

On a New Definition of a Stochastic-based Accuracy Concept of Data Reconciliation-Based Estimators

M. Bagajewicz
University of Oklahoma
100 E. Boyd St., Norman OK 73019, USA

Abstract

Traditionally, accuracy of an instrument is defined as the sum of the precision and the bias. Recently, this notion was generalized to estimators. However, the concept used a maximum undetected bias, as well as ignored the frequency of failures. In this paper the definition of accuracy is modified to include expected undetected biases and their frequency.

Keywords: Instrumentation Network Design, Data Reconciliation, Plant Monitoring.

1. Introduction

Traditionally, accuracy of an instrument is defined as the sum of the precision and the bias (Miller, 1996). In a recent paper (Bagajewicz, 2004) this notion was generalized to estimators arguing that the accuracy of an estimator is the sum of the precision and the maximum induced bias. This maximum induced is the maximum value of the bias of the estimator used, that is, a result of a certain specific number of biases in the network which have not been detected. This lead to a definition of accuracy that is dependent on the number of biases chosen. Aside from many other shortcomings of the definition, two stand out as the most important: The definition has no time horizon associated to it, nor states anything about the frequency at which each sensor will fail, or the time it will take to repair it. In addition, the definition could be more realistic if expected bias, instead of maximum bias is used.

In this paper, we review the definitions and discuss the results of a Montecarlo technique that can help determine an expected value of accuracy.

2. Background

Accuracy was defined for individual measurements as the sum of the absolute value of the systematic error plus the standard deviation of the meter (Miller, 1996). Since the bias is usually not known, the definition has little practical value. Bagajewicz (2004) introduced a new definition of accuracy of an estimator (or software accuracy) defined as the sum of the maximum undetected induced bias plus the precision of the estimator:

$$\hat{a}_i = \hat{\sigma}_i + \delta_i^* \quad (1)$$

where $\hat{a}_i$, δ_i^* and $\hat{\sigma}_i$ are the accuracy, the maximum undetected induced bias and the precision (square root) of the estimator's variance $\hat{S}_{ii}$, respectively. In turn, the accuracy of the system can be defined in various ways, for example making an average of all accuracies or taking the maximum among them. Since this involves comparing the accuracy of measurements of different magnitude, relative values are recommended.

The maximum undetected induced bias is obtained from the assumption that a particular gross error detection test is used. In the case of the maximum power measurement test, and under the assumption of one gross error being present in the system this value is given by:

$$\hat{\delta}_i^{(p,1)} = Z_{crit}^{(p)} \underset{\forall s}{Max} \frac{[(I - SW)_{is}]}{\sqrt{W_{ss}}} \quad (2)$$

where $Z_{crit}^{(p)}$ is the critical value for the test at confidence level p , S is the variance-covariance matrix of the measurements and $W = A^T (A S A^T)^{-1} A$ (A is the incidence matrix). When a larger number of gross errors are present in the system, an optimisation model is needed. Thus, for each set T we obtain the maximum induced and undetected bias by solving the following problem:

$$\left. \begin{aligned} \hat{\delta}_i^{(p)}(T) = \underset{\forall s \in T}{Max} & \left| \delta_{crit,i} - \sum_{s \in T} (SW)_{is} \delta_{crit,s} \right| \\ s.t. & \\ & \left| \sum_{\forall s \in T} W_{ks} \delta_{crit,s} \right| \leq Z_{crit}^{(p)} \sqrt{W_{kk}} \quad \forall k \end{aligned} \right\} \quad (3)$$

Therefore, considering all possible combinations of bias locations, we write

$$\hat{\delta}_i^{(p,n_T)} = \underset{\forall T}{Max} \delta_i^{(p)}(T) \quad (4)$$

As it was mentioned above, this definition states what the accuracy of the system is, *when and if* a certain number of gross errors are expected to take place. *In other words, it represents the worst case scenario and does not discuss the frequency of such scenario.* We now discuss a new definition and how to obtain an expected value next

3. Stochastic Based Accuracy

We define the stochastic based maximum induced biased as the sum over all possible n_T biases of the expected fraction of time (Γ_{n_T}) in which these biases are present.

$$\hat{\delta}_i^{(p)} = \sum_{n_T} \hat{\delta}_i^{(p,n_T)} E[\Gamma_{n_T}] \quad (5)$$

The formula assumes that a) when errors in a certain number of sensors occur they replace other existing set of undetected errors and that b) Sensors with detected errors are repaired instantaneously.

Sensors have their own failure frequency, which is independent of what happens with other sensors. For example, the probability of one sensor failing at time t , when all sensors were functioning correctly between time zero and time t is $\Phi_i^1 = f_i(t) \prod_{s \neq i} [1 - f_s(t)]$, where $f_i(t)$ is the service reliability function of sensor i if sensors are not repaired. When sensors are repaired, one can use availability and write $f_i(t) = r_i / (r_i + \mu_i)$, where r_i is the repair rate and μ_i is the failure rate. The second issue, the repair time, is more problematic because it also affects the value of $\hat{\sigma}_i$, which becomes the residual precision during that period of time. So, $E[\Gamma_{n_T}]$ can only be estimated by identifying the probability of the state with the frequency of the state in the case of negligible repair time. However, when repair time is significant $E[\Gamma_{n_T}]$ is more difficult to estimate and there are no expressions available.

In addition, multiple gross errors do not arise from a simultaneous event, but rather from a gross error occurring and adding to an existing set of undetected gross errors. In addition, problem (3) assumes the worst case in which all will flag at first, but it does not say what will happen if some are eliminated.

We now define the stochastic-based expected induced biased as the sum over all possible n_T biases of the expected fraction of time (Γ_{n_T}) in which these biases are present.

$$E[\tilde{\delta}_i^{(p)}] = \sum_{n_T} E[\tilde{\delta}_i^{(p, n_T)}] E[\Gamma_{n_T}] \quad (6)$$

To understand how the stochastic-based induced bias (and by extension, the stochastic-based accuracy) can be calculated. Assume that a system is bias free in the period $[0, t_1]$ and that sensor k fails at time t_1 . Thus, if the bias is not detected, then there is an expected induced bias that one can calculate as follows:

$$E[\tilde{\delta}_i^{(p, 1)}(k)] = \int_{-\delta_{k, crit}}^{\delta_{k, crit}} [I - SW]_{ik} \theta_k |h(\theta_k; \bar{\delta}_k, \rho_k)| d\theta_k \quad (7)$$

where $h(\theta_k; \bar{\delta}_k, \rho_k)$ is the pdf of the bias θ_k with mean value $\bar{\delta}_k$ and variance ρ_k . Note that we integrate over all values of θ_k , but we only count absolute values, as the accuracy definition requires. Thus, in between t_1 and the time of the next failure of some sensor t_2 , the system has an accuracy given by $\hat{\sigma}_i + E[\tilde{\delta}_{i, undet}^{(p, 1)}(k)]$.

In turn, if the bias is detected, the sensor is taken out of line for a duration of the repair time R_k . During this time (and assuming no new failure takes place), the system has no induced bias, but it has a lower precision, simply because the measurement is no longer used to perform data reconciliation. Thus, during repair time, the expectation of the accuracy due to detected biases is given by the residual precision $\hat{\sigma}_i^R(k)$. After a period of time R_k the accuracy returns to the normal value when no biases are present $\hat{\sigma}_i$. Thus, in the interval $[0, t_2)$, the accuracy is given by $[\hat{\sigma}_i t_1 + \hat{\sigma}_i^R(k) R_k + \hat{\sigma}_i * (t_2 -$

$t_1 - R_k)/t_2$ when bias k is detected and $[\hat{\sigma}_i t_1 + E[\tilde{\delta}_{i,undet}^{(p,1)}(k)](t_2 - t_1)]/t_2$ when bias k is undetected. The expectation is then given by multiplying the undetected portion by the corresponding probability

$$p_{undet}(k) = \int_{\delta_{k,crit}}^{\delta_{k,crit}} h(\theta_k; \bar{\delta}_k, \rho_k) d\theta_k \quad (8)$$

and the detected by its complement $[1 - p_{undet}(k)]$.

Assume now that the bias in sensor k is undetected at t_1 and another bias in some other sensor r occurs at t_2 , which can be in turn detected or not detected. If it is undetected, then the expected induced bias is given by:

$$E[\tilde{\delta}_i^{(p,2)}(k, r)] = \int_{\delta_{k,crit}}^{\delta_{k,crit}} \int_{\delta_{r,crit}}^{\delta_{r,crit}} |[I - SW]_{ik} \theta_k + [I - SW]_{ir} \theta_r| h(\theta_k; \bar{\delta}_k, \rho_k) h(\theta_r; \bar{\delta}_r, \rho_r) d\theta_k d\theta_r \quad (9)$$

where, for simplicity of presentation we have assumed that $\delta_{k,crit}$ and $\delta_{r,crit}$ can be used as integration limits. (in reality, the integration region is not a rectangle). We leave this detail for future work. In turn, if the error in sensor r is detected, then we assume that the induced bias remains.

Quite clearly, the scenario shown is one of many, and while one is able to obtain the expected induced errors in each case, the problem of calculating the expected fraction of time in each state persists. Thus, we resort to Montecarlo simulations to assess this.

3.1 Montecarlo simulations

Consider a scenario s , composed of a set of n_s values of time $(t_1, t_2, \dots, t_{ns})$ within the time horizon T^h . For each time t_i , one considers a sample of one sensor failing with one of two conditions: its bias is detected or undetected. Sensors that have been biased between t_{i-1} and t_i and where undetected at t_i , continue undetected. Thus, when bias in sensor k is detected, for the time between t_i and $t_i + R_k$ we write

$$E[a_i] = \hat{\sigma}_i^R(k) + E[\tilde{\delta}_i^{(p, m_{i-1})}(l_{1,i-1}, l_{2,i-1}, \dots, l_{m_i, i-1})] \quad (10)$$

where the second term is the expected bias due to the presence of m_{i-1} undetected errors.

$$E[\tilde{\delta}_i^{(p, m_{i-1})}(l_{1,i-1}, l_{2,i-1}, \dots, l_{m_i, i-1})] = \int_{\delta_{1,i-1}}^{\delta_{v,crit}} \dots \int_{\delta_{v,crit}}^{\delta_{v,crit}} \left| \sum_{v=1}^n z(v) [I - SW]_{iv} \theta_v \right| \prod_v h(\theta_v; \bar{\delta}_v, \rho_v) d\theta_v \quad (11)$$

For the interval $(t_i + R_k, t_{i+1})$, we write

$$E[a_i] = \hat{\sigma}_i + E[\tilde{\delta}_{i,undet}^{(p, m_{i-1})}(l_{1,i-1}, l_{2,i-1}, \dots, l_{m_i, i-1})] \quad (12)$$

In turn, if the error was not detected, then we write t_{i+1} , we write

$$E[a_i] = \hat{\sigma}_i(k) + E[\tilde{\delta}_i^{(p, m_{i-1})}(l_{1,i-1}, l_{2,i-1}, \dots, l_{m_i, i-1})] \quad (13)$$

The above formula is valid for $k \neq l_{v,i-1}, v=1, \dots, m_{i-1}$. Otherwise, the same formula is used, but k is removed from $\tilde{\delta}_i^{(p, m_{i-1})}(l_{1,i-1}, l_{2,i-1}, \dots, l_{m_{i-1},i-1})$.

To obtain an average accuracy of the system in the horizon T^h and for the scenario s , the accuracy in each interval or sub-interval is multiplied by the duration of such interval and divided by the time horizon T^h . Finally all the values are added to obtain the expectation for that scenario. The final accuracy is obtained using the average of all scenarios. Finally, scenarios are sampled the following way. For each sensor a set of failure times is obtained by sampling the reliability function repeatedly and assuming that sensors are as good as new after repair (AGAN maintenance). Of these, undetectability is sampled using a pdf given by $p_{undet}(k)$ and its complement.

4. Example

Consider the example of figure 1. Assume flowmeters with $\sigma_i=1, 2$ and 3 , respectively. We also assume that the biases have zero mean and standard deviation $\rho_k=2, 4$ and 6 respectively, failure rate of $0.025, 0.015, 0.005$ (1/day) and repair time of $0.5, 2$ and 1 day respectively. The system is barely redundant (Only one gross error can be determined, and when it is flagged by the measurement test, hardware inspection is needed to obtain its exact location. This is due to gross error equivalency (equivalency theory: Bagajewicz and Jiang, 1998).

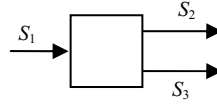


Figure 1. Example

The problem was run with scenarios containing 20 event samples. A portion of one such sample is for example depicted in Table 1. Convergence is achieved very quickly (see figure 2) to a value of accuracy of 1.89. (The solid line is the average value). Comparatively the accuracy defined for maximum bias of one bias present is 6.30. This highlights the fact that using a maximum expected undetected bias is too conservative

5. Discussion and Conclusions

The problems with an existing definition of accuracy have been highlighted and a new definition, which gives a more realistic value, has been presented. In addition a Montecarlo sampling technique was suggested to determine the value of the accuracy. Some shortcomings still remain: The expected value of existing undetected biases is determined using rectangular integration regions, when it is known these regions have other more complex forms. This can be addressed analytically somehow, but one can also resort to sample the bias sizes as well. All this is part of ongoing work.

Table 1. Example of one scenario (Portion)

Time	Bias in sensor	Bias detected
15.6	S_1	No
43.6	S_1	No
62	S_2	Yes
90	S_2	Yes
100	S_2	Yes
115	S_1	Yes
150	S_3	Yes
160	S_1	Yes
170	S_2	No
185	S_2	No
189	S_1	Yes
193	S_1	No
208	S_2	Yes

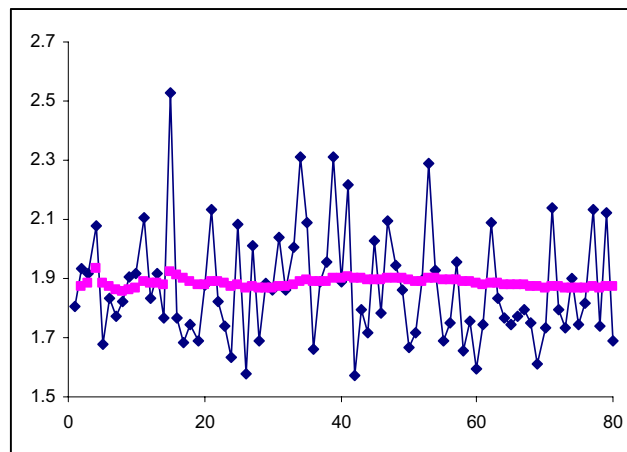


Figure 2. Montecarlo Iterations convergence.

References

- Bagajewicz, M., 2004, On the Definition of Software Accuracy in Redundant Measurement Systems. To appear. AIChE J., (available at <http://www.ou.edu/class/che-design/Unpublished-papers.htm>).
- Bagajewicz M. and Q. Jiang. *Gross Error Modelling and Detection in Plant Linear Dynamic Reconciliation*. Computers and Chemical Engineering, 22, 12, 1789-1810 (1998).
- Miller R. W. *Flow Measurement Engineering Handbook*. McGraw Hill, (1996)

Acknowledgements

Funding from the US-NSF, Grant CTS-0350501, is acknowledged.