

Total Least-Squares Collocation for Deformation Analysis

Kyle SNOW^{1,*} and Burkhard SCHAFFRIN²

¹ Polaris Geospatial Services, Westerville, Ohio, USA, (ksnow6378137@gmail.com)

² Division of Geodetic Science, School of Earth Sciences, The Ohio State University, Columbus, Ohio, USA, (baschaffrin@gmail.com)

* corresponding author

Abstract

For many situations in deformation analysis, data are collected that not only enter the usual observation vector $\mathbf{y}$, but also the coefficient matrix A after linearization. Such models fall into the category of Errors-In-Variables (EIV) Models and may be treated by Total Least-Squares (TLS) adjustment. Moreover, if the deformation as described by the parameters ξ follows certain “expectations” that can be quantified, the normally non-random parameters will turn into “random effects” $\mathbf{x}$, in which case the standard estimation procedure needs to be replaced by one that resembles collocation, but is here obviously based on the TLS principle. Earlier studies involved only a non-singular combination of the observation covariance matrices; these include Schaffrin (2009), Snow and Schaffrin (2012), and Schaffrin (2020). A more general treatment was developed by Snow (2012) in his PhD dissertation and has been further extended in Snow and Schaffrin (2025) for the case where the combination of covariance matrices turns out to be singular. Here, an algorithm is presented and applied to a problem in deformation analysis.

Keywords: Errors-In-Variables (EIV) model, Random Effects Model (REM), Prior information, Total Least-Squares (TLS) adjustment, TLS collocation, Singular covariance matrices, Deformation analysis

Received: 9th December 2024. Revised: 2nd March 2025. Accepted: 24th March 2025.

1 Introduction

Certain problems with measurement variables appearing in both the observation vector $\mathbf{y}$ and the coefficient matrix A , and that also include stochastic prior information for the unknown parameters, are perhaps best treated by the *errors-in-variables with random effects model* (EIV-REM). As the name suggests, the model is formed by combining the EIV model (cf. Snow, 2012 and Jazaeri et al., 2014) with the REM (cf. Schaffrin, 2001) in a suitable way.

It may be insightful to briefly reflect on certain signature characteristics of the REM and EIV models individually before focusing on their combination into a single model. That was done by Schaffrin (2020) who characterized respective REM and EIV models in relationship to the classical Gauss-Markov Model (GMM). There, he described the REM as a model that *strengthens* the parameters

through the introduction of stochastic prior information, and he called the least-squares solution within the REM *least-squares collocation* following Moritz (1970). In contrast, he characterized the EIV model as one that *weakens* the coefficient matrix of the parameters by allowing it to contain observed data rather than only fixed entries, hence the often used term “data matrix.” In this case, he called the associated least-squares estimation technique *total least-squares (TLS) estimation*. Schaffrin (2020) nicely illustrated these relationships through a diagram that had been presented in an earlier workshop (Schaffrin, 2009) and is repeated in Fig. 1 herein.

Building on the work of Schaffrin (2009), for example, Snow (2012, §5.2.1) developed a more general EIV-REM model that allowed correlation between data in the observation vector and the data matrix. Moreover, the cofactor (scaled covariance) matrices for those data were allowed to be singular (i.e., positive semidefinite) as long as a certain combination

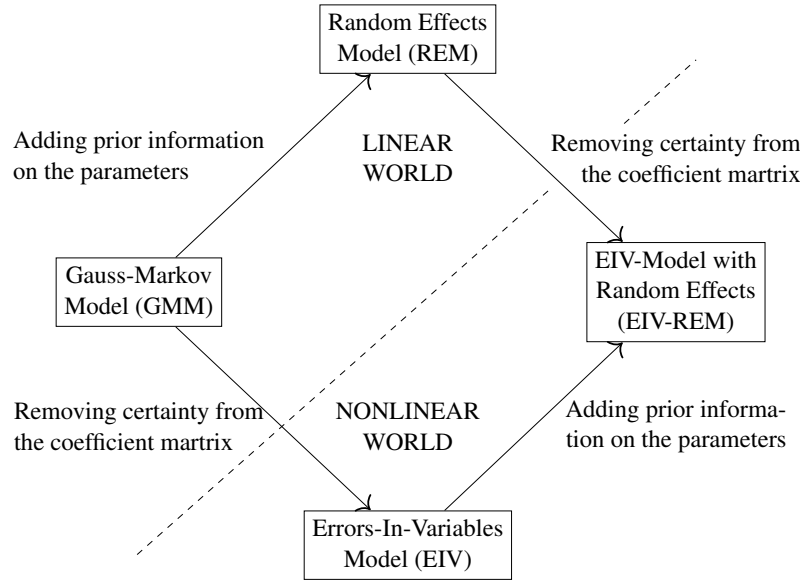


Figure 1. The most informative model (REM) is at the top; the least informative model (EIV) is at the bottom. Like the GMM, the EIV-REM is formed at the intermediate level (in terms of the information it includes), but it belongs to the “nonlinear world,” where nonlinear normal equations are formed and subsequently solved by iteration. The directions of the arrows are consistent with the statements adjacent to them.

of them—defined in (9b) below as Q_1 —turned out to be nonsingular. *In the following, we generalize Snow’s work further to handle the case where the matrix Q_1 is singular.* The Kronecker-product structure of Schaffrin (2020) for Q_A in (1c) below is also no longer required.

The remainder of the paper is organized as follows: Section 2 defines the EIV-REM in detail and presents a thorough collection of formulas for the prediction of the model parameters and the adjustment residuals based on Snow and Schaffrin (2025). It also provides a compact algorithm for computer code and briefly discusses some points about it before section 3 presents an example in a variety of case settings. Finally, section 4 summarizes the contributions of this paper and notes some outstanding questions that are the focus of our ongoing work.

2 Formula collection

As noted above, certain problems with measurement variables appearing in both the observation vector $\mathbf{y}$ and the coefficient matrix A are perhaps best treated by the *errors-in-variables with random effects model* (EIV-REM) when stochastic prior information is available for the parameters. Such a

model is formed by combining the EIV model (cf. Snow, 2012; Fang, 2013; and Jazaeri et al., 2014) with the REM (cf. Schaffrin, 2001), viz.

$$\mathbf{y}_{n \times 1} = \underbrace{(A - E_A)}_{n \times m} \mathbf{x} + \mathbf{e}_y, \quad (1a)$$

$$\beta_0 = \mathbf{x} + \mathbf{e}_0, \quad (1b)$$

$$\begin{bmatrix} \mathbf{e}_y \\ \mathbf{e}_A \\ \mathbf{e}_0 \end{bmatrix} \sim \left(\begin{bmatrix} \mathbf{0} \\ \mathbf{0} \\ \mathbf{0} \end{bmatrix}, \sigma_0^2 \begin{bmatrix} Q & 0 \\ 0 & Q_0 \end{bmatrix} \right) := \sigma_0^2 \left(\begin{array}{cc|c} Q_y & Q_{yA} & 0 \\ Q_{Ay} & Q_A & 0 \\ \hline 0 & 0 & Q_0 \end{array} \right). \quad (1c)$$

The equations in (1a) are often called observation equations; (1b) provides prior information for the unknown, random parameters $\mathbf{x}$, and (1c) expresses the distribution of the unknown, random errors $\mathbf{e}_y$, $\mathbf{e}_A := \text{vec } E_A$, and $\mathbf{e}_0$. The terms in (1) are defined more specifically as follows:

$\mathbf{y}$ is an $n \times 1$ vector of observations (“data vector”).

A is an $n \times m$ matrix of random variables (“data matrix”), possibly containing some fixed (known) elements, for which variances and covariances are set to zero.

E_A is an $n \times m$ matrix of unknown random errors associated with the data in A .

$\mathbf{x}$ is an $m \times 1$ vector of unknown random parameters (random effects, aka “signal”).

$\mathbf{e}_y$ is an $n \times 1$ vector of unknown random errors associated with the data vector $\mathbf{y}$.

β_0 is a given $m \times 1$ vector of prior information.

$\mathbf{e}_A$ is the vectorial form of E_A (i.e., $\mathbf{e}_A := \text{vec } E_A$) with size $nm \times 1$.

$\mathbf{e}_0$ is an $m \times 1$ vector of unknown random errors associated with the prior information.

σ_0^2 is an unknown variance component.

Q_y is a given $n \times n$ symmetric, positive (semi)definite cofactor matrix for $\mathbf{e}_y$.

Q_A is a given $nm \times nm$ symmetric, positive (semi)definite cofactor matrix for $\mathbf{e}_A$.

Q_{yA} is a given $n \times nm$ matrix that accounts for correlations, if any, between the data in $\mathbf{y}$ and A , with $Q_{yA}^T = Q_{Ay}$.

Q_0 is a given $m \times m$ symmetric, positive (semi)definite cofactor matrix for $\mathbf{e}_0$.

The model as presented does not allow for correlations between the observation equations and the prior information, as they commonly come from different, and entirely independent, sources. Nevertheless, it does admit much more general cofactor matrices than the form shown in Schaffrin and Snow (2013), who focused their attention there to mean squared error risk.

It is noted that both the coefficient (data) matrix A and the cofactor matrix Q_0 for the prior information could be rank deficient, as long as the matrix $[A^T | Q_0]$ has full row rank m . Thus, the model redundancy r is defined by

$$r := n - \text{rk}[A^T | Q_0] + \text{row-dim}(\mathbf{x}) = n - m + m = n, \quad (2)$$

assuming $\text{rk } A = \text{rk}(A - E_A)$.

Total least-squares collocation Assuming momentarily that the cofactor matrix Q is nonsingular, we define a *weight matrix* P for the observation equations (1a) as

$$P := Q^{-1} = \begin{bmatrix} Q_y & Q_{yA} \\ Q_{Ay} & Q_A \end{bmatrix}^{-1} = \begin{bmatrix} P_{11} & P_{12} \\ P_{21} & P_{22} \end{bmatrix}. \quad (3)$$

Further assuming, momentarily, that the cofactor matrix Q_0 is non-singular, the predictor $\tilde{\mathbf{x}}$ of the unknown vector of random effects $\mathbf{x}$ can be derived from the principle of Total Least Squares (TLS) as shown in Snow (2012). Starting with the statement

$$\mathbf{e}_y^T P_{11} \mathbf{e}_y + 2\mathbf{e}_y^T P_{12} \mathbf{e}_A + \mathbf{e}_A^T P_{22} \mathbf{e}_A + \mathbf{e}_0^T Q_0^{-1} \mathbf{e}_0 = \min \quad (4)$$

and subjecting it to the model (1), allows formation of the Lagrange target function

$$\begin{aligned} \Phi(\mathbf{e}_y, \mathbf{e}_A, \mathbf{e}_0, \boldsymbol{\lambda}) = & \mathbf{e}_y^T P_{11} \mathbf{e}_y + 2\mathbf{e}_y^T P_{12} \mathbf{e}_A + \\ & + \mathbf{e}_A^T P_{22} \mathbf{e}_A + \mathbf{e}_0^T Q_0^{-1} \mathbf{e}_0 + 2\boldsymbol{\lambda}^T (\mathbf{y} - A\beta_0 - \mathbf{e}_y + \\ & + (\beta_0^T \otimes I_n) \mathbf{e}_A + A\mathbf{e}_0 - E_A \mathbf{e}_0) = \text{stationary}, \end{aligned} \quad (5)$$

where $\boldsymbol{\lambda}$ is an $n \times 1$ vector of Lagrange multipliers to be estimated.

Forming the Euler-Lagrange (or first-order) necessary conditions results in the set of normal equations

$$P_{11} \tilde{\mathbf{e}}_y + P_{12} \tilde{\mathbf{e}}_A - \hat{\boldsymbol{\lambda}} = \mathbf{0}, \quad (6a)$$

$$P_{21} \tilde{\mathbf{e}}_y + P_{22} \tilde{\mathbf{e}}_A + (\beta_0 \otimes I_n) \hat{\boldsymbol{\lambda}} - (\tilde{\mathbf{e}}_0 \otimes I_n) \hat{\boldsymbol{\lambda}} = \mathbf{0}, \quad (6b)$$

$$Q_0^{-1} \tilde{\mathbf{e}}_0 + A^T \hat{\boldsymbol{\lambda}} - \tilde{E}_A^T \hat{\boldsymbol{\lambda}} = \mathbf{0}, \quad (6c)$$

$$\mathbf{y} - A\beta_0 - \tilde{\mathbf{e}}_y + \tilde{E}_A \beta_0 + (A - \tilde{E}_A) \tilde{\mathbf{e}}_0 = \mathbf{0}. \quad (6d)$$

After some algebraic manipulations, following Snow (2012, pp. 43–44), predictions of the unknown random errors $\mathbf{e}_y$ and $\mathbf{e}_A$ (residuals), expressed in terms of the cofactor matrices and the estimated vector of Lagrange multipliers $\hat{\boldsymbol{\lambda}}$, are provided by

$$\tilde{\mathbf{e}}_y = \{-Q_{yA}[(\beta_0 - \tilde{\mathbf{e}}_0) \otimes I_n] + Q_y\} \hat{\boldsymbol{\lambda}}, \quad (7)$$

$$\tilde{\mathbf{e}}_A = \{Q_{Ay} - Q_A[(\beta_0 - \tilde{\mathbf{e}}_0) \otimes I_n]\} \hat{\boldsymbol{\lambda}}, \quad (8)$$

with $\tilde{E}_A = \text{Invec } \tilde{\mathbf{e}}_A$. Here, the Invec operator reverses the operation of the vec operator by reshaping its argument back to the size of the argument in the original vec operation.

For compactness, we substitute $\tilde{\mathbf{x}}$ for $\beta_0 - \tilde{\mathbf{e}}_0$ and, following Snow and Schaffrin (2025), arrive at the estimated Lagrange multipliers

$$\hat{\boldsymbol{\lambda}} = Q_1^{-1} (\mathbf{y} - A\tilde{\mathbf{x}}) \quad \text{if } Q_1^{-1} \text{ exists}, \quad (9a)$$

with the $n \times n$ matrix Q_1 defined by

$$\begin{aligned} Q_1 := & [Q_y - Q_{yA}(\tilde{\mathbf{x}} \otimes I_n) - (\tilde{\mathbf{x}} \otimes I_n)^T Q_{Ay} + \\ & + (\tilde{\mathbf{x}} \otimes I_n)^T Q_A (\tilde{\mathbf{x}} \otimes I_n)]. \end{aligned} \quad (9b)$$

For the case where Q_1 turns out to be singular, we introduce an $m \times m$ symmetric, positive (semi)definite matrix S and use it to define the $n \times n$ matrix

$$Q_3 := [Q_1 + (A - \tilde{E}_A)S(A - \tilde{E}_A)^T] = Q_3(\tilde{\mathbf{x}}, \hat{\lambda}), \quad (10)$$

which is non-singular if the rank condition $\text{rk}[Q_1, (A - \tilde{E}_A)S] = n$ holds.

We note the relationship $\tilde{\mathbf{e}}_0 = \beta_0 - \tilde{\mathbf{x}}$ and write the following two equation for the prediction of $\mathbf{x}$, as in Snow and Schaffrin (2025):

$$\tilde{\mathbf{x}} = \beta_0 + [(A - \tilde{E}_A)^T Q_3^{-1} (A - \tilde{E}_A) (Q_0 - S) Q_0^{-1} + Q_0^{-1}]^{-1} (A - \tilde{E}_A)^T Q_3^{-1} (\mathbf{y} - A\beta_0 + \tilde{E}_A \tilde{\mathbf{e}}_0), \quad (11)$$

$$\tilde{\mathbf{x}} = \beta_0 + Q_0 [I_m + (A - \tilde{E}_A)^T Q_3^{-1} (A - \tilde{E}_A) \cdot (Q_0 - S)]^{-1} (A - \tilde{E}_A)^T Q_3^{-1} (\mathbf{y} - A\beta_0 + \tilde{E}_A \tilde{\mathbf{e}}_0). \quad (12)$$

Obviously, (11) requires Q_0 to be nonsingular, whereas (12) can be used for singular or nonsingular Q_0 . Incidentally, if $S = 0$, implying that Q_1 is nonsingular, then (10) to (12) reduce to equations (5.24)–(5.25b) of Snow (2012). On the other hand, if the choice $S := Q_0$ keeps

$$Q_{30} := [Q_1 + (A - \tilde{E}_A)Q_0(A - \tilde{E}_A)^T] \quad (13)$$

invertible, the formulas (11) and (12) can be further simplified to

$$\tilde{\mathbf{x}} = \beta_0 + Q_0 (A - \tilde{E}_A)^T Q_{30}^{-1} (\mathbf{y} - A\beta_0 + \tilde{E}_A \tilde{\mathbf{e}}_0), \quad (14)$$

while the corresponding residual vector would then read

$$\tilde{\mathbf{e}}_0 = -Q_0 (A - \tilde{E}_A)^T Q_{30}^{-1} (\mathbf{y} - A\beta_0 + \tilde{E}_A \tilde{\mathbf{e}}_0) = \beta_0 - \tilde{\mathbf{x}}. \quad (15)$$

After prediction of the random errors and estimation of the Lagrange multipliers, the *total sum of squared residuals* (TSSR), denoted herein by Ω , can be computed by

$$\Omega = [\tilde{\mathbf{e}}_y^T, \tilde{\mathbf{e}}_A^T] \begin{bmatrix} P_{11} & P_{12} \\ P_{21} & P_{22} \end{bmatrix} \begin{bmatrix} \tilde{\mathbf{e}}_y \\ \tilde{\mathbf{e}}_A \end{bmatrix} + \tilde{\mathbf{e}}_0^T Q_0^{-1} \tilde{\mathbf{e}}_0 = \quad (16a)$$

$$= \hat{\lambda}^T (Q_1 + (A - \tilde{E}_A)Q_0(A - \tilde{E}_A)^T) \hat{\lambda} = \quad (16b)$$

$$= \hat{\lambda}^T Q_{30} \hat{\lambda}. \quad (16c)$$

It is sometimes the case that columns of the data matrix A are fixed, resulting in a singular cofactor

matrix Q in the EIV model (1), in which case the partitioned weight matrix in (16a) would not even exist, though Q_1 always exists and may still be regular. Nevertheless, (16b) can still be used to compute the TSSR, even if the cofactor matrix Q_0 for the prior information is singular, too. In any case, a suitable formula for estimating the unknown variance component σ_0^2 shown in (1c) is provided by

$$\hat{\sigma}_0^2 = \Omega/n, \quad (17)$$

where n is the model redundancy according to (2).

An algorithm for numerical computations Algorithm 1 below can be considered for numerical computations. It is straight forward and can easily be compared to the formulas developed in the previous section. Note that the use of zero in superscripts denotes an initial value or a value from the previous iteration, and use of a left arrow denotes that the variable on the left is to be initialized or updated to the values contained in the variable on the right. A more detailed list of comments about the algorithm can be found in Snow and Schaffrin (2025).

3 Example

In order to highlight the utility of the presented algorithm, we apply it here to a problem of estimating the two-dimensional (2D) strain tensor of a body undergoing homogeneous deformation. Such a body may be natural, e.g., the crust of the earth, or man-made, e.g., a bridge or a dam. In any case, the deformation phenomenon is often realized from the change in estimated coordinates, of monuments anchored to the body, between two epochs of time sufficiently separated to allow for deformation during the time interval. The coordinates are most often obtained from a free network adjustment of surveying measurements.

In 2D, a *network datum* is specified by a scale, one orientation, and two origin parameters. In a free network adjustment, the observations do not carry sufficient information to completely define the network's datum, and thus the network has a *datum defect*. For example, if data for a 2D network is comprised of only distance and direction observations, a datum defect of three occurs, since those observations supply no information about the origin or orientation of the network.

Algorithm 1 TLS Collocation**Given:**

- 1: Model variables $\mathbf{y}, A, \beta_0, Q_y, Q_A, Q_0$;
- 2: Optionally provide $Q_{yA} = 0, S = I_m, \delta = 1.0 \times 10^{-12}$ (with reasonable default values shown)

Initialize:

- 3: $\tilde{\mathbf{x}}^{(0)} \leftarrow \beta_0, \tilde{\mathbf{e}}_0^{(0)} \leftarrow \mathbf{0}_m, \tilde{E}_A^{(0)} \leftarrow 0, \hat{\boldsymbol{\lambda}}^{(0)} \leftarrow \mathbf{0}_n$
- 4: $Q = \begin{bmatrix} Q_y & Q_{yA} \\ Q_{Ay} & Q_A \end{bmatrix}$
- 5: $\Delta\tilde{\mathbf{x}} = \mathbf{1}_m$
- 6: **while** $\|\Delta\tilde{\mathbf{x}}\| > \delta$ **do**
- 7: $B := [I_n \mid -(\tilde{\mathbf{x}}^{(0)} \otimes I_n)^T], \quad Q_1 = BQB^T$
- 8: **if** $\text{rank } Q_1 < n$ **then**
- 9: $Q_3 = Q_1 + (A - \tilde{E}_A^{(0)})S(A - \tilde{E}_A^{(0)})^T$
- 10: $R \leftarrow$ inverse of Cholesky factor of Q_3
- 11: **else**
- 12: $R \leftarrow$ inverse of Cholesky factor of Q_1
- 13: **end if**
- 14: $W := RR^T$, being the inverse of Q_1 or Q_3
- 15: $Z := (A - \tilde{E}_A^{(0)})^T W$
- 16: $\tilde{\mathbf{e}}_0 = -Q_0[I_m + Z(A - \tilde{E}_A^{(0)})(Q_0 - S)]^{-1}Z \cdot (\mathbf{y} - A\beta_0 + \tilde{E}_A^{(0)} \cdot \tilde{\mathbf{e}}_0^{(0)})$
- 17: $\tilde{\mathbf{x}} = \beta_0 - \tilde{\mathbf{e}}_0$
- 18: **if** $\text{rank } Q_1 = n$ **then** $\hat{\boldsymbol{\lambda}} = W(\mathbf{y} - A\tilde{\mathbf{x}})$
- 19: **else if** Q_0 is singular **then**
- 20: $\hat{\boldsymbol{\lambda}} = W\{\mathbf{y} - A\tilde{\mathbf{x}} + (A - \tilde{E}_A^{(0)})S[(A - \tilde{E}_A^{(0)})^T \hat{\boldsymbol{\lambda}}^{(0)}]\}$
- 21: **else** $\hat{\boldsymbol{\lambda}} = W[(\mathbf{y} - A\tilde{\mathbf{x}}) - (A - \tilde{E}_A^{(0)})SQ_0^{-1}\tilde{\mathbf{e}}_0]$
- 22: **end if**
- 23: $\tilde{\mathbf{e}}_A = [Q_{Ay} - Q_A(\tilde{\mathbf{x}} \otimes I_n)]\hat{\boldsymbol{\lambda}}, \tilde{E}_A = \text{Invec}(\tilde{\mathbf{e}}_A)$
- 24: **Update:** $\Delta\tilde{\mathbf{x}} = \tilde{\mathbf{x}} - \tilde{\mathbf{x}}^{(0)}, \tilde{\mathbf{x}}^{(0)} \leftarrow \tilde{\mathbf{x}}, \tilde{\mathbf{e}}_0^{(0)} \leftarrow \tilde{\mathbf{e}}_0, \tilde{E}_A^{(0)} \leftarrow \tilde{E}_A, \hat{\boldsymbol{\lambda}}^{(0)} \leftarrow \hat{\boldsymbol{\lambda}}$
- 25: **end while**
- 26: Perform a check of the model by confirming that equations (6c) and (6d) are satisfied.

Furthermore, a free network adjustment yields a minimum bias of estimated coordinates while also providing unique observation residuals. It does so by minimizing changes in some or all of the coordinates from their given values (minimum Euclidean norm). Consequently, it also yields a singular covariance matrix for the estimated coordinates that has a rank deficiency usually equal to the datum defect of the network. The singular covariance matrix is then to be used in the subsequent estimation of the deformation strain tensor as described below.

The observation equations of point coordinates from two independent epochs expressed as a function of

strain-tensor elements and other parameters have been presented by various authors, including Brunner et al. (1981), Caspary (2000), and Nkuite (1998). For the i th point of $n/2$ points, at each of two epochs denoted 1 and 2, having estimated coordinates (x_{1i}, y_{1i}) and (x_{2i}, y_{2i}) , respectively, the expectation of the observation equation in matrix form reads

$$E\left\{\begin{bmatrix} x_{2i} - x_{1i} \\ y_{2i} - y_{1i} \end{bmatrix}\right\} = E\left\{\begin{bmatrix} x_{1i} & 0 & y_{1i} & -y_{1i} & 1 & 0 \\ 0 & y_{1i} & x_{1i} & x_{1i} & 0 & 1 \end{bmatrix}\right\} \cdot E\{\mathbf{x}\}, \quad (18)$$

where E denotes expectation, and the elements in the 6×1 vector of unknown random effects $\mathbf{x} := [\epsilon_{xx} \ \epsilon_{yy} \ \epsilon_{xy} \ \omega \ t_x \ t_y]^T$ are defined as follows: ϵ_{xx} and ϵ_{yy} are extensional strains along the x and y axes, respectively; the effect of ϵ_{xy} is called shear strain; ω is a rotation angle; and t_x and t_y are translation parameters along the x and y axes, respectively. According to Caspary (2000, p. 137), the quantity $2 \cdot \epsilon_{xy}$ “is equivalent to the angular distortion of a right angle which was originally parallel to the axes of the coordinate system.”

For all $n/2$ points, the vector on the left side of (18) extends to size $n \times 1$ and the matrix on the right side to size $n \times m, m = 6$. These quantities are the observation vector $\mathbf{y}$ and data matrix A , resp., of (1a).

Let V_1 and V_2 denote the singular covariance matrices from free network adjustments at epoch 1 and 2 for the coordinates appearing in (18). Then the cofactor matrix Q appearing in the model (1) is defined as follows. Define the 2×2 matrices

$$B_1 := \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix}, B_2 := \begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix}, B_3 := \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}, \quad (19)$$

$$B_4 := \begin{bmatrix} 0 & -1 \\ 1 & 0 \end{bmatrix}, B_5 = B_6 := \begin{bmatrix} 0 & 0 \\ 0 & 0 \end{bmatrix}$$

and, with $p := n/2$, the following six $n \times n$ (sparse) matrices

$$C_i := I_p \otimes B_i, \quad i = 1, \dots, 6, \quad (20)$$

so that finally we can write

$$Q := \left[\begin{array}{c|c} V_2 & 0 \\ \hline 0 & 0 \end{array} \right]_{\substack{n \times n \\ 6n \times n}} + [-I_n | C_1^T \ C_2^T \ C_3^T \ C_4^T \ C_5^T \ C_6^T]^T \cdot V_1 \cdot [-I_n | C_1^T \ C_2^T \ C_3^T \ C_4^T \ C_5^T \ C_6^T]. \quad (21)$$

Note that the last $2n$ rows and columns are zero, owing to the last two columns of the data matrix A being fixed, and that Q_{yA} is nonzero.

Table 1. Simulated strain-field parameters for generating coordinates at epoch 2 shown in Table 2

Quantity	Value	$\pm$ Noise	Unit
ϵ_{xx}	-5.5×10^{-6}	1.0×10^{-6}	strain
ϵ_{yy}	4.5×10^{-6}	0.5×10^{-6}	strain
ϵ_{xy}	4.5×10^{-6}	0.9×10^{-6}	strain
ω	0.0	3.0×10^{-2}	grad
t_x	0.1	0.013	m
t_y	0.0	0.013	m

Numerical example Our numerical example is adapted from Caspary (2000, section 10.11), where we selected the following eight points from his network, being among those that belong to the part of the network considered to be undergoing deformation: 13, 15, 17, 19, 35, 37, 45, 47. Coordinates for the points at epoch 1 are taken from Caspary's Table 10.9. Coordinates for epoch 2 were obtained by simulating the deformation field using values in column 2 of Table 1 (also appearing in Caspary's Table 10.12). Noise was added to those values using a uniform pseudorandom distribution over the range of negative to positive values shown in column 3 of the same table. Additional pseudorandom noise over the interval ± 2 cm was also added to simulate variances associated with a second measurement campaign.

The resulting coordinates used for the analysis herein are shown in Table 2 below to the exact precision used. We obtained a covariance matrix, V_1 , for the points by adjusting all observations of Caspary's Tables 10.8a and 10.8b associated with the mentioned eight points. Those data involve only distance and direction observations. Thus the resulting 16×16 covariance matrix V_1 for the coordinates has a rank of 13, corresponding to a network datum defect of three. To keep things simple, we used the same covariance matrix for both sets of coordinates, so that $Q_y = 2V_1$, and Q_A and Q_{yA} are as reflected in (21). The upper-triangular portion of the covariance matrix V_1 is given in Table 5 to the same precision used in our analysis.

In practice, prior information may come from a previous estimation process, a publication by an authoritative agency, or some other source, depending on the problem at hand. In some cases, it might even be based on past experience or a hunch; especially for the variances if they are not readily available.

Table 2. Point numbers and coordinates in meters for epoch 1 and differences between epochs.

No.	x_1	y_1	$x_2 - x_1$	$y_2 - y_1$
13	-15 350.0270	15 659.9289	3.8120	3.4851
15	-18 220.0625	-7150.0583	-1.4687	4.0274
17	-46 450.0713	-15 850.0411	-3.3303	10.2741
19	-68 270.0590	2829.9345	1.0754	15.2358
35	-24 130.0146	33 609.9813	8.0043	5.5251
37	-35 500.0516	6059.9617	1.6909	7.9572
45	-35 850.0108	44 120.0219	10.5022	8.2131
47	-49 930.0410	25 210.0489	6.1709	11.2541

Cases of prior information (p.i.) Below we present the results of four cases of prior information. Except in case 4, the vector of prior information β_0 is taken from column 2 of Table 1. The covariance matrix Q_0 is diagonal and nonsingular, but it varies from case to case as described below. In contrast, the same matrix S was used in all cases, defining it as a diagonal matrix with the squares of the values in column 3 of Table 1 for its diagonal elements.

For numerical stability, we scaled the coordinates and the translation parameters of the prior information by 1.0×10^{-3} (essentially converting them to units of km) for the adjustment, while leaving their respective covariance matrices in units of squared meters. This scaling required the variances of the non-translation parameters appearing in Q_0 to be scaled by 1.0×10^6 to maintain a consistent ratio of variables to their variances. The residuals were then scaled by 1.0×10^3 after the adjustment (converting them back to meters) for computation of the TSSR according to (16b). Alternatively, we could have achieved a similar level of numerical stability by scaling the cofactor matrices Q and Q_0 by 1.0×10^6 and leaving the coordinates in units of meters.

Case 1 For comparison purposes, no prior information is used. The problem is solved by TLS within the EIV model using algorithm 3 of Snow (2012).

Case 2 The diagonal elements of Q_0 are the squares of the values listed under case 2 in Table 3. Here the variances are increased to very large values, thereby greatly reducing the influence of the prior information.

Case 3 The diagonal elements of Q_0 are the squares of the values listed under case 3 in Table 3.

Here the variances are decreased to very small values, thereby causing the prior information to dominate the prediction of $\mathbf{x}$.

Case 4 The diagonal elements of Q_0 are the squares of the values listed under case 4 in Table 3. However, in this case the prior information for the strain parameters is modified to $\epsilon_{xx} = -7 \times 10^{-6}$ and $\epsilon_{yy} = \epsilon_{xy} = 3 \times 10^{-6}$. This was done to see the relative influence of the observations and prior information different than that used to simulate the coordinates of epoch 2.

Table 3. Square roots of the six diagonal elements of Q_0 (i.e., SDs) for various cases of p.i. The units for the first three rows are strain. The units for row 4a are radians; they are converted to grad in row 4b for illustrative purposes, where $200 \text{ grad} = \pi \text{ radians}$. The units for rows 5 and 6 are meters. Scaling of the first four rows by 1.0×10^6 as noted in the discussion above about numerical stability is reflected here.

	Case 2	Case 3	Case 4
1	1.0×10^{-1}	1.0×10^{-8}	1.0×10^{-3}
2	1.0×10^{-1}	0.5×10^{-8}	1.0×10^{-3}
3	1.0×10^{-1}	0.9×10^{-8}	1.0×10^{-3}
4a	6.283 185	1.57×10^{-4}	1.570 796
4b	4.0×10^2	1.0×10^{-2}	1.0×10^2
5	1.0×10^3	1.3×10^{-2}	1.0
6	1.0×10^3	1.3×10^{-2}	1.0

The numerical results of the four cases are shown in Table 4. Again, case 1 is merely shown for comparison purposes.

Discussion Not surprisingly, case 2 matches case 1 almost exactly. This is to be expected, since the prior information has been drastically de-weighted, and thus it has virtually no influence on the prediction of $\mathbf{x}$. In contrast, case 3 completely replicates the prior information for the strain parameters. This is also to be expected, since the weights (inverses of variances) have been magnified so greatly that the prior information completely dominates the prediction of $\mathbf{x}$. The model check—norm of (6d)—is much larger than the other cases. Perhaps the tiny values in Q_0 make the system less stable. Nevertheless, the larger value may still be acceptable given the noise level of the data. Likewise, the tiny values in Q_0 have greatly magnified the TSSR, though the residuals themselves (not

shown herein) still look reasonable. Note that the other parameters can also be driven to their given values by further decreasing their variances a few orders of magnitude, but doing so drives the TSSR and model check to much higher values.

Case 4 differs from the others in that the prior information in β_0 has been modified from the simulated values as noted above. For context, we note that the typical value on the diagonal of Q_y is about $(0.035 \text{ m})^2$. Note that the observational data would tend to pull ϵ_{xx} towards $-4.54 \mu\text{-strain}$, whereas the prior information would pull it towards $-7 \mu\text{-strain}$. Table 4 does show a more-or-less middle value of $-6.12 \mu\text{-strain}$, reflecting a balance between the observational data and the prior information when reasonable weights are assigned.

4 Conclusions

The formulas developed by Snow and Schaffrin (2025) for predicting the vector of random effects $\mathbf{x}$ within the EIV model with prior information via total least-squares collocation represent a more general treatment than what has been presented previously, in that the combined cofactor matrix Q_1 can now be singular. An algorithm for the numerical solution of the problem has been presented and tested on a simulated deformation problem. It behaves as expected when the cofactor matrix Q_0 for the prior information takes on extreme values, and it appears to strike a good balance between the influence of the observational data and the prior information when those components of the model are more realistically weighted.

In any case, future work should consider the case where prior information is only available for part of the unknown parameters, so that only some elements of the parameter vector become random. Also, the mean square error matrix for the prediction of $\mathbf{x}$ needs to be developed so that the accuracy of the prediction can be better characterized, though this will likely prove difficult due to the non-linear nature of the normal equations, just as it has for deriving the dispersion matrix for the total least-squares estimate of the unknown (fixed) parameters within the EIV model.

Table 4. Solutions of four cases. Case 1 uses no prior information (p.i.). Cases 2 and 3 take their p.i. from column 2 of Table 1. For case 4, p.i. is modified so that $\epsilon_{xx} = -7 \times 10^{-6}$ and $\epsilon_{yy} = \epsilon_{xy} = 3 \times 10^{-6}$, while the other terms remain unchanged. Units of cc-grad mean 1.0×10^{-4} grad, where $200 \text{ grad} = \pi \text{ radians}$. The convergence criterion was set to $\delta = 1.0 \times 10^{-10}$. Model check refers to the norm of (6d).

Quantity	Case 1	Case 2	Case 3	Case 4
$\tilde{\epsilon}_{xx}$ [μ -strain]	-4.541	-4.541	-5.500	-6.122
$\tilde{\epsilon}_{yy}$ [μ -strain]	4.845	4.845	4.500	3.385
$\tilde{\epsilon}_{xy}$ [μ -strain]	4.036	4.036	4.500	3.794
$\tilde{\omega}$ [cc-grad]	-144.740600	-144.740599	-144.730133	-144.767811
$\tilde{t}_x$ [m]	0.118103	0.118097	0.077052	0.062656
$\tilde{t}_y$ [m]	-0.015506	-0.015503	0.006922	-0.006881
TSSR (Ω)	2.058678	2.060113	2.15×10^6	5.779827
Model check	1.23×10^{-19}	1.23×10^{-19}	1.80×10^{-11}	3.72×10^{-19}
Iterations	2	3	3	3

References

- Brunner, F., Coleman, R., and Hirsch, B. (1981). A comparison of computation methods for crustal strains from geodetic measurements. *Tectonophysics*, 71:281–298.
- Caspary, W. (2000). *Concepts of Network and Deformation Analysis*. Monograph 11. School of Geomatic Engineering, University of New South Wales, Australia.
- Fang, X. (2013). Weighted total least squares: Necessary and sufficient conditions, fixed and random parameters. *J. of Geodesy*, 87:733–749.
- Jazaeri, S., Schaffrin, B., and Snow, K. (2014). On Weighted Total Least-Squares adjustment with multiple constraints and singular dispersion matrices. *Z. für Vermessungswesen*, 139(4):229–240.
- Moritz, H. (1970). A generalized least-squares model. *Studia Geophysica et Geodaetica*, 14(2):353–362.
- Nkuite, G. (1998). *Adjustment with Singular Variance-Covariance Matrix Using an Example of Geometric Deformation Analysis* (in German). Reihe C, Heft Nr. 501. Deutsche Geodätische Kommission, Munich.
- Schaffrin, B. (2001). Softly unbiased prediction. Part 2: The Random Effects Model. *Boll. Geod. Sci. Aff.*, 60(1):49–62.
- Schaffrin, B. (2009). Total Least-Squares Collocation: The Total Least-Squares approach to EIV-Models with prior information. Presented at the Intl. Workshop on Matrices and Statistics, Smolenice Castle/Slovakia (available upon request from the first author).
- Schaffrin, B. (2020). Total Least-Squares Collocation: An Optimal Estimation Technique for the EIV-Model with Prior Information. *Mathematics*, 8(6):971.
- Schaffrin, B. and Snow, K. (2013). Total Least-Squares Adjustment with Prior Information vs. the Penalized Least-Squares Approach to EIV-Models. In *Proceedings of the 59th ISI World Statistics Congress*, pages 2000–2005, Hong Kong.
- Snow, K. (2012). *Topics in Total Least-Squares Adjustment within the Errors-In-Variables Model: Singular Cofactor Matrices and Prior Information*. PhD thesis, Report No. 502, Div. of Geodetic Science, School of Earth Sciences, The Ohio State University, Columbus, Ohio, USA.
- Snow, K. and Schaffrin, B. (2012). Weighted Total Least-Squares Collocation with Geodetic Applications. Presented at the 2012 SIAM Conference on Applied Linear Algebra, Valencia, Spain (available upon request from the first author).
- Snow, K. and Schaffrin, B. (2025). Total Least-Squares collocation: The general case. *J. of Geodesy*. Submitted in March 2025.

Table 5. The upper-triangular portion of the covariance matrix V_1 for the coordinates at epoch 1: row and column numbers followed by the associated numerical value

1	1	5.2857049251185705e-04	1	2	3.0112072439900002e-07	1	3	2.2163970015471102e-04
1	4	-3.3060012167780104e-04	1	5	-2.2446581495537803e-04	1	6	-2.5343445696036004e-04
1	7	-3.0974036938091802e-04	1	8	-7.4422501336440000e-05	1	9	1.2316417979071101e-04
1	10	2.9849301683632302e-04	1	11	6.1944192322230007e-06	1	12	-1.0276614593379901e-04
1	13	-1.8577329964890101e-04	1	14	3.6244201617832402e-04	1	15	-1.5958929463899500e-04
1	16	9.9987021169947011e-05	2	2	5.1645507048873602e-04	2	3	4.0728507326387003e-05
2	4	-1.0000823727611901e-04	2	5	-1.8278284645652202e-04	2	6	-1.6265542801598300e-04
2	7	-1.0458569283936001e-04	2	8	-8.9543289024250001e-05	2	9	5.9950930752997007e-05
2	10	-3.5258061736731003e-05	2	11	-6.3711277472731004e-05	2	12	4.3403859615880001e-05
2	13	1.8024977671021402e-04	2	14	-1.2951456512824702e-04	2	15	6.9850434616229999e-05
2	16	-4.2879336832345003e-05	3	3	6.1101176008653707e-04	3	4	-2.0315427619786602e-04
3	5	-2.6683137867546802e-04	3	6	-3.0461425870949700e-04	3	7	-4.0530634456603404e-04
3	8	8.7952477461938002e-05	3	9	2.7767431373654001e-05	3	10	1.2061066507825600e-04
3	11	-4.7659830667130001e-06	3	12	-2.0971626484055001e-05	3	13	-5.4391024593689005e-05
3	14	1.6558385890034301e-04	3	15	-1.2912411448790800e-04	3	16	1.1386424329115401e-04
4	4	8.2777222054439801e-04	4	5	6.5995597460920006e-05	4	6	2.6599285903412302e-04
4	7	1.2180558107243101e-04	4	8	-1.7127450086786001e-05	4	9	-4.9251029387330006e-05
4	10	-4.2253794234299304e-04	4	11	1.0022267062051001e-05	4	12	1.2390531393505200e-04
4	13	2.1887205868590702e-04	4	14	-5.0339456913838106e-04	4	15	1.6631092866995802e-04
4	16	-1.7460073595352302e-04	5	5	6.2760185074075300e-04	5	6	2.0791753467522600e-04
5	7	1.5485520136481000e-05	5	8	2.4436563047611101e-04	5	9	-7.8315474715017009e-05
5	10	-1.8418744339853901e-04	5	11	-2.5854139004645000e-05	5	12	-4.2386999902525006e-05
5	13	1.6543949322898002e-05	5	14	-1.3052198865143701e-04	5	15	-6.4166116435632005e-05
5	16	2.1602313059309000e-05	6	6	9.2599529493215009e-04	6	7	5.8146483162555007e-04
6	8	1.2268553799023402e-04	6	9	-1.8062344580547401e-04	6	10	-3.7805609355220303e-04
6	11	-3.0203512679857001e-05	6	12	4.6556596221795004e-05	6	13	-1.2232674215847600e-04
6	14	-5.1313452391676202e-04	6	15	1.0181780646886600e-04	6	16	-3.0737798988240002e-04
7	7	1.1058615997014982e-03	7	8	2.6221711403204902e-04	7	9	-2.4311905812814501e-04
7	10	-2.4347880305738902e-04	7	11	-8.4710620000211009e-05	7	12	4.7395236099771003e-05
7	13	-1.9898662308994402e-04	7	14	-3.8043259991024501e-04	7	15	1.2051571080840701e-04
7	16	-2.8438456224831702e-04	8	8	7.4494925733171308e-04	8	9	-2.2056116050899301e-04
8	10	-2.1338643681943602e-04	8	11	6.8620856530167000e-05	8	12	-3.9234683498115003e-05
8	13	-3.4250623741234203e-04	8	14	-3.3626438587133203e-04	8	15	-2.5668503190103001e-05
8	16	-1.7208265335987702e-04	9	9	3.9583795016671903e-04	9	10	2.3759992909340001e-05
9	11	-1.7705163148991002e-05	9	12	-4.6034011328802005e-05	9	13	-5.6876992206821006e-05
9	14	3.7598217036760101e-04	9	15	-1.5075206205812601e-04	9	16	3.6777162028395004e-05
10	10	6.6836796286931803e-04	10	11	-2.6055431838001000e-05	10	12	-1.4499412607699402e-04
10	13	1.0601286059567501e-04	10	14	4.1404458048523201e-04	10	15	-9.5150345192206005e-05
10	16	1.1181550558688800e-04	11	11	2.4007402828181703e-04	11	12	-2.7330509375470003e-06
11	13	-1.2298404847052301e-04	11	14	-4.1665119094890003e-06	11	15	9.7506354775860004e-06
11	16	4.8227331877930002e-05	12	12	2.6319739300808003e-04	12	13	5.2668686072597006e-05
12	14	-2.3914824677896801e-04	12	15	1.1482641320053501e-04	12	16	-5.3685133414396005e-05
13	13	6.5192974789769309e-04	13	14	-7.1800128988254008e-05	13	15	-4.9460374059992003e-05
13	16	-2.1175344006979000e-05	14	14	1.1194914159530060e-03	14	15	-3.1708698737387001e-04
14	16	1.8791442347010102e-04	15	15	4.2282565878049401e-04	15	16	-1.4898164708493001e-05
16	16	4.5089593497817202e-04						