t and $\sigma_{r,s,p,t}^{rep}$ are the associated standard deviations characterizing the presumed accuracy of the information.

After this consolidation of trade flows, aggregate imports and exports are computed for all regions, products, and years. In the next step of the data consolidation exercise, the resulting import and export totals are treated as fixed, permitting to estimate each country independently from the others. Apart from computational feasibility this may also permit division of labour in case of an implementation within FAO.

The four different cases discussed above are shown again in the table below again.

Table 1: Different cases in the trade matrix problem

Table 1

2.3 Selected results

To illustrate the procedure we will present a few typical results. To begin with, let us look at the easy case where notifications of exporter and importer do not differ a lot. Fortunately this is quite typical for major trade flows.

Figure 1

Figure 1: Consolidation of data on exports of wheat [t] by Country U to Country M

Apart from the year 1992 the two notifications were very close together such that any set of weights would have yielded a consolidated series similar to the one shown in Figure 1. At the 1992 observations we may detect that the consolidated series is closer to the importer notification provided by Country M than to the exporter notification. This is a consequence of the transformed data quality indicator for the exporter $TQI_{U,wheat}^{ex} = 0.56$ whereas Country M is among the 10 best reporters for trade in wheat and therefore $TQI_{M,wheat}^{im} = 1.00$. For 1996 Country M did not report any trade to the FAO and in this case the final estimate is very close to the partner notification.

The next example is on wheat exports from Country U to Country E with higher differences in a number of years. Because $TQI_{E,wheat}^{im} = 0.38$ the resulting consolidated series is closer to the exporter notification (remember $TQI_{U,wheat}^{ex} = 0.56$).

Figure 2

Figure 2: Consolidation of data on exports of wheat [t] by Country U to Country E

The above example is also illuminating for the difference between a zero quantity notified for a particular trade flow and zero notifications for all trade flows which is considered a case of missing data. The zero notification for 1993 by the importer evidently pulls the consolidated series towards zero whereas in 1997 the importer did not report any trade such that the consolidated series essentially equals the exporter notification.

Looking for more difficult examples the case of exports from Country C to Country I will be interesting. For this trade flow the difference in the data quality indicators is larger ($TQI_{C,wheat}^{ex} = 0.62$ and $TQI_{I,wheat}^{im} = 0.25$) such that the results are pulled more strongly towards the exporter notification considered more reliable.

Figure 3

Figure 3: Consolidation of data on exports of wheat [t] by County C to Country I

Furthermore the example shows backcasted and HP filtered estimates because Country I did not report any trade prior to 1997. However, their weighting with the forecast error reduces their influence on the final results to almost zero, because the partner notified official data.

Finally we may look at exports of maize from Country G to Country J which might be considered an odd example because the traded quantities are almost negligible. Furthermore nonzero trade flows are reported only for three years in the whole ‘series’. Nonetheless these are precisely the features which are characteristic for the vast majority of trade flows such that it is interesting to see how our procedure deals with these ‘odd’ cases.

Figure 4

Figure 4: Consolidation of data on exports of maize [t] by Country G to Country J

Apart from the year 2000 where the exporter is reporting exports of 3 tons to Country J there is only one notification by each country: The Country G notifies exports of about 50 tons in 1997 whereas Country J acknowledges receipt of a similar quantity in 1998. This is an example where the explanation of differing notifications with delays and losses during transport is very convincing. Nonetheless we require some consolidation and Figure 4 shows that the consolidated value is quite close to the importer notification. This is again a consequence of the transformed data quality indicator because the exporter's $TQI_{G,maize}^{ex} = 0.06$. Note that this low value is very likely related to the minor importance of maize exports in this country which has higher $TQIs$ (up to 0.80) for other cereals. On the contrary, maize imports are non negligible for the importer which probably explains why they are reported quite accurately in this country ($TQI_{J,maize}^{im} = 0.50$).

3 Estimating Supply Utilisation Accounts (SUAs)

The FAO Supply Utilization Accounts are an extended form of market balances. They try to depict gross production (before netting out losses and on farm consumption) and they include related data such as processing coefficients, harvested areas / herd sizes / raw product inputs. The problem to compile SUAs matching with the consolidated trade matrices involves several challenges. First of all the available hard raw data are usually only the aggregated trade flows and production data. Notified information on human consumption, feed use, processing and other sinks is usually missing. However some information on typical shares of these sinks in total supply (production + imports) is available in FAO and will be incorporated in the current

SUA data. Furthermore plausibility suggests that some series, say human consumption, should be more stable than others, say stock changes. Watching for different variability of series over time requires to tackle this data consolidation problem simultaneously for all years. Furthermore the estimation methodology had to complete any missing elements, as in the trade matrix problem, but for the SUA problem the available priori information differed a lot in its precision. Finally there are a number of constraints linking the various SUA elements together (see below), an aspect that was also missing in the trade matrix problem. The major data constraint ensures that market balances be closed (consistent). In addition we controlled for the behavior of stock levels and yearly changes in total calories intake per capita as well as in an aggregate conversion rate in the livestock sector. Finally, processing to sugar and oil cakes is handled explicitly over constraints.

Technically the SUA estimation was first set up as a Generalised Cross Entropy problem but several computational disadvantages rendered this approach infeasible given the size of the problem. Instead we used a Bayesian Highest Posterior Density (HPD) estimator giving a computationally more attractive objective function. A number of constraints are introduced to impose consistency of identities and compliance with plausible relationships and for the further exposition of the SUA problem will be useful to start exactly here.

3.1 Constraints defining the feasible region

The most important requirement is of course a ‘closed’ balance for any one year and product:

$$\begin{aligned}
& sua_{r,p,Production,t} + sua_{r,p,Imports,t} + sua_{r,p,StockChanges,t} \\
& = sua_{r,p,SeedUse,t} + sua_{r,p,FeedUse,t} + sua_{r,p,Processing,t} + sua_{r,p,HumanConsumption,t} \\
& \quad + sua_{r,p,OtherUses,t} + sua_{r,p,Losses,t} + sua_{r,p,Exports}
\end{aligned} \tag{7}$$

Note that the FAO SUAs define *positive* “stock changes” as a *decrease* in stocks, accordingly, stock changes are booked on the left hand side of the equation.

Seed use and next years’ harvested area are linked through the seeding rate which is planned to be an expert input. However, for the time being it has been specified according to linear trends on the old FAO database. This is quite typical in our application: We investigate the possibility to compile SUAs based on some hard data and certain coefficients. The coefficients are assumed expert inputs by FAO staff. However, to preserve comparability with the current FAO data we like to specify them broadly in line with the current data set. The current data set is used, therefore, to generate reasonable coefficients which may move over time along linear trends. The empirical question is then to see how the results of the new methodology would look like if only the hard data and these coefficients were known. The equation for seed use is thus:

$$sua_{r,p,Seed,t} = sua_{r,p,AreaHarvested,t+1} \overline{sua_{r,p,SeedRate,t}} \quad (8)$$

A similar equation applies to *losses* which are linked through a trend driven loss factor to the sum of production and imports:

$$sua_{r,p,Losses,t} = \overline{sua_{r,p,LossRate,t}} \cdot (sua_{r,p,Production,t} + sua_{r,p,Imports,t}) \quad (9)$$

For *primary products*, production is area harvested respectively herd size times yield respectively carcass weight.

$$sua_{r,p,Production,t} = sua_{r,p,Yield,t} sua_{r,p,Area,t} \quad (10)$$

For *derived products*, production is input times the conversion factor:

$$sua_{r,p,Production,t} = sua_{r,p,ConversionFactor,t} sua_{r,p,Input,t} \quad (11)$$

There are more restrictions related to processing trees but these are largely irrelevant now because the only processed items not consolidated with the raw product are sugar and vegetable oils.

Stock changes (defined as removals from stock) link stock levels of subsequent years:

$$sua_{r,p,Stocks,t} = sua_{r,p,Stocks,t-1} - sua_{r,p,StockChange,t-1} \quad (12)$$

In addition we define a moving three year average of the *stock to use ratio* because we want to incorporate an a priori expectation for this variable:

$$sua_{r,p,StocktoUse,t}^{avg} = \frac{1}{3} \sum_{s=t-1}^{t+1} \frac{sua_{r,p,Stocks,s}}{sua_{r,p,Production,s} + sua_{r,p,Imports,s}} \quad (13)$$

Most variables have a fairly wide permissible range and are only restricted to be nonnegative (apart from stock changes). However for stock levels there are also bounds relative to expected values for *Production* ($\mu_{r,p,Production,t}$), *Imports* ($\mu_{r,p,Imports,t}$), and *Exports* ($\mu_{r,p,Exports,t}$), data which are usually considered ‘hard’:

$$\begin{aligned} sua_{r,p,Stocks,t} &\geq 0.1(\mu_{r,p,Production,t} + \mu_{r,p,Imports,t}) \\ sua_{r,p,Stocks,t} &\leq 1.1(\mu_{r,p,Production,t} + \mu_{r,p,Imports,t} + \mu_{r,p,Exports,t}) \end{aligned} \quad (14)$$

Whereas all of the above constraints refer to variables from a single SUA (in rare cases extending into a product tree) there are two constraints linking together all products, if they turn out binding. These constraints are maximal changes of indicators of the calorie balances for humans and in the livestock sector, assuming that such changes are implausible if they exceed certain limits. Taking human calorie needs, for example, the daily intake per capita $sua_{r,Calories,Food,t}$ is:

$$sua_{r,Calories,Food,t} = \frac{\sum_p sua_{r,p,Food,t} \cdot \overline{cal_p}}{pop_{r,t} \cdot 365} \quad (15)$$

where the calorie contents of food items cal_p are assumed independent of time and region and $pop_{r,t}$ is population in region r and year t . Changes in daily calorie intake of more than 5% from year to year are considered implausible. Due to this constraint (and a similar one related to the livestock sector) the estimation has to occur for all products simultaneously.

3.2 Objective function

In contrast to many other data reconciliation projects this study did not rely on a Generalised Cross Entropy estimator but on a Bayesian Highest Posterior Density approach. This estimator is introduced in Heckelevi, Mittelhammer, Britz (2005) but shall be briefly motivated in this section as well. In general the Bayesian approach to the estimation of a parameter vector β treats these model parameters as stochastic variables with an associated prior density function, $PD(\beta)$, summarizing available prior information on these parameters, before any data y were observed. The Likelihood function, $L(\beta|y)$, represents the information obtained from the data in conjunction with the assumed model. The posterior density, $H(\beta|y)$, is the result of combining prior and data information based on probability rules, Bayes theorem. In our case the ‘parameters’ to be estimated are the SUA elements, **sua**:

$$H(\mathbf{sua}|y) \propto PD(\mathbf{sua}) L(\mathbf{sua}|y) \quad (16)$$

The constraints described in the previous section define a feasible space $\Psi(y)$ possibly conditional on data information y . In our case there is no model generating a nonuniform Likelihood over possible solution values for admissible SUA elements. Instead the likelihood function assumes the form of an indicator function giving a positive constant for feasible points and zero for infeasible points:

$$H(\mathbf{sua}|y) \propto PD(\mathbf{sua}) I_{\mathbf{sua} \in \Psi} \quad (17)$$

Heckelei, Mittelhammer, and Britz (2005) have shown that the prior density $PD(\mathbf{sua})$ may be chosen to yield exactly the same results as the Generalised Cross Entropy estimator. However, other prior densities are equally permissible and the appropriate choice depends on the available a priori information. In our case the a priori information is not already given in the form of ‘support points’ and associated a priori probabilities, but we may be able to specify expected means and standard deviations as a measure of uncertainty. In this case it is natural to use a Normal distribution as a prior density (with a diagonal covariance matrix to ease computational implementation):

$$PD = \prod_{r,p,i,t} \frac{1}{\sigma_{r,p,i,t} \sqrt{2\pi}} \exp\left(-\frac{(sua_{r,p,i,t} - \mu_{r,p,i,t})^2}{2\sigma_{r,p,i,t}^2}\right) \quad (18)$$

where $\mu_{r,p,i,t}$ is the mean and $\sigma_{r,p,i,t}$ the standard deviation of the a priori distribution related to element $sua_{r,p,i,t}$ (region r , product p , SUA item i , year t). In the next section it will be explained in more detail how the means and standard deviations have been derived from our initial information but here it is useful to take them as given. Taking the log of Equation (18) shows that at the bottom line the objective function will be a sum of normalized squared deviations because constants may be omitted in the implementation:

$$LPD = - \sum_{r,p,i,t} \left[\ln(\sigma_{r,p,i,t}) + 0.5 \ln(2\pi) + 0.5 \left(\frac{sua_{r,p,i,t} - \mu_{r,p,i,t}}{\sigma_{r,p,i,t}} \right)^2 \right] \quad (19)$$

Nonetheless some complications will be introduced in the objective function which are motivated by data availability. In the ideal case raw statistical data would be available to specify expected means $\mu_{r,p,i,t}$ for all elements of the SUA. However, a key motivation for this study has been to investigate whether the historically developed data completion routines relying on very specific personal experience could be replaced by a mechanical procedure

which relies only on hard data and a clearly defined set of input parameters or coefficients. The hard statistical data are typically restricted in our case to imports and exports, production, as well as yield or carcass weight along with areas and herd sizes. Data on sinks as human consumption or feed use have been estimated in the past based on reasonable but quite intransparent rules applied by individuals. For these sinks only some coefficients, for example fairly imprecise typical stock-to-use-ratios could be assumed given.

In cases of elements of the market balance where no hard statistical data are given, we like to impose nonetheless some stability of the resulting estimates over time in addition to approximate adherence to the levels defined by the related coefficients (shares, stock to use ratio). Overall this gives four groups of terms contributing to the objective function:

$$obj = -Standard\ term - HP\ filter\ term - Share\ term - Stock\ to\ use\ term \quad (20)$$

The standard term follows directly from Equation (19):

$$Standard\ term = \sum_{r, p, i, t} \left(\frac{suq_{r, p, i, t} - \mu_{r, p, i, t}}{\sigma_{r, p, i, t}} \right)^2 \quad (21)$$

where usually $i \in \{Input, ConversionFactor, Production, StockChanges\}$. The assigned standard deviations will determine the penalty for any deviation of the results from the specified means which are typically derived from the raw data (see the next section).

The HP filter term corresponds to zero a priori means for second differences of so called ‘balancing items’ where raw data, and therefore expected means and standard term are missing:

$$HP \text{ filter term} = \sum_{r,p,i,t} \left(\frac{(sua_{r,p,i,t+1} - sua_{r,p,i,t}) - (sua_{r,p,i,t} - sua_{r,p,i,t-1})}{\sigma_{r,p,i,t}} \right)^2 \quad (22)$$

where balancing items are usually $i \in \{Feed, Processing, Food, OtherUse\}$. The specified standard deviations will determine the extent of the resulting stabilization in the results. In many cases, stock changes are also used as balancing items but they are already stabilised with an expected mean of zero in the standard terms.

The share term in the objective, introduced for sinks, pulls the three year average of the implied shares towards their a priori means:

$$Share \text{ term} = \sum_{r,p,sh,t} \left(\frac{\frac{1}{3} \sum_{s=t-1}^{t+1} \frac{sua_{r,p,sh,s}}{sua_{r,p,Production,s} + sua_{r,p,Imports,s}} - \mu_{r,p,sh,t}}{\sigma_{r,p,sh,t}} \right)^2 \quad (23)$$

where usually $sh \in \{Feed, Processing, Food, OtherUse\}$. The share term expresses, depending on the standard deviation $\sigma_{r,p,sh,t}$ an imprecise expectation about the mean importance of the item in total use.

Finally the stock to use term tends to stabilise the three year average stock to use ratio from Equation (13):

$$Stock \text{ to use term} = \sum_{r,p,t} \left(\frac{sua_{r,p,StocktoUse,t}^{avg} - \mu_{r,p,StocktoUse,t}}{\sigma_{r,p,StocktoUse,t}} \right)^2 \quad (24)$$

this term closely resembles the share term.

The share term and stock to use term thus determine the long term development of certain time series without hard statistical data, whereas the HP-Filter term steers which elements of the market balance mainly adjust to fluctuations of production plus net trade .

3.3 Specifying the a priori information

First of all remember that trade totals are fixed for the SUA consolidation step. For other standards items with raw statistical information, usually $i \in \{\text{Input (or hectares / herd sizes, Conversion Factor / Yield, Production)}\}$ the expected means will be either the given data y , usually from notifications, or HP filtered trend forecasts or interpolations:

$$\mu_{r,p,i,t} = \begin{cases} y_{r,p,i,t} & \text{if } y_{r,p,i,t} \neq 0 \\ sua_{r,p,i,t}^{HP} & \text{if } y_{r,p,i,t} = 0 \end{cases} \quad (25)$$

The associated standard deviations are derived from the standard errors of simple trends which are calculated on all variables:

$$\sigma_{r,p,i,t} = \begin{cases} 0.01 \cdot \tau_{r,p,i,t} & \text{if } y_{r,p,i,t} > 0 \\ 0.1 \cdot \tau_{r,p,i,t} \cdot (1 + dist_{r,p,i,t}) & \text{if interpolated} \\ \tau_{r,p,i,t} \cdot \left(1 + \frac{1}{Nobs} + \frac{(t-\bar{t})^2}{\sum_{t \in notified} (t-\bar{t})^2} \right) & \text{if trend forecast} \end{cases} \quad (26)$$

The reasoning behind these assignments corresponds to the specification of the weights in section 2.2 on the trade matrix problem. It may be mentioned that data on notified production are maintained if this does not cause an infeasibility.

For share items, usually $i \in \{Feed, Processing, Food, OtherUse\}$ the level will not be steered over the standard term in the objective, but through the combination of the share term and the HP filter term. As a consequence there is no need to specify an expected mean for the level of these items, but a priori expectations and standard deviations are required for the associated shares on total market appearances.

$$\begin{aligned}\mu_{r,p,sink\ share,t} &= sua_{r,p,sink\ share,t}^{Trend} \\ \sigma_{r,p,sink\ share,t} &= \tau_{r,p,sink\ share,t}\end{aligned}\tag{27}$$

where $sink \in \{Feed, Processing, Human\ consumption, Other\ use, Residual\}$, $\mu_{r,p,sinkshare,t}$ is the expected mean of a share item, $\sigma_{r,p,sinkshare,t}$ the corresponding standard deviation, $sua_{r,p,sink\ share,t}^{Trend}$ the trend forecast, and $\tau_{r,p,sinkshare,t}$ the estimated standard error of the trend line. Specification of expected means for these shares according to trends is a preliminary solution. In the future these may also come from expert input.

As mentioned above there is no need to specify expected means for the associated levels of the sink items, but standard deviations are needed in the HP filter term of the objective. Given that the existing SUA estimates for these sinks were to be replaced it would have been inconsistent to derive the required standard deviations from trend estimates on the old SUA data. Instead they have been specified as a certain percentage of the sum of the standard deviations of production and imports (here derived from the underlying trade flow information):

$$\sigma_{r,p,sink,t} = \mu_{r,p,sink\ share,t} \cdot (\sigma_{r,p,Production,t} + \sigma_{r,p,Imports,t}) \cdot FAC_{sink}\tag{28}$$

where $FAC_{sink} = 0.1$ for *Food*, $FAC_{sink} = 0.4$ for *Stock Changes*, and $FAC_{sink} = 0.3$ for other share items to direct any fluctuation originating in the hard raw data in particular to stock changes and least to food.

Finally the *Stock to use term* in the objective also requires corresponding means and standard deviations. For the time being we set them to $\mu_{r,p,StocktoUse,t} = 0.3$ for all regions and products and $\sigma_{r,p,StocktoUse,t} = 0.0005$. The latter value may be used to permit or prevent medium term fluctuations of the Stock-to-use ratio.

It may appear that the latter choices on appropriate standards deviations for the HP filter term and the stock to use term are completely arbitrary and that the approach is flawed thereby in general. We have at least three answers to this objection.

First of all note that a Cross Entropy approach requires even more arbitrary choices which are critical to the results, namely the support points and the associated a priori probabilities. If the a priori information is not already ‘given’ in that format, but has to be ‘translated’ into it, the Cross Entropy approach involves even more ‘arbitrary choices’.

Note second that our a priori density clearly expresses our prior degrees of belief in certain results, decreasing in squares as the results move away from the a priori means. In a Cross Entropy or Maximum entropy framework, this prior weighting scheme is implied by the specification of supports, a priori probabilities *and the entropy objective*. It is well known, for example, that the highest objective value in a Cross Entropy framework is obtained if the result equals the a priori mean (as in the HPD case), but the penalty for deviations from this mean is less transparent. Preckel (2001) had to show that this penalty is close to a quadratic function in case of two supports with equal probabilities.

Note finally that the specification of high or low standard deviations for variables to be estimated involves already a considerable degree of standardization compared to the current situation where different FAO officials, responsible for particular regions and/or products, may have developed their own routines which differ a lot. The proposed methodology is

supposed to impose some common standard, but nonetheless to have sufficient input options to tailor the procedure to 200 countries and 140 products. The procedure should be amenable therefore to fine tuning incorporating whatever information FAO staff considers useful.

The different settings are shown in the table below again

Table 2: Different cases in the SUA problem

Table 2

3.4 *Some remarks on computational aspects*

Given the size of the problem computing time is crucial in this project which motivated the use of the HDP estimator as well. As shown above, the objective function is quadratic in the parameters to estimate which speeds up solution time compared to a CGE framework. Secondly, the CGE framework requires additional variables in case of more than 2 supports, namely the probabilities attached to the supports along with adding up constraints of the probabilities to unity. Additional constraints are usually introduced to define the parameters as the multiplication of supports and probabilities. Not at least, the zero to unity bounds on the probabilities will enlarge the estimation framework. And finally, the supports define the admissible range of the parameters to estimate which may often require to introduce large outer supports with very low a priori probabilities, which may in turn provoke the introduction of several inner supports to get a reasonable shape of the penalty function.

Despite the use of the HDP estimator, the whole estimation (data preparation, trade matrix estimation, SUA problems and some processing of results) takes about 48 hours but this still required a few ‘tricks’. These include a frequent release of unneeded memory, the

introduction of a quick ‘prestep’, a stepwise release of bounds in case of infeasibilites, and some variations of solver options.

The ‘prestep’ relies on the fact that the solution of single product trees are independent if the restrictions on maximal calorie changes for humans (see Equation (15)) and livestock are not binding. As a consequence the ‘prestep’ involves the solution of the SUA problem for a given country tree by tree, in a loop and ignoring these restrictions which is much quicker than the simultaneous solution for all products. If the results turn out to be in line with the restrictions the country is ready, otherwise, the simultaneous solution has to be tackled.

Finally it is useful to mention that bounds are widened in case of infeasibilities. In the case of production and explicit expert inputs, for example, this feature is used in an attempt to fix these variables. Only if this attempt turns out infeasible the associated bounds are widened.

3.5 Selected results

As an example for the methodology applied in this study the Figure 5 presents original Data and new estimates of important elements of the SUA for wheat in Country C, a significant exporter serving as an example in Section 2.3 already.

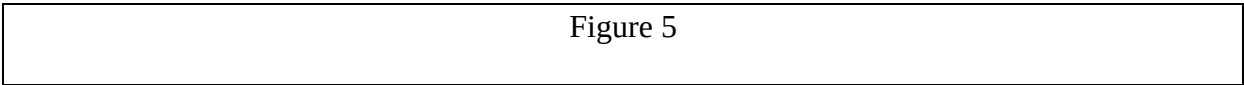


Figure 5: Original data and new estimates on selected SUA items [1000 t] for wheat in Country C

It may be recognised that the new results closely resemble those from the original SUA data, even though the uncertain data (stock changes, sinks) have been generated with an entirely different methodology. This is partly due to the fact that the key input data on production and trade have hardly changed. As mentioned above our estimation procedure tries to maintain the production data, if possible, and the exports are not significantly revised in our consolidation

of trade data (whereas imports are minor). An exception occurs in 1992 where the revised exports are about 5 mio tons lower than in the unconsolidated data. The additional quantities remaining on the domestic market are estimated to partly replenish stocks such that stock changes were strongly *negative* whereas the originally positive stock changes implied a decrease of stocks, see Equation (12). In addition feed use and food (less important and not depicted) are estimated higher than in the original data in the early 90ies. Both the original data as well as the new estimates do not use national data on feed use which provide an external check therefore. These data show that in this case feed use is better estimated by the original data than by the new estimates in the early 90ies. Apparently this balancing item should have been stabilised to a greater extent, at least in this example.

The next example is for Country Z. On average the country is a net exporter of Maize, but in case of bad harvests such as 1992 or 2003, the country needs significant imports. The new consolidated estimates of maize imports in these years are about 300.000 t smaller than the original data, requiring adjustments in some market positions. A similar adjustment need is implied by an upward correction of exports (not depicted) by about 300.000 t in 1994. Figure 6 shows that the most important sink, in this case food, is again absorbing an implausible portion of these “shocks”. The stock to use ratio is moderately fluctuating around 0.3 (as $\mu_{r,p,StocktoUse,t} = 0.3$) but more leeway for stock changes would permit greater stability for consumption. This stabilisation has been achieved in the original data, but the necessary stock changes are somewhat implausible, because their mean is close to + 200.000 t, implying that huge initial stocks were depleted over the decade.

Figure 6

Figure 6: Original data and new estimates on selected SUA items [1000 t] for maize in Country Z

The final example will also be for the SUA on maize in Country J. Production is negligible and maize is mainly stemming from imports. In 1992 there was a significant downward correction of imports as a result of the trade data consolidation. Nonetheless we have some increase in stocks in this year, whereas the original data generation process completely neglected this buffer. As a consequence the major sinks had to adjust to availability. Up to 1998 this was exclusively feed which strongly fluctuated in line with imports according to the original data. From 1999 onwards a significant but highly volatile share of availability was used for food consumption such that feed use fluctuated opposite to food.

Figure 7

Figure 7: Original data and new estimates on selected SUA items [1000 t] for maize in Country J

The new estimates differ a lot. First of all feed use fluctuates smoothly such that it does not follow the import peak in 1992 (releasing quantities for replenishment of stocks). Human consumption is almost stabilised to a straight line. Note that this follows from our preliminary solution to estimate the expected values for shares of sinks in total availability according to linear trends and from the HP Filter term favouring developments along straight lines. For feed use, the stabilisation may still be insufficient even though it appears to be more convincing than the original data.

4 Concluding remarks

A methodology has been developed to compile a complete and consolidated set of trade matrices together with matching domestic market balances for the products in question. The

first results on food products and a number of agricultural inputs are currently evaluated within FAO. This evaluation will determine which use is made of the procedure in the future.

Based on the first tests the approach has been shown to be sufficiently robust to work on more than 200 countries, around 140 products and years 1990-2004. Compared to the original FAO data the new results benefit from equation based safeguards to meet the constraints built into the procedure (closed balances, reasonable stock changes) whereas a certain percentage of erroneous or inconsistent data cannot be avoided if a large number of individuals shares responsibility in the data generating process.

Even though the new results will meet all constraints the generated time series on important sinks such as feed and food use are sometimes fluctuating more than appears plausible. The shape of the resulting series is strongly dependent on certain parameters ($\mu_{r,p,item,t}$, $\sigma_{r,p,item,t}$) which have been specified according to reasonable rules but so far neglecting national statistical information beyond that incorporated in FAO data. The empirical foundation could be strengthened if reliable data on market balances from selected counties and products were used to calibrate the key parameters. Subsequently these results could be applied, in modified form, to other countries and products.

5 References

- Heckelei, T., Mittelhammer, R., Britz, W. (2005): A Bayesian Alternative to Generalized Cross Entropy Solutions for Underdetermined Models, Paper presented at the 89th EAAE Seminar, February 3-5, Parma, Italy.
- Hodrick, R.J. & E.C. Prescott (1997): Postwar U.S. Business Cycles: An Empirical Investigation, *Journal of Money, Credit, and Banking*, 29, 1–16.
- Preckel (2001): Least Squares and Entropy: A Penalty Function Perspective. *American Journal of Agricultural Economics*, 83(2): 366-377.

Table 1: Different cases in the SUA problem

Element	Steered by	A priori knowledge
Production, yield/carcass weight, areas/herds	Standard term	μ : Given data or trend results; σ : based on error of trend line
Sinks excluding stock changes	Combination of (1) share term and (2) HP-Filter term	(1) μ : trend on share, σ : based on error of trend line on share (2) μ : trend on share times times expected mean of production plus imports; σ : subjective setting
Stock changes	Combination of (1) Stock-to-use term (2) Standard term	(1) $\mu = 0.3$ $\sigma = 0.0005$ (2) $\mu = 0.0$, σ large

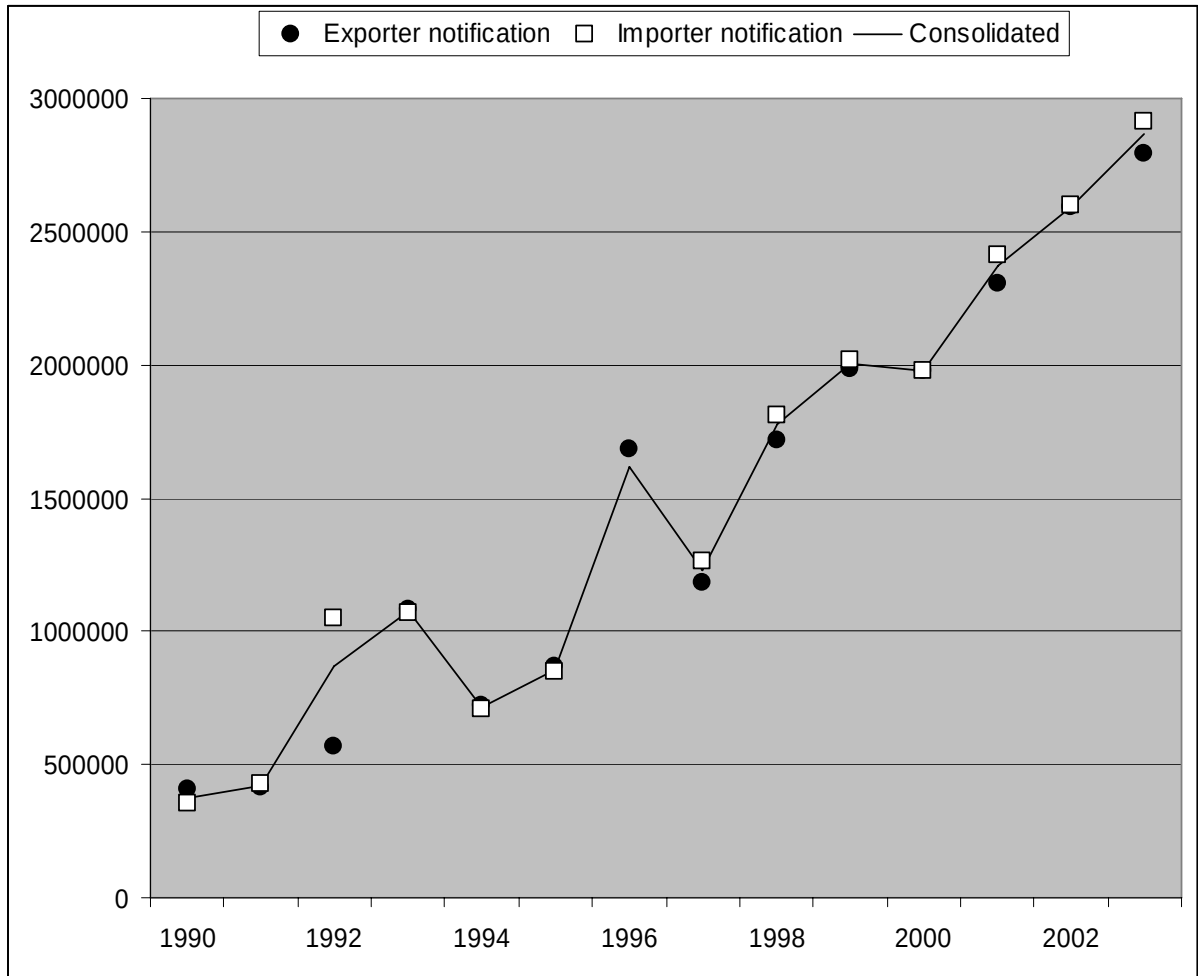


Figure 1: Consolidation of data on exports of wheat [t] by Country U to Country M

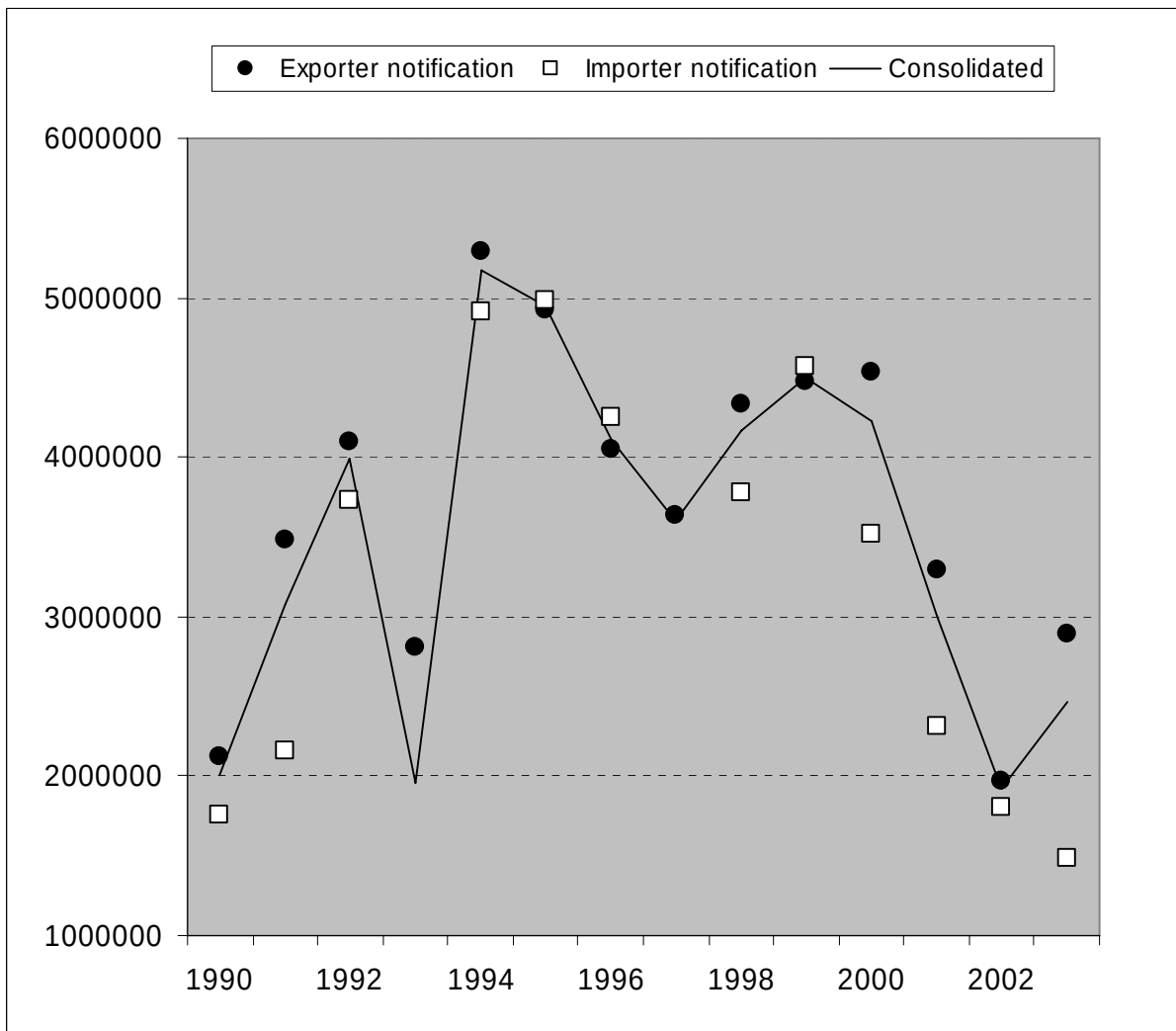


Figure 2: Consolidation of data on exports of wheat [t] by Country U to Country E

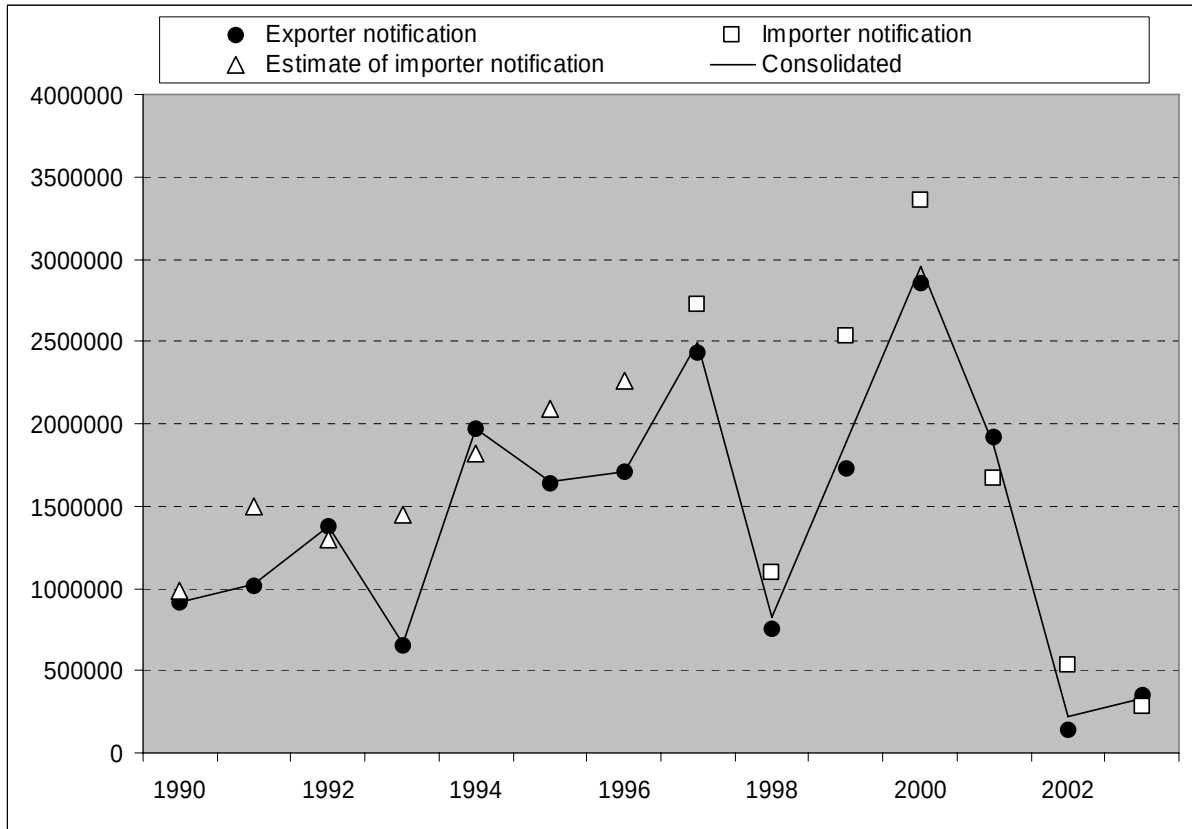


Figure 3: Consolidation of data on exports of wheat [t] by County C to Country I

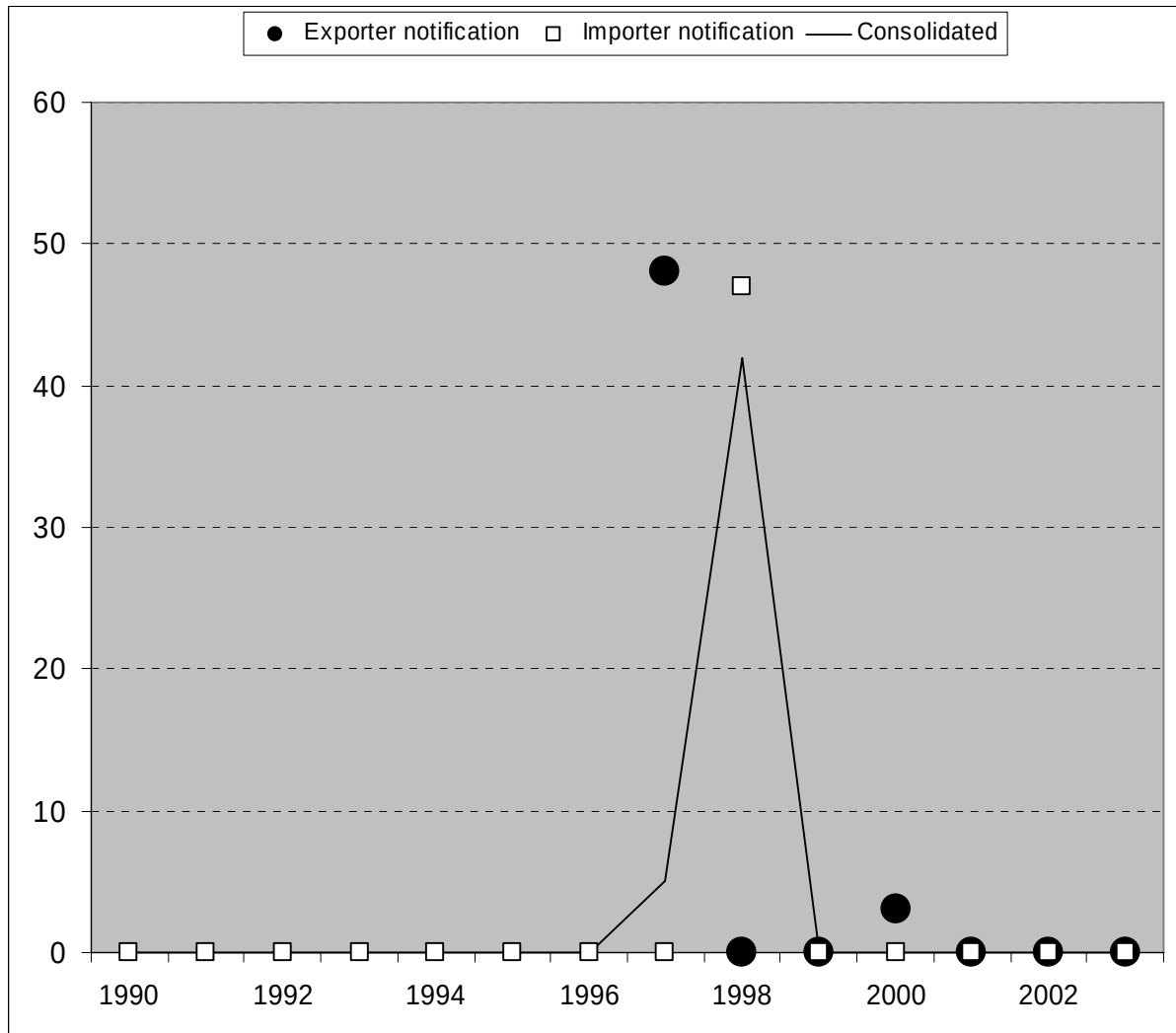


Figure 4: Consolidation of data on exports of maize [t] by Country G to Country J

Table 2: Different cases in the SUA problem

Element	Steered by	A priori knowledge
Production, yield/carcass weight, areas/herds	Standard term	μ : Given data or trend results; σ : based on error of trend line
Sinks excluding stock changes	Combination of (3) share term and (4) HP-Filter term	(1) μ : trend on share, σ : based on error of trend line on share (2) μ : trend on share times times expected mean of production plus imports; σ : subjective setting
Stock changes	Combination of (3) Stock-to-use term (4) Standard term	(1) $\mu = 0.3$ $\sigma = 0.0005$ (2) $\mu = 0.0$, σ large

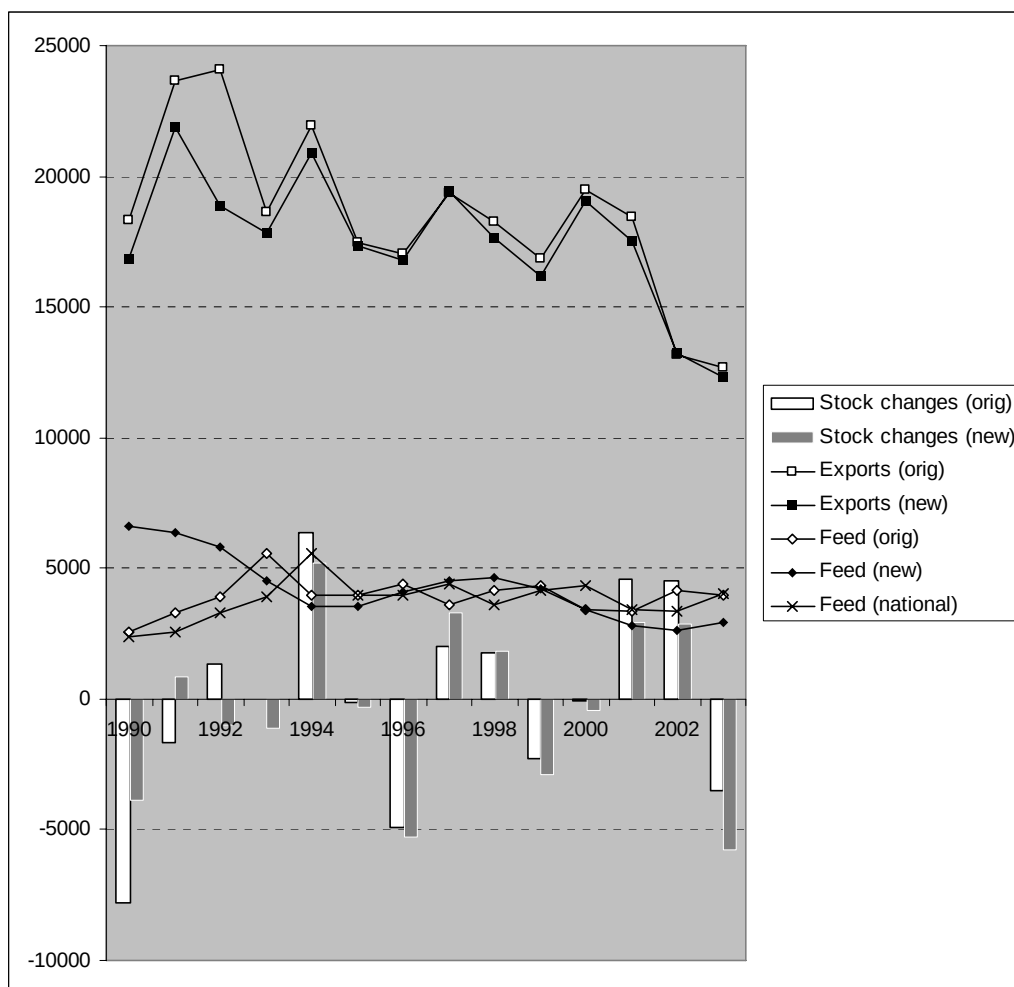


Figure 5: Original data and new estimates on selected SUA items [1000 t] for wheat in Country C

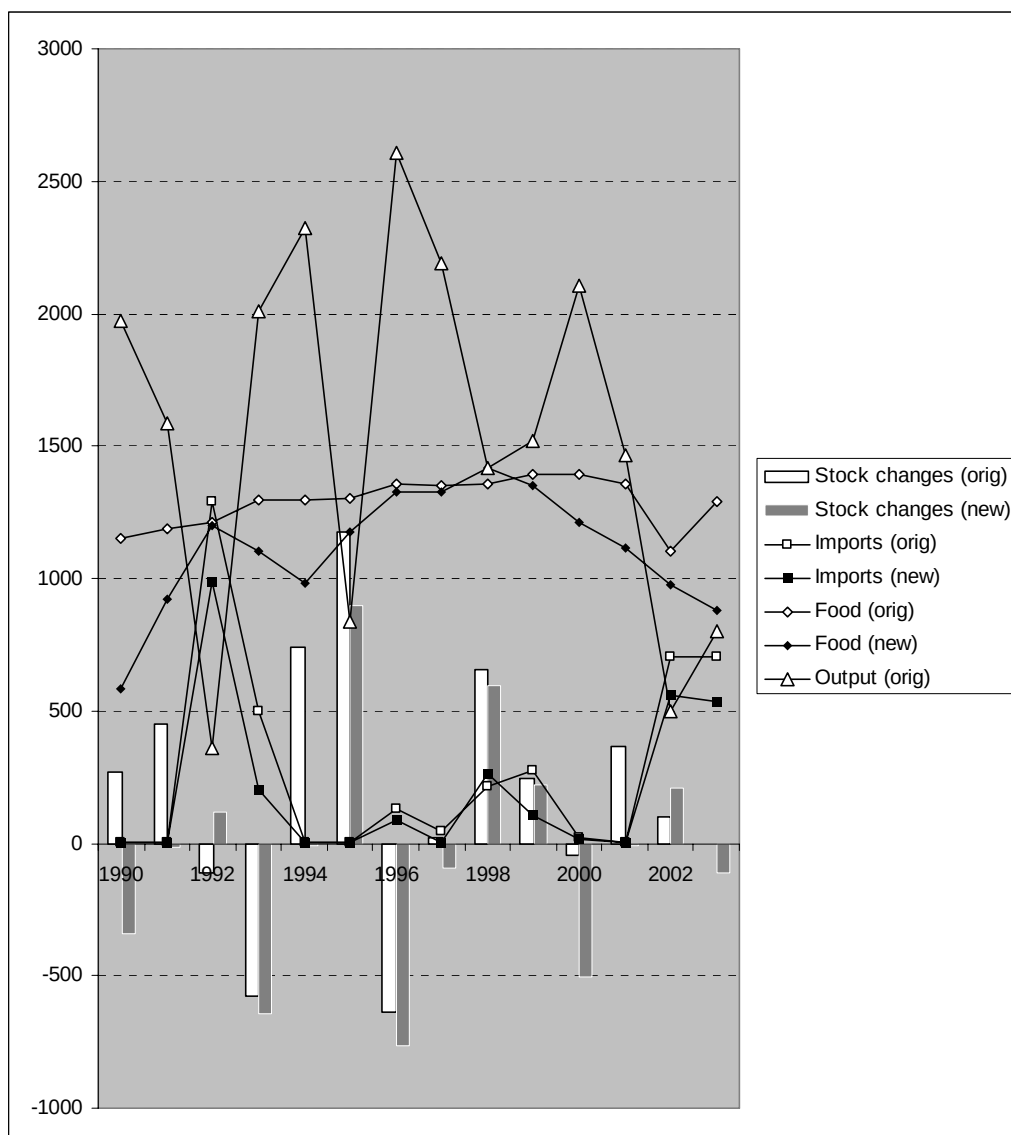


Figure 6: Original data and new estimates on selected SUA items [1000 t] for maize in Country Z

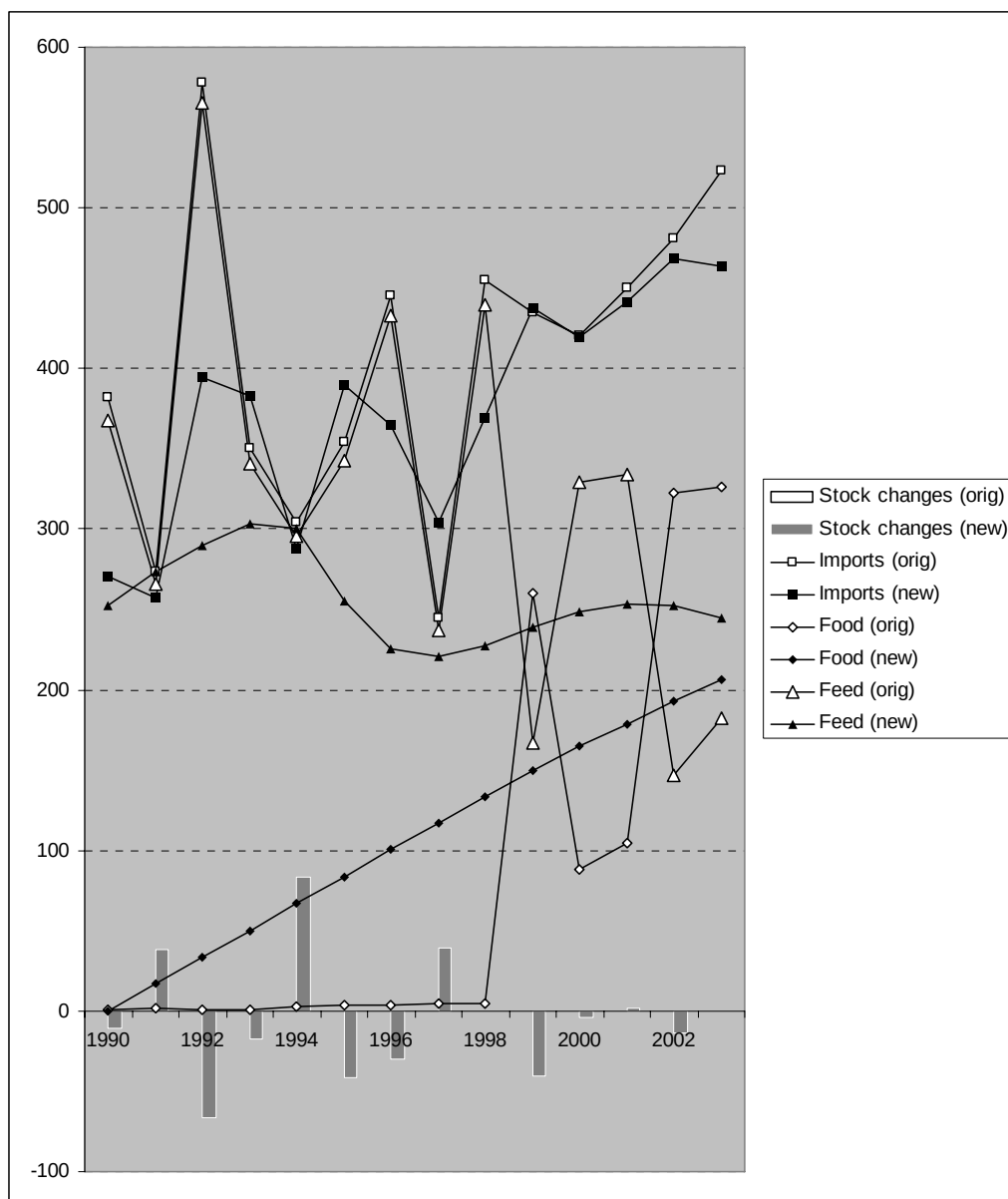


Figure 7: Original data and new estimates on selected SUA items [1000 t] for maize in Country J