



Delft University of Technology

Adjustment theory

An introduction

Teunissen, P.J.G.

DOI

[10.59490/tb.95](https://doi.org/10.59490/tb.95)

Publication date

2024

Document Version

Final published version

Citation (APA)

Teunissen, P. J. G. (2024). *Adjustment theory: An introduction*. TU Delft OPEN Publishing.
<https://doi.org/10.59490/tb.95>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

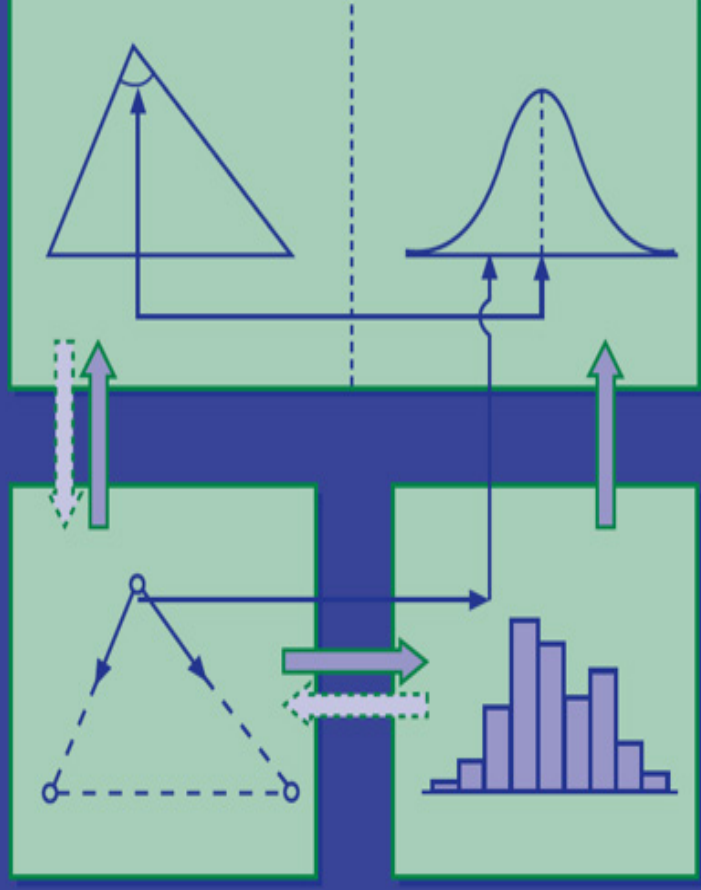
Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Adjustment theory: an introduction

Peter J.G. Teunissen



Adjustment theory

an introduction

Adjustment theory

an introduction

P.J.G.Teunissen



Delft University of Technology
Faculty of Civil Engineering and Geosciences
Department of Mathematical Geodesy and Positioning

Series on Mathematical Geodesy and Positioning

© VSSD, first edition 2000-2006

© TU Delft OPEN Publishing, second edition 2024

ISBN: 978-94-6366-884-2 (Ebook)

ISBN: 978-94-6366-885-9 (Paperback/softback)

Doi: <https://doi.org/10.59490/tb.95>



This work is licensed under a [Creative Commons Attribution 4.0 International license](https://creativecommons.org/licenses/by/4.0/)



Keywords: adjustment theory, geodesy

Preface 2nd Edition

To promote open access, this new edition of Adjustment Theory is published by TU Delft Open Publishing instead of Delft Academic Press.

Appendix G has been added to this edition. It places the origin of Adjustment Theory in its historical context and provides a brief account of the early methods of adjusting geodetic and astronomical observations.

May, 2024

Peter J.G. Teunissen

Preface

As in other physical sciences, empirical data are used in Geodesy and Earth Observation to make inferences so as to describe the physical reality. Many such problems involve the determination of a number of unknown parameters which bear a linear (or linearized) relationship to the set of data. In order to be able to check for errors or to reduce for the effect these errors have on the final result, the collected data often exceed the minimum necessary for a unique determination (redundant data). As a consequence of (measurement) uncertainty, redundant data are usually inconsistent in the sense that each sufficient subset yields different results from another subset. To obtain a unique solution, consistency needs to be restored by applying corrections to the data. This computational process of making the data consistent such that the unknown parameters can be determined uniquely, is referred to as *adjustment*. This book presents the material for an introductory course on Adjustment theory. The main goal of this course is to convey the knowledge necessary to perform an optimal adjustment of redundant data. Knowledge of linear algebra, statistics and calculus at the undergraduate level is required. Some prerequisites, repeatedly needed in the book, are summarized in the Appendices A-F.

Following the *Introduction*, the elementary adjustment principles of 'least-squares', unbiased minimum variance' and 'maximum likelihood' are discussed and compared in *Chapter 1*. These results are generalized to the multivariate case in chapters 2 and 3. In *Chapter 2* the model of observation equations is presented, while the model of condition equations is discussed in *Chapter 3*. Observation equations and condition equations are dual to each other in the sense that the first gives a parametric representation of the model, while the second gives an implicit representation. Observables that are either stochastically or functionally related to another set of observables are referred to as y^r -variates. Their adjustment is discussed in *Chapter 4*. In *Chapter 5* we consider mixed model representations and in *Chapter 6* the partitioned models. Models may be partitioned in different ways: partitioning of the observation equations, partitioning of the unknown parameters, partitioning of the condition equations, or all of the above. These different partitionings are discussed and their relevance is shown for various applications. Up to this point all the material is presented on the basis of the assumption that the data are linearly related. In geodetic applications however, there are only a few cases where this assumption holds true. In the last chapter, *Chapter 7*, we therefore start extending the theory from linear to nonlinear situations. In particular an introduction to the linearization and iteration of nonlinear models is given.

Many colleagues of the Department of Mathematical Geodesy and Positioning whose assistance made the completion of this book possible are greatly acknowledged. C.C.J.M. Tiberius took care of the editing, while the typing was done by Mrs. J. van der Bijl and Mrs. M.P.M. Scholtes, and the drawings by Mr. A. Smit. Various lecturers have taught the book's material over the past years. In particular the feedback and valuable recommendations of the lecturers C.C.J.M. Tiberius, F. Kenselaar, G.J. Husti and H.M. de Heus are acknowledged.

P.J.G. Teunissen
April, 2003

Contents

Introduction	1
1 Linear estimation theory: an introduction	5
1.1 Introduction	5
1.2 Least-squares estimation (orthogonal projection)	5
1.3 Weighted least-squares estimation	11
1.4 Weighted least-squares (also orthogonal projection)	20
1.5 The mean and covariance matrix of least-squares estimators	25
1.6 Best Linear Unbiased Estimators (BLUE's)	28
1.7 The Maximum Likelihood (ML) method	32
1.8 Summary	35
2 The model with observation equations	39
2.1 Introduction	39
2.2 Least-squares estimation	42
2.3 The mean and covariance matrix of least-squares estimators	45
2.4 Best Linear Unbiased Estimation	49
2.5 The orthogonal projector $A(A^*Q_y^{-1}A)^{-1}A^*Q_y^{-1}$	55
2.6 Summary	59
3 The model with condition equations	61
3.1 Introduction	61
3.2 Best Linear Unbiased Estimation	64
3.3 The case $B^*E\{\underline{y}\} = b^0, b^0 \neq 0$	67
3.4 Summary	68
4 $\underline{y}^R$ -Variates	71
4.1 Introduction	71
4.2 Free variates	75
4.3 Derived variates	77
4.4 Constituent variates	77
5 Mixed model representations	81
5.1 Introduction	81
5.2 The model representation $B^*E\{\underline{y}\}=Ax$	82
5.3 The model representation $E\{\underline{y}\}=Ax, B^*x=0$	84

6	Partitioned model representations	89
6.1	Introduction	89
6.2	The partitioned model $E\{\underline{y}\} = (A_1 x_1 + A_2 x_2)$	90
6.3	The partitioned model $E\{(\underline{y}_1^*, \underline{y}_2^*)^*\} = (A_1^* : A_2^*)^* x$ (recursive estimation)	100
6.4	The partitioned model $E\{(\underline{y}_1^*, \underline{y}_2^*)^*\} = (A_1^* : A_2^*)^* x$ (block estimation I)	109
6.5	The partitioned model $E\{(\underline{y}_1^*, \underline{y}_2^*)^*\} = (A_1^* : A_2^*)^* x$ (block estimation II)	113
6.6	The partitioned model $(B_1 : B_2)^* E\{\underline{y}\} = (0 \ 0)^*$ (estimation in phases)	123
7	Nonlinear models, linearization, iteration	137
7.1	The nonlinear A-model	137
7.1.1	Nonlinear observation equations	137
7.1.2	Linearization: Taylor's theorem	142
7.1.3	The linearized A-model	146
7.1.4	Least-squares iteration	150
7.2	The nonlinear B-model and its linearization	154
	Appendices	159
A	Taylor's theorem with remainder	159
B	Unconstrained optimization: optimality conditions	164
C	Linearly constrained optimization: optimality conditions	168
D	Constrained optimization: optimality conditions	176
E	Mean and variance of scalar random variables	178
F	Mean and variance of vector random variables	184
G	A brief account on the early history of adjusting geodetic and astronomical observations	192
	Literature	198
	Index	199

Introduction

The *calculus of observations* (in Dutch: *waarnemingsrekening*) can be regarded as the part of mathematical statistics that deals with the results of measurement. Mathematical statistics is often characterized by saying that it deals with the making of decisions in uncertain situations. In first instance this does not seem to relate to measurements. After all, we perform measurements in order to obtain sharp, objective and unambiguous quantitative information, information that is both correct and precise, in other words accurate information. Ideally we would like to have 'mathematical certainty'. But mathematical certainty pertains to statements like:

if $a=b$ and $b=c$ and $c=d$, then $a=d$.

If a , b , c , and d are measured terrain distances however, a surveyor to be sure will first compare d with a before stating that $a=d$. Mathematics is a product of our mind, whereas with measurements one is dealing with humans, instruments, and other materialistic things and circumstances, which not always follow the laws of logic. From experience we know that measurements produce results that are only sharp to a certain level. There is a certain vagueness attached to them and sometimes even, the results are simply wrong and not usable.

In order to eliminate uncertainty, it seems obvious that one seeks to improve the instruments, to educate people better and to try to control the circumstances better. This has been successfully pursued throughout the ages, which becomes clear when one considers the very successful development of geodetic instruments since 1600. However, there are still various reasons why measurements will always remain uncertain to some level:

1. 'Mathematical certainty' relates to abstractions and not to physical, technical or social matters. A mathematical model can often give a useful description of things and relations known from our world of experience, but still there will always remain an essential difference between a model and the real world.
2. The circumstances of measurement can be controlled to some level in the laboratory, but in nature they are uncontrollable and only partly descriptive.
3. When using improved methods and instrumentation, one still, because of the things mentioned under 1 and 2, will have uncertainty in the results of measurement, be it at another level.
4. Measurements are done with a purpose. Dependent on that purpose, some uncertainty is quite acceptable. Since better instruments and methods usually cost more, one will have to ask oneself the question whether the reduction in uncertainty is worth the extra costs.

Summarizing, it is fundamentally and practically impossible to obtain absolute certainty from measurements, while at the same time this is also often not needed on practical grounds. Measurements are executed to obtain quantitative information, information which is accurate enough for a specific purpose and at the same time cheap enough for that purpose. It is because of this intrinsic uncertainty in measurements, that we have to consider the calculus of observations. The calculus of observations deals with:

2 Adjustment theory

1. The description and analysis of measurement processes.
2. Computational methods that take care, in some sense, of the uncertainty in measurements.
3. The description of the quality of measurement and of the results derived therefrom.
4. Guidelines for the design of measurement set-ups, so as to reach optimal working procedures.

The calculus of observations is essential for a geodesist, its meaning comparable to what mechanics means for the civil-engineer or mechanical-engineer. In order to shed some more light on the type of problems the calculus of observations deals with, we take some surveying examples as illustration. Some experience in surveying will for instance already lead to questions such as:

1. Repeating a measurement does usually not give the same answer. How to describe this phenomenon?
2. Results of measurement can be corrupted by systematic or gross errors. Can these errors be traced and if so, how?
3. Geometric figures (e.g. a triangle) are, because of 2, usually measured with *redundant* (*overtallig*) observations (e.g. instead of two, all three angles of a triangle are measured). These redundant observations will usually not obey the mathematical rules that hold true for these geometric figures (e.g. the three angular measurements will usually not sum up to π). An *adjustment* (*vereffening*) is therefore needed to obtain consistent results again. What is the best way to perform such an adjustment?
4. How to describe the effect of 1 on the final result? Is it possible that the final result is still corrupted with non-detected gross errors?

Once we are able to answer questions like these, we will be able to determine the required measurement set-up from the given desired quality of the end product. In this book we mainly will deal with points 1 and 3, the elementary aspects of adjustment theory.

Functional and stochastic model

In order to analyze scientific or practical problems and/or make them accessible for a quantitative treatment, one usually tries to rid the problem from its unnecessary frills and describes its essential aspects by means of notions from an applicable mathematical theory. This is referred to as the making of a *mathematical model*. On the one hand the model should give a proper description of the situation, while, on the other hand, the model should not be needlessly complicated and difficult to apply. The determination of a proper model can be considered the 'art' of the discipline. Well-tried and well-tested experience, new experience from experiments, intuition and creativity, they all play a role in formulating the model.

Here we will restrict ourselves to an example from surveying. Assume that we want to determine the relative position of terrain-'points', whereby their differences in height are of no interest to us. If the distances between the 'points' are not too large, we are allowed, as experience tells us, to describe the situation by projecting the 'points' onto a level surface (*niveauvlak*) and by treating this surface as if it was a plane. For the description of the relative position of the points, we are then allowed to make use of the two-dimensional Euclidean geometry (*vlakke*

meetkunde), the theory of which goes back some two thousand years. This theory came into being as an abstraction from how our world was seen and experienced. But for us the theory is readily available, thus making it possible to anticipate on its use, for instance by using sharp and clear markings for the terrain-'points' such that they can be treated as if they were mathematically defined points. In this example, the two-dimensional Euclidean geometry acts as our *functional model*. That is, notions and rules from this theory (such as the sine- and cosine rule) are used to describe the relative position of the terrain-'points'.

Although the two-dimensional Euclidean geometry serves its purpose well for many surveying tasks, it is pointed out that the gratuitously use of a 'standard' model can be full of risks. Just like the laws of nature (natuurwetten), a model can only be declared 'valid' on the basis of experiments. We will come back to this important issue in some of our other courses.

When one measures some of the geometric elements of the above mentioned set of points, it means that one assigns a number (the observation) according to certain rules (the measurement procedure) to for instance an angle α (a mathematical notion). From experience we know that a measurement of a certain quantity does not necessarily give the same answer when repeated. In order to describe this phenomenon, one could repeat the measurements a great many times (e.g. 100). And on its turn the experiment itself could also be repeated. If it then turns out that the histograms of each experiment are sufficiently alike, that is, their shapes are sufficiently similar and their locations do not differ too much, one can describe the measurement variability by means of a stochastic variable. In this way notions from mathematical statistics are used in order to describe the inherent variability of measurement processes. This is referred to as the *stochastic model*.

In practice not every measurement will be repeated of course. In that case one will have to rely on the extrapolation of past experience. When measuring, not only the measurement result as a number is of interest, but also additional information such as the type of instrument used, the measurement procedure followed, the observer, the weather, etc. Taken together this information is referred to as the *registration (registratie)*. The purpose of the registration is to establish a link with past experience or experiments such that a justifiable choice for the stochastic model can be made.

In surveying and geodesy one often may assume as stochastic model that the results of measurement (the observations) of a certain measurement method can be described as being an independent sample drawn (aselecte trekking) from a normal (Gaussian) distribution with a certain standard deviation. The mathematical *expectation (verwachting)* of the normally distributed stochastic variable (the observable, in Dutch: de waarnemingsgrootheid), is then assumed to equal the unknown angle etc. From this it follows, if the three angles α , β and γ of a triangle are measured, that the expectations of the three angular stochastic variables obey the following relation of the functional model:

$$\alpha + \beta + \gamma = \pi$$

For the individual observations themselves this is usually not the case. A sketch of the relation between terrain situation, functional model, measurement, experiment and stochastic model is given in figure 0.1. The functional and stochastic model together form the mathematical model.

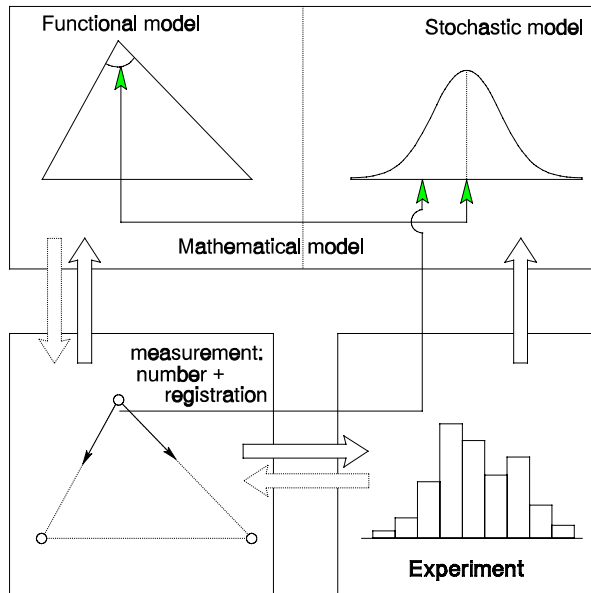


Figure 0.1: Diagram of the fundamental relations in adjustment theory

1 Linear estimation theory: an introduction

1.1 Introduction

Before stating the general problem, let us define a simple example that contains most of the elements that will require our consideration.

Suppose that an unknown distance, which we denote as x , has been measured twice. The two measurements are denoted as y_1 and y_2 , respectively. Due to all sorts of influences (e.g. instrumental effects, human effects, etc.) the two measurements will differ in general. That is $y_1 \neq y_2$ or $y_1 - y_2 \neq 0$. One of the two measurements, but probably both, will therefore also differ from the true but unknown distance x . Thus, $y_1 \neq x$ and $y_2 \neq x$. Let us denote the two differences $y_1 - x$ and $y_2 - x$, the so-called measurement errors, as e_1 and e_2 respectively. We can then write:

$$y_i = x + e_i, \quad i = 1, 2$$

or in vector notation:

(1)

$$\begin{array}{ccccc} y & = & a & x & + & e \\ 2 \times 1 & & 2 \times 1 & 1 \times 1 & & 2 \times 1 \end{array}$$

with:

$$y = (y_1, y_2)^*$$

$$a = (1, 1)^*$$

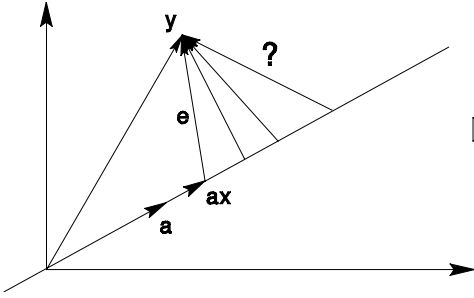
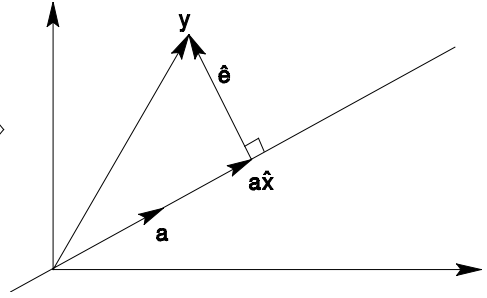
$$e = (e_1, e_2)^*$$

(* denotes the operation of transposition; thus y^* is the transpose of y).

Our problem is now, how to estimate the unknown x from the given measurement or observation vector y . It is to this and related questions that we address our attention. In this introductory chapter we will restrict ourselves to the case of two observations and one unknown. In the next section we will introduce the least-squares principle. Least-squares can be regarded as a deterministic approach to the estimation problem. In the sections following we will introduce additional assumptions regarding the probabilistic nature of the observations and measurement errors. This leads then to the so-called best linear unbiased estimators and the maximum likelihood estimators.

1.2 Least-squares estimation (orthogonal projection)

The equation $y = a x + e$, in which y and a are given, and x and e are unknown, can be visualized as in figure 1.1a.

Figure 1.1a: x and e unknownFigure 1.1b: $\hat{x}$ as estimate of x
 $\hat{e}$ as estimate of e

From the geometry of the figure it seems intuitively appealing to estimate x as $\hat{x}$, such that $a\hat{x}$ is as close as possible to the given observation vector y . From figure 1.1b follows that $a\hat{x}$ has minimum distance to y if $\hat{e} = y - a\hat{x}$ is orthogonal to a . That is if:

$a^*(y - a\hat{x}) = 0$	"orthogonality"
$\Downarrow$	
$(a^*a)\hat{x} = a^*y$	"normal equation"
$\Downarrow$	
$\hat{x} = (a^*a)^{-1}a^*y$	"LSQ-estimate of x "
$\Downarrow$	
(2) $\hat{y} = a\hat{x} = [a(a^*a)^{-1}a^*]y$	"LSQ-estimate of observations"
$\Downarrow$	
$\hat{e} = y - \hat{y} = [I - a(a^*a)^{-1}a^*]y$	"LSQ-estimate of measurement errors"
$\Downarrow$	
$\hat{e}^*\hat{e} = y^*[I - a(a^*a)^{-1}a^*]y$	"Sum of squares"

The name "*Least-Squares*" (*LSQ*) is given to the estimate $\hat{x} = (a^*a)^{-1}a^*y$ of x , since $\hat{x}$ minimizes the sum of squares $(y - ax)^*(y - ax)$. We will give two different proofs for this:

First proof: (algebra)

$$\begin{aligned}
 (y - ax)^*(y - ax) &= [y - a\hat{x} - a(x - \hat{x})]^*[y - a\hat{x} - a(x - \hat{x})] \\
 &= [\hat{e} - a(x - \hat{x})]^*[\hat{e} - a(x - \hat{x})] \\
 &= \hat{e}^*\hat{e} - \hat{e}^*a(x - \hat{x}) - (x - \hat{x})^*a^*\hat{e} + (x - \hat{x})^*a^*a(x - \hat{x}) \\
 &= \hat{e}^*\hat{e} + (x - \hat{x})^*a^*a(x - \hat{x}) \\
 &= \text{minimal for } x = \hat{x}.
 \end{aligned}$$

End of proof.

Second proof: (calculus)

$$\text{Call } E(x) \triangleq (y - ax)^*(y - ax) = y^*y - 2y^*ax + a^*ax^2.$$

From calculus we know that $\hat{x}$ is a solution of the minimization problem $\min_x E(x)$ if:

$$\frac{dE}{dx}(\hat{x}) = 0 \text{ and } \frac{d^2E}{dx^2}(\hat{x}) > 0.$$

With:

$$\frac{dE}{dx}(\hat{x}) = -2a^*y + 2a^*a\hat{x} = 0 \Rightarrow (a^*a)\hat{x} = a^*y$$

and:

$$\frac{d^2E}{dx^2}(\hat{x}) = 2a^*a > 0$$

the proof follows.

End of proof.

Example 1: Let us consider the problem defined in section 1.1:

$$\begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} 1 \\ 1 \end{pmatrix} x + \begin{pmatrix} e_1 \\ e_2 \end{pmatrix}.$$

With $a = (1, 1)^*$ the normal equation $(a^*a)\hat{x} = a^*y$ reads:

$$2\hat{x} = y_1 + y_2.$$

Hence the least-squares estimate $\hat{x} = (a^*a)^{-1}a^*y$ is:

$$\hat{x} = \frac{1}{2}(y_1 + y_2).$$

Thus, to estimate x , one adds the measurements and divides by the number of measurements. Hence the least-squares estimate equals in this case the arithmetic average. The least-squares estimates of the observations and observation errors follow from $\hat{y} = a\hat{x}$ and $\hat{e} = y - \hat{y}$ as:

$$\begin{pmatrix} \hat{y}_1 \\ \hat{y}_2 \end{pmatrix} = \frac{1}{2} \begin{pmatrix} y_1 + y_2 \\ y_1 + y_2 \end{pmatrix}, \text{ and } \begin{pmatrix} \hat{e}_1 \\ \hat{e}_2 \end{pmatrix} = \frac{1}{2} \begin{pmatrix} y_1 - y_2 \\ y_2 - y_1 \end{pmatrix}.$$

Finally the square of the length of $\hat{e}$ follows as:

$$\hat{e}^*\hat{e} = \frac{1}{2}(y_1 - y_2)^2.$$

End of example.

8 Adjustment theory

The matrices $a(a^*a)^{-1}a^*$ and $I - a(a^*a)^{-1}a^*$ that occur in (2) are known as *orthogonal projectors*. They play a central role in estimation theory. We will denote them as:

$$(3) \quad \boxed{P_a = a(a^*a)^{-1}a^* \quad , \quad P_a^\perp = I - a(a^*a)^{-1}a^*}$$

The matrix P_a projects *onto* a line spanned by a and *along* a line orthogonal to a . To see this, note that we can write $P_a y$ with the help of the cosine rule as (see figure 1.2):

$$P_a y = a(a^*a)^{-1}a^*y = a(a^*a)^{-1}\|a\|\|y\|\cos\alpha = \|y\|\cos\alpha \frac{a}{\|a\|} .$$

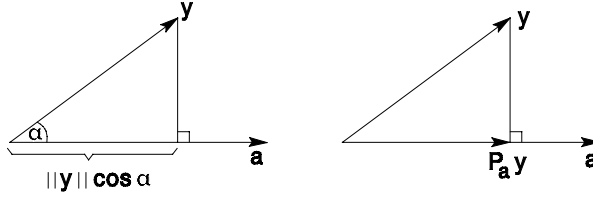


Figure 1.2

In a similar way we can show that $P_a^\perp$ projects *onto* a line orthogonal to a and *along* a line spanned by a (see figure 1.3).

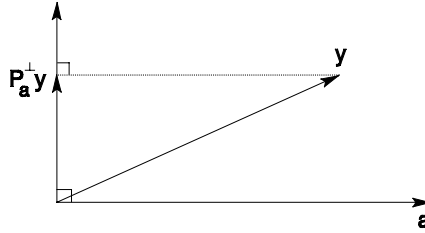


Figure 1.3

If we denote the vector that is orthogonal to a as b , thus $b^*a=0$, the projectors P_a and $P_a^\perp$ can alternatively be represented as:

$$(4) \quad \boxed{P_a = P_b^\perp = I - b(b^*b)^{-1}b^* \quad , \quad P_a^\perp = P_b = b(b^*b)^{-1}b^*}$$

To see this, consult figure 1.4 and note that with the help of the cosine rule we can write:

$$P_a^\perp y = \|y\| \cos \beta \frac{b}{\|b\|} = b \|b\|^{-2} b^* y = b(b^* b)^{-1} b^* y = P_b y.$$

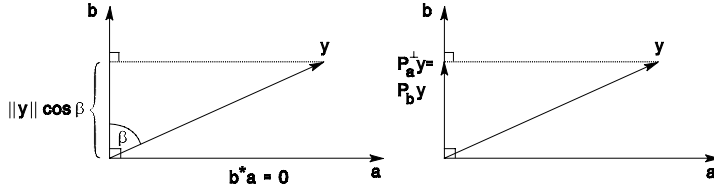


Figure 1.4

With (3) and (4) we have now two different representations of the orthogonal projectors $P_a = P_b^\perp$ and $P_a^\perp = P_b$. The representations of (3) are used in (2), which followed from solving equation (1) in a least-squares sense. We will now show that the representations of (4) correspond to solving the equation:

(5)

$$\begin{array}{c} b^* y = b^* e \\ 1 \times 2 \quad 2 \times 1 \quad 1 \times 2 \quad 2 \times 1 \end{array}$$

in a least-squares sense. Equation (5) follows from pre-multiplying equation (1) with b^* and noting that $b^* a = 0$. We know that the least-squares solution of (1) follows from solving the minimization problem:

$$\min_x \|y - ax\|^2.$$

We have used the notation $\|\cdot\|^2$ here to denote the square of the length of a vector. The minimization problem $\|y - ax\|^2$ can also be written as:

$$\min_z \|y - z\|^2 \quad \text{subject to} \quad z = ax.$$

The equation $z = ax$ is the *parametric representation* of a line through the origin with direction vector a . From linear algebra we know that this line can also be represented in *implicit form* with the help of the normal vector b . Thus we may replace $z = ax$ by $b^* z = 0$. This gives:

$$\min_z \|y - z\|^2 \quad \text{subject to} \quad b^* z = 0.$$

With $e = y - z$, this may also be written as:

$$\min_e \|e\|^2 \quad \text{subject to} \quad b^* y = b^* e.$$

The least-squares estimate $\hat{e}$ of the unknown e in (5) thus follows as the vector $\hat{e}$ which satisfies (5) and has minimal length. Figure 1.5 shows that $\hat{e}$ is obtained as the orthogonal projection of y onto the vector b . The least-squares estimates of (5) read therefore:

(6)

$$\hat{e} = P_b y = [b(b^*b)^{-1}b^*]y$$

$$\hat{y} = y - \hat{e} = P_b^\perp y = [I - b(b^*b)^{-1}b^*]y$$

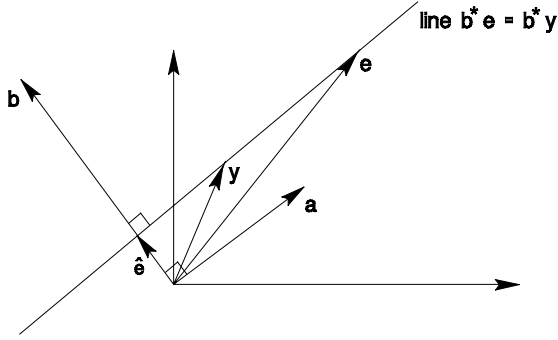


Figure 1.5

Example 2: Let us consider the problem defined in section 1.1:

$$\begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} 1 \\ 1 \end{pmatrix} x + \begin{pmatrix} e_1 \\ e_2 \end{pmatrix}.$$

The form corresponding to (5) is:

$$(1 \ -1) \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = (1 \ -1) \begin{pmatrix} e_1 \\ e_2 \end{pmatrix}.$$

This equation expresses the fact that the difference between the two observed distances y_1 and y_2 is unequal to zero in general because of the presence of measurement errors e_1 and e_2 .

Since the vector b is given in this case as:

$$b = (1 \ -1)^*,$$

we have:

$$P_b = \frac{1}{2} \begin{pmatrix} 1 & -1 \\ -1 & 1 \end{pmatrix} \quad \text{and} \quad P_b^\perp = \frac{1}{2} \begin{pmatrix} 1 & 1 \\ 1 & 1 \end{pmatrix}.$$

The least-squares estimates read therefore:

$$\begin{pmatrix} \hat{e}_1 \\ \hat{e}_2 \end{pmatrix} = \mathbf{P}_b \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \frac{1}{2} \begin{pmatrix} y_1 - y_2 \\ y_2 - y_1 \end{pmatrix} \quad \text{and} \quad \begin{pmatrix} \hat{y}_1 \\ \hat{y}_2 \end{pmatrix} = \mathbf{P}_b^\perp \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \frac{1}{2} \begin{pmatrix} y_1 + y_2 \\ y_1 + y_2 \end{pmatrix},$$

which is identical to the results obtained earlier.

End of example.

1.3 Weighted least-squares estimation

Up to now, we have regarded the two observations y_1 and y_2 as equally accurate, and solved:

$$\underset{2 \times 1}{y} = \underset{2 \times 11 \times 1}{a} \underset{2 \times 1}{x} + \underset{2 \times 1}{e}$$

by looking for the value $\hat{x}$ that minimized:

$$E(x) = (y_1 - a_1 x)^2 + (y_2 - a_2 x)^2 = (y - ax)^* I (y - ax) .$$

The least-squares principle can however be generalized by introducing a *symmetric, positive-definite weight matrix* W so that:

$$E(x) = w_{11}(y_1 - a_1 x)^2 + 2w_{12}(y_1 - a_1 x)(y_2 - a_2 x) + w_{22}(y_2 - a_2 x)^2 = (y - ax)^* W (y - ax) .$$

The elements of the 2×2 matrix W can be chosen to emphasize (or de-emphasize) the influence of specific observations upon the estimate $\hat{x}$. For instance, if $w_{11} > w_{22}$, then more importance is attached to the first observation, and the process of minimizing $E(x)$ tries harder to make $(y_1 - a_1 x)^2$ small. This seems reasonable if one believes that the first observation is more trustworthy than the second; for instance if observation y_1 is obtained from a more accurate instrument than observation y_2 .

The solution $\hat{x}$ of:

$$\min_x E(x) = \min_x (y - ax)^* W (y - ax),$$

is called the *weighted* least-squares (WLSQ) estimate of x . As we know the solution can be obtained by solving:

$$\frac{dE}{dx}(\hat{x}) = 0,$$

and checking whether:

$$\frac{d^2 E}{dx^2}(\hat{x}) > 0 .$$

From:

$$E(x) = y^* W y - 2y^* W a x + a^* W a x^2$$

follows:

$$\frac{dE}{dx}(\hat{x}) = -2a^* W y + 2a^* W a \hat{x} = 0$$

and:

$$\frac{d^2 E}{dx^2}(\hat{x}) = 2a^* W a .$$

Since the weight matrix W is assumed to be positive-definite we have $a^* W a > 0$ (recall from linear algebra that a matrix M is said to be *positive-definite* if it is symmetric and $z^* M z > 0 \forall z \neq 0$). From the first equation we therefore get:

$$\begin{aligned}
 & a^* W (y - a \hat{x}) = 0 \\
 & \Downarrow \\
 & (a^* W a) \hat{x} = a^* W y && \text{"normal equation"} \\
 & \Downarrow \\
 & \hat{x} = (a^* W a)^{-1} a^* W y && \text{"WLSQ-estimate of } x \text{"} \\
 & \Downarrow \\
 (7) \quad & \hat{y} = a \hat{x} = [a(a^* W a)^{-1} a^* W] y && \text{"WLSQ-estimate of } y \text{"} \\
 & \Downarrow \\
 & \hat{e} = y - \hat{y} = [I - a(a^* W a)^{-1} a^* W] y && \text{"WLSQ-estimate of } e \text{"} \\
 & \Downarrow \\
 & \hat{e}^* W \hat{e} = y^* [W - W a(a^* W a)^{-1} a^* W] y && \text{"weighted sum of squares"}
 \end{aligned}$$

Compare this result with (2).

Example 3: Consider again the problem defined in section 1.1:

$$\begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} 1 \\ 1 \end{pmatrix} x + \begin{pmatrix} e_1 \\ e_2 \end{pmatrix}.$$

As weight matrix we take:

$$W = \begin{pmatrix} w_{11} & 0 \\ 0 & w_{22} \end{pmatrix}.$$

With $a = (1 \ 1)^*$ the normal equation $(a^* W a) \hat{x} = a^* W y$ reads:

$$(w_{11} + w_{22})\hat{x} = w_{11}y_1 + w_{22}y_2 .$$

Hence, the weighted least-squares estimate $\hat{x}=(a^*Wa)^{-1}a^*Wy$ is:

$$\hat{x} = \frac{w_{11}y_1 + w_{22}y_2}{w_{11} + w_{22}} .$$

Thus instead of the arithmetic average of y_1 and y_2 , as we had with $W=I$, the above estimate is a *weighted* average of the data. This average is closer to y_1 than to y_2 if $w_{11}>w_{22}$. The WLSQ-estimate of the observations, $\hat{y}=a\hat{x}$, is:

$$\begin{pmatrix} \hat{y}_1 \\ \hat{y}_2 \end{pmatrix} = \frac{1}{w_{11} + w_{22}} \begin{pmatrix} w_{11}y_1 + w_{22}y_2 \\ w_{11}y_1 + w_{22}y_2 \end{pmatrix},$$

and the WLSQ-estimate of e , $\hat{e}=y - \hat{y}$, is:

$$\begin{pmatrix} \hat{e}_1 \\ \hat{e}_2 \end{pmatrix} = \frac{1}{w_{11} + w_{22}} \begin{pmatrix} w_{22}(y_1 - y_2) \\ w_{11}(y_2 - y_1) \end{pmatrix}$$

Note that $|\hat{e}_2|>|\hat{e}_1|$ if $w_{11}>w_{22}$.

Finally the weighted sum of squares $\hat{e}^*W\hat{e}$ is:

$$\hat{e}^*W\hat{e} = \frac{w_{11}w_{22}}{w_{11} + w_{22}} (y_1 - y_2)^2.$$

End of example.

In section 1.2 we introduced the least-squares principle by using the intuitively appealing geometric principle of orthogonal projection. In this section we introduced the weighted least-squares principle as the problem of minimizing the weighted sum of squares $(y-ax)^*W(y-ax)$. An interesting question is now, whether it is also possible to give a geometric interpretation to the *weighted* least-squares principle. This turns out to be the case. First note that the equation:

$$z^*Wz = \text{constant} = c$$

or:

$$w_{11}z_1^2 + 2w_{12}z_1z_2 + w_{22}z_2^2 = c$$

describes an *ellipse*. If $w_{11}=w_{22}$ and $w_{12}=0$, the ellipse is a circle. If $w_{11}\neq w_{22}$ and $w_{12}=0$, the ellipse has its major and minor axis parallel to the z_1 - and z_2 - axes, respectively. And if $w_{12}\neq 0$, the ellipse may have an arbitrary orientation (see figure 1.6).

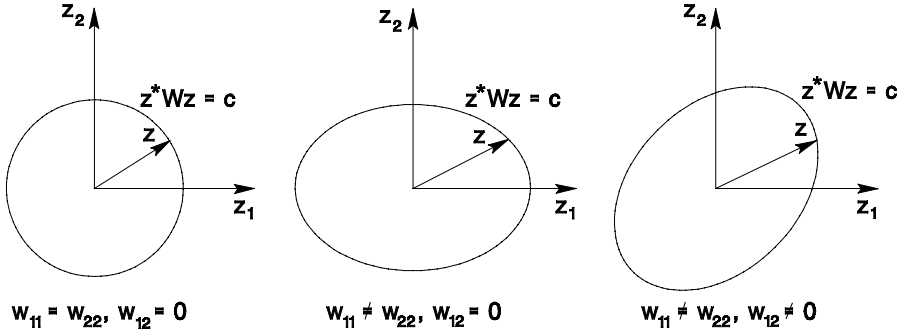


Figure 1.6

If we define a function $F(z)$ as $F(z) = z^* W z$, then we know from calculus that the *gradient* of $F(z)$ evaluated at z_0 is a vector which is *normal* to the ellipse $F(z) = c$ at z_0 . The gradient of $F(z)$ evaluated at z_0 is:

$$\begin{pmatrix} \frac{\partial F}{\partial z_1} \\ \frac{\partial F}{\partial z_2} \end{pmatrix}_{z_0} = \begin{pmatrix} 2w_{11}z_1 + 2w_{12}z_2 \\ 2w_{12}z_1 + 2w_{22}z_2 \end{pmatrix}_{z_0} = 2Wz_0 .$$

This vector is thus normal or orthogonal to the ellipse at z_0 (see figure 1.7).

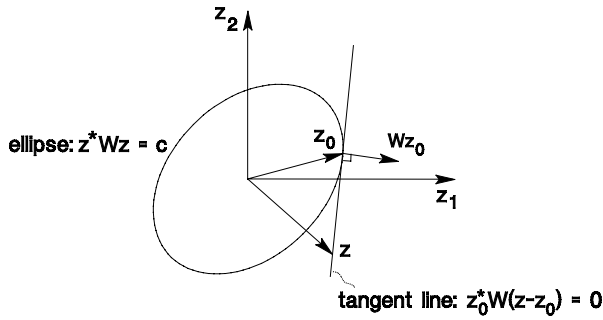


Figure 1.7

From this follows that the tangent line of the ellipse $z^* W z = c$ at z_0 is given by:

$$z_0^* W (z - z_0) = 0 .$$

Comparing this evaluation with the first equation of (7) shows that the WLSQ-estimate of x is given by that $\hat{x}$ for which $y - a\hat{x}$ is orthogonal to the normal Wa at a of the ellipse $z^*Wz = a^*Wa$. Hence, $y - a\hat{x}$ is parallel to the tangent line of the ellipse $z^*Wz = a^*Wa$ at a (see figure 1.8). Note that if $W=I$, the ellipse becomes a circle and the unweighted least-squares estimate is obtained through orthogonal projection (see figure 1.9).

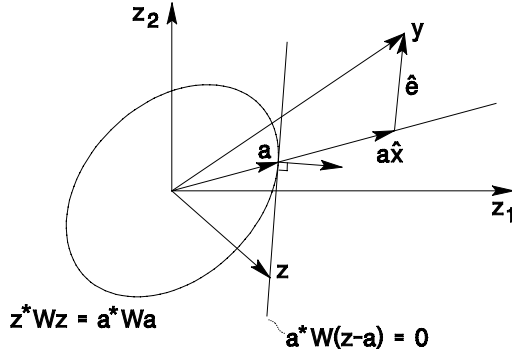


Figure 1.8

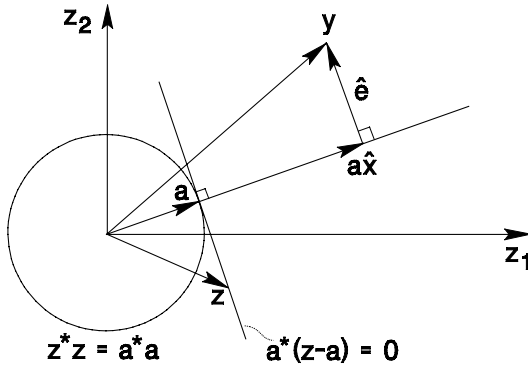


Figure 1.9

The matrices $a(a^*Wa)^{-1}a^*W$ and $I - a(a^*Wa)^{-1}a^*W$ that occur in (7) are known as *oblique projectors*.

We will denote them as:

$$(8) \quad \begin{aligned} P_{a,(Wa)^\perp} &= a(a^*Wa)^{-1}a^*W, \\ P_{a,(Wa)^\perp}^\perp &= P_{(Wa)^\perp,a} = I - a(a^*Wa)^{-1}a^*W \end{aligned}$$

The matrix $P_{a,(Wa)^\perp}$ projects *onto* a line spanned by a and *along* a line orthogonal to Wa .

To see this, note that we can write $P_{a,(Wa)^\perp}y$ with the help of the cosine rule as (see figure 1.10):

$$P_{a,(Wa)^\perp}y = a(a^*Wa)^{-1}a^*Wy = a \frac{\|Wa\|\|y\|\cos\alpha}{\|Wa\|\|a\|\cos\beta} = \|y\| \frac{\cos\alpha}{\cos\beta} \frac{a}{\|a\|}$$

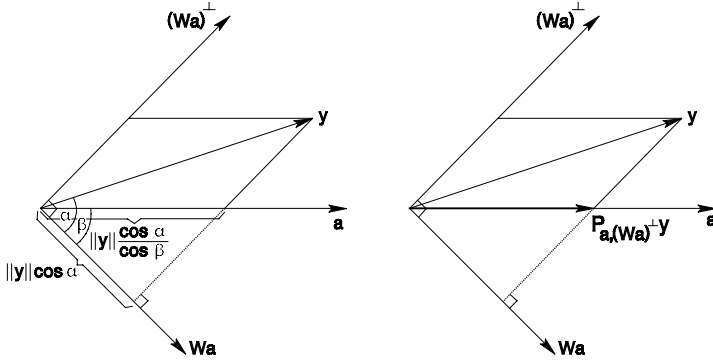


Figure 1.10

In a similar way we can show that $P_{a,(Wa)^\perp}^\perp$ projects *onto* a line orthogonal to Wa and *along* a line spanned by a (see figure 1.11).

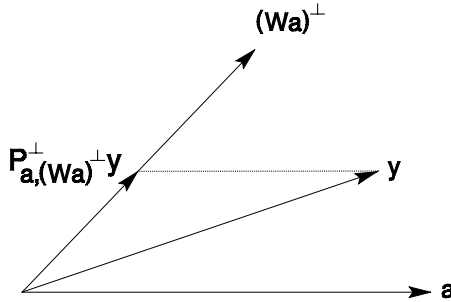


Figure 1.11

Above we have denoted the vector orthogonal to Wa as $(Wa)^\perp$; thus $((Wa)^\perp)^*(Wa)=0$. With the help of the vector b , which is defined in section 1.2 as being orthogonal to a , thus $b^*a=0$, we can write instead of $(Wa)^\perp$ also $W^{-1}b$, because $(W^{-1}b)^*(Wa)=0$ (the inverse W^{-1} of W exists, since we assumed W to be positive-definite). With the help of $W^{-1}b$ we can now represent the two projectors of (8) alternatively as:

(9)

$$\begin{aligned} P_{a,(Wb)^\perp} &= P_{a,W^{-1}b} = I - W^{-1}b(b^*W^{-1}b)^{-1}b^* , \\ P_{(Wb)^\perp,a} &= P_{W^{-1}b,a} = W^{-1}b(b^*W^{-1}b)^{-1}b^* \end{aligned}$$

To see this, consult figure 1.12 and note that with the help of the cosine rule we can write:

$$P_{W^{-1}b,a}y = \|y\| \frac{\cos \alpha}{\cos \beta} \frac{W^{-1}b}{\|W^{-1}b\|} = W^{-1}b \frac{\|y\| \cos \alpha}{\|W^{-1}b\| \cos \beta} = W^{-1}b \frac{b^*y}{(W^{-1}b)^*b} ,$$

or:

$$P_{W^{-1}b,a}y = W^{-1}b(b^*W^{-1}b)^{-1}b^*y .$$

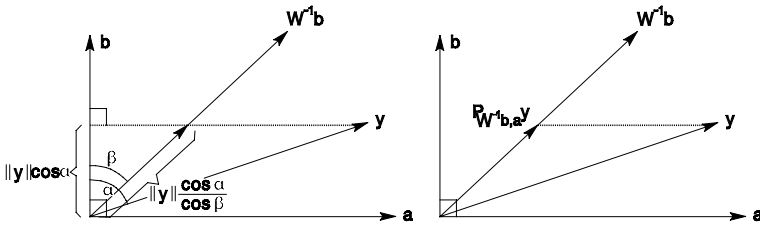


Figure 1.12

Note that if $W=I$, then $W^{-1}=I$ and equations (8) and (9) reduce to equations (3) and (4) respectively. In section 1.2 we showed that for $W=I$ the representations of (9) correspond to solving equation (5) in an unweighted least-squares sense. We will now show that the representations of (9) correspond to solving (5) in a weighted least-squares sense. From the equivalence of the following minimization problems:

$$\min_x (y - ax)^* W(y - ax),$$

$$\min_z (y - z)^* W(y - z) \quad \text{subject to} \quad z = ax,$$

$$\min_z (y - z)^* W(y - z) \quad \text{subject to} \quad b^*z = 0,$$

$$\min_e (e)^* W(e) \quad \text{subject to} \quad b^*y = b^*e,$$

follows that the weighted least-squares estimate of e is given by that $\hat{e}$ on the line $b^*y=b^*e$ for which e^*We is minimal. We know that $e^*We = \text{constant} = c$ is the equation of an ellipse. For different values of c we get different ellipses and the larger the value of c is, the larger the ellipse (see figure 1.13).

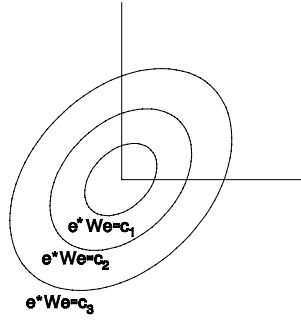


Figure 1.13 $c_3 > c_2 > c_1$

The solution $\hat{e}$ is therefore given by the *point of tangency* of the line $b^*y = b^*e$ with one of the ellipses $e^*We = c$ (see figure 1.14).

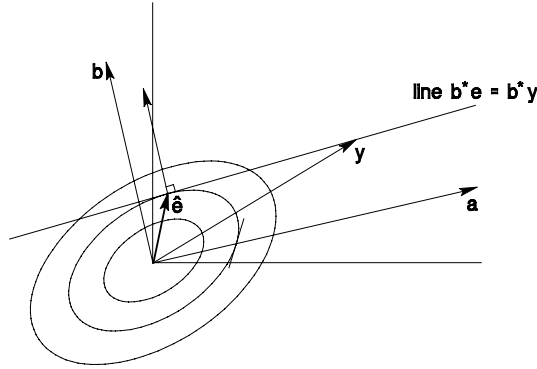


Figure 1.14

At this point of tangency the normal of the ellipse is parallel to the normal of the line $b^*e = b^*y$. But since the normal of the line $b^*e = b^*y$ is given by b , it follows that the normal of the ellipse at $\hat{e}$ is also given by b . From this we can conclude that $\hat{e}$ has to be parallel to $W^{-1}b$. Hence we may write:

$$\hat{e} = W^{-1}b\alpha .$$

The unknown scalar α may now be determined from the fact that $\hat{e}$ has to lie on the line $b^*e = b^*y$. Pre-multiplying $\hat{e} = W^{-1}b\alpha$ with b^* gives for α

$$\alpha = (b^*W^{-1}b)^{-1}b^*\hat{e} = (b^*W^{-1}b)^{-1}b^*y .$$

Substituting this in the equation $\hat{e} = W^{-1}b\alpha$ shows that the weighted least-squares estimates of equation (5) read:

$$(10) \quad \begin{aligned} \hat{e} &= P_{W^{-1}b,a}y = [W^{-1}b(b^*W^{-1}b)^{-1}b^*]y \\ \hat{y} &= y - \hat{e} = P_{W^{-1}b,a}^\perp y = [I - W^{-1}b(b^*W^{-1}b)^{-1}b^*]y \end{aligned}$$

Compare this result with equation (6).

Example 4: Consider again the problem defined in section 1.1:

$$\begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} 1 \\ 1 \end{pmatrix}x + \begin{pmatrix} e_1 \\ e_2 \end{pmatrix}$$

The form corresponding to (5) is:

$$(1-1) \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = (1-1) \begin{pmatrix} e_1 \\ e_2 \end{pmatrix}.$$

thus the vector b is given as:

$$b = (1-1)^*.$$

As weight matrix we take:

$$W = \begin{pmatrix} w_{11} & 0 \\ 0 & w_{22} \end{pmatrix}.$$

Its inverse reads therefore:

$$W^{-1} = \begin{pmatrix} w_{11}^{-1} & 0 \\ 0 & w_{22}^{-1} \end{pmatrix}.$$

The two projectors $P_{W^{-1}b,a}$ and $P_{W^{-1}b,a}^\perp$ can now be computed as:

$$P_{W^{-1}b,a} = W^{-1}b(b^*W^{-1}b)^{-1}b^* = \frac{w_{11}w_{22}}{w_{11} + w_{22}} \begin{pmatrix} w_{11}^{-1} & -w_{11}^{-1} \\ -w_{22}^{-1} & w_{22}^{-1} \end{pmatrix},$$

and

$$\mathbf{P}_{W^{-1}\mathbf{b},\mathbf{a}}^\perp = \mathbf{I} - \mathbf{W}^{-1}\mathbf{b}(\mathbf{b}^*\mathbf{W}^{-1}\mathbf{b})^{-1}\mathbf{b}^* = \frac{w_{11}w_{22}}{w_{11} + w_{22}} \begin{pmatrix} w_{22}^{-1} & w_{11}^{-1} \\ w_{22}^{-1} & w_{11}^{-1} \end{pmatrix}.$$

The weighted least-squares estimates $\hat{e}$ and $\hat{y}$ read therefore:

$$\begin{pmatrix} \hat{e}_1 \\ \hat{e}_2 \end{pmatrix} = \mathbf{P}_{W^{-1}\mathbf{b},\mathbf{a}}^\perp \mathbf{y} = \frac{1}{w_{11} + w_{22}} \begin{pmatrix} w_{22}(y_1 - y_2) \\ w_{11}(y_2 - y_1) \end{pmatrix},$$

and

$$\begin{pmatrix} \hat{y}_1 \\ \hat{y}_2 \end{pmatrix} = \mathbf{P}_{W^{-1}\mathbf{b},\mathbf{a}}^\perp \mathbf{y} = \frac{1}{w_{11} + w_{22}} \begin{pmatrix} w_{11}y_1 + w_{22}y_2 \\ w_{11}y_1 + w_{22}y_2 \end{pmatrix}.$$

These results are identical to the results obtained in example 3.

End of example.

1.4 Weighted least-squares is also orthogonal projection

In the previous two sections we have seen that unweighted and weighted least-squares estimation could be interpreted as *orthogonal* and *oblique* projection respectively (see figure 1.15).

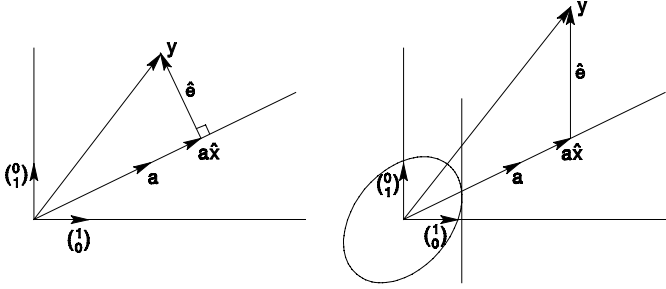


Figure 1.15

It is important to realize however that these geometric interpretations are very much dependent on the choice of base vectors with the help of which the vectors $\mathbf{y}$, $\mathbf{a}$ and $\mathbf{e}$ are constructed. In all the cases considered so far we have implicitly assumed to be dealing with the orthogonal base vectors $(1,0)^*$ and $(0,1)^*$. That is, the vectors $\mathbf{y}$, $\mathbf{a}$ and $\mathbf{e}$ of the equation:

$$\mathbf{y} = \mathbf{a}\mathbf{x} + \mathbf{e},$$

were constructed from the scalars y_1 , y_2 , a_1 , a_2 , e_1 and e_2 as:

$$\begin{aligned}
 y &= \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = y_1 \begin{pmatrix} 1 \\ 0 \end{pmatrix} + y_2 \begin{pmatrix} 0 \\ 1 \end{pmatrix}, \\
 a &= \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} = a_1 \begin{pmatrix} 1 \\ 0 \end{pmatrix} + a_2 \begin{pmatrix} 0 \\ 1 \end{pmatrix}, \\
 e &= \begin{pmatrix} e_1 \\ e_2 \end{pmatrix} = e_1 \begin{pmatrix} 1 \\ 0 \end{pmatrix} + e_2 \begin{pmatrix} 0 \\ 1 \end{pmatrix}.
 \end{aligned}
 \tag{11}$$

Let us now investigate what happens if instead of the orthonormal set $(1,0)^*$ and $(0,1)^*$, a different set, i.e. a non-orthonormal set, of base vectors is chosen. Let us denote these new base vectors as d_1 and d_2 . With the help of the base vectors d_1 and d_2 , the vectors y , a and e can then be constructed from the scalars y_1 , y_2 , a_1 , a_2 , e_1 and e_2 as:

$$\begin{aligned}
 y &= y_1 d_1 + y_2 d_2, \\
 a &= a_1 d_1 + a_2 d_2, \\
 e &= e_1 d_1 + e_2 d_2.
 \end{aligned}
 \tag{12}$$

Again the equation $y=ax+e$ holds. This is shown in figure 1.16a. Note that the vectors of (12) are not the same as those of (11). They do however have the same components.

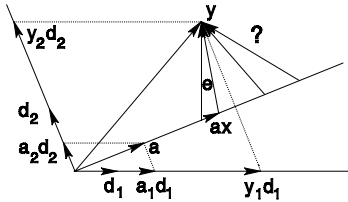


Figure 1.16a

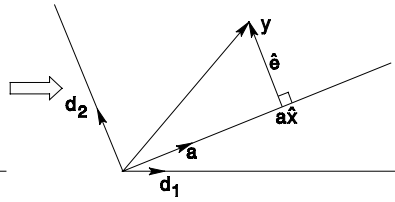


Figure 1.16b

Looking at figure 1.16 it seems again, like in section 1.2, intuitively appealing to estimate x as $\hat{x}$ such that $a\hat{x}$ is as close as possible to given vector y . This is the case if the vector $y-a\hat{x}$ is orthogonal to the vector a ; that is, if:

$$a^* (y - a\hat{x}) = 0. \quad \text{"orthogonality"}
 \tag{13}$$

With (12) this gives:

$$(a_1 d_1 + a_2 d_2)^* [y_1 d_1 + y_2 d_2 - (a_1 d_1 + a_2 d_2) \hat{x}] = 0,$$

or

$$\begin{pmatrix} a_1 \\ a_2 \end{pmatrix}^* \begin{pmatrix} d_1^* d_1 & d_1^* d_2 \\ d_2^* d_1 & d_2^* d_2 \end{pmatrix} \begin{bmatrix} y_1 \\ y_2 \end{bmatrix} - \begin{pmatrix} a_1 \\ a_2 \end{pmatrix}^* \hat{x} = 0,$$

or

(14)

$$\begin{pmatrix} a_1 \\ a_2 \end{pmatrix}^* D^* D \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} - \begin{pmatrix} a_1 \\ a_2 \end{pmatrix}^* \hat{x} = 0,$$

with

$$D = (d_1 \ d_2).$$

The matrix D^*D of (14) is symmetric, since $(D^*D)^* = D^*D$, and it is positive-definite, since $z^*D^*Dz = (Dz)^*(Dz) > 0$ for all $z \neq 0$. Hence, the matrix D^*D may be interpreted as a weight matrix. By interpreting the matrix D^*D of (14) as a weight matrix we see that the solution of (13) can in fact be interpreted as a *weighted least-squares* solution, i.e. a solution of the minimization problem:

$$\min_x \left[\begin{pmatrix} y_1 \\ y_2 \end{pmatrix} - \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} x \right]^* D^* D \left[\begin{pmatrix} y_1 \\ y_2 \end{pmatrix} - \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} x \right].$$

We have thus shown that if in a weighted least-squares problem the weight matrix W can be written as the product of a matrix D^* and its transpose D , then the weighted least-squares estimate $\hat{y} = a\hat{x}$ can be interpreted as being obtained through *orthogonal* projection of $y = y_1d_1 + y_2d_2$ onto the line with direction vector $a = a_1d_1 + a_2d_2$, where d_1 and d_2 are the column vectors of matrix D (see figure 1.17).

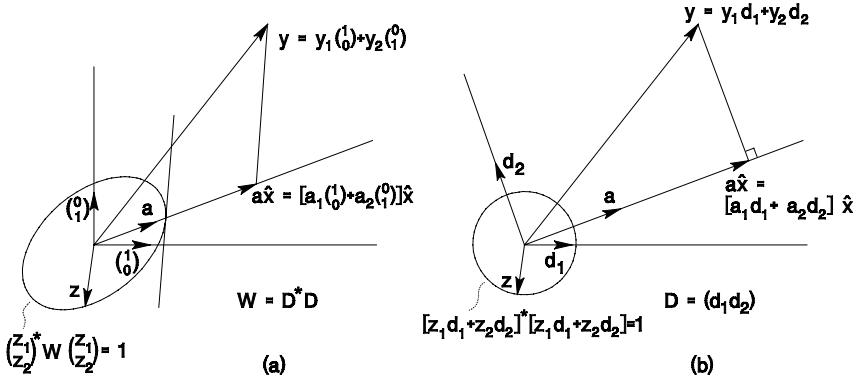


Figure 1.17

The crucial question is now whether *every* weight matrix W , i.e. every symmetric and positive-definite matrix, can be written as a product of a matrix and its transpose, i.e. as $W = D^*D$. If this is the case, then every weighted least-squares problem can be interpreted as a problem of orthogonal projection! The answer is: *yes*, every weight matrix W can be written as $W = D^*D$. For a proof, see the *intermezzo* at the end of this section.

The weight matrix $W=D^*D$ can be interpreted as defining an *inner product* (Strang, 1988). In the unweighted case, $W=I$, the inner product of two vectors u and v with components u_1, u_2 and v_1, v_2 respectively, is given as:

$$(u, v)_I = \begin{pmatrix} u_1 \\ u_2 \end{pmatrix}^* I \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} = \left[u_1 \begin{pmatrix} 1 \\ 0 \end{pmatrix} + u_2 \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right]^* \left[v_1 \begin{pmatrix} 1 \\ 0 \end{pmatrix} + v_2 \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right].$$

This is the so-called *standard inner product*, where the inner product is represented by the unit matrix I . The square of the length of the vector u reads then:

$$\|u\|_I^2 = u_1^2 + u_2^2.$$

In the weighted case, $W=D^*D$, the inner product of u and v is given as:

$$(15) \quad (u, v)_W = \begin{pmatrix} u_1 \\ u_2 \end{pmatrix}^* D^* D \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} = (u_1 \mathbf{d}_1 + u_2 \mathbf{d}_2)^* (v_1 \mathbf{d}_1 + v_2 \mathbf{d}_2).$$

The square of the length of the vector u reads then:

$$\|u\|_W^2 = u_1^2 \mathbf{d}_1^* \mathbf{d}_1 + 2u_1 u_2 \mathbf{d}_1^* \mathbf{d}_2 + u_2^2 \mathbf{d}_2^* \mathbf{d}_2.$$

Note that the two ways in which $(u, v)_W$ is written in (15) correspond with the two figures 1.17a and 1.17b respectively. In figure 1.17a orthonormal base vectors are used and the metric is *visualized* by the ellipse:

$$(z, z)_W = \begin{pmatrix} z_1 \\ z_2 \end{pmatrix}^* W \begin{pmatrix} z_1 \\ z_2 \end{pmatrix} = 1.$$

In figure 1.17b however, the ellipse is replaced by a circle:

$$(z, z)_W = (z_1 \mathbf{d}_1 + z_2 \mathbf{d}_2)^* I (z_1 \mathbf{d}_1 + z_2 \mathbf{d}_2) = 1,$$

and the metric is *visualized* by the non-orthonormal base vectors $\mathbf{d}_1$ and $\mathbf{d}_2$.

Now that we have established the fact that weighted least-squares estimation can be interpreted as orthogonal projection, it must also be possible to interpret the oblique projectors of equation (8) as orthogonal projectors. We will show that the matrix $P_{a, (Wa)^\perp}$ of (8) can be interpreted as the matrix which projects orthogonally onto a line spanned by the vector a , where orthogonality is measured with respect to the inner product defining weight matrix W . From figure 1.18 follows that:

$$\hat{y} = \|y\|_W \cos \alpha \frac{a}{\|a\|_W} = a \|a\|_W^{-2} (a, y)_W,$$

or

$$\hat{y}_1 \mathbf{d}_1 + \hat{y}_2 \mathbf{d}_2 = (\mathbf{a}_1 \mathbf{d}_1 + \mathbf{a}_2 \mathbf{d}_2) \left[\begin{pmatrix} \mathbf{a}_1 \\ \mathbf{a}_2 \end{pmatrix}^* W \begin{pmatrix} \mathbf{a}_1 \\ \mathbf{a}_2 \end{pmatrix} \right]^{-1} \begin{pmatrix} \mathbf{a}_1 \\ \mathbf{a}_2 \end{pmatrix}^* W \begin{pmatrix} y_1 \\ y_2 \end{pmatrix},$$

or

$$\begin{pmatrix} \hat{y}_1 \\ \hat{y}_2 \end{pmatrix} = \begin{pmatrix} \mathbf{a}_1 \\ \mathbf{a}_2 \end{pmatrix} \left[\begin{pmatrix} \mathbf{a}_1 \\ \mathbf{a}_2 \end{pmatrix}^* W \begin{pmatrix} \mathbf{a}_1 \\ \mathbf{a}_2 \end{pmatrix} \right]^{-1} \begin{pmatrix} \mathbf{a}_1 \\ \mathbf{a}_2 \end{pmatrix}^* W \begin{pmatrix} y_1 \\ y_2 \end{pmatrix},$$

and in this expression we recognize indeed the matrix $\mathbf{P}_{\mathbf{a},(W\mathbf{a})}^\perp$ of equation (8). From now on we will assume that it is clear from the context which inner-product is taken and therefore also denote the matrix $\mathbf{P}_{\mathbf{a},(W\mathbf{a})}^\perp$ as $\mathbf{P}_\mathbf{a}$.

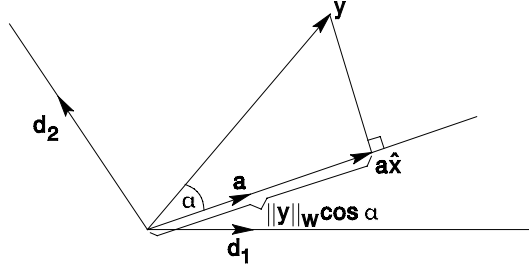


Figure 1.18

Intermezzo: Cholesky decomposition:

If W is an $n \times n$ symmetric positive definite matrix, there is a nonsingular lower triangular matrix L with positive diagonal elements such that $W = LL^*$.

Proof: The proof is by induction. Clearly, the result is true for $n=1$ and the induction hypothesis is that it is true for any $(n-1) \times (n-1)$ symmetric positive definite matrix.

Let

$$W = \begin{pmatrix} A & \mathbf{a} \\ \mathbf{a}^* & b_{nn} \end{pmatrix},$$

where A is a $(n-1) \times (n-1)$ matrix. Since W is symmetric positive definite, it will be clear that A is also symmetric positive definite. Thus, by the induction hypothesis, there is a nonsingular lower triangular matrix L_{n-1} with positive diagonal elements such that $A = L_{n-1} L_{n-1}^*$.

Let

$$L = \begin{pmatrix} L_{n-1} & \mathbf{0} \\ \mathbf{l}^* & \beta \end{pmatrix},$$

where $\mathbf{l}$ and β are to be obtained by equating LL^* to W :

$$LL^* = \begin{pmatrix} L_{n-1} & 0 \\ l^* & \beta \end{pmatrix} \begin{pmatrix} L_{n-1}^* & l \\ 0 & \beta \end{pmatrix} = \begin{pmatrix} A & a \\ a^* & b_{nn} \end{pmatrix} = W.$$

This implies that l and β must satisfy:

$$L_{n-1}l = a, \quad l^*l + \beta^2 = b_{nn}.$$

Since L_{n-1} is nonsingular, the first equation determines l as $l = L_{n-1}^{-1}a$. The second equation allows a positive β provided that $b_{nn} - l^*l > 0$. In order to prove that $b_{nn} - l^*l > 0$ we make use of the fact that W is positive definite, i.e.:

$$x^*Wx = (x_1^* x_2^*) \begin{pmatrix} A & a \\ a^* & b_{nn} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = (x_1^* x_2^*) \begin{pmatrix} L_{n-1}L_{n-1}^* & a \\ a^* & b_{nn} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} > 0 \quad \forall x \neq 0$$

Substitution of $x_1 = -(L_{n-1}^*)^{-1}l$ and $x_2 = 1$ shows that $b_{nn} - l^*l > 0$.

End of proof.

A simple example of a Cholesky decomposition is:

$$\begin{pmatrix} 4 & 2 \\ 2 & 4 \end{pmatrix} = \begin{pmatrix} 2 & 0 \\ 1 & \sqrt{3} \end{pmatrix} \begin{pmatrix} 2 & 1 \\ 0 & \sqrt{3} \end{pmatrix}.$$

1.5 The mean and covariance matrix of least-squares estimators

From now on we shall make no distinction between unweighted and weighted least-squares estimation, and just speak of least-squares estimation. From the context will be clear whether $W=I$ or not.

In our discussions so far no assumptions regarding the probabilistic nature of the observations or measurement errors were made. We started by writing the two equations:

$$\begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} a_1 \\ a_2 \end{pmatrix} x + \begin{pmatrix} e_1 \\ e_2 \end{pmatrix}$$

in vector notation as:

$$y = ax + e,$$

and showed that the least-squares estimates of x , y and e follow from solving the minimization problem:

$$\min_x (y - ax)^* W (y - ax)$$

as:

$$(16) \quad \begin{aligned} \hat{x} &= (a^* W a)^{-1} a^* W y, \\ \hat{y} &= a \hat{x}, \\ \hat{e} &= y - a \hat{x}. \end{aligned}$$

Vector $\hat{e}$ contains the least-squares residuals. Since functions of random variables are again random variables, it follows that if the observation vector is assumed to be a random vector $\underline{y}$ (note: the underscore indicates that we are dealing with random variables) and substituted for y in (16), the left-hand sides of (16) also become random variables:

$$(17) \quad \boxed{\begin{aligned} \hat{\underline{x}} &= (a^* W a)^{-1} a^* W \underline{y}, \\ \hat{\underline{y}} &= a \hat{\underline{x}}, \\ \hat{\underline{e}} &= \underline{y} - a \hat{\underline{x}} \end{aligned}}$$

These random variables are called least-squares *estimators*, whereas if $\underline{y}$ is replaced by its sample value y one speaks of least-squares *estimates*.

From your lecture notes in statistics you know how to derive the probability density functions of the estimators of (17) from the probability density function $p_{\underline{y}}(y)$ of $\underline{y}$. You also know that if the probability density function $p_{\underline{y}}(y)$ of $\underline{y}$ is "normal" or "Gaussian", then also $\hat{\underline{x}}$, $\hat{\underline{y}}$ and $\hat{\underline{e}}$ are normally distributed. This follows, since $\hat{\underline{x}}$, $\hat{\underline{y}}$ and $\hat{\underline{e}}$ are *linear functions of* $\underline{y}$. You know furthermore how to compute the means and covariance matrices of $\underline{y}$. This is particularly straightforward in case of linear functions (see also appendix E and F).

Application of the "propagation law of means" to (17) gives:

$$(18) \quad \begin{aligned} E\{\hat{\underline{x}}\} &= (a^* W a)^{-1} a^* W E\{\underline{y}\} \\ E\{\hat{\underline{y}}\} &= a E\{\hat{\underline{x}}\} = P_a E\{\underline{y}\} \\ E\{\hat{\underline{e}}\} &= E\{\underline{y}\} - a E\{\hat{\underline{x}}\} = (I - P_a) E\{\underline{y}\} = P_a^\perp E\{\underline{y}\} \end{aligned}$$

where $E\{\cdot\}$ denotes *mathematical expectation*. The result (18) holds true irrespective the type of probability distribution one assumes for $\underline{y}$. Let us now, without assuming anything about the type of the probability distribution of $\underline{y}$, be a bit more specific about the mean $E\{\underline{y}\}$ of $\underline{y}$. We will assume that the randomness of the observation vector $\underline{y}$ is a consequence of the random nature of the measurement errors. Thus we assume that $\underline{y}$ can be written as the sum of a deterministic, i.e. non-random, component ax and a random component $\underline{e}$:

$$(19) \quad \boxed{\underline{y} = ax + \underline{e}}$$

Furthermore we assume that the mean $E\{\underline{e}\}$ of $\underline{e}$ is zero. This assumption seems acceptable if the measurement errors are zero "on the average".

With (19) the assumption $E\{\underline{e}\}=0$ implies that:

$$(20) \quad E\{\underline{y}\} = \underline{a}x$$

Substitution of (20) into (18) gives then:

$$(21) \quad \begin{aligned} E\{\hat{x}\} &= x \\ E\{\hat{y}\} &= \underline{a}x = E\{\underline{y}\} \\ E\{\hat{e}\} &= E\{\underline{y}\} - \underline{a}x = 0 \end{aligned}$$

This important result shows that if (19) with (20) holds, the least-squares estimators $\hat{x}$, $\hat{y}$ and $\hat{e}$ are *unbiased* estimators. Note that this unbiasedness property of least-squares estimators is independent of the choice for the weight matrix W ! Let us now derive the covariance matrices of the estimators $\hat{x}$, $\hat{y}$ and $\hat{e}$. We will denote the covariance matrix of $\underline{y}$ as Q_y , the variance of $\hat{x}$ as $\sigma_{\hat{x}}^2$ and the covariance matrices of $\hat{y}$ and $\hat{e}$ as $Q_{\hat{y}}$ and $Q_{\hat{e}}$ respectively. Application of the "propagation law of covariances" to (17) gives:

$$(22) \quad \begin{aligned} \sigma_{\hat{x}}^2 &= (\underline{a}^* W \underline{a})^{-1} \underline{a}^* W Q_y W \underline{a} (\underline{a}^* W \underline{a})^{-1} \\ Q_{\hat{y}} &= \underline{a} \sigma_{\hat{x}}^2 \underline{a}^* = \underline{P}_a Q_y \underline{P}_a^* \\ Q_{\hat{e}} &= Q_y - \underline{a} Q_{\hat{y}} - Q_{\hat{y}} \underline{a}^* + \underline{a} \sigma_{\hat{x}}^2 \underline{a}^* \\ &= Q_y - \underline{P}_a Q_y - Q_y \underline{P}_a^* + \underline{P}_a Q_y \underline{P}_a^* \end{aligned}$$

As you know from your lecture notes in Statistics, the variance of a random variable is a measure for the spread of the probability density function around the mean value of the random variable. It is a measure of the variability one can expect to see when independent samples are drawn from the probability distribution. It is therefore in general desirable to have estimators with small variances. This is particularly true for our least-squares estimators. Together with the property of unbiasedness, a small variance would namely imply that there is a high probability that a sample from the distribution of $\hat{x}$ will be close to the true, but unknown, value x . Since we left the choice for the weight matrix W open up till now, let us investigate what particular weight matrix makes $\sigma_{\hat{x}}^2$ minimal. Using the decomposition $Q_y = D^* D$, it follows from (22) that $\sigma_{\hat{x}}^2$ can be written as:

$$\sigma_{\hat{x}}^2 = \frac{\underline{a}^* W D^* D W \underline{a}}{(\underline{a}^* D^{-1} D W \underline{a})^2} = \frac{\|D W \underline{a}\|^2}{\|D W \underline{a}\|^2 \|(D^*)^{-1} \underline{a}\|^2 \cos^2 \alpha} = \frac{1}{\|(D^*)^{-1} \underline{a}\|^2 \cos^2 \alpha}$$

where α is the angle between the two vectors $D W \underline{a}$ and $(D^*)^{-1} \underline{a}$. From this result follows that the variance $\sigma_{\hat{x}}^2$ is minimal when the angle α is zero; that is, when the two vectors $D W \underline{a}$ and $(D^*)^{-1} \underline{a}$ are parallel. The variance $\sigma_{\hat{x}}^2$ is therefore minimal when $D W \underline{a} = (D^*)^{-1} \underline{a}$ or $D^* D W \underline{a} = \underline{a}$ or $Q_y W \underline{a} = \underline{a}$. Hence, the variance is minimal when we choose the weight matrix W as:

(23)

$$W = Q_y^{-1}$$

With this choice for the weight matrix, the variance $\sigma_{\hat{x}}^2$ becomes:

(24)

$$\sigma_{\hat{x}}^2 = (a^* Q_y^{-1} a)^{-1}.$$

1.6 Best Linear Unbiased Estimators (BLUE's)

Up till now we have been dealing with least-squares estimation. The least-squares estimates were derived from the *deterministic* principles of "orthogonality" and "minimum distance". That is, no probabilistic considerations were taken into account in the derivation of least-squares estimates.

Nevertheless, it was established in the previous section that least-squares estimators do have some "optimal properties" in a probabilistic sense. They are *unbiased estimators* independent of the choice for the weight matrix W , and their *variance can be minimized* by choosing the weight matrix as the inverse of the covariance matrix of the observations. Moreover, they are *linear* estimators, i.e. linear functions of the observations.

Estimators which are linear, unbiased and have minimum variance are called best linear unbiased estimators or simply BLUE's. "Best" in this case means minimum variance. Thus least-squares estimators are BLUE's in case the weight matrix equals the inverse of the covariance matrix of the observations. In order to find out how large the class of best linear unbiased estimators is (is there only one? or are there more than one?), we will now derive this class. Thus instead of starting from the least-squares principle of "orthogonality" we start from the conditions of linearity, unbiasedness and best. Again we assume that the mean of y satisfies:

(25)

$$E\{y\} = ax.$$

Our problem is to find an estimator $\hat{x}$ of the unknown parameter x , where $\hat{x}$ is a function of y . Since we want the estimator $\hat{x}$ to be a *linear* function of y , $\hat{x}$ must be of the form:

(26)

$$\hat{x} = l^* y \quad \text{"L-property"}$$

Furthermore, we want $\hat{x}$ to be an *unbiased* estimator of x for all the values x may take. This means that:

(27)

$$E\{\hat{x}\} = x, \quad \forall x \quad \text{"U-property"}$$

And finally, we want to find those estimators $\hat{x}$, which in the class of estimators that satisfy (26) and (27), have *minimum variance*:

(28)

$$\sigma_{\hat{x}}^2 \text{ minimal in the class of LU-estimators} \quad \text{"B-property"}$$

It will be clear that the class of linear estimators is infinite. Let us therefore look at the class of linear and unbiased estimators.

From taking the expectation of (26) follows, with (25) and (27) that:

$$E\{\hat{x}\} = l^* E\{y\} = l^* a x = x \quad \forall x$$

or

(29)

$$l^* a = 1$$

Thus for LU-estimators the vector l has to satisfy (29) (see figure 1.19).

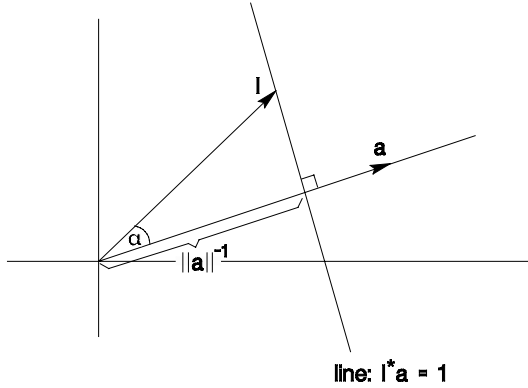


Figure 1.19: $l^* a = \|l\| \|a\| \cos \alpha = 1$

It will be clear that, although the class of permissible vectors l has been narrowed down to the line $l^* a = 1$, the class of LU-estimators is still infinite. For instance, every vector l of the form:

$$l = c(a^* c)^{-1},$$

where c is arbitrary, is permitted.

Let us now include the "B-property". The variance of the estimator $\hat{x}$ follows from applying the "propagation law of variances" to (26). This gives:

(30)

$$\sigma_{\hat{x}}^2 = l^* Q_y l$$

The "B-property" implies now that we have to find the vectors l , let us denote them by $\hat{l}$, that minimize $\sigma_{\hat{x}}^2$ of (30) and at the same time satisfy (29). We thus have the minimization problem:

(31)

$$\min_l l^* Q_y l \quad \text{subject to} \quad l^* a = 1$$

With the decomposition $Q_y = D^* D$, this can also be written as:

$$\min_l (Dl)^* (Dl) \quad \text{subject to} \quad (Dl)^* (D^*)^{-1} a = 1$$

or as:

$$(32) \quad \min_{\bar{l}} \bar{l}^* \bar{l} \quad \text{subject to} \quad \bar{l}^* (D^*)^{-1} a = 1 ,$$

where:

$$(33) \quad \bar{l} = D l .$$

Geometrically, the minimization problem (32) boils down to finding those vectors $\hat{\bar{l}}$ whose endpoints lie along the line $\bar{l}^* (D^*)^{-1} a = 1$ and who have minimum length (see figure 1.20).

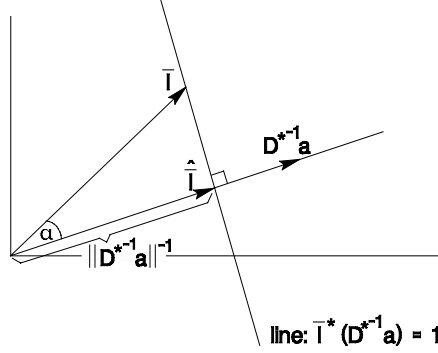


Figure 1.20: $\bar{l}^* (D^*)^{-1} a = \|\bar{l}\| \| (D^*)^{-1} a \| \cos \alpha = 1$

From figure 1.20 it is clear that there is only one such vector $\hat{\bar{l}}$. This vector lies along vector $D^{*-1} a$ and has length $\| (D^*)^{-1} a \|^{-1}$. Thus:

$$(34) \quad \hat{\bar{l}} = (D^*)^{-1} a \beta ,$$

where the positive scalar β follows from:

$$\hat{\bar{l}}^* \hat{\bar{l}} = \| (D^*)^{-1} a \|^2 \beta^2 = \| (D^*)^{-1} a \|^2$$

as:

$$(35) \quad \beta = \| (D^*)^{-1} a \|^2$$

Substitution of (35) into (34) gives with (33):

$$\hat{l} = D^{-1} (D^*)^{-1} a [a^* D^{-1} (D^*)^{-1} a]^{-1}$$

or:

$$(36) \quad \boxed{\hat{l} = Q_y^{-1} a (a^* Q_y^{-1} a)^{-1}}$$

The best linear unbiased estimator $\hat{x}$ of x follows therefore with (26) as:

$$(37) \quad \hat{x} = (a^* Q_y^{-1} a)^{-1} a^* Q_y^{-1} y$$

From this result we can draw the important conclusion that the best linear unbiased estimator is *unique* and equal to the least-squares estimator if $W=Q_y^{-1}$!

We will now give an *alternative* derivation of the solution (37). Equation (29) is the equation of a line with a direction vector orthogonal to a (see figure 1.20a). Since the vector b

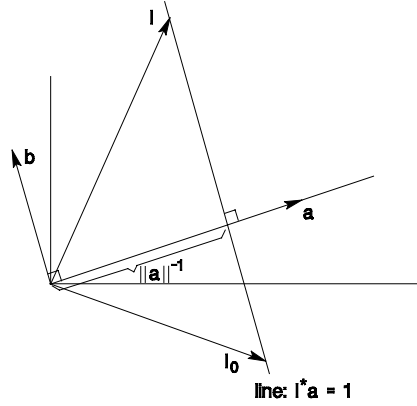


Figure 1.20a: $l^* a = 1$ is equivalent to $l = l_o + b\alpha$

was defined as being orthogonal to a , $b^* a = 0$, the line $a^* l = 1$ can be represented parametrically as:

(37a)

$$\underset{2 \times 1}{l} = \underset{2 \times 1}{l_o} + \underset{2 \times 1}{b} \underset{1 \times 1}{\alpha},$$

where l_o is a particular solution of (29). Hence, the condition $a^* l = 1$ in (31) may be replaced by $l = l_o + b\alpha$:

$$\min_l l^* Q_y l \quad \text{subject to} \quad l = l_o + b\alpha$$

This may also be written as:

$$\min_{\alpha} (l_o + b\alpha)^* Q_y (l_o + b\alpha).$$

The solution of this minimization problem is:

(37b)

$$\hat{\alpha} = -(b^* Q_y b)^{-1} b^* Q_y l_o.$$

This is easily verified by using an algebraic proof like the one given in section 1.2. Substitution of (37b) into (37a) gives:

$$\hat{l} = [I - b(b^* Q_y b)^{-1} b^* Q_y] l_o$$

The best linear unbiased estimator $\hat{x}$ of x follows therefore with (26) as:

$$(37c) \quad \hat{x} = l_o^* [I - Q_y b (b^* Q_y b)^{-1} b^*] y.$$

From equations (8) and (9) in section 1.3 we know that:

$$I - Q_y b (b^* Q_y b)^{-1} b^* = a (a^* Q_y^{-1} a)^{-1} a^* Q_y^{-1}.$$

Substitution in (37c) gives therefore:

$$\hat{x} = l_o^* a (a^* Q_y^{-1} a)^{-1} a^* Q_y^{-1} y$$

or:

$$\hat{x} = (a^* Q_y^{-1} a)^{-1} a^* Q_y^{-1} y,$$

since $l_o^* a = 1$. This concludes our alternative derivation of (37).

1.7 The Maximum Likelihood (ML) method

In the previous sections we have seen the two principles of *least-squares* estimation and *best linear unbiased* estimation at work. The least-squares principle is a deterministic principle where no probability assumptions concerning the vector of observations is made. In case $y = ax + e$, the least-squares principle is based on minimizing the quadratic form $(y - ax)^* W(y - ax)$, with the weight matrix W .

In contrast with the least-squares principle, the principle of best linear unbiased estimation does make use of some probability assumptions. In particular it is based on assumptions concerning the mean and covariance matrix of the random vector of observables y . However, only the mean and covariance matrix of y need to be specified, not its complete probability distribution. We will now present a method of estimation that in contrast with the BLUE's method needs the complete probability distribution of y .

Let $p(y;x)$ denote the probability density function of y . For a fixed value of x the function $p(y;x)$ is a function of y and for different values of x the function $p(y;x)$ may take different forms (see figure 1.21).

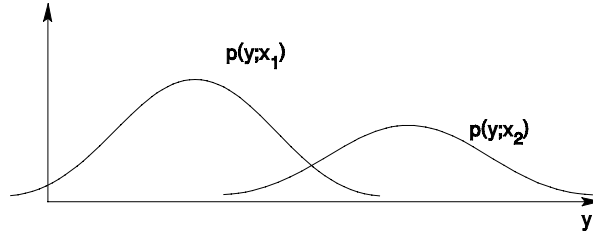


Figure 1.21: Two one-dimensional probability density functions $p(y;x_1)$, $p(y;x_2)$

Our objective is now to determine or estimate the parameter x on the basis of the observation vector y . That is, we wish to determine the correct value of x that produced the observed y . This suggests considering for each possible x how probable the observed y would be if x were the true value. The higher this probability, the more one is attracted to the explanation that the x in question produced y , and the more likely the value of x appears. Therefore, the expression $p(y;x)$ considered for fixed y as a function of x has been called the *likelihood* of x . Note that the likelihood of x is a probability density or a differential probability. The estimation principle is now to maximize the likelihood of x , given the observed value of y :

(38)

$$\max_x p(y;x)$$

Thus in the *method of maximum likelihood* we choose as an estimate of x the value which maximizes $p(y;x)$ for the given observed y . Note that in the ML-method one needs to know the complete mathematical expression for the probability density function $p(y;x)$ of $\underline{y}$. Thus it does not suffice, as is the case with the BLUE's method, to specify only the mean and covariance matrix of $\underline{y}$.

As an important example we consider the case that $p(y;x)$ is the probability density function of the normal distribution. Let us therefore assume that $\underline{y}$ is normally distributed with mean $E\{\underline{y}\} = \underline{a}x$ and covariance matrix $\underline{Q}_y$. The probability density function of $\underline{y}$ reads then in case it is a two-dimensional vector:

$$(39) \quad p(y;x) = (2\pi)^{-1} |\underline{Q}_y|^{-1/2} \exp\left[-\frac{1}{2} (y - \underline{a}x)^* \underline{Q}_y^{-1} (y - \underline{a}x)\right].$$

As it is often more convenient to work with $\ln p(y;x)$ then with $p(y;x)$ itself, the required value for x is found by solving:

(40)

$$\max_x \ln p(y;x).$$

this is permitted since $\ln p(y;x)$ and $p(y;x)$ have their maxima at the same values of x . Taking the logarithm of (39) gives:

$$(41) \quad \ln p(y;x) = -\ln(2\pi) - \frac{1}{2} \ln |Q_y| - \frac{1}{2} (y - ax)^* Q_y^{-1} (y - ax).$$

Since the first two terms on the right-hand side of (41) are constant, it follows that maximization of $\ln p(y;x)$ amounts to maximization of $-1/2(y-ax)^* Q_y^{-1}(y-ax)$. But this is equivalent to minimization of $(y-ax)^* Q_y^{-1}(y-ax)$. Hence we have obtained the important conclusion that maximum likelihood estimators, in case $p(y;x)$ is the normal distribution, are identical to least-squares estimators for which the weight matrix is $W=Q_y^{-1}$!

Let us now exemplify the above discussion geometrically. From (39) follows that the contours of constant probability density are ellipses:

$$p(y;x) = \text{constant} \Leftrightarrow (y-ax)^* Q_y^{-1} (y-ax) = \text{constant} = c.$$

For different values of c we get different ellipses. The *smaller* the probability density, the *larger* the value of c and the larger the ellipses are (see figure 1.22).

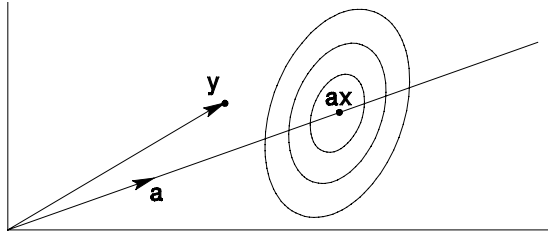


Figure 1.22: Contours of $p(y;x) = \text{constant}$

By varying the values of x the family of contour lines centred at ax translate along the line with direction vector a (see figure 1.23).

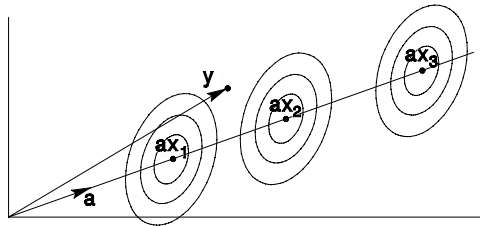
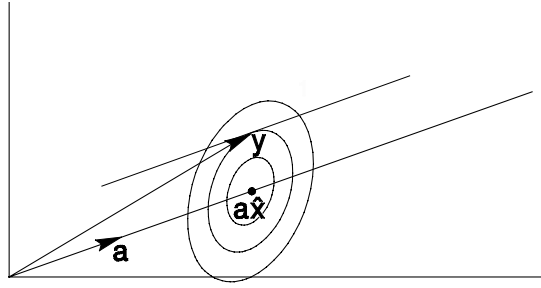
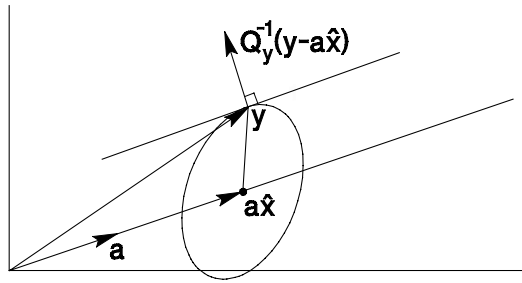


Figure 1.23

From the figure it is clear that the likelihood of x obtains its maximum there where the observation point y is a point of tangency of one of the ellipses (see figure 1.24).


 Figure 1.24: The ML-estimate $\hat{x}$

Hence, after translating the family of ellipses along the line with direction vector a such that y becomes a point of tangency, the corresponding centre $a\hat{x}$ of the ellipses gives the maximum likelihood estimate $\hat{x}$. At the point of tangency y , the normal to the ellipse is $Q_y^{-1}(y-a\hat{x})$ (see figure 1.25).


 Figure 1.25: The normal $Q_y^{-1}(y-a\hat{x})$

The normal $Q_y^{-1}(y-a\hat{x})$ is orthogonal to the tangent line at y . But this tangent line is parallel to the line with direction vector a . Thus we have $a^* Q_y^{-1}(y-a\hat{x})=0$, which shows once again that in case of the normal distribution, ML-estimators and LSQ-estimators with $W=Q_y^{-1}$ are identical!

1.8 Summary

In this chapter we have discussed three different methods of estimation: the *LSQ*-method, the *BLUE*-method and the *ML*-method. For an overview of the results see table 1. The LSQ-method, being a deterministic method, was applied to the model:

$$(42) \quad \underset{2 \times 1}{y} = \underset{2 \times 1}{a} \underset{1 \times 1}{x} + \underset{2 \times 1}{e} .$$

In case the measurement errors are considered to be random, $\underline{e}$, the observables also become random, $\underline{y}$, and (42) has to be written as:

$$(43) \quad \underline{y} = \underline{a}x + \underline{e}.$$

With the assumptions that $E\{\underline{e}\}=0$ and $E\{\underline{e}\underline{e}^*\}=\underline{Q}_y$, we get with (43) the model:

$$(44) \quad E\{\underline{y}\} = \underline{a}x ; E\{(\underline{y} - \underline{a}x)(\underline{y} - \underline{a}x)^*\} = \underline{Q}_y .$$

These two assumption were the starting point of the BLUE-method. If in addition to (44), it is also assumed that $\underline{y}$ is normally distributed then the model becomes:

$$(45) \quad \underline{y} \sim N(\underline{a}x, \underline{Q}_y)$$

It was shown that for this model the ML-estimator is identical to the BLU-estimator and the LSQ-estimator if $\underline{W}=\underline{Q}_y^{-1}$.

We have restricted ourselves in this introductory chapter to the case of two observations y_1, y_2 and one unknown x . In the next chapter we will start to generalize the theory to higher dimensions. That is, in the next chapter we will start developing the theory for m observations $y_1, y_2, \dots, y_m$ and n unknowns $x_1, x_2, \dots, x_n$. In that case model (45) generalizes to:

$$(46) \quad \underset{m \times 1}{\underline{y}} \sim N(\underset{m \times n}{\underline{A}} \underset{n \times 1}{x}, \underset{m \times m}{\underline{Q}_y}),$$

where $\underline{A}$ is an $m \times n$ matrix, i.e. it has m rows and n columns.

Least Squares Estimation (LSQ)	Best Linear Unbiased Estimation (BLUE)	Maximum Likelihood Estimation (ML)
<i>model:</i> $y = ax + e$ <i>estimation principle:</i> $\min_x (y - ax)^* W (y - ax)$ <i>solution:</i> $\hat{x} = (a^* W a)^{-1} a^* W y$ $\hat{y} = a \hat{x} = P_a y$ $\hat{e} = y - a \hat{x} = P_a^\perp y$ $\hat{e}^* W \hat{e} = y^* W P_a^\perp y$ $P_a = a(a^* W a)^{-1} a^* W$ $P_a^\perp = I - P_a$	<i>model:</i> $y = ax + e$ $E\{e\} = 0 ; E\{e e^*\} = Q_y$ or $E\{\hat{y}\} = ax ; E\{(y - ax)(y - ax)^*\} = Q_y$ <i>estimation principle:</i> $L: \hat{x} = l^* y ; U: E\{\hat{x}\} = x ; B: \sigma_{\hat{x}}^2 \min.$ <i>solution:</i> $\hat{x} = (a^* Q_y^{-1} a)^{-1} a^* Q_y^{-1} y ; \sigma_{\hat{x}}^2 = (a^* Q_y^{-1} a)^{-1}$ $\hat{y} = a \hat{x} = P_a y ; Q_{\hat{y}} = P_a Q_y$ $\hat{e} = y - a \hat{x} = P_a^\perp y ; Q_{\hat{e}} = P_a^\perp Q_y$ $\hat{e}^* Q_y^{-1} \hat{e} = y^* Q_y^{-1} P_a^\perp y ;$ $P_a = a(a^* Q_y^{-1} a)^{-1} a^* Q_y^{-1}$ $P_a^\perp = I - P_a$	<i>model:</i> $y \sim p(y; x)$ <i>estimation principle:</i> $\max_x p(y; x)$ In case $y \sim N(ax, Q_y)$ the solution is identical to BLUE and to LSQ if $W = Q_y^{-1}$

Table 1.1

2 The model with observation equations

2.1 Introduction

As an introduction consider the following example: Figure 2.1 shows a *levelling network*. The points 0,1,2,3,4 and 5 are connected with levelled height differences. The total number of levelled height differences is 9. The height of point 0 is assumed to be known in the *NAP-reference system* (NAP= Normaal Amsterdams Peil). The heights of the remaining points are unknown. Our objective is to determine their heights from the measured height differences.

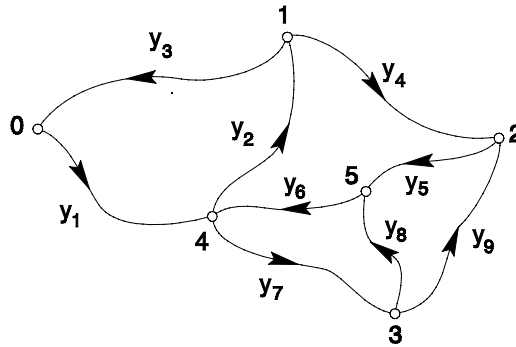


Figure 2.1: Levelling network

One way to determine their heights is to use the configuration of figure 2.2. One starts at the reference point 0 of which the height is known, and then determines the heights of the points 4,5,3,2 and 1 as shown in figure 2.2. In this case only 5 out of the 9 available levelled height differences are used, namely y_1 , y_6 , y_8 , y_9 and y_4 .

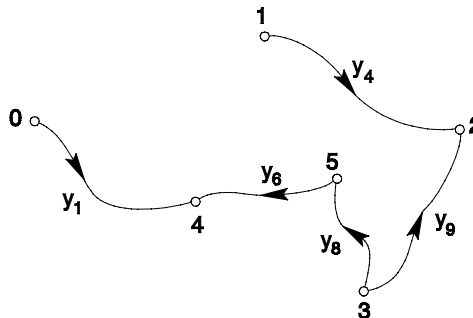


Figure 2.2

An alternative way to determine the heights of the points is to use the configuration of figure 2.3. One starts again at the reference point 0, and then determines the heights of the points 1,2,5,3 and 4 as shown in figure 2.3. In this case again 5 out of the 9 available levelled height differences are used. Note however that the height differences used differ from the ones used in the previous case.

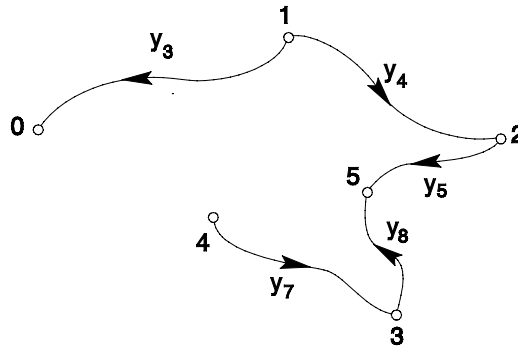


Figure 2.3

It will be clear that the two solutions of figure 2.2 and figure 2.3 are identical in case the height differences were *perfectly* levelled. However, perfect measurements do not exist! Hence, since one relies in the two cases considered on different observations with different measurement errors, the computed heights will also differ. Therefore, if one wants to make use of all the 9 levelled height differences, one has to take care in one way or another of the resulting inconsistencies or discrepancies.

One could also argue that in the case of figure 2.1 too many height differences are measured. This is true as far as the construction of the levelling network is concerned. That is, strictly speaking only 5 out of the 9 available height differences are needed to determine the 5 unknown heights. However, with 5 measured height differences one will *never* be able to detect *blunders or outliers* in the measurements. This is one of the reasons why always more measurements are taken than strictly necessary for determining the geometric figure of the network. In the course of your studies in *Mathematical Geodesy* you will learn to answer questions as: how many measurements are needed, how should they be distributed over the network etc. That is, in the course of your studies you will learn, given the purpose of a network, how to *design* a network with a certain *quality*.

In the above example of figure 2.1 we have 9 measured height differences and 5 unknown heights. As we did in the introductory section of chapter 1, we can write down the 9 equations that relate the measurements to the unknowns. If we denote the unknown heights as $x_1, x_2, \dots, x_5$ and the known height of point 0 as x_0 , we obtain with the help of figure 2.1:

$$\begin{aligned}
 y_1 &= x_4 - x_0 + e_1 \\
 y_2 &= x_1 - x_4 + e_2 \\
 y_3 &= x_0 - x_1 + e_3 \\
 y_4 &= x_2 - x_1 + e_4 \\
 y_5 &= x_5 - x_2 + e_5 \\
 y_6 &= x_4 - x_5 + e_6 \\
 y_7 &= x_3 - x_4 + e_7 \\
 y_8 &= x_5 - x_3 + e_8 \\
 y_9 &= x_2 - x_3 + e_9
 \end{aligned}
 \tag{1}$$

where the e_i , $i=1,\dots,9$, are the unknown random measurements errors. The above 9 equations can be written in matrix notation as:

$$\begin{pmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \\ y_5 \\ y_6 \\ y_7 \\ y_8 \\ y_9 \end{pmatrix} = \begin{pmatrix} & & & 1 & \\ & 1 & & -1 & \\ & -1 & & & \\ & -1 & 1 & & \\ & & -1 & & 1 \\ & & & 1 & -1 \\ & & 1 & -1 & \\ & & -1 & & 1 \\ 1 & -1 & & & \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ x_5 \end{pmatrix} + \begin{pmatrix} -x_0 \\ 0 \\ x_0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{pmatrix} + \begin{pmatrix} e_1 \\ e_2 \\ e_3 \\ e_4 \\ e_5 \\ e_6 \\ e_7 \\ e_8 \\ e_9 \end{pmatrix}
 \tag{2}$$

For simplicity we will assume that the known height x_0 of the reference point 0 is equal to zero. The matrix equation of (2) is then of the form:

$$\begin{matrix} (3) \\ \boxed{y = A \quad x + e,} \\ \begin{matrix} m \times 1 & m \times n & n \times 1 & m \times 1 \end{matrix} \end{matrix}$$

with $m=9$ and $n=5$. Matrix A thus has 9 rows and 5 columns. If the known height x_0 is unequal to zero, then (2) can still be brought into the form of (3) by bringing the known vector containing x_0 in (2) to the left-hand side.

Equation (3) is the multi dimensional generalization of equation (19) in section 1.5. In this chapter we will generalize the theory of chapter 1 for multi dimensional cases.

2.2 Least-squares estimation

The equations of the linear system:

$$(4) \quad \underset{m \times 1}{y} = \underset{m \times n}{A} \underset{n \times 1}{x} + \underset{m \times 1}{e},$$

are known as *observation equations* (waarnemingsvergelijkingen). Let us assume that an observation vector y is given, and that we want to solve for the unknown vectors x and e in the system:

$$(5) \quad \underset{m \times 1}{y} = \underset{m \times n}{A} \underset{n \times 1}{x} + \underset{m \times 1}{e}.$$

If we denote the column vectors of the matrix A by a_i , $i=1, \dots, n$, then $A=(a_1, a_2, \dots, a_n)$, and (5) may be written as:

$$(6) \quad \underset{m \times 1}{y} = \sum_{i=1}^n \underset{m \times 1}{a_i} \underset{m \times 1}{x_i} + \underset{m \times 1}{e}$$

In chapter 1 we solved equation (6) in a least-squares sense for the case $n=1$. We will now consider the general case that n may be arbitrary, provided that $n \leq m$. That is, the restriction is that the number of observations, m , has to be larger than or equal to the number of unknown parameters, n . For $m=3$ and $n=2$ we may visualize equation (6) as in figure 2.4.

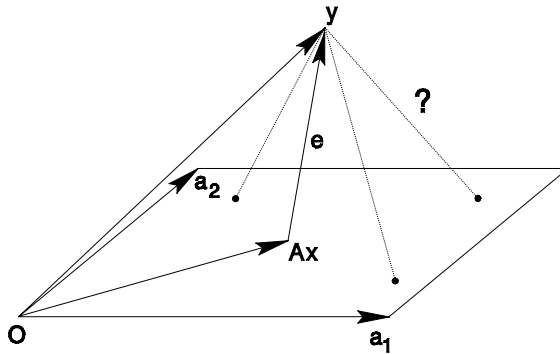


Figure 2.4: x and e unknown in $y=Ax+e$.

From the geometry of the figure it seems intuitively appealing to estimate x as $\hat{x}$, such that $\sum_{i=1}^n a_i \hat{x}_i$ or $A\hat{x}$ is as close as possible to the given observation vector y . This is the case if $\hat{e} = y - \sum_{i=1}^n a_i \hat{x}_i$ is *orthogonal* to the "plane" spanned by the column vectors a_i , $i=1, \dots, n$, of A .

That is, if:

$$a_j^* (y - \sum_{i=1}^n a_i \hat{x}_i) = 0, \quad j=1, \dots, n$$

or if:

$$\begin{aligned} a_1^* (y - \sum_{i=1}^n a_i \hat{x}_i) &= 0 \\ a_2^* (y - \sum_{i=1}^n a_i \hat{x}_i) &= 0 \\ &\vdots \\ a_n^* (y - \sum_{i=1}^n a_i \hat{x}_i) &= 0 \end{aligned}$$

or if:

$$A^* (y - A\hat{x}) = 0$$

or if:

$$(7) \quad \underset{n \times n}{A^* A} \underset{n \times 1}{\hat{x}} = \underset{n \times m}{A^*} \underset{m \times 1}{y}$$

This is a system of n linear equations with n unknowns. If the matrix A^*A is invertible, the solution of (7) is unique and reads:

$$(8) \quad \hat{x} = (A^*A)^{-1} A^* y.$$

Compare this result with equation (2) in section 1.2. The $n \times n$ matrix A^*A is invertible if and only if $\text{rank } A^*A = n$. And since $\text{rank } A^*A = \text{rank } A$ (for a proof see the end of this section), it follows that A^*A is invertible if $\text{rank } A = n$. We will therefore assume that matrix A has *full rank* n . Verify that the matrix of equation (2) has full rank; its rank is 5. Verify also that if the height of the reference point 0 is unknown, the matrix which follows from including x_0 in the vector of unknowns has a *rankdefect* of 1. The reason is of course that "absolute heights" cannot be determined from height differences only. The derivation of the estimate $\hat{x}$ of (8) was based on the principle of orthogonality. The estimate $\hat{x}$ however also minimizes the sum of squares $(y - Ax)^*(y - Ax)$. That is, $\hat{x}$ of (8) is the solution of:

$$(9) \quad \min_x (y - Ax)^*(y - Ax)$$

This is easily shown, by using the same algebraic derivation as was used in section 1.2. Along the same lines one can show that the solution of:

$$(10) \quad \boxed{\min_x (y - Ax)^* W (y - Ax)} .$$

where W is an $m \times m$ symmetric and positive definite weight matrix, is:

$$(11) \quad \hat{x} = (A^* W A)^{-1} A^* W y.$$

Proof: (algebra)

$$\begin{aligned}(y-Ax)^*W(y-Ax) &= [y-A\hat{x}-A(x-\hat{x})]^*W[y-Ax-A(x-\hat{x})] \\ &= (y-A\hat{x})^*W(y-A\hat{x}) - 2(y-A\hat{x})^*WA(x-\hat{x}) + (x-\hat{x})^*A^*WA(x-\hat{x}).\end{aligned}$$

$$\text{But } A^*W(y-A\hat{x}) = A^*W[(y-A(A^*WA)^{-1}A^*W)y] = 0, \text{ hence,}$$

$$\begin{aligned}(y-Ax)^*W(y-Ax) &= (y-A\hat{x})^*W(y-A\hat{x}) + (x-\hat{x})^*A^*WA(x-\hat{x}) \\ &= \text{minimal for } x=\hat{x}\end{aligned}$$

End of proof.

The estimate $\hat{x}$ of (11) is the *weighted least-squares estimate* of x . As a generalization of equation (7) in section 1.3, the weighted least-squares estimators read therefore for the multi dimensional case:

$$(12) \quad \begin{aligned}\hat{\underline{x}} &= (\underline{A}^* \underline{W} \underline{A})^{-1} \underline{A}^* \underline{W} \underline{y} \\ \hat{\underline{y}} &= \underline{A} \hat{\underline{x}} = \underline{A} (\underline{A}^* \underline{W} \underline{A})^{-1} \underline{A}^* \underline{W} \underline{y} \\ \hat{\underline{e}} &= \underline{y} - \hat{\underline{y}} = [\underline{I} - \underline{A} (\underline{A}^* \underline{W} \underline{A})^{-1} \underline{A}^* \underline{W}] \underline{y} \\ \hat{\underline{e}}^* \underline{W} \hat{\underline{e}} &= \underline{y}^* [\underline{W} - \underline{W} \underline{A} (\underline{A}^* \underline{W} \underline{A})^{-1} \underline{A}^* \underline{W}] \underline{y}\end{aligned}$$

Intermezzo:

$$(13) \quad \text{rank } \underline{A} = \text{rank } \underline{A}^* \underline{A}$$

From theory in Linear Algebra we know that:

$$\text{rank } \underline{A} = n - \dim.N(\underline{A}),$$

where $N(\underline{A})$ denotes the *nullspace* (kern) of the $m \times n$ matrix $\underline{A}$. For $\underline{A}^* \underline{A}$ we have:

$$\text{rank } \underline{A}^* \underline{A} = n - \dim.N(\underline{A}^* \underline{A}).$$

Hence, $\text{rank } \underline{A} = \text{rank } \underline{A}^* \underline{A}$ is equivalent to:

$$(14) \quad \dim.N(\underline{A}) = \dim.N(\underline{A}^* \underline{A})$$

The nullspace of $\underline{A}$ is defined as the set of vectors x that satisfy:

$$\underline{A}x = 0.$$

Since $\underline{A}^* \underline{A}x = 0$ if $\underline{A}x = 0$, it follows that:

$$(15) \quad N(\underline{A}) \subset N(\underline{A}^* \underline{A})$$

But also:

$$(16) \quad N(A^*A) \subset N(A)$$

holds. Because if $x \in N(A^*A)$ then $A^*Ax=0$, and this implies that $Ax=0$ or $x \in N(A)$. $A^*Ax=0$ can never imply that $Ax \neq 0$. From (15) and (16) follows that $N(A)=N(A^*A)$ and thus that $\dim. N(A) = \dim N(A^*A)$ or $\text{rank } A = \text{rank } A^*A$.

End of proof.

2.3 The mean and covariance matrix of least-squares estimators

In section 5 of chapter 1 we derived the means and covariance matrices of the least-squares estimators of the model $\underline{y}=ax+\underline{e}$. This was done through application of the propagation laws for means and variances to (17) in section 1.5. In the present section we want to generalize these results. That is, we want to derive the mean and covariance matrix of the estimators of (12). This is again done through application of the propagation laws for means and variances. For completeness sake we first state and prove three laws for the general multi-dimensional case.

Let $\underline{y}$ be a $p \times 1$ dimensional random vector and let $\underline{u}$ be a $q \times 1$ dimensional random vector. Furthermore let $\underline{u}$ and $\underline{y}$ be functionally related by:

$$(17) \quad \underline{y} = F(\underline{u}),$$

where $F(\cdot)$ is a vector function.

The mean of $\underline{y}$, $E\{\underline{y}\}$, is then *defined* as (see also appendices E and F):

$$(18) \quad E\{\underline{y}\} = \int F(u) \underline{p}_u(u) du.$$

where $\underline{p}_u(u)$ is the probability density function of $\underline{u}$. Equation (18) shows that the mean of $\underline{y}$ can be computed once the vector function $F(\cdot)$ and probability density function $\underline{p}_u(\cdot)$ are known. We will now use (18) to derive the two propagation laws of means and variances for the case of *linear* functions.

The propagation law of means

Let the two random vectors $\underline{u}$ and $\underline{y}$ be related as:

$$(19) \quad \begin{array}{c} \underline{y} = M \underline{u} + m \\ p \times 1 \quad p \times q \quad q \times 1 \quad p \times 1 \end{array}$$

where M is a constant $p \times q$ matrix and m is a constant $p \times 1$ vector. Using (18) we get for the mean of $\underline{v}$:

$$\begin{aligned} E\{\underline{v}\} &= \int (Mu + m) \underline{p}_{\underline{u}}(u) d\mathbf{u} = \int Mu \underline{p}_{\underline{u}}(u) d\mathbf{u} + \int m \underline{p}_{\underline{u}}(u) d\mathbf{u} \\ (20) \quad &= M \int u \underline{p}_{\underline{u}}(u) d\mathbf{u} + m \int \underline{p}_{\underline{u}}(u) d\mathbf{u} \end{aligned}$$

According to (18), $\int u \underline{p}_{\underline{u}}(u) d\mathbf{u}$ equals the mean of $\underline{u}$, $E\{\underline{u}\}$. Furthermore, since the "area" or "volume" of a probability density function equals 1, we have that $\int \underline{p}_{\underline{u}}(u) d\mathbf{u} = 1$. Hence, (20) reduces to:

$$(21) \quad \boxed{\begin{matrix} E\{\underline{v}\} = M E\{\underline{u}\} + m \\ p \times 1 \quad p \times q \quad q \times 1 \quad p \times 1 \end{matrix}}$$

This result shows that in case $\underline{u}$ and $\underline{v}$ are *linearly* related, the mean of $\underline{v}$, $E\{\underline{v}\}$, can be computed once the mean of $\underline{u}$, $E\{\underline{u}\}$, is known! This is generally *not* the case for nonlinear functions.

The propagation law of variances and covariances

The covariance matrix of the random vector $\underline{v}$ is *defined* as:

$$\underline{Q}_v = E\{(\underline{v} - E\{\underline{v}\})(\underline{v} - E\{\underline{v}\})^*\}.$$

If $\underline{u}$ and $\underline{v}$ are related as (19), it follows with (21) that:

$$\begin{aligned} \underline{Q}_v &= E\{(\underline{v} - E\{\underline{v}\})(\underline{v} - E\{\underline{v}\})^*\} = E\{(M\underline{u} + m - ME\{\underline{u}\} - m)(M\underline{u} + m - ME\{\underline{u}\} - m)^*\} \\ &= E\{[M(\underline{u} - E\{\underline{u}\})][M(\underline{u} - E\{\underline{u}\})]^*\} \\ &= E\{M(\underline{u} - E\{\underline{u}\})(\underline{u} - E\{\underline{u}\})^* M^*\} \end{aligned}$$

Using (18) this gives:

$$\begin{aligned} \underline{Q}_v &= \int M(\underline{u} - E\{\underline{u}\})(\underline{u} - E\{\underline{u}\})^* M^* \underline{p}_{\underline{u}}(u) d\mathbf{u} \\ (22) \quad &= M \int (\underline{u} - E\{\underline{u}\})(\underline{u} - E\{\underline{u}\})^* \underline{p}_{\underline{u}}(u) d\mathbf{u} M^* \\ &= ME\{(\underline{u} - E\{\underline{u}\})(\underline{u} - E\{\underline{u}\})^*\} M^* \end{aligned}$$

But $E\{(\underline{u} - E\{\underline{u}\})(\underline{u} - E\{\underline{u}\})^*\}$ is by definition the covariance matrix of $\underline{u}$, $\underline{Q}_u$. Hence, (22) reduces to:

(23)

$$\begin{array}{c} \mathbf{Q}_v = \mathbf{M} \mathbf{Q}_u \mathbf{M}^* \\ p \times p \quad p \times q \quad q \times q \quad q \times p \end{array}$$

This result shows that in case $\underline{u}$ and $\underline{v}$ are *linearly* related, the covariance matrix of $\underline{v}$, $\mathbf{Q}_v$, can be computed once the covariance matrix of $\underline{u}$, $\mathbf{Q}_u$, is known! Thus in case of linear functions one does not need to know the complete probability density function of $\underline{u}$.

If we denote the covariance between $\underline{u}$ and $\underline{v}$ as $\mathbf{Q}_{uv}$, then by *definition*:

$$\mathbf{Q}_{uv} = E\{(\underline{u} - E\{\underline{u}\})(\underline{v} - E\{\underline{v}\})^*\}.$$

With a derivation similar to the one given above follows then that:

(24)

$$\begin{array}{c} \mathbf{Q}_{uv} = \mathbf{Q}_u \mathbf{M}^* \\ q \times p \quad q \times q \quad q \times p \end{array}$$

Prove this yourself!

It will now be clear since the first three estimators of (12) are all *linear* functions of $\underline{y}$, and thus of the form (19), that their means and covariance matrices can be derived from the results of (21), (23) and (24).

The model used for deriving the estimators of (12) was:

(25)

$$\begin{array}{c} \underline{y} = \mathbf{A} \ x + \underline{e} \\ m \times 1 \quad m \times n \ n \times 1 \quad m \times 1 \end{array}$$

We will now assume that $E\{\underline{e}\}=0$ and $E\{\underline{e} \underline{e}^*\}=\mathbf{Q}_y$. With (25) this gives:

(26)

$$E\{\underline{y}\} = \mathbf{A}x ; E\{(\underline{y} - \mathbf{A}x)(\underline{y} - \mathbf{A}x)^*\} = \mathbf{Q}_y$$

With the first equation of (26), application of the *propagation law for means*, (21), to (12) gives:

(27)

$$\begin{array}{c} E\{\hat{x}\} = x \\ E\{\hat{y}\} = \mathbf{A}x \\ E\{\hat{e}\} = 0 \end{array}$$

Compare this result with equation (21) in section 1.5. This important result shows that if (26) holds, the least-squares estimators are *unbiased* estimators. And this property of unbiasedness is independent of the choice for W .

With the second equation of (26), application of the *propagation law for variances*, (23), to (12) gives:

$$(28) \quad \begin{aligned} Q_{\hat{x}} &= (A^*WA)^{-1}A^*WQ_yWA(A^*WA)^{-1} \\ Q_{\hat{y}} &= AQ_xA^* \\ Q_{\hat{e}} &= [I - A(A^*WA)^{-1}A^*W]Q_y[I - WA(A^*WA)^{-1}A^*] \end{aligned}$$

In a similar way, application of the *propagation law of covariances*, (24), to (12) gives:

$$(29) \quad \begin{aligned} Q_{\hat{x}\hat{y}} &= (A^*WA)^{-1}A^*WQ_yWA(A^*WA)^{-1}A^* = Q_{\hat{x}}A^* \\ Q_{\hat{x}\hat{e}} &= (A^*WA)^{-1}A^*WQ_y - (A^*WA)^{-1}A^*WQ_yWA(A^*WA)^{-1}A^* \\ Q_{\hat{y}\hat{e}} &= A(A^*WA)^{-1}A^*WQ_y - A(A^*WA)^{-1}A^*WQ_yWA(A^*WA)^{-1}A^* \end{aligned}$$

In order to derive the mean of the random variable $\hat{e}^*W\hat{e}$ of (12), we cannot make use of (21) directly. The random variable $\hat{e}^*W\hat{e}$ is namely *not* a linear function of $\underline{y}$. It is possible however to write $\hat{e}^*W\hat{e}$ as a linear combination of products of the elements of $\hat{e}$, in which case (21) can be applied again. If we denote the i th-element of $\hat{e}$ by $\hat{e}_i$, and the (i,j) th-element of the matrix W by $(W)_{ij}$, then:

$$\hat{e}^*W\hat{e} = \sum_{i=1}^m \sum_{j=1}^m (W)_{ij} \hat{e}_i \hat{e}_j$$

This shows that $\hat{e}^*W\hat{e}$ is a *linear* combination of the products $\hat{e}_i\hat{e}_j$. Hence:

$$E\{\hat{e}^*W\hat{e}\} = \sum_{i=1}^m \sum_{j=1}^m (W)_{ij} E\{\hat{e}_i \hat{e}_j\},$$

where $E\{\hat{e}_i \hat{e}_j\}$ is the covariance between $\hat{e}_i$ and $\hat{e}_j$. But the covariance between $\hat{e}_i$ and $\hat{e}_j$ is the (i,j) th-element of the covariance matrix $Q_{\hat{e}}$. Thus:

$$E\{\hat{e}^*W\hat{e}\} = \sum_{i=1}^m \sum_{j=1}^m (W)_{ij} (Q_{\hat{e}})_{ij}$$

But since the covariance matrix $Q_{\hat{e}}$ is symmetric and thus $(Q_{\hat{e}})_{ij} = (Q_{\hat{e}})_{ji}$, we may write:

$$E\{\hat{e}^*W\hat{e}\} = \sum_{i=1}^m \sum_{j=1}^m (W)_{ij} (Q_{\hat{e}})_{ji}$$

For fixed i , the term $\sum_{j=1}^m (W)_{ij} (Q_{\hat{e}})_{ji}$ equals the product of the i th-row of W with the i th-column of $Q_{\hat{e}}$:

$$\sum_{j=1}^m (W)_{ij} (Q_{\hat{e}})_{ji} = \begin{array}{|c|} \hline \text{ } \\ \hline \end{array} \cdot \begin{array}{|c|} \hline \text{ } \\ \hline \end{array}$$

Thus $E\{\underline{\hat{e}}^* W \underline{\hat{e}}\}$ is the sum of these m products:

$$E\{\underline{\hat{e}}^* W \underline{\hat{e}}\} = \begin{array}{c} 1 \\ \hline \square \end{array} \cdot \begin{array}{c} 1 \\ | \\ \square \end{array} + \begin{array}{c} 2 \\ \hline \square \end{array} \cdot \begin{array}{c} 2 \\ | \\ \square \end{array} + \dots + \begin{array}{c} m \\ \hline \square \end{array} \cdot \begin{array}{c} m \\ | \\ \square \end{array}$$

This shows that $E\{\underline{\hat{e}}^* W \underline{\hat{e}}\}$ equals the *trace* (spoor) of the matrix product $W Q_{\hat{e}}$:

(30)

$$E\{\underline{\hat{e}}^* W \underline{\hat{e}}\} = \text{trace } (W Q_{\hat{e}})$$

For a definition of the trace of a square matrix consult your lecture notes on Linear Algebra.

It should be noted that all the results concerning the means, variances and covariances of the least-squares estimators, were derived without making any assumptions on the probability distribution of $\underline{y}$! This is in principle also possible for the variance of $\underline{\hat{e}}^* W \underline{\hat{e}}$. The derivation is however rather complicated and therefore omitted. Even for the case that $\underline{y}$ is normally distributed, the derivation of the variance of $\underline{\hat{e}}^* W \underline{\hat{e}}$ is complex. We therefore state without proof that if $\underline{y}$ is *normally* distributed with mean and covariance matrix as given in (26), the variance of $\underline{\hat{e}}^* W \underline{\hat{e}}$ reads:

(31)

$$\sigma_{\hat{e}^* W \hat{e}}^2 = 2 \text{ trace } [(W - W A (A^* W A)^{-1} A^* W) Q_y (W - W A (A^* W A)^{-1} A^* W) Q_y]$$

2.4 Best Linear Unbiased Estimation

We consider again the model:

(32)

$$\begin{array}{c} E\{\underline{y}\} = \underline{A} \quad \underline{x} \quad ; \quad E\{(\underline{y} - \underline{A}\underline{x})(\underline{y} - \underline{A}\underline{x})^*\} = \underline{Q}_y \\ m \times 1 \quad m \times n \quad n \times 1 \quad \quad \quad m \times m \quad \quad \quad m \times m \end{array}$$

Let us assume that we are interested in estimating a parameter θ , which is a linear function of $\underline{x}$:

(33)

$$\begin{array}{c} \underline{\theta} = \underline{f}^* \underline{x} \\ 1 \times 1 \quad 1 \times n \quad n \times 1 \end{array}$$

For instance, in case of the levelling network of figure 2.1 we could be interested in estimating the height difference between the points 1 and 5. Then θ of (33) plays the role of the height difference between the points 1 and 5, and the vector $\underline{f}$ of (33) takes the form:

$$\underline{f} = (-1 \ 0 \ 0 \ 0 \ 1)^*$$

Let $\underline{\hat{\theta}}$ be the estimator of θ . Then according to the criteria of best linear unbiased estimation, we want the estimator $\underline{\hat{\theta}}$ to be a linear function of $\underline{y}$:

$$(34) \quad \underset{1 \times 1}{\hat{\theta}} = \underset{1 \times m}{l^*} \underset{m \times 1}{y} \quad \text{"L-property"}$$

such that

$$(35) \quad E\{\hat{\theta}\} = \theta, \quad \text{"U-property"}$$

and:

$$(36) \quad \sigma_{\hat{\theta}}^2 \text{ is minimal} \quad \text{"B-property"}$$

Our objective is now to find a vector l such that with (34), the conditions (35) and (36) are fulfilled. Substitution of (34) into (35) gives:

$$E\{\hat{\theta}\} = l^* E\{y\} = \theta,$$

and with (32) and (33) this gives:

$$l^* A x = f^* x$$

Since we want this condition to hold for all x , it follows that:

$$(37)$$

$$\underset{n \times 1}{f} = \underset{n \times m}{A^*} \underset{m \times 1}{l}$$

Compare this with equation (29) in section 1.6. Equation (37) is a system of n linear equations with m unknowns. The unknowns are the elements of the vector l . Every vector l that satisfies (37) gives with (34) a linear unbiased estimator of θ . If C is an $m \times n$ matrix such that $(A^* C)^{-1}$ exists, then:

$$(38) \quad l = C(A^* C)^{-1} f$$

is a solution of (37). This means that:

$$(39) \quad \hat{\theta} = f^* (C^* A)^{-1} C^* y$$

is an unbiased estimator of θ . That this is indeed the case, follows if we take the expectation of:

$$E\{\hat{\theta}\} = f^* (C^* A)^{-1} C^* E\{y\} = f^* (C^* A)^{-1} C^* A x = f^* x = \theta.$$

The $m \times n$ matrix C of (38) is not unique. One choice for matrix C would be:

$$C = W A$$

Substitution into (39) gives:

$$\hat{\theta} = f^*(A^*WA)^{-1}A^*Wy$$

This shows once again that the weighted least-squares estimator is an unbiased estimator. From your lecture notes in Linear Algebra (Strang, 1988) you know that the general solution of the *inhomogeneous* system (37) is given by the sum of a *particular solution* of (37) and the solution of the *homogeneous* system:

$$(40) \quad \underset{n \times 1}{0} = \underset{n \times m}{A^*} \underset{m \times 1}{l}$$

The solution of this homogeneous system is the *nullspace*, $N(A^*)$, of the matrix A^* :

$$N(A^*) = \{l \in \mathbb{R}^m \mid A^*l = 0\}$$

The dimension of this space follows from the dimension theorem (Strang, 1988) as:

$$\dim. N(A^*) = m - \dim. R(A^*),$$

where $R(A^*)$ is the *range space* or column space (beeldruimte of kolomruimte) of A^* . Since $\dim. R(A^*) = \text{rank } A^* = \text{rank } A$ and since we assumed that $\text{rank } A = n$, it follows that:

$$\dim. N(A^*) = m - n.$$

Since the dimension of the solution space of (40) equals $m-n$, there exist $m-n$ linearly independent vectors b_i , $i=1, \dots, m-n$, that form a basis of $N(A^*)$. Each of these vectors satisfy:

$$(41) \quad 0 = A^*b_i, \quad i = 1, \dots, m-n.$$

If we collect the $m-n$ vectors b_i in a matrix B :

$$\underset{m \times (m-n)}{B} = (b_1, b_2, \dots, b_{m-n}),$$

then matrix B is of order $m \times (m-n)$ with full rank, $\text{rank } B = m-n$, and satisfies because of (41):

(42)

$$\underset{n \times (m-n)}{0} = \underset{n \times m}{A^*} \underset{m \times (m-n)}{B}$$

This shows that every solution of (40) can be represented as:

$$\underset{m \times 1}{l} = \underset{m \times (m-n)(m-n) \times 1}{B\alpha}$$

Hence, if l_o is a particular solution of (37), the general solution of (37) takes the form:

$$(43) \quad l = l_o + B\alpha$$

Compare this result with equation (37a) in section 1.6. Equation (43) shows that all vectors l that satisfy (37) lie on a *linear variety or linear manifold* (lineaire variëteit) parallel to the range space of matrix B .

Let us now continue our derivation of the best linear unbiased estimator of (33). With (37) or (43) we have found the class of linear unbiased estimators. We now want to find within this class, the estimator that has minimum variance. Application of the propagation law of variances to (34) gives with (32):

$$\sigma_{\hat{\theta}}^2 = l^* Q_y l.$$

Hence, our minimization problem reads:

(44)

$$\min_l l^* Q_y l \text{ subject to } f = A^* l$$

Compare this with equation (31) in section 1.6. Since (43) is the solution of (37), (44) may also be written as:

$$\min_l l^* Q_y l \text{ subject to } l = l_o + B\alpha$$

or as:

(45)

$$\min_{\alpha} (l_o + B\alpha)^* Q_y (l_o + B\alpha).$$

The solution of this minimization problem is:

(46)

$$\hat{\alpha} = -(B^* Q_y B)^{-1} B^* Q_y l_o.$$

This is easily verified by using an algebraic proof like the one given in section 1.2. Substitution of (46) into (43) gives:

(47)

$$\hat{l} = [I - B(B^* Q_y B)^{-1} B^* Q_y] l_o.$$

The best linear unbiased estimator $\hat{\theta}$ of θ follows therefore with (34) as:

(48)

$$\hat{\theta} = l_o^* [I - Q_y B(B^* Q_y B)^{-1} B^*] y.$$

This solution is expressed in terms of the matrix B , which is related to matrix A of (32) through (42). Since our starting point was the model of (32), it would be nice if we could express the solution (48) in terms of the matrix A . In chapter 1 (see section 1.6) this was possible by making use of the identity:

$$I - Q_y b(b^* Q_y b)^{-1} b^* = a(a^* Q_y^{-1} a)^{-1} a^* Q_y^{-1}$$

Compare with equation (8) and (9) in section 1.3. For the multi-dimensional case, this *suggests* the identity:

(49)

$$I - Q_y B (B^* Q_y B)^{-1} B^* = A (A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1}$$

We will prove this identity, by proving that:

(50)

$$\underset{m \times m}{I} = \underset{m \times m}{A (A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1}} + \underset{m \times m}{Q_y B (B^* Q_y B)^{-1} B^*}$$

Since $Ix=x$, $\forall x$, we have proved (50) once we can show that:

(51)

$$[A (A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1} + Q_y B (B^* Q_y B)^{-1} B^*] x = x \quad \forall x \in \mathbb{R}^m$$

If R is a square and *invertible* matrix, then (51) is equivalent to:

(52)

$$[A (A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1} + Q_y B (B^* Q_y B)^{-1} B^*] R z = R z \quad \forall z \in \mathbb{R}^m$$

If we take for R the $m \times m$ partitioned matrix:

(53)

$$R = \begin{pmatrix} A & Q_y B \end{pmatrix}, \quad \begin{matrix} m \times m & m \times n & m \times (m-n) \end{matrix}$$

then:

(54)

$$\begin{aligned} [A (A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1} + Q_y B (B^* Q_y B)^{-1} B^*] R &= [A (A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1} + Q_y B (B^* Q_y B)^{-1} B^*] (A : Q_y B) \\ &= A (A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1} (A : Q_y B) + Q_y B (B^* Q_y B)^{-1} B^* (A : Q_y B) \\ &= (A : 0) + (0 : Q_y B) = (A : Q_y B) = R, \end{aligned}$$

where we have made use of the fact that $A^* B = 0$, see (42). Equation (54) shows that with the choice (53) for R , the equation (52) holds. What remains to be shown is that matrix R of (53) is invertible. This square matrix is invertible if it has full rank m , and it has full rank only if the solution of the system:

(55)

$$\begin{pmatrix} A & Q_y B \end{pmatrix} \begin{pmatrix} \alpha \\ \beta \end{pmatrix} = 0, \quad \begin{matrix} m \times n & m \times (m-n) & m \times 1 \end{matrix}$$

is the *trivial* solution $(\alpha^* \beta^*)^* = 0$.

Since $B^* A = 0$, premultiplication of (55) with B^* gives:

(56)

$$\begin{matrix} B^* Q_y B & \beta & = & 0 \\ (m-n) \times (m-n) & (m-n) \times 1 & & (m-n) \times 1 \end{matrix}$$

But matrix B has full rank $m-n$ and therefore $\beta=0$ is the only solution of (56). With $\beta=0$ equation (55) reduces to:

$$(57) \quad \underset{m \times n}{A} \underset{n \times 1}{\alpha} = \underset{m \times 1}{0}$$

But matrix A has full rank n and therefore $\alpha=0$ is the only solution of (57). Hence $(\alpha^* \beta^*)^*=0$ is the only solution of (55), which shows that $(A : Q_y B)$ is invertible. This concludes the proof of the identity (49).

Substitution of (49) into (48) gives:

$$\hat{\theta} = l_o^* A (A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1} y$$

or since l_o is a particular solution of (37):

$$(58) \quad \hat{\theta} = f^* (A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1} y.$$

This important result tells us two things. First of all it tells us that the best linear unbiased estimator of x is given by:

$$(59) \quad \boxed{\hat{x} = (A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1} y}$$

This follows from (58) if we take for the vector f respectively $(1,0,\dots,0)^*$, $(0,1,0,\dots,0)^*$, ..., $(0,0,\dots,1)^*$. Compare (59) with equation (37) in section 1.6. Equation (59) shows that the BLUE of x is identical to the weighted least-squares estimator of x if $W=Q_y^{-1}$. The result (58) tells us also that the BLUE of $\theta=f^*x$ is:

$$(60) \quad \boxed{\hat{\theta} = f^* \hat{x}}$$

where $\hat{x}$ is the BLUE of x .

Now that we know that the best linear unbiased estimators are identical to the least-squares estimators if $W=Q_y^{-1}$, we can use the results of the previous section 2.3 to obtain the covariance matrices of the BLUE's. Substitution of $W=Q_y^{-1}$ into equation (28) gives:

$$(61) \quad \boxed{\begin{aligned} Q_{\hat{x}} &= (A^* Q_y^{-1} A)^{-1} \\ Q_{\hat{y}} &= A Q_{\hat{x}} A^* \\ Q_{\hat{e}} &= Q_y - Q_{\hat{y}} \end{aligned}}$$

Substitution of $W=Q_y^{-1}$ into equation (29) gives:

$$(62) \quad \boxed{\begin{aligned} Q_{\hat{x}\hat{y}} &= Q_{\hat{x}} A^* \\ Q_{\hat{x}\hat{e}} &= 0 \\ Q_{\hat{y}\hat{e}} &= 0 \end{aligned}}$$

This shows that in case of best linear unbiased estimation, the estimator $\hat{\underline{e}}$ is *not* correlated with either $\hat{\underline{x}}$ or $\hat{\underline{y}}$.

Substitution of $W=Q_y^{-1}$ into (30) and (31) gives:

$$\begin{aligned} E\{\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}}\} &= \text{trace} (Q_y^{-1} Q_{\hat{\underline{e}}}) = m - \text{trace} (Q_y^{-1} Q_{\hat{\underline{y}}}) \\ \sigma_{\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}}}^2 &= 2 \text{trace} (Q_y^{-1} Q_{\hat{\underline{e}}}) = 2m - 2 \text{trace} (Q_y^{-1} Q_{\hat{\underline{y}}}) \end{aligned}$$

Since $\text{trace} (Q_y^{-1} Q_{\hat{\underline{y}}}) = \text{trace} (Q_{\hat{\underline{y}}} Q_y^{-1}) = n$ if $\text{rank } A = n$ (for a proof see end of next section) we get:

$$(63) \quad \boxed{\begin{aligned} E\{\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}}\} &= m - n \\ \sigma_{\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}}}^2 &= 2(m - n) \end{aligned}}$$

To conclude this section let us investigate what happens with our results if we assume that the number of observations, m , equals the number of unknowns, n . If $m=n$, then matrix A is a square matrix. We assumed that $\text{rank } A = n$. This implies that if $m=n$, the matrix A is invertible. That is, A^{-1} exists. If A^{-1} exists, the matrix $(A^* Q_y^{-1} A)^{-1}$ can be written as:

$$(64) \quad (A^* Q_y^{-1} A)^{-1} = A^{-1} Q_y A^{*-1}$$

Substitution into (59) gives then:

$$\hat{\underline{x}} = A^{-1} \underline{y},$$

and from this follows that:

$$\hat{\underline{y}} = A \hat{\underline{x}} = \underline{y} \quad \text{and} \quad \hat{\underline{e}} = \underline{y} - \hat{\underline{y}} = \underline{0}.$$

Similarly we find from substituting (64) into (61) that:

$$\begin{aligned} Q_{\hat{\underline{x}}} &= A^{-1} Q_y A^{*-1} \\ Q_{\hat{\underline{y}}} &= Q_y \\ Q_{\hat{\underline{e}}} &= 0 \end{aligned}.$$

These results show that if $m=n$ (there is *no redundancy*), the estimate $\hat{\underline{x}}$ is simply the unique solution of the system $\underline{y} = A\underline{x}$ and the best linear unbiased estimator of $E\{\underline{y}\} = A\underline{x}$, $\hat{\underline{y}}$, is simply $\underline{y}$. Thus if $m=n$ we have just enough observations available to determine the unknown parameter vector $\underline{x}$ uniquely, but not enough observations to get a better estimator for $E\{\underline{y}\} = A\underline{x}$ than $\underline{y}$.

2.5 The orthogonal projector

As a generalization of the theory developed in chapter 1 we will prove in this section that the matrix $A(A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1}$ is an *orthogonal projector* that projects onto the rangespace of A and along its orthogonal complement.

Let the m -dimensional linear vectorspace $\mathbb{R}^m$ have the *inner product*:

$$(65) \quad (u, v)_{Q_y^{-1}} = u^* Q_y^{-1} v \quad \forall u, v \in \mathbb{R}^m.$$

Matrix Q_y^{-1} is thus the *matrix representation* of the inner product (65).

The *range space* of the $m \times n$ matrix A is defined as:

$$(66) \quad R(A) = \{z \in \mathbb{R}^m \mid z = A\alpha \text{ for some } \alpha \in \mathbb{R}^n\}.$$

The range space $R(A)$ is a linear subspace of $\mathbb{R}^m$, $R(A) \subset \mathbb{R}^m$, and its dimension equals the rank of A , which in our case is n :

$$(67) \quad \dim. R(A) = \text{rank } A = n.$$

The *orthogonal complement* of $R(A)$ is defined as the set of vectors that are orthogonal to $R(A)$ (Strang, 1988):

$$(68) \quad R(A)^\perp = \{z \in \mathbb{R}^m \mid A^* Q_y^{-1} z = 0\}.$$

According to your lecture notes in Linear Algebra, $\dim. R(A) + \dim. R(A)^\perp = \dim \mathbb{R}^m$. With $\dim. R(A) = n$ and $\dim. \mathbb{R}^m = m$ this gives:

$$(69) \quad \dim. R(A)^\perp = m - n.$$

In section 2.4 a full rank matrix B of order $m \times (m-n)$ was introduced that satisfied:

$$(70) \quad \begin{matrix} A^* & B & = & 0 \\ n \times m & m \times (m-n) & & n \times (m-n) \end{matrix}.$$

With this matrix B we can represent the orthogonal complement $R(A)^\perp$ of (68) as:

$$(71) \quad R(A)^\perp = \{z \in \mathbb{R}^m \mid z = Q_y B \beta \text{ for some } \beta \in \mathbb{R}^{m-n}\}.$$

According to theory in Linear Algebra, the space $\mathbb{R}^m$ is the *direct sum* of $R(A)$ and $R(A)^\perp$:

$$(72) \quad \mathbb{R}^m = R(A) \oplus R(A)^\perp$$

This means that $R(A) \cap R(A)^\perp = \{0\}$, i.e. the two subspaces $R(A)$ and $R(A)^\perp$ only have the null vector in common, and that a basis of $R(A)$ together with a basis of $R(A)^\perp$ forms a basis of $\mathbb{R}^m$. Since the $m \times n$ matrix A has full rank n , its column vectors form a basis of $R(A)$. Similarly, since the $m \times (m-n)$ matrix B has full rank $m-n$ and the $m \times m$ matrix Q_y has full rank m , the $m \times (m-n)$ matrix $Q_y B$ has full rank $m-n$ and its column vectors form a basis of $R(A)^\perp$. From this follows that the column vectors of the $m \times m$ matrix:

$$(73) \quad \underset{m \times n}{(A : Q_y B)} \underset{m \times (m-n)}{(A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1} \atop (B^* Q_y B)^{-1} B^*}$$

form a basis of $\mathbb{R}^m$. This implies that the $m \times m$ matrix of (73) is square and invertible (this was already shown in the previous section). If we write equation (50) of the previous section as:

$$(A : Q_y B) \begin{pmatrix} (A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1} \\ (B^* Q_y B)^{-1} B^* \end{pmatrix} = I,$$

it follows that the inverse of (73) is given as:

$$(74) \quad (A : Q_y B)^{-1} = \begin{pmatrix} (A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1} \\ (B^* Q_y B)^{-1} B^* \end{pmatrix}.$$

According to theory in Linear Algebra, a vector $z \in \mathbb{R}^m$ can be decomposed in only one way along $R(A)$ and $R(A)^\perp$. If:

$$(75) \quad z = u + v \text{ with } u \in R(A) \text{ and } v \in R(A)^\perp,$$

then u is called the *orthogonal projection* of z on $R(A)$. The matrix P_A , defined as:

$$(76) \quad P_A z = u,$$

is called the *orthogonal projector* on $R(A)$. We now want to find the matrix representation of P_A . Since the column vectors of the matrix of (73) form a basis of $\mathbb{R}^m$, every $z \in \mathbb{R}^m$ can be represented as:

$$z = (A : Q_y B) \begin{pmatrix} \alpha \\ \beta \end{pmatrix} = A \alpha + Q_y B \beta.$$

In this representation $u = A \alpha$ is the component of z along $R(A)$. With (76) this gives:

$$P_A (A : Q_y B) \begin{pmatrix} \alpha \\ \beta \end{pmatrix} = A \alpha$$

or:

$$(77) \quad P_A (A : Q_y B) \begin{pmatrix} \alpha \\ \beta \end{pmatrix} = (A : 0) \begin{pmatrix} \alpha \\ \beta \end{pmatrix}.$$

Since (76) holds for all $z \in \mathbb{R}^m$, (77) holds for all α and β :

$$P_A(A \vdash Q_y B) = (A \vdash 0).$$

Post-multiplication with $(A \vdash Q_y B)^{-1}$ gives:

$$P_A = (A \vdash 0)(A \vdash Q_y B)^{-1}.$$

And substitution of (74) gives then finally:

$$(78) \quad \boxed{P_A = A(A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1}}$$

According to theory in Linear Algebra, P_A has the following properties:

$$(79) \quad \begin{cases} R(P_A) = R(A) \\ N(P_A) = R(A)^\perp = R(Q_y B) \\ P_A P_A = P_A \text{ (idempotence).} \end{cases}$$

From Linear Algebra we know that $I - P_A$ is also an orthogonal projector. $I - P_A$ projects onto $R(A)^\perp$ and along $R(A)$. $I - P_A$ has the properties:

$$(80) \quad \begin{cases} R(I - P_A) = R(A)^\perp \\ N(I - P_A) = R(A) \\ (I - P_A)(I - P_A) = I - P_A \text{ (idempotence).} \end{cases}$$

Instead of $I - P_A$ we will sometimes use the notation $P_A^\perp$. It follows from equation (49) that:

$$(81) \quad \boxed{P_A^\perp = Q_y B(B^* Q_y B)^{-1} B^*}$$

To conclude this section we will prove that:

$$(82) \quad \text{trace } (Q_y Q_y^{-1}) = n$$

Substitution of $Q_y = A(A^* Q_y^{-1} A)^{-1} A^*$ (see (61)) gives:

$$\text{trace } \left[(A(A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1}) \right] = n$$

or with (78):

$$(83) \quad \text{trace } (P_A) = n$$

Since the trace of a matrix is the sum of its eigenvalues (Strang, 1988 and Lay, 1997), we need the eigenvalues of P_A in order to prove (83). The eigenvalue problem for P_A reads:

$$(84) \quad P_A z = \lambda z$$

with $P_A P_A z = \lambda P_A z = \lambda^2 z$ and $P_A P_A z = P_A z = \lambda z$ follows that $\lambda^2 z = \lambda z$ or for $z \neq 0$:

$$(85) \quad \lambda (\lambda - 1) = 0.$$

This shows that the eigenvalues of P_A are 1 or 0. In order to prove (83) we thus need to show that P_A has n -number of eigenvalues equal to 1. This corresponds to showing that for $\lambda=1$ the solution space of (84) has dimension n . For $\lambda=1$ the solution space of (84) is the solution space of $(I-P_A)z=0$ and thus the nullspace of $I-P_A$. According to (80) the nullspace of $I-P_A$ is the rangespace of A . Hence, $\dim. N(I-P_A)=\dim. R(A)=\text{rank } A=n$. This concludes the proof.

2.6 Summary

In the table on the next page an overview is given of the results of Best Linear Unbiased Estimation.

Best Linear Unbiased Estimation

Model

$$E\{\underline{y}\} = \underset{m \times 1}{A} \underset{m \times n \ n \times 1}{x} ; E\{(\underline{y} - A\underline{x})(\underline{y} - A\underline{x})^*\} = \underset{m \times m}{Q_y} \underset{m \times m}{}$$

Estimators

$$\begin{aligned} \hat{\underline{x}} &= (A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1} \underline{y} ; \quad \hat{\underline{\theta}} = f^* \hat{\underline{x}} \\ \hat{\underline{y}} &= A \hat{\underline{x}} = P_A \underline{y} ; \quad \hat{\underline{e}} = \underline{y} - \hat{\underline{y}} = P_A^\perp \underline{y} \\ \hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}} &= \underline{y}^* Q_y^{-1} P_A^\perp \underline{y} \end{aligned}$$

Mean

$$\begin{aligned} E\{\hat{\underline{x}}\} &= x ; \quad E\{\hat{\underline{\theta}}\} = f^* x = \underline{\theta} \\ E\{\hat{\underline{y}}\} &= A x = E\{\underline{y}\} ; \quad E\{\hat{\underline{e}}\} = 0 \\ E\{\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}}\} &= m - n \end{aligned}$$

Variances and covariances

$$\begin{aligned} Q_{\hat{\underline{x}}} &= (A^* Q_y^{-1} A)^{-1} ; \quad \sigma_{\hat{\underline{\theta}}}^2 = f^* Q_{\hat{\underline{x}}} f \\ Q_{\hat{\underline{y}}} &= A Q_{\hat{\underline{x}}} A^* = P_A Q_y P_A^* ; \quad Q_{\hat{\underline{e}}} = Q_y - Q_{\hat{\underline{y}}} = P_A^\perp Q_y P_A^{\perp *} \\ &= P_A^\perp Q_y = Q_y P_A^{\perp *} \\ \sigma_{\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}}}^2 &= 2(m - n) \\ Q_{\hat{\underline{x}}\hat{\underline{y}}} &= Q_{\hat{\underline{x}}} A^* ; \quad Q_{\hat{\underline{x}}\hat{\underline{e}}} = 0 ; \quad Q_{\hat{\underline{y}}\hat{\underline{e}}} = 0 \end{aligned}$$

Orthogonal projectors

$$\begin{aligned} P_A &= A(A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1} ; \quad R(P_A) = R(A) \quad , \quad N(P_A) = R(A)^\perp \\ P_A^\perp &= I - A(A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1} ; \quad R(P_A^\perp) = R(A)^\perp \quad , \quad N(P_A^\perp) = R(A) \end{aligned}$$

3 The model with condition equations

3.1 Introduction

Figure 3.1 shows the levelling network of section 2.1.

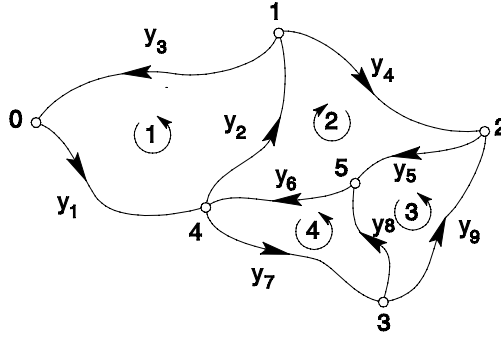


Figure 3.1: Levelling network

In section 2.1 it was shown that the *observation equations* of the levelling network of figure 3.1 take the form:

$$(1) \quad \begin{pmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \\ y_5 \\ y_6 \\ y_7 \\ y_8 \\ y_9 \end{pmatrix} = \begin{pmatrix} & & & & 1 \\ & 1 & & & -1 \\ -1 & & & & \\ -1 & 1 & & & \\ & -1 & & 1 & \\ & & & 1 & -1 \\ & & 1 & -1 & \\ & & -1 & & 1 \\ 1 & -1 & & & \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ x_5 \end{pmatrix} + \begin{pmatrix} e_1 \\ e_2 \\ e_3 \\ e_4 \\ e_5 \\ e_6 \\ e_7 \\ e_8 \\ e_9 \end{pmatrix}$$

For simplicity it is assumed here that the known height of the reference point 0 is equal to zero. The matrix equation of (1) is of the form:

$$(2) \quad \boxed{\begin{matrix} y = A x + e \\ m \times 1 \quad m \times n \quad n \times 1 \quad m \times 1 \end{matrix}}$$

with $m=9$ and $n=5$. Thus since $\text{rank } A=n=5$, the number of *redundant* observations is $m-n=4$.

Instead of writing the observed height differences as functions of the unknown heights, as is done in (1), it is also possible to write down *condition equations* (voorwaarde vergelijkingen) which the height differences have to fulfill. In figure 3.1 four loops (kringen) are shown. With perfect measurements the height differences of each loop would sum up to zero. Due however to the presence of measurement errors these sums will differ from zero. That is, misclosures will exist. If we denote these four *misclosures* (tegenspraken) by respectively $\underline{t}_1$, $\underline{t}_2$, $\underline{t}_3$ and $\underline{t}_4$, the four condition equations become (see figure 3.1):

$$(3) \quad \begin{cases} y_1 + y_2 + y_3 &= \underline{t}_1 \\ y_2 + y_4 + y_5 + y_6 &= \underline{t}_2 \\ y_9 + y_5 - y_8 &= \underline{t}_3 \\ y_6 + y_7 + y_8 &= \underline{t}_4 \end{cases}$$

These four condition equations can be written in matrix notation as:

$$(4) \quad \begin{pmatrix} 1 & 1 & 1 & & & & & & \\ & 1 & & 1 & 1 & 1 & & & \\ & & & 1 & & & -1 & 1 & \\ & & & & 1 & 1 & 1 & & \end{pmatrix} \begin{pmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \\ y_5 \\ y_6 \\ y_7 \\ y_8 \\ y_9 \end{pmatrix} = \begin{pmatrix} \underline{t}_1 \\ \underline{t}_2 \\ \underline{t}_3 \\ \underline{t}_4 \end{pmatrix}$$

This matrix equation is of the form:

$$(5) \quad \boxed{\begin{matrix} \mathbf{B}^* \mathbf{y} = \mathbf{t} \\ \mathbf{b} \times \mathbf{m} \quad \mathbf{m} \times \mathbf{1} \quad \mathbf{b} \times \mathbf{1} \end{matrix}}$$

with $b=4$ and $m=9$. Thus matrix B has $m=9$ rows and $b=4$ columns.

Note that in the above example of the levelling network the number of redundant observations, $m-n=4$, equals the number of condition equations, $b=4$. This is generally true, because each additional observation on top of the n -number of observations which are needed for the unique determination of x in (2), gives rise to an extra condition equation. Thus:

$$(6) \quad \boxed{b = m - n}$$

Also note that if the matrix of (1) is pre-multiplied with the matrix of (4), the result equals the zero-matrix. In terms of the matrices A and B this means that:

$$(7) \quad \boxed{\begin{matrix} \mathbf{B}^* \mathbf{A} = \mathbf{0} \\ b \times m \quad m \times n \quad b \times n \end{matrix}}$$

Also this relation between the two matrices A and B is generally true. This can be seen as follows. Let us start from equation (2). If we pre-multiply equation (2) with a $b \times m$ matrix B^* we get:

$$(8) \quad \begin{matrix} \mathbf{B}^* \mathbf{y} \\ b \times m \quad m \times 1 \end{matrix} = \begin{matrix} \mathbf{B}^* \mathbf{A} \quad \mathbf{x} \\ b \times m \quad m \times n \quad n \times 1 \end{matrix} + \begin{matrix} \mathbf{B}^* \mathbf{e} \\ b \times m \quad m \times 1 \end{matrix}$$

Our objective is now to investigate whether a full rank matrix B^* , with rank $B^* = b = m - n$, exists such that $B^*A = 0$. If such a matrix exists, then the unknown parameter vector x gets eliminated from (8), and (8) reduces to the form of (5) with:

$$(9) \quad \boxed{\underline{t} = \mathbf{B}^* \mathbf{e}}$$

If we denote the b -number of columnvectors of matrix B by $b_i, i=1, \dots, b$, then $B^*A=0$ is equivalent to:

$$(10) \quad \begin{matrix} \mathbf{A}^* \mathbf{b}_i \\ n \times m \quad m \times 1 \end{matrix} = \begin{matrix} \mathbf{0} \\ n \times 1 \end{matrix}, i=1, \dots, b$$

The question is thus if and how many linearly independent vectors exist that satisfy (10). From Linear Algebra we know that for an $n \times m$ matrix A^* :

$$(11) \quad \begin{cases} \dim. R(A^*) = \text{rank} A^* \\ \text{rank} A^* = \text{rank} A \\ \dim. R(A^*) + \dim. N(A^*) = \dim. \mathbb{R}^m = m \end{cases}$$

Since rank $A = n$, it follows from (11) that:

$$(12) \quad \dim. N(A^*) = m - n.$$

Thus the dimension of the nullspace of A^* is $m - n$. This shows that indeed $m - n$ linearly independent vectors b_i exist that satisfy (10). Note that if there are no redundant observations $m = n$ and (10) has only the trivial zero vector as solution. Thus if $m = n$, no matrix B exists for which (7) holds and no condition equations exist.

With (2) and (5) we now have two equivalent representations of the same model. Equation (2) is a parametric representation. The equations of (2) are known as *observation equations*. Equation (5) is an implicit representation. The equations of (5) are known as *condition equations*. In

chapter 2 we developed the theory of linear estimation for (2). In this chapter we will develop the theory for the model with condition equations (5).

3.2 Best Linear Unbiased Estimation

From equations (5) and (9) follows that:

$$(13) \quad \underset{b \times m}{B}^* \underset{m \times 1}{y} = \underset{b \times 1}{t} \text{ with } \underset{b \times 1}{t} = \underset{b \times m}{B}^* \underset{m \times 1}{e}$$

This shows that the misclosure vector $\underline{t}$ is due to the random measurement error vector $\underline{e}$. We will assume, like we did in chapter 2, that:

$$(14) \quad E\{\underline{e}\} = \underline{0} \quad \text{and} \quad E\{\underline{e}\underline{e}^*\} = \underline{Q}_y$$

This, together with (13) gives:

$$(15) \quad \boxed{\underset{b \times m}{B}^* E\{\underset{m \times 1}{y}\} = \underset{b \times 1}{0} ; E\{(\underset{m \times 1}{y} - E\{\underset{m \times 1}{y}\})(\underset{m \times 1}{y} - E\{\underset{m \times 1}{y}\})^*\} = \underset{m \times m}{Q}_y}$$

This model will be our starting point for best linear unbiased estimation. Let us assume that we are interested in estimating a parameter ϕ , which is a linear function of $E\{\underline{y}\}$:

$$(16) \quad \underset{1 \times 1}{\phi} = \underset{1 \times m}{g}^* \underset{m \times 1}{E\{\underline{y}\}}$$

For instance, in case of the levelling network of figure 3.1 we could be interested in estimating the height difference between the points 1 and 5. Then ϕ of (16) plays the role of the height difference between the points 1 and 5, and the vector g of (16) takes the form:

$$g = (0 \ 0 \ 0 \ 1 \ 1 \ 0 \ 0 \ 0)^*$$

Let $\hat{\phi}$ be the estimator of ϕ . Then according to the criteria of best linear unbiased estimation, we want the estimator $\hat{\phi}$ to be a linear function of $\underline{y}$:

$$(17) \quad \underset{1 \times 1}{\hat{\phi}} = \underset{1 \times m}{l}^* \underset{m \times 1}{y} \quad \text{"L-property"}$$

such that

$$(18) \quad E\{\hat{\phi}\} = \phi, \quad \text{"U-property"}$$

and

$$(19) \quad \sigma_{\hat{\phi}}^2 = \text{minimal} \quad \text{"B-property"}$$

Our objective is now to find a vector l such that with (17), the conditions (18) and (19) are fulfilled.

Substitution of (17) into (18) gives:

$$(20) \quad E\{\hat{\phi}\} = l^* E\{y\} = \phi.$$

If we subtract (16) from (20) we get:

$$(21) \quad \underset{1 \times m}{(l - g)^*} \underset{m \times 1}{E\{y\}} = \underset{1 \times 1}{0}.$$

This equation states that the vector $l - g$ is orthogonal to the mean $E\{y\}$. From the first equation of (15) we learn that the subspace orthogonal to $E\{y\}$ is spanned by the b linearly independent column vectors of matrix B . This, together with (21) shows that the vector $l - g$ is an element of the range space of the matrix B :

$$l - g \in R(B).$$

This means that a $b \times 1$ vector α exists such that:

$$l - g = B\alpha$$

or:

$$(22) \quad \boxed{\underset{m \times 1}{l} = \underset{m \times 1}{g} + \underset{m \times b}{B} \underset{b \times 1}{\alpha}}$$

The meaning of this result is that every vector l that can be represented as (22), gives with (17) an unbiased estimator of (16). To verify this, we substitute (22) into (17):

$$\hat{\phi} = g^* y + \alpha^* B^* y$$

If we take the expectation we get:

$$E\{\hat{\phi}\} = g^* E\{y\} + \alpha^* B^* E\{y\},$$

and with (15):

$$E\{\hat{\phi}\} = g^* E\{y\},$$

which according to (16) is equal to ϕ .

Let us now introduce condition (19), the condition of minimum variance. Application of the *propagation law of variances* to (17) gives with (15):

$$\sigma_{\hat{\phi}}^2 = l^* Q_y l$$

Hence, our minimization problem reads:

$$(23) \quad \min_l l^* Q_y l \text{ subject to } l = g + B\alpha$$

This minimization problem may also be written as:

$$\min_{\alpha} (g + B\alpha)^* Q_y (g + B\alpha).$$

The solution of this minimization problem is:

$$(24) \quad \hat{\alpha} = -(B^* Q_y B)^{-1} B^* Q_y g.$$

This is easily verified by using an algebraic proof like the one given in section 1.2. Substitution of (24) into (22) gives:

$$(25) \quad \hat{l} = [I - B(B^* Q_y B)^{-1} B^* Q_y] g.$$

The best linear unbiased estimator $\hat{\phi}$ of ϕ follows therefore with (17) as:

$$(26) \quad \hat{\phi} = g^* [I - Q_y B(B^* Q_y B)^{-1} B^*] y$$

If we take for the vector g respectively $(1, 0, \dots, 0)^*$, $(0, 1, 0, \dots, 0)^*$, ..., $(0, \dots, 0, 1)^*$, the BLUE of $E\{y\}$ follows as:

$$(27) \quad \hat{y} = [I - Q_y B(B^* Q_y B)^{-1} B^*] y$$

This, together with (26) shows that:

$$(28) \quad \hat{\phi} = g^* \hat{y}$$

If we pre-multiply (27) with B^* , we get:

$$(29) \quad B^* \hat{y} = 0,$$

which shows that the estimator $\hat{y}$ satisfies the condition equations. Since $\underline{t} = B^* \underline{y}$ (see (13)), application of the propagation law of variances gives $Q_t = B^* Q_y B$. This shows that (27) can also be written as:

$$(30) \quad \hat{y} = y - Q_y B Q_t^{-1} \underline{t}$$

This equation shows perhaps more clearly than (27), that the estimator $\hat{y}$ is obtained from y by subtracting a part which depends on the misclosure vector $\underline{t}$ such that $\hat{y}$ satisfies (29). Since:

$$(31) \quad I - Q_y B(B^* Q_y B)^{-1} B^* = A(A^* Q_y^{-1} A)^{-1} A^* Q_y^{-1},$$

(see equation (49) in section 2.2), the estimator of (27) is *identical* to the estimator of $E\{y\}$ which was derived in the previous chapter (see for instance the summary in section 2.6). This is of course understandable, because we used the *same* principle of estimation, namely BLUE, and the *same* model. Model (15) is namely equivalent to the model in section 2.6. The two models only differ in their representation.

The results of estimation for model (15) are therefore identical to the results of the previous chapter. The results only differ in their appearance. In chapter 2 everything was expressed in terms of the matrix A , whereas in this chapter everything is expressed in terms of the matrix B .

3.3 The case $B^* E\{y\} = b^0, b^0 \neq 0$

The model of condition equations (4) was based on the fact that the observed height difference of each levelling loop in figure 3.1 should sum up to zero. Now consider the triangulation network of figure 3.2.

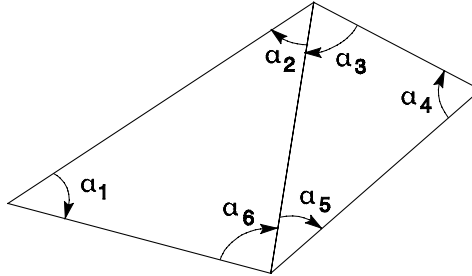


Figure 3.2: Small triangulation network

From the figure it is clear that the six angle observables should fulfill two conditions: in each triangle the expectation of the three observed angles sums up to π . Due however to the presence of measurement errors these condition equations will result in small misclosures. The two condition equations become in matrix notation:

$$(32) \quad \begin{pmatrix} \alpha_1 \\ \alpha_2 \\ \alpha_3 \\ \alpha_4 \\ \alpha_5 \\ \alpha_6 \end{pmatrix} = \begin{pmatrix} \pi \\ \pi \end{pmatrix} + \begin{pmatrix} t_1 \\ t_2 \end{pmatrix}.$$

Compared to (5) this matrix equation is of the form:

$$(33) \quad \underset{\substack{\mathbf{B}^* \\ b \times m}}{\mathbf{B}^*} \underset{\substack{\mathbf{y} \\ m \times 1}}{\mathbf{y}} = \underset{\substack{\mathbf{b}^0 \\ b \times 1}}{\mathbf{b}^0} + \underset{\substack{\mathbf{t} \\ b \times 1}}{\mathbf{t}}$$

with $b=2$, $m=6$ and $\mathbf{b}^0 = (\pi, \pi)^*$.

For a non-zero vector $\mathbf{b}^0$ straightforward application of the estimation formulae of the previous section is no longer possible. At first sight we can rewrite (32) e.g. as:

$$\begin{pmatrix} 1 & 1 & 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 1 & 1 & 0 \end{pmatrix} \begin{pmatrix} \underline{\alpha}_1 - \pi \\ \underline{\alpha}_2 \\ \underline{\alpha}_3 \\ \underline{\alpha}_4 - \pi \\ \underline{\alpha}_5 \\ \underline{\alpha}_6 \end{pmatrix} = \begin{pmatrix} \underline{t}_1 \\ \underline{t}_2 \end{pmatrix},$$

so that a model of the form of (5) is obtained. However, in general it is difficult to redefine the vector of observables so that it is of the form of (5). Therefore we will adapt the estimation formulae so that the non-zero vector $\mathbf{b}^0$ can be taken into account. From (33) the vector of misclosures and its variance matrix are obtained as:

$$(34) \quad \mathbf{t} = \mathbf{B}^* \mathbf{y} - \mathbf{b}^0; \quad \mathbf{Q}_t = \mathbf{B}^* \mathbf{Q}_y \mathbf{B}.$$

Then with (30) the estimators $\hat{\mathbf{y}}$ and $\hat{\mathbf{e}}$ become:

$$(35) \quad \begin{cases} \hat{\mathbf{y}} = [\mathbf{I} - \mathbf{Q}_y \mathbf{B} (\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{B}^*] \mathbf{y} + \mathbf{Q}_y \mathbf{B} (\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{b}^0 \\ \hat{\mathbf{e}} = \mathbf{Q}_y \mathbf{B} (\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{B}^* \mathbf{y} - \mathbf{Q}_y \mathbf{B} (\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{b}^0 \end{cases}$$

Comparing this result with (27) shows that the non-zero vector $\mathbf{b}^0$ results in an additional term for the least-squares estimators. Since $\mathbf{b}^0$ is a vector of constants the extension of the model of condition equations with a non-zero $\mathbf{b}^0$ -vector does not affect the variance matrices $\mathbf{Q}_p$, $\mathbf{Q}_s$ and $\mathbf{Q}_e$.

3.4 Summary

Since the results of estimation for model (15) are identical to the results of the previous chapter, we merely need to represent the results of section 2.6 in terms of the matrix $\mathbf{B}$ to obtain a summary of the results for model (15). This is done on the next page, there upon the main results of this and the previous chapter are compared.

Best Linear Unbiased Estimation

Model

$$\underset{b \times m}{B}^* \underset{m \times 1}{E\{\underline{y}\}} = \underset{b \times 1}{0} ; \underset{m \times m}{E\{(\underline{y} - E\{\underline{y}\})(\underline{y} - E\{\underline{y}\})^*\}} = \underset{m \times m}{Q_y}$$

Estimators

$$\begin{aligned} \hat{\underline{y}} &= P_{Q_y B}^\perp \underline{y} ; \quad \hat{\underline{e}} = \underline{y} - \hat{\underline{y}} = Q_y B Q_t^{-1} \underline{t} = P_{Q_y B} \underline{y} \\ \hat{\underline{\phi}} &= \underline{g}^* \hat{\underline{y}} \\ \hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}} &= \underline{t}^* Q_t^{-1} \underline{t} = \underline{y}^* Q_y^{-1} P_{Q_y B} \underline{y} \end{aligned}$$

Mean

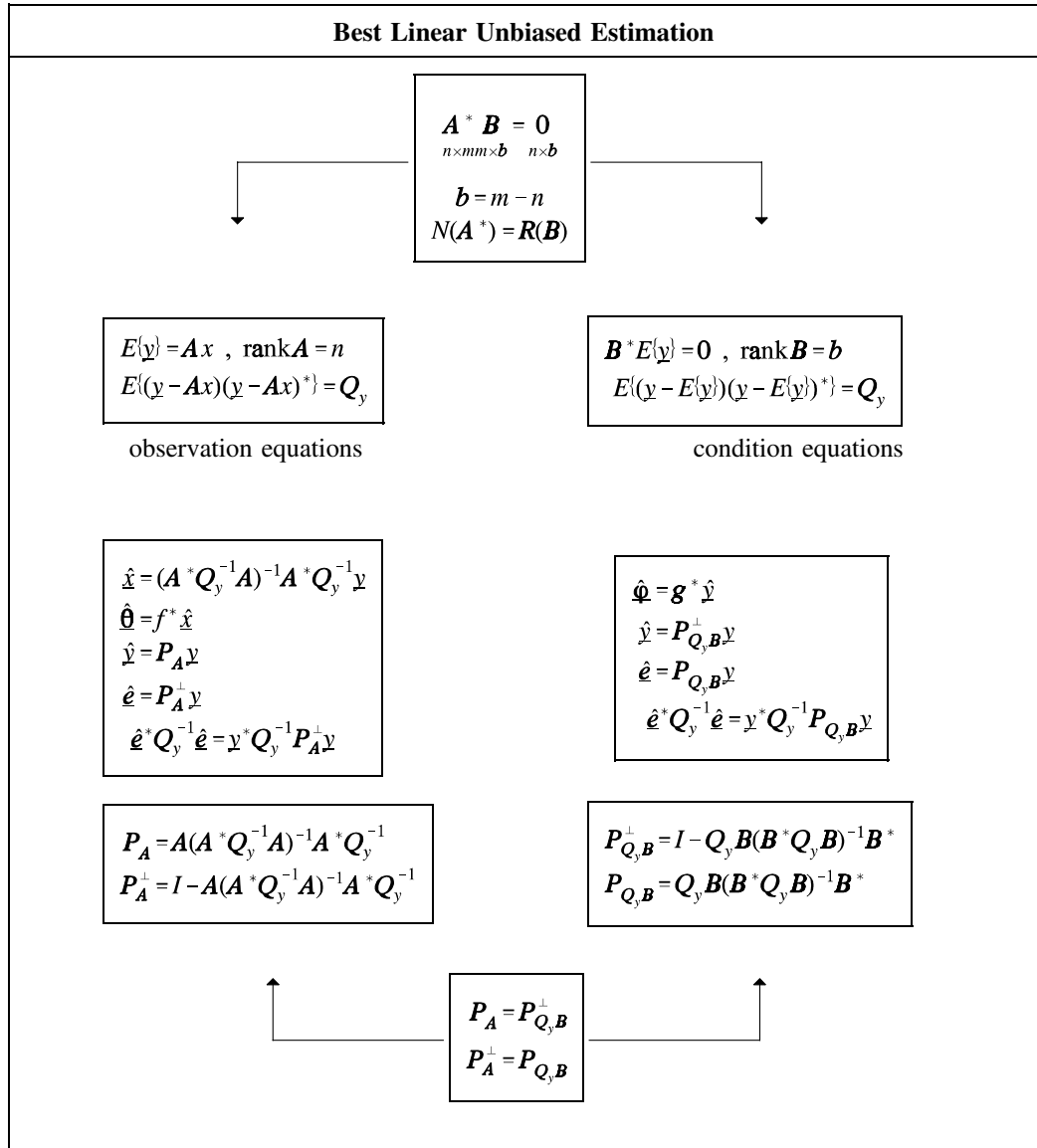
$$\begin{aligned} E\{\hat{\underline{y}}\} &= E\{\underline{y}\} ; \quad E\{\hat{\underline{e}}\} = 0 \\ E\{\hat{\underline{\phi}}\} &= \underline{g}^* E\{\underline{y}\} \\ E\{\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}}\} &= E\{\underline{t}^* Q_t^{-1} \underline{t}\} = b \end{aligned}$$

Variances and covariances

$$\begin{aligned} Q_{\hat{\underline{y}}} &= P_{Q_y B}^\perp Q_y P_{Q_y B}^{\perp *} ; \quad Q_{\hat{\underline{e}}} = Q_y - Q_{\hat{\underline{y}}} = Q_y B Q_t^{-1} B^* Q_y \\ &= P_{Q_y B} Q_y P_{Q_y B}^* \\ \sigma_{\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}}}^2 &= \sigma_{\underline{t}^* Q_t^{-1} \underline{t}}^2 = 2b \\ Q_{\hat{\underline{y}} \hat{\underline{e}}} &= 0 \end{aligned}$$

Orthogonal projectors

$$\begin{aligned} P_{Q_y B}^\perp &= I - Q_y B (B^* Q_y B)^{-1} B^* ; \quad R(P_{Q_y B}^\perp) = R(Q_y B)^\perp ; \quad N(P_{Q_y B}^\perp) = R(Q_y B) \\ P_{Q_y B} &= Q_y B (B^* Q_y B)^{-1} B^* ; \quad R(P_{Q_y B}) = R(Q_y B) ; \quad N(P_{Q_y B}) = R(Q_y B)^\perp \end{aligned}$$



4 y^R -Variates

4.1 Introduction

For the case that $m=2$ and $n=1$ the model with *observation equations* reads:

$$(1) \quad \begin{matrix} E\{\underline{y}\} = & \mathbf{a} & \mathbf{x} & ; & E\{(\underline{y}-\mathbf{a}\mathbf{x})(\underline{y}-\mathbf{a}\mathbf{x})^*\} = & \mathbf{Q}_y \end{matrix} \quad \begin{matrix} 2 \times 1 & 2 \times 1 & 1 \times 1 & & 2 \times 2 & 2 \times 2 \end{matrix}$$

The corresponding model with *condition equations* reads:

$$(2) \quad \begin{matrix} \mathbf{b}^* & E\{\underline{y}\} = & \mathbf{0} & ; & E\{(\underline{y}-E\{\underline{y}\})(\underline{y}-E\{\underline{y}\})^*\} = & \mathbf{Q}_y \end{matrix} \quad \begin{matrix} 1 \times 2 & 2 \times 1 & 1 \times 1 & & 2 \times 2 & 2 \times 2 \end{matrix}$$

In chapter 1 (see figure 1.8 or figure 1.14 in section 1.3) it was shown that the estimator $\hat{\underline{y}}$ of $E\{\underline{y}\}$ is the oblique projection of $\underline{y}$ onto the line with direction vector $\mathbf{a}$. The direction of projection is parallel to the tangent line of the ellipse $\mathbf{z}^* \mathbf{Q}_y^{-1} \mathbf{z} = \mathbf{a}^* \mathbf{Q}_y^{-1} \mathbf{a}$ at $\mathbf{a}$ (see figure 4.1).

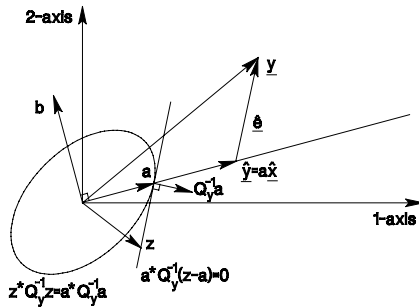


Figure 4.1: Oblique projection of $\underline{y}$

If we rotate the direction vector $\mathbf{a}$ so that it becomes parallel to the 1-axis, figure 4.1 transforms into figure 4.2.

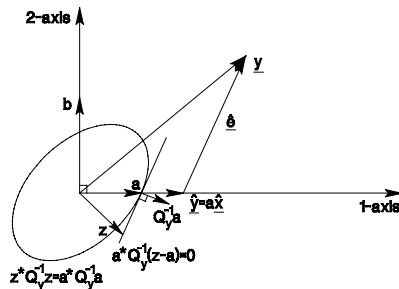


Figure 4.2: Oblique projection of $\underline{y}$ in case $\mathbf{a} = (a_1, 0)^*$

In this case the second component a_2 of the vector a is zero, and the first component b_1 of the vector b is zero. The corresponding two models (1) and (2) then take the form:

$$(3) \quad E\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = \begin{pmatrix} a_1 \\ 0 \end{pmatrix} x ; \quad E(y - ax)(y - ax)^* = \begin{pmatrix} \sigma_1^2 & \sigma_{12} \\ \sigma_{21} & \sigma_2^2 \end{pmatrix}$$

and:

$$(4) \quad \begin{pmatrix} 0 \\ 1 \end{pmatrix}^* E\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = 0 ; \quad E\{(y - E\{y\})(y - E\{y\})^*\} = \begin{pmatrix} \sigma_1^2 & \sigma_{12} \\ \sigma_{21} & \sigma_2^2 \end{pmatrix}.$$

Equation (4) shows that the coefficient of $E\{y_1\}$ equals zero. This implies that the observable y_1 does *not* appear in the condition equation. Observables that do not appear in the condition equations are called *free variates*. Thus $E\{y_1\}$ is a free variate. If one writes (4) as:

$$(5) \quad E\{y_2\} = 0 ; \quad E\{(y_2 - E\{y_2\})(y_2 - E\{y_2\})^*\} = \sigma_2^2 ,$$

one might be inclined to conclude, since $E\{y_1\}$ does not appear in (5), that the estimator $\hat{y}_1$ of $E\{y_1\}$ equals y_1 and thus that $\hat{\epsilon}_1 = y_1 - \hat{y}_1$ is zero. This however is *not true*. Figure 4.2 shows namely clearly that both $\hat{\epsilon}_1$ and $\hat{\epsilon}_2$ are unequal to zero. The reason for this is that the major and minor axes of the ellipse of figure 4.2 make a non-zero angle with the 1-axis and 2-axis. If the major and minor axes of the ellipse are parallel to the 1-axis and the 2-axis, then figure 4.2 transforms into figure 4.3.

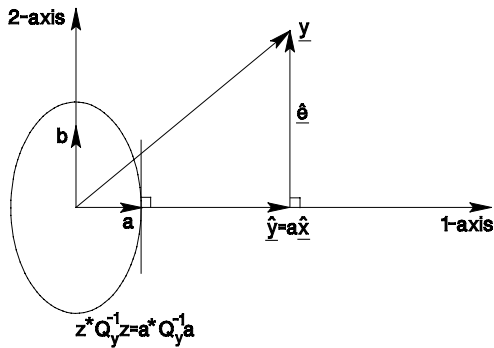


Figure 4.3: The covariance matrix Q_y is diagonal

In this case $\sigma_{12}=0$, Q_y is diagonal and $\hat{\epsilon}_1=0$. It thus seems that the covariance, σ_{12} , between y_1 and y_2 determines to what extent $\hat{\epsilon}_1 = y_1 - \hat{y}_1$ differs from zero. If $\sigma_{12}=0$ then $\hat{\epsilon}_1=0$, and if

$\sigma_{12} \neq 0$ then also $\hat{e}_1 \neq 0$. In order to find out how $\hat{e}_1$ depends on σ_{12} consult figure 4.4. Figure 4.4 shows that:

$$(6) \quad \frac{\hat{e}_1}{\hat{e}_2} = \tan \beta ,$$

and that β is the angle between the vector $Q_y^{-1}a$ and the 1-axis.

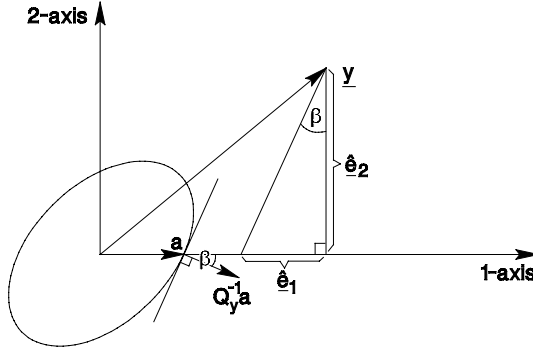


Figure 4.4: $\tan \beta = \frac{\hat{e}_1}{\hat{e}_2}$

Since the inverse of Q_y is:

$$Q_y^{-1} = (\sigma_1^2 \sigma_2^2 - \sigma_{12}^2)^{-1} \begin{pmatrix} \sigma_2^2 & -\sigma_{12} \\ -\sigma_{12} & \sigma_1^2 \end{pmatrix}$$

it follows with $a = (a_1 \ 0)^*$ that:

$$Q_y^{-1}a = (\sigma_1^2 \sigma_2^2 - \sigma_{12}^2)^{-1} \begin{pmatrix} \sigma_2^2 a_1 \\ -\sigma_{12} a_1 \end{pmatrix} .$$

From this it follows that:

$$(7) \quad \tan \beta = \frac{\sigma_{12}}{\sigma_2^2} ,$$

which with (6) gives that:

$$(8) \quad \boxed{\hat{e}_1 = \sigma_{12} \sigma_2^{-2} \hat{e}_2 .}$$

This result shows how $\hat{e}_1 = y_1 - \hat{y}_1$, and thus the estimator $\hat{y}_1 = y_1 - \hat{e}_1$ of free variates, can be determined from $\hat{e}_2 = y_2 - \hat{y}_2$. In the next section we will generalize (8) to the multi-dimensional case.

Free variates are examples of so-called $\underline{y}^R$ -variates. The *definition* of $\underline{y}^R$ -variates reads as follows:

$\underline{y}^R$ -variates are observables that are either stochastically or functionally Related to another set of observables $\underline{y}$.

There are three types of $\underline{y}^R$ -variates:

1. $\underline{y}^R$ -variates that correlate with $\underline{y}$ -variates. These are the *free variates* (vrije grootheden);
2. $\underline{y}^R$ -variates that are functions of $\underline{y}$ -variates. These are the *derived variates* (afgeleide grootheden);
3. $\underline{y}^R$ -variates of which the $\underline{y}$ -variates are functions. These are the *constituent variates* (samenstellende grootheden).

ad 1. These variates occur for instance when new measurements need to be included in existing geodetic networks. Consider for instance the levelling network of figure 3.1 of Chapter 3. Assume that the five heights have been estimated with the available nine levelled height differences. This gives the five estimators of the heights, $\hat{x}_1, \hat{x}_2, \hat{x}_3, \hat{x}_4, \hat{x}_5$. Now assume that a new height difference has been measured, say between the points 1 and 5. If we denote this height difference by y_{10} , we can write the condition equation as:

$$(9) \quad E\{y_{10}\} - E\{\hat{x}_5 - \hat{x}_1\} = 0.$$

On the basis of this condition equation we can compute $\hat{y}_{10}$ and $\hat{x}_5, \hat{x}_1$, the improved estimators for the heights of the points 1 and 5. However, since the estimators $\hat{x}_2, \hat{x}_3$ and $\hat{x}_4$ are correlated with $\hat{x}_1$ and $\hat{x}_5$, also the estimators $\hat{x}_2, \hat{x}_3$ and $\hat{x}_4$ can be improved. The variates $\hat{x}_2, \hat{x}_3$ and $\hat{x}_4$ are in this case the free variates since they do not appear in condition equation (9).

- ad 2. Examples of derived variates in the case of the levelling network are height differences which are not directly measured, but which can be computed as functions of the estimated heights.
- ad 3. Consider the levelling network of figure 4.5. Assume that the height difference between the points 1 and 2 has not been measured directly, but that instead it equals the sum of a number of measured height differences.

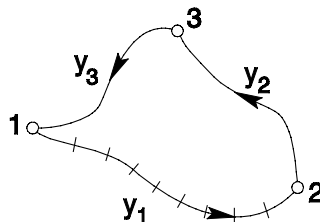


Figure 4.5: Levelling network

One can now include either the individual measured height differences in the condition equation or their sum $\underline{y}_1$. In the last case the individual measured height differences are the constituent variates.

In the following sections it will be shown that formula (8) holds for all three types of $\underline{y}^R$ -variates.

4.2 Free variates

In this section we will generalize formula (8) to the multi-dimensional case. We will give two derivations. One based on the model with *condition equations*, and one based on the model with *observation equations*.

We know that the solution of the model with condition equations

$$(10) \quad \mathbf{B}^* E\{\underline{y}\} = \mathbf{0} \quad ; \quad E\{(\underline{y} - E\{\underline{y}\})(\underline{y} - E\{\underline{y}\})^*\} = \mathbf{Q}_y ,$$

reads:

$$(11) \quad \begin{cases} \hat{\underline{y}} = [\mathbf{I} - \mathbf{Q}_y \mathbf{B}(\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{B}^*] \underline{y} \\ \hat{\underline{e}} = \underline{y} - \hat{\underline{y}} = \mathbf{Q}_y \mathbf{B}(\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{B}^* \underline{y}. \end{cases}$$

Now let us assume that instead of (10) we have the model:

$$(12) \quad \begin{pmatrix} \mathbf{B} \\ \mathbf{0} \end{pmatrix}^* E\left\{\begin{pmatrix} \underline{y} \\ \underline{y}^R \end{pmatrix}\right\} = \mathbf{0}; E\left\{\left[\begin{pmatrix} \underline{y} \\ \underline{y}^R \end{pmatrix} - E\left\{\begin{pmatrix} \underline{y} \\ \underline{y}^R \end{pmatrix}\right\}\right] \left[\begin{pmatrix} \underline{y} \\ \underline{y}^R \end{pmatrix} - E\left\{\begin{pmatrix} \underline{y} \\ \underline{y}^R \end{pmatrix}\right\}\right]^*\right\} = \begin{pmatrix} \mathbf{Q}_y & \mathbf{Q}_{y^R} \\ \mathbf{Q}_{y^R} & \mathbf{Q}_R \end{pmatrix}.$$

In this case the coefficients of $\underline{y}^R$ are zero; thus $\underline{y}^R$ is a vector of *free variates*. When we compare (12) with (10) we see that $\mathbf{B}^*$ of (10) is replaced by $(\mathbf{B}^* \mathbf{0})$ and $\mathbf{Q}_y$ of (10) by:

$$\begin{pmatrix} \mathbf{Q}_y & \mathbf{Q}_{y^R} \\ \mathbf{Q}_{y^R} & \mathbf{Q}_R \end{pmatrix}.$$

Using the second equation of (11), we therefore get for model (12):

$$\begin{pmatrix} \hat{\underline{e}} \\ \hat{\underline{e}}^R \end{pmatrix} = \begin{pmatrix} \underline{y} \\ \underline{y}^R \end{pmatrix} - \begin{pmatrix} \hat{\underline{y}} \\ \hat{\underline{y}}^R \end{pmatrix} = \begin{pmatrix} \mathbf{Q}_y & \mathbf{Q}_{y^R} \\ \mathbf{Q}_{y^R} & \mathbf{Q}_R \end{pmatrix} \begin{pmatrix} \mathbf{B} \\ \mathbf{0} \end{pmatrix} (\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{B}^* \underline{y}$$

or:

$$(13) \quad \begin{pmatrix} \hat{\underline{e}} \\ \hat{\underline{e}}^R \end{pmatrix} = \begin{pmatrix} \mathbf{Q}_y \mathbf{B}(\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{B}^* \underline{y} \\ \mathbf{Q}_{y^R} \mathbf{B}(\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{B}^* \underline{y} \end{pmatrix}.$$

From the first equation of (13) it follows that:

$$\mathbf{Q}_y^{-1} \hat{\underline{e}} = \mathbf{B}(\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{B}^* \underline{y} .$$

If we substitute this into the second equation of (13) we finally get:

(14)

$$\underline{\hat{e}}^R = Q_{Ry} Q_y^{-1} \underline{\hat{e}}$$

This is the multi-dimensional generalization of equation (8) of section 4.1.

We will now derive formula (14) using the model with observation equations. The equivalent of model (12) in terms of observation equations reads:

$$(15) \quad E \begin{pmatrix} y \\ y^R \end{pmatrix} = \begin{pmatrix} A & 0 \\ 0 & I \end{pmatrix} \begin{pmatrix} x \\ E(y^R) \end{pmatrix}; \quad E \left[\begin{pmatrix} y \\ y^R \end{pmatrix} - E \begin{pmatrix} y \\ y^R \end{pmatrix} \right] \left[\begin{pmatrix} y \\ y^R \end{pmatrix} - E \begin{pmatrix} y \\ y^R \end{pmatrix} \right]^* = \begin{pmatrix} Q_y & Q_{yR} \\ Q_{Ry} & Q_R \end{pmatrix}.$$

Note that:

$$\begin{pmatrix} B \\ 0 \end{pmatrix}^* \begin{pmatrix} A & 0 \\ 0 & I \end{pmatrix} = \begin{pmatrix} 0 & 0 \end{pmatrix}$$

holds.

If we denote the inverse of the covariance matrix of the observables by:

$$(16) \quad \begin{pmatrix} Q_y & Q_{yR} \\ Q_{Ry} & Q_R \end{pmatrix}^{-1} = \begin{pmatrix} G_y & G_{yR} \\ G_{Ry} & G_R \end{pmatrix},$$

the *normal equations* for model (15) take the form:

$$(17) \quad \begin{pmatrix} A^* G_y A & A^* G_{yR} \\ G_{Ry} A & G_R \end{pmatrix} \begin{pmatrix} \hat{x} \\ \hat{y}^R \end{pmatrix} = \begin{pmatrix} A^* G_y & A^* G_{yR} \\ G_{Ry} & G_R \end{pmatrix} \begin{pmatrix} y \\ y^R \end{pmatrix}.$$

From the last equation of (17) it follows that:

$$G_R \hat{y}^R = G_R y^R + G_{Ry} (y - A \hat{x})$$

or:

$$\underline{\hat{e}}^R = y^R - \hat{y}^R = -G_R^{-1} G_{Ry} (y - A \hat{x})$$

or:

$$(18) \quad \underline{\hat{e}}^R = -G_R^{-1} G_{Ry} \underline{\hat{e}}.$$

This results looks already very similar to (14). However, equation (14) is expressed in terms of covariance matrices, whereas (18) is expressed in terms of weight matrices.

Since:

$$\begin{pmatrix} G_y & G_{yR} \\ G_{Ry} & G_R \end{pmatrix} \begin{pmatrix} Q_y & Q_{yR} \\ Q_{Ry} & Q_R \end{pmatrix} = \begin{pmatrix} I & 0 \\ 0 & I \end{pmatrix}$$

(see (16)), it follows that $G_{Ry} Q_y + G_R Q_{Ry} = 0$ and thus that:

$$(19) \quad -G_R^{-1} G_{Ry} = Q_{Ry} Q_y^{-1} .$$

This, together with (18) proves (14).

4.3 Derived variates

In this section we will prove that formula (14) also holds for *derived variates*. y^R -variates are derived variates if they are functions of the observables y . Let us assume that this functional relationship takes the form:

$$(20) \quad \underset{p \times 1}{y^R} = \underset{p \times m}{\Lambda} \underset{m \times 1}{y} .$$

Then also:

$$(21) \quad \hat{y}^R = \Lambda \hat{y} .$$

Subtraction of (21) from (20) gives:

$$(22) \quad \underline{\hat{e}}^R = \Lambda \underline{\hat{e}} .$$

What remains to be shown is now that Λ equals $Q_{Ry} Q_y^{-1}$. Application of the *propagation law of covariances* to (20) gives:

$$Q_{Ry} = \Lambda Q_y$$

or:

$$\Lambda = Q_{Ry} Q_y^{-1} .$$

This, together with (22) shows that formula (14) also holds for derived variates.

4.4 Constituent variates

In this section we will prove that formula (14) holds for *constituent variates*. Constituent variates are y^R -variates of which the y -variates are functions. Let us assume that this functional relationship takes the form:

$$(23) \quad \underset{m \times 1}{y} = \underset{m \times p}{\Lambda} \underset{p \times 1}{y^R} .$$

Application of the propagation law of variances gives:

$$(24) \quad Q_y = \Lambda Q_R \Lambda^* ,$$

and application of the propagation law of covariances gives:

$$(25) \quad Q_{Ry} = Q_R \Lambda^* .$$

As we know $\hat{\underline{e}} = \underline{y} - \hat{\underline{y}}$ follows from solving the model:

$$(26) \quad \mathbf{B}^* E\{\underline{y}\} = \mathbf{0} ; E\{(\underline{y} - E\{\underline{y}\})(\underline{y} - E\{\underline{y}\})^*\} = Q_y$$

as:

$$(27) \quad \hat{\underline{e}} = Q_y \mathbf{B} (\mathbf{B}^* Q_y \mathbf{B})^{-1} \mathbf{B}^* \underline{y} .$$

With (23) we may formulate the model with condition equations also as:

$$(28) \quad \mathbf{B}^* \Lambda E\{\underline{y}^R\} = \mathbf{0} ; E\{(\underline{y}^R - E\{\underline{y}^R\})(\underline{y}^R - E\{\underline{y}^R\})^*\} = Q_R .$$

When compared with (26) this means that $\mathbf{B}$ of (26) is replaced by $\Lambda^* \mathbf{B}$, that $\underline{y}$ is replaced by $\underline{y}^R$ and that Q_y is replaced by Q_R . Instead of (27) we therefore get for model (28):

$$(29) \quad \hat{\underline{e}}^R = Q_R \Lambda^* \mathbf{B} (\mathbf{B}^* \Lambda Q_R \Lambda^* \mathbf{B})^{-1} \mathbf{B}^* \Lambda \underline{y}^R .$$

With (23), (24) and (25) this can also be written as:

$$\hat{\underline{e}}^R = Q_{Ry} \mathbf{B} (\mathbf{B}^* Q_y \mathbf{B})^{-1} \mathbf{B}^* \underline{y}$$

or with (27) as:

$$\hat{\underline{e}}^R = Q_{Ry} Q_y^{-1} \hat{\underline{e}} ,$$

which is identical to (14).

Let us, as an example of the above, "rederive" the solution of the model with condition equations. The model with condition equations reads:

$$(30) \quad \mathbf{B}^* E\{\underline{y}\} = \mathbf{0} ; E\{(\underline{y} - E\{\underline{y}\})(\underline{y} - E\{\underline{y}\})^*\} = Q_y .$$

If we define the misclosure vector as:

$$(31) \quad \underline{t} = \mathbf{B}^* \underline{y} ,$$

then (30) can also be written as:

$$(32) \quad I^* E\{\underline{t}\} = \mathbf{0} ; E\{(\underline{t} - E\{\underline{t}\})(\underline{t} - E\{\underline{t}\})^*\} = Q_t .$$

This model is also in the form of condition equations. The coefficient matrix of the condition equations is in this case the unit matrix I . If we solve for model (32) we get for $\hat{\underline{e}}_t = \underline{t} - \hat{\underline{t}}$:

$$(33) \quad \hat{\underline{e}}_t = Q_t I (I^* Q_t I)^{-1} I^* \underline{t} = \underline{t} .$$

We can now use our formula (14) in order to derive $\hat{\underline{e}} = \underline{y} - \hat{\underline{y}}$. This is done by interpreting $\underline{y}$ of (31) as an $\underline{y}^R$ -variate and $\underline{t}$ of (31) as an $\underline{y}$ -variate. Application of formula (14) gives then:

$$(34) \quad \hat{\underline{e}} = \underline{Q}_{yt} \underline{Q}_t^{-1} \hat{\underline{e}}_t .$$

Application of the propagation law of variances to (31) gives:

$$(35) \quad \underline{Q}_t = \underline{B}^* \underline{Q}_y \underline{B} ,$$

and application of the propagation law of covariances to (31) gives:

$$(36) \quad \underline{Q}_{yt} = \underline{Q}_y \underline{B} .$$

Substitution of (33), (35) and (36) into (34) gives with (31) then finally:

$$\hat{\underline{e}} = \underline{Q}_y \underline{B} (\underline{B}^* \underline{Q}_y \underline{B})^{-1} \underline{B}^* \underline{y} .$$

Hence, we have obtained our well-known and familiar result again.

5 Mixed model representations

5.1 Introduction

In chapter 2 we developed the theory of linear estimation for the model with *observation equations*:

$$(1) \quad E\{\underline{y}\} = \underline{A}x \quad ; \quad E\{(\underline{y}-\underline{A}x)(\underline{y}-\underline{A}x)^*\} = \underline{Q}_y.$$

It was shown in chapter 2 that the estimator $\hat{\underline{x}}$ of x follows from solving the system of *normal equations*:

$$(2) \quad (\underline{A}^* \underline{Q}_y^{-1} \underline{A}) \hat{\underline{x}} = \underline{A}^* \underline{Q}_y^{-1} \underline{y}$$

as:

$$(3) \quad \hat{\underline{x}} = (\underline{A}^* \underline{Q}_y^{-1} \underline{A})^{-1} \underline{A}^* \underline{Q}_y^{-1} \underline{y}.$$

From this result the estimators $\hat{\underline{y}} = \underline{A} \hat{\underline{x}}$ and $\hat{\underline{e}} = \underline{y} - \hat{\underline{y}}$ follow as:

$$(4) \quad \hat{\underline{y}} = \underline{A} (\underline{A}^* \underline{Q}_y^{-1} \underline{A})^{-1} \underline{A}^* \underline{Q}_y^{-1} \underline{y}$$

and:

$$(5) \quad \hat{\underline{e}} = [I - \underline{A} (\underline{A}^* \underline{Q}_y^{-1} \underline{A})^{-1} \underline{A}^* \underline{Q}_y^{-1}] \underline{y}.$$

In chapter 3 we developed the theory of linear estimation for the model with *condition equations*:

$$(6) \quad \underline{B}^* E\{\underline{y}\} = \underline{0} \quad ; \quad E\{(\underline{y} - E\{\underline{y}\})(\underline{y} - E\{\underline{y}\})^*\} = \underline{Q}_y.$$

In this case the estimators $\hat{\underline{y}}$ and $\hat{\underline{e}} = \underline{y} - \hat{\underline{y}}$ can be represented in terms of the matrix $\underline{B}$ as:

$$(7) \quad \hat{\underline{y}} = [I - \underline{Q}_y \underline{B} (\underline{B}^* \underline{Q}_y \underline{B})^{-1} \underline{B}^*] \underline{y}$$

and:

$$(8) \quad \hat{\underline{e}} = \underline{Q}_y \underline{B} (\underline{B}^* \underline{Q}_y \underline{B})^{-1} \underline{B}^* \underline{y}.$$

In the present chapter we will introduce two additional model representations. The first new model representation we will consider takes the form:

$$(9) \quad \underset{\substack{\underline{b} \times m \quad m \times 1}}{\underline{B}^*} E\{\underset{\substack{\underline{b} \times n \quad n \times 1}}{\underline{y}\}} = \underset{\substack{\underline{b} \times n \quad n \times 1}}{\underline{A}} \underset{\substack{m \times 1}}{x} \quad ; \quad E\{(\underline{y} - E\{\underline{y}\})(\underline{y} - E\{\underline{y}\})^*\} = \underset{\substack{m \times m}}{\underline{Q}_y}.$$

In this model representation x is an $n \times 1$ vector of unknown parameters and $\underline{y}$ is an $m \times 1$ vector of observables. Note that (9) is a mixture of (1) and (6). If matrix $\underline{B}$ of (9) equals the identity matrix, then (9) reduces to the model with observation equations (1). If on the other hand matrix

A of (9) equals the zero matrix, then (9) reduces to the model with condition equations (6). In the next section the solution of (9) will be derived.

The second new model representation that we will consider takes the form:

$$(10) \quad \begin{array}{c} E\{\underline{y}\} = \underset{m \times 1}{A} \underset{m \times n}{x} \underset{n \times 1}{}, \underset{b \times n}{B}^* \underset{n \times 1}{x} = \underset{b \times 1}{0} ; E\{(\underline{y} - E\{\underline{y}\})(\underline{y} - E\{\underline{y}\})^*\} = \underset{m \times m}{Q_y}. \end{array}$$

This representation is in the form of observation equations with conditions on the parameter vector x . If $B=0$ then (10) reduces simply to (1). The solution of (10) will be derived in section 5.3.

5.2 The model representation $B^*E\{\underline{y}\}=Ax$

Our starting point is the representation:

$$(11) \quad \begin{array}{c} B^* E\{\underline{y}\} = \underset{b \times m}{A} \underset{m \times 1}{x} ; E\{(\underline{y} - E\{\underline{y}\})(\underline{y} - E\{\underline{y}\})^*\} = \underset{m \times m}{Q_y}. \end{array}$$

We will assume that both matrices A and B have full column rank. Thus:

$$(12) \quad \text{rank } A = n \quad \text{and} \quad \text{rank } B = b .$$

Note that matrix A can only have full column rank if $n \leq b$: the number of unknown parameters should not exceed the number of conditions on the observations.

There are different derivations possible for deriving the estimators $\hat{x}$, $\hat{y}$ and $\hat{\underline{e}}$ for model (11). One possibility would be to transform (11) into a form with *observation equations only*. This can be done as follows. Consider the system:

$$(13) \quad \begin{array}{c} B^* E\{\underline{y}\} = \underset{b \times m}{A} \underset{m \times 1}{x} \end{array}$$

as an inhomogeneous system of linear equations with the unknown vector $E\{\underline{y}\}$. Vector $E\{\underline{y}\}$ can then be written as the sum of a particular solution of (13) and the homogeneous solution of:

$$(14) \quad \begin{array}{c} B^* E\{\underline{y}\} = \underset{b \times m}{0} . \end{array}$$

If we denote the $m \times (m-b)$ matrix of which the column vectors are orthogonal to the column vectors of matrix B by $B^\perp$, then:

$$\begin{array}{c} B^* B^\perp = \underset{b \times (m-b)}{0} . \end{array}$$

From this it follows that the solution of (14) takes the form:

$$(15) \quad E\{\underline{y}\}_h = \begin{matrix} \mathbf{B}^\perp & \lambda \\ m \times 1 & m \times (m-b) \quad (m-b) \times 1 \end{matrix}.$$

This is the homogeneous solution. A particular solution of (13) is:

$$(16) \quad E\{\underline{y}\}_p = \begin{matrix} \mathbf{B}(\mathbf{B}^* \mathbf{B})^{-1} \mathbf{A} & x \\ m \times 1 & m \times n \quad n \times 1 \end{matrix} \quad (\text{verify !})$$

Hence, the solution for $E\{\underline{y}\}$ of (13) takes the form of the sum of the particular solution (16) and the homogeneous solution (15):

$$(17) \quad E\{\underline{y}\} = \mathbf{B}(\mathbf{B}^* \mathbf{B})^{-1} \mathbf{A} x + \mathbf{B}^\perp \lambda = [\mathbf{B}(\mathbf{B}^* \mathbf{B})^{-1} \mathbf{A} ; \mathbf{B}^\perp] \begin{pmatrix} x \\ \lambda \end{pmatrix}.$$

Note that this representation is in terms of observation equations only. Since (17) is completely equivalent to the first equation of (11), we can now apply our solution method for observation equations to (17) in order to get the solution of (11).

Another possibility to solve for (11) would be to transform (11) into a form with *condition equations only*. This can be done as follows. Consider again the system (13). If we denote the $b \times (b-n)$ matrix of which the column vectors are orthogonal to the column vectors of matrix $\mathbf{A}$ by $\mathbf{A}^\perp$, then

$$\begin{matrix} \mathbf{A}^{\perp *} & \mathbf{A} & = & \mathbf{0} \\ (b-n) \times b & b \times n & & (b-n) \times n \end{matrix}.$$

From this it follows that pre-multiplication of (13) with $\mathbf{A}^{\perp *}$ gives:

$$(18) \quad \begin{matrix} (\mathbf{B} \mathbf{A}^\perp)^* & E\{\underline{y}\} & = & \mathbf{0} \\ (b-n) \times m & m \times 1 & & (b-n) \times 1 \end{matrix}.$$

Note that this representation is in terms of condition equations only. Since (18) is completely equivalent to the first equation of (11), we can now apply our solution method for condition equations to (18) in order to get the solution of (11).

Instead of following the above two mentioned approaches for solving (11), we will make use of a third approach. This approach makes use of the concept of $\underline{y}^R$ -variates. Define a $b \times 1$ vector $\underline{t}$ as:

$$(19) \quad \begin{matrix} \underline{t} & = & \mathbf{B}^* & \underline{y} \\ b \times 1 & & b \times m & m \times 1 \end{matrix}.$$

Substitution into (11) gives the representation:

$$(20) \quad \begin{matrix} E\{\underline{t}\} & = & \mathbf{A} & x \\ b \times 1 & & b \times n & n \times 1 \end{matrix} ; \quad \begin{matrix} E\{(\underline{t} - E\{\underline{t}\})(\underline{t} - E\{\underline{t}\})^*\} & = & \mathbf{Q}_t \\ & & b \times b & b \times b \end{matrix}.$$

This representation is in the form of observation equations only. Hence, we know how to compute the estimators $\hat{\underline{x}}$, $\hat{\underline{t}}$ and $\hat{\underline{e}}$. The estimator $\hat{\underline{x}}$ follows as:

$$(21) \quad \hat{\underline{x}} = (\mathbf{A}^* \mathbf{Q}_t^{-1} \mathbf{A})^{-1} \mathbf{A}^* \mathbf{Q}_t^{-1} \underline{t} .$$

Application of the propagation law for variances to (19) gives for $\mathbf{Q}_t$:

$$(22) \quad \mathbf{Q}_t = \mathbf{B}^* \mathbf{Q}_y \mathbf{B} .$$

Substitution of (19) and (22) into (21) gives then for the estimator $\hat{\underline{x}}$ of $\underline{x}$:

$$(23) \quad \hat{\underline{x}} = [\mathbf{A}^* (\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{A}]^{-1} \mathbf{A}^* (\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{B}^* \underline{y} .$$

The estimators $\hat{\underline{t}}$ and $\hat{\underline{e}}$ follow simply as:

$$(24) \quad \begin{cases} \hat{\underline{t}} = \mathbf{A} \hat{\underline{x}} \\ \hat{\underline{e}}_t = \underline{t} - \hat{\underline{t}} = \underline{t} - \mathbf{A} \hat{\underline{x}} \end{cases} .$$

Now, in order to obtain $\hat{\underline{e}}_y = \underline{y} - \hat{\underline{y}}$, note that $\underline{y}$ of (19) can be interpreted as a constituent $\underline{y}^R$ -variate.

Application of formula (14) of chapter 4 gives then:

$$(25) \quad \hat{\underline{e}}_y = \mathbf{Q}_{yt} \mathbf{Q}_t^{-1} \hat{\underline{e}}_t .$$

The matrix $\mathbf{Q}_{yt}$ follows from application of the propagation law of covariances to (19) as:

$$(26) \quad \mathbf{Q}_{yt} = \mathbf{Q}_y \mathbf{B} .$$

The estimators $\hat{\underline{e}}_y$ and $\hat{\underline{y}} = \underline{y} - \hat{\underline{e}}_y$ follow then finally with (22), (23), (24), (25), and (26) as:

$$(27) \quad \begin{cases} \hat{\underline{e}}_y = \mathbf{Q}_y \mathbf{B} (\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} [I - \mathbf{A} [\mathbf{A}^* (\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{A}]^{-1} \mathbf{A}^* (\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1}] \mathbf{B}^* \underline{y} \\ \hat{\underline{y}} = \underline{y} - \mathbf{Q}_y \mathbf{B} (\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} [I - \mathbf{A} [\mathbf{A}^* (\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{A}]^{-1} \mathbf{A}^* (\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1}] \mathbf{B}^* \underline{y} \end{cases} .$$

5.3 The model representation $E\{\underline{y}\} = \mathbf{A}\underline{x}$, $\mathbf{B}^*\underline{x} = 0$

Our starting point is the representation:

$$(29) \quad \begin{matrix} E\{\underline{y}\} = \mathbf{A} & \underline{x} & , & \mathbf{B}^* & \underline{x} & = & \mathbf{0} & ; & E\{(\underline{y} - E\{\underline{y}\})(\underline{y} - E\{\underline{y}\})^*\} = & \mathbf{Q}_y . \\ m \times 1 & m \times n & n \times 1 & b \times n & n \times 1 & b \times 1 & & m \times m & m \times m \end{matrix}$$

This representation is in the form of observation equations with conditions on the parameter vector $\underline{x}$. In order to derive the estimators $\hat{\underline{x}}$, $\hat{\underline{y}}$ and $\hat{\underline{e}}$ we will first transform (29) into a form with observation equations only. This is done by finding the parametric representation for:

$$(30) \quad \begin{matrix} \mathbf{B}^* & \mathbf{x} & = & \mathbf{0} \\ b \times n & n \times 1 & & b \times 1 \end{matrix} .$$

Let us denote the $n \times (n-b)$ matrix of which the column vectors are orthogonal to the column vectors of matrix B by $B^\perp$. Then:

$$(31) \quad \begin{matrix} \mathbf{B}^* & \mathbf{B}^\perp & = & \mathbf{0} \\ b \times n & n \times (n-b) & & b \times (n-b) \end{matrix} .$$

With (31) the parametric representation of the homogeneous system, (30) becomes:

$$(32) \quad \begin{matrix} \mathbf{x} & = & \mathbf{B}^\perp & \boldsymbol{\lambda} \\ n \times 1 & & n \times (n-b) & (n-b) \times 1 \end{matrix} .$$

Thus (32) is completely equivalent to (30). It is remarked that we assume that both the matrices A and B are of full column rank; thus $\text{rank } A = n$ and $\text{rank } B = b$. Note that matrix B can only have full column rank if $b \leq n$: one cannot impose more than n independent conditions on the n unknown parameters.

If we substitute (32) into the first equation of (29) we get:

$$(33) \quad \begin{matrix} E\{\mathbf{y}\} & = & \mathbf{A}\mathbf{B}^\perp & \boldsymbol{\lambda} & ; & E\{(\mathbf{y} - E\{\mathbf{y}\})(\mathbf{y} - E\{\mathbf{y}\})^*\} & = & \mathbf{Q}_y \\ m \times 1 & & m \times (n-b) & (n-b) \times 1 & & m \times m & & m \times m \end{matrix} .$$

This representation is completely equivalent to (29) and it is in the form of observation equations only. Thus we know how to compute the estimator $\hat{\boldsymbol{\lambda}}$ of $\boldsymbol{\lambda}$. This estimator reads:

$$(34) \quad \hat{\boldsymbol{\lambda}} = (\mathbf{B}^{\perp*} \mathbf{A}^* \mathbf{Q}_y^{-1} \mathbf{A} \mathbf{B}^\perp)^{-1} \mathbf{B}^{\perp*} \mathbf{A}^* \mathbf{Q}_y^{-1} \mathbf{y} .$$

Now let us consider the solution of (29) *without* the conditions on the parameter vector. This is our standard model representation with observation equations. As you know the solution of (29) without the conditions on the parameter vector satisfies the normal equations:

$$(35) \quad \mathbf{A}^* \mathbf{Q}_y^{-1} \mathbf{A} \hat{\mathbf{x}}_A = \mathbf{A}^* \mathbf{Q}_y^{-1} \mathbf{y} .$$

We have given $\hat{\mathbf{x}}_A$ of (35) the subscript A to emphasize that $\hat{\mathbf{x}}_A$ is *not* the solution of (29) but of:

$$(36) \quad E\{\mathbf{y}\} = \mathbf{A}\mathbf{x} \quad ; \quad E\{(\mathbf{y} - E\{\mathbf{y}\})(\mathbf{y} - E\{\mathbf{y}\})^*\} = \mathbf{Q}_y .$$

The covariance matrix of $\hat{\mathbf{x}}_A$ reads:

$$(37) \quad Q_{\hat{x}_A} = (A^* Q_y^{-1} A)^{-1} .$$

With (35) and (37) we can now write (34) as:

$$(38) \quad \hat{\underline{x}} = (B^{\perp*} Q_{\hat{x}_A}^{-1} B^{\perp})^{-1} B^{\perp*} Q_{\hat{x}_A}^{-1} \hat{\underline{x}}_A .$$

This result together with (32) gives the estimator of x for model (29):

$$(39) \quad \hat{\underline{x}} = B^{\perp} (B^{\perp*} Q_{\hat{x}_A}^{-1} B^{\perp})^{-1} B^{\perp*} Q_{\hat{x}_A}^{-1} \hat{\underline{x}}_A ,$$

If we look at the two formulae (38) and (39) we note that they take a very familiar form. And indeed, (38) and (39) are the solution of the model:

$$(40) \quad E\{\hat{\underline{x}}_A\} = B^{\perp} \lambda \quad ; \quad E\{(\hat{\underline{x}}_A - E\{\hat{\underline{x}}_A\})(\hat{\underline{x}}_A - E\{\hat{\underline{x}}_A\})^*\} = Q_{\hat{x}_A} .$$

This model representation is in terms of observation equations. The equivalent of (40) in terms of condition equations follows with (31) as:

$$(41) \quad B^* E\{\hat{\underline{x}}_A\} = 0 \quad ; \quad E\{(\hat{\underline{x}}_A - E\{\hat{\underline{x}}_A\})(\hat{\underline{x}}_A - E\{\hat{\underline{x}}_A\})^*\} = Q_{\hat{x}_A} .$$

As you know the solution of (41) is:

$$(42) \quad \hat{\underline{x}}_A = [I - Q_{\hat{x}_A} B(B^* Q_{\hat{x}_A} B)^{-1} B^*] \hat{\underline{x}}_A .$$

But since (41) is equivalent to (40) and (39) is the solution to (40), (42) must be identical to (39). This implies that the estimator $\hat{\underline{x}}$ of x for model (29) can be written as:

$$(43) \quad \hat{\underline{x}} = [I - Q_{\hat{x}_A} B(B^* Q_{\hat{x}_A} B)^{-1} B^*] \hat{\underline{x}}_A .$$

We thus have now found an expression for the estimator $\hat{\underline{x}}$ of x for model (29), which is in terms of the matrices A and B of (29). It follows then with (35), (37) and (43) that we can compute the estimators for model (29) as:

$$(44) \quad \boxed{\begin{aligned} \hat{\underline{x}}_A &= Q_{\hat{x}_A} A^* Q_y^{-1} \underline{y} \quad ; \quad Q_{\hat{x}_A} = (A^* Q_y^{-1} A)^{-1} \\ \hat{\underline{x}} &= [I - Q_{\hat{x}_A} B(B^* Q_{\hat{x}_A} B)^{-1} B^*] \hat{\underline{x}}_A \\ \hat{\underline{y}} &= A \hat{\underline{x}} \\ \hat{\underline{e}} &= \underline{y} - \hat{\underline{y}} \end{aligned}}$$

This result shows that the solution of (29) can be obtained in *two steps*. In the first step one considers (29) without the conditions on the parameters. That is, one solves for the model with *observation equations* (36). This gives the estimator $\hat{\underline{x}}_A$. Then in a second step one includes the conditions on the parameters. That is, one solves for the model with *condition equations* (41). This then finally gives the estimator $\hat{\underline{x}}$.

It follows from (44) that:

$$\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}} = \|\hat{\underline{e}}\|^2 = \|\underline{y} - A\hat{\underline{x}}_A + A Q_{\hat{\underline{x}}_A} B (B^* Q_{\hat{\underline{x}}_A} B)^{-1} B^* \hat{\underline{x}}_A\|^2,$$

or that:

$$\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}} = \|P_A^\perp \underline{y} + A Q_{\hat{\underline{x}}_A} B (B^* Q_{\hat{\underline{x}}_A} B)^{-1} B^* \hat{\underline{x}}_A\|^2.$$

Since $P_A^{\perp*} Q_y^{-1} A = 0$, this may also be written as:

$$\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}} = \|P_A^\perp \underline{y}\|^2 + \|A Q_{\hat{\underline{x}}_A} B (B^* Q_{\hat{\underline{x}}_A} B)^{-1} B^* \hat{\underline{x}}_A\|^2.$$

This gives finally:

(45)

$$\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}} = \underline{y}^* Q_y^{-1} P_A^\perp \underline{y} + \hat{\underline{x}}_A^* B (B^* Q_{\hat{\underline{x}}_A} B)^{-1} B^* \hat{\underline{x}}_A.$$

Note that the first term on the right hand side of (45) corresponds with the squared norm of the least-squares residual vector of the model $E\{\underline{y}\}=A\underline{x}$, whereas the second term on the right hand side of (45) corresponds with the squared norm of the least-squares residual vector of the model $B^* E\{\hat{\underline{x}}\} = 0$.

6 Partitioned model representations

6.1 Introduction

In chapters 2 and 3 the theory of linear estimation was developed for the model with *observation equations* and for the model with *condition equations*. Recall from chapter 2 that the estimators of the model:

$$(1) \quad \begin{matrix} E\{\underline{y}\} &= \mathbf{A} & \underline{x} &; \mathbf{D}\{\underline{y}\} &= \mathbf{Q}_y & \text{ } 1) \\ m \times 1 & m \times n & n \times 1 & m \times m & m \times m & \end{matrix}$$

are given as:

$$(2) \quad \left\{ \begin{array}{ll} \hat{\underline{x}} &= (\mathbf{A}^* \mathbf{Q}_y^{-1} \mathbf{A})^{-1} \mathbf{A}^* \mathbf{Q}_y^{-1} \underline{y} ; \quad \mathbf{Q}_{\hat{\underline{x}}} &= (\mathbf{A}^* \mathbf{Q}_y^{-1} \mathbf{A})^{-1} \\ \hat{\underline{y}} &= \mathbf{A} \hat{\underline{x}} = \mathbf{P}_A \underline{y} &; \quad \mathbf{Q}_{\hat{\underline{y}}} &= \mathbf{P}_A \mathbf{Q}_y \\ \hat{\underline{e}} &= \underline{y} - \hat{\underline{y}} = \mathbf{P}_A^\perp \underline{y} &; \quad \mathbf{Q}_{\hat{\underline{e}}} &= \mathbf{Q}_y - \mathbf{Q}_{\hat{\underline{y}}} = \mathbf{P}_A^\perp \mathbf{Q}_y \\ \hat{\underline{e}}^* \mathbf{Q}_y^{-1} \hat{\underline{e}} &= \underline{y}^* \mathbf{Q}_y^{-1} \mathbf{P}_A^\perp \underline{y} &; \quad \sigma_{\hat{\underline{e}}^* \mathbf{Q}_y^{-1} \hat{\underline{e}}}^2 &= 2(m-n). \end{array} \right.$$

And recall from chapter 3 that the estimators of the model:

$$(3) \quad \begin{matrix} \mathbf{B}^* E\{\underline{y}\} &= \mathbf{0} &; \mathbf{D}\{\underline{y}\} &= \mathbf{Q}_y \\ b \times m & m \times 1 & b \times 1 & m \times m & m \times m \end{matrix}$$

are given as:

$$(4) \quad \left\{ \begin{array}{ll} \hat{\underline{y}} &= [\mathbf{I} - \mathbf{Q}_y \mathbf{B} (\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{B}^*] \underline{y} = \mathbf{P}_{\mathbf{Q}_y \mathbf{B}}^\perp \underline{y} ; \quad \mathbf{Q}_{\hat{\underline{y}}} &= \mathbf{P}_{\mathbf{Q}_y \mathbf{B}}^\perp \mathbf{Q}_y \\ \hat{\underline{e}} &= \underline{y} - \hat{\underline{y}} &= \mathbf{P}_{\mathbf{Q}_y \mathbf{B}} \underline{y} ; \quad \mathbf{Q}_{\hat{\underline{e}}} &= \mathbf{Q}_y - \mathbf{Q}_{\hat{\underline{y}}} = \mathbf{P}_{\mathbf{Q}_y \mathbf{B}} \mathbf{Q}_y . \\ \hat{\underline{e}}^* \mathbf{Q}_y^{-1} \hat{\underline{e}} &= \underline{y}^* \mathbf{Q}_y^{-1} \mathbf{P}_{\mathbf{Q}_y \mathbf{B}} \underline{y} &; \quad \sigma_{\hat{\underline{e}}^* \mathbf{Q}_y^{-1} \hat{\underline{e}}}^2 &= 2b \end{array} \right.$$

In this chapter we will partition the matrices $\mathbf{A}$ and $\mathbf{B}$ of (1) and (3) respectively, *columnwise* and *rowwise*. The following four cases will therefore be considered in this chapter:

$$1^\circ \quad E\{\underline{y}\} = (\mathbf{A}_1 : \mathbf{A}_2) \begin{pmatrix} x_1 \\ x_2 \end{pmatrix},$$

¹⁾ The notation " $\mathbf{D}\{\underline{y}\}$ " will be used instead of the notation " $E\{(\underline{y}-\mathbf{A}\underline{x})(\underline{y}-\mathbf{A}\underline{x})^*\}$ ". The kernel " $\mathbf{D}$ " stands for "*Dispersion*".

$$2^{\circ} \quad E\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = \begin{pmatrix} A_1 \\ A_2 \end{pmatrix} x ,$$

$$3^{\circ} \quad \begin{pmatrix} B_1^* \\ B_2^* \end{pmatrix} E\{y\} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} ,$$

$$4^{\circ} \quad (B_1^* \ B_2^*) E\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = 0 .^{2)}$$

In this chapter it will be shown how the above partitioning of A and B can be carried through to the estimators of (2) and (4) respectively.

6.2 The partitioned model $E\{y\} = (A_1 : A_2) \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$.

Let us consider the partitioned model:

$$(5) \quad \begin{matrix} E\{y\} & = & (A_1 : A_2) & \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \\ m \times 1 & & m \times n_1 \ m \times n_2 & (n_1 + n_2) \times 1 \end{matrix} ; \quad \begin{matrix} D\{y\} & = & Q_y \\ m \times m & & m \times m \end{matrix} .$$

The corresponding partitioned system of *normal equations* reads:

$$(6) \quad \begin{pmatrix} A_1^* Q_y^{-1} A_1 & A_1^* Q_y^{-1} A_2 \\ A_2^* Q_y^{-1} A_1 & A_2^* Q_y^{-1} A_2 \end{pmatrix} \begin{pmatrix} \hat{x}_1 \\ \hat{x}_2 \end{pmatrix} = \begin{pmatrix} A_1^* Q_y^{-1} y \\ A_2^* Q_y^{-1} y \end{pmatrix} .$$

Now let us assume that for a particular application we are only interested in the estimator $\hat{x}_1$. There are two ways for finding $\hat{x}_1$. Either we solve the complete system of normal equations (6) and then look at the first n_1 -elements of the estimator $\hat{x}$. Or, we eliminate $\hat{x}_2$ from the first n_1 equations of (6) and solve for $\hat{x}_1$. This second approach is more efficient since it gets round the necessity of solving the complete system of normal equations. The estimator $\hat{x}_2$ gets eliminated from the first n_1 equations of (6), if we premultiply (6) with the *square* and *full rank* matrix:

$$(7) \quad \begin{pmatrix} I & -A_1^* Q_y^{-1} A_2 (A_2^* Q_y^{-1} A_2)^{-1} \\ 0 & I \end{pmatrix} .$$

Premultiplication of (6) with (7) gives:

²⁾ This case will not be dealt with explicitly in this chapter.

$$(8) \begin{pmatrix} \mathbf{A}_1^* \mathbf{Q}_y^{-1} [I - \mathbf{A}_2 (\mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_2)^{-1} \mathbf{A}_2^* \mathbf{Q}_y^{-1}] \mathbf{A}_1 & 0 \\ \mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_1 & \mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_2 \end{pmatrix} \begin{pmatrix} \hat{\mathbf{x}}_1 \\ \hat{\mathbf{x}}_2 \end{pmatrix} = \begin{pmatrix} \mathbf{A}_1^* \mathbf{Q}_y^{-1} [I - \mathbf{A}_2 (\mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_2)^{-1} \mathbf{A}_2^* \mathbf{Q}_y^{-1}] \mathbf{y} \\ \mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{y} \end{pmatrix}$$

In the first n_1 equations of (8) we recognize the expression for the orthogonal projector:

$$(9) \quad \mathbf{P}_{\mathbf{A}_2}^\perp = I - \mathbf{A}_2 (\mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_2)^{-1} \mathbf{A}_2^* \mathbf{Q}_y^{-1}.$$

Using the properties:

$$(10) \quad \mathbf{Q}_y^{-1} \mathbf{P}_{\mathbf{A}_2}^\perp = \mathbf{P}_{\mathbf{A}_2}^{\perp*} \mathbf{Q}_y^{-1} = \mathbf{P}_{\mathbf{A}_2}^{\perp*} \mathbf{Q}_y^{-1} \mathbf{P}_{\mathbf{A}_2}^\perp,$$

and the abbreviation:

$$(11) \quad \bar{\mathbf{A}}_1 = \mathbf{P}_{\mathbf{A}_2}^\perp \mathbf{A}_1,$$

we may therefore write (8) also as:

$$(12) \quad \begin{pmatrix} \bar{\mathbf{A}}_1^* \mathbf{Q}_y^{-1} \bar{\mathbf{A}}_1 & 0 \\ \mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_1 & \mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_2 \end{pmatrix} \begin{pmatrix} \hat{\mathbf{x}}_1 \\ \hat{\mathbf{x}}_2 \end{pmatrix} = \begin{pmatrix} \bar{\mathbf{A}}_1^* \mathbf{Q}_y^{-1} \mathbf{y} \\ \mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{y} \end{pmatrix}.^3$$

The solution for $\hat{\mathbf{x}}_1$ follows therefore as:

$$(13) \quad \hat{\mathbf{x}}_1 = (\bar{\mathbf{A}}_1^* \mathbf{Q}_y^{-1} \bar{\mathbf{A}}_1)^{-1} \bar{\mathbf{A}}_1^* \mathbf{Q}_y^{-1} \mathbf{y}.$$

Application of the propagation law of variances gives:

$$(14) \quad \mathbf{Q}_{\hat{\mathbf{x}}_1} = (\bar{\mathbf{A}}_1^* \mathbf{Q}_y^{-1} \bar{\mathbf{A}}_1)^{-1}.$$

Once $\hat{\mathbf{x}}_1$ is known, $\hat{\mathbf{x}}_2$ can be found from the last n_2 equations of (12). This gives:

$$(15) \quad \hat{\mathbf{x}}_2 = (\mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_2)^{-1} \mathbf{A}_2^* \mathbf{Q}_y^{-1} (\mathbf{y} - \mathbf{A}_1 \hat{\mathbf{x}}_1).$$

Application of the propagation law of variances gives with (13) and (14):

$$\mathbf{Q}_{\hat{\mathbf{x}}_2} = (\mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_2)^{-1} \mathbf{A}_2^* \mathbf{Q}_y^{-1} [\mathbf{Q}_y - \bar{\mathbf{A}}_1 \mathbf{Q}_{\hat{\mathbf{x}}_1} \bar{\mathbf{A}}_1^* - \mathbf{A}_1 \mathbf{Q}_{\hat{\mathbf{x}}_1} \bar{\mathbf{A}}_1^* + \mathbf{A}_1 \mathbf{Q}_{\hat{\mathbf{x}}_1} \mathbf{A}_1^*] \mathbf{Q}_y^{-1} \mathbf{A}_2 (\mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_2)^{-1}.$$

Since $\mathbf{A}_2^* \mathbf{Q}_y^{-1} \bar{\mathbf{A}}_1 = 0$, this reduces to:

$$(16) \quad \mathbf{Q}_{\hat{\mathbf{x}}_2} = (\mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_2)^{-1} + (\mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_2)^{-1} \mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_1 \mathbf{Q}_{\hat{\mathbf{x}}_1} \mathbf{A}_1^* \mathbf{Q}_y^{-1} \mathbf{A}_2 (\mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_2)^{-1}.$$

Summarizing, we thus have shown that:

³⁾ The normal matrix of this system is called a "reduced normal matrix". It is said to be reduced for $\mathbf{x}_2$.

$$\begin{aligned}
 (17) \quad & \hat{x}_1 = (\bar{A}_1^* Q_y^{-1} \bar{A}_1)^{-1} \bar{A}_1^* Q_y^{-1} y, \text{ with } \bar{A}_1 = P_{A_2}^\perp A_1, \\
 & Q_{\hat{x}_1} = (\bar{A}_1^* Q_y^{-1} \bar{A}_1)^{-1}, \\
 & \hat{x}_2 = (A_2^* Q_y^{-1} A_2)^{-1} A_2^* Q_y^{-1} (y - A_1 \hat{x}_1), \\
 & Q_{\hat{x}_2} = (A_2^* Q_y^{-1} A_2)^{-1} + (A_2^* Q_y^{-1} A_2)^{-1} A_2^* Q_y^{-1} A_1 Q_{\hat{x}_1} A_1^* Q_y^{-1} A_2 (A_2^* Q_y^{-1} A_2)^{-1}
 \end{aligned}$$

Instead of eliminating the estimator $\hat{x}_2$ from the first n_1 equations of (6), one may also eliminate the estimator $\hat{x}_1$ from the last n_2 equations of (6). This is achieved by premultiplying (6) with the *square* and *full rank* matrix:

$$(18) \quad \begin{pmatrix} I & 0 \\ -A_2^* Q_y^{-1} A_1 & (A_1^* Q_y^{-1} A_1)^{-1} I \end{pmatrix}.$$

This gives analogous to (12) the equations:

$$(19) \quad \begin{pmatrix} A_1^* Q_y^{-1} A_1 & A_1^* Q_y^{-1} A_2 \\ 0 & \bar{A}_2^* Q_y^{-1} \bar{A}_2 \end{pmatrix} \begin{pmatrix} \hat{x}_1 \\ \hat{x}_2 \end{pmatrix} = \begin{pmatrix} A_1^* Q_y^{-1} y \\ \bar{A}_2^* Q_y^{-1} y \end{pmatrix},$$

with

$$(20) \quad \bar{A}_2 = P_{A_1}^\perp A_2.$$

From (19) the estimators $\hat{x}_1$ and $\hat{x}_2$ and their variance matrices follow as:

$$\begin{aligned}
 (21) \quad & \hat{x}_1 = (A_1^* Q_y^{-1} A_1)^{-1} A_1^* Q_y^{-1} (y - A_2 \hat{x}_2) \\
 & Q_{\hat{x}_1} = (A_1^* Q_y^{-1} A_1)^{-1} + (A_1^* Q_y^{-1} A_1)^{-1} A_1^* Q_y^{-1} A_2 Q_{\hat{x}_2} A_2^* Q_y^{-1} A_1 (A_1^* Q_y^{-1} A_1)^{-1} \\
 & \hat{x}_2 = (\bar{A}_2^* Q_y^{-1} \bar{A}_2)^{-1} \bar{A}_2^* Q_y^{-1} y, \text{ with } \bar{A}_2 = P_{A_1}^\perp A_2 \\
 & Q_{\hat{x}_2} = (\bar{A}_2^* Q_y^{-1} \bar{A}_2)^{-1}
 \end{aligned}$$

Compare (21) with (17) and note that (21) follows from (17) by interchanging the role of x_1 and x_2 . The above results can now be used to investigate what happens to the precision of the estimators when a model with observation equations is enlarged with additional parameters. Let us assume that our original model is given as:

$$(22) \quad \begin{matrix} E\{y\} = & A_1 & x_1 & ; & D\{y\} = & Q_y \\ m \times 1 & m \times n_1 & n_1 \times 1 & & m \times m & m \times m \end{matrix}$$

The corresponding estimator and its variance matrix will be denoted by $\hat{x}_1^1$ and $Q_{\hat{x}_1^1}$ respectively. Then:

$$(23) \quad \begin{cases} \hat{x}_1^1 = (A_1^* Q_y^{-1} A_1)^{-1} A_1^* Q_y^{-1} y \\ Q_{\hat{x}_1^1} = (A_1^* Q_y^{-1} A_1)^{-1} \end{cases}.$$

Suppose now that model (22) is enlarged to :

$$(24) \quad E\{\underline{y}\} = \begin{matrix} m \times 1 \\ \end{matrix} = \begin{matrix} m \times n_1 & m \times n_2 \\ \mathbf{A}_1 & \mathbf{A}_2 \end{matrix} \begin{matrix} n_1 \times 1 \\ n_2 \times 1 \\ \end{matrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} ; \quad \begin{matrix} m \times m \\ \mathbf{D}\{\underline{y}\} = \mathbf{Q}_y \end{matrix} .$$

Then according to (17):

$$(25) \quad \begin{cases} \hat{x}_1 = (\bar{\mathbf{A}}_1^* \mathbf{Q}_y^{-1} \bar{\mathbf{A}}_1)^{-1} \bar{\mathbf{A}}_1^* \mathbf{Q}_y^{-1} \underline{y} \\ \mathbf{Q}_{\hat{x}_1} = (\bar{\mathbf{A}}_1^* \mathbf{Q}_y^{-1} \bar{\mathbf{A}}_1)^{-1} \end{cases} .$$

Note the resemblance between (23) and (25). The two estimators $\hat{x}_1^1$ and $\hat{x}_1$ are identical if and only if $\mathbf{A}_1 = \bar{\mathbf{A}}_1$, that is if:

$$(26) \quad \mathbf{A}_1 = \mathbf{P}_{\mathbf{A}_2}^\perp \mathbf{A}_1 .$$

This is the case if the range space of $\mathbf{A}_1$ is *orthogonal* to the range space of $\mathbf{A}_2$, that is if $R(\mathbf{A}_1) \perp R(\mathbf{A}_2)$ or $\mathbf{A}_1^* \mathbf{Q}_y^{-1} \mathbf{A}_2 = 0$. Note that in this case the partitioned normal matrix of (6) becomes *blockdiagonal*. We may thus conclude that the estimator of x_1 is not affected by an enlargement of the model if and only if $R(\mathbf{A}_1)$ is orthogonal to $R(\mathbf{A}_2)$. In order to find out what happens when $R(\mathbf{A}_1)$ and $R(\mathbf{A}_2)$ are not orthogonal to each other, we take instead of the expression $\mathbf{Q}_{\hat{x}_1} = (\bar{\mathbf{A}}_1^* \mathbf{Q}_y^{-1} \bar{\mathbf{A}}_1)^{-1}$, the following expression from (21):

$$\mathbf{Q}_{\hat{x}_1} = (\mathbf{A}_1^* \mathbf{Q}_y^{-1} \mathbf{A}_1)^{-1} + (\mathbf{A}_1^* \mathbf{Q}_y^{-1} \mathbf{A}_1)^{-1} \mathbf{A}_1^* \mathbf{Q}_y^{-1} \mathbf{A}_2 \mathbf{Q}_{\hat{x}_2} \mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_1 (\mathbf{A}_1^* \mathbf{Q}_y^{-1} \mathbf{A}_1)^{-1} .$$

With (23) this can be written as:

$$(27) \quad \mathbf{Q}_{\hat{x}_1} = \mathbf{Q}_{\hat{x}_1^1} + \mathbf{Q}_{\hat{x}_1^1} \mathbf{A}_1^* \mathbf{Q}_y^{-1} \mathbf{A}_2 \mathbf{Q}_{\hat{x}_2} \mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_1 \mathbf{Q}_{\hat{x}_1^1} .$$

Now suppose we are interested in a particular linear function, $\theta = f^* x_1$, of x_1 . With model (22) the estimator of θ reads $\hat{\theta}^1 = f^* \hat{x}_1^1$ and with model (24) we get as estimator of θ , $\hat{\theta} = f^* \hat{x}_1$. Their variances read:

$$\sigma_{\hat{\theta}^1}^2 = f^* \mathbf{Q}_{\hat{x}_1^1} f \text{ and } \sigma_{\hat{\theta}}^2 = f^* \mathbf{Q}_{\hat{x}_1} f .$$

With (27) this gives:

$$\sigma_{\hat{\theta}}^2 = \sigma_{\hat{\theta}^1}^2 + f^* \mathbf{Q}_{\hat{x}_1^1} \mathbf{A}_1^* \mathbf{Q}_y^{-1} \mathbf{A}_2 \mathbf{Q}_{\hat{x}_2} \mathbf{A}_2^* \mathbf{Q}_y^{-1} \mathbf{A}_1 \mathbf{Q}_{\hat{x}_1^1} f .$$

Since the second term on the right hand side is always non-negative, it follows that:

$$(28) \quad \sigma_{\hat{\theta}}^2 \geq \sigma_{\hat{\theta}^1}^2 .$$

This important result shows that the precision of the estimators generally gets *poorer* if more parameters are added to the linear model.

Let us now investigate the variance matrix of the estimator of the added parameter vector x_2 . According to (21), $Q_{\hat{x}_2}$ is the inverse of the matrix $\bar{A}_2^* Q_y^{-1} \bar{A}_2$; see also (19). With $\bar{A}_2 = P_{A_1}^\perp A_2$, we may write:

$$(29) \quad \bar{A}_2^* Q_y^{-1} \bar{A}_2 = A_2^* Q_y^{-1} A_2 - A_2^* Q_y^{-1} P_{A_1} A_2 .$$

Since $P_{A_1} A_2 = A_2$ if and only if $R(A_2) \subset R(A_1)$, it follows that $\bar{A}_2^* Q_y^{-1} \bar{A}_2 = 0$ if the column vectors of A_2 are linear dependent of the column vectors of A_1 . Thus, $Q_{\hat{x}_2}$ does *not* exist and x_2 is *not* estimable if $R(A_2) \subset R(A_1)$. In this case matrix $(A_1 : A_2)$ also fails to be of full rank. Now assume that x_2 is a scalar. Then $n_2=1$ and A_2 is a vector. With the cosine rule it follows then that:

$$(30) \quad A_2^* Q_y^{-1} P_{A_1} A_2 = \|A_2\| \|P_{A_1} A_2\| \cos \alpha ,$$

see figure 6.1. But we also have that:

$$(31) \quad A_2^* Q_y^{-1} P_{A_1} A_2 = (P_{A_1} A_2)^* Q_y^{-1} (P_{A_1} A_2) = \|P_{A_1} A_2\|^2 .$$

From (30) and (31) it follows that:

$$(32) \quad A_2^* Q_y^{-1} P_{A_1} A_2 = \|A_2\|^2 \cos^2 \alpha = A_2^* Q_y^{-1} A_2 \cos^2 \alpha .$$

This together with (29) shows that:

$$(33) \quad \bar{A}_2^* Q_y^{-1} \bar{A}_2 = A_2^* Q_y^{-1} A_2 \sin^2 \alpha .$$

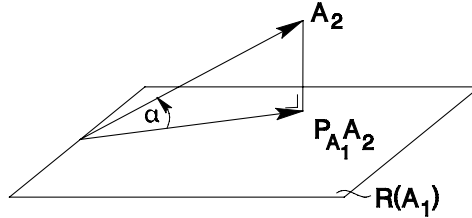


Figure 6.1: $A_2^* Q_y^{-1} P_{A_1} A_2 = \|A_2\| \|P_{A_1} A_2\| \cos \alpha$

From this it follows that:

$$(34) \quad \sigma_{\hat{x}_2}^2 = (A_2^* Q_y^{-1} A_2)^{-1} \sin^2 \alpha .$$

The angle α measures the *degree of dependence* between A_2 and A_1 . If $\alpha = \frac{1}{2}\pi$, then A_2 is orthogonal to $R(A_1)$ and $\sigma_{\hat{x}_2}^2 = (A_2^* Q_y^{-1} A_2)^{-1}$. If $\alpha=0$, then $A_2 \in R(A_1)$ and $\sigma_{\hat{x}_2}^2 = \infty$. If α is small, then $\sigma_{\hat{x}_2}^2$ is large and x_2 is poorly estimable. Since (see(33)):

$$(35) \quad \sin^2 \alpha = \frac{\bar{A}_2^* Q_y^{-1} \bar{A}_2}{A_2^* Q_y^{-1} A_2} ,$$

the degree of dependence between A_2 and A_1 can be computed in a rather straight forward manner from the ratio of the diagonal elements of the normal matrix *after* and *before* reduction. Compare (19) and (6). The computation of the scalar $\sin^2\alpha$ is usually included in the software in order to diagnose and detect possible singularities or rankdefects in the normal matrix.

So far we have been concerned with the estimators $\hat{x}_1$ and $\hat{x}_2$. Let us now consider the estimator $\hat{y}$. Since $\hat{y} = A_1 \hat{x}_1 + A_2 \hat{x}_2$, it follows from (17) that:

$$\begin{aligned}\hat{y} &= A_1 \hat{x}_1 + A_2 (A_2^* Q_y^{-1} A_2)^{-1} A_2^* Q_y^{-1} (y - A_1 \hat{x}_1) \\ &= A_1 \hat{x}_1 + P_{A_2} (y - A_1 \hat{x}_1) \\ &= P_{A_2}^\perp A_1 \hat{x}_1 + P_{A_2} y \\ &= \bar{A}_1 \hat{x}_1 + P_{A_2} y \\ &= \bar{A}_1 (\bar{A}_1^* Q_y^{-1} \bar{A}_1)^{-1} \bar{A}_1^* Q_y^{-1} y + P_{A_2} y\end{aligned}$$

or:

(36)

$$\hat{y} = P_{\bar{A}_1} y + P_{A_2} y$$

Since $\hat{y} = P_A y$, with $A = (A_1 : A_2)$, we have with (36) the following orthogonal decomposition of the orthogonal projector P_A :

(37)

$$P_A = P_{\bar{A}_1} + P_{A_2}$$

In a completely analogous way it follows from (21) that:

(38)

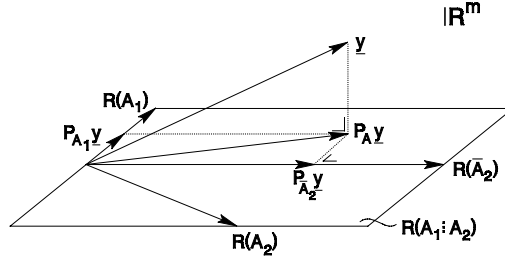
$$\hat{y} = P_{A_1} y + P_{\bar{A}_2} y$$

and:

(39)

$$P_A = P_{A_1} + P_{\bar{A}_2}$$

The above orthogonal decomposition is shown in figure 6.2.

Figure 6.2: $P_A y = P_{A_1} y + P_{A_2} y$

Since $\hat{e} = (I - P_A)y = P_A^\perp y$, it follows with (39) that:

$$(40) \quad \hat{e} = P_A^\perp y = (I - P_{A_1} - P_{A_2})y = P_{A_1}^\perp y - P_{A_2} y = \hat{e}^1 - P_{A_2} y,$$

where $\hat{e}^1$ is the least-squares residual vector of model (22). From (40) it follows that:

$$(41) \quad \hat{e}^* Q_y^{-1} \hat{e} = \hat{e}^{1*} Q_y^{-1} \hat{e}^1 - 2 \hat{e}^* Q_y^{-1} P_{A_2} y + y^* Q_y^{-1} P_{A_2} y.$$

$$\text{But: } \hat{e}^{1*} Q_y^{-1} P_{A_2} y = y^* P_{A_1}^\perp Q_y^{-1} P_{A_2} y = y^* Q_y^{-1} P_{A_1}^\perp P_{A_2} y = y^* Q_y^{-1} P_{A_2} y.$$

Hence, (41) can be written as:

$$\hat{e}^* Q_y^{-1} \hat{e} = \hat{e}^{1*} Q_y^{-1} \hat{e}^1 - y^* Q_y^{-1} P_{A_2} y,$$

or as:

$$(42) \quad \boxed{\|\hat{e}\|^2 = \|\hat{e}^1\|^2 - \|P_{A_2} y\|^2}$$

This important results shows that the length of the residual vector generally gets *smaller* if more parameters are added to the linear model. In the extreme case that the number of parameters equals the number of observations, that is $m=n_1+n_2$, matrix $A=(A_1:A_2)$ becomes square and regular, and P_A reduces to the identity matrix I . In this case (see(39)) $I = P_{A_1} + P_{A_2}$ or $P_{A_2} = P_{A_1}^\perp$, and therefore $\hat{e}^1 = P_{A_1}^\perp y = P_{A_2} y$. With (42) this gives $\|\hat{e}\|^2=0$. Thus the least-squares residual vector is identically zero if the *redundancy* is zero.

The results of this section will play an important role and are used repeatedly in testing theory.

Example 1:

The model that will be considered first is given as:

$$(43) \quad E \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_m \end{pmatrix}_{m \times 1} = \begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_m \end{pmatrix}_{m \times 1} x_1^1; D\{\underline{y}\} = \sigma^2 I_m \text{ }_{1 \times 1}.$$

Since the observation equations are of the form $E\{y_i\} = a_i x_1^1$, they describe the equation of a straight line through the origin with *slope* x_1^1 . This is shown in figure 6.3.

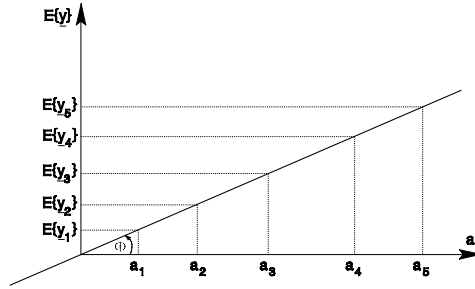


Figure 6.3: The line $E\{\underline{y}\} = ax_1^1$ with slope $x_1^1 = \tan \phi$

The least-squares estimate of x_1^1 follows from the minimization problem:

$$\min_{x_1^1} \frac{1}{\sigma^2} \sum_{i=1}^m (y_i - a_i x_1^1)^2.$$

Since $|y_i - a_i x_1^1|$ is the *vertical* distance from the point (a_i, y_i) to the straight line $E\{\underline{y}\} = ax_1^1$, the least-squares estimate $\hat{x}_1^1$ follows from a minimization of the sum of the squares of these vertical distances. The least-squares estimator $\hat{x}_1^1$ reads:

$$(44) \quad \hat{x}_1^1 = \sum_{i=1}^m a_i y_i / \sum_{i=1}^m a_i^2.$$

Its variance is given as:

$$(45) \quad \sigma_{\hat{x}_1^1}^2 = \left(\sum_{i=1}^m a_i^2 \right)^{-1} \sigma^2.$$

Let us now consider the enlarged linear model:

$$(46) \quad E \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_m \end{pmatrix} = \begin{pmatrix} a_1 & 1 \\ a_2 & 1 \\ \vdots & \vdots \\ a_m & 1 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}; \quad D\{\underline{y}\} = \sigma^2 I_m .$$

In this case the observation equations $E\{\underline{y}\} = ax_1 + x_2$ describe a straight line with *intercept* x_2 and *slope* x_1 . This is shown in figure 6.4.

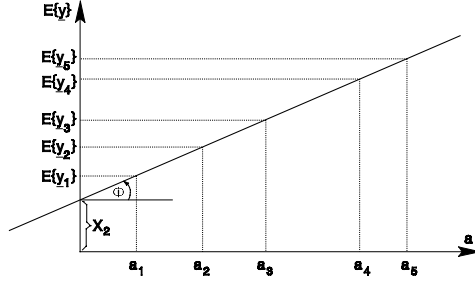


Figure 6.4: The line $E\{\underline{y}\} = ax_1 + x_2$ with slope $x_1 = \tan \phi$ and intercept x_2

With the definitions $A_1 = (a_1, a_2, \dots, a_m)^*$ and $A_2 = (1 \dots 1)^*$ it follows that:

$$(47) \quad \begin{cases} P_{A_2}^\perp = \begin{pmatrix} 1 & 0 \\ & 1 \\ & \ddots \\ 0 & 1 \end{pmatrix} - \begin{pmatrix} 1 \\ 1 \\ \vdots \\ 1 \end{pmatrix} \frac{1}{m} (1 \dots 1) , \\ \bar{A}_1 = P_{A_2}^\perp A_1 = \begin{pmatrix} a_1 \\ \vdots \\ a_m \end{pmatrix} - \begin{pmatrix} 1 \\ \vdots \\ 1 \end{pmatrix} \frac{1}{m} \sum_{i=1}^m a_i . \end{cases}$$

In order to simplify the expression for $\bar{A}$ somewhat, we define:

$$(48) \quad \begin{cases} a_c = \frac{1}{m} \sum_{i=1}^m a_i & \text{(average value of the } a_i \text{'s)} \\ \bar{a}_i = a_i - a_c & \text{(centered with respect to } a_c) \end{cases}$$

We now may write $\bar{A}$ as:

$$(49) \quad \bar{A}_1 = (\bar{a}_1, \bar{a}_2, \dots, \bar{a}_m)^* .$$

Using (17) and (49) the least-squares estimator $\hat{x}_1$ becomes:

$$(50) \quad \hat{x}_1 = \sum_{i=1}^m \bar{a}_i y_i / \sum_{i=1}^m \bar{a}_i^2 .$$

Compare (50) with (44). The variance of $\hat{x}_1$ is given as:

$$(51) \quad \sigma_{\hat{x}_1}^2 = (\sum_{i=1}^m \bar{a}_i^2)^{-1} \sigma^2 .$$

According to the theory of this section $\sigma_{\hat{x}_1}^2 \geq \sigma_{\hat{x}_1^*}^2$, or:

$$(52) \quad (\sum_{i=1}^m \bar{a}_i^2)^{-1} \geq (\sum_{i=1}^m a_i^2)^{-1}$$

should hold. This is easily verified. From (52) it follows that:

$$\sum_{i=1}^m a_i^2 \geq \sum_{i=1}^m \bar{a}_i^2$$

or:

$$\sum_{i=1}^m (a_i^2 - \bar{a}_i^2) \geq 0$$

or:

$$\sum_{i=1}^m (a_i^2 - (a_i - a_c)^2) \geq 0$$

or:

$$\sum_{i=1}^m (2a_i a_c - a_c^2) \geq 0$$

or:

$$2ma_c^2 - ma_c^2 \geq 0$$

or:

$$ma_c^2 \geq 0 .$$

Hence, inequality (52) holds indeed. Since $A_1 = (a_1, a_2, \dots, a_m)^*$ and $A_2 = (1 \dots 1)^*$, it follows with (50) from (17) that the least-squares estimator $\hat{x}_2$ is given as:

$$\hat{x}_2 = \frac{1}{m} \sum_{i=1}^m (y_i - a_i \hat{x}_1) ,$$

or as:

$$(53) \quad \hat{x}_2 = \frac{1}{m} \sum_{i=1}^m y_i - a_c \sum_{i=1}^m \bar{a}_i y_i / \sum_{i=1}^m \bar{a}_i^2 .$$

The variance of $\hat{x}_2$ follows from (17) as:

$$\sigma_{\hat{x}_2}^2 = \frac{\sigma^2}{m} + \frac{1}{m} \sum_{i=1}^m a_i^2 \sigma_{\hat{x}_1}^2 \sum_{i=1}^m a_i^2 \frac{1}{m} .$$

With (48) and (51) this gives:

$$(54) \quad \sigma_{\hat{x}_2}^2 = \sigma^2 \left(\frac{1}{m} + \frac{\mathbf{a}_c^2}{\sum_{i=1}^m \bar{\mathbf{a}}_i^2} \right).$$

6.3 The partitioned model $E \begin{pmatrix} \mathbf{y}_1 \\ \mathbf{y}_2 \end{pmatrix} = \begin{pmatrix} \mathbf{A}_1 \\ \mathbf{A}_2 \end{pmatrix} x$

Recursive estimation

Consider the partitioned model:

$$(55) \quad E \begin{pmatrix} \mathbf{y}_1 \\ \mathbf{y}_2 \end{pmatrix} = \begin{pmatrix} \mathbf{A}_1 \\ \mathbf{A}_2 \end{pmatrix} x \quad ; \quad D \begin{pmatrix} \mathbf{y}_1 \\ \mathbf{y}_2 \end{pmatrix} = \begin{pmatrix} \mathbf{Q}_1 & 0 \\ 0 & \mathbf{Q}_2 \end{pmatrix}$$

$(m_1+m_2) \times 1$ $(m_1+m_2) \times n \quad n \times 1$ $(m_1+m_2) \times (m_1+m_2)$

Note that it is assumed that $\mathbf{y}_1$ and $\mathbf{y}_2$ are *uncorrelated*. This is usually the case in practical applications. The least-squares estimator of x of model (55) will be denoted as $\hat{x}_{(2)}$. It reads:

$$(56) \quad \hat{x}_{(2)} = (\mathbf{A}_1^* \mathbf{Q}_1^{-1} \mathbf{A}_1 + \mathbf{A}_2^* \mathbf{Q}_2^{-1} \mathbf{A}_2)^{-1} (\mathbf{A}_1^* \mathbf{Q}_1^{-1} \mathbf{y}_1 + \mathbf{A}_2^* \mathbf{Q}_2^{-1} \mathbf{y}_2).$$

Let us now consider the partial model:

$$(57) \quad E \begin{pmatrix} \mathbf{y}_1 \end{pmatrix} = \mathbf{A}_1 x \quad ; \quad D \begin{pmatrix} \mathbf{y}_1 \end{pmatrix} = \mathbf{Q}_1$$

$m_1 \times 1$ $m_1 \times n \quad n \times 1$ $m_1 \times m_1$ $m_1 \times m_1$

Its solution will be denoted as $\hat{x}_{(1)}$. It reads:

$$(58) \quad \begin{cases} \hat{x}_{(1)} &= (\mathbf{A}_1^* \mathbf{Q}_1^{-1} \mathbf{A}_1)^{-1} \mathbf{A}_1^* \mathbf{Q}_1^{-1} \mathbf{y}_1, \\ \mathbf{Q}_{\hat{x}_{(1)}} &= (\mathbf{A}_1^* \mathbf{Q}_1^{-1} \mathbf{A}_1)^{-1}. \end{cases}$$

Using this result, we may write (56) as:

$$(59) \quad \hat{x}_{(2)} = (\mathbf{Q}_{\hat{x}_{(1)}}^{-1} + \mathbf{A}_2^* \mathbf{Q}_2^{-1} \mathbf{A}_2)^{-1} (\mathbf{Q}_{\hat{x}_{(1)}}^{-1} \hat{x}_{(1)} + \mathbf{A}_2^* \mathbf{Q}_2^{-1} \mathbf{y}_2).$$

But this is also the solution of:

$$(60) \quad E \begin{pmatrix} \hat{x}_{(1)} \\ \mathbf{y}_2 \end{pmatrix} = \begin{pmatrix} \mathbf{I} \\ \mathbf{A}_2 \end{pmatrix} x \quad ; \quad D \begin{pmatrix} \hat{x}_{(1)} \\ \mathbf{y}_2 \end{pmatrix} = \begin{pmatrix} \mathbf{Q}_{\hat{x}_{(1)}} & 0 \\ 0 & \mathbf{Q}_2 \end{pmatrix}.$$

Hence we have proven that the solution of the partitioned model (55) can be found in two steps: First one solves for the partial model (57). This gives $\hat{x}_{(1)}$ and $\mathbf{Q}_{\hat{x}_{(1)}}$. Then in the second step one uses this result together with $\mathbf{y}_2$ to find $\hat{x}_{(2)}$ via model (60). This two step procedure

has an important practical implication. It implies that if new observables, say $\underline{y}_2$, become available one does not need to save the old observables $\underline{y}_1$ to compute $\hat{\underline{x}}_{(2)}$. One can compute $\hat{\underline{x}}_{(2)}$ from the solution $\hat{\underline{x}}_{(1)}$ of the first step and the new observables $\underline{y}_2$. In this way one can recursively determine $\hat{\underline{x}}_{(k)}$ from $\hat{\underline{x}}_{(k-1)}$ and $\underline{y}_k$. This is shown in table 6.1. Expression (59) for $\hat{\underline{x}}_{(2)}$ shows that a matrix of the order n needs to be inverted. One can however also derive an expression for $\hat{\underline{x}}_{(2)}$ in which a matrix of the order m_2 needs to be inverted. This expression is found if we solve (60) via the model of condition equations. Model (60) in terms of condition equations reads:

$$(61) \quad (-\mathbf{A}_2 \ I) E \begin{pmatrix} \hat{\underline{x}}_{(1)} \\ \underline{y}_2 \end{pmatrix} = \mathbf{0} ; \quad \mathbf{D} \begin{pmatrix} \hat{\underline{x}}_{(1)} \\ \underline{y}_2 \end{pmatrix} = \begin{pmatrix} \mathbf{Q}_{\hat{\underline{x}}_{(1)}} & \mathbf{0} \\ \mathbf{0} & \mathbf{Q}_2 \end{pmatrix}$$

Its solution reads:

$$\begin{pmatrix} \hat{\underline{x}}_{(2)} \\ \hat{\underline{y}}_2 \end{pmatrix} = \begin{bmatrix} \mathbf{I} & \mathbf{0} \\ \mathbf{0} & \mathbf{I} \end{bmatrix} - \begin{pmatrix} -\mathbf{Q}_{\hat{\underline{x}}_{(1)}} \mathbf{A}_2^* \\ \mathbf{Q}_2 \end{pmatrix} (\mathbf{Q}_2 + \mathbf{A}_2 \mathbf{Q}_{\hat{\underline{x}}_{(1)}} \mathbf{A}_2^*)^{-1} (-\mathbf{A}_2 \ I) \begin{pmatrix} \hat{\underline{x}}_{(1)} \\ \underline{y}_2 \end{pmatrix}.$$

Hence for $\hat{\underline{x}}_{(2)}$ we get:

$$(62) \quad \hat{\underline{x}}_{(2)} = \hat{\underline{x}}_{(1)} + \mathbf{Q}_{\hat{\underline{x}}_{(1)}} \mathbf{A}_2^* (\mathbf{Q}_2 + \mathbf{A}_2 \mathbf{Q}_{\hat{\underline{x}}_{(1)}} \mathbf{A}_2^*)^{-1} (\underline{y}_2 - \mathbf{A}_2 \hat{\underline{x}}_{(1)}).$$

Application of the propagation law of variances gives:

$$(63) \quad \mathbf{Q}_{\hat{\underline{x}}_{(2)}} = \mathbf{Q}_{\hat{\underline{x}}_{(1)}} - \mathbf{Q}_{\hat{\underline{x}}_{(1)}} \mathbf{A}_2^* (\mathbf{Q}_2 + \mathbf{A}_2 \mathbf{Q}_{\hat{\underline{x}}_{(1)}} \mathbf{A}_2^*)^{-1} \mathbf{A}_2 \mathbf{Q}_{\hat{\underline{x}}_{(1)}}.$$

Both the expressions (59) and (62) give identical results. But expression (62) is more advantageous if m_2 is small compared to n . In particular if $m_2=1$, only a scalar needs to be inverted in (62), whereas in (59) all three matrices $\mathbf{Q}_{\hat{\underline{x}}_{(1)}}$, $\mathbf{Q}_2$ and $(\mathbf{Q}_{\hat{\underline{x}}_{(1)}}^{-1} + \mathbf{A}_2^* \mathbf{Q}_2^{-1} \mathbf{A}_2)$ need to be inverted.

The equations (59) and (62) are called *measurement update equations*. Expression (62) clearly shows how $\hat{\underline{x}}_{(2)}$ is found from updating $\hat{\underline{x}}_{(1)}$. The correction to $\hat{\underline{x}}_{(1)}$ depends on the difference $\underline{y}_2 - \mathbf{A}_2 \hat{\underline{x}}_{(1)}$. Since $\mathbf{A}_2 \hat{\underline{x}}_{(1)}$ can be interpreted as the prediction of $E\{\underline{y}_2\}$ based on $\underline{y}_1$, the difference $\underline{y}_2 - \mathbf{A}_2 \hat{\underline{x}}_{(1)}$ is called the *predicted residual* of $E\{\underline{y}_2\}$. Note that the predicted residual is *not* equal to the least-squares residual of $E\{\underline{y}_2\}$. This least-squares residual reads namely $\underline{y}_2 - \hat{\underline{y}}_2 = \underline{y}_2 - \mathbf{A}_2 \hat{\underline{x}}_{(2)}$.

Expression (62) shows that the correction to $\hat{\underline{x}}_{(1)}$ is small if the predicted residual is small, and that the correction to $\hat{\underline{x}}_{(1)}$ is large if also the predicted residual is large. This is also what one would expect. Also note that the correction to $\hat{\underline{x}}_{(1)}$ is small if the variance of $\hat{\underline{x}}_{(1)}$ is small. This is also understandable, because if the variance of $\hat{\underline{x}}_{(1)}$ is small one has more confidence in $\hat{\underline{x}}_{(1)}$ and therefore would like to give more weight to $\hat{\underline{x}}_{(1)}$ than to $\underline{y}_2$.

Partitioned model

$$E \begin{Bmatrix} y_1 \\ y_2 \\ \vdots \\ y_k \end{Bmatrix} = \begin{Bmatrix} A_1 \\ A_2 \\ \vdots \\ A_k \end{Bmatrix} x ; D \begin{Bmatrix} y_1 \\ y_2 \\ \vdots \\ y_k \end{Bmatrix} = \begin{pmatrix} Q_1 & & 0 \\ & Q_2 & \\ 0 & & \ddots \\ & & & Q_k \end{pmatrix}$$

$$\hat{x}_{(k)} = \left(\sum_{i=1}^k A_i^* Q_i^{-1} A_i \right)^{-1} \left(\sum_{i=1}^k A_i^* Q_i^{-1} y_i \right)$$

$$Q_{\hat{x}_{(k)}} = \left(\sum_{i=1}^k A_i^* Q_i^{-1} A_i \right)^{-1}$$

Recursive estimation

$$1^\circ \quad E\{y_1\} = A_1 x ; D\{y_1\} = Q_1$$

$$\hat{x}_{(1)} = (A_1^* Q_1^{-1} A_1)^{-1} A_1^* Q_1^{-1} y_1$$

$$Q_{\hat{x}_{(1)}} = (A_1^* Q_1^{-1} A_1)^{-1}$$

$$k^\circ \quad E \begin{Bmatrix} \hat{x}_{(k-1)} \\ y_k \end{Bmatrix} = \begin{pmatrix} I \\ A_k \end{pmatrix} x ; D \begin{Bmatrix} \hat{x}_{(k-1)} \\ y_k \end{Bmatrix} = \begin{pmatrix} Q_{\hat{x}_{(k-1)}} & 0 \\ 0 & Q_k \end{pmatrix}$$

$$\hat{x}_{(k)} = (Q_{\hat{x}_{(k-1)}}^{-1} + A_k^* Q_k^{-1} A_k)^{-1} (Q_{\hat{x}_{(k-1)}}^{-1} \hat{x}_{(k-1)} + A_k^* Q_k^{-1} y_k)$$

$$Q_{\hat{x}_{(k)}} = (Q_{\hat{x}_{(k-1)}}^{-1} + A_k^* Q_k^{-1} A_k)^{-1}$$

$$k = 2, 3, \dots$$

Table 6.1: Recursive estimation (A-form)

Expression (63) shows how the variance of the estimator gets updated. Because of the minus sign in (63) the precision of the estimator gets better. This is understandable, because by including the additional observable y_2 more information is available to estimate x . Table 6.2 shows the recursive estimation scheme based on (62) and (63).

Apart from the recursive scheme for the estimator of x , it is also possible to derive a recursive scheme for $\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}}$. Consider again the partitioned model:

$$E \begin{Bmatrix} y_1 \\ y_2 \end{Bmatrix} = \begin{Bmatrix} A_1 \\ A_2 \end{Bmatrix} x ; D \begin{Bmatrix} y_1 \\ y_2 \end{Bmatrix} = \begin{pmatrix} Q_1 & 0 \\ 0 & Q_2 \end{pmatrix} .$$

Then:

$$(64) \quad (\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}})_{(2)} = (y_1 - A_1 \hat{x}_{(2)})^* Q_1^{-1} (y_1 - A_1 \hat{x}_{(2)}) + (y_2 - A_2 \hat{x}_{(2)})^* Q_2^{-1} (y_2 - A_2 \hat{x}_{(2)}) .$$

If we use the abbreviation:

$$(65) \quad v_{(2)} = y_2 - A_2 \hat{x}_{(1)} ; Q_{v_{(2)}} = Q_2 + A_2 Q_{\hat{x}_{(1)}} A_2^* ,$$

we may write the measurements update equation (62) as:

$$(66) \quad \hat{x}_{(2)} = \hat{x}_{(1)} + Q_{\hat{x}_{(1)}} A_2^* Q_{v_{(2)}}^{-1} v_{(2)} .$$

Substitution of (66) into the first term on the right-hand side of (64) gives:

$$(67) \quad (y_1 - A_1 \hat{x}_{(2)})^* Q_1^{-1} (y_1 - A_1 \hat{x}_{(2)}) = \left[(y_1 - A_1 \hat{x}_{(1)}) - A_1 Q_{\hat{x}_{(1)}} A_2^* Q_{v_{(2)}}^{-1} v_{(2)} \right]^* Q_1^{-1} \left[(y_1 - A_1 \hat{x}_{(1)}) - A_1 Q_{\hat{x}_{(1)}} A_2^* Q_{v_{(2)}}^{-1} v_{(2)} \right]$$

Since $A_1^* Q_1^{-1} (y_1 - A_1 \hat{x}_{(1)}) = 0$ and $Q_{\hat{x}_{(1)}}^{-1} = A_1^* Q_1^{-1} A_1$, equation (67) can be written as:

$$(68) \quad (y_1 - A_1 \hat{x}_{(2)})^* Q_1^{-1} (y_1 - A_1 \hat{x}_{(2)}) = (y_1 - A_1 \hat{x}_{(1)})^* Q_1^{-1} (y_1 - A_1 \hat{x}_{(1)}) + v_{(2)}^* Q_{v_{(2)}}^{-1} A_2 Q_{\hat{x}_{(1)}} A_2^* Q_{v_{(2)}}^{-1} v_{(2)}$$

Since $y_1 - A_1 \hat{x}_{(1)}$ is the least-squares residual vector of the partial model:

$$(69) \quad E \{ y_1 \} = A_1 x ; D \{ y_1 \} = Q_1 ,$$

we have:

$$(70) \quad (\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}})_{(1)} = (y_1 - A_1 \hat{x}_{(1)})^* Q_1^{-1} (y_1 - A_1 \hat{x}_{(1)}) .$$

Substitution of (70) into (68) gives:

$$(71) \quad (y_1 - A_1 \hat{x}_{(2)})^* Q_1^{-1} (y_1 - A_1 \hat{x}_{(2)}) = (\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}})_{(1)} + v_{(2)}^* Q_{v_{(2)}}^{-1} A_2 Q_{\hat{x}_{(1)}} A_2^* Q_{v_{(2)}}^{-1} v_{(2)}$$

Partitioned model

$$E \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_k \end{pmatrix} = \begin{pmatrix} A_1 \\ A_2 \\ \vdots \\ A_k \end{pmatrix} x ; D \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_k \end{pmatrix} = \begin{pmatrix} Q_1 & 0 \\ & Q_2 \\ & & \ddots \\ 0 & & & Q_k \end{pmatrix}$$

$$\hat{x}_{(k)} = (\sum_{i=1}^k A_i^* Q_i^{-1} A_i)^{-1} (\sum_{i=1}^k A_i^* Q_i^{-1} y_i)$$

$$Q_{\hat{x}_{(k)}} = (\sum_{i=1}^k A_i^* Q_i^{-1} A_i)^{-1}$$

Recursive estimation

$$1^o. \quad E\{y_1\} = A_1 x ; D\{y_1\} = Q_1$$

$$\hat{x}_{(1)} = (A_1^* Q_1^{-1} A_1)^{-1} A_1^* Q_1^{-1} y_1$$

$$Q_{\hat{x}_{(1)}} = (A_1^* Q_1^{-1} A_1)^{-1}$$

$$k^o. \quad (-A_k I) E \begin{pmatrix} \hat{x}_{(k-1)} \\ y_k \end{pmatrix} = 0 ; D \begin{pmatrix} \hat{x}_{(k-1)} \\ y_k \end{pmatrix} = \begin{pmatrix} Q_{\hat{x}_{(k-1)}} & 0 \\ 0 & Q_k \end{pmatrix}$$

$$\hat{x}_{(k)} = \hat{x}_{(k-1)} + Q_{\hat{x}_{(k-1)}} A_k^* (Q_k + A_k Q_{\hat{x}_{(k-1)}} A_k^*)^{-1} (y_k - A_k \hat{x}_{(k-1)})$$

$$Q_{\hat{x}_{(k)}} = Q_{\hat{x}_{(k-1)}} - Q_{\hat{x}_{(k-1)}} A_k^* (Q_k + A_k Q_{\hat{x}_{(k-1)}} A_k^*)^{-1} A_k Q_{\hat{x}_{(k-1)}}$$

$$k = 2, 3, \dots$$

Table 6.2: Recursive estimation (B-form)

Substitution of (66) into the second term on the right-hand side of (64) gives together with (65):

$$\begin{aligned}
 (72) \quad (\mathbf{y}_2 - \mathbf{A}_2 \hat{\mathbf{x}}_{(2)})^* \mathbf{Q}_2^{-1} (\mathbf{y}_2 - \mathbf{A}_2 \hat{\mathbf{x}}_{(2)}) &= \hat{\mathbf{y}}_{(2)}^* \left[\mathbf{I} - \mathbf{A}_2 \mathbf{Q}_{\hat{\mathbf{x}}_{(1)}} \mathbf{A}_2^* \mathbf{Q}_{\mathbf{v}_{(2)}}^{-1} \right]^* \mathbf{Q}_2^{-1} \left[\mathbf{I} - \mathbf{A}_2 \mathbf{Q}_{\hat{\mathbf{x}}_{(1)}} \mathbf{A}_2^* \mathbf{Q}_{\mathbf{v}_{(2)}}^{-1} \right] \mathbf{y}_{(2)} \\
 &= \mathbf{v}_{(2)}^* \mathbf{Q}_{\mathbf{v}_{(2)}} \mathbf{Q}_2 \mathbf{Q}_{\mathbf{v}_{(2)}}^{-1} \mathbf{v}_{(2)}
 \end{aligned}$$

Substitution of (71) and (72) into (64) shows finally that:

$$(73) \quad (\hat{\mathbf{e}}^* \mathbf{Q}_y^{-1} \hat{\mathbf{e}})_{(2)} = (\hat{\mathbf{e}}^* \mathbf{Q}_y^{-1} \hat{\mathbf{e}})_{(1)} + \mathbf{v}_{(2)}^* \mathbf{Q}_{\mathbf{v}_{(2)}}^{-1} \mathbf{v}_{(2)}.$$

This equation shows how the norm of the least-squares residual vector should be updated when the observable $\mathbf{y}_2$ becomes available. The update depends on the predicted residual. Expression (73) can be generalized to more steps. This is shown in table 6.3.

Partitioned model	
$E \begin{Bmatrix} \mathbf{y}_1 \\ \mathbf{y}_2 \\ \vdots \\ \mathbf{y}_k \end{Bmatrix} = \begin{Bmatrix} \mathbf{A}_1 \\ \mathbf{A}_2 \\ \vdots \\ \mathbf{A}_k \end{Bmatrix} \mathbf{x}; \quad D \begin{Bmatrix} \mathbf{y}_1 \\ \mathbf{y}_2 \\ \vdots \\ \mathbf{y}_k \end{Bmatrix} = \begin{pmatrix} \mathbf{Q}_1 & & 0 \\ & \mathbf{Q}_2 & \\ 0 & & \ddots \\ & & & \mathbf{Q}_k \end{pmatrix}$	
$(\hat{\mathbf{e}}^* \mathbf{Q}_y^{-1} \hat{\mathbf{e}})_{(k)} = \sum_{i=1}^k (\mathbf{y}_i - \mathbf{A}_i \hat{\mathbf{x}}_{(k)})^* \mathbf{Q}_i^{-1} (\mathbf{y}_i - \mathbf{A}_i \hat{\mathbf{x}}_{(k)})$	
Recursive update	
$ \begin{aligned} (\hat{\mathbf{e}}^* \mathbf{Q}_y^{-1} \hat{\mathbf{e}})_{(k)} &= (\hat{\mathbf{e}}^* \mathbf{Q}_y^{-1} \hat{\mathbf{e}})_{(k-1)} + \mathbf{v}_{(k)}^* \mathbf{Q}_{\mathbf{v}_{(k)}}^{-1} \mathbf{v}_{(k)} \\ \mathbf{v}_{(k)} &= \mathbf{y}_k - \mathbf{A}_k \hat{\mathbf{x}}_{(k-1)}; \quad \mathbf{Q}_{\mathbf{v}_{(k)}} = (\mathbf{Q}_k + \mathbf{A}_k \mathbf{Q}_{\hat{\mathbf{x}}_{(k-1)}} \mathbf{A}_k^*) \end{aligned} $	

Table 6.3: Recursive update of $(\hat{\mathbf{e}}^* \mathbf{Q}_y^{-1} \hat{\mathbf{e}})_{(k)}$

So far it was assumed that the observables $\mathbf{y}_1$ and $\mathbf{y}_2$ are uncorrelated. Let us now assume that they are *correlated*. Hence, we consider the partitioned model:

$$(74) \quad E \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} A_1 \\ A_2 \end{pmatrix} x ; D \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} Q_1 & Q_{12} \\ Q_{21} & Q_2 \end{pmatrix} ,$$

With $Q_{12} \neq 0$, $Q_{21} \neq 0$. In this case it is *not* possible to apply the above derived recursive schemes directly. In order to be able to apply the above derived recursive schemes, the observables y_1 and y_2 first need to be *de-correlated*. This is possible if we transform model (74) such that the transformed observables are uncorrelated. For instance, if we define new observables as:

$$(75) \quad \begin{pmatrix} y_1 \\ \bar{y}_2 \end{pmatrix} = \begin{pmatrix} I & 0 \\ -Q_{21}Q_1^{-1} & I \end{pmatrix} \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} ,$$

then it follows with (74) that:

$$(76) \quad E \begin{pmatrix} y_1 \\ \bar{y}_2 \end{pmatrix} = \begin{pmatrix} A_1 \\ \bar{A}_2 \end{pmatrix} x ; D \begin{pmatrix} y_1 \\ \bar{y}_2 \end{pmatrix} = \begin{pmatrix} Q_1 & 0 \\ 0 & \bar{Q}_2 \end{pmatrix} ,$$

with:

$$(77) \quad \bar{A}_2 = A_2 - Q_{21}Q_1^{-1}A_1 ; \bar{Q}_2 = Q_2 - Q_{21}Q_1^{-1}Q_{12} .$$

Similarly, if we define new observables as:

$$(78) \quad \begin{pmatrix} \bar{y}_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} I & -Q_{12}Q_2^{-1} \\ 0 & I \end{pmatrix} \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} ,$$

we get with (74) that:

$$(79) \quad E \begin{pmatrix} \bar{y}_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} \bar{A}_1 \\ A_2 \end{pmatrix} x ; D \begin{pmatrix} \bar{y}_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} \bar{Q}_1 & 0 \\ 0 & Q_2 \end{pmatrix} ,$$

with:

$$(80) \quad \bar{A}_1 = A_1 - Q_{12}Q_2^{-1}A_2 ; \bar{Q}_1 = Q_1 - Q_{12}Q_2^{-1}Q_{21}$$

With (76) or (79) it is now possible again to apply the recursive schemes. In this case however the recursive schemes *lose* their practical implication. The transformed observable $\bar{y}_2$ of (76) reads namely:

$$\bar{y}_2 = y_2 - Q_{21}Q_1^{-1}y_1 .$$

This shows that one still needs to save the old observable y_1 for computing $\hat{x}_{(2)}$ from $\hat{x}_{(1)}$ and y_2 . The conclusion reads therefore that the recursive schemes of this section are only of

practical value if the observables are uncorrelated. Fortunately this is the case in most practical applications.

The theory as developed in this section plays an important role in the estimation theory of dynamic systems and in the theory of Kalman filtering.

Example 2

Consider the linear model:

$$(81) \quad E \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_k \end{pmatrix} = \begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_k \end{pmatrix} x ; D \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_k \end{pmatrix} = \sigma^2 I_k .$$

Its solution reads:

$$(82) \quad \begin{cases} \hat{x}_{(k)} = \sum_{i=1}^k a_i y_i / \sum_{i=1}^k a_i^2 , \\ \sigma_{\hat{x}_{(k)}}^2 = \sigma^2 (\sum_{i=1}^k a_i^2)^{-1} . \end{cases}$$

The solution of the partial model:

$$(83) \quad E \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_{k-1} \end{pmatrix} = \begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_{k-1} \end{pmatrix} x ; D \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_{k-1} \end{pmatrix} = \sigma^2 I_{k-1} ,$$

reads:

$$(84) \quad \begin{cases} \hat{x}_{(k-1)} = \sum_{i=1}^{k-1} a_i y_i / \sum_{i=1}^{k-1} a_i^2 , \\ \sigma_{\hat{x}_{(k-1)}}^2 = \sigma^2 (\sum_{i=1}^{k-1} a_i^2)^{-1} . \end{cases}$$

According to table 6.2, the solution of (81) can also be written as:

$$(85) \quad \begin{cases} \hat{x}_{(k)} = \hat{x}_{(k-1)} + \sigma_{\hat{x}_{(k-1)}}^2 a_k (\sigma^2 + \sigma_{\hat{x}_{(k-1)}}^2 a_k^2)^{-1} (y_k - a_k \hat{x}_{(k-1)}) , \\ \sigma_{\hat{x}_{(k)}}^2 = \sigma_{\hat{x}_{(k-1)}}^2 - \sigma_{\hat{x}_{(k-1)}}^2 a_k (\sigma^2 + \sigma_{\hat{x}_{(k-1)}}^2 a_k^2)^{-1} a_k \sigma_{\hat{x}_{(k-1)}}^2 . \end{cases}$$

We will now show that (85) is indeed identical to (82). We will first consider the first equation of (85). This equation can be written as:

$$\hat{x}_{(k)} = \left(1 - \frac{\sigma_{\hat{x}_{(k-1)}}^2 a_k^2}{\sigma^2 + \sigma_{\hat{x}_{(k-1)}}^2 a_k^2} \right) \hat{x}_{(k-1)} + \frac{\sigma_{\hat{x}_{(k-1)}}^2 a_k}{\sigma^2 + \sigma_{\hat{x}_{(k-1)}}^2 a_k^2} y_k$$

or as:

$$\hat{x}_{(k)} = \frac{\sigma_{\hat{x}_{(k-1)}}^2 \hat{x}_{(k-1)} + \sigma_{\hat{x}_{(k-1)}}^2 a_k y_k}{\sigma^2 + \sigma_{\hat{x}_{(k-1)}}^2 a_k^2}$$

or with (84) as:

$$\hat{x}_{(k)} = \frac{\sigma^2 (\sum_{i=1}^{k-1} a_i y_i + a_k y_k) / \sum_{i=1}^{k-1} a_i^2}{\sigma^2 (1 + a_k^2 / \sum_{i=1}^{k-1} a_i^2)}$$

or as:

$$\hat{x}_{(k)} = \frac{\sigma^2 (\sum_{i=1}^{k-1} a_i y_i + a_k y_k) / \sum_{i=1}^{k-1} a_i^2}{\sigma^2 (\sum_{i=1}^{k-1} a_i^2 + a_k^2) / \sum_{i=1}^{k-1} a_i^2}$$

which indeed is identical to the first equation of (82). Now consider the second equation of (85). This can be written as:

$$\sigma_{\hat{x}_{(k)}}^2 = \left(1 - \frac{\sigma_{\hat{x}_{(k-1)}}^2 a_k^2}{\sigma^2 + \sigma_{\hat{x}_{(k-1)}}^2 a_k^2} \right) \sigma_{\hat{x}_{(k-1)}}^2$$

or as:

$$\sigma_{\hat{x}_{(k)}}^2 = \frac{\sigma^2 \sigma_{\hat{x}_{(k-1)}}^2}{\sigma^2 + \sigma_{\hat{x}_{(k-1)}}^2 a_k^2}$$

or with (84) as:

$$\sigma_{\hat{x}_{(k)}}^2 = \sigma^2 \frac{\sigma^2 (\sum_{i=1}^{k-1} a_i^2)^{-1}}{\sigma^2 + \sigma^2 (\sum_{i=1}^{k-1} a_i^2)^{-1} a_k^2}$$

or as:

$$\sigma_{\hat{x}_{(k)}}^2 = \sigma^2 \frac{1}{\sum_{i=1}^{k-1} a_i^2 + a_k^2} ,$$

which is indeed identical to the second equation of (82).

6.4 The partitioned model $E\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = \begin{pmatrix} A_1 \\ A_2 \end{pmatrix} x$

Block estimation I

Consider again the partitioned model:

$$(86) \quad E\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = \begin{pmatrix} A_1 \\ A_2 \end{pmatrix} x \quad ; \quad D\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = \begin{pmatrix} Q_1 & 0 \\ 0 & Q_2 \end{pmatrix}$$

Its system of normal equations reads:

$$(87) \quad \sum_{i=1}^2 A_i^* Q_i^{-1} A_i \hat{x} = \sum_{i=1}^2 A_i^* Q_i^{-1} y_i .$$

Now consider the two partial models:

$$(88) \quad \begin{cases} E\{y_1\} = A_1 x & ; \quad D\{y_1\} = Q_1 \\ E\{y_2\} = A_2 x & ; \quad D\{y_2\} = Q_2 \end{cases}$$

Their systems of normal equations read: ⁴⁾

$$(89) \quad \begin{cases} A_1^* Q_1^{-1} A_1 \hat{x}_{(1)} = A_1^* Q_1^{-1} y_1 \\ A_2^* Q_2^{-1} A_2 \hat{x}_{(2)} = A_2^* Q_2^{-1} y_2 \end{cases}$$

Comparison of (89) with (87) shows that the normal matrix of model (86) can be found from the *sum* of the two normal matrices of the two partial models in (88), and that the righthand side of (87) can be found from the *sum* of the two right hand sides of the normal equations in (89).

Let us now generalize model (86) to:

$$(90) \quad E\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = \begin{pmatrix} A_{11} & A_{12} & 0 \\ 0 & A_{22} & A_{23} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} ; \quad D\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = \begin{pmatrix} Q_1 & 0 \\ 0 & Q_2 \end{pmatrix}$$

The corresponding partitioned system of normal equations reads:

⁴⁾ Do not confuse $\hat{x}_{(2)}$ of (89) with $\hat{x}_{(2)}$ of the previous section.

$$(91) \quad \begin{pmatrix} \mathbf{A}_{11}^* \mathbf{Q}_1^{-1} \mathbf{A}_{11} & \mathbf{A}_{11}^* \mathbf{Q}_1^{-1} \mathbf{A}_{12} & 0 \\ \mathbf{A}_{12}^* \mathbf{Q}_1^{-1} \mathbf{A}_{11} & \sum_{i=1}^2 \mathbf{A}_{i2}^* \mathbf{Q}_i^{-1} \mathbf{A}_{i2} & \mathbf{A}_{22}^* \mathbf{Q}_2^{-1} \mathbf{A}_{23} \\ 0 & \mathbf{A}_{23}^* \mathbf{Q}_2^{-1} \mathbf{A}_{22} & \mathbf{A}_{23}^* \mathbf{Q}_2^{-1} \mathbf{A}_{23} \end{pmatrix} \begin{pmatrix} \hat{\mathbf{x}}_1 \\ \hat{\mathbf{x}}_2 \\ \hat{\mathbf{x}}_3 \end{pmatrix} = \begin{pmatrix} \mathbf{A}_{11}^* \mathbf{Q}_1^{-1} \mathbf{y}_1 \\ \sum_{i=1}^2 \mathbf{A}_{i2}^* \mathbf{Q}_i^{-1} \mathbf{y}_i \\ \mathbf{A}_{23}^* \mathbf{Q}_2^{-1} \mathbf{y}_2 \end{pmatrix}.$$

The estimator $\hat{\mathbf{x}}_1$ gets eliminated from the second set of normal equations if we premultiply (91) with the square and full rank matrix:

$$(92) \quad \begin{pmatrix} I & 0 & 0 \\ -\mathbf{A}_{12}^* \mathbf{Q}_1^{-1} \mathbf{A}_{11} (\mathbf{A}_{11}^* \mathbf{Q}_1^{-1} \mathbf{A}_{11})^{-1} & I & 0 \\ 0 & 0 & I \end{pmatrix}.$$

This gives:

$$(93) \quad \begin{pmatrix} \mathbf{A}_{11}^* \mathbf{Q}_1^{-1} \mathbf{A}_{11} & \mathbf{A}_{11}^* \mathbf{Q}_1^{-1} \mathbf{A}_{12} & 0 \\ 0 & (\bar{\mathbf{A}}_{12}^* \mathbf{Q}_1^{-1} \bar{\mathbf{A}}_{12} + \mathbf{A}_{22}^* \mathbf{Q}_2^{-1} \mathbf{A}_{22}) & \mathbf{A}_{22}^* \mathbf{Q}_2^{-1} \mathbf{A}_{23} \\ 0 & \mathbf{A}_{23}^* \mathbf{Q}_2^{-1} \mathbf{A}_{22} & \mathbf{A}_{23}^* \mathbf{Q}_2^{-1} \mathbf{A}_{23} \end{pmatrix} \begin{pmatrix} \hat{\mathbf{x}}_1 \\ \hat{\mathbf{x}}_2 \\ \hat{\mathbf{x}}_3 \end{pmatrix} = \begin{pmatrix} \mathbf{A}_{11}^* \mathbf{Q}_1^{-1} \mathbf{y}_1 \\ \bar{\mathbf{A}}_{12}^* \mathbf{Q}_1^{-1} \mathbf{y}_1 + \mathbf{A}_{22}^* \mathbf{Q}_2^{-1} \mathbf{y}_2 \\ \mathbf{A}_{23}^* \mathbf{Q}_2^{-1} \mathbf{y}_2 \end{pmatrix},$$

where we have used the abbreviation:

$$(94) \quad \bar{\mathbf{A}}_{12} = \mathbf{P}_{\mathbf{A}_{11}}^\perp \mathbf{A}_{12} . \quad (\text{see also section 6.2})$$

The estimator $\hat{\mathbf{x}}_3$ gets eliminated from the second set of normal equations if we premultiply (93) with the square and full rank matrix:

$$(95) \quad \begin{pmatrix} I & 0 & 0 \\ 0 & I & -\mathbf{A}_{22}^* \mathbf{Q}_2^{-1} \mathbf{A}_{23} (\mathbf{A}_{23}^* \mathbf{Q}_2^{-1} \mathbf{A}_{23})^{-1} \\ 0 & 0 & I \end{pmatrix}.$$

This gives:

$$(96) \quad \begin{pmatrix} \mathbf{A}_{11}^* \mathbf{Q}_1^{-1} \mathbf{A}_{11} & \mathbf{A}_{11}^* \mathbf{Q}_1^{-1} \mathbf{A}_{12} & 0 \\ 0 & \sum_{i=1}^2 \bar{\mathbf{A}}_{i2}^* \mathbf{Q}_i^{-1} \bar{\mathbf{A}}_{i2} & 0 \\ 0 & \mathbf{A}_{23}^* \mathbf{Q}_2^{-1} \mathbf{A}_{22} & \mathbf{A}_{23}^* \mathbf{Q}_2^{-1} \mathbf{A}_{23} \end{pmatrix} \begin{pmatrix} \hat{\mathbf{x}}_1 \\ \hat{\mathbf{x}}_2 \\ \hat{\mathbf{x}}_3 \end{pmatrix} = \begin{pmatrix} \mathbf{A}_{11}^* \mathbf{Q}_1^{-1} \mathbf{y}_1 \\ \sum_{i=1}^2 \bar{\mathbf{A}}_{i2}^* \mathbf{Q}_i^{-1} \mathbf{y}_i \\ \mathbf{A}_{23}^* \mathbf{Q}_2^{-1} \mathbf{y}_2 \end{pmatrix},$$

where we have used the abbreviation:

$$(97) \quad \bar{\mathbf{A}}_{22} = \mathbf{P}_{\mathbf{A}_{23}}^\perp \mathbf{A}_{22} .$$

The partially reduced system of normal equations (96) can now be solved in the following way. First one uses the second set of normal equations in (96) to solve for $\hat{x}_2$. Once $\hat{x}_2$ is known, it can be substituted into the first and third set of normal equations in (96) to solve for $\hat{x}_1$ and $\hat{x}_3$ respectively. Thus the solution of (96) is found as:

$$(98) \quad \begin{cases} a) \quad \hat{x}_2 = (\sum_{i=1}^2 \bar{A}_{i2}^* Q_i^{-1} \bar{A}_{i2})^{-1} (\sum_{i=1}^2 \bar{A}_{i2}^* Q_i^{-1} y_i) , \\ b) \quad \hat{x}_1 = (A_{11}^* Q_1^{-1} A_{11})^{-1} A_{11}^* Q_1^{-1} (y_1 - A_{12} \hat{x}_2) , \\ c) \quad \hat{x}_3 = (A_{23}^* Q_2^{-1} A_{23})^{-1} A_{23}^* Q_2^{-1} (y_2 - A_{22} \hat{x}_2) . \end{cases}$$

What is the practical implication of the above solution method? Let us assume that two countries, say the Netherlands and Germany, are covered with a large geodetic network. (See figure 6.5.)

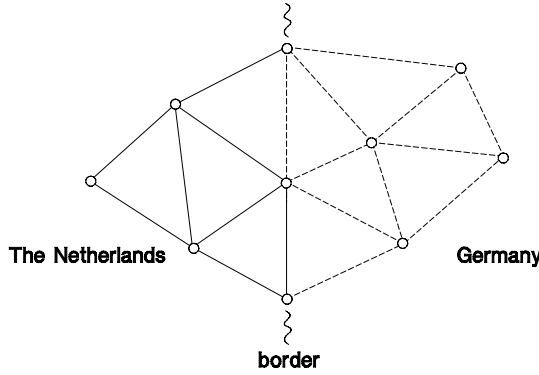


Figure 6.5: Geodetic network covering the Netherlands and Germany

The Dutch part of the network (full lines in figure 6.5) has been measured by the Dutch Triangulation Department (Rijksdriehoeksmeting), and the German part of the network (dashed lines in figure 6.5) has been measured by the German Triangulation Department. We will assume that the vector x_1 contains the unknown coordinates of the network points in the Netherlands, that the vector x_2 contains the unknown coordinates of the network points on the border, and that the vector x_3 contains the unknown coordinates of the network points in Germany.

The partial model of the Dutch part of the network reads then:

$$(99) \quad E\{y_1\} = (A_{11} \ A_{12}) \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} ; \quad D\{y_1\} = Q_1 .$$

Similarly, the partial model of the German part of the network reads:

$$(100) \quad E\{y_2\} = (A_{22} \ A_{23}) \begin{pmatrix} x_2 \\ x_3 \end{pmatrix} ; \quad D\{y_2\} = Q_2 .$$

There are now two ways in which a Computing Centre can perform the adjustment of the geodetic network. The *first* method consists of the following: Both the Dutch and the Germans send their partial models, (99) and (100), together with the measured data to the Computing Centre. The Computing Centre merges the two partial models to get (90), the complete model. This model is then solved by the Computing Centre for $\hat{x}_1$, $\hat{x}_2$ and $\hat{x}_3$. The solution $\hat{x}_1$, $\hat{x}_2$ is then sent by the Computing Centre back to the Dutch, and similarly the solution $\hat{x}_2$, $\hat{x}_3$ is sent back to the Germans. The *second* method consists of the following: based on their partial model, the Dutch form their partial system of normal equations and reduce it for x_1 . This gives:

$$(101) \quad \begin{pmatrix} A_{11}^* Q_1^{-1} A_{11} & A_{11}^* Q_1^{-1} A_{12} \\ 0 & \bar{A}_{12}^* Q_1^{-1} \bar{A}_{12} \end{pmatrix} \begin{pmatrix} \hat{x}_1 \\ \hat{x}_2 \end{pmatrix} = \begin{pmatrix} A_{11}^* Q_1^{-1} y_1 \\ \bar{A}_{12}^* Q_1^{-1} y_1 \end{pmatrix}.$$

Similarly, the Germans form, on the basis of their partial model, their partial system of normal equations and reduce it for x_3 . This gives:

$$(102) \quad \begin{pmatrix} \bar{A}_{22}^* Q_2^{-1} \bar{A}_{22} & 0 \\ A_{23}^* Q_2^{-1} A_{22} & A_{23}^* Q_2^{-1} A_{23} \end{pmatrix} \begin{pmatrix} \hat{x}_2 \\ \hat{x}_3 \end{pmatrix} = \begin{pmatrix} \bar{A}_{22}^* Q_2^{-1} y_2 \\ A_{23}^* Q_2^{-1} y_2 \end{pmatrix}.$$

The Dutch then sent their reduced normal block $\bar{A}_{12}^* Q_1^{-1} \bar{A}_{12}$ and the reduced right-hand side $\bar{A}_{12}^* Q_1^{-1} y_1$ to the Computing Centre. And the Germans do the same for their reduced normal block $A_{22}^* Q_2^{-1} A_{22}$ and their reduced right-hand side $A_{22}^* Q_2^{-1} y_2$. The reduced normal blocks and the reduced right-hand sides are then added by the Computing Centre to form the system:

$$(103) \quad \left(\sum_{i=1}^2 \bar{A}_{i2}^* Q_i^{-1} \bar{A}_{i2} \right) \hat{x}_2 = \sum_{i=1}^2 \bar{A}_{i2}^* Q_i^{-1} y_i.$$

This system is solved by the Computing Centre for $\hat{x}_2$ (see (98a)). The solution $\hat{x}_2$ is then sent back to both the Dutch and the Germans. With $\hat{x}_2$ and (101) the Dutch form the system:

$$(104) \quad A_{11}^* Q_1^{-1} A_{11} \hat{x}_1 = A_{11}^* Q_1^{-1} (y_1 - A_{12} \hat{x}_2).$$

This system is solved by the Dutch for $\hat{x}_1$ (see (98b)). In a similar way the Germans use $\hat{x}_2$ and (102) to form the system:

$$(105) \quad A_{23}^* Q_2^{-1} A_{23} \hat{x}_3 = A_{23}^* Q_2^{-1} (y_2 - A_{22} \hat{x}_2).$$

This system is solved by the Germans for $\hat{x}_3$ (see (98c)).

When we compare the above two methods we note that:

- 1°. The computational load for the Computing Centre is less with the second method than with the first method. With the second method some of the computational load stays namely at the two national Triangulation Departments. The order of the system to be solved by the Computing Centre is then equal to the dimension of x_2 (see (103)).

- 2°. With the second method the Computing Centre never gets to know the original measured data, since $\bar{A}_{i2}^* Q_i^{-1} y_i$, $i=1,2$, is provided and not y_i , $i=1,2$. Hence, the Computing Centre will never be able to compute the complete network. This may be an advantage in case secrecy of data and coordinates is asked for.

The above developed second method was originally introduced by the famous German geodesist F.R. Helmert (1843-1917). The method is known as *Helmert's Block method*. The method has been used with great success for the readjustment of the European triangulation network (RETrig: Réseau Européen Trigonometrique).

6.5 The partitioned model $E\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = \begin{pmatrix} A_1 \\ A_2 \end{pmatrix} x$

Block estimation II

Consider again the partioned model:

$$(106) \quad E\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = \begin{pmatrix} A_1 \\ A_2 \end{pmatrix} x \quad ; \quad D\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = \begin{pmatrix} Q_1 & 0 \\ 0 & Q_2 \end{pmatrix}.$$

Its solution reads:

$$(107) \quad \hat{x} = (A_1^* Q_1^{-1} A_1 + A_2^* Q_2^{-1} A_2)^{-1} (A_1^* Q_1^{-1} y_1 + A_2^* Q_2^{-1} y_2).$$

Now consider the two partial models:

$$(108) \quad \begin{cases} E\{y_1\} = A_1 x & ; \quad D\{y_1\} = Q_1, \\ E\{y_2\} = A_2 x & ; \quad D\{y_2\} = Q_2. \end{cases}$$

Their solutions read:

$$(109) \quad \begin{cases} \hat{x}^{(1)} = (A_1^* Q_1^{-1} A_1)^{-1} A_1^* Q_1^{-1} y_1 & ; \quad Q_{\hat{x}^{(1)}} = (A_1^* Q_1^{-1} A_1)^{-1}, \\ \hat{x}^{(2)} = (A_2^* Q_2^{-1} A_2)^{-1} A_2^* Q_2^{-1} y_2 & ; \quad Q_{\hat{x}^{(2)}} = (A_2^* Q_2^{-1} A_2)^{-1}. \end{cases}$$

Using this results we may write (107) also as:

$$(110) \quad \hat{x} = (Q_{\hat{x}^{(1)}} + Q_{\hat{x}^{(2)}})^{-1} (Q_{\hat{x}^{(1)}}^{-1} \hat{x}^{(1)} + Q_{\hat{x}^{(2)}}^{-1} \hat{x}^{(2)}).$$

But this is also the solution of:

$$(111) \quad E\left\{\begin{pmatrix} \hat{x}^{(1)} \\ \hat{x}^{(2)} \end{pmatrix}\right\} = \begin{pmatrix} I \\ I \end{pmatrix} x \quad ; \quad D\left\{\begin{pmatrix} \hat{x}^{(1)} \\ \hat{x}^{(2)} \end{pmatrix}\right\} = \begin{pmatrix} Q_{\hat{x}^{(1)}} & 0 \\ 0 & Q_{\hat{x}^{(2)}} \end{pmatrix}.$$

This model is in terms of observation equations. Its equivalent form in terms of condition equations reads:

$$(112) \quad (I - I)E\left\{\begin{pmatrix} \hat{\underline{x}}^{(1)} \\ \hat{\underline{x}}^{(2)} \end{pmatrix}\right\} = \mathbf{0} ; \quad D\left\{\begin{pmatrix} \hat{\underline{x}}^{(1)} \\ \hat{\underline{x}}^{(2)} \end{pmatrix}\right\} = \begin{pmatrix} Q_{\hat{x}^{(1)}} & \mathbf{0} \\ \mathbf{0} & Q_{\hat{x}^{(2)}} \end{pmatrix}$$

And its solution reads:

$$(113) \quad \begin{pmatrix} \hat{\underline{x}} \\ \hat{\underline{x}} \end{pmatrix} = \begin{bmatrix} I & \mathbf{0} \\ \mathbf{0} & I \end{bmatrix} - \begin{pmatrix} Q_{\hat{x}^{(1)}} \\ -Q_{\hat{x}^{(2)}} \end{pmatrix} (Q_{\hat{x}^{(1)}} + Q_{\hat{x}^{(2)}})^{-1} (I - I) \begin{bmatrix} \hat{\underline{x}}^{(1)} \\ \hat{\underline{x}}^{(2)} \end{bmatrix}.$$

Hence we have the following two equivalent expressions for the solution of (112):

$$(114) \quad \begin{cases} a) \quad \hat{\underline{x}} &= \hat{\underline{x}}^{(1)} - Q_{\hat{x}^{(1)}} (Q_{\hat{x}^{(1)}} + Q_{\hat{x}^{(2)}})^{-1} (\hat{\underline{x}}^{(1)} - \hat{\underline{x}}^{(2)}), \\ b) \quad Q_{\hat{x}} &= Q_{\hat{x}^{(1)}} - Q_{\hat{x}^{(1)}} (Q_{\hat{x}^{(1)}} + Q_{\hat{x}^{(2)}})^{-1} Q_{\hat{x}^{(1)}} \end{cases}$$

and:

$$(115) \quad \begin{cases} a) \quad \hat{\underline{x}} &= \hat{\underline{x}}^{(2)} - Q_{\hat{x}^{(2)}} (Q_{\hat{x}^{(1)}} + Q_{\hat{x}^{(2)}})^{-1} (\hat{\underline{x}}^{(2)} - \hat{\underline{x}}^{(1)}), \\ b) \quad Q_{\hat{x}} &= Q_{\hat{x}^{(2)}} - Q_{\hat{x}^{(2)}} (Q_{\hat{x}^{(1)}} + Q_{\hat{x}^{(2)}})^{-1} Q_{\hat{x}^{(2)}}. \end{cases}$$

Note that (115) follows from (114) by interchanging the role of $\hat{\underline{x}}^{(1)}$ and $\hat{\underline{x}}^{(2)}$. The estimators of (114) and (115) are identical to (110) and therefore also identical to (107). Hence, we have shown how the solution of the partitioned model (106) can be obtained in three separate steps. This is shown in table 6.4.

Partitioned model	
$E\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = \begin{pmatrix} A_1 \\ A_2 \end{pmatrix} x ; D\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = \begin{pmatrix} Q_1 & 0 \\ 0 & Q_2 \end{pmatrix}$ $\hat{x} = (A_1^* Q_1^{-1} A_1 + A_2^* Q_2^{-1} A_2)^{-1} (A_1^* Q_1^{-1} y_1 + A_2^* Q_2^{-1} y_2)$ $Q_{\hat{x}} = (A_1^* Q_1^{-1} A_1 + A_2^* Q_2^{-1} A_2)^{-1} .$	
Block estimation	
1°.	$E\{y_1\} = A_1 x ; D\{y_1\} = Q_1 .$ $\hat{x}^{(1)} = (A_1^* Q_1^{-1} A_1)^{-1} A_1^* Q_1^{-1} y_1 ; Q_{\hat{x}^{(1)}} = (A_1^* Q_1^{-1} A_1)^{-1} .$
2°.	$E\{y_2\} = A_2 x ; D\{y_2\} = Q_2$ $\hat{x}^{(2)} = (A_2^* Q_2^{-1} A_2)^{-1} A_2^* Q_2^{-1} y_2 ; Q_{\hat{x}^{(2)}} = (A_2^* Q_2^{-1} A_2)^{-1}$
3°.	$(I - I) E\left\{\begin{pmatrix} \hat{x}^{(1)} \\ \hat{x}^{(2)} \end{pmatrix}\right\} = 0 ; D\left\{\begin{pmatrix} \hat{x}^{(1)} \\ \hat{x}^{(2)} \end{pmatrix}\right\} = \begin{pmatrix} Q_{\hat{x}^{(1)}} & 0 \\ 0 & Q_{\hat{x}^{(2)}} \end{pmatrix}$
or	$\begin{cases} \hat{x} &= \hat{x}^{(1)} - Q_{\hat{x}^{(1)}} (Q_{\hat{x}^{(1)}} + Q_{\hat{x}^{(2)}})^{-1} (\hat{x}^{(1)} - \hat{x}^{(2)}) \\ Q_{\hat{x}} &= Q_{\hat{x}^{(1)}} - Q_{\hat{x}^{(1)}} (Q_{\hat{x}^{(1)}} + Q_{\hat{x}^{(2)}})^{-1} Q_{\hat{x}^{(1)}} \end{cases}$ $\begin{cases} \hat{x} &= \hat{x}^{(2)} - Q_{\hat{x}^{(2)}} (Q_{\hat{x}^{(1)}} + Q_{\hat{x}^{(2)}})^{-1} (\hat{x}^{(2)} - \hat{x}^{(1)}) \\ Q_{\hat{x}} &= Q_{\hat{x}^{(2)}} - Q_{\hat{x}^{(2)}} (Q_{\hat{x}^{(1)}} + Q_{\hat{x}^{(2)}})^{-1} Q_{\hat{x}^{(2)}} \end{cases}$

Table 6.4: Block estimation

We will now generalize the above results. Consider therefore the partitioned model:

$$(116) \quad E\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = \begin{pmatrix} A_{11} & A_{12} & 0 \\ 0 & A_{22} & A_{23} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix}; \quad D\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = \begin{pmatrix} Q_1 & 0 \\ 0 & Q_2 \end{pmatrix}.$$

Its solution reads:

$$(117) \quad \begin{pmatrix} \hat{x}_1 \\ \hat{x}_2 \\ \hat{x}_3 \end{pmatrix} = \begin{pmatrix} A_{11}^* Q_1^{-1} A_{11} & A_{11}^* Q_1^{-1} A_{12} & 0 \\ A_{12}^* Q_1^{-1} A_{11} & \sum_{i=1}^2 A_{i2} Q_i^{-1} A_{i2} & A_{22}^* Q_2^{-1} A_{23} \\ 0 & A_{23}^* Q_2^{-1} A_{22} & A_{23}^* Q_2^{-1} A_{23} \end{pmatrix}^{-1} \begin{pmatrix} A_{11}^* Q_1^{-1} y_1 \\ \sum_{i=1}^2 A_{i2}^* Q_i^{-1} y_i \\ A_{23}^* Q_2^{-1} y_2 \end{pmatrix}.$$

Now consider the two partial models:

$$(118) \quad \begin{cases} E\{y_1\} = (A_{11} \ A_{12}) \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}; \quad D\{y_1\} = Q_1, \\ E\{y_2\} = (A_{22} \ A_{23}) \begin{pmatrix} x_2 \\ x_3 \end{pmatrix}; \quad D\{y_2\} = Q_2. \end{cases}$$

Their solutions read:

$$(119) \quad \begin{cases} \begin{pmatrix} \hat{x}_1^{(1)} \\ \hat{x}_2^{(1)} \end{pmatrix} = \begin{pmatrix} A_{11}^* Q_1^{-1} A_{11} & A_{11}^* Q_1^{-1} A_{12} \\ A_{12}^* Q_1^{-1} A_{11} & A_{12}^* Q_1^{-1} A_{12} \end{pmatrix}^{-1} \begin{pmatrix} A_{11}^* Q_1^{-1} y_1 \\ A_{12}^* Q_1^{-1} y_1 \end{pmatrix}; \quad Q_{\hat{x}^{(1)}} = \begin{pmatrix} A_{11}^* Q_1^{-1} A_{11} & A_{11}^* Q_1^{-1} A_{12} \\ A_{12}^* Q_1^{-1} A_{11} & A_{12}^* Q_1^{-1} A_{12} \end{pmatrix}^{-1}, \\ \begin{pmatrix} \hat{x}_2^{(2)} \\ \hat{x}_3^{(2)} \end{pmatrix} = \begin{pmatrix} A_{22}^* Q_2^{-1} A_{22} & A_{22}^* Q_2^{-1} A_{23} \\ A_{23}^* Q_2^{-1} A_{22} & A_{23}^* Q_2^{-1} A_{23} \end{pmatrix}^{-1} \begin{pmatrix} A_{22}^* Q_2^{-1} y_2 \\ A_{23}^* Q_2^{-1} y_2 \end{pmatrix}; \quad Q_{\hat{x}^{(2)}} = \begin{pmatrix} A_{22}^* Q_2^{-1} A_{22} & A_{22}^* Q_2^{-1} A_{23} \\ A_{23}^* Q_2^{-1} A_{22} & A_{23}^* Q_2^{-1} A_{23} \end{pmatrix}^{-1}. \end{cases}$$

If we use the notation:

$$Q_{\hat{x}^{(1)}}^{-1} = \begin{pmatrix} A_{11}^* Q_1^{-1} A_{11} & A_{11}^* Q_1^{-1} A_{12} \\ A_{12}^* Q_1^{-1} A_{11} & A_{12}^* Q_1^{-1} A_{12} \end{pmatrix} = \begin{pmatrix} (Q_{\hat{x}^{(1)}}^{-1})_{11} & (Q_{\hat{x}^{(1)}}^{-1})_{12} \\ (Q_{\hat{x}^{(1)}}^{-1})_{21} & (Q_{\hat{x}^{(1)}}^{-1})_{22} \end{pmatrix}$$

and a similar notation for $Q_{\hat{x}^{(2)}}$, it follows with (119) that (117) may also be written as:

$$(120) \quad \begin{pmatrix} \hat{x}_1 \\ \hat{x}_2 \\ \hat{x}_3 \end{pmatrix} = \begin{pmatrix} (Q_{\hat{x}^{(1)}}^{-1})_{11} & (Q_{\hat{x}^{(1)}}^{-1})_{12} & 0 \\ (Q_{\hat{x}^{(1)}}^{-1})_{21} & \sum_{i=1}^2 (Q_{\hat{x}^{(1)}}^{-1})_{22} & (Q_{\hat{x}^{(2)}}^{-1})_{23} \\ 0 & (Q_{\hat{x}^{(2)}}^{-1})_{32} & (Q_{\hat{x}^{(2)}}^{-1})_{33} \end{pmatrix}^{-1} \begin{pmatrix} \sum_{i=1}^2 (Q_{\hat{x}^{(1)}}^{-1})_{1i} \hat{x}_i^{(1)} \\ \sum_{i=1}^2 (Q_{\hat{x}^{(1)}}^{-1})_{2i} \hat{x}_i^{(1)} + \sum_{i=2}^3 (Q_{\hat{x}^{(2)}}^{-1})_{2i} \hat{x}_i^{(2)} \\ \sum_{i=2}^3 (Q_{\hat{x}^{(2)}}^{-1})_{3i} \hat{x}_i^{(2)} \end{pmatrix}$$

But this is also the solution of:

$$(121) \quad E \begin{pmatrix} \hat{x}_1^{(1)} \\ \hat{x}_2^{(1)} \\ \hat{x}_2^{(2)} \\ \hat{x}_3^{(2)} \end{pmatrix} = \begin{pmatrix} I & 0 & 0 \\ 0 & I & 0 \\ 0 & I & 0 \\ 0 & 0 & I \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} ; \quad D \begin{pmatrix} \hat{x}_1^{(1)} \\ \hat{x}_2^{(1)} \\ \hat{x}_2^{(2)} \\ \hat{x}_3^{(2)} \end{pmatrix} = \begin{pmatrix} (Q_{\hat{x}^{(1)}})_{11} & (Q_{\hat{x}^{(1)}})_{12} & 0 & 0 \\ (Q_{\hat{x}^{(1)}})_{21} & (Q_{\hat{x}^{(1)}})_{22} & 0 & 0 \\ 0 & 0 & (Q_{\hat{x}^{(2)}})_{22} & (Q_{\hat{x}^{(2)}})_{23} \\ 0 & 0 & (Q_{\hat{x}^{(2)}})_{32} & (Q_{\hat{x}^{(2)}})_{33} \end{pmatrix}.$$

This model is in terms of observation equations. Its equivalent form in terms of condition equations reads:

$$(122) \quad (0 \quad I - I \quad 0) E \begin{pmatrix} \hat{x}_1^{(1)} \\ \hat{x}_2^{(1)} \\ \hat{x}_2^{(2)} \\ \hat{x}_3^{(2)} \end{pmatrix} = 0 ; \quad D \begin{pmatrix} \hat{x}_1^{(1)} \\ \hat{x}_2^{(1)} \\ \hat{x}_2^{(2)} \\ \hat{x}_3^{(2)} \end{pmatrix} = \begin{pmatrix} (Q_{\hat{x}^{(1)}})_{11} & (Q_{\hat{x}^{(1)}})_{12} & 0 & 0 \\ (Q_{\hat{x}^{(1)}})_{21} & (Q_{\hat{x}^{(1)}})_{22} & 0 & 0 \\ 0 & 0 & (Q_{\hat{x}^{(2)}})_{22} & (Q_{\hat{x}^{(2)}})_{23} \\ 0 & 0 & (Q_{\hat{x}^{(2)}})_{32} & (Q_{\hat{x}^{(2)}})_{33} \end{pmatrix}$$

Note that in this model, $\hat{x}_1^{(1)}$ and $\hat{x}_3^{(2)}$ are free $\underline{y}^R$ -variates. The solution of (122) reads:

$$(123) \quad \begin{pmatrix} \hat{x}_1 \\ \hat{x}_2 \\ \hat{x}_2 \\ \hat{x}_3 \end{pmatrix} = \begin{pmatrix} I & 0 & 0 & 0 \\ 0 & I & 0 & 0 \\ 0 & 0 & I & 0 \\ 0 & 0 & 0 & I \end{pmatrix} - \begin{pmatrix} (Q_{\hat{x}^{(1)}})_{12} \\ (Q_{\hat{x}^{(1)}})_{22} \\ -(Q_{\hat{x}^{(2)}})_{22} \\ -(Q_{\hat{x}^{(2)}})_{32} \end{pmatrix} [(Q_{\hat{x}^{(1)}})_{22} + (Q_{\hat{x}^{(2)}})_{22}]^{-1} (0 \quad I - I \quad 0) \begin{pmatrix} \hat{x}_1^{(1)} \\ \hat{x}_2^{(1)} \\ \hat{x}_2^{(2)} \\ \hat{x}_3^{(2)} \end{pmatrix}$$

Note that the solution (123) generalizes solution (113). The above results are summarized in table 6.5.

$$(1) \quad E\begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} A_{11} & A_{12} & 0 \\ 0 & A_{22} & A_{23} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix}; \quad D\begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} Q_1 & 0 \\ 0 & Q_2 \end{pmatrix}.$$

Solving the *partial model*

$$(2) \quad E\{y_1\} = (A_{11} \ A_{12}) \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}; \quad D\{y_1\} = Q_1$$

gives:

$$(3) \quad \begin{pmatrix} \hat{x}_1^{(1)} \\ \hat{x}_2^{(1)} \end{pmatrix} \text{ and } \begin{pmatrix} (Q_{\hat{x}^{(1)}})_{11} & (Q_{\hat{x}^{(1)}})_{12} \\ (Q_{\hat{x}^{(1)}})_{21} & (Q_{\hat{x}^{(1)}})_{22} \end{pmatrix}.$$

Solving the *partial model*

$$E\{y_2\} = (A_{22} \ A_{23}) \begin{pmatrix} x_2 \\ x_3 \end{pmatrix}; \quad D\{y_2\} = Q_2$$

gives:

$$(4) \quad \begin{pmatrix} \hat{x}_2^{(2)} \\ \hat{x}_3^{(2)} \end{pmatrix} \text{ and } \begin{pmatrix} (Q_{\hat{x}^{(2)}})_{22} & (Q_{\hat{x}^{(2)}})_{23} \\ (Q_{\hat{x}^{(2)}})_{32} & (Q_{\hat{x}^{(2)}})_{33} \end{pmatrix}.$$

With (3) and (4) we formulate the model of *condition equations*:

$$(5) \quad \begin{pmatrix} 0 & I & -I & 0 \end{pmatrix} E\begin{pmatrix} \hat{x}_1^{(1)} \\ \hat{x}_2^{(1)} \\ \hat{x}_2^{(2)} \\ \hat{x}_3^{(2)} \end{pmatrix} = 0; \quad D\begin{pmatrix} \hat{x}_1^{(1)} \\ \hat{x}_2^{(1)} \\ \hat{x}_2^{(2)} \\ \hat{x}_3^{(2)} \end{pmatrix} = \begin{pmatrix} (Q_{\hat{x}^{(1)}})_{11} & (Q_{\hat{x}^{(1)}})_{12} & 0 & 0 \\ (Q_{\hat{x}^{(1)}})_{21} & (Q_{\hat{x}^{(1)}})_{22} & 0 & 0 \\ 0 & 0 & (Q_{\hat{x}^{(2)}})_{22} & (Q_{\hat{x}^{(2)}})_{23} \\ 0 & 0 & (Q_{\hat{x}^{(2)}})_{32} & (Q_{\hat{x}^{(2)}})_{33} \end{pmatrix}$$

Its solution reads:

$$(6) \quad \begin{cases} \hat{x}_1 = \hat{x}_1^{(1)} - (Q_{\hat{x}^{(1)}})_{12} [(Q_{\hat{x}^{(1)}})_{22} + (Q_{\hat{x}^{(2)}})_{22}]^{-1} (\hat{x}_2^{(1)} - \hat{x}_2^{(2)}) \\ \hat{x}_2 = \hat{x}_2^{(1)} - (Q_{\hat{x}^{(1)}})_{22} [(Q_{\hat{x}^{(1)}})_{22} + (Q_{\hat{x}^{(2)}})_{22}]^{-1} (\hat{x}_2^{(1)} - \hat{x}_2^{(2)}) \\ \hat{x}_2 = \hat{x}_2^{(2)} - (Q_{\hat{x}^{(2)}})_{22} [(Q_{\hat{x}^{(1)}})_{22} + (Q_{\hat{x}^{(2)}})_{22}]^{-1} (\hat{x}_2^{(2)} - \hat{x}_2^{(1)}) \\ \hat{x}_3 = \hat{x}_3^{(2)} - (Q_{\hat{x}^{(2)}})_{32} [(Q_{\hat{x}^{(1)}})_{22} + (Q_{\hat{x}^{(2)}})_{22}]^{-1} (\hat{x}_2^{(2)} - \hat{x}_2^{(1)}) \end{cases}$$

This is also the solution of (1).

Table 6.5: Block estimation

Example 3

Consider the levelling network of figure 6.6. Point 0 is taken as reference point.

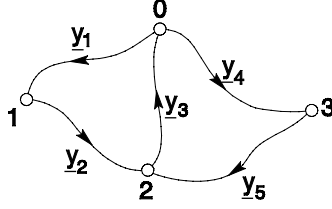


Figure 6.6: Levelling network

Its height is assumed to be known and equal to zero, that is $x_0=0$. The corresponding model of observation equations reads then:

$$(124) \quad E \begin{Bmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \\ y_5 \end{Bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ -1 & 1 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \\ 0 & 1 & -1 \end{bmatrix} \begin{Bmatrix} x_1 \\ x_2 \\ x_3 \end{Bmatrix}; \quad D \begin{Bmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \\ y_5 \end{Bmatrix} = \sigma^2 I_5.$$

The normal equations read:

$$(125) \quad \begin{pmatrix} 2 & -1 & 0 \\ -1 & 3 & -1 \\ 0 & -1 & 2 \end{pmatrix} \begin{pmatrix} \hat{x}_1 \\ \hat{x}_2 \\ \hat{x}_3 \end{pmatrix} = \begin{pmatrix} y_1 - y_2 \\ y_2 - y_3 + y_5 \\ y_4 - y_5 \end{pmatrix}.$$

In order to solve (125), we first eliminate $\hat{x}_1$ from the second normal equation. This is achieved if we premultiply (125) with the square and regular matrix:

$$\begin{pmatrix} 1 & 0 & 0 \\ \frac{1}{2} & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

This gives:

$$(126) \quad \begin{pmatrix} 2 & -1 & 0 \\ 0 & 2\frac{1}{2} & -1 \\ 0 & -1 & 2 \end{pmatrix} \begin{pmatrix} \hat{x}_1 \\ \hat{x}_2 \\ \hat{x}_3 \end{pmatrix} = \begin{pmatrix} y_1 - y_2 \\ \frac{1}{2}y_1 + \frac{1}{2}y_2 - y_3 + y_5 \\ y_4 - y_5 \end{pmatrix}.$$

The estimators $\hat{x}_2$ and $\hat{x}_3$ follow then from:

$$\begin{pmatrix} \hat{x}_2 \\ \hat{x}_3 \end{pmatrix} = \begin{pmatrix} 2\frac{1}{2} & -1 \\ -1 & 2 \end{pmatrix}^{-1} \begin{pmatrix} \frac{1}{2}y_1 + \frac{1}{2}y_2 - y_3 + y_5 \\ y_4 - y_5 \end{pmatrix}$$

as

$$(127) \quad \begin{pmatrix} \hat{x}_2 \\ \hat{x}_3 \end{pmatrix} = \frac{1}{8} \begin{pmatrix} 2y_1 + 2y_2 - 4y_3 + 2y_4 + 2y_5 \\ y_1 + y_2 - 2y_3 + 5y_4 - 3y_5 \end{pmatrix}.$$

Substitution of $\hat{x}_2$ into the first equation of (126) gives for $\hat{x}_1$:

$$(128) \quad \hat{x}_1 = \frac{1}{8} (5y_1 - 3y_2 - 2y_3 + y_4 + y_5).$$

With (127) and (128) we have found the solution of (124). Let us now try to solve (124) using the block estimation scheme of table 6.6. The two partial models read:

$$(129) \quad E \begin{pmatrix} y_1 \\ y_2 \\ y_3 \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ -1 & 1 \\ 0 & -1 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}; \quad D \begin{pmatrix} y_1 \\ y_2 \\ y_3 \end{pmatrix} = \sigma^2 I_3,$$

and:

$$(130) \quad E \begin{pmatrix} y_4 \\ y_5 \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ 1 & -1 \end{pmatrix} \begin{pmatrix} x_2 \\ x_3 \end{pmatrix}; \quad D \begin{pmatrix} y_4 \\ y_5 \end{pmatrix} = \sigma^2 I_2.$$

The system of normal equations for partial model (129) reads:

$$(131) \quad \begin{pmatrix} 2 & -1 \\ -1 & 2 \end{pmatrix} \begin{pmatrix} \hat{x}_{1^{(1)}} \\ \hat{x}_{2^{(1)}} \end{pmatrix} = \begin{pmatrix} y_1 - y_2 \\ y_2 - y_3 \end{pmatrix}.$$

From this it follows that:

$$(132) \quad \begin{pmatrix} \hat{x}_{1^{(1)}} \\ \hat{x}_{2^{(1)}} \end{pmatrix} = \frac{1}{3} \begin{pmatrix} 2y_1 - y_2 - y_3 \\ y_1 + y_2 - 2y_3 \end{pmatrix}; \quad Q_{\hat{x}^{(1)}} = \frac{1}{3} \sigma^2 \begin{pmatrix} 2 & 1 \\ 1 & 2 \end{pmatrix}.$$

The system of normal equations for partial model (130) reads:

$$(133) \quad \begin{pmatrix} 1 & -1 \\ -1 & 2 \end{pmatrix} \begin{pmatrix} \hat{x}_2^{(2)} \\ \hat{x}_3^{(2)} \end{pmatrix} = \begin{pmatrix} y_5 \\ y_4 - y_5 \end{pmatrix}.$$

From this it follows that:

$$(134) \quad \begin{pmatrix} \hat{x}_2^{(2)} \\ \hat{x}_3^{(2)} \end{pmatrix} = \begin{pmatrix} y_4 + y_5 \\ y_4 \end{pmatrix}; \quad Q_{x^{(2)}} = \sigma^2 \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}.$$

It will be clear that (132) is the solution of the partial levelling network shown in figure 6.7a, and that (134) is the solution of the partial levelling network shown in figure 6.7b.

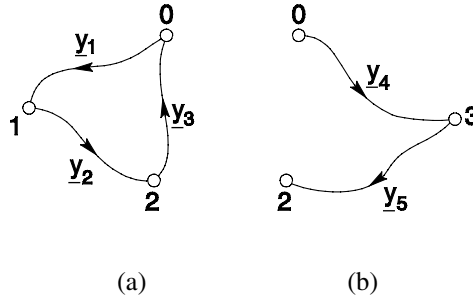


Figure 6.7: Two partial levelling networks

Now that the solutions of the two partial networks have been found, we can *connect* the two networks by formulating the model of condition equations:

$$(135) \quad \begin{pmatrix} 0 & 1 & -1 & 0 \end{pmatrix} E \left\{ \begin{pmatrix} \hat{x}_1^{(1)} \\ \hat{x}_2^{(1)} \\ \hat{x}_2^{(2)} \\ \hat{x}_3^{(2)} \end{pmatrix} \right\} = 0; \quad D \left\{ \begin{pmatrix} \hat{x}_1^{(1)} \\ \hat{x}_2^{(1)} \\ \hat{x}_2^{(2)} \\ \hat{x}_3^{(2)} \end{pmatrix} \right\} = \sigma^2 \begin{pmatrix} \frac{2}{3} & \frac{1}{3} & 0 & 0 \\ \frac{1}{3} & \frac{2}{3} & 0 & 0 \\ 0 & 0 & 2 & 1 \\ 0 & 0 & 1 & 1 \end{pmatrix}.$$

Note that the number of condition equations in (135) equals one. Thus the redundancy equals one. The redundancy in partial model (130) equals zero and the redundancy of partial model (129) equals one. Hence the total redundancy equals $1+0+1=2$. And this is of course equal to the number of linear independent condition equations that can be found in the levelling network of figure 6.6.

The solution of (135) follows as (see also (123)):

$$(136) \quad \begin{cases} \hat{x}_1 = \hat{x}_1^{(1)} - \frac{1}{3} \left(\frac{2}{3} + 2\right)^{-1} (\hat{x}_2^{(1)} - \hat{x}_2^{(2)}) , \\ \hat{x}_2 = \hat{x}_2^{(1)} - \frac{2}{3} \left(\frac{2}{3} + 2\right)^{-1} (\hat{x}_2^{(1)} - \hat{x}_2^{(2)}) , \\ \hat{x}_2 = \hat{x}_2^{(2)} - 2 \left(\frac{2}{3} + 2\right)^{-1} (\hat{x}_2^{(2)} - \hat{x}_2^{(1)}) , \\ \hat{x}_3 = \hat{x}_3^{(2)} - 1 \left(\frac{2}{3} + 2\right)^{-1} (\hat{x}_2^{(2)} - \hat{x}_2^{(1)}) . \end{cases}$$

We will now verify that solution (136) is indeed identical to (127) and (128). Substitution of $\hat{x}_1^{(1)}$, $\hat{x}_2^{(1)}$ and $\hat{x}_2^{(2)}$ from (132) and (134) into the first equation of (136) gives:

$$(137) \quad \begin{aligned} \hat{x}_1 &= \frac{1}{3}(2y_1 - y_2 - y_3) - \frac{1}{3} \frac{3}{8} [(\frac{1}{3}y_1 + \frac{1}{3}y_2 - \frac{2}{3}y_3) - (y_4 + y_5)] \\ &= \frac{1}{8}(5y_1 - 3y_2 - 2y_3 + y_4 + y_5). \end{aligned}$$

And this is indeed identical to (128). The verification of the remaining equations in (136) is left to the reader.

The block estimation scheme of table 6.5 that has been developed in this section, can also be derived from the mixed model representation of section 5.3. To see this, note that (116) is equivalent to:

$$(138) \quad E\left\{\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right\} = \begin{pmatrix} A_{11} & A_{12} & 0 & 0 \\ 0 & 0 & A_{22} & A_{23} \end{pmatrix} \begin{pmatrix} x_1^{(1)} \\ x_2^{(1)} \\ x_2^{(2)} \\ x_3^{(2)} \end{pmatrix} ; \quad (0 \quad I \quad -I \quad 0) \begin{pmatrix} x_1^{(1)} \\ x_2^{(1)} \\ x_2^{(2)} \\ x_3^{(2)} \end{pmatrix} = 0 .$$

But this representation is of the form:

$$(139) \quad E\{y\} = Ax ; B^*x = 0 .$$

Show for yourself that the block estimation scheme of table 6.5 follows when the results of section 5.3 are applied to model (138). According to (45) of section 5.3 we have

$$\hat{\underline{e}}^* Q_y^{-1} \hat{\underline{e}} = y^* Q_y^{-1} P_A^\perp y + \hat{x}_A^* B(B^* Q_{\hat{x}_A} B)^{-1} B^* \hat{x}_A .$$

If we "translate" this result to model (138) we get with $\hat{\underline{e}}_{1^{(1)}} = y_1 - \hat{y}_{1^{(1)}} = y_1 - A_{11}\hat{x}_{1^{(1)}} - A_{12}\hat{x}_{2^{(1)}}$ and $\hat{\underline{e}}_{2^{(2)}} = y_2 - \hat{y}_{2^{(2)}} = y_2 - A_{22}\hat{x}_{2^{(2)}} - A_{23}\hat{x}_{3^{(2)}}$:

$$(140) \quad \begin{aligned} \hat{\mathbf{e}}^* \mathbf{Q}_y^{-1} \hat{\mathbf{e}} &= \hat{\mathbf{e}}_1^{(1)*} \mathbf{Q}_1^{-1} \hat{\mathbf{e}}_1^{(1)} \\ &+ \hat{\mathbf{e}}_2^{(2)*} \mathbf{Q}_2^{-1} \hat{\mathbf{e}}_2^{(2)} + (\hat{\mathbf{x}}_2^{(1)} - \hat{\mathbf{x}}_2^{(2)})^* [(\mathbf{Q}_{\hat{\mathbf{x}}^{(1)}})_{22} + (\mathbf{Q}_{\hat{\mathbf{x}}^{(2)}})_{22}]^{-1} (\hat{\mathbf{x}}_2^{(1)} - \hat{\mathbf{x}}_2^{(2)}) \end{aligned}$$

This shows that the squared norm of the residuals of model (138) and thus of model (116), can be found from the sum of the squared norm of the residuals of the three models in (118) and (122).

6.6 The partitioned model $\begin{pmatrix} \mathbf{B}_1^* \\ \mathbf{B}_2^* \end{pmatrix} E\{\mathbf{y}\} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$

Estimation in phases

In this section we consider the partitioned model of *condition equations*:

$$(141) \quad \begin{pmatrix} \mathbf{B}_1^* \\ \mathbf{B}_2^* \end{pmatrix} E\{\mathbf{y}\} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} ; \quad D\{\mathbf{y}\} = \mathbf{Q}_y .$$

It will be shown that the solution of this model, which will be denoted by $\hat{\mathbf{y}}_B$, can be found by first solving the partial model:

$$(142) \quad \mathbf{B}_1^* E\{\mathbf{y}\} = 0 ; \quad D\{\mathbf{y}\} = \mathbf{Q}_y ,$$

which gives $\hat{\mathbf{y}}_{B_1}$ and $\mathbf{Q}_{\hat{\mathbf{y}}_{B_1}}$, and then solving for the model:

$$(143) \quad \mathbf{B}_2^* E\{\hat{\mathbf{y}}_{B_1}\} = 0 ; \quad D\{\mathbf{y}_{B_1}\} = \mathbf{Q}_{\hat{\mathbf{y}}_{B_1}} .$$

This solution method is known as the *method of estimation in phases* (fase vereffening). This method was originally introduced by the famous Dutch geodesist J.M. Tienstra (1895-1951). Tienstra was appointed professor in Mathematical Geodesy in 1935 at the Delft University of Technology.

With $B=(B_1; B_2)$, the solution of (141) reads:

$$(144) \quad \hat{\mathbf{y}}_B = [I - \mathbf{Q}_y \mathbf{B} (\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{B}^*] \mathbf{y} .$$

If we define the vector random variable $\hat{\mathbf{z}}$ as:

$$(145) \quad \hat{\mathbf{z}} = (\mathbf{B}^* \mathbf{Q}_y \mathbf{B})^{-1} \mathbf{B}^* \mathbf{y} ,$$

then $\hat{\mathbf{z}} = (\hat{\mathbf{z}}_1^*, \hat{\mathbf{z}}_2^*)^*$ is the solution of the partitioned system of these "normal equations":

$$(146) \quad \begin{pmatrix} \mathbf{B}_1^* \mathbf{Q}_y \mathbf{B}_1 & \mathbf{B}_1^* \mathbf{Q}_y \mathbf{B}_2 \\ \mathbf{B}_2^* \mathbf{Q}_y \mathbf{B}_1 & \mathbf{B}_2^* \mathbf{Q}_y \mathbf{B}_2 \end{pmatrix} \begin{pmatrix} \hat{\mathbf{z}}_1^* \\ \hat{\mathbf{z}}_2^* \end{pmatrix} = \begin{pmatrix} \mathbf{B}_1^* \mathbf{y} \\ \mathbf{B}_2^* \mathbf{y} \end{pmatrix} .$$

The estimator $\hat{\underline{z}}_1$ can be eliminated from the second set of normal equations if we premultiply (146) with the square and full rank matrix:

$$\begin{pmatrix} I & 0 \\ -\mathbf{B}_2^* \mathbf{Q}_y \mathbf{B}_1 (\mathbf{B}_1^* \mathbf{Q}_y \mathbf{B}_1)^{-1} & I \end{pmatrix}.$$

This gives:

$$(147) \quad \begin{pmatrix} \mathbf{B}_1^* \mathbf{Q}_y \mathbf{B}_1 & \mathbf{B}_1^* \mathbf{Q}_y \mathbf{B}_2 \\ 0 & \mathbf{B}_2^* [I - \mathbf{Q}_y \mathbf{B}_1 (\mathbf{B}_1^* \mathbf{Q}_y \mathbf{B}_1)^{-1} \mathbf{B}_1^*] \mathbf{Q}_y \mathbf{B}_2 \end{pmatrix} \begin{pmatrix} \hat{\underline{z}}_1 \\ \hat{\underline{z}}_2 \end{pmatrix} = \begin{pmatrix} \mathbf{B}_1^* \underline{y} \\ \mathbf{B}_2^* [I - \mathbf{Q}_y \mathbf{B}_1 (\mathbf{B}_1^* \mathbf{Q}_y \mathbf{B}_1)^{-1} \mathbf{B}_1^*] \underline{y} \end{pmatrix}$$

In (147) we recognize the orthogonal projector:

$$\mathbf{P}_{\mathbf{Q}_y \mathbf{B}_1}^\perp = [I - \mathbf{Q}_y \mathbf{B}_1 (\mathbf{B}_1^* \mathbf{Q}_y \mathbf{B}_1)^{-1} \mathbf{B}_1^*] .$$

Since:

$$(148) \quad \mathbf{P}_{\mathbf{Q}_y \mathbf{B}_1}^\perp \mathbf{Q}_y \mathbf{P}_{\mathbf{Q}_y \mathbf{B}_1}^{\perp*} = \mathbf{P}_{\mathbf{Q}_y \mathbf{B}_1}^\perp \mathbf{Q}_y = \mathbf{Q}_y \mathbf{P}_{\mathbf{Q}_y \mathbf{B}_1}^{\perp*} ,$$

we may write (147), using the abbreviation:

$$(149) \quad \bar{\mathbf{B}}_2 = \mathbf{P}_{\mathbf{Q}_y \mathbf{B}_1}^{\perp*} \mathbf{B}_2 ,$$

also as:

$$(150) \quad \begin{pmatrix} \mathbf{B}_1^* \mathbf{Q}_y \mathbf{B}_1 & \mathbf{B}_1^* \mathbf{Q}_y \mathbf{B}_2 \\ 0 & \bar{\mathbf{B}}_2^* \mathbf{Q}_y \bar{\mathbf{B}}_2 \end{pmatrix} \begin{pmatrix} \hat{\underline{z}}_1 \\ \hat{\underline{z}}_2 \end{pmatrix} = \begin{pmatrix} \mathbf{B}_1^* \underline{y} \\ \bar{\mathbf{B}}_2^* \underline{y} \end{pmatrix} .$$

The solution of this reduced system of normal equations follows then as:

$$(151) \quad \begin{cases} \hat{\underline{z}}_1 = (\mathbf{B}_1^* \mathbf{Q}_y \mathbf{B}_1)^{-1} \mathbf{B}_1^* (\underline{y} - \mathbf{Q}_y \mathbf{B}_2 \hat{\underline{z}}_2) , \\ \hat{\underline{z}}_2 = (\bar{\mathbf{B}}_2^* \mathbf{Q}_y \bar{\mathbf{B}}_2)^{-1} \bar{\mathbf{B}}_2^* \underline{y} . \end{cases}$$

With $\mathbf{B} = (\mathbf{B}_1 : \mathbf{B}_2)$ it follows from (144) and (145) that $\hat{\underline{y}}_{\mathbf{B}}$ can be written as:

$$\hat{\underline{y}}_{\mathbf{B}} = \underline{y} - \mathbf{Q}_y \mathbf{B}_1 \hat{\underline{z}}_1 - \mathbf{Q}_y \mathbf{B}_2 \hat{\underline{z}}_2 .$$

Substitution of (151) gives:

$$\hat{\underline{y}}_{\mathbf{B}} = [I - \mathbf{Q}_y \mathbf{B}_1 (\mathbf{B}_1^* \mathbf{Q}_y \mathbf{B}_1)^{-1} \mathbf{B}_1^*] (\underline{y} - \mathbf{Q}_y \mathbf{B}_2 \hat{\underline{z}}_2) ,$$

or:

$$\hat{\underline{y}}_{\mathbf{B}} = [I - \mathbf{Q}_y \mathbf{B}_1 (\mathbf{B}_1^* \mathbf{Q}_y \mathbf{B}_1)^{-1} \mathbf{B}_1^*] [I - \mathbf{Q}_y \mathbf{B}_2 (\bar{\mathbf{B}}_2^* \mathbf{Q}_y \bar{\mathbf{B}}_2)^{-1} \bar{\mathbf{B}}_2^*] \underline{y}$$

or, with (149):

$$\hat{\underline{y}}_{\mathbf{B}} = [\mathbf{P}_{\mathbf{Q}_y \mathbf{B}_1}^\perp - \mathbf{P}_{\mathbf{Q}_y \mathbf{B}_1}^\perp \mathbf{Q}_y \mathbf{B}_2 (\mathbf{B}_2^* \mathbf{P}_{\mathbf{Q}_y \mathbf{B}_1}^\perp \mathbf{Q}_y \mathbf{P}_{\mathbf{Q}_y \mathbf{B}_1}^{\perp*} \mathbf{B}_2)^{-1} \mathbf{B}_2^* \mathbf{P}_{\mathbf{Q}_y \mathbf{B}_1}^\perp] \underline{y}$$

or:

$$(152) \quad \hat{\underline{y}}_{\mathbf{B}} = [I - \mathbf{P}_{\mathbf{Q}_y \mathbf{B}_1}^\perp \mathbf{Q}_y \mathbf{B}_2 (\mathbf{B}_2^* \mathbf{P}_{\mathbf{Q}_y \mathbf{B}_1}^\perp \mathbf{Q}_y \mathbf{P}_{\mathbf{Q}_y \mathbf{B}_1}^{\perp*} \mathbf{B}_2)^{-1} \mathbf{B}_2^*] \mathbf{P}_{\mathbf{Q}_y \mathbf{B}_1}^\perp \underline{y} .$$

Clearly:

$$(153) \quad \hat{y}_{B_1} = P_{Q, B_1}^+ y ,$$

is the solution of (142). The variance matrix of $\hat{y}_{B_1}$ reads:

$$(154) \quad Q_{\hat{y}_{B_1}} = P_{Q, B_1}^+ Q_y P_{Q, B_1}^{+*} = P_{Q, B_1}^+ Q_y = Q_y P_{Q, B_1}^{+*} .$$

Substitution of (153) and (154) into (152) gives:

$$(155) \quad \hat{y}_B = [I - Q_{\hat{y}_{B_1}} B_2 (B_2^* Q_{\hat{y}_{B_1}} B_2)^{-1} B_2^*] \hat{y}_{B_1}$$

And this is clearly the solution of (143). Hence we have shown that the solution of the partitioned model of condition equations (141) can indeed be found by first solving for the partial model (142) and then solving for the model (143).

Let us now consider the least-squares residual vector $\hat{e}_B = y - \hat{y}_B$. From (155) it follows that:

$$(156) \quad \hat{e}_B = y - \hat{y}_{B_1} + Q_{\hat{y}_{B_1}} B_2 (B_2^* Q_{\hat{y}_{B_1}} B_2)^{-1} B_2^* \hat{y}_{B_1} .$$

With:

$$\hat{e}_{B_1} = y - \hat{y}_{B_1} ,$$

and:

$$\hat{e}_{B_2} = Q_{\hat{y}_{B_1}} B_2 (B_2^* Q_{\hat{y}_{B_1}} B_2)^{-1} B_2^* \hat{y}_{B_1} ,$$

this shows that:

$$(157) \quad \hat{e}_B = \hat{e}_{B_1} + \hat{e}_{B_2} .$$

Hence, the least-squares residual vector of the partitioned model (141) can be found from the sum of the least-squares residual vectors of the two separate models (142) and (143). From (157) it follows that:

$$(158) \quad \hat{e}_B^* Q_y^{-1} \hat{e}_B = \hat{e}_{B_1}^* Q_y^{-1} \hat{e}_{B_1} + 2 \hat{e}_{B_1}^* Q_y^{-1} \hat{e}_{B_2} + \hat{e}_{B_2}^* Q_y^{-1} \hat{e}_{B_2} .$$

Since:

$$\begin{aligned} \hat{e}_{B_1}^* Q_y^{-1} \hat{e}_{B_2} &= (y - \hat{y}_{B_1})^* Q_y^{-1} Q_{\hat{y}_{B_1}} B_2 (B_2^* Q_{\hat{y}_{B_1}} B_2)^{-1} B_2^* \hat{y}_{B_1} \\ &= y^* P_{Q, B_1}^* Q_y^{-1} Q_y P_{Q, B_1}^{+*} B_2 (B_2^* Q_{\hat{y}_{B_1}} B_2)^{-1} B_2^* \hat{y}_{B_1} \\ &= 0 , \end{aligned}$$

because:

$$P_{Q, B_1}^+ P_{Q, B_1} = 0 ,$$

it follows that (158) reduces to:

(159)

$$\hat{\underline{e}}_{\underline{B}}^* Q_y^{-1} \hat{\underline{e}}_{\underline{B}} = \hat{\underline{e}}_{\underline{B}_1}^* Q_y^{-1} \hat{\underline{e}}_{\underline{B}_1} + \hat{\underline{e}}_{\underline{B}_2}^* Q_y^{-1} \hat{\underline{e}}_{\underline{B}_2}$$

From the result it seems that one needs Q_y for computing the scalar $\hat{\underline{e}}_{\underline{B}_1}^* Q_y^{-1} \hat{\underline{e}}_{\underline{B}_1}$. If this would be the case, it would disrupt the method of estimation in phases, because it would imply that Q_y is also needed in the second phase, that is, when model (143) is solved for. Fortunately one can show that Q_y is *not* required for computing the scalar $\hat{\underline{e}}_{\underline{B}_2}^* Q_y^{-1} \hat{\underline{e}}_{\underline{B}_2}$. In order to show this we define the vectors of misclosures:

(160)

$$\begin{cases} \underline{t}_{\underline{B}_1} = \underline{B}_1^* \underline{y} & , \quad Q_{\underline{t}_{\underline{B}_1}} = \underline{B}_1^* Q_y \underline{B}_1 & , \\ \underline{t}_{\underline{B}_2} = \underline{B}_2^* \hat{\underline{y}}_{\underline{B}_1} & , \quad Q_{\underline{t}_{\underline{B}_2}} = \underline{B}_2^* Q_{\hat{\underline{y}}_{\underline{B}_1}} \underline{B}_2 & . \end{cases}$$

With (160) we have:

$$\hat{\underline{e}}_{\underline{B}_1} = Q_y \underline{B}_1 (\underline{B}_1^* Q_y \underline{B}_1)^{-1} \underline{B}_1^* \underline{y} = Q_y \underline{B}_1 Q_{\underline{t}_{\underline{B}_1}}^{-1} \underline{t}_{\underline{B}_1} .$$

From this it follows that:

(161)

$$\hat{\underline{e}}_{\underline{B}_1}^* Q_y^{-1} \hat{\underline{e}}_{\underline{B}_1} = \underline{t}_{\underline{B}_1}^* Q_{\underline{t}_{\underline{B}_1}}^{-1} \underline{t}_{\underline{B}_1} .$$

With (160) we also have:

(162)

$$\hat{\underline{e}}_{\underline{B}_2} = Q_{\hat{\underline{y}}_{\underline{B}_1}} \underline{B}_2 (\underline{B}_2^* Q_{\hat{\underline{y}}_{\underline{B}_1}} \underline{B}_2)^{-1} \underline{B}_2^* \hat{\underline{y}}_{\underline{B}_1} = Q_{\hat{\underline{y}}_{\underline{B}_1}} \underline{B}_2 Q_{\underline{t}_{\underline{B}_2}}^{-1} \underline{t}_{\underline{B}_2} .$$

Since:

$$Q_{\hat{\underline{y}}_{\underline{B}_1}} Q_y^{-1} Q_{\hat{\underline{y}}_{\underline{B}_1}} = P_{Q_y \underline{B}_1}^\perp Q_y Q_y^{-1} Q_y P_{Q_y \underline{B}_1}^{\perp *} = Q_{\hat{\underline{y}}_{\underline{B}_1}} \quad (\text{see (154)}),$$

it follows from (162) that:

(163)

$$\hat{\underline{e}}_{\underline{B}_2}^* Q_y^{-1} \hat{\underline{e}}_{\underline{B}_2} = \underline{t}_{\underline{B}_2}^* Q_{\underline{t}_{\underline{B}_2}}^{-1} \underline{t}_{\underline{B}_2} .$$

Substitution of (161) and (163) into (159) finally gives:

(164)

$$\hat{\underline{e}}_{\underline{B}}^* Q_y^{-1} \hat{\underline{e}}_{\underline{B}} = \underline{t}_{\underline{B}_1}^* Q_{\underline{t}_{\underline{B}_1}}^{-1} \underline{t}_{\underline{B}_1} + \underline{t}_{\underline{B}_2}^* Q_{\underline{t}_{\underline{B}_2}}^{-1} \underline{t}_{\underline{B}_2} .$$

A summary of the above results is given in table 6.6.

Partitioned model

$$\begin{pmatrix} B_1^* \\ B_2^* \end{pmatrix} E\{\underline{y}\} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} ; D\{\underline{y}\} = Q_y .$$

$$\hat{\underline{y}}_B = [I - Q_y B (B^* Q_y B)^{-1} B^*] \underline{y}$$

$$\hat{\underline{e}}_B = \underline{y} - \hat{\underline{y}}_B = Q_y B (B^* Q_y B)^{-1} B^* \underline{y}$$

$$\hat{\underline{e}}_B^* Q_y^{-1} \hat{\underline{e}}_B = \underline{y}^* B (B^* Q_y B)^{-1} B^* \underline{y} = \underline{t}_B^* Q_{\underline{t}_B}^{-1} \underline{t}_B, \text{ with } \underline{t}_B = B^* \underline{y} .$$

$$B_1^* E\{\underline{y}\} = 0 ; D\{\underline{y}\} = Q_y$$

$$\hat{\underline{y}}_{B_1} = [I - Q_y B_1 (B_1^* Q_y B_1)^{-1} B_1^*] \underline{y}$$

1.

$$\hat{\underline{e}}_{B_1} = \underline{y} - \hat{\underline{y}}_{B_1} = Q_y B_1 (B_1^* Q_y B_1)^{-1} B_1^* \underline{y}$$

$$\hat{\underline{e}}_{B_1}^* Q_y^{-1} \hat{\underline{e}}_{B_1} = \underline{y}^* B_1 (B_1^* Q_y B_1)^{-1} B_1^* \underline{y} = \underline{t}_{B_1}^* Q_{\underline{t}_{B_1}}^{-1} \underline{t}_{B_1}, \text{ with } \underline{t}_{B_1} = B_1^* \underline{y}$$

$$B_2^* E\{\hat{\underline{y}}_{B_1}\} = 0 ; D\{\hat{\underline{y}}_{B_1}\} = Q_{\hat{\underline{y}}_{B_1}}$$

$$\hat{\underline{y}}_B = [I - Q_{\hat{\underline{y}}_{B_1}} B_2 (B_2^* Q_{\hat{\underline{y}}_{B_1}} B_2)^{-1} B_2^*] \hat{\underline{y}}_{B_1}$$

2.

$$\hat{\underline{e}}_{B_2} = \hat{\underline{y}}_{B_1} - \hat{\underline{y}}_B = Q_{\hat{\underline{y}}_{B_1}} B_2 (B_2^* Q_{\hat{\underline{y}}_{B_1}} B_2)^{-1} B_2^* \hat{\underline{y}}_{B_1}$$

$$\hat{\underline{e}}_{B_2}^* Q_y^{-1} \hat{\underline{e}}_{B_2} = \hat{\underline{y}}_{B_1}^* B_2 (B_2^* Q_{\hat{\underline{y}}_{B_1}} B_2)^{-1} B_2^* \hat{\underline{y}}_{B_1} = \underline{t}_{B_2}^* Q_{\underline{t}_{B_2}}^{-1} \underline{t}_{B_2}, \text{ with } \underline{t}_{B_2} = B_2^* \hat{\underline{y}}_{B_1}$$

$$\begin{aligned} \hat{\underline{e}}_B &= \hat{\underline{e}}_{B_1} + \hat{\underline{e}}_{B_2} ; \quad \hat{\underline{e}}_B^* Q_y^{-1} \hat{\underline{e}}_B = \hat{\underline{e}}_{B_1}^* Q_y^{-1} \hat{\underline{e}}_{B_1} + \hat{\underline{e}}_{B_2}^* Q_y^{-1} \hat{\underline{e}}_{B_2} \\ &= \underline{t}_{B_1}^* Q_{\underline{t}_{B_1}}^{-1} \underline{t}_{B_1} + \underline{t}_{B_2}^* Q_{\underline{t}_{B_2}}^{-1} \underline{t}_{B_2} \end{aligned}$$

Table 6.6: Estimation in phases

Example 4

Consider the levelling network of figure 6.8.

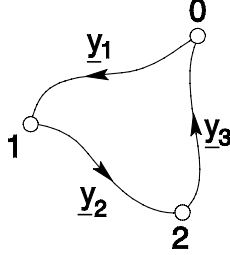


Figure 6.8: A one loop levelling network

The corresponding model of condition equations reads:

$$(165) \quad (1 \ 1 \ 1)E\begin{pmatrix} y_1 \\ y_2 \\ y_3 \end{pmatrix} = 0 ; D\begin{pmatrix} y_1 \\ y_2 \\ y_3 \end{pmatrix} = \sigma^2 I_3 .$$

Its solution reads:

$$\begin{pmatrix} \hat{y}_1 \\ \hat{y}_2 \\ \hat{y}_3 \end{pmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} - \sigma^2 I_3 \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \frac{1}{3} \sigma^{-2} (1 \ 1 \ 1) \begin{pmatrix} y_1 \\ y_2 \\ y_3 \end{pmatrix} ,$$

or:

$$(166) \quad \begin{pmatrix} \hat{y}_1 \\ \hat{y}_2 \\ \hat{y}_3 \end{pmatrix} = \frac{1}{3} \begin{pmatrix} 2y_1 - y_2 - y_3 \\ -y_1 + 2y_2 - y_3 \\ -y_1 - y_2 + 2y_3 \end{pmatrix} .$$

The variance matrix of (166) reads:

$$(167) \quad Q_{\hat{y}} = \frac{1}{3} \sigma^2 \begin{pmatrix} 2 & -1 & -1 \\ -1 & 2 & -1 \\ -1 & -1 & 2 \end{pmatrix} .$$

Now let us assume that the levelling network of figure 6.8 is enlarged with two additional height difference observables, such that the two loop levelling network of figure 6.9 results.

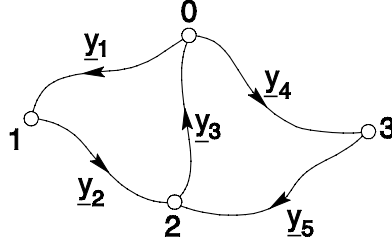


Figure 6.9: A two loop levelling network

The corresponding model of condition equations reads:

$$(168) \quad \begin{pmatrix} 1 & 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 & 1 \end{pmatrix} E \begin{pmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \\ y_5 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} ; \quad D \begin{pmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \\ y_5 \end{pmatrix} = \sigma^2 I_5 .$$

According to table 6.6, this model can be solved in two phases. The *first phase* would then correspond to solving the partial model:

$$(169) \quad (1 \ 1 \ 1 \ 0 \ 0) E \begin{pmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \\ y_5 \end{pmatrix} = 0 ; \quad D \begin{pmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \\ y_5 \end{pmatrix} = \sigma^2 I_5 .$$

Its solution reads:

$$\begin{pmatrix} \hat{y}_1 \\ \hat{y}_2 \\ \hat{y}_3 \\ \hat{y}_4 \\ \hat{y}_5 \end{pmatrix}_{B_1} = \begin{bmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} - \sigma^2 I_5 \begin{pmatrix} 1 \\ 1 \\ 1 \\ 0 \\ 0 \end{pmatrix} \frac{1}{3} \sigma^{-2} (1 \ 1 \ 1 \ 0 \ 0) \begin{pmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \\ y_5 \end{pmatrix} ,$$

or:

$$(170) \quad \begin{pmatrix} \hat{y}_1 \\ \hat{y}_2 \\ \hat{y}_3 \\ \hat{y}_4 \\ \hat{y}_5 \end{pmatrix}_{B_1} = \begin{pmatrix} \begin{pmatrix} 2y_1 - y_2 - y_3 \\ -y_1 + 2y_2 - y_3 \\ -y_1 - y_2 + 2y_3 \end{pmatrix} \\ \begin{pmatrix} y_4 \\ y_5 \end{pmatrix} \end{pmatrix}.$$

The variance matrix of (170) reads:

$$(171) \quad Q_{\hat{y}_{B_1}} = \sigma^2 \begin{pmatrix} \frac{1}{3} \begin{pmatrix} 2 & -1 & -1 \\ -1 & 2 & -1 \\ -1 & -1 & 2 \end{pmatrix} & \mathbf{0} \\ \mathbf{0} & \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \end{pmatrix}.$$

Comparison of (170) and (171) with (166) and (167) shows that the results of the first phase for the first three observables is identical to the solution of (165). Also note that the last two observables do not get a correction in the first phase. Hence, they do not play a role in the first phase. It should be recognized however that this is a consequence of the assumption that y_4 and y_5 are not correlated with y_1 , y_2 and y_3 . This is fortunately the case in most practical applications.

In order to obtain the final solution of (168) we have to solve in the *second phase* the partial model:

$$(172) \quad (\mathbf{0} \ \mathbf{0} \ 1 \ 1 \ 1) E \left\{ \begin{pmatrix} \hat{y}_1 \\ \hat{y}_2 \\ \hat{y}_3 \\ \hat{y}_4 \\ \hat{y}_5 \end{pmatrix}_{B_1} \right\} = \mathbf{0} \quad ; \quad D \left\{ \begin{pmatrix} \hat{y}_1 \\ \hat{y}_2 \\ \hat{y}_3 \\ \hat{y}_4 \\ \hat{y}_5 \end{pmatrix}_{B_1} \right\} = Q_{\hat{y}_{B_1}}.$$

With (170) and (171) the solution of (172) reads:

$$\begin{pmatrix} \hat{y}_1 \\ \hat{y}_2 \\ \hat{y}_3 \\ \hat{y}_4 \\ \hat{y}_5 \end{pmatrix} = \begin{bmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix} - \sigma^2 \begin{pmatrix} -\frac{1}{3} \\ -\frac{1}{3} \\ \frac{2}{3} \\ 1 \\ 1 \end{pmatrix} \frac{3}{8} \sigma^{-2} (0 \ 0 \ 1 \ 1 \ 1) \begin{bmatrix} \frac{1}{3} \begin{pmatrix} 2y_1 - y_2 - y_3 \\ -y_1 + 2y_2 - y_3 \\ -y_1 - y_2 + 2y_3 \end{pmatrix} \\ \begin{pmatrix} y_4 \\ y_5 \end{pmatrix} \end{bmatrix}.$$

or:

$$(173) \quad \begin{pmatrix} \hat{y}_1 \\ \hat{y}_2 \\ \hat{y}_3 \\ \hat{y}_4 \\ \hat{y}_5 \end{pmatrix} = \frac{1}{8} \begin{pmatrix} 5y_1 - 3y_2 - 2y_3 + y_4 + y_5 \\ -3y_1 + 5y_2 - 2y_3 + y_4 + y_5 \\ -2y_1 - 2y_2 + 4y_3 - 2y_4 - 2y_5 \\ y_1 + y_2 - 2y_3 + 5y_4 - 3y_5 \\ y_1 + y_2 - 2y_3 - 3y_4 + 5y_5 \end{pmatrix}.$$

It will be clear that if point 0 in figure 6.9 is taken as reference point with a height equal to zero, then $\hat{x}_1 = \hat{y}_1$, $\hat{x}_2 = -\hat{y}_3$ and $\hat{x}_3 = \hat{y}_4$. Compare now (173) with (127), (128) and also (137).

There exists a close correspondence between the results of this section and the results of section 6.2. In section 6.2 we considered the partitioned model:

$$(174) \quad E\{\underline{y}\} = \begin{pmatrix} A_1 & A_2 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} ; D\{\underline{y}\} = Q_y, \\ m \times 1 \quad m \times n_1 \quad m \times n_2 \quad (n_1 + n_2) \times 1 \quad m \times m \quad m \times m$$

and the partial model:

$$(175) \quad E\{\underline{y}\} = A_1 x_1 ; D\{\underline{y}\} = Q_y. \\ m \times 1 \quad m \times n_1 \quad n_1 \times 1 \quad m \times m \quad m \times m$$

In the present section we considered the partitioned model:

$$(176) \quad \begin{pmatrix} B_1^* \\ B_2^* \end{pmatrix} E\{\underline{y}\} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} ; D\{\underline{y}\} = Q_y,$$

and the partial model:

$$(177) \quad B_1^* E\{\underline{y}\} = 0 ; D\{\underline{y}\} = Q_y.$$

Since model (175) follows from model (174) by setting x_2 equal to zero, the *redundancy* of model (175) is higher than that of (174). In fact the redundancy of (174) is $m-n_1-n_2$, whereas the redundancy of (175) is $m-n_1$. Since the redundancy of model (175) is higher than that of (174), more condition equations can be formed for model (175) than for model (174). Hence, one may consider (176) to be the equivalent of model (175), and model (177) to be the equivalent of model (174). With $A=(A_1 : A_2)$ and $B=(B_1 : B_2)$ this gives:

$$(178) \quad \begin{cases} B^* A_1 = 0 & \text{for (176) and (175), and} \\ B_1^* A = 0 & \text{for (177) and (174).} \end{cases}$$

This implies that the solution of (176) is identical to the solution of (175), and that the solution of (177) is identical to the solution of (174). For the two models (174) and (175) we have the orthogonal decomposition:

$$(179) \quad \boxed{P_A = P_{A_1} + P_{\bar{A}_2}, \text{ with } \bar{A}_2 = P_{A_1}^\perp A_2} \quad (\text{see (39)})$$

In a similar way we have for the two models (176) and (177) the orthogonal decomposition:

$$(180) \quad \boxed{P_{Q,B} = P_{Q,B_1} + P_{Q,\bar{B}_2}, \text{ with } Q_y \bar{B}_2 = P_{Q,B_1}^\perp Q_y B_2}$$

With (179) and (180) it is now possible to relate the various orthogonal projectors to each other. Clearly:

$$(181) \quad \boxed{P_A = P_{Q,B_1}^\perp \text{ and } P_{A_1} = P_{Q,B}^\perp}$$

And with (179) and (180) this shows that:

$$(182) \quad \boxed{P_{\bar{A}_2} = P_{Q,\bar{B}_2}}$$

The above results can now be used to show the correspondence between the solutions of the four models (174), (175), (176) and (177). An overview of this correspondence is given in table 6.7.

With the results of table 6.7 we can now also give an alternative proof of the *method of estimation in phases*. Since:

$$P_A = P_{A_1} + P_{\bar{A}_2} \text{ and } P_{A_1} P_{\bar{A}_2} = 0 ,$$

it follows that:

$$P_{A_1} = P_{A_1} (P_{A_1} + P_{\bar{A}_2}) = P_{A_1} P_A .$$

$\begin{pmatrix} \mathbf{B}_1^* \\ \mathbf{B}_2^* \end{pmatrix} E\{\mathbf{y}\} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} \quad \mathbf{B}^* \mathbf{A}_1 = 0 \quad E\{\mathbf{y}\} = \mathbf{A}_1 x_1$		
$\begin{aligned} \hat{\mathbf{y}}_{\mathbf{B}} &= \mathbf{P}_{Q,\mathbf{B}}^\perp \mathbf{y} & \hat{\mathbf{y}}_{\mathbf{B}} &= \hat{\mathbf{y}}_{\mathbf{A}_1} & \hat{\mathbf{y}}_{\mathbf{A}_1} &= \mathbf{P}_{\mathbf{A}_1} \mathbf{y} \\ \hat{\mathbf{e}}_{\mathbf{B}} &= \mathbf{P}_{Q,\mathbf{B}} \mathbf{y} & \hat{\mathbf{e}}_{\mathbf{B}} &= \hat{\mathbf{e}}_{\mathbf{A}_1} & \hat{\mathbf{e}}_{\mathbf{A}_1} &= \mathbf{P}_{\mathbf{A}_1}^\perp \mathbf{y} \\ \ \hat{\mathbf{e}}_{\mathbf{B}}\ ^2 &= \ \mathbf{P}_{Q,\mathbf{B}} \mathbf{y}\ ^2 & \ \hat{\mathbf{e}}_{\mathbf{B}}\ ^2 &= \ \hat{\mathbf{e}}_{\mathbf{A}_1}\ ^2 & \ \hat{\mathbf{e}}_{\mathbf{A}_1}\ ^2 &= \ \mathbf{P}_{\mathbf{A}_1}^\perp \mathbf{y}\ ^2 \end{aligned}$		
$\mathbf{B}_1^* E\{\mathbf{y}\} = 0 \quad \mathbf{B}_1^* \mathbf{A} = 0 \quad E\{\mathbf{y}\} = (\mathbf{A}_1 : \mathbf{A}_2) \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$		
$\begin{aligned} \hat{\mathbf{y}}_{\mathbf{B}_1} &= \mathbf{P}_{Q,\mathbf{B}_1}^\perp \mathbf{y} & \hat{\mathbf{y}}_{\mathbf{B}_1} &= \hat{\mathbf{y}}_{\mathbf{A}} & \hat{\mathbf{y}}_{\mathbf{A}} &= \mathbf{P}_{\mathbf{A}} \mathbf{y} \\ \hat{\mathbf{e}}_{\mathbf{B}_1} &= \mathbf{P}_{Q,\mathbf{B}_1} \mathbf{y} & \hat{\mathbf{e}}_{\mathbf{B}_1} &= \hat{\mathbf{e}}_{\mathbf{A}} & \hat{\mathbf{e}}_{\mathbf{A}} &= \mathbf{P}_{\mathbf{A}}^\perp \mathbf{y} \\ \ \hat{\mathbf{e}}_{\mathbf{B}_1}\ ^2 &= \ \mathbf{P}_{Q,\mathbf{B}_1} \mathbf{y}\ ^2 & \ \hat{\mathbf{e}}_{\mathbf{B}_1}\ ^2 &= \ \hat{\mathbf{e}}_{\mathbf{A}}\ ^2 & \ \hat{\mathbf{e}}_{\mathbf{A}}\ ^2 &= \ \mathbf{P}_{\mathbf{A}}^\perp \mathbf{y}\ ^2 \end{aligned}$		
$\begin{aligned} \mathbf{P}_{Q,\mathbf{B}} &= \mathbf{P}_{Q,\mathbf{B}_1} + \mathbf{P}_{Q,\bar{\mathbf{B}}_2} & \mathbf{Q}_y \bar{\mathbf{B}}_2 &= \mathbf{P}_{Q,\mathbf{B}_1}^\perp \mathbf{Q}_y \mathbf{B}_2 & \mathbf{P}_{\mathbf{A}} &= \mathbf{P}_{\mathbf{A}_1} + \mathbf{P}_{\bar{\mathbf{A}}_2} & \bar{\mathbf{A}}_2 &= \mathbf{P}_{\mathbf{A}_1}^\perp \mathbf{A}_2 \\ \mathbf{P}_{\mathbf{A}} &= \mathbf{P}_{Q,\mathbf{B}_1}^\perp & \mathbf{P}_{\mathbf{A}_1} &= \mathbf{P}_{Q,\mathbf{B}}^\perp & \mathbf{P}_{\bar{\mathbf{A}}_2} &= \mathbf{P}_{Q,\bar{\mathbf{B}}_2} \end{aligned}$		
$\begin{aligned} \hat{\mathbf{y}}_{\mathbf{B}_1} &= \hat{\mathbf{y}}_{\mathbf{B}} + \mathbf{P}_{Q,\bar{\mathbf{B}}_2} \mathbf{y} & \hat{\mathbf{y}}_{\mathbf{A}} &= \hat{\mathbf{y}}_{\mathbf{A}_1} + \mathbf{P}_{\bar{\mathbf{A}}_2} \mathbf{y} \\ \hat{\mathbf{e}}_{\mathbf{B}_1} &= \hat{\mathbf{e}}_{\mathbf{B}} - \mathbf{P}_{Q,\bar{\mathbf{B}}_2} \mathbf{y} & \hat{\mathbf{e}}_{\mathbf{A}} &= \hat{\mathbf{e}}_{\mathbf{A}_1} - \mathbf{P}_{\bar{\mathbf{A}}_2} \mathbf{y} \\ \ \hat{\mathbf{e}}_{\mathbf{B}_1}\ ^2 &= \ \hat{\mathbf{e}}_{\mathbf{B}}\ ^2 - \ \mathbf{P}_{Q,\bar{\mathbf{B}}_2} \mathbf{y}\ ^2 & \ \hat{\mathbf{e}}_{\mathbf{A}}\ ^2 &= \ \hat{\mathbf{e}}_{\mathbf{A}_1}\ ^2 - \ \mathbf{P}_{\bar{\mathbf{A}}_2} \mathbf{y}\ ^2 \end{aligned}$		

Table 6.7: Correspondence between partitioned condition equations and partitioned observation equations

Substitution of $P_{A_1} = P_A - P_{A_2}^\perp$ in the right-hand side gives:

$$P_{A_1} = (P_A - P_{A_2}^\perp)P_A = (P_{A_2}^\perp - P_A^\perp)P_A .$$

But $P_A^\perp P_A = 0$. Hence:

$$(183) \quad P_{A_1} = P_{A_2}^\perp P_A .$$

With:

$$P_{A_1} = P_{Q_y B}^\perp , \quad P_{A_2}^\perp = P_{Q_y \bar{B}_2}^\perp , \quad P_A = P_{Q_y B_1}^\perp ,$$

equation (183) becomes:

$$(184) \quad \boxed{P_{Q_y B}^\perp = P_{Q_y \bar{B}_2}^\perp P_{Q_y B_1}^\perp}$$

The geometry of this relation is shown in figure 6.10.

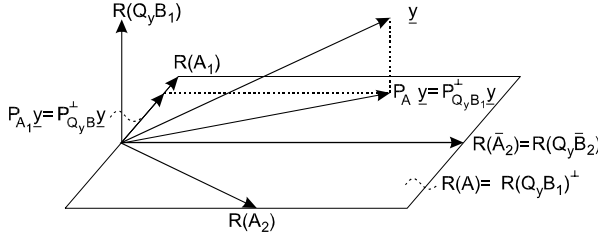


Figure 6.10: Geometry of $P_{Q_y B}^\perp = P_{Q_y \bar{B}_2}^\perp P_{Q_y B_1}^\perp$

Relation (184) already shows the two phases of the method of estimation in phases. Substitution of:

$$Q_y \bar{B}_2 = P_{Q_y B_1}^\perp Q_y B_2 \quad (\text{see (180)}),$$

into:

$$P_{Q_y \bar{B}_2}^\perp P_{Q_y B_1}^\perp = [I - Q_y \bar{B}_2 (\bar{B}_2^* Q_y \bar{B}_2)^{-1} \bar{B}_2^*] P_{Q_y B_1}^\perp ,$$

using:

$$P_{Q_y B_1}^\perp Q_y P_{Q_y B_1}^{\perp*} = P_{Q_y B_1}^\perp Q_y = Q_y P_{Q_y B_1}^{\perp*} ,$$

and:

$$P_{Q_y B_1}^\perp P_{Q_y B_1}^\perp = P_{Q_y B_1}^\perp ,$$

gives:

$$(185) \quad P_{Q_y \bar{B}_2}^\perp P_{Q_y B_1}^\perp = [I - P_{Q_y B_1}^\perp Q_y B_2 (B_2^* P_{Q_y B_1}^\perp Q_y B_2)^{-1} B_2^*] P_{Q_y B_1}^\perp .$$

Since the variance matrix of $\hat{\underline{y}}_{B_1} = P_{Q_y B_1}^\perp \underline{y}$ is given by:

$$Q_{\hat{\underline{y}}_{B_1}} = P_{Q_y B_1}^\perp Q_y ,$$

it follows that (185) may be written as:

$$P_{Q_y \bar{B}_2}^\perp P_{Q_y B_1}^\perp = [I - Q_{\hat{\underline{y}}_{B_1}} B_2 (B_2^* Q_{\hat{\underline{y}}_{B_1}} B_2)^{-1} B_2^*] P_{Q_y B_1}^\perp .$$

And this shows that:

$$(186) \quad P_{Q_y B}^\perp = [I - Q_{\hat{\underline{y}}_{B_1}} B_2 (B_2^* Q_{\hat{\underline{y}}_{B_1}} B_2)^{-1} B_2^*] [I - Q_y B_1 (B_1^* Q_y B_1)^{-1} B_1^*]$$

This concludes the proof of the method of estimation in phases.

7 Nonlinear models, linearization, iteration

7.1. The nonlinear A-model

7.1.1 Nonlinear observation equations

Up to this point the theoretical development was based on the assumption that the m -vector $E\{y\}$ is *linearly* related to the n -vector of unknown parameters x . In geodetic applications there are however only a few cases where this assumption truly holds. A typical example is levelling. In the majority of applications however the m -vector $E\{y\}$ is *nonlinearly* related to the n -vector of unknown parameters x . This implies that instead of the linear A-model (1) of section 6.1, we are generally dealing with a nonlinear model of observation equations:

$$(1) \quad \boxed{E\{y\} = A(x) ; D\{y\} = Q_y}$$

where $A(\cdot)$ is a nonlinear vector function from $\mathbb{R}^n$ into $\mathbb{R}^m$. The following definition makes clear when a vector function $A(\cdot) : \mathbb{R}^n \rightarrow \mathbb{R}^m$ is linear or nonlinear.

Definition: The vector function $A(\cdot) : \mathbb{R}^n \rightarrow \mathbb{R}^m$ is *linear* if and only if:

$$(2) \quad \begin{aligned} \text{a.} \quad & A(\alpha u) = \alpha A(u), \quad \forall \alpha \in \mathbb{R}, u \in \mathbb{R}^n \\ \text{b.} \quad & A(u+v) = A(u) + A(v), \quad \forall u, v \in \mathbb{R}^n \end{aligned}$$

The vector function $A(\cdot)$ is said to be *nonlinear* if it does not satisfy (2).

Example 1

Verify yourself that the vector function:

$$A(x) = \begin{pmatrix} x_1 + 2x_2 \\ 4x_2 \\ 3x_1 + x_2 \end{pmatrix} = \begin{pmatrix} 1 & 2 \\ 0 & 4 \\ 3 & 1 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix},$$

is linear and from $\mathbb{R}^2$ into $\mathbb{R}^3$.

Example 2

Verify yourself that the vector function:

$$A(x) = \begin{pmatrix} x_1 + (x_2)^2 \\ x_2 \\ (x_1^2) + \tan x_2 \end{pmatrix},$$

is nonlinear and from $\mathbb{R}^2$ into $\mathbb{R}^3$.

Example 3

Consider the levelling network of figure 7.1.

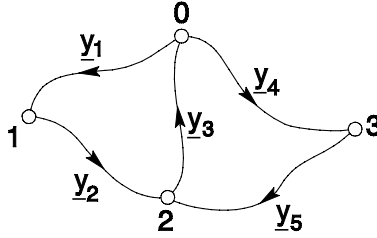


Figure 7.1: Levelling network

Point 0 is taken as reference point. Its height is assumed to be known and equal to zero, that is, $x_0=0$. The observed height differences are assumed to be uncorrelated and to have equal variances σ^2 . The corresponding model of observation equations reads then:

$$(3) \quad \underbrace{E\left\{ \begin{pmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \\ y_5 \end{pmatrix} \right\}}_y = \underbrace{\begin{pmatrix} 1 & 0 & 0 \\ -1 & 1 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \\ 0 & 1 & -1 \end{pmatrix}}_{A(x)} \underbrace{\begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix}}_{x} ; \quad D\left\{ \begin{pmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \\ y_5 \end{pmatrix} \right\} = \sigma^2 I_5.$$

Verify yourself that the vector function $A(x)$ of (3) is linear and from $\mathbb{R}^3$ into $\mathbb{R}^5$.

Example 4

Consider the configuration of figure 7.2a. The Cartesian x, y coordinates of the three points 1, 2 and 3 are known and the two Cartesian coordinates x_4 and y_4 of point 4 are unknown. The observables consist of the three distance variates L_{14} , L_{24} and L_{34} . These observables are assumed to be uncorrelated and to have equal variances σ^2 . Since distance and coordinates are related as (see figure 7.2b):

$$l_{ij} = [(x_j - x_i)^2 + (y_j - y_i)^2]^{1/2} = (x_{ij}^2 + y_{ij}^2)^{1/2},$$

the model of observation equations for the configuration of figure 7.2a reads:

$$(4) \quad \underbrace{E \begin{pmatrix} L_{14} \\ L_{24} \\ L_{34} \end{pmatrix}}_y = \underbrace{\begin{pmatrix} (x_{14}^2 + y_{14}^2)^{1/2} \\ (x_{24}^2 + y_{24}^2)^{1/2} \\ (x_{34}^2 + y_{34}^2)^{1/2} \end{pmatrix}}_{A(x)} ; \quad D \begin{pmatrix} L_{14} \\ L_{24} \\ L_{34} \end{pmatrix} = \sigma^2 I_3.$$

This model consists of 3 nonlinear observation equations in the 2 unknown parameters x_4 and y_4 . Hence, the vector function $A(x)$ of (4) is nonlinear and from $\mathbb{R}^2$ into $\mathbb{R}^3$.

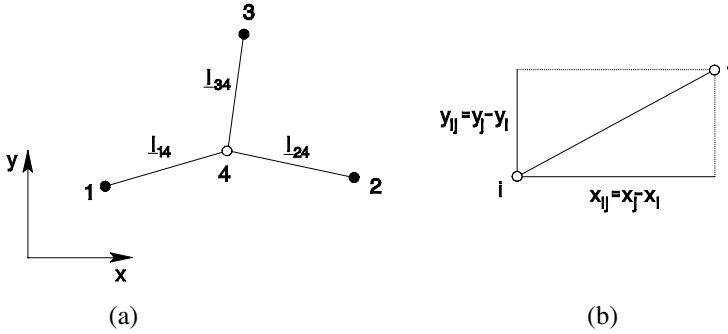


Figure 7.2: Distance resection

Example 5

Consider the parabola of figure 7.3a.

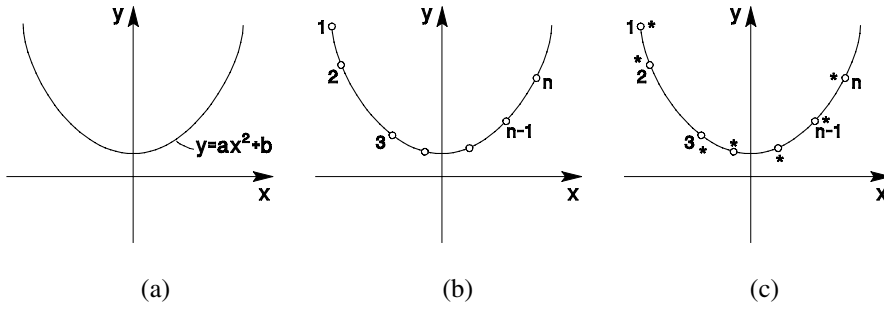


Figure 7.3

(a) The parabola $y=ax^2+b$
 (b) The points (o) on the parabola that need to be digitized.
 (c) The sample values (*) of the digitized points

We are asked to determine the values of the parameters a and b of the parabola $y=ax^2+b$. In order to solve this problem we decide to measure the x, y coordinates of an n -number of points on the parabola. The coordinates of the points $1, 2, \dots, n$ (see figure 7.3b) are measured with a digitizer. These coordinate observables $(\underline{x}_i, \underline{y}_i)$, $i=1, \dots, n$, are assumed to be uncorrelated and to have equal variances σ^2 . The observed or sample values of the coordinates are shown in figure 7.3c. The model of observation equations for this problem reads:

$$(5) \quad \underbrace{E\left\{\begin{pmatrix} \underline{x}_1 \\ \vdots \\ \underline{x}_n \\ \underline{y}_1 \\ \vdots \\ \underline{y}_n \end{pmatrix}\right\}}_{\underline{y}} = \underbrace{\begin{pmatrix} x_1 \\ \vdots \\ x_n \\ ax_1^2+b \\ \vdots \\ ax_n^2+b \end{pmatrix}}_{A(x)} ; \quad D\left\{\begin{pmatrix} \underline{x}_1 \\ \vdots \\ \underline{x}_n \\ \underline{y}_1 \\ \vdots \\ \underline{y}_n \end{pmatrix}\right\} = \sigma^2 I_{2n}.$$

This model consists of $2n$ observation equations in $(n+2)$ unknown parameters, namely $x_1, \dots, x_n$ and a, b . Note that the first n observation equations are linear, but that the second set of n observation equations are nonlinear. Hence, the vector function $A(x)$, with $x=(x_1, \dots, x_n, a, b)^*$, of (5) is nonlinear and from $\mathbb{R}^{n+2}$ into $\mathbb{R}^{2n}$.

Example 6

Consider figure 7.4. It shows an equilateral triangle. The position, orientation and size of this triangle are unknown. In order to determine these parameters, the x , y coordinates of the three points 1, 2 and 3 are observed.

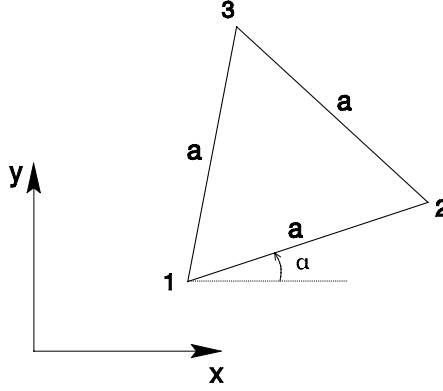


Figure 7.4: An equilateral triangle

The coordinate observables (x_i, y_i) , $i = 1, 2, 3$, are assumed to be uncorrelated and to have equal variances σ^2 . The corresponding model of observation equations reads:

$$(6) \quad \underbrace{E\left\{\begin{pmatrix} x_1 \\ y_1 \\ x_2 \\ y_2 \\ x_3 \\ y_3 \end{pmatrix}\right\}}_y = \underbrace{\begin{pmatrix} x_1 \\ y_1 \\ x_1 + a \cos \alpha \\ y_1 + a \sin \alpha \\ x_1 + a \cos(\alpha + \frac{1}{3}\pi) \\ y_1 + a \sin(\alpha + \frac{1}{3}\pi) \end{pmatrix}}_{A(x)} ; \quad D\left\{\begin{pmatrix} x_1 \\ y_1 \\ x_2 \\ y_2 \\ x_3 \\ y_3 \end{pmatrix}\right\} = \sigma^2 I_6.$$

This model consists of 6 observation equations in 4 unknown parameters, namely x_1 , y_1 , a and α . The vector function $A(x)$ of (6), with $x=(x_1, y_1, a, \alpha)^*$, is nonlinear and from $\mathbb{R}^4$ into $\mathbb{R}^6$.

7.1.2 Linearization: Taylor's theorem

In the previous section a number of examples were given of nonlinear A-models. We do know how to compute least-squares estimators in case of a linear A-model. But what about a nonlinear A-model? How should we compute the least-squares estimators if the model of observation equations is nonlinear? For the majority of non-linear problems the solution is to approximate the originally nonlinear A-model with a *linear(ized)* one. This approximation is made possible with the theorem of Taylor (see also appendix A). This section is therefore devoted to a discussion of Taylor's theorem.

Taylor's theorem for $f: \mathbb{R} \rightarrow \mathbb{R}$

Let $f: \mathbb{R} \rightarrow \mathbb{R}$ be a function of which all its q th-order derivatives exist and for which $\frac{d^q}{dx^q} f(x)$ is continuous. Then a scalar θ between x and x^0 exists such that:

(7)

$$f(x) = f(x^0) + \frac{d}{dx} f(x^0) \Delta x + \frac{1}{2!} \frac{d^2}{dx^2} f(x^0) \Delta x^2 + \dots + \frac{1}{(q-1)!} \frac{d^{q-1}}{dx^{q-1}} f(x^0) \Delta x^{q-1} + R_q(\theta, \Delta x),$$

with $\Delta x = x - x^0$ and with the remainder:

(8)

$$R_q(\theta, \Delta x) = \frac{1}{q!} \frac{d^q}{dx^q} f(\theta) \Delta x^q$$

It follows from (7) and (8) that for $q = 2$ we have:

$$(9) \quad f(x) = f(x^0) + \left(\frac{d}{dx} f(x^0) + \frac{1}{2} \frac{d^2}{dx^2} f(\theta) \Delta x \right) \Delta x.$$

Note that the second term within the brackets can be made arbitrarily small with respect to the first term within the brackets, by letting x^0 approach x . If this second term is negligible with respect to the first term, we may decide to ignore it, and approximate (9) as:

(10)

$$f(x) \doteq f(x^0) + \frac{d}{dx} f(x^0) \Delta x$$

In this case we speak of a *linear approximation* of $f(x)$ or a *linearization* of $f(x)$ at x^0 . A geometric interpretation of (10) is given in figure 7.5.

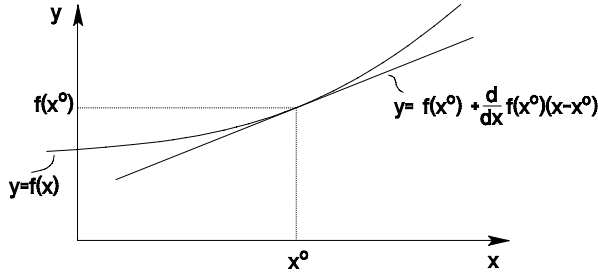


Figure 7.5: The curve $y=f(x)$ and its tangent line $y=f(x^0)+\frac{d}{dx}f(x^0)(x-x^0)$ at the point $(x^0, f(x^0))$

Example 7

The linearization of $f(x) = \tan x$ at x^0 reads:

$$f(x) = \tan x^0 + \frac{1}{\cos^2 x^0}(x-x^0) \quad .$$

Example 8

The linearization of $f(x)=ax^2+b$ at x^0 reads:

$$f(x) = [a(x^0)^2+b] + (2ax^0)(x-x^0) \quad .$$

We will now consider functions of more than one variable. A function $f(x_1, \dots, x_n)$ of the variables $x_1, \dots, x_n$ will be denoted as $f(x)$ with the n -vector $x=(x_1, \dots, x_n)^*$. The first-order partial derivatives of $f(x)$, $\frac{\partial}{\partial x_\alpha} f(x)$, $\alpha=1, \dots, n$, will be denoted as $\partial_\alpha f(x)$, $\alpha=1, \dots, n$. Similarly, we denote the second-order partial derivatives of $f(x)$, $\frac{\partial^2}{\partial x_\alpha \partial x_\beta} f(x)$, as $\partial_{\alpha\beta}^2 f(x)$, $\alpha, \beta=1, \dots, n$. And so on.

Taylor's Theorem for $f: \mathbb{R}^n \rightarrow \mathbb{R}$:

Let $f: \mathbb{R}^n \rightarrow \mathbb{R}$ be a function of which all its q th-order partial derivatives exist and for which $\partial_{\alpha_1 \dots \alpha_q}^q f(x)$ is continuous. Let $\Delta x = x - x^0$ and $\theta = x^0 + t(x - x^0)$ with $t \in \mathbb{R}$. Then a scalar $t \in (0, 1)$ exists such that:

(11)

$$f(x) = f(x^0) + \sum_{\alpha=1}^n \partial_{\alpha} f(x^0) \Delta x_{\alpha} + \frac{1}{2!} \sum_{\alpha=1}^n \sum_{\beta=1}^n \partial_{\alpha\beta}^2 f(x^0) \Delta x_{\alpha} \Delta x_{\beta} + \dots$$

$$\dots + \frac{1}{(q-1)!} \sum_{\alpha_1=1}^n \dots \sum_{\alpha_{q-1}=1}^n \partial_{\alpha_1 \dots \alpha_{q-1}}^{q-1} f(x^0) \Delta x_{\alpha_1} \dots \Delta x_{\alpha_{q-1}} + R_q(\theta, \Delta x)$$

with the remainder:

(12)

$$R_q(\theta, \Delta x) = \frac{1}{q!} \sum_{\alpha_1=1}^n \dots \sum_{\alpha_q=1}^n \partial_{\alpha_1 \dots \alpha_q}^q f(\theta) \Delta x_{\alpha_1} \dots \Delta x_{\alpha_q}$$

See appendix A for a proof of this theorem. Note that (11) and (12) reduce to (7) and (8) for $n = 1$.

For the case $q = 2$, it follows from (11) and (12) that:

(13)

$$f(x) = f(x^0) + \sum_{\alpha=1}^n \partial_{\alpha} f(x^0) \Delta x_{\alpha} + \frac{1}{2} \sum_{\alpha=1}^n \sum_{\beta=1}^n \partial_{\alpha\beta}^2 f(\theta) \Delta x_{\alpha} \Delta x_{\beta}.$$

If we introduce the *gradient vector* and *Hessian matrix* of $f(x)$ respectively as:

$$\partial_x f(x) = \begin{pmatrix} \partial_1 f(x) \\ \vdots \\ \partial_n f(x) \end{pmatrix}_{n \times 1} \quad \text{and} \quad \partial_{xx}^2 f(x) = \begin{pmatrix} \partial_{11}^2 f(x) & \dots & \partial_{1n}^2 f(x) \\ \vdots & & \vdots \\ \partial_{n1}^2 f(x) & \dots & \partial_{nn}^2 f(x) \end{pmatrix}_{n \times n},$$

then equation (13) may be written in the more compact matrix-vector form as:

$$f(x) = f(x^0) + \partial_x f(x^0)^* \Delta x + \frac{1}{2} \Delta x^* \partial_{xx}^2 f(\theta) \Delta x,$$

or as:

(14)

$$f(x) = f(x^0) + \left(\partial_x f(x^0) + \frac{1}{2} \partial_{xx}^2 f(\theta) \Delta x \right)^* \Delta x.$$

Note that this equation reduces to (9) in case $n = 1$. Equation (14) shows that a nonlinear function $f: \mathbb{R}^n \rightarrow \mathbb{R}$ can be written as the sum of three terms. The first term in this sum is the *zero-order term* $f(x^0)$. The zero-order term depends on x^0 but is independent of x . The second term in the sum of (14) is the *first-order term* $\partial_x f(x^0)^* \Delta x$. It depends on x^0 and is *linearly* dependent on x . Finally, the third term in the sum is the *second-order remainder* $R_2(\theta, \Delta x)$.

An important consequence of Taylor's theorem is that the remainder $R_2(\theta, \Delta x)$ can be made arbitrarily small by choosing x^0 close enough to x . Now assume that x^0 is indeed chosen close enough to x , such that the second-order remainder can be neglected. Then, instead of (14) we may write to a sufficient approximation:

(15)

$$f(x) \doteq f(x^0) + \partial_x f(x^0)^* \Delta x$$

Hence, if x^0 is sufficiently close to x , the *nonlinear* function $f: \mathbb{R}^n \rightarrow \mathbb{R}$ can be approximated to a sufficient degree by the function $f(x^0) + \partial_x f(x^0)^* \Delta x$ which is *linear* in x . This function is called the *linearization of $f(x)$ at x^0* . Note that (15) is the n -dimensional generalization of (10).

Example 9

The linearization of $f(x_1, x_2) = (x_1^2 + x_2^2)^{1/2}$ at (x_1^0, x_2^0) reads:

$$\underbrace{f(x_1, x_2)}_{f(x)} = \underbrace{((x_1^0)^2 + (x_2^0)^2)^{1/2}}_{f(x^0)} + \underbrace{\left(\frac{x_1^0}{((x_1^0)^2 + (x_2^0)^2)^{1/2}}, \frac{x_2^0}{((x_1^0)^2 + (x_2^0)^2)^{1/2}} \right)}_{\partial_x f(x^0)^*} \underbrace{\begin{pmatrix} \Delta x_1 \\ \Delta x_2 \end{pmatrix}}_{\Delta x}$$

Example 10

The linearization of $f(x_1, x_2) = x_1^2 + \sin x_2$ at (x_1^0, x_2^0) reads:

$$\underbrace{f(x_1, x_2)}_{f(x)} = \underbrace{(x_1^0)^2 + \sin x_2^0}_{f(x^0)} + \underbrace{(2x_1^0, \cos x_2^0)}_{\partial_x f(x^0)^*} \underbrace{\begin{pmatrix} \Delta x_1 \\ \Delta x_2 \end{pmatrix}}_{\Delta x}$$

7.1.3 The linearized A-model

In the previous section we have seen how Taylor's theorem can be applied to obtain a linear approximation of a nonlinear function. In this section we will see how the results of the previous section can be used to approximate the originally nonlinear A-model:

(16)

$$E\{\underline{y}\} = A(x); \quad D\{\underline{y}\} = Q_y$$

with its *linearized* version. Consider the nonlinear vector function $A(\cdot): \mathbb{R}^n \rightarrow \mathbb{R}^m$:

(17)

$$A(x) = \begin{pmatrix} a_1(x) \\ \vdots \\ a_m(x) \end{pmatrix}.$$

Each function $a_i(x)$, $i=1, \dots, m$, of (17) is a function from $\mathbb{R}^n$ into $\mathbb{R}$. Hence, each function $a_i(x)$, $i=1, \dots, m$ may be linearized according to (15). This gives:

(18)

$$A(x) = \begin{pmatrix} a_1(x^0) \\ \vdots \\ a_m(x^0) \end{pmatrix}_{m \times 1} + \begin{pmatrix} \partial_x a_1(x^0)^* \\ \vdots \\ \partial_x a_m(x^0)^* \end{pmatrix}_{m \times n} \Delta x_{n \times 1}.$$

If we denote the $m \times n$ matrix of (18) by $\partial_x A(x^0)$ and substitute (18) into (16) we get:

(19)

$$\begin{matrix} E\{\underline{y}\} & = & A(x^0) & + & \partial_x A(x^0) \Delta x & ; & D\{\underline{y}\} & = & Q_y . \\ m \times 1 & & m \times 1 & & m \times n & & n \times 1 & & m \times m & & m \times m \end{matrix}$$

If we bring the constant m -vector $A(x^0)$ of (19) to the left hand side of the equation and define $\Delta \underline{y} = \underline{y} - A(x^0)$ we finally obtain our linearized model of observation equations:

(20)

$$E\{\Delta \underline{y}\} = \partial_x A(x^0) \Delta x ; \quad D\{\Delta \underline{y}\} = Q_y$$

This is the *linearized A-model*. Compare (20) with (16) and with (1) of subsection 7.1.1. Note when comparing (20) with (1) that in the linearized A-model $\Delta \underline{y}$ takes the place of $\underline{y}$, $\partial_x A(x^0)$ takes the place of A and Δx takes the place of x . Since the linearized A-model is linear in $\Delta x = x - x^0$, our standard formulae for least-squares estimation can be applied again. This gives for the least-squares estimator $\hat{x} = x^0 + \Delta \hat{x}$

(21)

$$\hat{x} = x^0 + \left(\partial_x A(x^0)^* Q_y^{-1} \partial_x A(x^0) \right)^{-1} \left[\partial_x A(x^0)^* Q_y^{-1} (\underline{y} - A(x^0)) \right].$$

And application of the propagation law of variances to (21) gives:

$$(22) \quad Q_{\hat{x}} = \left(\partial_x A(x^0)^* Q_y^{-1} \partial_x A(x^0) \right)^{-1}.$$

It will be clear that the above results, (21) and (22), are approximate in the sense that the second-order remainder has been neglected in (20). But these approximations are good enough if the second-order remainder can be neglected to a sufficient degree. In this case also the optimality conditions of least-squares (e.g. unbiasedness, minimal variance) hold to a sufficient degree. A summary of the *linearized least-squares estimators* is given in table 7.1:

The non linear A-model	
$E\{\underline{y}\} = A(x) ; D\{\underline{y}\} = Q_y ; A(\cdot) : \mathbb{R}^n \rightarrow \mathbb{R}^m$	
The linearized A-model	
$E\{\Delta \underline{y}\} = \partial_x A(x^0) \Delta x ; D\{\Delta \underline{y}\} = Q_y$ $m \times 1 \quad m \times n \quad n \times 1$	
Least-squares estimators	
$\hat{x} = x^0 + \left(\partial_x A(x^0)^* Q_y^{-1} \partial_x A(x^0) \right)^{-1} \left(\partial_x A(x^0)^* Q_y^{-1} (y - A(x^0)) \right)$ $\hat{y} = A(\hat{x}) ; \hat{e} = y - \hat{y}$	
Variances	
$Q_{\hat{x}} = \left(\partial_x A(x^0)^* Q_y^{-1} \partial_x A(x^0) \right)^{-1}$ $Q_{\hat{y}} = \partial_x A(x^0) Q_{\hat{x}} \partial_x A(x^0)^*$ $Q_{\hat{e}} = Q_y - Q_{\hat{y}}$	

Table 7.1: Linearized least-squares estimation

Example 11

Consider the *nonlinear* A-model (4) of example 4:

$$(23) \quad \underbrace{E\left\{\begin{pmatrix} l_{14} \\ l_{24} \\ l_{34} \end{pmatrix}\right\}}_{\underline{y}} = \underbrace{\begin{pmatrix} (x_{14}^2 + y_{14}^2)^{1/2} \\ (x_{24}^2 + y_{24}^2)^{1/2} \\ (x_{34}^2 + y_{34}^2)^{1/2} \end{pmatrix}}_{A(x)} ; \quad D\left\{\begin{pmatrix} l_{14} \\ l_{24} \\ l_{34} \end{pmatrix}\right\} = \sigma^2 I_3.$$

The corresponding *linearized* A-model reads:

$$(24) \quad \underbrace{E\left\{\begin{pmatrix} \Delta l_{14} \\ \Delta l_{24} \\ \Delta l_{34} \end{pmatrix}\right\}}_{\Delta \underline{y}} = \underbrace{\begin{pmatrix} x_{14}^0/l_{14}^0 & y_{14}^0/l_{14}^0 \\ x_{24}^0/l_{24}^0 & y_{24}^0/l_{24}^0 \\ x_{34}^0/l_{34}^0 & y_{34}^0/l_{34}^0 \end{pmatrix}}_{\partial_x A(x^0)} \underbrace{\begin{pmatrix} \Delta x_4 \\ \Delta y_4 \end{pmatrix}}_{\Delta x} ; \quad D\left\{\begin{pmatrix} \Delta l_{14} \\ \Delta l_{24} \\ \Delta l_{34} \end{pmatrix}\right\} = \sigma^2 I_3.$$

with $x_{i4}^0 = x_4^0 - x_i$, $y_{i4}^0 = y_4^0 - y_i$, $l_{i4}^0 = ((x_{i4}^0)^2 + (y_{i4}^0)^2)^{1/2}$, $i=1,2,3$.

Up to this point nothing has been said about how one can obtain the approximate values of the parameters, which are needed to perform the linearization. In the majority of geodetic applications one can compute these approximate values of the parameters directly from the observations. In the above problem for instance, the approximate coordinates of point 4, x_4^0 and y_4^0 , can be computed from the two observed distances l_{14} , l_{24} and the known coordinates of the two points 1 and 2, x_1 , y_1 , x_2 and y_2 .

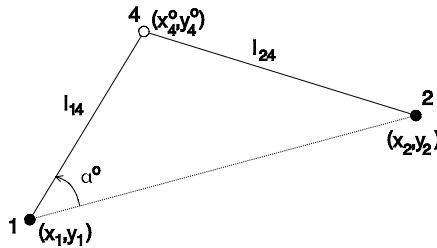


Figure 7.6: Computation of the approximate coordinates x_4^0 and y_4^0

The approximate angle α^0 (see figure 7.6) satisfies the cosine-rule:

$$l_{24}^2 = l_{14}^2 + (x_{12}^2 + y_{12}^2) - 2l_{14}(x_{12}^2 + y_{12}^2)^{1/2} \cos \alpha^0.$$

Hence:

$$(25) \quad \alpha^0 = \arccos \left(\frac{l_{14}^2 + (x_{12}^2 + y_{12}^2) - l_{24}^2}{2l_{14}(x_{12}^2 + y_{12}^2)^{1/2}} \right).$$

With this result the approximate coordinates of point 4 can be computed as:

$$(26) \quad \begin{pmatrix} x_4^0 \\ y_4^0 \end{pmatrix} = \begin{pmatrix} x_1 \\ y_1 \end{pmatrix} + \frac{l_{14}}{(x_{12}^2 + y_{12}^2)^{1/2}} \begin{pmatrix} \cos \alpha^0 & -\sin \alpha^0 \\ \sin \alpha^0 & \cos \alpha^0 \end{pmatrix} \begin{pmatrix} x_2 - x_1 \\ y_2 - y_1 \end{pmatrix}.$$

These approximate values x_4^0, y_4^0 can now be used to linearize (23) to obtain (24). Of course, there is still an ambiguity in the above computation. It is namely not clear whether the angle α^0 should be counted clockwise or counter clockwise. In expression (26) it has been assumed that α^0 should be counted counter clockwise. In practical applications this ambiguity usually does not occur, since one usually knows whether point 4 lies left or right from the line connecting the two points 1 and 2.

Example 12

Consider the *nonlinear* A-model (6) of example 6:

$$(27) \quad \underbrace{\begin{pmatrix} x_1 \\ y_1 \\ x_2 \\ y_2 \\ x_3 \\ y_3 \end{pmatrix}}_y = \underbrace{\begin{pmatrix} x_1 \\ y_1 \\ x_1 + a \cos \alpha \\ y_1 + a \sin \alpha \\ x_1 + a \cos(\alpha + \frac{1}{3}\pi) \\ y_1 + a \sin(\alpha + \frac{1}{3}\pi) \end{pmatrix}}_{A(x)} ; \underbrace{\begin{pmatrix} x_1 \\ y_1 \\ x_2 \\ y_2 \\ x_3 \\ y_3 \end{pmatrix}}_{D} = \sigma^2 I_6$$

The corresponding *linearized* A-model reads:

$$\underbrace{E\left\{\begin{pmatrix} \Delta x_1 \\ \Delta y_1 \\ \Delta x_2 \\ \Delta y_2 \\ \Delta x_3 \\ \Delta y_3 \end{pmatrix}\right\}}_{\Delta y} = \underbrace{\begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 1 & 0 & \cos\alpha^0 & -a^0\sin\alpha^0 \\ 0 & 1 & \sin\alpha^0 & a^0\cos\alpha^0 \\ 1 & 0 & \cos(\alpha^0 + \frac{1}{3}\pi) & -a^0\sin(\alpha^0 + \frac{1}{3}\pi) \\ 0 & 1 & \sin(\alpha^0 + \frac{1}{3}\pi) & a^0\cos(\alpha^0 + \frac{1}{3}\pi) \end{pmatrix}}_{\partial_x A(x^0)} \underbrace{\begin{pmatrix} \Delta x_1 \\ \Delta y_1 \\ \Delta x_2 \\ \Delta y_2 \\ \Delta x_3 \\ \Delta y_3 \end{pmatrix}}_{\Delta x} ; D\left\{\begin{pmatrix} \Delta x_1 \\ \Delta y_1 \\ \Delta x_2 \\ \Delta y_2 \\ \Delta x_3 \\ \Delta y_3 \end{pmatrix}\right\} = \sigma^2 I_6$$

(28)

The approximate values x_1^0 , y_1^0 , a^0 and α^0 of the four parameters can be computed from the observed coordinates as follows:

$$\begin{aligned}
 x_1^0 &= x_1, \quad a^0 = \left((x_2 - x_1)^2 + (y_2 - y_1)^2\right)^{1/2} \\
 y_1^0 &= y_1, \quad \alpha^0 = \arccos \frac{x_2 - x_1}{\left((x_2 - x_1)^2 + (y_2 - y_1)^2\right)^{1/2}}.
 \end{aligned}$$

(29)

7.1.4 Least-squares iteration

Up to this point it was assumed that x^0 was a good enough approximation such that the second-order remainder could be neglected. If this is not the case however, then $\hat{x}$ as computed by (21) is not the least-squares estimator and hence an unacceptable error is made. In order to repair this situation, we need to improve upon the approximation x^0 . It seems reasonable to expect that the estimate:

$$x^1 = x^0 + (\partial_x A(x^0)^* Q_y^{-1} \partial_x A(x^0))^{-1} \partial_x A(x^0)^* Q_y^{-1} (y - A(x^0))$$

is a better approximation than x^0 . That is, it seems reasonable to expect that x^1 is closer to the true least-squares estimate than x^0 . In fact one can show that this is indeed the case for most practical applications. But if x^1 is a better approximation than x^0 , a further improvement can be expected if we replace x^0 by x^1 in the linearization of the nonlinear A-model. The recomputed linearized least-squares estimate reads then:

$$x^2 = x^1 + (\partial_x A(x^1)^* Q_y^{-1} \partial_x A(x^1))^{-1} \partial_x A(x^1)^* Q_y^{-1} (y - A(x^1)).$$

Now x^2 can be expected to be a better approximation than x^1 . But this implies that a further improvement can be expected if we replace x^1 by x^2 in the linearization of the nonlinear A-

model. The recomputed linearized least-squares estimate, based on the approximation x^2 reads then:

$$x^3 = x^2 + (\partial_x A(x^2)^* Q_y^{-1} \partial_x A(x^2))^{-1} \partial_x A(x^2)^* Q_y^{-1} (y - A(x^2)) \quad .$$

By repeating this process a number of times, one can expect that finally the solution *converges* to the actual least-squares estimate $\hat{x}$. This process is called the *least-squares iteration* process. The iteration is usually terminated if the difference between successive solutions is negligible. A flow diagram of the least-squares iteration process is shown in the table on the following page.

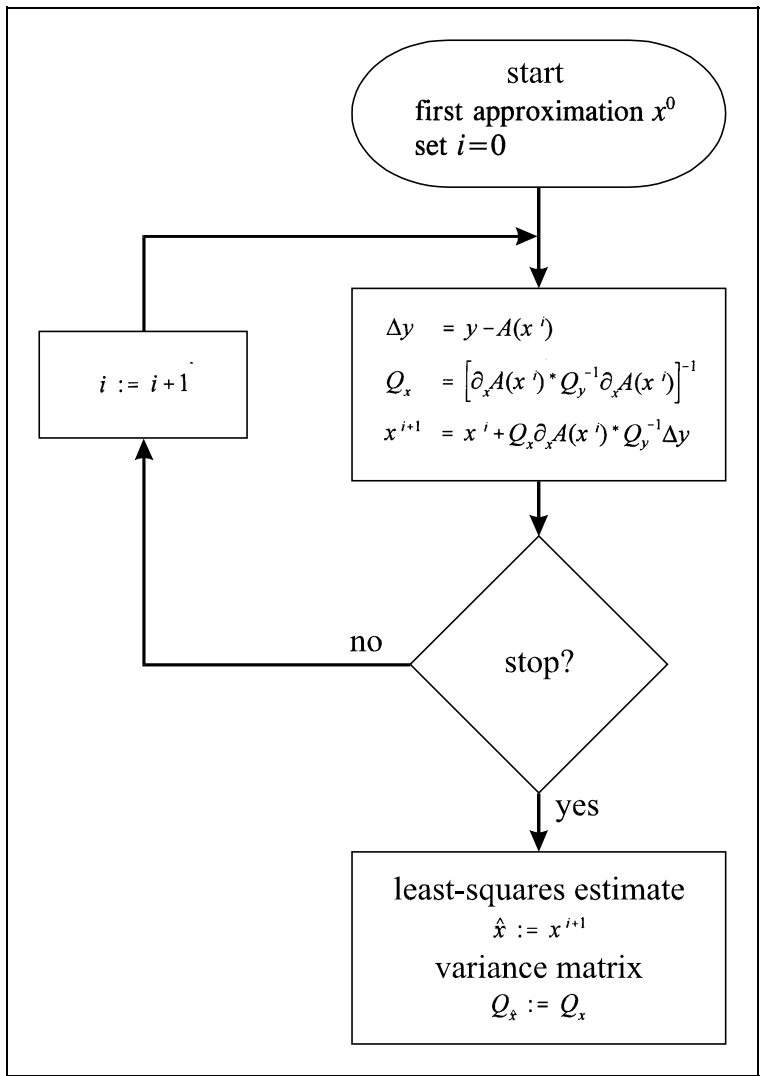


Table 7.2: Least-squares iteration

Example 13

Consider the *nonlinear* A-model:

$$(30) \quad \underbrace{E\left(\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right)}_{\underline{y}} = \underbrace{\begin{pmatrix} x^2 \\ x \end{pmatrix}}_{\mathbf{A}(x)} ; \underbrace{D\left(\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}\right)}_{\mathbf{Q}_y} = \sigma^2 I_2$$

It consists of two observation equations in one unknown parameter. Hence, $\mathbf{A}(\cdot): \mathbb{R} \rightarrow \mathbb{R}^2$. The corresponding *linearized* A-model reads:

$$(31) \quad \underbrace{E\left(\begin{pmatrix} \Delta y_1 \\ \Delta y_2 \end{pmatrix}\right)}_{\Delta \underline{y}} = \underbrace{\begin{pmatrix} 2x^0 \\ 1 \end{pmatrix} \Delta x}_{\partial_x \mathbf{A}(x^0)} ; \underbrace{D\left(\begin{pmatrix} \Delta y_1 \\ \Delta y_2 \end{pmatrix}\right)}_{\mathbf{Q}_y} = \sigma^2 I_2$$

Let us assume that we are asked to solve model (30). The two observations are given as:

$$(32) \quad y_1 = 4.1, \quad y_2 = 1.9.$$

In order to find the least-squares estimate $\hat{x}$ of x , we need the linearized model (31), a first approximation x^0 , and the iteration scheme:

$$(33) \quad x^{i+1} = x^i + (\partial_x \mathbf{A}(x^i)^* \mathbf{Q}_y^{-1} \partial_x \mathbf{A}(x^i))^{-1} \partial_x \mathbf{A}(x^i)^* \mathbf{Q}_y^{-1} (y - \mathbf{A}(x^i)), \text{ for } i=0,1,2,\dots$$

In our case, expression (33) becomes:

$$(34) \quad x^{i+1} = x^i + (1 + 4(x^i)^2)^{-1} [2x^i(y - (x^i)^2) + (y_2 - x^i)]$$

As a first approximation x^0 we take y_2 . Thus:

$$(35) \quad x^0 = y_2 = 1.9.$$

With (32), (35) and (34) we are now able to compute x^{i+1} for $i = 0, 1, 2, \dots$. The results of this least squares iteration process are given in the following table:

iteration step i	solution x^i
0	1.9
1	2.0205959
2	2.0176464
3	2.0176324
4	2.0176344
5	2.0176344

Note that the solution does not change much after iteration step $i = 2$. Hence, depending on the required *numerical* precision, any of these solutions can be taken as the least-squares estimate $\hat{x}$. For instance, if the solution is required to have 5 significant digits, then the least-squares solution can be taken as $\hat{x} = 2.01763$.

7.2 The nonlinear B-model and its linearization

Just like in case of the A-model, there are very few geodetic applications for which the model of condition equations is linear. In most cases the model of condition equations is nonlinear. The nonlinear B-model reads:

$$(36) \quad \boxed{B^*(E\{y\}) = 0 ; D\{y\} = Q_y},$$

where $B^*(.)$ is a nonlinear vector function from $\mathbb{R}^m$ into $\mathbb{R}^{m-n}$. The relation between the nonlinear B-model and the nonlinear A-model is given by:

$$(37) \quad \boxed{B^*(A(x)) = 0, \forall x \in \mathbb{R}^n}.$$

This is the *nonlinear generalization* of (7) of section 3.1. If we take the partial derivative with respect to x of (37) and apply the chain rule (kettingregel), we get:

$$(38) \quad \boxed{(\partial_y B(y^0))^* (\partial_x A(x^0)) = 0, \quad y^0 = A(x^0)}.$$

This is the *linearized* version of (37). Compare (38) with (7) of section 3.1. With (38) we are now in the position to construct the linearized B-model from the linearized A-model (20). Pre-multiplication of (20) with the matrix $(\partial_y B(y^0))^*$ gives together with (38) the result:

(39)

$$(\partial_y \mathbf{B}(y^0))^* E\{\Delta y\} = 0 ; D\{\Delta y\} = Q_y.$$

This is the *linearized B-model*. With (39) we are now in the position again to apply our standard least-squares estimation formulae.

As with the nonlinear A-model, also the solution of the nonlinear B-model needs an *iteration process*. But since the iteration process needed for solving the nonlinear B-model is quite involved, it will not be discussed here.

Example 14

Consider the *nonlinear A-model* (5) of example 5:

$$(40) \quad \underbrace{E\left\{\begin{pmatrix} x_1 \\ \vdots \\ x_n \\ y_1 \\ \vdots \\ y_n \end{pmatrix}\right\}}_y = \underbrace{\begin{pmatrix} x_1 \\ \vdots \\ x_n \\ ax_1^2 + b \\ \vdots \\ ax_n^2 + b \end{pmatrix}}_{A(x)} ; D\left\{\begin{pmatrix} x_1 \\ \vdots \\ x_n \\ y_1 \\ \vdots \\ y_n \end{pmatrix}\right\} = \sigma^2 I_{2n}.$$

This model consists of $2n$ observation equations in $(n+2)$ unknowns. Hence, the *redundancy* equals $2n - (n+2) = n-2$. Therefore, $(n-2)$ independent condition equations can be formulated. These $(n-2)$ independent nonlinear condition equations are:

$$\frac{E\{y_1\} - E\{y_j\}}{E\{y_2\} - E\{y_j\}} - \frac{E\{x_1\}^2 - E\{x_j\}^2}{E\{x_2\}^2 - E\{x_j\}^2} = 0, \quad j = 3, \dots, n,$$

or:

$$(41) \quad (E\{y_1\} - E\{y_j\})(E\{x_2\}^2 - E\{x_j\}^2) - (E\{y_2\} - E\{y_j\})(E\{x_1\}^2 - E\{x_j\}^2) = 0 \quad \text{for } j = 3, \dots, n.$$

In order to obtain the linearized B-model, we need to linearize (41). Linearization of (41) gives:

$$\underbrace{\begin{pmatrix} -2x_1^0 y_{32}^0 & 2x_2^0 y_{31}^0 & 2x_3^0 y_{12}^0 & & (x_2^{0^2} - x_3^{0^2}) & (x_3^{0^2} - x_1^{0^2}) & (x_1^{0^2} - x_2^{0^2}) \\ \vdots & \vdots & \ddots & & \vdots & \vdots & \ddots \\ \vdots & \vdots & & \ddots & \vdots & \vdots & \ddots \\ -2x_1^0 y_{n2}^0 & 2x_2^0 y_{n1}^0 & & 2x_n^0 y_{12}^0 & (x_2^{0^2} - x_n^{0^2}) & (x_n^{0^2} - x_1^{0^2}) & (x_1^{0^2} - x_2^{0^2}) \end{pmatrix}}_{\partial_y B(y^0)^*} E \left\{ \underbrace{\begin{pmatrix} \Delta x_1 \\ \Delta x_2 \\ \Delta x_3 \\ \vdots \\ \Delta x_n \\ \Delta y_1 \\ \Delta y_2 \\ \Delta y_3 \\ \vdots \\ \Delta y_n \end{pmatrix}}_{\Delta y} \right\} = \begin{pmatrix} 0 \\ \vdots \\ 0 \end{pmatrix}, \quad (42)$$

with $y_{i2}^0 = y_2^0 - y_i^0$, $y_{i1}^0 = y_1^0 - y_i^0$, $i=3, \dots, n$. This is the *linearized B-model*. The *linearized A-model* follows from linearizing (40) as:

$$\underbrace{E \left\{ \begin{pmatrix} \Delta x_1 \\ \Delta x_2 \\ \Delta x_3 \\ \vdots \\ \Delta x_n \\ \Delta y_1 \\ \Delta y_2 \\ \Delta y_3 \\ \vdots \\ \Delta y_n \end{pmatrix} \right\}}_{\Delta y} = \underbrace{\begin{pmatrix} 1 & & & & & & & & \\ & 1 & & & & & & & \\ & & 1 & & & & & & \\ & & & \ddots & & & & & \\ & & & & 1 & & & & \\ & 2a^0 x_1^0 & & & & (x_1^0)^2 & 1 & & \\ & & 2a^0 x_2^0 & & & (x_2^0)^2 & 1 & & \\ & & & 2a^0 x_3^0 & & (x_3^0)^2 & 1 & & \\ & & & & \ddots & \vdots & \vdots & & \\ & & & & & 2a^0 x_n^0 & (x_n^0)^2 & 1 \end{pmatrix}}_{\partial_x A(x^0)} \underbrace{\begin{pmatrix} \Delta x_1 \\ \Delta x_2 \\ \Delta x_3 \\ \vdots \\ \Delta x_n \\ \Delta a \\ \Delta b \end{pmatrix}}_{\Delta x}. \quad (43)$$

Verify yourself that $\left(\partial_y B(y^0) \right)^* \left(\partial_x A(x^0) \right) = 0$ indeed holds, if $y_i^0 = a^0 (x_i^0)^2 + b^0$

Example 15

Consider the *nonlinear* A-model (27) of example 12:

$$(44) \quad \underbrace{E\left\{\begin{pmatrix} x_1 \\ y_1 \\ x_2 \\ y_2 \\ x_3 \\ y_3 \end{pmatrix}\right\}}_{\mathbf{y}} = \underbrace{\begin{pmatrix} x_1 \\ y_1 \\ x_1 + a \cos \alpha \\ y_1 + a \sin \alpha \\ x_1 + a \cos(\alpha + \frac{1}{3}\pi) \\ y_1 + a \sin(\alpha + \frac{1}{3}\pi) \end{pmatrix}}_{\mathbf{A}(x)} ; \quad D\left\{\begin{pmatrix} x_1 \\ y_1 \\ x_2 \\ y_2 \\ x_3 \\ y_3 \end{pmatrix}\right\} = \sigma^2 I_6$$

This model consists of 6 observation equations in the four unknown parameters x_1 , y_1 , a , and α . Hence, the *redundancy* equals $6-4 = 2$. Therefore two independent condition equations can be formulated. These two nonlinear condition equations read:

$$(45) \quad \begin{cases} \left[(E\{x_2\} - E\{x_1\})^2 + (E\{y_2\} - E\{y_1\})^2 \right] - \left[(E\{x_3\} - E\{x_1\})^2 + (E\{y_3\} - E\{y_1\})^2 \right] = 0 \\ \left[(E\{x_2\} - E\{x_1\})^2 + (E\{y_2\} - E\{y_1\})^2 \right] - \left[(E\{x_3\} - E\{x_2\})^2 + (E\{y_3\} - E\{y_2\})^2 \right] = 0 \end{cases}$$

Note that these condition equations state that the sides of the equilateral triangle should be equal.

Linearization of (45) gives:

$$(46) \quad \underbrace{\begin{pmatrix} 2x_{23}^0 & 2y_{23}^0 & 2x_{12}^0 & 2y_{12}^0 & 2x_{31}^0 & 2y_{31}^0 \\ 2x_{21}^0 & 2y_{21}^0 & 2x_{13}^0 & 2y_{13}^0 & 2x_{32}^0 & 2y_{32}^0 \end{pmatrix}}_{\partial_y \mathbf{B}^*(y^0)} E \left\{ \underbrace{\begin{pmatrix} \Delta x_1 \\ \Delta y_1 \\ \Delta x_2 \\ \Delta y_2 \\ \Delta x_3 \\ \Delta y_3 \end{pmatrix}}_{\Delta y} \right\} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}.$$

Verify yourself that $(\partial_y \mathbf{B}(y^0))^* (\partial_x \mathbf{A}(x^0)) = 0$ indeed holds (see (28) of subsection 7.1.3), if the approximate values used in (46) satisfy the nonlinear condition equations.

Appendix A

Taylor's theorem with remainder

In this appendix we review the basic theorem of Taylor for functions of one or more variables. We will start with functions of one variable.

Definition: Let $f: D \subset \mathbb{R} \rightarrow \mathbb{R}$ be a function whose domain D is an open set $\mathbb{R}$. If all of the q th-order derivatives of $f(x)$ exist at every $x \in D$ and each $\frac{d^q}{dx^q} f(x)$ is a continuous function on D , then $f(x)$ is said to be a *function of class $C^{(q)}$ on D* .

Taylor's Theorem

Let $f(x)$ be a function of class $C^{(q)}$ on $D \subset \mathbb{R}$, $x, x^0 \in D$, and $\Delta x = x - x^0$. Then a scalar θ between x and x_0 exists such that:

$$(1) \quad \begin{aligned} f(x) = & f(x^0) + \frac{d}{dx} f(x^0) \Delta x + \frac{1}{2!} \frac{d^2}{dx^2} f(x^0) \Delta x^2 + \dots \\ & \dots + \frac{1}{(q-1)!} \frac{d^{q-1}}{dx^{q-1}} f(x^0) \Delta x^{q-1} + R_q(\theta, \Delta x) . \end{aligned}$$

with the remainder:

$$(2) \quad R_q(\theta, \Delta x) = \frac{1}{q!} \frac{d^q}{dx^q} f(\theta) \Delta x^q .$$

If we ignore the remainder $R_q(\theta, \Delta x)$, the right-hand side of (1) is a polynomial in Δx of degree $q-1$. If the remainder is small this polynomial furnishes an approximation to $f(x)$. We speak of a *linear approximation* or a *linearization of $f(x)$ at x^0* if $f(x)$ is approximated as:

$$(3) \quad f(x) \approx f(x^0) + \frac{d}{dx} f(x^0) (x - x^0) .$$

This approximation is acceptable if the second term within the square brackets of:

$$\left(\frac{d}{dx} f(x^0) + \frac{1}{2} \frac{d^2}{dx^2} f(\theta) \Delta x \right) \Delta x ,$$

is negligible with respect to the first term. Note that this second term can be made arbitrarily small by letting x^0 approach x . A geometric interpretation of (3) is given in figure A.1.

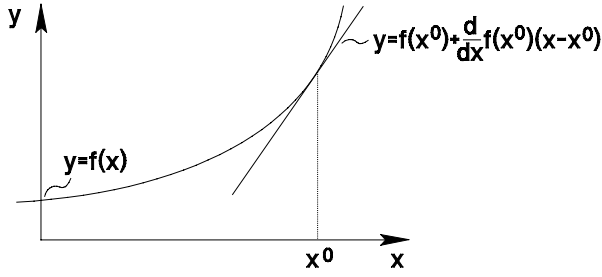


Figure A.1: The curve $y=f(x)$ and its tangentline $y=f(x^0)+\frac{d}{dx}f(x^0)(x-x^0)$ at $f(x^0)$

Examples

1. The linearization of $f(x)=\cos x$ at x^0 reads:

$$f(x) \doteq \cos x^0 - \sin x^0 (x - x^0) .$$

2. The linearization of $f(x)=(x^2+a^2)^{1/2}$ at x_0 reads:

$$f(x) \doteq ((x^0)^2 + a^2)^{1/2} + \frac{x^0}{((x^0)^2 + a^2)^{1/2}} (x - x^0) .$$

3. The linearization of $f(x)=x^3$ at x^0 reads:

$$f(x) \doteq (x^0)^3 + 3(x^0)^2 (x - x^0) .$$

We will now consider functions of more than one variable. A function $f(x_1, x_2, \dots, x_n)$ of the variables $x_1, x_2, \dots, x_n$ will be denoted as $f(x)$ with the n -vector $x=(x_1, x_2, \dots, x_n)^*$. The first-order partial derivatives of $f(x)$, $\frac{\partial}{\partial x_\alpha} f(x)$, $\alpha=1, \dots, n$, will be denoted as $\partial_\alpha f(x)$, $\alpha=1, \dots, n$. Similarly, we denote the second order partial derivatives of $f(x)$, $\frac{\partial^2}{\partial x_\alpha \partial x_\beta} f(x)$, $\alpha, \beta=1, \dots, n$, as $\partial_{\alpha\beta}^2 f(x)$, $\alpha, \beta=1, \dots, n$. And so on.

Definition: Let $f(x): D \subset \mathbb{R}^n \rightarrow \mathbb{R}$ be a function whose domain D is an open set of $\mathbb{R}^n$.

If all of the q th-order partial derivatives of $f(x)$ exist at every $x=(x_1, x_2, \dots, x_n)^* \in D$ and each $\partial_{\alpha_1 \dots \alpha_q}^q f(x)$ is a continuous function on D , then $f(x)$ is said to be a *function of class $C^{(q)}$* on D .

Taylor's Theorem: Let $f(x)$ be a function of class $C^{(q)}$ on $D \subset \mathbb{R}^n$ and $x, x_0 \in D$ such that the line segment joining x and x^0 is contained in D . In particular, if D is convex then x and x^0 can be any pair of points in D . Let $\Delta x = x - x^0$ and $\theta = x^0 + t(x - x^0)$ with $t \in \mathbb{R}$. Then a scalar $t \in (0, 1)$ exists such that:

(4)

$$\begin{aligned}
 f(x) = & f(x^0) + \sum_{\alpha=1}^n \partial_{\alpha} f(x^0) \Delta x_{\alpha} + \frac{1}{2!} \sum_{\alpha=1}^n \sum_{\beta=1}^n \partial_{\alpha\beta}^2 f(x^0) \Delta x_{\alpha} \Delta x_{\beta} + \dots \\
 & \dots + \frac{1}{(q-1)!} \sum_{\alpha_1=1}^n \dots \sum_{\alpha_{q-1}=1}^n \partial_{\alpha_1 \dots \alpha_{q-1}}^{q-1} f(x^0) \Delta x_{\alpha_1} \dots \Delta x_{\alpha_{q-1}} + R_q(\theta, \Delta x)
 \end{aligned}$$

with the remainder:

(5)

$$R_q(\theta, \Delta x) = \frac{1}{q!} \sum_{\alpha_1=1}^n \dots \sum_{\alpha_q=1}^n \partial_{\alpha_1 \dots \alpha_q}^q f(\theta) \Delta x_{\alpha_1} \dots \Delta x_{\alpha_q}$$

Proof: Define the functions $\theta(t)$ and $\phi(t)$ by:

$$\theta(t) = x^0 + t(x - x^0) \quad \text{and} \quad \phi(t) = f(\theta(t)) .$$

The domain of $\phi(t)$ is $\{t | x^0 + t(x - x^0) \in D\}$ which is an open set of $\mathbb{R}$ containing the closed interval $[0, 1]$. By repeated application of the *chain rule* we find that:

(6)

$$\begin{cases}
 \frac{d}{dt} \phi(t) &= \sum_{\alpha=1}^n \partial_{\alpha} f(\theta(t)) \Delta x_{\alpha} \\
 \frac{d^2}{dt^2} \phi(t) &= \sum_{\alpha=1}^n \sum_{\beta=1}^n \partial_{\alpha\beta}^2 f(\theta(t)) \Delta x_{\alpha} \Delta x_{\beta} \\
 &\vdots \\
 \frac{d^q}{dt^q} \phi(t) &= \sum_{\alpha_1=1}^n \dots \sum_{\alpha_q=1}^n \partial_{\alpha_1 \dots \alpha_q}^q f(\theta(t)) \Delta x_{\alpha_1} \dots \Delta x_{\alpha_q} .
 \end{cases}$$

By Taylor's theorem for functions of one variable there exists a $t \in (0, 1)$ such that:

(7)

$$\phi(1) = \phi(0) + \frac{d}{dt} \phi(0) + \dots + \frac{1}{(q-1)!} \frac{d^{q-1}}{dt^{q-1}} \phi(0) + \frac{1}{q!} \frac{d^q}{dt^q} \phi(t)$$

But $\phi(1) = f(x)$, $\phi(0) = f(x^0)$, and we have, by substitution of (6) into (7), Taylor's formula with remainder.

If we use the vector and matrix notation:

$$(8) \quad \begin{cases} \Delta x &= (\Delta x_1, \Delta x_2, \dots, \Delta x_n)^* \\ \partial_x f(x^0) &= (\partial_1 f(x^0), \partial_2 f(x^0), \dots, \partial_n f(x^0))^* \\ \partial_{xx}^2 f(x^0) &= \begin{pmatrix} \partial_{11}^2 f(x^0) & \dots & \partial_{1n}^2 f(x^0) \\ \vdots & \ddots & \vdots \\ \partial_{n1}^2 f(x^0) & \dots & \partial_{nn}^2 f(x^0) \end{pmatrix} \end{cases}$$

Taylor's formula with third-order remainder becomes:

$$(9) \quad \boxed{f(x) = f(x^0) + \partial_x f(x^0)^* \Delta x + \frac{1}{2} \Delta x^* \partial_{xx}^2 f(x^0) \Delta x + R_3(\theta, \Delta x)}$$

The vector $\partial_x f(x^0)$ is known as the *gradient* of $f(x)$ at x^0 . The matrix $\partial_{xx}^2 f(x^0)$ is known as the *Hessian matrix* of $f(x)$ at x^0 .

Examples:

1. The linearization of $f(x, y) = x \cos y$ at x^0, y^0 reads:

$$f(x, y) \doteq x^0 \cos y^0 + (\cos y^0 - x^0 \sin y^0) \begin{pmatrix} x - x^0 \\ y - y^0 \end{pmatrix}.$$

2. The quadratic approximation of $f(x, y) = x \cos y$ at x^0, y^0 reads:

$$f(x, y) \doteq x^0 \cos y^0 + (\cos y^0 - x^0 \sin y^0) \begin{pmatrix} x - x^0 \\ y - y^0 \end{pmatrix} + \frac{1}{2} \begin{pmatrix} x - x^0 \\ y - y^0 \end{pmatrix}^* \begin{pmatrix} 0 & -\sin y^0 \\ -\sin y^0 & -x^0 \cos y^0 \end{pmatrix} \begin{pmatrix} x - x^0 \\ y - y^0 \end{pmatrix}.$$

3. The quadratic approximation of $f(x, y) = (x^2 + y^2)^{1/2}$ at x^0, y^0 reads:

$$\begin{aligned} f(x, y) &\doteq (x^0)^2 + (y^0)^2)^{1/2} + ((x^0)^2 + (y^0)^2)^{-1/2} (x^0 \ y^0) \begin{pmatrix} x - x^0 \\ y - y^0 \end{pmatrix} \\ &+ \frac{1}{2} ((x^0)^2 + (y^0)^2)^{-3/2} \begin{pmatrix} x - x^0 \\ y - y^0 \end{pmatrix}^* \begin{pmatrix} (y^0)^2 & -x^0 y^0 \\ -x^0 y^0 & (x^0)^2 \end{pmatrix} \begin{pmatrix} x - x^0 \\ y - y^0 \end{pmatrix}. \end{aligned}$$

It is possible to give an explicit estimate of the remainder in Taylor's formula, in terms of bounds for the q th-order partial derivatives of $f(x)$.

Theorem: Suppose that K is a convex subset of D , such that $x, x^0 \in K$. Then also $\theta = x^0 + t(x - x^0) \in K$ for $t \in [0, 1]$. Moreover, suppose that all q th-order partial derivatives of $f(x)$ satisfy:

$$(10) \quad |\partial_{\alpha_1 \dots \alpha_q}^q f(x)| \leq M \text{ for all } x \in K.$$

Then the remainder in Taylor's formula satisfies:

$$(11) \quad |R_q(\theta, \Delta x)| \leq M \frac{n^{q/2}}{q!} \|\Delta x\|^q$$

Proof: From (5) and (10) it follows that:

$$|R_q(\theta, \Delta x)| \leq M \frac{1}{q!} \sum_{\alpha_1=1}^n \dots \sum_{\alpha_q=1}^n |\Delta x_{\alpha_1}| \dots |\Delta x_{\alpha_q}|$$

or that:

$$(12) \quad |R_q(\theta, \Delta x)| \leq M \frac{1}{q!} \left(\sum_{\alpha=1}^n |\Delta x_{\alpha}| \right)^q.$$

For the inner product of two vectors x and y we have:

$$x^* y = \|x\| \cdot \|y\| \cdot \cos \alpha.$$

From this expression the Cauchy-Schwarz inequality follows:

$$|x^* y| \leq \|x\| \cdot \|y\|.$$

This gives for $y = (1, 1, \dots, 1)^*$:

$$(13) \quad \sum_{\alpha=1}^n |x_{\alpha}| \leq n^{1/2} \|x\|.$$

If we use this inequality in (12) we obtain the estimate of (11).

Appendix B

Unconstrained optimization: optimality conditions

In this appendix we will formulate necessary and sufficient conditions for finding (local or global) solutions to the minimization problem:

$$(1) \quad \min_x F(x), \quad x \in \mathbb{R}^n, \quad F: \mathbb{R}^n \rightarrow \mathbb{R}.$$

In the investigation of the minimization problem (1) we distinguish two kinds of solution points: local minimizers and global minimizers.

Definition: The vector $\hat{x}$ is said to be a *local minimum* of $F(x)$ if $F(\hat{x}) \leq F(x) \quad \forall x \in B(\hat{x}, \varepsilon)$. The ε -ball, $B(\hat{x}, \varepsilon)$, is defined as $B(\hat{x}, \varepsilon) = \{x \in \mathbb{R}^n \mid \|x - \hat{x}\| < \varepsilon\}$. The local minimum is said to be *isolated* if $F(\hat{x}) < F(x) \quad \forall x \in B(\hat{x}, \varepsilon)$.

Definition: The vector $\hat{x}$ is said to be a *global minimum* of $F(x)$ if $F(\hat{x}) \leq F(x) \quad \forall x \in \mathbb{R}^n$. The global minimum is *unique* if $F(\hat{x}) < F(x) \quad \forall x \in \mathbb{R}^n$.

In the following we will assume that a (local or global) minimum of $F(x)$ exists. The problem of computing the minimum of $F(x)$ can be facilitated by deriving certain properties that must be satisfied by the minimizing vector $\hat{x}$. The following theorem states necessary conditions for $\hat{x}$ to be a minimum of $F(x)$.

Theorem (necessary conditions):

Let $F(x)$ be a class $C^{(3)}$. If $\hat{x}$ is a (local or global) minimum of $F(x)$, then:

$$(2) \quad \boxed{\begin{array}{l} 1. \quad \partial_x F(\hat{x}) = 0 \\ 2. \quad \partial_{xx}^2 F(\hat{x}) \geq 0 \end{array}}$$

Proof:

First we shall prove (2.1). Consider the vector $x = \hat{x} + td$ where t is a scalar and d an n -vector. Expansion of $F(x)$ in a Taylor series at $\hat{x}$ gives:

$$(3) \quad F(x) = F(\hat{x}) + t \partial_x F(\hat{x})^* d + O(t).$$

The order term $O(t)$ indicates the remainder in Taylor's formula. It has the property:

$$(4) \quad \lim_{t \rightarrow 0} \frac{O(t)}{t} = 0.$$

Since $\hat{x}$ is a (local or global) minimum of $F(x)$ by assumption, it follows that for sufficiently small t :

$$(5) \quad F(\hat{x}) \leq F(x) = F(\hat{x} + td).$$

This, together with (3) gives after dividing by t :

$$\begin{cases} \partial_x F(\hat{x})^* d + \frac{O(t)}{t} \geq 0 & \text{for } t > 0 \\ \partial_x F(\hat{x})^* d + \frac{O(t)}{t} \leq 0 & \text{for } t < 0 . \end{cases}$$

Taking the limit as $t \rightarrow 0$ and using (4) shows that (2.1) must be true.

Next we shall prove (2.2). Expanding $F(x)$ in a Taylor series at $\hat{x}$, but now retaining the quadratic terms, gives with (2.1):

$$(6) \quad F(x) = F(\hat{x}) + \frac{1}{2} t^2 d^* \partial_{xx}^2 F(\hat{x}) d + O(t^2) ,$$

where $O(t^2)$ has the property:

$$(7) \quad \lim_{t \rightarrow 0} \frac{O(t^2)}{t^2} = 0 .$$

Equations (5) and (6) imply that:

$$\frac{1}{2} d^* \partial_{xx}^2 F(\hat{x}) d + \frac{O(t^2)}{t^2} \geq 0 .$$

Taking the limit as $t \rightarrow 0$ and using (7) shows that:

$$(8) \quad \frac{1}{2} d^* \partial_{xx}^2 F(\hat{x}) d \geq 0 .$$

Since the n -vector d is arbitrary, (8) must hold for all $d \in \mathbb{R}^n$. This means by definition that the Hessian matrix $\partial_{xx}^2 F(\hat{x})$ is *positive semi-definite*. Positive semi-definite matrices A are denoted as " $A \geq 0$ ". This proves (2.2).

Examples

1. The function $F(x,y) = x^2 + xy + y^2$ has a minimum at $(\hat{x}, \hat{y}) = (0,0)$. According to the theorem, the gradient of $F(x,y)$ must be identical to the zero-vector at this point, and the Hessian matrix of $F(x,y)$ must be positive semi-definite at this point. This is easily verified, since:

$$\begin{pmatrix} \partial_x F(x,y) \\ \partial_y F(x,y) \end{pmatrix} = \begin{pmatrix} 2x + y \\ 2y + x \end{pmatrix} ; \quad \begin{pmatrix} \partial_{xx}^2 F(x,y) & \partial_{xy}^2 F(x,y) \\ \partial_{yx}^2 F(x,y) & \partial_{yy}^2 F(x,y) \end{pmatrix} = \begin{pmatrix} 2 & 1 \\ 1 & 2 \end{pmatrix} .$$

2. Let the $n \times n$ matrix Q be positive definite. A matrix Q is said to be *positive definite* if:

$$(9) \quad x^* Q x > 0 \quad \forall x \in \mathbb{R}^n \setminus \{0\}.$$

Since Q is positive definite it follows that the function $F: \mathbb{R}^n \rightarrow \mathbb{R}$

$$(10) \quad F(x) = x^* Q x,$$

has a minimum at $\hat{x}=0$. According to the theorem it must hold that $\partial_x F(\hat{x})=0$ and $\partial_{xx}^2 F(\hat{x}) \geq 0$. This is easily verified, since it follows from (10) that:

$$\partial_x F(x) = 2Qx \quad \text{and thus} \quad \partial_x F(0) = 2Q0 = 0,$$

and

$$\partial_{xx}^2 F(x) = 2Q \quad \text{and thus with (9):} \quad \partial_{xx}^2 F(\hat{x}=0) \geq 0.$$

The above theorem gives necessary conditions for $\hat{x}$ to be a minimum of $F(x)$. The stated conditions are however *not* sufficient. This is easily illustrated by the following simple example. Consider the function $F_1(x)=x^4$. Clearly it has a (global) minimum $\hat{x}=0$ and $\partial_x F_1(0)=0$ and $\partial_{xx}^2 F_1(0)=0$ hold. Consider now the function $F_2(x)=-x^4$. Then once again $\partial_x F_2(0)=0$ and $\partial_{xx}^2 F_2(0)=0$ hold. But now, however, $\hat{x}=0$ is a (global) maximum of $F_2(x)$. The following theorem gives sufficient conditions for $\hat{x}$ to be a minimum of $f(x)$.

Theorem (sufficient conditions):

Let $F(x)$ be of class $C^{(3)}$. If:

$$(11) \quad \boxed{\begin{array}{l} 1. \partial_x F(\hat{x}) = 0 \\ 2. \partial_{xx}^2 F(\hat{x}) > 0 \end{array}}$$

then $\hat{x}$ is a (local or global) minimum of $F(x)$.

Proof:

A Taylor expansion of $F(x)$ at $\hat{x}$ gives with $x=\hat{x}+td$ and (11.1):

$$(12) \quad F(x) = F(\hat{x}) + \left(\frac{1}{2} d^* \partial_{xx}^2 F(\hat{x}) d + \frac{O(t^2)}{t^2} \right) t^2.$$

Since $\partial_{xx}^2 F(\hat{x})$ is positive definite by assumption, the first term in the bracket on the right hand side of (12) is strictly positive. Hence, in view of (7), we can conclude that the bracketed term

of (12) is strictly positive for sufficiently small t . Thus from (12) follows that $F(x) > F(\hat{x})$ for all x near $\hat{x}$, that is, for t small.

It should be noted that $\hat{x}$ can be a minimum of $F(x)$ and still violate the sufficient conditions of the above theorem (for example $F(x) = x^4$). Nevertheless the above two theorems facilitate the problem of computing the minimum of $F(x)$. The idea is that one first solves the system of equations $\partial_x F(x) = 0$. The solutions of $\partial_x F(x) = 0$ are called *stationary or critical points* of $F(x)$. Once the stationary points of $F(x)$ are found, one can check whether $\partial_{xx}^2 F(\hat{x}) > 0$ holds in order to show that the corresponding stationary point must be a local minimum.

If $\partial_{xx}^2 F(\hat{x}) \geq 0$, one first computes the vectors d that satisfy $\partial_{xx}^2 F(\hat{x})d = 0$ and then checks whether $F(\hat{x}) \leq F(x)$ for all those x that lie on the ray $x = \hat{x} + td$ that emanates from $\hat{x}$ with direction vector d . After the local minima are found, the global minima follow from comparison of the function values at the local minima.

Example (least-squares)

The function $F: \mathbb{R}^n \rightarrow \mathbb{R}$ is given as:

$$(13) \quad \begin{aligned} F(x) &= (y - Ax)^*(y - Ax) \\ &= y^*y - 2y^*Ax + x^*A^*Ax . \end{aligned}$$

From this it follows that:

$$(14) \quad \partial_x F(x) = -2A^*y + 2A^*Ax ,$$

and:

$$(15) \quad \partial_{xx}^2 F(x) = 2A^*A .$$

If we assume that the $m \times n$ matrix A has full rank n , then A^*A has full rank and $(A^*A)^{-1}$ exists. The solution of $\partial_x F(x) = 0$ follows from (14) as:

$$(16) \quad \hat{x} = (A^*A)^{-1}A^*y .$$

Thus the function $F(x)$ has only one stationary point. Since $x^*A^*Ax > 0 \ \forall x \in \mathbb{R}^n \setminus \{0\}$ it follows that the Hessian matrix of (15) is positive-definite. Hence, $\hat{x}$ of (16) is the unique global minimum of $F(x)$.

Appendix C

Linearly constrained optimization: optimality conditions

In the following appendix we consider general minimization problems. However, there are special classes of constrained minima which are important and can be studied with the help of the theory of appendix B. These are the *linearly* constrained minima. In problems of this type we restrict our points $y \in \mathbb{R}^m$ to $\mathbb{R}^n$ in an n -plane in $\mathbb{R}^m$. An $\mathbb{R}^n$ -plane can be considered to be an n -dimensional Euclidean space embedded in $\mathbb{R}^m$. It need not pass through the origin $\mathbb{R}^m$. If y^0 is a point in the n -plane and $a_1, a_2, \dots, a_n$ are n linearly independent vectors in this plane, every point y in the n -plane is expressible in the form:

$$(1) \quad y = y^0 + x_1 a_1 + x_2 a_2 + \dots + x_n a_n$$

with suitably chosen constants $x_i \in \mathbb{R} \ i=1, \dots, n$. Equation (1) is sketched in figure C.1 for the case $m=3$ and $n=2$.

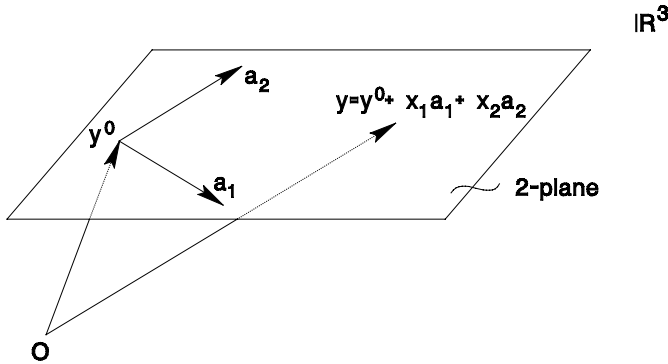


Figure C.1: A 2-plane embedded in $\mathbb{R}^3$

An n -plane is uniquely determined by a point y^0 and a set of n linearly independent vectors $a_1, a_2, \dots, a_n$. If $n=1$, we have a line; if $n=2$, we have a two-dimensional plane; and so on. To relate the ideas here presented to those given for more general constrained minima on arbitrary surfaces, it will be convenient to introduce the concept of the tangent space and the normal space to an n -plane. A tangent vector a is a vector parallel to the n -plane. If the n -plane is represented in the *parametric form* (1), a tangent vector a is any vector expressible as a linear combination:

$$(2) \quad a = x_1 a_1 + \dots + x_n a_n$$

of the vectors $a_1, \dots, a_n$. The *tangent space* of the n -plane is the set of all vectors that can be written as (2). It is the n -plane through the origin that is parallel to the set $\{a_1, \dots, a_n\}$. Any vector b that is orthogonal to the n -plane and hence to the tangent space will be called a normal vector to the n -plane. The set of all normal vectors to the n -plane is the orthogonal complement of the tangent space and will be called *normal space*. It is of dimension $m-n$. If $b_1, b_2, \dots, b_{m-n}$ form a basis for the normal space, every normal vector b is uniquely expressible as a linear combination:

$$(3) \quad \mathbf{b} = \lambda_1 \mathbf{b}_1 + \dots + \lambda_{m-n} \mathbf{b}_{m-n}.$$

We may use the normal vectors $\mathbf{b}_1, \dots, \mathbf{b}_{m-n}$ to express the n -plane as:

$$(4) \quad \mathbf{b}_i^* (\mathbf{y} - \mathbf{y}^0) = 0, \quad i = 1, \dots, m-n.$$

This representation of the n -plane is in contrast to the parametric form (1), known as the *implicit form*. Equation (4) is sketched in figure C.2 for the case $m=3$ and $n=2$.

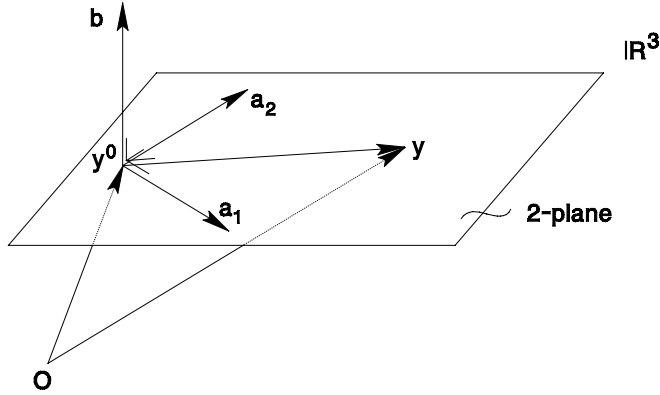


Figure C.2: A 2-plane embedded in $\mathbb{R}^3$ with normal vector $\mathbf{b}$

If we collect the vectors $\mathbf{a}_1, \dots, \mathbf{a}_n$ into an $m \times n$ matrix $A = (\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n)$ and define the n -vector $\mathbf{x}$ as $\mathbf{x} = (x_1, \dots, x_n)^*$, we may write (1) in matrix form as:

$$(5) \quad \underset{m \times 1}{\mathbf{y}} = \underset{m \times 1}{\mathbf{y}^0} + \underset{m \times n}{\mathbf{A}} \underset{n \times 1}{\mathbf{x}}$$

Similarly, if we collect the vectors $\mathbf{b}_1, \dots, \mathbf{b}_{m-n}$ into an $m \times (m-n)$ matrix $B = (\mathbf{b}_1, \dots, \mathbf{b}_{m-n})$, we may write (4) in matrix form as:

$$(6) \quad \underset{(m-n) \times m}{\mathbf{B}^*} \underset{m \times 1}{(\mathbf{y} - \mathbf{y}^0)} = \underset{(m-n) \times 1}{\mathbf{0}}$$

Observe that the two matrices A and B satisfy the relation:

$$(7) \quad \boxed{\underset{(m-n) \times m}{\mathbf{B}^*} \underset{m \times n}{\mathbf{A}} = \underset{(m-n) \times n}{\mathbf{0}}.}$$

We now turn to the characterization of the minimum $\hat{\mathbf{y}}$ of a function $F(\mathbf{y})$ on an n -plane. The following theorem gives necessary conditions for $\hat{\mathbf{y}}$ to be a solution to the linearly constrained minimization problem:

$$(8) \quad \left\{ \begin{array}{l} \min_y F(y), \quad y \in \mathbb{R}^m, \quad F : \mathbb{R}^m \rightarrow \mathbb{R} \\ \text{with } y \text{ constrained to an } n\text{-plane} \end{array} \right.$$

Theorem (necessary conditions):

Let $F(y)$ be of class $C^{(3)}$. If $\hat{y}$ is a point on the n -plane that affords a (local or global) minimum of $F(y)$ on the n -plane, then:

$$(9) \quad \boxed{\begin{array}{l} 1. \quad \partial_y F(\hat{y})^* A \quad = \quad 0 \\ 2. \quad A^* \partial_{yy}^2 F(\hat{y}) A \quad \geq \quad 0 \end{array}}$$

where A is an $m \times n$ matrix of which the column vectors form a basis of the tangent space of the n -plane.

Proof:

The vector $\hat{x}=0$ minimizes:

$$f(x) = F(\hat{y} + Ax) \quad .$$

Hence, it follows from the first theorem of appendix B that:

$$\partial_x f(0)^* = \partial_y F(\hat{y})^* A = 0 \quad ,$$

and:

$$\partial_{xx}^2 f(0) = A^* \partial_{yy}^2 F(\hat{y}) A \geq 0 \quad .$$

The following theorem gives sufficient conditions for $\hat{y}$ to be a solution to (8).

Theorem (sufficient conditions):

Let $F(y)$ be of class $C^{(3)}$. If:

$$(10) \quad \boxed{\begin{array}{l} 1. \quad \partial_y F(\hat{y})^* A \quad = \quad 0 \\ 2. \quad A^* \partial_{yy}^2 F(\hat{y}) A \quad > \quad 0 \end{array}}$$

where A is an $m \times n$ matrix of which the column vectors form a basis of the tangent space of the n -plane, then $\hat{y}$ is a (local or global) minimum of $F(y)$ on this n -plane.

Proof:

The proof follows with $f(x)=F(\hat{y}+Ax)$, $\partial_x f(x)^* = \partial_y F(\hat{y}+Ax)^* A$ and $\partial_{xx}^2 f(x) = A^* \partial_{yy}^2 F(\hat{y}+Ax) A$ from the second theorem of appendix B.

The above two theorems can be used if the n -plane is given in the parametric form (5). In case the n -plane is defined by the implicit form (6), the results of the above two theorems can be restated as in the following theorem.

Theorem (Lagrange multiplier rule):

Let $\hat{y}$ be a point satisfying the linear constraints:

$$(11) \quad \underset{(m-n) \times m}{\mathbf{B}^*} \underset{m \times 1}{(y - y^0)} = \underset{(m-n) \times 1}{\mathbf{0}}$$

If $\hat{y}$ minimizes $F(y)$ subject to these constraints, there exists a multiplier vector $\lambda \in \mathbb{R}^{m-n}$, such that, if we set:

$$(12) \quad L(y) = F(y) + \lambda^* \mathbf{B}^* (y - y^0),$$

then:

$$(13) \quad \boxed{\begin{array}{ll} 1. \partial_y L(\hat{y}) & = \mathbf{0} \\ 2. \mathbf{a}^* \partial_{yy}^2 L(\hat{y}) \mathbf{a} & \geq 0 \end{array}}$$

for all vectors $\mathbf{a}$ that satisfy:

$$(14) \quad \underset{(m-n) \times m}{\mathbf{B}^*} \underset{m \times 1}{\mathbf{a}} = \underset{(m-n) \times 1}{\mathbf{0}}.$$

Conversely, if there exists for $\hat{y}$ a function $L(y)$ of form (12) such that:

$$(15) \quad \boxed{\begin{array}{ll} 1. \partial_y L(\hat{y}) & = \mathbf{0} \\ 2. \mathbf{a}^* \partial_{yy}^2 L(\hat{y}) \mathbf{a} & > 0 \end{array}}$$

for solutions $\mathbf{a} \neq \mathbf{0}$ of (14), then $\hat{y}$ affords a (local or global) minimum to $F(y)$ subject to the constraints (11).

The result here given is known as the *Lagrange multiplier rule*. The function $L(y)$ of (12) is known as the *Lagrangian* and the vector λ as the *Lagrange multiplier vector*.

Proof:

We first prove (13.1). Condition (9.1) states that the gradient $\partial_y F(\hat{y})$ is orthogonal to the range space of $\mathbf{A}$. This implies that $\partial_y F(\hat{y})$ can be written as a linear combination of the base vectors of the normal space of the n -plane. Hence, there exists a vector $\lambda \in \mathbb{R}^{m-n}$ such that:

$$\partial_y F(\hat{y}) + \mathbf{B} \lambda = \mathbf{0}$$

But this implies that the gradient of $L(y)$ must vanish at $\hat{y}$. Condition (13.2) follows simply from (9.2) by noting that $\partial_{yy}^2 L(\hat{y}) = \partial_{yy}^2 F(\hat{y})$.

The sufficient conditions (15) follow in a similar way from (10).

Example 1 (observation equations)

Let y^0 be a fixed given vector in $\mathbb{R}^m$, and let $F: \mathbb{R}^m \rightarrow \mathbb{R}$ be defined as:

$$(16) \quad F(y) = (y^0 - y)^*(y^0 - y) = \|y^0 - y\|^2.$$

Note that the function $F(y)$ measures the square of the distance between y and the fixed vector y^0 . Hence, the $(m-1)$ -dimensional surface $F(y) = \text{constant}$ is an hypersphere embedded in $\mathbb{R}^m$ (see figure C.3).

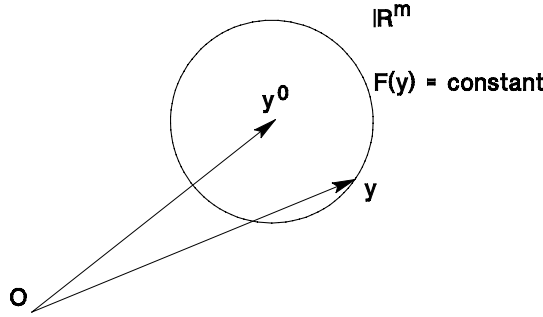


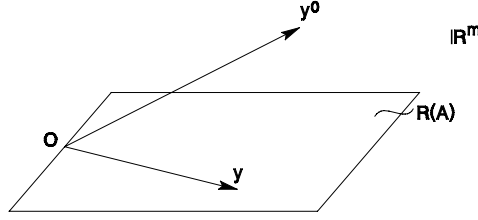
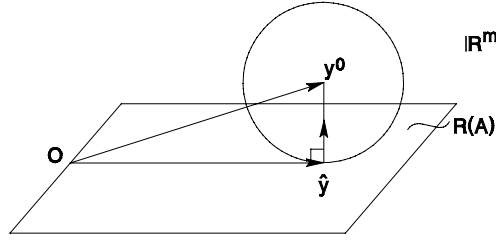
Figure C.3: $F(y) = \|y^0 - y\|^2 = \text{constant}$

Our objective is to minimize the function $F(y)$ of (16) subject to the linear restrictions:

$$(17) \quad y = A x, \quad \text{rank } A = n$$

$m \times 1 \quad m \times n \quad n \times 1$

Those vectors that satisfy (17) lie on an n -plane through the origin with tangent vectors given by the column vectors of the matrix A (see figure C.4). The problem of minimizing $F(y)$ subject to (17) can now be visualized as the problem of finding the proper radius of the hypersphere of figure C.3 such that the sphere just touches the n -plane of figure C.4. This is shown in figure C.5. The point of contact between sphere and n -plane is the solution $\hat{y}$.


 Figure C.4: The n -plane: $y=Ax$, rank $A=n$

 Figure C.5: Hypersphere $\|y^0 - y\|^2 = \text{constant}$, and n -plane $y=Ax$

The solution $\hat{y}$ can be found using (10). Since $\partial_y F(y) = -2(y^0 - y)$, it follows that:

$$\partial_y F(y)^* A = -2(y^0 - y)^* A.$$

Substitution of (17) and setting the results equal to zero gives:

$$A^* A x = A^* y^0.$$

This system of equations has the unique solution $\hat{x} = (A^* A)^{-1} A^* y^0$. Hence:

$$(18) \quad \hat{y} = A \hat{x} = A(A^* A)^{-1} A^* y^0.$$

Since $\partial_{yy}^2 F(y) = 2I_m$ and $A^* A$ is positive definite it follows that:

$$A^* \partial_{yy}^2 F(\hat{y}) A = 2A^* A > 0.$$

Hence, $\hat{y}$ is a minimum of $F(y)$ subject to (17).

Example 2 (condition equations)

Let y^0 be a fixed given vector in $\mathbb{R}^m$, and let $F: \mathbb{R}^m \rightarrow \mathbb{R}$ be defined as:

$$(19) \quad F(y) = (y^0 - y)^* (y^0 - y) = \|y^0 - y\|^2.$$

Our objective is to minimize the function $F(y)$ of (19) subject to the linear restrictions.

$$(20) \quad \underset{(m-n) \times m}{B^*} \underset{m \times 1}{y} = \underset{(m-n) \times 1}{0}, \text{ rank } B = m - n.$$

Since these restrictions are in *implicit form*, we have to make use of the Lagrange multiplier rule. We construct the Lagrangian:

$$(21) \quad L(y) = (y^0 - y)^*(y^0 - y) + 2\lambda^* B^* y.$$

The factor 2 in the multiplier vector has merely been introduced for convenience. The stationary points of the Lagrangian have to satisfy:

$$(22) \quad \partial_y L(y) = -2(y^0 - y) + 2B\lambda = 0.$$

This shows that:

$$(23) \quad \hat{y} = y^0 - B\lambda.$$

In order to determine $\hat{y}$, we first have to determine λ . Pre-multiplication of (22) with B^* gives together with (20):

$$(24) \quad B^* B \lambda = B^* y^0.$$

Since matrix B has full rank, the inverse of $B^* B$ exists. Hence λ follows from (24) as:

$$(25) \quad \lambda = (B^* B)^{-1} B^* y^0.$$

Substitution of this result into (23) gives:

$$(26) \quad \hat{y} = [I - B(B^* B)^{-1} B^*] y^0.$$

Thus the Lagrangian (21) has only one stationary point. Since $\partial_{yy}^2 L(y) = 2I_m$, it follows that condition (15.2) is satisfied. Hence, $\hat{y}$ of (26) is the solution of our linearly constrained minimization problem.

Example 3 (Best Linear Unbiased Estimation)

In Chapter 2, section 4 the principle of "*Best Linear Unbiased Estimation*" has been discussed. This principle leads to the following linearly constrained minimization problem:

$$(27) \quad \min_l l^* Q_y l \text{ subject to } A^* l = f.$$

This problem may be solved by the Lagrange multiplier rule. We consider the following Lagrangian:

$$(28) \quad L(l) = l^* Q_y l + 2\lambda^* (A^* l - f).$$

The stationary points of the Lagrangian have to satisfy:

$$(29) \quad \partial_l L(l) = 2Q_y l + 2A\lambda = 0.$$

This shows that:

$$(30) \quad \hat{l} = -Q_y^{-1} A \lambda .$$

In order to determine $\hat{l}$ we first have to determine λ . Pre-multiplication of (29) with $A^* Q_y^{-1}$ gives together with the constraints $A^* l = f$:

$$(31) \quad A^* Q_y^{-1} A \lambda = -f$$

Since A is of full rank and Q_y is positive definite, also $A^* Q_y^{-1} A$ is of full rank and positive definite. Hence, the inverse of $A^* Q_y^{-1} A$ exists and λ follows from (31) as:

$$(32) \quad \lambda = -(A^* Q_y^{-1} A)^{-1} f .$$

Substitution into (30) gives:

$$(33) \quad \hat{l} = Q_y^{-1} A (A^* Q_y^{-1} A)^{-1} f .$$

Since $\partial_{ll}^2 L(l) = 2Q_y$ and Q_y is positive definite, it follows that $\hat{l}$ of (33) is the unique solution of (27).

Appendix D

Constrained optimization: optimality conditions

In the previous appendix we were concerned with the problem of minimizing a function $F(y)$ with linear equality constraints. In this appendix we seek to minimize a function $F(y)$ with *nonlinear* equality constraints. Hence, we want to solve the constrained minimization problem:

$$(1) \quad \begin{cases} \min_y F(y) , F: \mathbb{R}^m \rightarrow \mathbb{R} \\ \text{subject to } B(y)=0, B: \mathbb{R}^m \rightarrow \mathbb{R}^{m-n} \end{cases}$$

The constraints in (1) are in *implicit form* and they are nonlinear if the vector function $B: \mathbb{R}^m \rightarrow \mathbb{R}^{m-n}$ is nonlinear. Except in abnormal situations, the constraints of (1) describe an n -dimensional surface (or manifold) embedded in $\mathbb{R}^m$. Examples of $B(y) = 0$ are:

$$\begin{aligned} 1. \quad & x^2 + y^2 + z^2 - R^2 = 0 , \\ 2. \quad & \begin{cases} x^2 + y^2 + z^2 - R^2 = 0 \\ y - x = 0 . \end{cases} \end{aligned}$$

The first equation describes the surface of a sphere with radius R embedded in $\mathbb{R}^3$. The last two equations together describe a circle embedded in $\mathbb{R}^3$.

By virtue of the implicit function theorem the constraints:

$$(2) \quad B(y) = 0 , B: \mathbb{R}^m \rightarrow \mathbb{R}^{m-n}$$

are expressible locally in *parametric form* as:

$$(3) \quad y = A(x) , A: \mathbb{R}^n \rightarrow \mathbb{R}^m .$$

For instance, example 1 is expressible in parametric form as:

$$1. \quad \begin{cases} x = R \cos \varphi \cos \lambda \\ y = R \cos \varphi \sin \lambda \\ z = R \sin \varphi \end{cases} .$$

And the equations of example 2 as:

$$2. \quad \begin{cases} x = \frac{1}{2}\sqrt{2} R \cos \varphi \\ y = \frac{1}{2}\sqrt{2} R \cos \varphi \\ z = R \sin \varphi \end{cases} .$$

Since (3) puts the same restrictions on y as (2) we may write the minimization problem (1) also as:

$$(4) \quad \begin{cases} \min_y F(y) , & F: \mathbb{R}^m \rightarrow \mathbb{R} \\ \text{subject to } y = A(x) , & A: \mathbb{R}^n \rightarrow \mathbb{R}^m . \end{cases}$$

Consequently, the problem of minimizing $F(y)$ subject to $B(y) = 0$ is equivalent to the unconstrained problem of minimizing $F(A(x))$. In some cases when it is simple to transform (2) into (3) this is a suitable procedure. However, in many cases it is easier to use the *Lagrange multiplier rule*. The following two theorems, which will be given without proof, give necessary and sufficient conditions for $\hat{y}$ to be a (local or global) solution to (1).

Theorem (necessary conditions):

Suppose that $\hat{y}$ affords a (local or global) minimum to:

$$F(y) , F: \mathbb{R}^m \rightarrow \mathbb{R} ,$$

subject to the constraints:

$$B(y) = 0 , B: \mathbb{R}^m \rightarrow \mathbb{R}^{m-n} .$$

Suppose further that the $m \times (m-n)$ matrix of partial derivatives $\partial_y B(\hat{y})$ has full rank $m-n$. Then there exists a unique multiplier vector $\lambda \in \mathbb{R}^{m-n}$ such that, if we set:

$$L(y) = F(y) + \lambda^* B(y) ,$$

then:

$$(5) \quad \begin{array}{l} 1. \quad \partial_y L(\hat{y}) = \partial_y F(\hat{y}) + \partial_y B(\hat{y}) \lambda = 0 \\ 2. \quad a^* \partial_{yy}^2 L(\hat{y}) a \geq 0 \end{array}$$

for all vectors a that satisfy:

$$\begin{array}{ccc} \partial_y B(\hat{y})^* a & = & 0 \\ (m-n) \times m & m \times 1 & (m-n) \times 1 \end{array} .$$

Theorem (sufficient conditions):

Suppose that for a point $\hat{y}$ satisfying $B(y) = 0$ there is a Lagrangian $L(y) = F(y) + \lambda^* B(y)$ such that:

$$(6) \quad \begin{array}{l} 1. \quad \partial_y L(\hat{y}) = 0 \\ 2. \quad a^* \partial_{yy}^2 L(\hat{y}) a > 0 \end{array}$$

for all vectors $a \neq 0$ satisfying $\partial_y B(\hat{y})^* a = 0$, then $\hat{y}$ is a (local or global) minimum of $F(y)$ subject to $B(y) = 0$.

Appendix E

Mean and variance of scalar random variables

Mean of a scalar random variable

Let $\underline{x}$ be a continuous scalar random variable with probability density function $p_{\underline{x}}(x)$. The *expectation* or *mean* of $\underline{x}$ is by definition the integral:

$$(1) \quad E\{\underline{x}\} = \int_{-\infty}^{+\infty} x p_{\underline{x}}(x) dx .$$

This number will also be denoted by m_x .

Mean of a function of a scalar random variable

Given a scalar random variable $\underline{x}$ and a function $f(x)$, we form the random variable $\underline{y} = f(\underline{x})$. As we see from (1), the mean of $\underline{y}$ is given by:

$$(2) \quad E\{\underline{y}\} = \int_{-\infty}^{+\infty} y p_{\underline{y}}(y) dy .$$

It appears, therefore, that to determine the mean of $\underline{y}$, we must first find its probability density function $p_{\underline{y}}(y)$. This, however, is not necessary. As the next theorem shows, $E\{\underline{y}\}$ can be expressed directly in terms of the function $f(x)$ and the density $p_{\underline{x}}(x)$ of $\underline{x}$.

Theorem:

$$(3) \quad E\{f(\underline{x})\} = \int_{-\infty}^{+\infty} f(x) p_{\underline{x}}(x) dx .$$

Proof:

We shall sketch a proof using the curve $f(x)$ of figure E.1. With $y = f(x_1) = f(x_2) = f(x_3)$ as in the figure, we see that:

$$P(y \leq \underline{y} \leq y + dy) = P(x_1 \leq \underline{x} \leq x_1 + dx_1) + P(x_2 - dx_2 \leq \underline{x} \leq x_2) + P(x_3 \leq \underline{x} \leq x_3 + dx_3)$$

or:

$$p_{\underline{y}}(y) dy = p_{\underline{x}}(x_1) dx_1 + p_{\underline{x}}(x_2) dx_2 + p_{\underline{x}}(x_3) dx_3 .$$

Multiplying by y , we obtain:

$$y p_y(y) dy = f(x_1) p_x(x_1) dx_1 + f(x_2) p_x(x_2) dx_2 + f(x_3) p_x(x_3) dx_3.$$

Thus, to each differential in (2) there corresponds one or more differentials in (3). As dy covers the y -axis, the corresponding dx 's are nonoverlapping and they cover the entire x -axis. Hence, the integrals in (2) and (3) are equal.

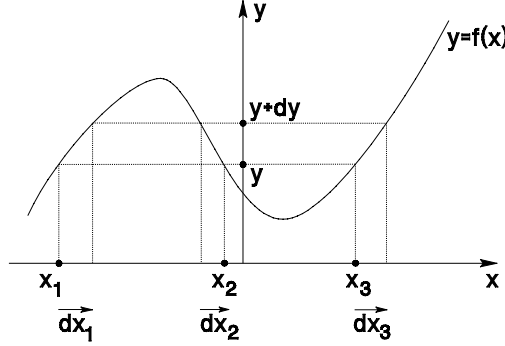


Figure E.1: The curve $y = f(x)$

It appears from (3) that to determine the mean of $\underline{y}$, we must know the probability density function $p_x(x)$. In general this is true. However, in the special case that the function $f(x)$ is *linear* not the complete density of $\underline{x}$ needs to be known but only the mean m_x of $\underline{x}$.

Theorem (propagation law of the mean):

Given a scalar random variable $\underline{x}$ and a function $f(x)$, we form the random variable $\underline{y} = f(\underline{x})$. If the function $f(x)$ is *linear*:

$$(4) \quad f(x) = ax + b,$$

then:

$$(5) \quad \boxed{m_y = am_x + b}$$

Proof:

Substitution of (4) into (3) gives:

$$\begin{aligned} m_y &= \int_{-\infty}^{+\infty} (ax + b) p_x(x) dx \\ &= a \int_{-\infty}^{+\infty} x p_x(x) dx + b \int_{-\infty}^{+\infty} 1 p_x(x) dx \\ &= am_x + b. \end{aligned}$$

If the function $f(x)$ is nonlinear we need strictly speaking the density $p_x(x)$ of $\underline{x}$ in order to compute the mean of $\underline{y} = f(\underline{x})$. However, by using Taylor's formula an approximation to the

mean of $\underline{y}$ can be derived that gets around the difficulty of having to know the density $p_x(x)$ of $\underline{x}$.

Theorem (linearized propagation law of the mean):

Given a scalar random variable $\underline{x}$ and a nonlinear function $f(x)$, we form the random variable $\underline{y} = f(\underline{x})$. Let x^0 be an approximation to a sample of $\underline{x}$ and define $\Delta \underline{y} = \underline{y} - f(x^0)$ and $\Delta \underline{x} = \underline{x} - x^0$. Then a first-order approximation to the mean of $\Delta \underline{y}$ is:

$$(6) \quad m_{\Delta y} \doteq \frac{d}{dx} f(x^0) m_{\Delta x} .$$

Proof:

Application of Taylor's formula gives:

$$\underline{y} = f(x^0) + \frac{d}{dx} f(x^0) (\underline{x} - x^0) + \frac{1}{2} \frac{d^2}{dx^2} f(x^0) (\underline{x} - x^0)^2 + \dots$$

If we take the expectation we get:

$$E\{\underline{y}\} = f(x^0) + \frac{d}{dx} f(x^0) E\{\underline{x} - x^0\} + \frac{1}{2} \frac{d^2}{dx^2} f(x^0) E\{(\underline{x} - x^0)^2\} + \dots$$

This may be written with:

$$m_{\Delta y} = E\{\underline{y}\} - f(x^0) , \quad m_{\Delta x} = E\{\underline{x} - x^0\} ,$$

and:

$$\begin{aligned} E\{(\underline{x} - x^0)^2\} &= E\{[(\underline{x} - m_x) - (x^0 - m_x)]^2\} \\ &= E\{(\underline{x} - m_x)^2\} - 2E\{(\underline{x} - m_x)(x^0 - m_x)\} + E\{(x^0 - m_x)^2\} \\ &= E\{(\underline{x} - m_x)^2\} + m_{\Delta x}^2 \\ &= \sigma_x^2 + m_{\Delta x}^2 , \end{aligned}$$

as:

$$m_{\Delta y} = \frac{d}{dx} f(x^0) m_{\Delta x} + \frac{1}{2} \frac{d^2}{dx^2} f(x^0) (\sigma_x^2 + m_{\Delta x}^2) + \dots$$

If we neglect the second and higher order terms, the result (6) follows.

Examples

1. With $f(x) = \cos x$, we define the random variable $\underline{y} = f(\underline{x})$. Since $\frac{d}{dx}f(x^0) = -\sin x^0$, we find to a first order:

$$E\{\underline{y} - \cos x^0\} \doteq -\sin x^0 E\{\underline{x} - x^0\} .$$

2. With $f(x) = x^3$, we define the random variable $\underline{y} = f(\underline{x})$. Since $\frac{d}{dx}f(x^0) = 3(x^0)^2$ we find to a first order:

$$E\{\underline{y} - (x^0)^3\} \doteq 3(x^0)^2 E\{\underline{x} - x^0\} .$$

Variance of a scalar random variable

The *variance* of a scalar random variable is by definition the integral:

$$(7) \quad E\{(\underline{x} - E\{\underline{x}\})^2\} = \int_{-\infty}^{+\infty} (x - m_x)^2 p_x(x) dx .$$

This number will also be denoted by σ_x^2 . The positive constant σ_x is called the *standard deviation* of $\underline{x}$. From the definition it follows that:

$$\sigma_x^2 = E\{(\underline{x} - m_x)^2\} = E\{\underline{x}^2\} - 2E\{\underline{x}\}m_x + m_x^2 ,$$

or

$$(8) \quad \sigma_x^2 = E\{\underline{x}^2\} - m_x^2 .$$

Variance of a function of a scalar random variable

Theorem (propagation law of the variance):

Given a scalar random variable $\underline{x}$ and a function $f(x)$, we form the random variable $\underline{y} = f(\underline{x})$. If the function $f(x)$ is *linear*:

$$(9) \quad f(x) = ax + b ,$$

then:

$$(10) \quad \sigma_y^2 = a^2 \sigma_x^2$$

Proof:

According to definition (7) we have:

$$\sigma_y^2 = E\{(y - E\{y\})^2\} = E\{(f(\underline{x}) - m_y)^2\}.$$

With (3) this gives:

$$\sigma_y^2 = \int_{-\infty}^{+\infty} (f(x) - m_y)^2 p_{\underline{x}}(x) dx.$$

Substitution of (5) and (9) gives:

$$\begin{aligned}\sigma_y^2 &= \int_{-\infty}^{+\infty} [(ax + b) - (am_x + b)]^2 p_{\underline{x}}(x) dx \\ &= a^2 \int_{-\infty}^{+\infty} (x - m_x)^2 p_{\underline{x}}(x) dx \\ &= a^2 \sigma_x^2.\end{aligned}$$

The above result shows that if $f(x)$ is linear, knowledge of σ_x^2 is sufficient for computing the variance of $y = f(\underline{x})$. For nonlinear functions $f(x)$ this is generally not true. If the function $f(x)$ is nonlinear one will generally need to know the complete density $p_{\underline{x}}(x)$ of $\underline{x}$. However, by using Taylor's formula an approximation to the variance of y can be derived that gets around the difficulty of having to know $p_{\underline{x}}(x)$.

Theorem (linearized propagation law of the variance):

Given a scalar random variable $\underline{x}$ and a nonlinear function $f(x)$, we form the random variable $y = f(\underline{x})$. Let x^0 be an approximation to a sample of $\underline{x}$. Then a first-order approximation to the variance of y is:

$$(11) \quad \sigma_y^2 \doteq \left(\frac{d}{dx} f(x^0) \right)^2 \sigma_x^2$$

Proof:

Substitution of:

$$\begin{cases} f(x) &= f(x^0) + \frac{d}{dx} f(x^0)(x - x^0) + \dots \\ m_y &= f(x^0) + \frac{d}{dx} f(x^0)(m_x - x^0) + \dots \end{cases}$$

into:

$$\sigma_y^2 = \int_{-\infty}^{+\infty} (f(x) - m_y)^2 p_{\underline{x}}(x) dx,$$

gives after neglecting second and higher order terms:

$$\begin{aligned}
 \sigma_y^2 &= \int_{-\infty}^{+\infty} \left(\frac{d}{dx} f(x^0)(x-m_x) \right)^2 p_x(x) dx \\
 &= \left(\frac{d}{dx} f(x^0) \right)^2 \int_{-\infty}^{+\infty} (x-m_x)^2 p_x(x) dx \\
 &= \left(\frac{d}{dx} f(x^0) \right)^2 \sigma_x^2
 \end{aligned}$$

Examples

1. A first order approximation to the variance of $y = \cos x$ is:

$$\sigma_y^2 = \sin^2 x^0 \sigma_x^2.$$

2. A first order approximation to the variance of $y = x^3$ is:

$$\sigma_y^2 = 9(x^0)^4 \sigma_x^2.$$

Appendix F

Mean and variance of vector random variables

Mean of a vector random variable

Let $\underline{x}_i, i = 1, 2, \dots, n$ be n continuous scalar random variables with *joint* probability density function $p_{\underline{x}}(x_1, x_2, \dots, x_n)$. The expectation or mean of the random n -vector $(\underline{x}_1, \underline{x}_2, \dots, \underline{x}_n)^*$ is by definition the integral:

$$(1) \quad E \left\{ \begin{pmatrix} \underline{x}_1 \\ \underline{x}_2 \\ \vdots \\ \underline{x}_n \end{pmatrix} \right\} = \int_{-\infty}^{+\infty} \dots \int_{-\infty}^{+\infty} \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} p_{\underline{x}}(x_1, x_2, \dots, x_n) dx_1 \dots dx_n$$

If we use the notation:

$$\underline{x} = (\underline{x}_1, \underline{x}_2, \dots, \underline{x}_n)^*, \quad x = (x_1, \dots, x_n)^*, \quad p_{\underline{x}}(x_1, \dots, x_n) = p_{\underline{x}}(x), \quad dx_1, \dots, dx_n = dx \quad \text{and} \quad \int_{-\infty}^{+\infty} \dots \int_{-\infty}^{+\infty} = \int$$

we may write (1) in the compact form:

$$(2) \quad E\{\underline{x}\} = \int x p_{\underline{x}}(x) dx.$$

From (1) it follows that:

$$(3) \quad E\{\underline{x}_i\} = \int_{-\infty}^{+\infty} \dots \int_{-\infty}^{+\infty} x_i p_{\underline{x}}(x_1, \dots, x_i, \dots, x_n) dx_1 \dots dx_n.$$

This result seems to contradict definition (1) of appendix E. Note however that (3) reduces to (1) of appendix E, since:

$$\int_{-\infty}^{+\infty} p_{\underline{x}}(x_1, \dots, x_j, \dots, x_n) dx_j = p_{\underline{x}}(x_1, \dots, x_{j-1}, x_{j+1}, \dots, x_n).$$

Hence, (1) may also be written as:

$$(4) \quad E\left\{\begin{pmatrix} \underline{x}_1 \\ \underline{x}_2 \\ \vdots \\ \underline{x}_n \end{pmatrix}\right\} = \begin{pmatrix} \int_{-\infty}^{+\infty} x_1 p_{\underline{x}_1}(x_1) dx_1 \\ \int_{-\infty}^{+\infty} x_2 p_{\underline{x}_2}(x_2) dx_2 \\ \vdots \\ \int_{-\infty}^{+\infty} x_n p_{\underline{x}_n}(x_n) dx_n \end{pmatrix}.$$

Thus in order to compute $E\{\underline{x}_i\}$ one only needs the *marginal* density $p_{\underline{x}_i}(x_i)$.

Mean of a vector function of a random vector

Given a random n -vector $\underline{x}$ and a vector function $F(x)$, $F: \mathbb{R}^n \rightarrow \mathbb{R}^m$, we form the random m -vector $\underline{y} = F(\underline{x})$. As we see from (4) the mean of $\underline{y}_i$ is given by:

$$(5) \quad E\{\underline{y}_i\} = \int_{-\infty}^{+\infty} y_i p_{\underline{y}_i}(y_i) dy_i.$$

It appears, therefore, that to determine the mean of $\underline{y}_i$, we must first find its *marginal* probability density function $p_{\underline{y}_i}(y_i)$. This, however, is not necessary. As the next theorem shows, $E\{\underline{y}_i\}$ can be expressed directly in terms of the vector function $F(x)$ and the *joint* density of $\underline{x}$.

Theorem:

$$(6) \quad E\{F(\underline{x})\} = \int_{-\infty}^{+\infty} \dots \int_{-\infty}^{+\infty} F(x_1, \dots, x_n) p_{\underline{x}}(x_1, \dots, x_n) dx_1 \dots dx_n$$

or:

$$(7) \quad E\{F(\underline{x})\} = \int F(x) p_{\underline{x}}(x) dx$$

Proof:

The proof is similar to the proof given in appendix E.

The following theorem is an extremely important one, and it is used frequently.

Theorem (propagation law of the mean):

Given a random n -vector $\underline{x}$ and a vector function $F(x)$, $F: \mathbb{R}^n \rightarrow \mathbb{R}^m$, we form the random m -vector $\underline{y} = F(\underline{x})$. If the vector function $F(x)$ is linear:

$$(8) \quad F(x) = \underset{m \times 1}{A} \underset{m \times n}{x} + \underset{m \times 1}{b}$$

then:

$$(9) \quad \boxed{\begin{matrix} m_y & = & A & m_x & + & b \\ m \times 1 & & m \times n & n \times 1 & & m \times 1 \end{matrix}} .$$

Proof:

If we denote the n column vectors of matrix A by a_i , $i=1, \dots, n$, we may write (8) as:

$$F(x) = \sum_{i=1}^n a_i x_i + b.$$

If we substitute this into (6) we get:

$$E\{F(\underline{x})\} = \sum_{i=1}^n a_i \int_{-\infty}^{+\infty} \dots \int_{-\infty}^{+\infty} x_i p_{\underline{x}}(x_1, \dots, x_n) dx_1 \dots dx_n + b \int_{-\infty}^{+\infty} \dots \int_{-\infty}^{+\infty} p_{\underline{x}}(x_1, \dots, x_n) dx_1 \dots dx_n$$

or:

$$\begin{aligned} E\{F(\underline{x})\} &= \sum_{i=1}^n a_i \int_{-\infty}^{+\infty} x_i p_{\underline{x}}(x_i) dx_i + b \\ &= \sum_{i=1}^n a_i m_{x_i} + b \\ &= A m_x + b. \end{aligned}$$

Without proof we also give the linearized version of the propagation law of the mean.

Theorem (linearized propagation law of the mean):

Given a random n -vector $\underline{x}$ and a nonlinear vector function $F(x)$, $F: \mathbb{R}^n \rightarrow \mathbb{R}^m$ we form the random m -vector $\underline{y} = F(\underline{x})$. Let $x^0 \in \mathbb{R}^n$ be an approximation to a sample of $\underline{x}$ and define $\Delta \underline{y} = \underline{y} - F(x^0)$ and $\Delta \underline{x} = \underline{x} - x^0$. Then we have to a first-order:

$$(10) \quad \boxed{m_{\Delta y} \doteq \partial_x F(x^0) m_{\Delta x} .}$$

Example 1

Let the two random variables y_1 and y_2 be defined as:

$$\begin{cases} y_1 = 1x_1 + 3x_2 + 5x_3 + 2 \\ y_2 = 4x_1 + 2x_2 - 1x_3 + 3 \end{cases}$$

Then:

$$E \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \begin{pmatrix} 1 & 3 & 5 \\ 4 & 2 & -1 \end{pmatrix} E \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} + \begin{pmatrix} 2 \\ 3 \end{pmatrix}.$$

Example 2

Let the two random variables y_1 and y_2 be defined as:

$$\begin{cases} y_1 = \sin(x_1 x_2) + x_3 \\ y_2 = x_1^2 + x_2 + 4. \end{cases}$$

With the approximate values x_1^0 , x_2^0 and x_3^0 we have to a first-order:

$$E \begin{pmatrix} y_1 - \sin(x_1^0 x_2^0) - x_3^0 \\ y_2 - (x_1^0)^2 - x_2^0 - 4 \end{pmatrix} = \begin{pmatrix} x_2^0 \cos(x_1^0 x_2^0) & x_1^0 \cos(x_1^0 x_2^0) & 1 \\ 2x_1^0 & 1 & 0 \end{pmatrix} E \begin{pmatrix} x_1 - x_1^0 \\ x_2 - x_2^0 \\ x_3 - x_3^0 \end{pmatrix}.$$

Variance matrix of a random vector

Let $x_i, i=1,2,\dots,n$ be n continuous scalar random variables with joint probability density function $p_x(x_1, \dots, x_n)$. The *variance matrix* of the random n -vector $(x_1, x_2, \dots, x_n)^*$ is by definition the integral:

$$(11) \quad E \left\{ \begin{pmatrix} x_1 - E\{x_1\} \\ \vdots \\ x_n - E\{x_n\} \end{pmatrix} \cdot \begin{pmatrix} x_1 - E\{x_1\} \\ \vdots \\ x_n - E\{x_n\} \end{pmatrix}^* \right\} = \int_{-\infty}^{+\infty} \dots \int_{-\infty}^{+\infty} \begin{pmatrix} x_1 - E\{x_1\} \\ \vdots \\ x_n - E\{x_n\} \end{pmatrix} \cdot \begin{pmatrix} x_1 - E\{x_1\} \\ \vdots \\ x_n - E\{x_n\} \end{pmatrix}^* p_x(x_1, \dots, x_n) dx_1 \dots dx_n.$$

This variance matrix will be denoted by Q_x . Using vector notation we may write (11) in the compact form:

$$(12) \quad E\{(x - m_x)(x - m_x)^*\} = \int (x - m_x)(x - m_x)^* p_x(x) dx.$$

From (11) it follows that:

$$E\{(\underline{x}_i - m_{x_i})(\underline{x}_j - m_{x_j})\} = \int_{-\infty}^{+\infty} \dots \int_{-\infty}^{+\infty} (\underline{x}_i - m_{x_i})(\underline{x}_j - m_{x_j}) p_{\underline{x}}(x_1, \dots, x_n) dx_1 \dots dx_n .$$

Integration gives:

$$(13) \quad E\{(\underline{x}_i - m_{x_i})(\underline{x}_j - m_{x_j})\} = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} (\underline{x}_i - m_{x_i})(\underline{x}_j - m_{x_j}) p_{\underline{x}}(x_i, x_j) dx_i dx_j ,$$

in which $p_{\underline{x}}(x_i, x_j)$ is the joint density function of the two random variables $\underline{x}_i$ and $\underline{x}_j$. The scalar (13) is called the *covariance* of the two random variables $\underline{x}_i$ and $\underline{x}_j$. This covariance will be denoted as $\sigma_{x_i x_j}$. Note that the off-diagonal elements of the variance matrix Q_x consist of the covariances between the elements of the random vector $\underline{x}$. If $i=j$ it follows from (13) that:

$$E\{(\underline{x}_i - m_{x_i})^2\} = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} (\underline{x}_i - m_{x_i})^2 p_{\underline{x}}(x_i, x_j) dx_i dx_j .$$

Integration gives:

$$E\{(\underline{x}_i - m_{x_i})^2\} = \int_{-\infty}^{+\infty} (\underline{x}_i - m_{x_i})^2 p_{\underline{x}_i}(x_i) dx_i .$$

This is the *variance* $\sigma_{x_i}^2$ of $\underline{x}_i$. Note that the variance of $\sigma_{x_i}^2$ is the i th-diagonal element of the variance matrix Q_x . Hence, the variance matrix Q_x can be written as:

$$(14) \quad Q_x = \begin{matrix} & \begin{matrix} \sigma_{x_1}^2 & \sigma_{x_1 x_2} & \dots & \sigma_{x_1 x_n} \end{matrix} \\ \begin{matrix} \sigma_{x_2 x_1} & \sigma_{x_2}^2 & & \vdots \end{matrix} \\ \begin{matrix} \vdots & & \ddots & \vdots \end{matrix} \\ \begin{matrix} \sigma_{x_n x_1} & \dots & \dots & \sigma_{x_n}^2 \end{matrix} \end{matrix} \quad .$$

Note that the variance matrix is a *symmetric* matrix.

Variance matrix of a vector function of a random vector

Theorem (propagation law of variances):

Given a random n -vector $\underline{x}$ and a vector function $F(x)$, $F: \mathbb{R}^n \rightarrow \mathbb{R}^m$, we form the random m -vector $\underline{y} = F(\underline{x})$. If the vector function $F(x)$ is *linear*:

$$(15) \quad F(x) = \underset{m \times 1}{A} \underset{m \times n}{x} + \underset{n \times 1}{b} , \quad \underset{m \times 1}{} ,$$

then:

(16)

$$\boxed{\begin{matrix} Q_y & = & A & Q_x & A^* \\ m \times m & & m \times n & n \times n & n \times m \end{matrix}}$$

Proof:

If we denote the n column vectors of matrix A by $a_i, i=1, \dots, n$, we may write (15) as:

$$(17) \quad F(x) = \sum_{i=1}^n a_i x_i + b.$$

In a similar way we may write (9) as:

$$(18) \quad m_y = \sum_{i=1}^n a_i m_{x_i} + b.$$

Substitution of (17) and (18) into:

$$Q_y = E\{(y - m_y)(y - m_y)^*\} \int_{-\infty}^{+\infty} \dots \int_{-\infty}^{+\infty} (F(x) - m_y)(F(x) - m_y)^* p_{\Delta}(x_1, \dots, x_n) dx_1 \dots dx_n$$

gives:

$$\begin{aligned} Q_y &= \sum_{i=1}^n \sum_{j=1}^n a_i a_j^* \int_{-\infty}^{+\infty} \dots \int_{-\infty}^{+\infty} (x_i - m_{x_i})(x_j - m_{x_j}) p_{\Delta}(x_1, \dots, x_n) dx_1 \dots dx_n \\ &= \sum_{i=1}^n \sum_{j=1}^n a_i a_j^* \sigma_{x_i x_j}. \end{aligned}$$

This can be written as:

$$Q_y = (a_1, \dots, a_n) \begin{pmatrix} \sum_{j=1}^n \sigma_{x_1 x_j} a_j^* \\ \vdots \\ \sum_{j=1}^n \sigma_{x_n x_j} a_j^* \end{pmatrix}$$

or as:

$$\begin{aligned} Q_y &= (a_1, \dots, a_n) \begin{pmatrix} \sigma_{x_1 x_1} & \dots & \sigma_{x_1 x_n} \\ \vdots & & \vdots \\ \sigma_{x_n x_1} & \dots & \sigma_{x_n x_n} \end{pmatrix} \begin{pmatrix} a_1^* \\ \vdots \\ a_n^* \end{pmatrix} \\ &= A Q_x A^*. \end{aligned}$$

Without proof we also give the linearized version of the propagation law of variances.

Theorem (linearized propagation law of variances):

Given a random n -vector $\underline{x}$ and a nonlinear vector function $F(\underline{x})$, $F: \mathbb{R}^n \rightarrow \mathbb{R}^m$, we form the random m -vector $\underline{y} = F(\underline{x})$. Let $\underline{x}^0 \in \mathbb{R}^n$ be an approximation to a sample of $\underline{x}$. Then we have to a first-order:

$$(19) \quad \boxed{\begin{matrix} \mathbf{Q}_y & \doteq & \partial_x F(\underline{x}^0) & \mathbf{Q}_x & \partial_x F(\underline{x}^0)^* \\ m \times m & & m \times n & n \times n & n \times m \end{matrix}}.$$

Example 1

Let the two random variables y_1 and y_2 be defined as:

$$\begin{cases} y_1 = 1x_1 + 3x_2 + 5x_3 + 2 \\ y_2 = 4x_1 + 2x_2 - 1x_3 + 3 \end{cases}$$

The variance matrix of $\underline{x} = (x_1, x_2, x_3)^*$ is given as:

$$\mathbf{Q}_x = \begin{pmatrix} 1 & 1 & 0 \\ 1 & 2 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Then:

$$\begin{aligned} \mathbf{Q}_y &= \begin{pmatrix} 1 & 3 & 5 \\ 4 & 2 & -1 \end{pmatrix} \begin{pmatrix} 1 & 1 & 0 \\ 1 & 2 & 0 \\ 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 4 \\ 3 & 2 \\ 5 & -1 \end{pmatrix} \\ &= \begin{pmatrix} 50 & 25 \\ 25 & 41 \end{pmatrix}. \end{aligned}$$

Example 2

Let the two random variables y_1 and y_2 be defined as:

$$(20) \quad \begin{cases} y_1 = x_1 + x_2^2 + x_1 x_3 \\ y_2 = x_1^3 + \sin x_2. \end{cases}$$

The variance matrix of $\underline{x} = (x_1, x_2, x_3)^*$ is given as:

$$Q_x = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 2 \end{pmatrix}.$$

The 2-by-3 matrix of partial derivatives $\partial_x F(x^0)$ follows from (20) as:

$$(22) \quad \begin{aligned} \partial_x F(x^0)_{2 \times 3} &= \begin{pmatrix} \partial_{x_1} F^1(x^0) & \partial_{x_2} F^1(x^0) & \partial_{x_3} F^1(x^0) \\ \partial_{x_1} F^2(x^0) & \partial_{x_2} F^2(x^0) & \partial_{x_3} F^2(x^0) \end{pmatrix} \\ &= \begin{pmatrix} 1+x_3^0 & 2x_2^0 & x_1^0 \\ 3(x_1^0)^2 & \cos x_2^0 & 0 \end{pmatrix}. \end{aligned}$$

If we take as approximate values $x_1^0=1$, $x_2^0=0$, $x_3^0=0$, (22) becomes:

$$(23) \quad \partial_x F(x^0) = \begin{pmatrix} 1 & 0 & 1 \\ 3 & 1 & 0 \end{pmatrix}.$$

The variance matrix Q_y follows now from (21) and (23) to a first-order as:

$$\begin{aligned} Q_y &\doteq \begin{pmatrix} 1 & 0 & 1 \\ 3 & 1 & 0 \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 2 \end{pmatrix} \begin{pmatrix} 1 & 3 \\ 0 & 1 \\ 1 & 0 \end{pmatrix} \\ &\doteq \begin{pmatrix} 3 & 3 \\ 3 & 10 \end{pmatrix}. \end{aligned}$$

A BRIEF ACCOUNT ON THE EARLY HISTORY OF ADJUSTING GEODETIC AND ASTRONOMICAL OBSERVATIONS¹

Prof. dr ir P.J.G. Teunissen, Netherlands Geodetic Commission (KNAW)

I. Introduction

Adjustment theory can be regarded as the part of mathematical geodesy² that deals with the optimal combination of redundant observations for the purpose of determining values for parameters of interest. It is essential for a geodesist, its meaning comparable to what mechanics means to a civil engineer or a mechanical engineer. The two main reasons for performing redundant measurements are the wish to increase the accuracy of the results computed and the requirement to be able to check for errors. Due to the intrinsic uncertainty in observations, redundancy generally leads to an inconsistent system of equations. Without additional criteria, such a system is not uniquely solvable.

The problem of solving an inconsistent system of equations has attracted the attention of leading scientists in the middle of the 18th century. Historically, the first methods of combining redundant observations originate from studies in geodesy and astronomy, namely from the problem of determining the size and shape of the Earth, and the problem of finding a mathematical representation of the motions of the moon. Since its discovery almost 200 years ago, least-squares has been the most popular method of adjustment. Although the method of least squares may seem 'natural' for a modern student of adjustment theory, its discovery evolved only slowly from earlier methods of combining redundant observations³. In this contribution we sketch the historical line of development of these adjustment methods in the second half of the 18th century.

II. The method of selected points

It is convenient to cast the problem of combining redundant observations in terms of vectors and matrices. Suppose we are given a set of linear equation of the form

$$\underset{m \times 1}{y} = \underset{m \times n}{A} \underset{n \times 1}{x}$$

where y is a vector of observations, A is a given matrix of full rank and x is the vector of unknown parameters. This set of linear equations is said to be overdetermined when there are more observations than unknowns, $m > n$. The problem is to combine the m observations so that one can

¹ Based on a presentation given at the occasion of the official opening of the Geodetic-Astronomical Observatory Westerbork, 24 September 1999.

² Mathematical geodesy covers the development of theory and its implementation as is needed in order to process, analyse, integrate and validate the various geodetic data. It concerns itself with the calculus of observations (adjustment and estimation theory), with the validation of mathematical models (testing and reliability theory) and with the analysis for spatial and temporal phenomena (interpolation and prediction theory). Founders of the Dutch School of Mathematical Geodesy, internationally known as the 'Delft School', are the professors J.M. Tienstra (1895-1951) and W. Baarda (1917-2005).

³ Stigler, S.M. (1986): *The History of Statistics*, Belknap, Harvard. Hald, A. (1998): *A History of Mathematical Statistics*, Wiley-Interscience.

solve for the n unknown parameters. If we restrict ourselves to linear combinations of the observations, we can write the general solution in the following form

$$\underbrace{x}_{n \times 1} = \underbrace{B}_{n \times m} \underbrace{y}_{m \times 1} \quad \text{with} \quad \underbrace{B}_{n \times m} = \underbrace{(LA)^{-1}}_{n \times n} \underbrace{L}_{n \times m}$$

and where matrix L is a suitably chosen matrix defining the linear combinations. Different choices of L give different linear combinations and therefore different solutions. In modern terminology, matrix B is called a left-inverse of A , since B times A equals the identity matrix. Before 1750 a popular, albeit subjective, method of solving an overdetermined set of linear equations was *the method of selected points*. It consists of choosing n out of m observations (referred to as the selected points) and using their equations to solve for x . If the choice falls on the first n observations, the corresponding L matrix takes the form

$$L = \begin{bmatrix} \underbrace{I}_{n \times n} & \underbrace{0}_{n \times (m-n)} \end{bmatrix}$$

where I is the identity matrix. For the method of selected points, n residuals (the difference between the observed and adjusted observations) are by definition equal to zero. Many scientists using this method calculated the remaining $m - n$ residuals and studied their sign and size to get an impression of the goodness of fit between observations and the proposed law. The method is subjective because no clear rule is given which observations to select and which to throw out. Selecting another set of n observations leads to a different solution for x . Although the disadvantage of not using all observations was recognized, no simple method existed to tackle this shortcoming. Sometimes all possible combinations of n observations were considered and then averaged to obtain the final result. But since this approach requires handling m over n combinations, it was only practical for problems of low dimensions.

III. The method of averages

Tobias Mayer (1723-1762), professor of mathematics and head of the Göttingen observatory, made numerous observations of the Moon with the purpose of determining the characteristics of the Moon's orbit. In 1750⁴ Mayer proposed a new method for adjusting his Moon data, a method which solved the above-mentioned pitfall of the method of selected points. Apart from his adjustment method, Mayer is also known for his other contributions to surveying and navigation. In 1752 he invented the Repeating or Reflecting Circle, an instrument for observing the angle between two celestial bodies. The accuracy of Mayer's instrument was comparable to John Hadley's reflecting octant (1731), but had the advantage that it could be used to measure angles of over 90 degrees⁵. Mayer also contributed to solving the mariner's 'longitude problem'. It was the British Parliament, which in 1714, offered the 'Longitude Prize' to those who could find a 'useful and practical' method for determining longitude at sea. To determine longitude of a ship at sea, the mariner needs to know both his local time and the time at some standard location. Local time was readily determined, but the determination of standard time at sea was more complicated. Mayer's detailed lunar tables (1755) made it possible to translate the instrument readings into longitude positions. The use of Mayer's lunar tables was later superseded

⁴ Mayer, T. (1750): Abhandlung ueber die Umwalzung des Monds um seine Axe und die scheinbare Bewegung der Mondsflexen, *Kosmog. Nachr. Samml. Auf das Jahr 1748*, 1, 52-183.

⁵ Forbes, E.G. (1974): *The Birth of Scientific Navigation*, Maritime Monographs and Reports, No. 10, National Maritime Museum, Greenwich, London.

by John Harrison's marine chronometer H-4 (1759). In recognition of their contributions, both men were awarded part of the 'Longitude Prize', with the larger sum going to Harrison⁶.

Mayer studied the libration of the Moon by observing the changing position of the crater Manilius as seen from the Earth. Using spherical geometry, he found a linearized relationship between his observables and some location parameters of Manilius and the Moon's pole. This gave him an inconsistent system of $m = 27$ linear equations in 3 unknown parameters

$$\underbrace{\begin{bmatrix} y_1 \\ \vdots \\ y_m \end{bmatrix}}_y = \underbrace{\begin{bmatrix} 1 & a_{12} & a_{13} \\ \vdots & \vdots & \vdots \\ 1 & a_{m2} & a_{m3} \end{bmatrix}}_A \underbrace{\begin{bmatrix} x_1 \\ \vdots \\ x_3 \end{bmatrix}}_x$$

Mayer proposed to divide the 27 equations into 3 groups of 9 each, to sum the equations within each group, and to solve the resulting 3 equations in 3 unknowns. For a general set of m equations in n unknowns, this approach amounts to a separation of the m equations into n groups, followed by a groupwise summation. For example, in case $n = 2$, the corresponding L matrix takes the form

$$L = \begin{bmatrix} e_1 & 0 \\ 0 & e_2 \end{bmatrix}$$

where the e_1 and e_2 are row vectors having only 1's as their entries. Since one may use averages instead of sums, the method became later known as *the method of averages*. Mayer's method of averages soon became popular. It used all observations and it was very simple to apply. However, due to the lack of an objective criterion on how to group the observations, the method was still a subjective one.

IV. The method of least absolute deviations

To determine the Earth's flattening as predicted by Newton's theory of gravitation (*Principia*⁷, 1687), the French Academy of Sciences organized arc-measurement expeditions to Peru, Lapland and the Cape of Good Hope in the period 1735-1754. These expeditions aroused the interest in other countries and in 1750 Pope Benedict XIV commissioned the Jesuit and professor of mathematics Roger Joseph Boscovich (1711-1787) to perform a similar geodetic survey near Rome, the results of which were published in 1755. In a summary of this report, published in 1757⁸, Boscovich formulated his new method, now known as *the method of least absolute deviations*, and applied it to the data of the French and Italian arc measurements.

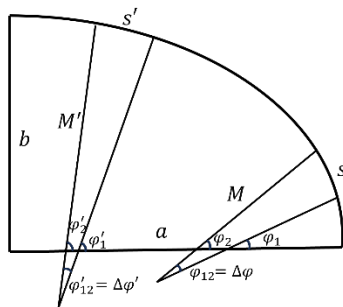


Figure 1: Latitude arc measurements along a meridian.

In order to understand the equations used by Boscovich, we first need to introduce some elements from ellipsoidal geodesy. For short meridian arcs, the arc length s (see figure 1) can be written as $s = M(\varphi)\Delta\varphi$, with $M(\varphi)$ the meridian radius of curvature, φ the geodetic latitude of the midpoint of the arc and $\Delta\varphi$ the latitude difference of the two arc endpoints.

⁶ Sobel, D. (1995): *Longitude*. Walker, New York, Dutch translation by E. van Altena (1997): Dava Sobel: *Lengtegraad*, Ambo.

⁷ Newton, I. (1687): *Philosophiae Naturalis Principia Mathematica*.

⁸ Boscovich, R.J., Maire, C. (1770): *Voyage astronomique et géographique dans l'état de l'église*. French translation of original 1755 publication. For Boscovich' contributions, see also Sheynin, O.B. (1973): *Arch. History Exact Sci.*, 9: 306-324.

The meridian curvature and its expansion are given as

$$M = \frac{a(1 - e^2)}{(1 - e^2 \sin^2 \varphi)^{3/2}} = a(1 - e^2) \left\{ 1 + \frac{3}{2} e^2 \sin^2 \varphi + \dots \right\}$$

with $e^2 = (a^2 - b^2)/a^2$ the eccentricity, a and b the half lengths of the major and minor axis and $a(1 - e^2)\Delta\varphi$ the length of an arc at the equator. Using only the first two terms in the expansion, the lengths of an $\Delta\varphi$ -arc can be written as

$$s = x_1 + \sin^2 \varphi x_2 \quad \text{with} \quad x_1 = a(1 - e^2)\Delta\varphi \quad \text{and} \quad x_2 = \frac{3}{2}ae^2(1 - e^2)\Delta\varphi$$

This is one equation in two unknowns, x_1 and x_2 . The arc length s and geodetic latitude φ are determined from astronomical and geodetic measurements, while x_1 and x_2 contain the unknown dimensions of the ellipsoid of revolution. Although a minimum of two arcs is needed to solve for the unknowns, it is preferable to use more than two $\Delta\varphi$ -arcs. As a result, one obtains the following system of linear equations

$$\underbrace{\begin{bmatrix} s_1 \\ \vdots \\ s_m \end{bmatrix}}_y = \underbrace{\begin{bmatrix} 1 & \sin^2 \varphi_1 \\ \vdots & \vdots \\ 1 & \sin^2 \varphi_m \end{bmatrix}}_A \underbrace{\begin{bmatrix} x_1 \\ x_2 \end{bmatrix}}_x$$

This is the system of equations which formed the start of Boscovich' analysis. Note that the A -matrix becomes near rank defect, when all arcs are chosen to the same latitude. For an accurate determination of the two unknowns, it is therefore preferable to choose arcs at widely different latitudes. From the data available, Boscovich chose five such arcs ($m = 5$). In his first analyses, Boscovich used the method of selected points. He chose the two arcs with the largest difference in latitude. Not satisfied with the result obtained (the 3 residuals were considered too large), he considers all possible pairs of measured arcs. This gave him 10 selected points to solve, but again he is not satisfied with the results obtained. After having struggled for some time on how to proceed, Boscovich finally formulates his new method of solution in 1757. He states that the parameters x_1 and x_2 should be chosen in such a way that the residuals sum up to zero and have minimum absolute sum. In formula form these two conditions read

$$\sum_{i=1}^m (s_i - x_1 - x_2 \sin^2 \varphi_i) = 0 \quad \text{and} \quad \sum_{i=1}^m |s_i - x_1 - x_2 \sin^2 \varphi_i| = \min$$

The first condition (although not essential) was motivated by the assumed symmetry in the error distribution, while the second was chosen to get the adjusted values 'as close as possible' to the observed ones. Boscovich gave the graphical algorithm for solving his problem, but no analytical one. The analytical proof of the solution was first given by Laplace in 1793. Using his principle, Boscovich first determined the two parameters x_1 and x_2 and from them the flattening as $f = x_2/3x_1$. Here he only used the first term of the expansion

$$f = \frac{(a - b)}{a} = \frac{1}{2}e^2 + \frac{1}{8}e^4 + \frac{1}{16}e^6 + \dots$$

The value obtained by Boscovich equals $f = 1/246$, which was smaller than the flattening predicted by Newton. Based on the rotational ellipsoid as an equilibrium figure for a homogeneous, fluid, rotating Earth, Newton obtained the value $f = 1/230$. Boscovich value is however larger than the value known today (International Association of Geodesy (1980): $f = 1/298.257$). Boscovich' method was the first adjustment method that started from the principle of minimizing a function of the

residuals. However, although the method is objective and uses all the observations, it did not reach the same level of popularity as Mayer's method. The method, being nonlinear, was difficult to apply, while the at that time available algorithm could only handle a system of equations with a maximum of two unknowns. In the second half of the 20th century the method gained in popularity due to its property of being resistant (robust) against outliers. Nowadays Boscovich adjustment method is usually referred to as an L_1 adjustment, since the L_1 -norm of a vector is the sum of absolute values of its entries.

V. The method of least squares

Adrien-Marie Legendre (1752-1833), a professor of mathematics at the École Militaire in Paris, was appointed by the French Academy of Sciences as a member of various committees on astronomical and geodetic projects, among them the committee on the standardization of weights and measures. The committee proposed to define the meter as 10^{-7} times the length of the terrestrial meridian quadrant through Paris at mean sea-level. The arc-measurements took place in the period 1792-1795 and were analysed, among others⁹, by both Laplace and Legendre. Legendre's 1805¹⁰ publication on the determination of the orbits of comets, contains a nine-page appendix in which for the first time *the method of least-squares*¹¹ is described, together with an application of the method to the arc measurements. Legendre used a different equation than Boscovich. Boscovich used the equation

$$s = \left[a(1 - e^2) + \frac{3}{2}ae^2(1 - e^2)\sin^2\varphi \right] \Delta\varphi$$

while Legendre, having the determination of the meter in mind, used a parametrization which differed from the one used by Boscovich. With $\sin^2\varphi = \frac{1}{2}(1 - \cos 2\varphi)$, $d = a(1 - e^2) + \frac{1}{2}\frac{3}{2}ae^2(1 - e^2)$, $\sin\Delta\varphi \approx \Delta\varphi$, and $fd \approx \frac{1}{2}a(1 - e^2)e^2$, the above equation can be written as $s = d\Delta\varphi - \frac{3}{2}fd\sin\Delta\varphi \cos 2\varphi$. As a result, we may write for the two endpoints of an arc,

$$\varphi_{12} = s_{12}d^{-1} + \frac{3}{2}f\sin\varphi_{12}\cos(\varphi_1 + \varphi_2).$$

This is one equation in two unknowns, d and f . Note that $d\frac{\pi}{180}$ is the length of a 1-degree arc at 45 degree latitude.

Legendre understood that the observed latitude differences would correlate in case the arcs were connected. He therefore transformed the above equation of differences into an equivalent undifferenced form. This can be achieved by introducing an appropriate additional equation with an additional unknown. As a result, we obtain Legendre's linear system of equations as

⁹ In 1795 the French government invited other European governments to delegate scientists to Paris for completing and checking the computations for the standardization of weights and measures. The Dutch delegates were professor Jan Hendrik van Swinden and the navy officer Henricus Aeneae. The European scientists met in Paris in 1798 and reported on their findings in 1799. The report on the meter was given by van Swinden, the standard meter, a bar of platinum with rectangular section of 254 millimetres, was placed in the French State Archives, Mètre et Kilogramme des Archives. In 1983 the standard meter was defined as the length travelled by light in vacuum in $1/299.792.458$ seconds.

¹⁰ Legendre, A.M. (1805): *Nouvelles méthodes pour la détermination des orbites des comètes*. (Appendix: *Sur la méthode des moindres carrés*).

¹¹ When Carl Friedrich Gauss published his first probabilistic version of the method of least-squares in 1809, he claimed that he had been using the method ('our principle') already since 1795. This claim resulted in a priority dispute between Legendre and Gauss, for a discussion see e.g. Plackett, R.L. (1972): The discovery of the method of least-squares. *Biometrika* 59: 239-251.

$$\underbrace{\begin{bmatrix} \varphi_1 \\ \vdots \\ \varphi_l \\ \vdots \\ \varphi_m \end{bmatrix}}_y = \underbrace{\begin{bmatrix} 1 & s_{l1} & \frac{3}{2} \sin \varphi_{l1} \cos(\varphi_1 + \varphi_l) \\ \vdots & \vdots & \vdots \\ 1 & 0 & 0 \\ \vdots & \vdots & \vdots \\ 1 & s_{lm} & \frac{3}{2} \sin \varphi_{lm} \cos(\varphi_m + \varphi_l) \end{bmatrix}}_A \underbrace{\begin{bmatrix} \varphi_l \\ d^{-1} \\ f \end{bmatrix}}_x$$

Since Legendre had four connected arcs ($m = 5$) at his disposal, he had to solve 5 equations in 3 unknowns. In order to solve this overdetermined linear system of equations, Legendre proposed to determine x such that the sum of the squares of the residuals is minimized. In a vector-matrix form $(y - Ax)^T(y - Ax) = \min$.

By setting the derivatives of this quadratic form equal to zero, he shows that the solution satisfies the consistent system of linear equations (nowadays referred to as ‘the normal equations’) $A^T y = A^T Ax$. Note that this result corresponds to the following choice of the L matrix: $L = A^T$.

After solving his set of equations for the three unknowns, φ_3 , d and f , Legendre obtained for the flattening the value $f = 1/148$, which he recognizes as being too large. He therefore recomputed his least-squares adjustment, but now with f constrained to the at that time adopted value for the Earth’s flattening. As a result, he obtained a value for d , which was now close to the value obtained earlier by Laplace and on which the actual definition of the meter was based. The reason for Legendre having to constrain f lies in the poor resolution of his data. The total arc length of his data covered only about 10 degrees.

Although Legendre did not give a clear motivation for his ‘least-squares’ criterion, he did realize its potential. His method used all the observations, had an objective criterion and most importantly, resulted – as opposed to Boscovich’ method – in a solvable *linear* system of equations. The method met with almost immediate success. Within ten years after Legendre’s publication, the method of least-squares became a standard tool in astronomy and geodesy in various European countries. And within twenty years, also the probabilistic foundations of the method were largely completed, the main contributors being Laplace and Gauss.

Literature

1. Calculus

Almering, J.H.J. Revised by H. Bavinck and R.W. Goldbach; *Analyse*
VSSD, 6th ed. (1996)

2. Linear Algebra

Braber, C.A. den, H. van Iperen and M.A. Vlietgever; *Matrixrekening*
VSSD, 2nd ed. (1989)

Lay, D.C.; *Linear algebra and its applications*
Reading, Addison-Wesley, 2nd ed. (1997)

Ortega, J.M.; *Matrix theory, a second course*
New York, Plenum Press (1987)

Strang, G.; *Linear algebra and its applications*
San Diego, Harcourt Brace Jovanovich Publishers, 3rd ed. (1988)

3. Probability Theory

Grimmett, G. and D. Welsh; *Probability; an introduction*
Oxford, Clarendon Press (1986)

Lopuhaä, H.P.; *Lecture notes: Statistiek voor geodeten*
Delft University of Technology, Faculty of Information Technology and Systems
(1994)

Papoulis, A.; *Probability, random variables and stochastic processes*
New York, McGraw-Hill Book Company (1991)

Soest, J. van; *Elementaire statistiek*
VSSD, 7th ed. (1997)

4. Estimation Theory

Cooper, M.A.R.; *Control surveys in civil engineering*
London, Collins Publishers (1987)

Koch, K.R.; *Parameter estimation and hypothesis testing in linear models*
Berlin Springer Verlag, 2nd (1999)

Mikhail, E.M., Gracie, G.; *Analysis and adjustment of survey measurements*
New York, van Nostrand Reinhold Company (1981)

Tienstra, J.M.; *Theory of the adjustment of normally distributed observations*
Amsterdam, Argus, (1956)

Index

a

A-model; see model with observation equations 39

b

Best Linear Unbiased Estimation 28, 49, 64, 174

block estimation 109, 113

B-model; see model with condition equations 61

c

Cholesky 24

condition equation 62

constituent variates 77

constrained optimization 168, 176

covariance 188

covariance matrix; see variance matrix 187

d

degree of dependence 94

derived variates 77

dispersion 89; see also variance matrix

e

eigenvalues 58

ellipse 14, 18, 34

estimable 94

estimate 26

estimation in phases 123

estimator 26

expectation 26, 178

f

free variates 72, 75

functional model 3

h

Helmert 113

Helmert's block method 113

i

inner product 23, 56

intercept 98

iteration 150

l

Lagrange 171

Lagrange multiplier rule 171

least-squares 6

least-squares residuals 26

levelling 39, 61

linearized A-model 146

linearization 142

m

maximum likelihood 33

mean 178, 185

measurement error 5

measurement update 101

minimization 29, 52, 164, 168

minimum variance 28, 52

misclosures 62

mixed model 81

model with condition equations 61

model with observation equations 39

n

nonlinear observation equation 137

nonlinear A-model 137

nonlinear B-model 154

normal equation 90

o

observation (measurement) 2

observation equation 42

optimization 164

optimization constrained 168, 176

orthogonal projection 15

orthogonal projector 8, 55, 95

p

partitioned model 89

phases, estimation in 123

predicted residual 101

propagation law 26, 45, 46, 47, 179, 185

r

rank defect 43, 95

random variable 178

random variable, vector 184

recursive estimation 100

reduction 91

redundancy 55, 61

remainder 163

s

slope 97, 98

stochastic model 3

t

Taylor's theorem 142, 159

Tienstra 123

trace 58

u

unbiased 27, 49

v

variance 181, 188

variance matrix 187

w

weight matrix 11, 43

weighted least-squares 11, 44

y

y^R -variates 74

Adjustment theory: an introduction

Peter J.G. Teunissen

Adjustment theory can be regarded as the part of mathematical geodesy that deals with the optimal combination of redundant measurements together with the estimation of unknown parameters. It is essential for a geodesist, its meaning comparable to what mechanics means to a civil engineer or a mechanical engineer. Historically, the first methods of combining redundant measurements originate from the study of three problems in geodesy and astronomy, namely to determine the size and shape of the Earth, to explain the long-term inequality in the motions of Jupiter and Saturn, and to find a mathematical representation of the motions of the Moon. Nowadays, the methods of adjustment are used for a much greater variety of geodetic applications, ranging from, for instance, surveying and navigation to remote sensing and global positioning.

The two main reasons for performing redundant measurements are the wish to increase the accuracy of the results computed and the requirement to be able to check for errors. Due to the intrinsic uncertainty in measurements, measurement redundancy generally leads to an inconsistent system of equations. Without additional criteria, such a system of equations is not uniquely solvable. In this introductory course on adjustment theory, methods are developed and presented for solving inconsistent systems of equations. The leading principle is that of least-squares adjustment together with its statistical properties.

The inconsistent systems of equations can come in many different guises. They could be given in parametric form, in implicit form, or as a combination of these two forms. In each case the same principle of least-squares applies. The algorithmic realizations of the solution will differ however. Depending on the application at hand, one could also wish to choose between obtaining the solution in one single step or in a step-wise manner. This leads to the need of formulating the system of equations in partitioned form. Different partitions exist, measurement partitioning, parameter partitioning, or a partitioning of both measurements and parameters. The choice of partitioning also affects the algorithmic realization of the solution. In this introductory text the methodology of adjustment is emphasized, although various samples are given to illustrate the theory. The methods discussed form the basis for solving different adjustment problems in geodesy.



P.J.G. Teunissen

Delft University of Technology
Faculty of Civil Engineering and Geosciences

Dr (Peter) Teunissen is Professor of Geodesy at Delft University of Technology (DUT) and an elected member of the Royal Netherlands Academy of Arts and Sciences. He is research-active in various fields of Geodesy, with current research focused on the development of theory, models, and algorithms for high-accuracy applications of satellite navigation and remote sensing systems. His past DUT positions include Head of the Delft Earth Observation Institute, Education Director of Geomatics Engineering and Vice-Dean of Civil Engineering and Geosciences. His books at TUDelft Open are Adjustment Theory, Testing Theory, Dynamic Data Processing and Network Quality Control.

© 2024 TU Delft Open
ISBN 978-94-6366-884-2
DOI <https://doi.org/10.59490/tb.95>
textbooks.open.tudelft.nl

Cover image: J.E. Alberda

Series on Mathematical
Geodesy and Positioning