

10.4. PARAMETRE ÇÖZÜNÜRLÜĞÜ

10.4.1. Veri Ayrımlılık Dizeyi

Veri çekirdek dizeyinin sütunları, veri noktalarının parametrelerden etkileniş biçimini gösterir. Örneğin, j numaralı sütunun bireyleri göreceli olarak yeterli sayıda yüksek değerler kapsamakta ise j numaralı parametre duyarlı şekilde çözülebilir. Eğer, belirli bir sütun veya sütunların sayısal değerleri küçük veya büyük değerli birey sayısı yetersiz ise, veri grubu bu sütunlara karşılık gelen parametreyi çözmek için kullanılamaz ve çözüm için ek bilgiye gerek vardır.

Aşırı tanımlı problemlerin çözümü,

$$\mathbf{G}_L^{-1} = (\mathbf{G}^T \mathbf{G})^{-1} \mathbf{G}^T \quad (10.4.1)$$

gösterimi ile

$$\mathbf{p} = \mathbf{G}_L^{-1} \mathbf{d} \quad (10.4.2)$$

şeklinde yazılabilir. Burada, $\mathbf{G}_L^{-1}$ dizeyi Lanczos tersini göstermektedir. Hesaplanan parametreler kullanılarak, verinin kestirilen değerleri (kuramsal veri) hesaplanabilir. İzleyen

$$\mathbf{G} \mathbf{G}_L^{-1} = \mathbf{I} \quad (10.4.3)$$

koşulunu sağlayan, $\mathbf{G}$ veri çekirdek dizeyi yardımı ile

$$\mathbf{d}^{\text{kestirilen}} = \mathbf{G} \mathbf{p}^{\text{hesap}} \quad (10.4.4)$$

yazılabilir. Burada, $\mathbf{p}^{\text{hesap}}$ (10.4.2) denkleminde hesaplandığından, (10.4.4) bağıntısında yerine konulması ile

$$\mathbf{d}^{\text{kestirilen}} = \mathbf{G} \mathbf{G}_L^{-1} \mathbf{d}^{\text{ölçülen}} \quad (10.4.5)$$

elde edilir. Bu denklem, $\mathbf{G} \mathbf{G}_L^{-1}$ çarpımının birim dizeye eşit olması halinde ölçülen ve kestirilen verilerin birbirine eşit olacağını göstermektedir. Bu çarpım,

$$\mathbf{N}_{n \times n} = \mathbf{G}_{n \times m} (\mathbf{G}_L^{-1})_{m \times n} \quad (10.4.6)$$

dizeyi ile gösterilir ve veri ayrımlılık dizeyi (data resolution matrix) adını alır. Boyutu $(n \times n)$ dir. $\mathbf{N}$ dizeyi birim dizeyi ise, yani köşegen bireyleri birim değere yakınsa, ayrımlılık iyidir ve veri parametreleri çözmek için yeterli bilgiyi kapsamaktadır. Köşegen bireyler birim "eğerden uzak ve köşegene yakın bireyler sıfırdan farklı değerler almakta ise, verinin parametrelerin çözümünde yeterli ayrımlılığı sağlayamadığı söylenebilir.

10.4.2. Parametre Ayrımlılık Dizeyi

Ters çözüm, ölçülen veriden parametere hesaplanması olarak tanımlanmıştır:

$$\mathbf{p}^{\text{hesap}} = \mathbf{G}_L^{-1} \mathbf{d}^{\text{ölçülen}} \quad (10.4.10)$$

Ters çözüm tekniğinin genel kuralı gereği, ölçülen veri aynı zamanda gerçek parametrelerden hesaplanabilmelidir:

$$\mathbf{d}^{\text{ölçülen}} = \mathbf{G} \mathbf{p}^{\text{gerçek}} \quad (10.4.9)$$

Bu denklem, (10.3.10) denkleminde yerine konur ise

$$\mathbf{p}^{\text{hesap}} = \mathbf{G}_L^{-1} \mathbf{G} \mathbf{p}^{\text{gerçek}} \quad (10.4.10)$$

denklemi elde edilir. İzleyen,

$$\mathbf{R}_{mxm} = (\mathbf{G}_L^{-1})_{mxn} \mathbf{G}_{n \times m} \quad (10.4.11)$$

dizeyi ($m \times m$) boyutundadır ve parametre ayrımlılık dizeyi (parameter resolution matrix) olarak adlandırılır. $\mathbf{R}$ dizeyinin R_{ij} bireyi, i nci parametrenin hesaplanmasına, j inci parametre tarafından yapılan katkıyı temsil eder. Köşegen bireylerde $i = j$ ve diğer bireylerde i ve j birbirinden farklıdır. Köşegen dışı bireylerin sıfıra yakınlığı, parametrelerin birbirinden bağımsız çözülebileceğinin bir göstergesidir. Bu nedenle, $\mathbf{R}$ birim dizeye yakınsa ters çözüm işleminden elde edilen parametrelerin modeli temsil ettiği kabul edilebilir ve bütün parametreler tam olarak çözülür. Bu dizeyin bireylerinin birim dizeye yakınlığı, parametrelerin gerçek değerlerine yakınlığının bir ölçüsüdür.

Geophysical data analysis: Discrete inverse theory, William Menke, Academic Press

4.2 The Data Resolution Matrix

Suppose we have found a generalized inverse that in some sense solves the inverse problem $\mathbf{Gm} = \mathbf{d}$, yielding an estimate of the model parameters $\mathbf{m}^{\text{est}} = \mathbf{G}^{-*} \mathbf{d}$ (for the sake of simplicity we assume that there is no additive vector). We can then retrospectively ask how well this estimate of the model parameters fits the data. By plugging our estimate into the equation $\mathbf{Gm} = \mathbf{d}$ we conclude

$$\mathbf{d}^{\text{pre}} = \mathbf{Gm}^{\text{est}} = \mathbf{G}[\mathbf{G}^{-*} \mathbf{d}^{\text{obs}}] = [\mathbf{G}\mathbf{G}^{-*}] \mathbf{d}^{\text{obs}} = \mathbf{N} \mathbf{d}^{\text{obs}} \quad (4.1)$$

Here the superscripts obs and pre mean observed and predicted, respectively. The $N \times N$ square matrix $\mathbf{N} = \mathbf{G}\mathbf{G}^{-*}$ is called the *data resolution matrix*. This matrix describes how well the predictions match the data. If $\mathbf{N} = \mathbf{I}$, then $\mathbf{d}^{\text{pre}} = \mathbf{d}^{\text{obs}}$ and the prediction error is zero. On the other hand, if the data resolution matrix is not an identity matrix, the prediction error is nonzero.

If the elements of the data vector $\mathbf{d}$ possess a natural ordering, then the data resolution matrix has a simple interpretation. Consider, for example, the problem of fitting a straight line to (z, d) points, where the data have been ordered according to the value of the auxiliary variable z . If $\mathbf{N}$ is not an identity matrix but is close to an identity matrix (in the sense that its largest elements are near its main diagonal), then the configuration of the matrix signifies that averages of neighboring data can be predicted, whereas individual data cannot. Consider the i th row of $\mathbf{N}$. If this row contained all zeros except for a

one in the i th column, then d_i would be predicted exactly. On the other hand, suppose that the row contained the elements

$$[\cdots 0 \ 0 \ 0 \ 0.1 \ 0.8 \ 0.1 \ 0 \ 0 \ 0 \cdots] \quad (4.2)$$

where the 0.8 is in the i th column. Then the i th datum is given by

$$d_i^{\text{pre}} = \sum_{j=1}^N N_{ij} d_j^{\text{obs}} = 0.1 d_{i-1}^{\text{obs}} + 0.8 d_i^{\text{obs}} + 0.1 d_{i+1}^{\text{obs}} \quad (4.3)$$

The predicted value is a weighted average of three neighboring observed data. If the true data vary slowly with the auxiliary variable, then such an average might produce an estimate reasonably close to the observed value.

The rows of the data resolution matrix N describe how well neighboring data can be independently predicted, or *resolved*. If the data have a natural ordering, then a graph of the elements of the rows of N against column indices illuminates the sharpness of the resolution (Fig. 4.1). If the graphs have a single sharp maximum centered about the main diagonal, then the data are well resolved. If the graphs are very broad, then the data are poorly resolved. Even in cases where there is no natural ordering of the data, the resolution matrix still shows how much weight each observation has in influencing the

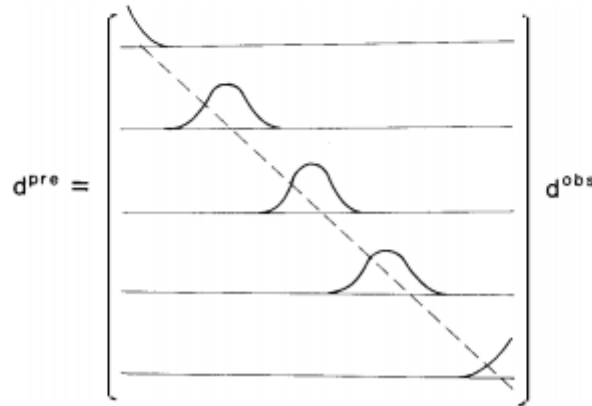


Fig. 4.1. Plots of selected rows of the data resolution matrix N indicate how well the data can be predicted. Narrow peaks occurring near the main diagonal of the matrix (dashed line) indicate that the resolution is good.

predicted value. There is then no special significance to whether large off-diagonal elements fall near to or far from the main diagonal.

Because the diagonal elements of the data resolution matrix indicate how much weight a datum has in its own prediction, these diagonal elements are often singled out and called the *importance* n of the data [Ref. 15]:

$$n = \text{diag}(N) \quad (4.4)$$

The data resolution matrix is not a function of the data but only of the data kernel G (which embodies the model and experimental geometry) and any a priori information applied to the problem. It can therefore be computed and studied without actually performing the experiment and can be a useful tool in experimental design.

4.3 The Model Resolution Matrix

The data resolution matrix characterizes whether the data can be independently predicted, or resolved. The same question can be asked about the model parameters. To explore this question we imagine that there is a true, but unknown set of model parameters $\mathbf{m}^{\text{true}}$ that solve $G\mathbf{m}^{\text{true}} = \mathbf{d}^{\text{obs}}$. We then inquire how closely a particular estimate of the model parameters $\mathbf{m}^{\text{est}}$ is to this true solution. Plugging the expression for the observed data $G\mathbf{m}^{\text{true}} = \mathbf{d}^{\text{obs}}$ into the expression for the estimated model $\mathbf{m}^{\text{est}} = G^{-1}\mathbf{d}^{\text{obs}}$ gives

$$\mathbf{m}^{\text{est}} = G^{-1}\mathbf{d}^{\text{obs}} = G^{-1}[G\mathbf{m}^{\text{true}}] = [G^{-1}G]\mathbf{m}^{\text{true}} = \mathbf{R}\mathbf{m}^{\text{true}} \quad (4.5)$$

[Ref. 20] Here $\mathbf{R}$ is the $M \times M$ *model resolution matrix*. If $\mathbf{R} = \mathbf{I}$, then each model parameter is uniquely determined. If $\mathbf{R}$ is not an identity matrix, then the estimates of the model parameters are really weighted averages of the true model parameters. If the model parameters have a natural ordering (as they would if they represented a discretized version of a continuous function), then plots of the rows of the resolution matrix can be useful in determining to what scale features in the model can actually be resolved (Fig. 4.2). Like the data resolution matrix, the model resolution is a function of only the data kernel and the a priori information added to the problem. It is therefore independent of the actual values of the data and can therefore be another important tool in experimental design.

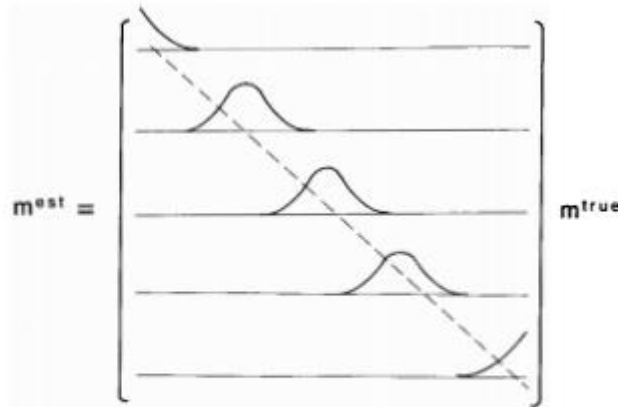


Fig. 4.2. Plots of selected rows of the model resolution matrix R indicate how well the true model parameters can be resolved. Narrow peaks occurring near the main diagonal of the matrix (dashed line) indicate that the model is well resolved.