

On the Definition of Software Accuracy in Redundant Measurement Systems

Miguel J. Bagajewicz

School of Chemical Engineering and Materials Science, University of Oklahoma, Norman, OK 73019

DOI 10.1002/aic.10379

Published online March 3, 2005 in Wiley InterScience (www.interscience.wiley.com).

Accuracy of measurements is defined as the sum of precision (standard deviation) and bias. Unfortunately, this definition is of little use unless bias is independently assessed. In this article, we extend this classical definition of accuracy of measurements to one of estimators obtained through data reconciliation techniques. Connection to gross error robustness measures, defined earlier by the author, is studied. © 2005 American Institute of Chemical Engineers AIChE J, 51: 1201–1206, 2005

Keywords: data reconciliation, gross error detection, accuracy

Introduction

Accuracy has been defined for individual measurements as the sum of the absolute value of the systematic error plus the standard deviation of the meter (Miller, 1996). However, it is well known that in the absence of hardware redundancy, that is, another measurement obtained using another meter, or software redundancy, that is, through data reconciliation, systematic errors cannot be detected. The definition is thus good only conceptually and has little practical value.

Because hardware redundancy is very costly, software redundancy is the most popular way of identifying biases. In particular, linear data reconciliation renders estimators of process variables that satisfy balance constraints and is capable of filtering biases to some extent. Hypothesis testing has been extensively used for detection and identification of gross errors. Literature surveys of this issue can be found in Mah (1990), Madron (1992), Sánchez and Romagnoli (2000), and Narasimhan and Jordache (2001).

The object of the work described in this article is to redefine accuracy, not in terms of the actual systematic error that an instrument has, but rather with respect to the ability of a gross error detection scheme to detect such systematic errors. In other words, accuracy should be reported in relation to the ability of a gross error detection technique to detect and eliminate gross errors. Moreover, because such ability depends on the number of gross errors present in the measurement, the

report on accuracy should be made contingent on that number, as discussed in detail in this article. Thus the aim herein is to distinguish the existing definition of hardware accuracy (that of the instrument) from software accuracy.

With respect to software accuracy, but not addressing it directly, some measures related to the ability of a set of instruments to detect gross errors or to control the impact of undetected ones exist. Bagajewicz (1997) introduced the concept of gross error detectability, which relates the structure and precision of an instrumentation network to the size of the bias that can be detected in a particular instrument. He also introduced the concept of resilience, which determines the effect of undetected gross errors throughout the instrumentation network on a particular estimator. Both concepts are of course strongly linked to data reconciliation techniques. In this article, it will be shown how these two concepts are related to the proposed concept of software accuracy.

The article is organized as follows. The definition of hardware or instrument accuracy is reviewed first, followed by the proposed definition of software accuracy (Bagajewicz, 2003). The paper extends the conditions to multiple gross errors and illustrates the concept through an example. Finally, connections to previous definitions of gross error robustness are discussed.

Accuracy of Measurements

Miller (1996) defines the accuracy of an instrument as the sum of the systematic error plus the precision of the instrument:

$$a_i = \delta_i + \sigma_i \quad (1)$$

M. J. Bagajewicz's e-mail address is bagajewicz@ou.edu.

where a_i , δ_i , and σ_i are the accuracy, systematic error, and precision (square root of variance) of the mean of a certain number of repeated measurements made by a meter on variable i .

The problem with this definition is that it is useless in practice. Indeed, on one hand, if one knows the systematic error, either the instrument is immediately recalibrated and therefore the systematic error is eliminated or the measurement is adjusted. On the other hand, if one does not know the value of the systematic error (the usual case), such systematic error cannot be inferred directly from the measurements of this instrument; only the precision can. Thus, although the definition has conceptual value, it has no practical implications.

Induced Bias

Let the measurement vector be described by

$$y = \mu + \delta + \varepsilon \quad (2)$$

where y is the vector of measurements, μ represents the true values, δ is the vector of biases, and ε is the vector of random errors. With the common assumption that $\varepsilon \cong N_p(0, S)$ and with the conservation laws given by $A\mu = 0$, the maximum likelihood estimate of μ is given by

$$\hat{\mu} = y - SA^T(ASA^T)^{-1}Ay \quad (3)$$

The expected value of the vector of estimators is then

$$E[\hat{\mu}] = \mu + \delta - SA^T(ASA^T)^{-1}A\delta \quad (4)$$

We note that $E[\hat{\mu}] = \mu$ only when $\delta = 0$. In other words, the estimator is unbiased only in the absence of systematic errors. This expression allows us to define the induced bias as the difference between the expectation when gross errors are present and the true values, that is

$$\hat{\delta} = [I - SW]\delta \quad (5)$$

where $W = A^T(ASA^T)^{-1}A$, which is the variance-covariance matrix of $S^{-1}(\hat{\mu} - y)_k$. We immediately notice that when there is one biased variable, the smearing effect of data reconciliation reduces the absolute size of the induced error in the variable. Indeed, assume that the only bias is δ_i . Then $\hat{\delta}_i = [I - SA^T(ASA^T)^{-1}A]_{ii}\delta_i$, which is smaller than δ_i , because SW has a nonnegative diagonal.

There are a variety of available techniques to detect and eliminate gross errors in a system. Mah and Tamhane (1982) proposed the measurement test. Later the maximum power measurement test was introduced (Tamhane and Mah, 1985). This test has been used in gross error hypothesis testing in common applications. Nodal methods, testing the residuals of the constraints using the normal distribution, have also been proposed. In turn, Narasimhan and Mah (1987) introduced the GLR (generalized likelihood ratio) identification test, which can also provide an estimate of the gross error. Narasimhan and Mah (1987) and Crowe (1988) showed that these two are equivalent and have maximum power for the case of one gross error. Later, principal-component tests (PCT) were proposed

by Tong and Crowe (1995, 1996) as an alternative for multiple bias and leak identification. Industrial applications of PCMT were reported by Tong and Bluck (1998), who indicated that tests based on principal-component analysis (PCA) are more sensitive to subtle gross errors than others, and have greater power to correctly identify the variables in error than the conventional nodal, measurement, and global tests. Principal-component tests, however, have proven to be less efficient than the measurement test when used for multiple gross error identification (Bagajewicz et al., 1999, 2000). Powerful and statistically correct combinations of nodal tests have been introduced by Rollins et al. (1996). We cite only the above methods because they are relevant to this article, but many other tests have been proposed and they can be found in the aforementioned books.

For the case of multiple gross errors, serial elimination was originally proposed by Ripps (1965). The method proposes to eliminate the measurement that renders the largest reduction in a test statistics until no test fails. Several authors proposed different variants of this procedure (see, for example, Nogita, 1972; Romagnoli and Stephanopoulos, 1981; Rosenberg et al., 1987). Of all these strategies, the one that has been implemented in commercial software more often is the one where the largest measurement test (MT) is used to determine the measurement to eliminate. Most current commercial software use serial elimination.

The aforementioned set of techniques are effective only in detecting gross errors of a size that is above a certain threshold. Below such a threshold, they are not detected and they remain as induced biases in the reported estimators. Thus, one must not be fooled to think that a system has no biases, even when the most powerful gross error detection scheme has been applied. Small induced biases will always persist—which leads to a new definition of accuracy, presented next.

Accuracy of Estimators

For consistency, it is proposed to extend Eq. 1 to define accuracy of an estimate in the same form. Thus, the accuracy of an estimator (or software accuracy) is defined as the sum of the maximum undetected induced bias plus the precision of the estimator, that is

$$\hat{a}_i = \hat{\sigma}_i + \delta_i^* \quad (6)$$

where $\hat{a}_i$, δ_i^* , and $\hat{\sigma}_i$ are the accuracy, the maximum undetected induced bias, and the precision (square root of variance $\hat{S}_{ii}$) of the estimator, respectively. The accuracy of the system can be defined in various ways, such as by making an average of all accuracies or taking the maximum among them. Because this involves comparing the accuracy of measurements of different magnitudes, relative values are recommended. Herein the following is used

$$\tilde{\alpha}_S = \max_{\forall i} \left\{ \frac{\hat{a}_i}{F_i} \right\} \quad (7)$$

where $\tilde{\alpha}_S$ is the accuracy of the system.

Maximum Induced Bias

Assume now that the maximum power measurement is used. The maximum power test statistics for stream k under the presence of gross error in stream s ($Z_{k,s}^{MP}$) is given by the following expression

$$Z_{k,s}^{MP} = \frac{|S^{-1}(\hat{\mu} - y)_k|}{\sqrt{W_{kk}}} = \frac{|e_k^T W(\mu + \delta_s e_s + \varepsilon)|}{\sqrt{W_{kk}}} \quad (8)$$

The hypothesis test consists of comparing the statistic with a critical value associated with the chosen degree of confidence. Thus, the usual assumption of white noise and confidence $p = 95\%$ renders a critical value. One can, of course, add the Bonferroni correction, which consists of dividing the level of confidence by the number of measurements. Thus the expected value of the maximum power test can be obtained directly from Eq. 8 by recognizing that

$$E[Z_{k,s}^{MP}] = \frac{|W_{ks}\delta_s|}{\sqrt{W_{kk}}} \quad (9)$$

This test is consistent; that is, that under the assumption of one gross error, it points to the right location. In other words,

$$E[Z_{s,s}^{MP}] > E[Z_{k,s}^{MP}] \quad \forall k \neq s \quad (10)$$

Therefore, a gross error in variable s (δ_s) will be detected if its value is larger than a threshold value $\delta_{crit,s}^{(p)}$, given by

$$\delta_{crit,s}^{(p)} = \frac{Z_{crit}^{(p)}}{\sqrt{W_{ss}}} \quad (11)$$

In turn, the corresponding maximum undetectable induced bias in variable i ($\hat{\delta}_{crit,i,s}^{(p)}$)—always under the assumption of one gross error—attributed to the undetected gross error in variable s is

$$\begin{aligned} \hat{\delta}_{crit,i,s}^{(p)} &= \{[I - SW]e_s \delta_{crit,s}^{(p)}\}_i = [(I - SW)_{is}] \delta_{crit,s}^{(p)} \\ &= \frac{[(I - SW)_{is}]}{\sqrt{W_{ss}}} Z_{crit}^{(p)} \end{aligned} \quad (12)$$

Thus, the maximum induced bias that will be undetected is given by (Bagajewicz, 2003)

$$\hat{\delta}_i^{(p,1)} = \max_{\forall s} \hat{\delta}_{crit,i,s}^{(p)} = Z_{crit}^{(p)} \max_{\forall s} \frac{[(I - SW)_{is}]}{\sqrt{W_{ss}}} \quad (13)$$

We note here that, although the maximum power test is consistent—that is, it will point to the right location of the bias—this does not mean that the maximum undetectable induced value from a location of a gross error in some other variable cannot be larger.

Maximum Power Test–Based Accuracy of Order One

One can now complete the definition of accuracy as follows:

The accuracy of the estimator of a variable i , under the assumption of the use of the maximum power test with confidence p and the presence of only one gross error, is given by

$$\tilde{\alpha}_i^{MP(p,1)} = \sqrt{\hat{\delta}_{ii}} + Z_{crit}^{(p)} \max_{\forall s} \frac{[I - (SW)_{is}]}{\sqrt{W_{ss}}} \quad (14)$$

or in relative terms

$$\tilde{\alpha}_i^{MP(p,1)} = \frac{\sqrt{\hat{\delta}_{ii}}}{F_i} + \frac{Z_{crit}^{(p)}}{F_i} \max_{\forall s} \frac{[I - (SW)_{is}]}{\sqrt{W_{ss}}} \quad (15)$$

Maximum Power Test–Based Accuracy of Higher Order

Consider now the presence of n_T gross errors. The maximum power statistics will render

$$E[Z_{k,T}^{MP}] = \frac{\left| \sum_{\forall s \in T} W_{ks} \delta_s \right|}{\sqrt{W_{kk}}} \quad (16)$$

where T is the set of n_T gross errors located in specific variables.

It is well known that under certain conditions the test points to the right set of gross errors, but it can sometimes fail. However, the inconsistency of the MP test is of small concern here. Indeed, we define accuracy in terms of undetected biases. When the measurement test points to an error, we assume that these errors, in their right location, are identified. It is of little use to define accuracy in terms of wrongly applied techniques. Thus, we assume that the practitioner is aware of the inconsistency and will take care of the problem by properly identifying the right location of the gross errors. The same can be said by the uncertainty of the gross error location arising from equivalency theory (Bagajewicz and Jiang, 1998). According to this theory, when a test (any test) points to a set of gross errors, there are sets of gross errors that are equally equivalent, in the sense that they produce exactly the same numerical effect in the data.

The MP test will flag for any combination of gross errors such that

$$Z_{crit}^{(p)} \leq \frac{\left| \sum_{\forall s \in T} W_{ks} \delta_s \right|}{\sqrt{W_{kk}}} \quad (17)$$

Therefore the sets of critical values of a particular set of gross errors T is not unique and cannot be uniquely determined as in the case of a single gross error. Consider one such set of critical values. We then write

$$\left| \sum_{\forall s \in T} W_{ks} \delta_{crit,s}^{(p)} \right| = Z_{crit}^{(p)} \sqrt{W_{kk}} \quad \forall k \quad (18)$$

The corresponding induced bias in variable i is

$$\hat{\delta}_i^{(p)} = \{[I - SW] \delta_{crit,i}^{(p)}\} = \delta_{crit,i}^{(p)} - \sum_{s \in T} (SW)_{is} \delta_{crit,s}^{(p)} \quad (19)$$

where $\delta_{crit}^{(p)}$ is the vector containing a critical value of the gross error size in the selected positions corresponding to the set T at the confidence level p . To find the maximum possible undetected induced bias, one has to explore all possible values of gross errors in the set. Thus, for each set T we obtain the maximum induced and undetected bias by solving the following problem

$$\left. \begin{aligned} \hat{\delta}_i^{(p)}(T) = \max_{\forall s \in T} & \left| \delta_{crit,i} - \sum_{s \in T} (SW)_{is} \delta_{crit,s} \right| \\ \text{s.t.} & \\ \left| \sum_{\forall s \in T} W_{ks} \delta_{crit,s} \right| & \leq Z_{crit}^{(p)} \sqrt{W_{kk}} \quad \forall k \end{aligned} \right\} \quad (20)$$

This problem is solved for each set of n_T gross errors located in specific variables given by the set T and the unknowns are the components of the vector δ_{crit} , where the inequality is written for convenience (the constraint is binding at the optimum).

We formally define the location of the gross errors through the use of a vector of binary variables q_T , that is, $q_{T,s} = 1$ if a gross error of set T is located in variable s and zero otherwise. Clearly $\sum_{\forall s} q_{T,s} = n_T$. In addition, we can drop the absolute value from the objective because one can invert the signs of all gross errors and obtain the same answer. Thus, we can write

$$\left. \begin{aligned} \hat{\delta}_i^{(p)}(T) = \max & \left\{ \delta_{crit,i} q_{T,i} - \sum_{\forall s} (SW)_{is} \delta_{crit,s} q_{T,s} \right\} \\ \text{s.t.} & \\ \left| \sum_{\forall s} \frac{W_{ks}}{\sqrt{W_{kk}}} \delta_{crit,s} q_{T,s} \right| & \leq Z_{crit}^{(p)} \quad \forall k \\ q_{T,k} |\delta_{crit,k}| & \geq 0 \quad \forall k \end{aligned} \right\} \quad (21)$$

After the absolute value is replaced by the following two standard inequalities

$$\begin{aligned} \sum_{\forall s} \frac{W_{ks}}{\sqrt{W_{kk}}} \delta_{crit,s} q_{T,s} & \leq Z_{crit}^{(p)} \quad \forall k \\ - \sum_{\forall s} \frac{W_{ks}}{\sqrt{W_{kk}}} \delta_{crit,s} q_{T,s} & \leq Z_{crit}^{(p)} \quad \forall k \end{aligned} \quad (22)$$

the problem is linear for a given instrumentation network (the vector q_T is a constant).

Therefore, because one has to consider all possible combinations of bias locations, we write

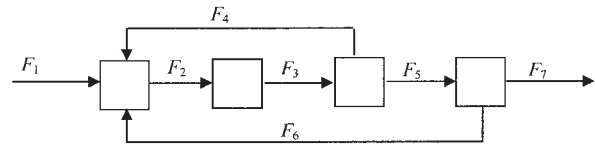


Figure 1. Example flow sheet.

$$\hat{\delta}_i^{(p,n_T)} = \max_{\forall T} \hat{\delta}_i^{(p)}(T) \quad (23)$$

which implies that one has to solve the problem (Eq. 21) for all combinations of n_T gross errors in the system. One can convert this into a simple MILP (mixed integer linear programming) problem by recognizing that it contains products of integer and continuous variables for which standard transformations are available.

Therefore, one can now complete the definition of accuracy as follows:

The accuracy of the estimator of a variable i , under the assumption of the use of the maximum power test and the presence of t gross errors is given by

$$\bar{a}_i^{MP(p,n_T)} = \sqrt{\bar{\delta}_{ii}} + \hat{\delta}_i^{(p,n_T)} \quad (24)$$

Example

We illustrate the above concepts using the following process, as illustrated in Figure 1, that was used by Rollins et al. (1996) but originally taken from Narasimhan and Mah (1987).

For this process

$$A = \begin{bmatrix} 1 & -1 & & 1 & & 1 \\ & 1 & -1 & & & \\ & & 1 & -1 & -1 & \\ & & & 1 & -1 & -1 \end{bmatrix} \quad (25)$$

Consider the following variance-covariance matrix

$$S = \begin{bmatrix} 1 & & & & & \\ & 0.2 & & & & \\ & & 1 & & & \\ & & & 1 & & \\ & & & & 1 & \\ & & & & & 1 \end{bmatrix} \quad (26)$$

The values of precision and accuracy of the different variables for order one are given in Table 1. The last column contains the sizes of the gross errors and their locations that cause the largest undetected bias in the corresponding variable.

The table illustrates that flows F_4 and F_6 have a reasonable precision (3.3 and 3.6%, respectively) but a very bad accuracy (of the order of 10%), whereas others do not. If, for example, one is interested only in F_1 or F_7 , then the accuracy is an acceptable 1.56%. That is, only induced gross errors of size < 0.95 will be undetected.

The accuracy of a variable cannot sometimes be improved by locating a better instrument in the variable. Indeed, if one locates an instrument of 0.1% precision measuring flow rate F_4 ,

Table 1. Accuracies of Order One for Example 1

Stream	σ_i	σ_i/F_i (%)	$\hat{\sigma}_i$	$\hat{\sigma}_i/F_i$ (%)	$\hat{a}_i^{MP(0.95,1)}$	$\hat{\alpha}_i^{MP(0.95,1)}$ (%)	$\delta_{crit,s}^{(p)}(s)$
F_1	1	1	0.614	0.614	1.56	1.56	2.509 (1)
F_2	0.447	0.319	0.389	0.278	1.74	1.24	1.788 (2)
F_3	1	0.714	0.389	0.278	1.74	1.24	1.788 (2)
F_4	1	5	0.659	3.294	1.80	9.00	2.632 (4)
F_5	1	0.833	0.629	0.525	1.64	1.37	2.548 (5)
F_6	1	5	0.713	3.569	2.15	10.77	2.827 (6)
F_7	1	1	0.614	0.614	1.56	1.56	2.509 (7)

its accuracy only decreases to be 8.77%, which is not a significant drop. To fix the accuracy of stream F_4 , actually, one needs to identify which undetected measurement bias in some other stream induces such a big bias in F_4 . For example, changing the precision of F_2 , F_3 , and F_5 to 0.1% reduces the accuracy value of F_4 to 1.72%. In fact, measuring this flow does not really help at all.

Table 2 shows the results for the accuracy of order two. The last column shows the values and locations of gross errors that produce the largest undetected bias.

We first note that F_1 and F_7 have the same accuracy because they are part of an equivalent set (Bagajewicz and Jiang, 1998). Second, we note that all streams but one have the largest undetected induced bias contributed by biases in their own measurements. Finally, we confirm that the smearing of data reconciliation renders smaller induced errors.

We now turn to analyze the case of three gross errors. In this case the definition fails to produce a bounded value of induced bias. Indeed, all streams in this flowsheet participate in a loop of three streams, which raises the possibility of a set of three gross errors in these streams to go completely undetected. Indeed, take streams F_1 , F_6 , and F_7 , for example. In this set, consider a set of three gross errors of the same size but different sign as follows: δ_1 , $\delta_6 = -\delta_1$ and $\delta_7 = \delta_1$. Such a set renders the maximum power test unable to find gross errors, simply because they are consistent with the balance equations.

Connections to Reliability and Estimation Availability

The above example raises some issues regarding the ability of the maximum power test to identify the presence of gross errors. We recognize, however, the extremely low probability of the cases that are raised. The way to handle this theoretically is to determine the failure distribution density through time and the size distribution of these biases. Once these are known, one can obtain a probabilistic-based definition of accuracy. This is left for future work.

Table 2. Accuracies of Orders Two and Three for Example 1

Stream	$\hat{a}_i^{MP(0.95,2)}$	$\hat{\alpha}_i^{MP(0.95,2)}$ (%)	$\delta_{crit,s_1}^{(p)}$, $\delta_{crit,s_2}^{(p)}$ (s_1 , s_2)
F_1	2.748	2.748	3.696, 1.958 (1, 7)
F_2	2.771	1.980	2.317, 4.208 (2, 3)
F_3	2.771	1.980	2.317, 4.208 (2, 3)
F_4	3.374	16.869	1.891, 4.208 (2, 4)
F_5	2.709	2.258	-1.903, 3.619 (4, 5)
F_6	3.145	15.724	1.841, 3.818 (5, 6)
F_7	2.748	2.748	1.958, 3.696 (1, 7)

Connections to Gross Error Detectability and Resilience

Accuracy embeds the notion that gross errors might be present together with white noise. In turn, gross error detectability is defined as the ability of a network to detect a gross error of a certain minimum size or larger. This concept was introduced by Bagajewicz (1997, 2000), who proposed that the minimum size be given by

$$\hat{\delta}_k(m, \alpha, \beta) = \rho(m, \alpha, \beta) \frac{\sigma_k^2}{(\sigma_k^2 - \tilde{\sigma}_k^2)^{1/2}} \quad (27)$$

which is the smallest size of gross error that can be detected with probability β . Typical values of β are 50 and 90% and tables for $\rho(m, \alpha, \beta)$, as well as an empirical expression for large number of degrees of freedom (m), are given by Madron (1992). This expression is obtained by realizing that the formula above is the relation between the minimum value $\hat{\delta}_i^*(m, \alpha, \beta)$ and the noncentrality parameter $\rho(m, \alpha, \beta)$ of a chi-squared distribution that is followed by $r(C_R Q_R C_R^T)^{-1} r$ in the presence of gross error. Expressions for larger numbers of gross errors have not yet been developed.

The expression used for gross error detectability is based on the idea that one is using the global test, whereas the definition of accuracy above assumes that the maximum power measurement test is used. It is therefore claimed that the accuracy definition proposed here is much more appropriate to be used as a measure of robustness than gross error detectability. Moreover, the former was extended to the presence of multiple gross errors, whereas the latter is limited in this regard.

Finally, resilience can be expressed as the difference between accuracy and precision. Indeed, resilience is the ability of a network to limit the smearing effect of gross errors. In other words, it is a number that sets an upper bound on the effect of undetected gross errors in estimators obtained using data reconciliation. Therefore, although the concept in itself can continue to be used, we see little purpose in using it, given that accuracy contains everything one might need in robustness terms.

Conclusions

Accuracy has been defined in the context of data reconciliation and related to one popular test, the maximum power measurement test. The definition is an extension of one used for measurements to estimators. We define the induced bias as the bias obtained in a stream after data reconciliation is performed in the presence of gross errors. We show that the smearing effect of data reconciliation renders induced biases

that are smaller than the actual measurement biases. Finally, we show that loops initiate unlimited accuracy and we claim that these are nonetheless events with low probability so we anticipate working on a probabilistic-based definition of accuracy.

Acknowledgments

Financial support from the National Science Foundation through Grant CTS-0075814 is gratefully acknowledged.

Notation

a = matrix for material balances
 a_i = hardware accuracy
 $\hat{a}_i$ = accuracy of estimators
 p = confidence level for tests
 q_i = binary vector to indicate bias location
 S = variance-covariance matrix of measurements
 y = vector of measurements
 $Z_{k,s}^{MP}$ = maximum power statistics for stream k with a gross error in stream s

Greek letters

$\bar{\alpha}_s$ = accuracy of the system
 δ_i = measurement bias
 δ_i^* = maximum undetected induced bias
 $\hat{\delta}_i$ = induced bias vector
 $\delta_{crit,s}^{(p)}$ = maximum undetected bias in stream s
 $\hat{\delta}_{crit,i,s}^{(p)}$ = maximum undetected induced bias in stream i for a gross error in stream s
 ε = vector of random errors
 μ = true values of process variables
 $\hat{\mu}$ = maximum likelihood estimator of μ
 σ_i = measurement precision
 $\hat{\sigma}_i$ = the precision of estimators

Literature Cited

- Bagajewicz, M., "Optimal Sensor Location in Process Plants," *AIChE J.*, **43**(9), 2300 (1997).
- Bagajewicz, M., *Design and Upgrade of Process Plant Instrumentation*, Technomic, Lancaster, PA (now distributed by CRC Press, Boca Raton, FL) (2000).
- Bagajewicz, M., "Consider Accuracy as a Robustness Measure in Instrumentation Design," Proc. of the Topical Conference on Sensors, AIChE Annual Meeting, San Francisco, CA, November (2003).
- Bagajewicz, M., and Q. Jiang, "Gross Error Modeling and Detection in Plant Linear Dynamic Reconciliation," *Comput. Chem. Eng.*, **22**(12), 1789 (1998).
- Bagajewicz, M., Q. Jiang, and M. Sánchez, "Performance Evaluation of PCA Tests for Multiple Gross Error Identification," *Comput. Chem. Eng. Suppl.*, **23**, S589 (1999).
- Bagajewicz, M., Q. Jiang, and M. Sánchez, "Performance Evaluation of PCA Tests in Serial Elimination Strategies for Gross Error Identification," *Chem. Eng. Commun.*, **183**, 119 (2000).
- Crowe, C. M., "Recursive Identification of Gross Errors in Linear Data Reconciliation," *AIChE J.*, **34**, 541 (1988).
- Madron, F., *Process Plant Performance, Measurement Data Processing for Optimization and Retrofits*, Ellis Horwood, West Sussex, UK (1992).
- Mah, R. S. H., *Chemical Process Structures and Information Flows*, Butterworths, Stoneham, MA (1990).
- Mah, R. S. H., and A. C. Tamhane, "Detection of Gross Errors in Process Data," *AIChE J.*, **28**, 828 (1982).
- Miller, R. W., *Flow Measurement Engineering Handbook*, McGraw-Hill, New York (1996).
- Narasimhan, S., and C. Jordache, *Data Reconciliation and Gross Error Detection. An Intelligent Use of Process Data*, Gulf Publishing, Houston, TX (2000).
- Narasimhan, S., and R. S. H. Mah, "Generalized Likelihood Ratio Method for Gross Error Detection," *AIChE J.*, **33**, 1514 (1987).
- Nogita, S., "Statistical Test and Adjustment of Process Data," *Ind. Eng. Chem. Process Des. Dev.*, **2**, 197 (1972).
- Ripps, D. L., "Adjustment of Experimental Data," *Chem. Eng. Prog. Symp. Ser.*, **61**, 8 (1965).
- Rollins, D. K., Y. Cheng, and S. Devanathan, "Intelligent Selection of Hypothesis Tests to Enhance Gross Error Identification," *Comput. Chem. Eng.*, **20**(5), 517530 (1996).
- Romagnoli, J., and G. Stephanopoulos, "On the Rectification of Measurement Errors for Complex Chemical Plants," *Chem. Eng. Sci.*, **35**(5), 1067 (1980).
- Rosenberg, J., R. S. H. Mah, and C. Jordache, "Evaluation of Schemes for Detecting and Identifying Gross Errors in Process Data," *Ind. Eng. Chem. Res.*, **26**, 555 (1987).
- Sánchez, M., and J. Romagnoli, *Data Processing and Reconciliation for Chemical Process Operations*, Academic Press, San Diego, CA (2000).
- Tamhane, A. C., and R. S. H. Mah, "Data Reconciliation and Gross Error Detection in Chemical Process Networks," *Technometrics*, **27**(4), 409 (1985).
- Tong, H., and D. Bluck, "An Industrial Application of Principal Component Test to Fault Detection and Identification," Proc. of IFAC Workshop on On-Line-Fault Detection and Supervision in the Chemical Process Industries, Solaize (Lyon), France (1998).
- Tong, H., and C. M. Crowe, "Detection of Gross Errors in Data Reconciliation by Principal Component Analysis," *AIChE J.*, **41**(7), 1712 (1995).
- Tong, H., and C. M. Crowe, "Detecting Persistent Gross Errors by Sequential Analysis of Principal Components," *Comput. Chem. Eng. Suppl.*, **20**, S733 (1996).

Manuscript received Dec. 2, 2003, and revision received Aug. 18, 2004.