

ON THE LINEAR LEAST SQUARES PROBLEM
WITH A QUADRATIC CONSTRAINT

by

Walter Gander

STAN-CS-78-697
November 19 78

COMPUTER SCIENCE DEPARTMENT
School of Humanities and Sciences
STANFORD UNIVERSITY



On the Linear Least Squares Problem
with a Quadratic Constraint

On the Linear Least Squares Problem
with a Quadratic Constraint

by
Walter Gander ^{*/}

Abstract.

In this paper we present the theory and practical computational aspects of the linear least squares problem with a quadratic constraint. New theorems characterizing properties of the solutions are given and extended for the problem of minimizing a general quadratic function subject to a quadratic constraint. For two important regularization methods we formulate dual equations which proved to be very useful for the application of smoothing of datas. The resulting algorithm is a numerically stable version of an algorithm proposed by Rutishauser. We show also how to choose a third order iteration method to solve the secular equation. However we are still far away from a foolproof machine independent algorithm.

Walter Gander*

by

* Neu-Technikum **Buchs, CH-9470 Buchs**, Switzerland.
Research supported in part by National Science Foundation Grant
No. **MCs78-17697**.

^{*/} Computer Science Department, Stanford University, Stanford, Calif. 94305.
On leave from Neu-Technikum Ruchs, Switzerland.
The author has been supported by the Swiss National Science Foundation,
Grant No. **5'521'330'61517**.

Acknowledgement.

This work was done during my Stanford year 1977/78. I am greatly indebted to Prof. G. Golub for his valuable comments and suggestions, for his hospitality and for creating the extraordinary spirit of Serra House that attracts in one year most numerical analysts of the world. My stay would not have been possible without the support of Prof. P. Henrici, one of my "godfathers" for the Swiss NSF. I would like to thank him for his continuous interest in my work and for his encouragement. I owe him much of my mathematical education. Also in Switzerland I have to thank Prof. J. Marti, my second "godfather", who happened to work at a similar problem. I am indebted to Prof. Ch. van Loan, Cornell University, for his careful reading of the manuscript and for his various suggestions helping to improve my english. Thanks also to the "**Forschungskommission**" of **ETHZ** and the Swiss National Science Foundation who gave me the grant that enabled my stay at Stanford. Finally I have to thank my wife Heidi and my parents-in-law, A. and E. Wolf, for their non-mathematical assistance.

Stanford, Fall 1978
Walter Gander

To Maurice

Table of Contents

0.	Introduction	1
1.	Basic Definitions and Remarks	5
2.	The Least Square Problem with Quadratic Constraint	9
2.1	Characterization of the Solution	15
2.2	The Solutions of the Normal Equations	23
2.3	The Solution for the Equality Constraint (PIE).	24
2.4	The Solution for Inequality Constraint (PI)	27
3.	The Relaxed Least Squares Problem	28
3.1	Results from the General Theory	29
3.2	The Dual Normal Equations	33
3.3	Elden's Transformation	35
3.4	Rutishauser's Relaxed and Doubly Relaxed Least Squares Problem	42
4.	Minimum Norm Solution with Given Norm of the Residual	43
4.1	The Dual Normal Equations	44
4.2	Representation as Least Squares Problems	45
5.	Computational Aspects	46
5.1	Solution of a Relaxed Least Squares Problem with Band Matrix	53
5.2	Solution of a Least Squares Problem with Two Band Matrices	56
5.3	Bidiagonalization	62
5.4	Computation of the Derivatives of the Length Function	71
6.	One Point Iterative Methods to Solve the Secular Equation	72
6.1	Convergence Factors	77
6.2	Third Order Iterative Methods	82
6.3	The Convergence Factors for a Third Order Method	84
6.4	Geometrical Interpretation of the Convergence Factors	90
6.5	Solving the Secular Equation	95
7.	Generalizations	99
8.	Smoothing of Datas	112
	References	

0. Introduction.

In this paper we consider least squares problems with a quadratic constraint. The matrices and vectors will be real and we use capital letters $A, B, \dots$ for matrices and small letters $a, b, \dots$ for vectors. Π will be used for the euclidian vectornorm.

Given A, C, b, d and a number $\alpha > 0$ we consider the problem to find x such that

$$\left. \begin{aligned} \|Ax - b\| &= \min \\ \text{subject to } \|Cx - d\| &\leq \alpha \end{aligned} \right\} \quad (0.1)$$

This problem is a generalization of the least square problem with equality constraints

$$\left. \begin{aligned} \|Ax - b\| &= \min \\ \text{subject to } Cx &= d \end{aligned} \right\} \quad (0.2)$$

since for $\alpha = 0$ every solution of (0.1) is also a solution of (0.2).

The motivation why to consider (0.1) rather than (0.2) is explained best by the following example. Let's assume we are given m values of a function f

$$y_i = f(t_i), \quad i = 1, \dots, m$$

and we seek the coefficients of a polynomial

$$p(t) = \sum_{i=0}^{n-1} x_i t^i$$

that approximates f .

If we insist that p interpolate the given data, then we have to solve the system (m equations and n unknowns)

$$\underline{\underline{A}}\underline{\underline{x}} = \underline{\underline{y}} \quad (0.3)$$

with $a_{ij} = t_i^j$. If $m > n$, (0.3) may have no solution. If $m < n$, the solution is not unique. If $m = n$, there is a unique solution if $t_i \neq t_j$, $i \neq j$. However it is well known that $\underline{\underline{A}}$ is ill conditioned which means that $\underline{\underline{x}}$ is difficult to compute accurately and that usually the norm of $\underline{\underline{x}}$ will be large. The polynomial p with large coefficient will be useless since cancellation will affect its evaluation. It may therefore be better not to interpolate but to approximate the datas in the least squares sense. This leads to the unconstrained least squares problem

$$\|\underline{\underline{A}}\underline{\underline{x}} - \underline{\underline{y}}\| = \min . \quad (0.4)$$

This problem has for $m \geq n$ and if $\underline{\underline{A}}$ has full rank a unique solution. If $\underline{\underline{A}}$ is rank deficient (0.4) has infinitely many solutions and therefore one usually looks for the solution with minimum norm

$$\left. \begin{array}{l} \min \|\underline{\underline{x}}\| \\ \text{subject to } \|\underline{\underline{A}}\underline{\underline{x}} - \underline{\underline{y}}\| = \min . \end{array} \right\} \quad (0.5)$$

The solution of (0.5) is given explicitly by

$$\underline{\underline{x}} = \underline{\underline{A}}^+ \underline{\underline{y}}$$

where $\underline{\underline{A}}^+$ is the pseudoinverse of $\underline{\underline{A}}$. However the solution of (0.4) [35] and (0.5) may also suffer having a large norm. The problem is then said to be ill posed and we consider two ways to regularize [42] the solution. We can prescribe a bound for $\|\underline{\underline{x}}\|$ and thus look for the solution of

$$\left. \begin{array}{l} \|\underline{\underline{A}}\underline{\underline{x}} - \underline{\underline{y}}\| = \min \\ \text{subject to } \|\underline{\underline{x}}\| \leq \alpha . \end{array} \right\} \quad (0.6)$$

Alternatively we may prefer that the deviation from the given data points be bounded. This leads to

$$\left. \begin{array}{l} \|\underline{\underline{x}}\| = \min \\ \text{subject to } \|\underline{\underline{A}}\underline{\underline{x}} - \underline{\underline{y}}\| \leq \alpha . \end{array} \right\} \quad (0.7)$$

The question how to choose α is not simple to answer. It may be necessary to use some statistical tools to estimate α [17]. We observe that (0.6) and (0.7) are least squares problems with a quadratic constraint. The degree of the polynomial (n-1) may be chosen independently of the number of given data points. We may wish to compute a low degree polynomial ($m \gg n$) or a polynomial with large degree $m < n$ but with small coefficients. Problems (0.6) and (0.7) have a unique solution (if it exists) and we will show how to compute it efficiently.

Returning to our example, let's consider a partitioning of the data points in two sets. We may ask for a polynomial that interpolates the datas of the first set exactly. This leads to a least squares problem with equality constraints:

$$\left. \begin{array}{l} \|\underline{\underline{C}}\underline{\underline{x}} - \underline{\underline{y}}_2\| = \min \\ \text{subject to } \underline{\underline{B}}\underline{\underline{x}} = \underline{\underline{y}}_1 \end{array} \right\} \quad (0.8)$$

where $\underline{\underline{y}}_1$ contains the function values of the first set. Instead of interpolating on the first set we could ask for a bound of the deviation and get again problem (0.1):

$$\left. \begin{array}{l} \|\underline{\underline{C}}\underline{\underline{x}} - \underline{\underline{y}}_2\| = \min \\ \text{subject to } \|\underline{\underline{B}}\underline{\underline{x}} - \underline{\underline{y}}_1\| \leq \alpha . \end{array} \right\} \quad (0.9)$$

In contrast to (0.6) and (0.7), (0.9) may not have a unique solution. As we shall see the solution is unique if and only if the nullspaces of $\underline{\underline{C}}$ and $\underline{\underline{B}}$ intersect trivially.

1. Basic Definitions and Remarks.

The following notation will be used:

Problem (P1):

$$\|\underline{\underline{A}}\underline{\underline{x}} - \underline{\underline{b}}\| = \min \quad (1.1)$$

subject to

$$\|\underline{\underline{C}}\underline{\underline{x}} - \underline{\underline{d}}\| \leq \alpha \quad (1.2)$$

If we have $\|\underline{\underline{C}}\underline{\underline{x}} - \underline{\underline{d}}\| = \alpha$ instead of (1.2) we will refer to the problem as (P1E). Two special cases of (P1) will be considered:

Problem (P2): like (P1) but $\underline{\underline{C}} = \underline{\underline{I}}$, $\underline{\underline{d}} = \underline{\underline{0}}$.

Problem (P3): like (P1) but $\underline{\underline{A}} = \underline{\underline{I}}$, $\underline{\underline{b}} = \underline{\underline{0}}$.

Finally (P2E) and (P3E) will be the corresponding problems with equality sign in the constraint.

The solution of (P1) is a stationary point of the Lagrange function (with the Lagrange multiplier λ)

$$L(\underline{\underline{x}}, \lambda) = \|\underline{\underline{A}}\underline{\underline{x}} - \underline{\underline{b}}\|^2 + \lambda \{\|\underline{\underline{C}}\underline{\underline{x}} - \underline{\underline{d}}\|^2 - \alpha^2\}$$

and therefore a solution of $\frac{\partial L}{\partial \underline{\underline{x}}} = \underline{\underline{0}}$ and $\frac{\partial L}{\partial \lambda} = 0$, which are the "normal equations":

$$(\underline{\underline{A}}^T \underline{\underline{A}} + \lambda \underline{\underline{C}}^T \underline{\underline{C}}) \underline{\underline{x}} = \underline{\underline{A}}^T \underline{\underline{b}} + \lambda \underline{\underline{C}}^T \underline{\underline{d}} \quad (1.3)$$

$$\|\underline{\underline{C}}\underline{\underline{x}} - \underline{\underline{d}}\|^2 = \alpha^2 \quad (1.4)$$

If the matrix $\underline{\underline{A}}^T \underline{\underline{A}} + \lambda \underline{\underline{C}}^T \underline{\underline{C}}$ is nonsingular, then we can define

$$\underline{\underline{x}}(\lambda) := \|\underline{\underline{C}} \underline{\underline{x}}(\lambda) - \underline{\underline{d}}\|^2 \quad (1.5)$$

where $\underline{\underline{x}}(\lambda)$ is the solution of (1.3). We will call f the "length

function". To determine a solution we have to solve the "secular equation"

$$f(h) = \alpha^2. \quad (1.6)$$

Finally we observe that for $\lambda > 0$ equations (1.3) and (1.4) are the normal equations of the least squares problem

$$\left\| \begin{pmatrix} \frac{A}{\sqrt{\lambda}} & \frac{C}{\sqrt{\lambda}} \end{pmatrix} \bar{x} - \begin{pmatrix} \frac{b}{\sqrt{\lambda}} \\ \frac{d}{\sqrt{\lambda}} \end{pmatrix} \right\| = \min.$$

A useful tool for the analysis of problems (P2) and (P3) is the singular value decomposition (SVD) [14] and its generalization (BSVD) [45] for (P1). These decompositions can also be used for the practical computations. However for the problems (P1E), (P2E) and (P3E) there are less expensive ways [8]. In some applications [33], $\bar{A}$ and $\bar{C}$ are band matrices and (P1) may be solved efficiently without transformations as we shall show.

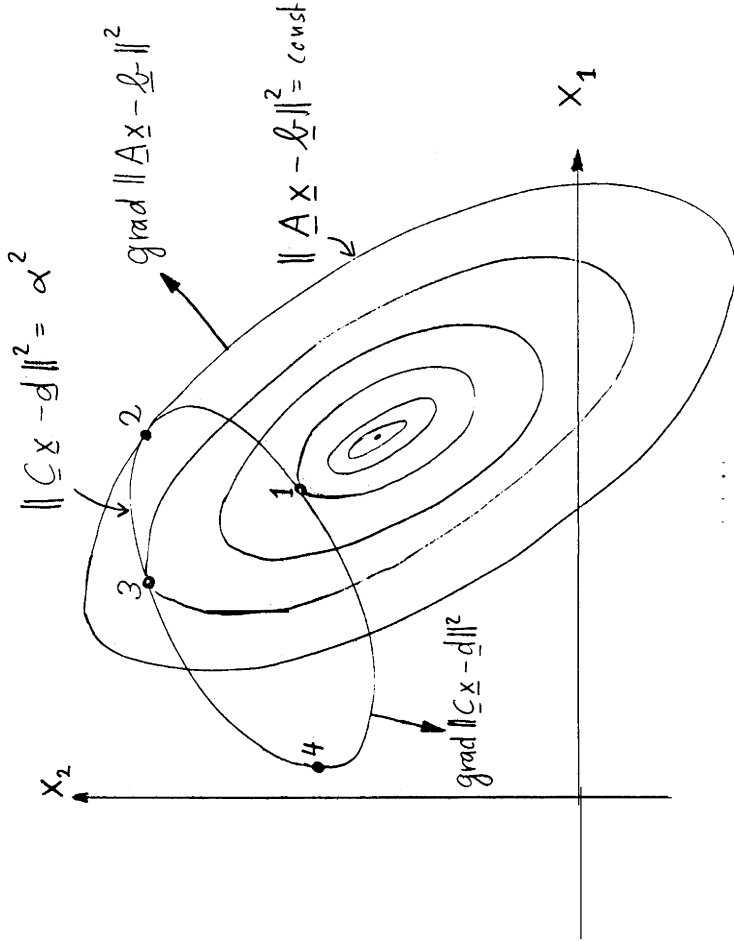
2. The Least Squares Problem with a Quadratic Constraint.

Let A be an $(m \times n)$ -matrix, C a $(p \times n)$ -matrix, b an m -vector, d a p -vector, and α a positive number.

We consider the problem to find an n -vector x so that

$$\left. \begin{aligned} \|\bar{A}\bar{x} - \bar{b}\| &= \min \\ \text{subject to } \|\bar{C}\bar{x} - \bar{d}\| &= \alpha \end{aligned} \right\} \quad (P1E)$$

For $n = 2$ we can interpret this problem geometrically. The level lines of $\|\bar{A}\bar{x} - \bar{b}\|^2 = \text{const}$ are ellipses centered at $\bar{A}^+ \bar{b}$. The constraints $\|\bar{C}\bar{x} - \bar{d}\|^2 = \alpha^2$ is also an ellipse:



We are looking for a point $\underline{\mathbf{x}}$ on the ellipse $\|\underline{\mathbf{C}}\underline{\mathbf{x}} - \underline{\mathbf{d}}\|^2 = \alpha^2$ which has the **smallest** value of $\|\underline{\mathbf{A}}\underline{\mathbf{x}} - \underline{\mathbf{b}}\|^2$. Clearly it is the point 1. At that point the **gradients** of the two functions

$$\begin{aligned} \|\underline{\mathbf{C}}\underline{\mathbf{x}} - \underline{\mathbf{d}}\|^2, \quad \|\underline{\mathbf{A}}\underline{\mathbf{x}} - \underline{\mathbf{b}}\|^2 \\ \underline{\text{grad}} \|\underline{\mathbf{A}}\underline{\mathbf{x}} - \underline{\mathbf{b}}\|^2 = -\lambda \underline{\text{grad}} \|\underline{\mathbf{C}}\underline{\mathbf{x}} - \underline{\mathbf{d}}\|^2. \end{aligned}$$

If we rearrange this equation we have

$$(\underline{\mathbf{A}}^T \underline{\mathbf{A}} + \lambda \underline{\mathbf{C}}^T \underline{\mathbf{C}}) \underline{\mathbf{x}} = \underline{\mathbf{A}}^T \underline{\mathbf{b}} + \lambda \underline{\mathbf{C}}^T \underline{\mathbf{d}}$$

which is equation (1.3) that we obtained using the technique of the Lagrange multipliers.

We observe that there are three other points for which the gradients of the two functions are parallel (2,3,4). However for these three points λ is negative since the gradients have the same directions.

We can think of problem (P1E) as blowing up a balloon centered at $\underline{\mathbf{A}}^+ \underline{\mathbf{b}}$ which has an ellipsoid shape until it touches the ellipsoid $\|\underline{\mathbf{C}}\underline{\mathbf{x}} - \underline{\mathbf{d}}\|^2 = \alpha^2$.

If $\underline{\mathbf{A}}^+ \underline{\mathbf{b}}$ is outside of $\|\underline{\mathbf{C}}\underline{\mathbf{x}} - \underline{\mathbf{d}}\|^2$ then at the point of touch the gradients **will** have opposite sign, i.e., $\lambda > 0$. However if $\underline{\mathbf{A}}^+ \underline{\mathbf{b}}$ is inside the gradients have the **same** sign and $\lambda < 0$.

For the problems with inequality constraints (P1), (P2) and (P3) we have only to consider the case that $\underline{\mathbf{A}}^+ \underline{\mathbf{b}}$ is outside of $\|\underline{\mathbf{C}}\underline{\mathbf{x}} - \underline{\mathbf{d}}\|^2$. If it is inside then $\underline{\mathbf{x}} = \underline{\mathbf{A}}^+ \underline{\mathbf{b}}$ solves the problem.

2.1 Characterization of the Solution.

The solution of (P1E) is **among** the solutions $(\underline{\mathbf{x}}, \lambda)$ of the normal equations (see Section 1):

$$\left. \begin{aligned} (\underline{\mathbf{A}}^T \underline{\mathbf{A}} + \lambda \underline{\mathbf{C}}^T \underline{\mathbf{C}}) \underline{\mathbf{x}} &= \underline{\mathbf{A}}^T \underline{\mathbf{b}} + \lambda \underline{\mathbf{C}}^T \underline{\mathbf{d}} \\ \|\underline{\mathbf{C}}\underline{\mathbf{x}} - \underline{\mathbf{d}}\|^2 &= \alpha^2 \end{aligned} \right\} \quad (2.1)$$

The following theorem compares two solutions of these equations:

Theorem 2.1. If $(\underline{\mathbf{x}}_1, \lambda_1)$ and $(\underline{\mathbf{x}}_2, \lambda_2)$ are solutions of the normal equations (2.1), then

$$\|\underline{\mathbf{A}}\underline{\mathbf{x}}_2 - \underline{\mathbf{b}}\|^2 - \|\underline{\mathbf{A}}\underline{\mathbf{x}}_1 - \underline{\mathbf{b}}\|^2 = \frac{\lambda_1 - \lambda_2}{2} \|\underline{\mathbf{C}}(\underline{\mathbf{x}}_1 - \underline{\mathbf{x}}_2)\|^2. \quad (2.2)$$

Proof. Since $(\underline{\mathbf{x}}_1, \lambda_1)$, $(\underline{\mathbf{x}}_2, \lambda_2)$ are solutions of (2.1) we have

$$\underline{\mathbf{A}}^T \underline{\mathbf{A}} \underline{\mathbf{x}}_1 - \underline{\mathbf{A}}^T \underline{\mathbf{b}} = -\lambda_1 \underline{\mathbf{C}}^T \underline{\mathbf{C}} \underline{\mathbf{x}}_1 + \lambda_1 \underline{\mathbf{C}}^T \underline{\mathbf{d}} \quad (2.3)$$

$$\underline{\mathbf{A}}^T \underline{\mathbf{A}} \underline{\mathbf{x}}_2 - \underline{\mathbf{A}}^T \underline{\mathbf{b}} = -\lambda_2 \underline{\mathbf{C}}^T \underline{\mathbf{C}} \underline{\mathbf{x}}_2 + \lambda_2 \underline{\mathbf{C}}^T \underline{\mathbf{d}} \quad (2.4)$$

$\underline{\mathbf{x}}_2^T$ (2.4) - $\underline{\mathbf{x}}_1^T$ (2.3) gives

$$\begin{aligned} \|\underline{\mathbf{A}}\underline{\mathbf{x}}_2\|^2 - \|\underline{\mathbf{A}}\underline{\mathbf{x}}_1\|^2 - \underline{\mathbf{b}}^T \underline{\mathbf{A}}(\underline{\mathbf{x}}_2 - \underline{\mathbf{x}}_1) \\ = \lambda_1 (\|\underline{\mathbf{C}}\underline{\mathbf{x}}_1\|^2 - \underline{\mathbf{d}}^T \underline{\mathbf{C}} \underline{\mathbf{x}}_1) - \lambda_2 (\|\underline{\mathbf{C}}\underline{\mathbf{x}}_2\|^2 - \underline{\mathbf{d}}^T \underline{\mathbf{C}} \underline{\mathbf{x}}_2) \end{aligned} \quad (2.5)$$

$\underline{\mathbf{x}}_2^T$ (2.3) - $\underline{\mathbf{x}}_1^T$ (2.4) gives

$$-\underline{\mathbf{b}}^T (\underline{\mathbf{A}}\underline{\mathbf{x}}_2 - \underline{\mathbf{x}}_1) = \lambda_1 (-\underline{\mathbf{x}}_2^T \underline{\mathbf{C}}^T \underline{\mathbf{C}} \underline{\mathbf{x}}_1 + \underline{\mathbf{d}}^T \underline{\mathbf{C}} \underline{\mathbf{x}}_1) - \lambda_2 (-\underline{\mathbf{x}}_1^T \underline{\mathbf{C}}^T \underline{\mathbf{C}} \underline{\mathbf{x}}_2 + \underline{\mathbf{d}}^T \underline{\mathbf{C}} \underline{\mathbf{x}}_2). \quad (2.6)$$

Observe that

$$\|\underline{Ax}_2 - \underline{b}\|^2 - \|\underline{Ax}_1 - \underline{b}\|^2 = \|\underline{Ax}_2\|^2 - \|\underline{Ax}_1\|^2 - 2\underline{b}^T \underline{A}(\underline{x}_2 - \underline{x}_1) .$$

So that if we add (2.5) - (2.6) we get

$$\begin{aligned} & \|\underline{Ax}_2 - \underline{b}\|^2 - \|\underline{Ax}_1 - \underline{b}\|^2 \\ &= \lambda_1 \{ \|\underline{Cx}_1\|^2 - \underline{d}^T \underline{Cx}_1 - \underline{x}_1^T \underline{C}^T \underline{Cx}_2 + \underline{d}^T \underline{Cx}_2 \} \\ & \quad - \lambda_2 \{ \|\underline{Cx}_2\|^2 - \underline{d}^T \underline{Cx}_2 - \underline{x}_1^T \underline{C}^T \underline{Cx}_1 + \underline{d}^T \underline{Cx}_1 \} . \end{aligned} \quad (2.7)$$

Now we have

$$\begin{aligned} \|\underline{Cx}_1 - \underline{d}\|^2 &= \|\underline{Cx}_2 - \underline{d}\|^2 = \alpha^2 \\ \Rightarrow \|\underline{Cx}_1\|^2 - 2\underline{d}^T \underline{Cx}_1 + \|\underline{d}\|^2 &= \|\underline{Cx}_2\|^2 - 2\underline{d}^T \underline{Cx}_2 + \|\underline{d}\|^2 \\ \Rightarrow \|\underline{Cx}_1\|^2 - \underline{d}^T \underline{Cx}_1 + \underline{d}^T \underline{Cx}_2 &= \|\underline{Cx}_2\|^2 - \underline{d}^T \underline{Cx}_2 + \underline{d}^T \underline{Cx}_1 . \end{aligned} \quad (2.8)$$

From (2.8) we conclude that the factors of λ_1 and λ_2 in (2.7) are the same. Therefore they also equal their arithmetic mean which is

$$\begin{aligned} & \frac{1}{2} \{ \|\underline{Cx}_1\|^2 - 2\underline{x}_1^T \underline{C}^T \underline{Cx}_2 + \|\underline{Cx}_2\|^2 \} \\ &= \frac{1}{2} \|\underline{C}(\underline{x}_1 - \underline{x}_2)\|^2 . \quad \square \end{aligned}$$

Corollary 2.1. The solution of (PIE) is the solution $\underline{x}(\lambda)$ of the normal equations (2.1) with the largest λ .

Proof. From (2.2) we have that if $\lambda_1 > \lambda_2$, then

$$\|\underline{Ax}_2 - \underline{b}\|^2 > \|\underline{Ax}_1 - \underline{b}\|^2 . \quad \square$$

The next theorem gives a result very similar to Theorem 2.1.

Theorem 2.2. Assume $(\underline{x}_1, \lambda_1)$ and $(\underline{x}_2, \lambda_2)$ are solutions of the normal equations (2.1). Assume that $|\lambda_1| + |\lambda_2| \neq 0$. Then

$$\|\underline{Ax}_2 - \underline{b}\|^2 - \|\underline{Ax}_1 - \underline{b}\|^2 = \frac{\lambda_2^{-A}}{\lambda_2 + \lambda_1} \|\underline{A}(\underline{x}_1 - \underline{x}_2)\|^2 .$$

Proof. From $(\underline{A}^T \underline{A} + \lambda \underline{C}^T \underline{C})\underline{x} = \underline{A}^T \underline{b} + \lambda \underline{C}^T \underline{d}$ we have

$$\lambda_1 \underline{C}^T \underline{Cx}_1 - \lambda_1 \underline{C}^T \underline{d} = -\underline{A}^T \underline{Ax}_1 + \underline{A}^T \underline{b} \quad (2.8)$$

$$\lambda_2 \underline{C}^T \underline{Cx}_2 - \lambda_2 \underline{C}^T \underline{d} = -\underline{A}^T \underline{Ax}_2 + \underline{A}^T \underline{b} . \quad (2.9)$$

$\lambda_1 \underline{x}_1^T$ (2.9) - $\lambda_2 \underline{x}_2^T$ (2.8) gives

$$\lambda_1 \lambda_2 \{ (\underline{x}_2 - \underline{x}_1)^T \underline{C}^T \underline{d} \} = (\lambda_2 - \lambda_1) \underline{x}_1^T \underline{A}^T \underline{Ax}_2 + (\lambda_1 \underline{x}_1 - \lambda_2 \underline{x}_2)^T \underline{A}^T \underline{b} . \quad (2.10)$$

$\lambda_1 \underline{x}_1^T$ (2.9) - $\lambda_2 \underline{x}_1^T$ (2.8) gives

$$\begin{aligned} & \lambda_1 \lambda_2 \{ \|\underline{Cx}_2\|^2 - \|\underline{Cx}_1\|^2 + (\underline{x}_1 - \underline{x}_2)^T \underline{C}^T \underline{d} \} \\ &= \lambda_2 \|\underline{Ax}_1\|^2 - \lambda_1 \|\underline{Ax}_2\|^2 + (\lambda_1 \underline{x}_2 - \lambda_2 \underline{x}_1)^T \underline{A}^T \underline{b} . \end{aligned} \quad (2.11)$$

Observe that

$$\begin{aligned} 0 &= \|\underline{Cx}_2 - \underline{d}\|^2 - \|\underline{Cx}_1 - \underline{d}\|^2 \\ &= \|\underline{Cx}_2\|^2 - \|\underline{Cx}_1\|^2 + 2(\underline{x}_1 - \underline{x}_2)^T \underline{C}^T \underline{d} . \end{aligned}$$

So that if we subtract (2.11) - (2.10) we get

But by Corollary 2.1,

$$\lambda_1 = A > -A = \lambda_2 \Rightarrow \|\underline{Ax}_2 - \underline{b}\|^2 > \|\underline{Ax}_1 - \underline{b}\|^2 .$$

Therefore $\lambda_1 + \lambda_2 \neq 0$ and we may divide in (2.13). $\square$

If we combine both results from Theorem 2.1 and Theorem 2.2 we have

Corollary 2.2. Let $(\underline{x}_1, \lambda_1)$ and $(\underline{x}_2, \lambda_2)$ be two solutions of the normal equations (2.1). If $|\lambda_1| + |\lambda_2| \neq 0$, then

$$-\frac{\lambda_1 + \lambda_2}{2} \|\underline{c}(\underline{x}_1 - \underline{x}_2)\|^2 = \|\underline{A}(\underline{x}_1 - \underline{x}_2)\|^2 . \quad (2.14)$$

From (2.14) we see immediately:

Corollary 2.3. The normal equations (2.1) have at most one solution $(\underline{x}^*, \lambda^*)$ with $A^* > 0$. For every other solution $(\underline{x}, \lambda)$ we have $A < -\lambda^*$.

The next theorem gives conditions for a unique solution of (PLE).

Theorem 2.3. The solution $\underline{x}$ of (PLE) is unique (if it exists) if

$$NS(\underline{A}) \cap NS(\underline{c}) = \{\underline{0}\}$$

and

$$\lambda \neq -\mu_i$$

where $(\underline{x}, \lambda)$ is a solution of the normal equation (2.1) and μ_i is an eigenvalue of the generalized eigenvalue problem

$$\det(\underline{A}^T \underline{A} - \mu \underline{c}^T \underline{c}) = 0 .$$

.

$$\begin{aligned} 0 &= \lambda_2 \|\underline{Ax}_1\|^2 - \lambda_1 \|\underline{Ax}_2\|^2 + (\lambda_1 \underline{x}_2 - \lambda_2 \underline{x}_1 - \lambda_1 \underline{x}_1 + \lambda_2 \underline{x}_2)^T \underline{A}^T \underline{b} \\ &\quad + (\lambda_1 - \lambda_2) \underline{x}_1^T \underline{A}^T \underline{Ax}_2 , \end{aligned}$$

or by rearranging

$$\begin{aligned} \lambda_1 \{ \|\underline{Ax}_2\|^2 - \underline{x}_2^T \underline{A}^T \underline{b} + \underline{x}_1^T \underline{A}^T \underline{b} - \underline{x}_1^T \underline{A}^T \underline{Ax}_2 \} \\ = \lambda_2 \{ \|\underline{Ax}_1\|^2 - \underline{x}_1^T \underline{A}^T \underline{b} + \underline{x}_2^T \underline{A}^T \underline{b} - \underline{x}_1^T \underline{A}^T \underline{Ax}_2 \} . \end{aligned} \quad (2.12)$$

Now the { } on the left hand side of (2.12) is

$$\begin{aligned} \frac{1}{2} \|\underline{Ax}_2\|^2 + \frac{1}{2} \{ \|\underline{Ax}_2\|^2 - 2 \underline{x}_1^T \underline{A}^T \underline{Ax}_2 + \|\underline{Ax}_1\|^2 \} \\ - \frac{1}{2} \|\underline{Ax}_1\|^2 + \underline{x}_1^T \underline{A}^T \underline{b} - \underline{x}_2^T \underline{A}^T \underline{b} + \frac{1}{2} \|\underline{b}\|^2 - \frac{1}{2} \|\underline{b}\|^2 \\ = \frac{1}{2} \{ \|\underline{Ax}_2 - \underline{b}\|^2 - \|\underline{Ax}_1 - \underline{b}\|^2 + \|\underline{A}(\underline{x}_2 - \underline{x}_1)\|^2 \} . \end{aligned}$$

The right hand side of (2.12) simplifies analogously and by rearranging we obtain:

$$(\lambda_1 + \lambda_2) \{ \|\underline{Ax}_2 - \underline{b}\|^2 - \|\underline{Ax}_1 - \underline{b}\|^2 \} = (\lambda_2 - \lambda_1) \|\underline{A}(\underline{x}_2 - \underline{x}_1)\|^2 . \quad (2.13)$$

We now prove by contradiction that $\lambda_1 + \lambda_2 \neq 0$. Assume $(\underline{x}_1, \lambda)$ and $(\underline{x}_2, -\lambda)$ were solutions of the normal equations (2.1) with $A > 0$. Then by (2.13) we would have

$$\begin{aligned} \|\underline{A}(\underline{x}_2 - \underline{x}_1)\|^2 &= 0 \\ \Rightarrow \underline{Ax}_2 &= \underline{Ax}_1 \\ \Rightarrow \|\underline{Ax}_2 - \underline{b}\|^2 &= \|\underline{Ax}_1 - \underline{b}\|^2 . \end{aligned}$$

Proof. The proof is by contradiction. Assume $(\underline{x}_1, \lambda_1)$ and $(\underline{x}_2, \lambda_2)$ are solutions of the normal equations (2.1) which also solve (PLE).

If $\lambda_1 \neq \lambda_2$, then we must have

$$\|\underline{A}\underline{x}_1 - \underline{b}\|^2 = \|\underline{A}\underline{x}_2 - \underline{b}\|^2 = \min \|\underline{A}\underline{x} - \underline{b}\|^2.$$

Theorem 2.1 implies

$$\|\underline{C}(\underline{x}_1 - \underline{x}_2)\|^2 = 0 \Rightarrow \underline{C}(\underline{x}_1 - \underline{x}_2) = \underline{0}. \quad (2.15)$$

While Theorem 2.2 gives

$$\|\underline{A}(\underline{x}_1 - \underline{x}_2)\|^2 = 0 \Rightarrow \underline{A}(\underline{x}_1 - \underline{x}_2) = \underline{0}. \quad (2.16)$$

Now if $\underline{x}_1 \neq \underline{x}_2$ then (2.15) and (2.16) show that $\underline{A}$ and $\underline{C}$ have non-trivially intersecting nullspaces, which is a contradiction.

Therefore we must have $\lambda_1 = \lambda_2 = \lambda$. But in this case we then have

$$(\underline{A}^T \underline{A} + \lambda \underline{C}^T \underline{C})(\underline{x}_1 - \underline{x}_2) = \underline{0}.$$

If $\underline{x}_1 \neq \underline{x}_2$ then $\lambda = -\mu_1$ which is also a contradiction. Therefore we must have $\lambda_1 = \lambda_2$ and $\underline{x}_1 = \underline{x}_2$. $\square$

$\underline{A}$ and $\underline{C}$ have a trivial intersection of their nullspaces if and only if

$$\text{rank} \begin{pmatrix} \underline{A} \\ \underline{C} \end{pmatrix} = n.$$

Therefore a necessary condition for a unique solution is

$$m+p > n$$

which means that we must have "enough" equations to determine $\underline{x}$.

We shall assume in the following (till the end of the paper) that

$\text{NS}(\underline{A}) \cap \text{NS}(\underline{C}) = \{\underline{0}\}$. In Section 2.2 we shall analyze the existence of

solutions of the normal equations. We shall see that (PLE) has a solution if the constraint

$$\|\underline{C}\underline{x} - \underline{d}\| = \alpha^2$$

is feasible, i.e., if

$$\min_{\underline{x}} \|\underline{C}\underline{x} - \underline{d}\|^2 = \|(\underline{C}^+ \underline{C}^+ - \underline{I})\underline{d}\|^2 < \alpha^2.$$

2.2 The Solutions of the Normal Equations.

In this section we shall discuss the solutions of

$$(\underline{A}^T \underline{A} + \lambda \underline{C}^T \underline{C})\underline{x} = \underline{A}^T \underline{b} + \lambda \underline{C}^T \underline{d} \quad (2.17a)$$

$$\|\underline{C}\underline{x} - \underline{d}\|^2 = \alpha^2. \quad (2.17b)$$

We note that we have assumed that $\text{NS}(\underline{A}) \cap \text{NS}(\underline{C}) = \{\underline{0}\}$.

If $\lambda \neq -\mu_1$ where $\mu_1 \geq 0$ is an eigenvalue of the eigenvalue problem $\det(\underline{A}^T \underline{A} - \mu \underline{C}^T \underline{C}) = 0$ then

$$\underline{x}(\lambda) = (\underline{A}^T \underline{A} + \lambda \underline{C}^T \underline{C})^{-1} (\underline{A}^T \underline{b} + \lambda \underline{C}^T \underline{d}). \quad (2.18)$$

If we now choose λ so that the secular equation is satisfied:

$$f(\lambda) := \|\underline{C} \underline{x}(\lambda) - \underline{d}\|^2 = \alpha^2 \quad (2.19)$$

then $(\underline{x}(\lambda), \lambda)$ is a solution of (2.17). We first discuss some properties of the length function f :

Lemma 2.1. If f is defined by (2.19), (2.18), then

- (1) f is a rational function defined on $\mathbb{R} - \{-\mu_1 \mid \det(\underline{A}^T \underline{A} - \mu_1 \underline{C}^T \underline{C}) = 0\}$
- (2) $f(\lambda) \equiv \|\underline{d}\|^2 = \text{const}$ if $\underline{A}^T \underline{b} = \underline{0}$ and $\underline{C}^T \underline{d} = \underline{0}$,
- (3) f has at least one and at most n poles for same $\lambda = -\mu_1$,

$$(4) \quad f(\lambda) > 0 \quad . \quad \lim_{\lambda \rightarrow +\infty} f(\lambda) = \|(\underline{C} \; \underline{C}^+ - \underline{I})\underline{d}\|^2 \quad .$$

$$(5) \quad f'(A) < 0 \text{ for } 0 < \lambda < \infty \quad .$$

proof. We use the generalized singular value decomposition BSVD [45] which gives the **decomposition**

$$\left. \begin{aligned} \underline{U}^T \underline{A} \underline{X} &= \underline{D}_A = \text{diag}(\alpha_1, \dots, \alpha_n), \quad \alpha_i \geq 0 \quad , \\ \underline{V}^T \underline{C} \underline{X} &= \underline{D}_C = \text{diag}(\gamma_1, \dots, \gamma_q), \quad \gamma_i \geq 0 \quad , \quad q = \min(n, p) \end{aligned} \right\} \quad (2.20)$$

where $\underline{U} \ (m \times m)$ $\underline{V} \ (p \times p)$ are orthogonal, $\underline{X} \ (n \times n)$ is **nonsingular**, and $\underline{D}_A, \underline{D}_C$ are **diagonal** matrices with $\gamma_1 \geq \dots \geq \gamma_q$. The decomposition exists only if $m > n$ which we can assume since we can add in

$$\|\underline{Ax} - \underline{b}\|^2 = \left\| \begin{pmatrix} \underline{A} \\ \underline{0} \end{pmatrix} \underline{x} - \begin{pmatrix} \underline{b} \\ \underline{0} \end{pmatrix} \right\|^2 \quad \text{zero rows in } \underline{A} \text{ and zero elements in } \underline{b}$$

without changing the solutions.

If $\gamma_1 \geq \gamma_2 \geq \dots \geq \gamma_r > \gamma_{r+1} = \dots = \gamma_q = 0$, then the eigenvalue of $\det(\underline{A}^T \underline{A} - \mu \underline{C}^T \underline{C}) = 0$ are given by

$$\mu_i = \frac{\alpha_i^2}{\gamma_i^2}, \quad i = 1, \dots, r \quad .$$

Because we are assuming that $\text{NS}(\underline{A}) \cap \text{NS}(\underline{C}) = \{\emptyset\}$ we have

$$\alpha_i > 0 \quad \text{for} \quad i = r+1, \dots, n \quad . \quad [45]$$

By substituting (2.20) into (2.16) we have

$$\underline{X}^{-T} (\underline{D}_A^T \underline{D}_A + \lambda \underline{D}_C^T \underline{D}_C) \underline{X}^{-1} \underline{x} = \underline{X}^{-T} \underline{D}_A^T \underline{U} \underline{U}^T \underline{D}_A + \lambda \underline{X}^{-T} \underline{D}_C^T \underline{V} \underline{V}^T \underline{D}_C$$

and putting

$$\underline{y} := \underline{X}^{-1} \underline{x} \quad , \quad \underline{c} := \underline{U}^T \underline{b} \quad , \quad \underline{e} := \underline{V}^T \underline{d}$$

we get

$$(\underline{D}_A^T \underline{D}_A + \lambda \underline{D}_C^T \underline{D}_C) \underline{y} = \underline{D}_A^T \underline{c} + \lambda \underline{D}_C^T \underline{e} \quad . \quad (2.21)$$

Now

$$f(\lambda) = \|\underline{c} - \underline{x}(\lambda)\|^2 = \|\underline{D}_C \underline{y}(\lambda) - \underline{e}\|^2$$

which is in components:

$$\begin{aligned} f(A) &= \sum_{i=1}^r \left(\gamma_i \frac{\alpha_i c_i + \lambda \gamma_i e_i}{\alpha_i^2 + \lambda \gamma_i^2} - e_i \right)^2 + \sum_{i=r+1}^p e_i^2 \\ f(\lambda) &= \sum_{i=1}^r \left(\alpha_i \frac{\gamma_i c_i - \alpha_i e_i}{c_i^2 + \lambda \gamma_i^2} \right)^2 + \sum_{i=r+1}^p e_i^2 \quad . \end{aligned} \quad (2.22)$$

Obviously $f(A)$ is a rational function in λ with sane poles. If $\alpha_i (\gamma_i c_i - \alpha_i e_i) \neq 0$ for some $1 \leq i \leq r$ then

$$\lambda = -\frac{\alpha_i^2}{\gamma_i^2} = -\mu_i$$

is a pole of f .

If $\underline{A}^T \underline{b} = \underline{0}$ and $\underline{C}^T \underline{d} = 0$ then from the definition of f (2.19), (2.18) we have

$$f(A) = \|\underline{d}\|^2 = \text{const.}$$

From (2.19) we see that $f(A) > 0$.

To prove the limit property in (4) we observe that the solution $\underline{x}(\lambda)$ of (2.17a) converges for $\lambda \rightarrow \infty$ to a solution of $\underline{C}^T \underline{C} \underline{x} = \underline{C}^T \underline{d}$. Though $\underline{x}(\infty)$ may not be $\underline{C}^+ \underline{d}$ it has the same residual. Therefore:

$$\|\underline{\bar{C}} \mathbf{x}(\infty) - \underline{\bar{d}}\|^2 = \|\underline{\bar{C}} \underline{\bar{C}}^T \underline{\bar{d}} - \underline{\bar{d}}\|^2 .$$

This property can of course also be seen from the representation (2.22).

To prove (5) we differentiate (2.19) and (2.17a)

$$\mathbf{f}'(\mathbf{h}) = 2\mathbf{x}'(\lambda)^T \underline{\bar{C}}^T (\underline{\bar{C}} \mathbf{x}(\lambda) - \underline{\bar{d}}) \quad (2.23)$$

$$(\underline{\bar{A}}^T \underline{\bar{A}} + \lambda \underline{\bar{C}}^T \underline{\bar{C}}) \mathbf{x}'(\lambda) = -\underline{\bar{C}}^T (\underline{\bar{C}} \mathbf{x}(\mathbf{h}) - \underline{\bar{d}}) . \quad (2.24)$$

Now since $\lambda > 0$ we can perform the Cholesky decomposition

$$(\underline{\bar{A}}^T \underline{\bar{A}} + \lambda \underline{\bar{C}}^T \underline{\bar{C}}) = \underline{\mathbf{R}}_\lambda^T \underline{\mathbf{R}}_\lambda . \quad (2.25)$$

Using (2.25) and (2.24) in (2.23) gives

$$\mathbf{f}'(\mathbf{A}) = -2\|\underline{\mathbf{R}}_\lambda^T \underline{\bar{C}}^T (\underline{\bar{C}} \mathbf{x}(\lambda) - \underline{\bar{d}})\|^2 < 0 . \quad (2.26)$$

Every solution λ of the secular equation (2.19) with $\mathbf{A} \neq -\mu_i$ defines $\mathbf{x}(\lambda)$ (2.18) and $(\mathbf{x}(\lambda), \lambda)$ is a solution of the normal equations (2.17a, b).

But there might be more solutions for some $\mathbf{A} = -\mu_i$. In this case the matrix of (2.17a) is singular. If a solution should exist, the system must be consistent. This is especially true if $\underline{\bar{A}}^T \underline{\bar{b}} = \underline{\bar{C}}^T \underline{\bar{d}} = 0$ then $\mathbf{x}(\lambda)$ is an eigenvector to $\lambda = -\mu_i$. The next lemma describes the condition for consistency:

Lemma 2.2. Let μ_i be an eigenvalue of $\det(\underline{\bar{A}}^T \underline{\bar{A}} - \mu \underline{\bar{C}}^T \underline{\bar{C}}) = 0$. Then

$$(\underline{\bar{A}}^T \underline{\bar{A}} - \mu \underline{\bar{C}}^T \underline{\bar{C}}) \mathbf{x} = \underline{\bar{A}}^T \underline{\bar{b}} - \mu_i \underline{\bar{C}}^T \underline{\bar{d}} \quad (2.27)$$

is a consistent system if one of the conditions holds:

$$(i) \quad \lim_{\mathbf{A} + -\mu_i} f(\mathbf{A}) = \lim_{\mathbf{A} \rightarrow -\mu_i} \|\underline{\bar{C}} \mathbf{x}(\lambda) - \underline{\bar{d}}\|^2 \text{ exists, i.e., } -\mu_i \text{ is not a pole of } f .$$

$$(ii) \quad \text{Let } J = \{j \mid 1 \leq j \leq k, \frac{\alpha_j^2}{\gamma_j^2} = \mu_i\} \text{ then } \alpha_j(\mathbf{c}_j, \gamma_j - \alpha_j \mathbf{e}_j) = 0$$

for $j \in J$, where all variables are defined in the proof of Lemma 2.1.

Proof. We prove first (ii) using the BSVD to transform (2.27).

$$\underline{\bar{A}} = \underline{\bar{U}} \underline{\bar{D}} \underline{\bar{X}}^{-1}, \quad \underline{\bar{D}} \underline{\bar{A}} = \text{diag}(\alpha_1, \dots, \alpha_n)$$

$$\underline{\bar{C}} = \underline{\bar{V}} \underline{\bar{D}} \underline{\bar{X}}^{-1}, \quad \underline{\bar{D}} \underline{\bar{C}} = \text{diag}(\gamma_1, \dots, \gamma_r, 0, \dots, 0) .$$

$$\text{With } \underline{\bar{y}} := \underline{\bar{X}}^{-1} \underline{\bar{x}}, \quad \underline{\bar{e}} = \underline{\bar{U}}^T \underline{\bar{b}}, \quad \underline{\bar{e}}_j = \underline{\bar{V}}^T \underline{\bar{d}}, \quad (\underline{\bar{A}}^T \underline{\bar{A}} + \mu \underline{\bar{C}}^T \underline{\bar{C}}) \underline{\bar{x}} = \underline{\bar{A}}^T \underline{\bar{b}} + \mu \underline{\bar{C}}^T \underline{\bar{d}}$$

becomes

$$\left. \begin{aligned} (\alpha_k^2 + \mu \gamma_k^2) y_k &= \alpha_k c_k + \lambda \gamma_k e_k, \quad k = 1, \dots, r \\ \alpha_k^2 y_k &= \alpha_k c_k, \quad k = r+1, \dots, n . \end{aligned} \right\} \quad (2.28)$$

Now observe that $r = \text{rank of } \underline{\bar{C}}$ and because $\text{NS}(\underline{\bar{A}}) \cap \text{NS}(\underline{\bar{C}}) = \{0\}$, $\alpha_k \neq 0$, $k = r+1, \dots, n$. If $\lambda = -\mu_i = -\alpha_i^2/\gamma_i^2$, then for every equation j of (2.28) with

$$\alpha_j^2 - \mu_j \gamma_j^2 = 0$$

$$\text{i.e., for all } j \in J = \{j \mid 1 \leq j \leq k, \mu_i = \frac{\alpha_j^2}{\gamma_j^2} = \frac{\alpha_i^2}{\gamma_i^2}\}, \text{ the right hand side must be zero:}$$

$$\begin{aligned} \alpha_j c_j - \frac{\alpha_j^2}{\gamma_j^2} \gamma_j e_j &= \alpha_j c_j - \frac{\alpha_j^2}{\gamma_j^2} e_j = 0 \\ &\Rightarrow \alpha_j (\gamma_j e_j - \gamma_j e_j) = 0 \quad \text{for } j \in J . \end{aligned}$$

We note that the solution of (2.28) is given by

$$\mathbf{y}_k = \begin{cases} (\alpha_k^c \gamma_k^{-\mu_1} \gamma_k^{e_k}) / (\alpha_k^2 - \mu_1 \gamma_k^2), & k \neq J \\ \text{arbitrary}, & k \in J \\ c_k / \alpha_k, & k = r+1, \dots, n \end{cases} \quad (2.29)$$

Furthermore

$$\tilde{\mathbf{y}}_k := \lim_{\lambda \rightarrow -\mu_1} \mathbf{y}_k(\lambda) = \begin{cases} (\alpha_k^c \gamma_k^{-\mu_1} \gamma_k^{e_k}) / (\alpha_k^2 - \mu_1 \gamma_k^2), & k \neq J \\ e_k / \gamma_k, & k \in J \\ c_k / \alpha_k, & k = r+1, \dots, n \end{cases} \quad (2.30)$$

From (2.30) we conclude that

$$\lim_{A \rightarrow D^{-\mu_1}} \mathbf{x}(\lambda) = \lim_{A \rightarrow -\mu_1} \mathbf{x}_- \mathbf{y}(\lambda) = \mathbf{x} \tilde{\mathbf{y}}$$

exists and therefore

$$\lim_{\lambda \rightarrow -\mu_1} \|\underline{\mathbf{C}} \mathbf{x}(\lambda) - \underline{\mathbf{d}}\|^2$$

exists. This implies that f has no pole for $A = -\mu_1$.

Conversely if f has not a pole for $\lambda = -\mu_1$ then also

$$\lim_{A \rightarrow -\mu_1} \mathbf{y}(\lambda)$$

must exist, which means that the system is consistent. $\square$

We can now characterize the solutions of the normal equations precisely:

Theorem 2.4. Let f be the length function defined by (2.19), (2.18).

If $f(A) = \alpha^2$ and $\det(\bar{A} + \lambda \bar{C}^T \bar{C}) \neq 0$ then there exists a unique $\mathbf{x}(\lambda)$ so that $(\mathbf{x}(\lambda), \lambda)$ solves the normal equations. If

$$\det(\bar{A} - \mu_1 \bar{C}^T \bar{C}) = 0 \text{ and } \lim_{\lambda \rightarrow -\mu_1} f(A) \leq \alpha^2 \quad (2.31)$$

then there exist a $\mathbf{x}(-\mu_1)$ so that $(\mathbf{x}(-\mu_1), -\mu_1)$ solves the normal equations, but $\mathbf{x}(-\mu_1)$ is only unique if $\lim_{A \rightarrow -\mu_1} f(A) = \alpha^2$.

Proof. We have only to prove (2.31). We use the terminology defined in the proof of Lemma 2.2. For $A = -\mu_1$ the general solution $\mathbf{y}$ of the transformed normal equations (2.28) is given by (2.29). We have to see if we can satisfy the equation

$$\|\underline{\mathbf{D}} \mathbf{y} - \underline{\mathbf{e}}\| = \alpha^2 \quad (2.32)$$

with this solution $\mathbf{y}$. In components (2.32) is

$$\sum_{k \neq J} (\gamma_k^y - e_k)^2 + \sum_{k \in J} (\gamma_k^y - e_k)^2 + \sum_{k=r+1}^n e_k^2 = \alpha^2. \quad (2.33)$$

We can choose the components of $\mathbf{y}_k$, $k \in J$ arbitrarily. From (2.30) we see that

$$\begin{aligned} \lim_{\lambda \rightarrow -\mu_1} f(A) &= \lim_{\lambda \rightarrow -\mu_1} \|\underline{\mathbf{D}} \mathbf{y}(\lambda) - \underline{\mathbf{e}}\|^2 = \|\underline{\mathbf{D}} \tilde{\mathbf{y}} - \underline{\mathbf{e}}\|^2 \\ \lim_{\lambda \rightarrow -\mu_1} f(A) &= \sum_{k \neq J} (\gamma_k^y - e_k)^2 + \sum_{k=r+1}^n e_k^2. \end{aligned} \quad (2.34)$$

Therefore $\|\underline{\mathbf{D}} \tilde{\mathbf{y}} - \underline{\mathbf{e}}\|^2$ is minimized for $\mathbf{y} = \tilde{\mathbf{y}}$.

From (2.34) we see that if

$$\lim_{A \rightarrow -\mu_1} f(A) > \alpha^2, \dots$$

we can not determine a solution $\underline{y}(\lambda)$ that satisfies (2.32). If

$$\lim_{h \rightarrow -\mu_1} f(A) = \alpha^2$$

then the unique solution is $\tilde{y}$ (2.30). If

$$\lim_{A \rightarrow -\mu_1} f(A) < \alpha^2$$

then we can choose y_k , $k \in J$ so that

$$\sum_{k \in J} (y_k y_k - e_k)^2 = \alpha^2 - \lim_{A \rightarrow -\mu_1} f(A) \quad (2.35)$$

and $\underline{y}(-\mu_1)$ is not unique. $\square$

We remark that the components of $\underline{y} : y_k$, $k \in J$ are the (possibly) non-zero components of the eigenvector $\underline{y}_{-\mu_1}$ of

$$(\underline{D}_A^T \underline{D}_A - \mu_1 \underline{D}_C^T \underline{D}_C) \underline{y}_{-\mu_1} = \underline{0}.$$

Therefore a solution for the case where $\lim_{A \rightarrow -\mu_1} f(A) < \alpha^2$ can also be

written in the following way. Take any eigenvector $\underline{y}_{-\mu_1}$ belonging to μ_1

and determine ρ such that

$$\|\underline{D}_C(\tilde{y} + \rho \underline{y}_{-\mu_1}) - \underline{d}\| = \alpha^2.$$

Then $(\tilde{y} + \rho \underline{y}_{-\mu_1})$ is the solution of the normal equations. /

2.3 The Solution for the Equality Constraint.

We now consider problem (PLE)

$$\begin{aligned} \|\underline{Ax} - \underline{b}\|^2 &= \min \\ \text{subject to } \|\underline{Cx} - \underline{d}\| &= \alpha^2. \end{aligned}$$

We continue to assume that the nullspaces of $\underline{A}$ and $\underline{C}$ intersect trivially. To find the solution we first have to compute the rightmost solution A of the secular equation. If $A^* > 0$ then $\underline{x}(\lambda^*)$ is the unique solution of (PLE). If $A^* < 0$ then we have to compute the smallest eigenvalue μ_r of the generalized eigenvalue problem

$$\det(\underline{A}^T \underline{A} - \mu \underline{C}^T \underline{C}) = 0$$

and an eigenvector $\underline{x}_r$. If $A^* > -\mu_r$ then again by Theorem 2.1 $\underline{x}(\lambda^*)$ is the unique solution. However if $A^* \leq -\mu_r$ then a solution has the form

$$\underline{x}(\rho) = \lim_{\lambda \rightarrow -\mu_r} \underline{x}(\lambda) + \rho \underline{x}_r$$

where we have to determine ρ such that

$$\|\underline{C} \underline{x}(\rho) - \underline{d}\|^2 = \alpha^2.$$

If $\underline{A}^T \underline{b} = \underline{C}^T \underline{d} = \underline{0}$ then $f(A) \equiv \|\underline{d}\|^2$ and λ^* does not exist. Then $(\rho \cdot \underline{x}_r, -\mu_r)$ solves the problem if $\|\underline{d}\|^2 \leq \alpha^2$ where ρ has to be chosen so that $\|\rho \underline{C} \underline{x}_r - \underline{d}\| = \alpha^2$.

2.4 The Solution for the Inequality Constraint.

We are now ready to solve (Pl):

$$\|\underline{Ax} - \underline{b}\|^2 = \min \quad (2.36)$$

subject to

$$\|\underline{Cx} - \underline{d}\|^2 \leq \alpha^2 \quad (2.37)$$

The weaker condition (2.37) simplifies the problem. Let

$\mathbf{M} = \{\underline{x} \mid \|\underline{Ax} - \underline{b}\| = \min\}$ denote the set of solutions of the unconstrained problem (2.36). If for **some** $\mathbf{x} \in \mathbf{M}$ we have $\|\underline{Cx} - \underline{d}\|^2 \leq \alpha^2$ then this $\mathbf{x}$ is a solution.

If $\mathbf{M} \neq \{\underline{A}^T \underline{b}\}$ then $\mu_1 = 0$ is an eigenvalue of $\det(\underline{A}^T \underline{A} - \mu \underline{C}^T \underline{C}) = 0$. Since 0 is never a pole of f (the **normal** equations (2.27) are consistent for $\lambda = 0$) we can define the number

$$\alpha_{\max}^2 := \lim_{\lambda \rightarrow 0} f(\lambda).$$

Let $\alpha_{\min}^2 := \|(\underline{C}^T \underline{C} - \underline{I})\underline{d}\|^2 = \lim_{\lambda \rightarrow \infty} f(\lambda)$. Then (Pl) has no solution if

$\alpha < \alpha_{\min}$. If $\alpha \geq \alpha_{\max}$ then

$$\underline{x}_0 = \lim_{\lambda \rightarrow 0} (\lambda \underline{A}^T \underline{A} + \lambda \underline{C}^T \underline{C})^{-1} (\lambda \underline{A}^T \underline{b} + \lambda \underline{C}^T \underline{d})$$

is the solution. Observe that $\underline{x}_0$ is in general not $\underline{A}^T \underline{b}$ [40]. Similarly if $\alpha = \alpha_{\min}$ then we have to determine among the solutions $\mathbf{x}$ that minimize $\|\underline{Cx} - \underline{d}\|^2$ the solution that minimizes $\|\underline{Ax} - \underline{b}\|^2$ which is

$$\underline{x}_{-\infty} = \lim_{\mu \rightarrow 0} (\mu \underline{A}^T \underline{A} + \underline{C}^T \underline{C})^{-1} (\mu \underline{A}^T \underline{b} + \underline{C}^T \underline{d}).$$

In general we will have

$$\alpha_{\min} < \alpha < \alpha_{\max}.$$

This means that we have to determine the unique positive solution $\mathbf{A}^*$ of the secular equation

$$f(\lambda) = \alpha^2.$$

And $\mathbf{x}(\lambda) = (\underline{A}^T \underline{A} + \lambda \underline{C}^T \underline{C})^{-1} (\underline{A}^T \underline{b} + \lambda \underline{C}^T \underline{d})$ will solve (Pl).

The extreme cases $\alpha = \alpha_{\max}$ and $\alpha = \alpha_{\min}$ correspond to the two problems

$$\left. \begin{array}{l} \min \|\underline{Cx} - \underline{d}\| \\ \text{subject to} \\ \|\underline{Ax} - \underline{b}\| = \min \end{array} \right\} \quad \text{if } \alpha = \alpha_{\max}$$

$$\left. \begin{array}{l} \min \|\underline{Ax} - \underline{b}\|^2 \\ \text{subject to} \\ \|\underline{Cx} - \underline{d}\| = \min \end{array} \right\} \quad \text{if } \alpha = \alpha_{\min}$$

Both are of course only nontrivial if $\underline{A}$ respectively $\underline{C}$ is rank deficient.

To compute the solution in the case $\alpha_{\min} < \alpha < \alpha_{\max}$ we proceed iteratively. The following algorithm shows the basic idea.

```
(a) start with  $\lambda > 0$ 
(b) while not converged do
    begin
        solve the least square problem
        
$$\left( \frac{\underline{A}}{\sqrt{\lambda} \underline{C}} \right) \underline{x}_{\lambda} \approx \begin{pmatrix} \underline{b} \\ \sqrt{\lambda} \underline{d} \end{pmatrix}$$

        correct  $\underline{A}$  to solve  $f(\underline{A}) = \alpha^2$  where
        
$$f(\lambda) = \|\underline{Cx}_{\lambda} - \underline{d}\|^2.$$

    end;
```

At every step of the iteration we have to solve a least squares problem. If we use an orthogonalization procedure [48] it is more expensive than solving the normal equations by the Cholesky decomposition [38,28,33]. But it is well known [26] that it is numerically preferable to solve the least square problem by orthogonal transformations.

We shall show in Section 5 how to reduce the amount of work using a structure of the matrices A and C or by using an orthogonal decomposition of them.

In the next two sections we shall consider the special cases $\underline{A} = \underline{I}$ and $\underline{C} = \underline{I}$. It is interesting that we can formulate dual equations for these special cases. Furthermore Eldén [8] has showed how to reduce the general problem to a problem with $\underline{C} = \underline{I}$.

3. The Relaxed Least Squares Problem.

We consider in this section the special case of a least squares problem with a quadratic constraint where $\underline{C} = \underline{I}$:

$$\begin{aligned} \|\underline{Ax} - \underline{b}\| &= \min \\ \text{subject to} \quad \|\underline{x}\| &= \alpha. \end{aligned} \quad (\text{P2E})$$

The problem with the inequality constraint $\|\underline{x}\| \leq \alpha$ (P2) was called by Rutishauser [35] the relaxed least squares problem. Problem (P2E) can be interpreted to find the minimum of a quadratic form $\|\underline{Ax} - \underline{b}\|^2$ on the sphere $\|\underline{x}\|^2 = \alpha^2$, which is a special case of finding the stationary values of a quadratic form on a sphere. This problem has been analyzed by Forsythe and Golub [10] in 1965. The two authors proved that if (λ_1, λ_1) and (λ_2, λ_2) are two solutions of the normal equations then

$$\lambda_1 > \lambda_2 \Rightarrow \|\underline{Ax}_2 - \underline{b}\|^2 > \|\underline{Ax}_1 - \underline{b}\|^2.$$

Their proof is rather complicated. Kahan [23] gave an elementary proof and stated the theorem

$$\|\underline{Ax}_2 - \underline{b}\|^2 - \|\underline{Ax}_1 - \underline{b}\|^2 = \frac{\lambda_1 - \lambda_2}{2} \|\underline{x}_2 - \underline{x}_1\| \quad (3.1)$$

which is Theorem 2.1 for $\underline{C} = \underline{I}$. Unfortunately this theorem was never published. In 1972 Spjøtvoll [39] gave a simpler proof than Forsythe - Golub but did not quite obtain Kahan's result. He showed that

$$\|\underline{Ax}_2 - \underline{b}\|^2 - \|\underline{Ax}_1 - \underline{b}\|^2 = (\lambda_1 - \lambda_2)(\alpha^2 - \underline{x}_1^T \underline{x}_2)$$

but did not see that $\alpha^2 = \frac{1}{2} (\|\underline{x}_1\|^2 + \|\underline{x}_2\|^2)$.

3.1 Results from the General Theory.

Considering the theory of Section 2 and putting $\underline{\mathbf{c}} = \underline{\mathbf{I}}$ and $\underline{\mathbf{d}} = \underline{\mathbf{0}}$ we get the normal equations

$$(\underline{\mathbf{A}}^T \underline{\mathbf{A}} + \underline{\mathbf{A}}_- \underline{\mathbf{I}}) \underline{\mathbf{x}} = \underline{\mathbf{A}}^T \underline{\mathbf{b}} \quad (3.2)$$

$$\|\underline{\mathbf{x}}\|^2 = \alpha^2. \quad (3.3)$$

Theorem 2.1 gives us equation (3.1). This means that the solution of

(P2E) is the solution $(\underline{\mathbf{x}}, \lambda)$ of (3.2), (3.3) with the largest λ .

Since $\text{NS}(\underline{\mathbf{A}}) \cap \text{NS}(\underline{\mathbf{I}}) = \{\underline{\mathbf{0}}\}$ we have from Theorem 2.3 that the solution of (P2E) is unique if $\lambda \neq -\sigma_1^2$ (where σ_1 is a singular value of $\underline{\mathbf{A}}$).

The length function f can be written

$$f(\lambda) = \|(\underline{\mathbf{A}}^T \underline{\mathbf{A}} + \underline{\mathbf{A}}_- \underline{\mathbf{I}})^{-1} \underline{\mathbf{A}}^T \underline{\mathbf{b}}\|^2. \quad (3.4)$$

Using the singular value decomposition of $\underline{\mathbf{A}}$ [14]

$$\underline{\mathbf{A}} = \underline{\mathbf{U}} \underline{\Sigma} \underline{\mathbf{V}}^T$$

where

$\underline{\mathbf{U}}, \underline{\mathbf{V}}$ orthogonal

$$\underline{\Sigma} = \text{diag}(\sigma_1, \dots, \sigma_p), \quad p = \min(m, n).$$

The right hand side of (3.4) becomes

$$f(\lambda) = \|(\underline{\Sigma}^+ \underline{\Sigma} + \lambda \underline{\mathbf{I}})^{-1} \underline{\Sigma}^T \underline{\mathbf{c}}\| \quad \text{with } \underline{\mathbf{c}} := \underline{\mathbf{U}}^T \underline{\mathbf{b}}.$$

In components this is

$$f(\lambda) = \sum_{i=1}^p \left(\frac{\sigma_i c_i}{\sigma_i^2 + \lambda} \right)^2. \quad (3.5)$$

If $\underline{\mathbf{A}}^T \underline{\mathbf{b}} = \underline{\mathbf{0}}$ then $c = \underline{\mathbf{0}}$ and $f(\lambda) = 0$. Theorem 2.4 and the fact that $\lim_{\lambda \rightarrow -\infty} f(\lambda) = 0$ implies that for every α we have a solution of (P2E) which is not unique if $\lambda = -\sigma_1^2$ and $\lim_{\lambda \rightarrow -\sigma_1^2} f(\lambda) < \alpha^2$.

This was already stated by Forsythe and Golub in 1965 [10].

For the problem with the inequality constraint $\|\underline{\mathbf{x}}\|^2 \leq \alpha^2$ (P2)

we have either

$$\|\underline{\mathbf{A}}^+ \underline{\mathbf{b}}\|^2 \leq \alpha^2$$

then the solution is $\underline{\mathbf{x}} = \underline{\mathbf{A}}^+ \underline{\mathbf{b}}$ or if $\|\underline{\mathbf{A}}^+ \underline{\mathbf{b}}\|^2 > \alpha^2$ then we have to determine the unique solution λ^* of the secular equation

$$f(\lambda) = \alpha^2 \quad \text{in } (0, \infty).$$

The solution is in that case $\mathbf{x}(\lambda^*)$. We note that we can reduce the length of $\underline{\mathbf{x}}$ arbitrarily by choosing α small.

3.2 The Dual Normal Equations.

Theorem 3.1.

(i) Let $(\underline{\mathbf{x}}, \lambda)$ be a solution of the primal normal equation

$$(\underline{\mathbf{A}}^T \underline{\mathbf{A}} + \underline{\mathbf{A}}_- \underline{\mathbf{I}}) \underline{\mathbf{x}} = \underline{\mathbf{A}}^T \underline{\mathbf{b}} \quad (3.6)$$

$$\|\underline{\mathbf{x}}\|^2 = \alpha^2 \quad (3.7)$$

with $\lambda \neq 0$. Then (z, λ) with

$$z = \frac{1}{\lambda} (\underline{\mathbf{A}} \underline{\mathbf{x}} - \underline{\mathbf{b}}) \quad (3.8)$$

is a solution of the dual normal equations

...

$$(\underline{\underline{AA}}^T + \lambda \underline{\underline{I}}) \underline{\underline{z}} = -\underline{\underline{b}} \quad (3.9)$$

$$\|\underline{\underline{A}}^T \underline{\underline{z}}\| = \alpha^2 . \quad (3.10)$$

(ii) Let $(\underline{\underline{z}}, \lambda)$ be a solution of the dual normal equations (3.8),

(3.9) . Then $(\underline{\underline{x}}, \lambda)$ with

$$\underline{\underline{x}} = -\underline{\underline{A}}^T \underline{\underline{z}} \quad (3.11)$$

is a solution of the prime normal equations (3.6), (3.7) .

Proof.

$$\begin{aligned} (i) \quad (\underline{\underline{AA}}^T + \lambda \underline{\underline{I}}) \frac{1}{\lambda} (\underline{\underline{Ax}} - \underline{\underline{b}}) &= \frac{1}{\lambda} \underline{\underline{A}} (\underline{\underline{A}}^T \underline{\underline{A}} + \lambda \underline{\underline{I}}) \underline{\underline{x}} - \underbrace{\frac{1}{\lambda} \underline{\underline{AA}}^T \underline{\underline{b}} - \underline{\underline{b}}}_{\underline{\underline{A}}^T \underline{\underline{b}}} \\ &= -\underline{\underline{b}} . \end{aligned}$$

Furthermore

$$\|\underline{\underline{A}}^T \underline{\underline{z}}\|^2 = \|\underline{\underline{A}}^T \frac{1}{\lambda} (\underline{\underline{Ax}} - \underline{\underline{b}})\|^2 = \|\underline{\underline{x}}\|^2 = \alpha^2 .$$

$$(ii) \quad (\underline{\underline{A}}^T \underline{\underline{A}} + \lambda \underline{\underline{I}}) (-\underline{\underline{A}}^T \underline{\underline{z}}) = -\underline{\underline{A}}^T (\underbrace{\underline{\underline{AA}}^T + \lambda \underline{\underline{I}}}_{-\underline{\underline{b}}} \underline{\underline{z}}) .$$

And

$$\|-\underline{\underline{A}}^T \underline{\underline{z}}\|^2 = \|\underline{\underline{x}}\|^2 = \alpha^2 . \quad \square$$

Theorem 3.1 shows that we may simplify computations to solve (P2) or (P2E) . Since $\underline{\underline{A}}$ is an $(m \times n)$ matrix one of the linear systems (3.6) or (3.8) will be smaller and it may be more economical to iterate with the smaller .

But Theorem 3.1 also gives theoretical equations .

Corollary 3.1. Let $\underline{\underline{A}}$ be an $(m \times n)$ matrix and $\underline{\underline{A}} \neq -\sigma_1^2$ where σ_1

is a singular value of $\underline{\underline{A}}$. Then

$$(\underline{\underline{A}}^T \underline{\underline{A}} + \lambda \underline{\underline{I}})^{-1} \underline{\underline{A}}^T = \underline{\underline{A}}^T (\underline{\underline{AA}}^T + \lambda \underline{\underline{I}})^{-1} \quad (3.12)$$

$$(\underline{\underline{AA}}^T + \lambda \underline{\underline{I}})^{-1} = \frac{1}{\lambda} (\underline{\underline{I}} - \underline{\underline{A}} (\underline{\underline{A}}^T \underline{\underline{A}} + \lambda \underline{\underline{I}})^{-1} \underline{\underline{A}}^T) . \quad (3.13)$$

Proof. From Theorem 3.1 (3.12) follows by equating $\underline{\underline{x}}$ and (3.13) by equating $\underline{\underline{z}}$ for the dual and primal equations. $\square$

As an application we consider (3.13) for $n = 1$ (i.e., $\underline{\underline{A}}$ is a vector) and $\lambda = 1$:

$$(\underline{\underline{I}} + \underline{\underline{aa}}^T)^{-1} = \underline{\underline{I}} - \frac{1}{(\underline{\underline{a}}^T \underline{\underline{a}} + 1)} \underline{\underline{aa}}^T . \quad (3.14)$$

We observe that Corollary 3.1 is a special case of the **Shermann - Morrison-Woodbury** formula:

Let $\underline{\underline{A}}$ be $(n \times n)$, $\underline{\underline{U}}$, $\underline{\underline{V}}$ $(n \times p)$ matrices, then every solution $\underline{\underline{x}}$ of

$$(\underline{\underline{A}} + \underline{\underline{U}} \underline{\underline{V}}^T) \underline{\underline{x}} = \underline{\underline{b}} , \quad (3.15)$$

is also the solution of the augmented system

$$\underline{\underline{Ax}} + \underline{\underline{Uy}} = \underline{\underline{b}} \quad (3.16a)$$

$$\underline{\underline{V}}^T \underline{\underline{x}} - \underline{\underline{y}} = \underline{\underline{0}} . \quad (3.16b)$$

Now assume that $\underline{\underline{A}}$ is nonsingular. Then we have from (3.16a)

$$\underline{\underline{x}} = \underline{\underline{A}}^{-1} \underline{\underline{b}} - \underline{\underline{A}}^{-1} \underline{\underline{Uy}} .$$

Introducing this in (3.16b) we get

3.3 Eldén's Transformation.

Eldén [8] has shown how to transform a problem (P1) to a problem (P2) using orthogonal transformations. To illustrate his idea consider the least square representation of (P1)

$$\begin{pmatrix} \underline{\underline{A}} \\ \sqrt{\lambda} \underline{\underline{C}} \end{pmatrix} \underline{\underline{x}} \approx \begin{pmatrix} \underline{\underline{b}} \\ \sqrt{\lambda} \underline{\underline{d}} \end{pmatrix} \quad (3.17)$$

$$\|\underline{\underline{C}}\underline{\underline{x}} - \underline{\underline{d}}\|^2 = \alpha^2 \quad (3.18)$$

We may assume that the rank of $C = p < n$. (C is a $(p \times n)$ matrix.) If not then we can perform a QR decomposition with column pivoting

$$\underline{\underline{C}} = \underline{\underline{Q}} \begin{pmatrix} \underline{\underline{R}} \\ 0 \end{pmatrix}^T, \quad \underline{\underline{R}} = \begin{pmatrix} \square \\ 0 \end{pmatrix}.$$

Observe that $\|\underline{\underline{C}}\underline{\underline{x}} - \underline{\underline{d}}\| = \|\underline{\underline{R}}^T \underline{\underline{P}}^T \underline{\underline{x}} - \underline{\underline{Q}}^T \underline{\underline{d}}\|$. Therefore we can replace C by $\underline{\underline{R}}$, $\underline{\underline{d}}$ by $\underline{\underline{Q}}^T \underline{\underline{d}}$, $\underline{\underline{A}}$ by $\underline{\underline{A}} \underline{\underline{P}}$, and $\underline{\underline{x}}$ by $\underline{\underline{y}} = \underline{\underline{P}}^T \underline{\underline{x}}$ in (3.17), (3.18). The problem for $\underline{\underline{y}}$ has now a matrix $\underline{\underline{Q}}$ ($p \times n$) with rank of $\underline{\underline{Q}} = p$.

After it is solved we obtain $\underline{\underline{x}} = \underline{\underline{P}} \underline{\underline{y}}$. The same preprocessing we can apply to $\underline{\underline{A}}$ and $\underline{\underline{b}}$ and so assume that $\text{rank}(\underline{\underline{A}}) = m \leq n$.

If $p = n$ then we can make the change of variables

$$\underline{\underline{x}}' := \underline{\underline{C}} \underline{\underline{x}}, \quad \underline{\underline{A}}' := \underline{\underline{A}} \underline{\underline{C}}^{-1}$$

which transforms the problem to

$$\begin{pmatrix} \underline{\underline{A}}' \\ \sqrt{\lambda} \underline{\underline{I}} \end{pmatrix} \underline{\underline{x}}' \approx \begin{pmatrix} \underline{\underline{b}} \\ \sqrt{\lambda} \underline{\underline{d}} \end{pmatrix}$$

$$\|\underline{\underline{x}}' - \underline{\underline{d}}\|^2 = \alpha^2.$$

$$\begin{aligned} \underline{\underline{y}} &= (\underline{\underline{I}} + \underline{\underline{V}}^T \underline{\underline{A}}^{-1} \underline{\underline{U}})^{-1} \underline{\underline{V}}^T \underline{\underline{A}}^{-1} \underline{\underline{b}} \\ \Rightarrow \underline{\underline{x}} &= \underline{\underline{A}}^{-1} \underline{\underline{b}} - \underline{\underline{A}}^{-1} \underline{\underline{U}} (\underline{\underline{I}} + \underline{\underline{V}}^T \underline{\underline{A}}^{-1} \underline{\underline{U}})^{-1} \underline{\underline{V}}^T \underline{\underline{A}}^{-1} \underline{\underline{b}} \end{aligned}$$

But from (3.15) and (3.16b) we have also

$$\begin{aligned} \underline{\underline{x}} &= (\underline{\underline{A}} + \underline{\underline{U}} \underline{\underline{V}}^T)^{-1} \underline{\underline{b}} \\ \underline{\underline{y}} &= \underline{\underline{V}}^T (\underline{\underline{A}} + \underline{\underline{U}} \underline{\underline{V}}^T)^{-1} \underline{\underline{b}} \end{aligned}$$

Equating both expressions for $\underline{\underline{x}}$ and $\underline{\underline{y}}$ we get the matrix equations

$$\begin{aligned} \underline{\underline{V}}^T (\underline{\underline{A}} + \underline{\underline{U}} \underline{\underline{V}}^T)^{-1} &= (\underline{\underline{I}} + \underline{\underline{V}}^T \underline{\underline{A}}^{-1} \underline{\underline{U}})^{-1} \underline{\underline{V}}^T \underline{\underline{A}}^{-1} \\ (\underline{\underline{A}} + \underline{\underline{U}} \underline{\underline{V}}^T)^{-1} &= \underline{\underline{A}}^{-1} (\underline{\underline{I}} - \underline{\underline{U}} (\underline{\underline{I}} + \underline{\underline{V}}^T \underline{\underline{A}}^{-1} \underline{\underline{U}})^{-1} \underline{\underline{V}}^T \underline{\underline{A}}^{-1}) \end{aligned}$$

(Shermann - Morrison - Woodbury formulas)

we note that for $\underline{\underline{A}} := \lambda \underline{\underline{I}}$ and $\underline{\underline{U}} = \underline{\underline{V}} := \underline{\underline{A}}$ we get the equations (3.12) and (3.13).

We finally remark that if $\underline{\underline{A}} > 0$ the dual equations are the normal equations of the least squares problem

$$\begin{pmatrix} \underline{\underline{A}}' \\ \sqrt{\lambda} \underline{\underline{I}} \end{pmatrix} \underline{\underline{z}} \approx \begin{pmatrix} 0 \\ -\frac{1}{\sqrt{\lambda}} \underline{\underline{b}} \end{pmatrix}.$$

The dual equations have no solution if $\underline{\underline{A}} = 0$. We see this clearly because of the factor $-\frac{1}{\sqrt{\lambda}}$, but also from (3.9) since it is clear that for $\lambda = 0$ a solution $\underline{\underline{z}}$ of (3.9) may not exist for arbitrary $\underline{\underline{b}}$ and $m > n$.

We assume now $p < n$ and C has rank p . Then we make a QR decomposition of C^T :

$$C^T = \begin{pmatrix} V_1 & V_2 \end{pmatrix} \begin{pmatrix} R \\ 0 \end{pmatrix} \quad \begin{matrix} \text{p} \\ \text{n-p} \end{matrix} \quad (3.19)$$

with nonsingular triangular matrix R .

We change variables:

$$\underline{X} =: \underline{V}_1 \underline{Y}_1 + \underline{V}_2 \underline{Y}_2$$

and (3.17) becomes:

$$\begin{pmatrix} \underline{A} \underline{V}_1 & \underline{A} \underline{V}_2 \\ \sqrt{\lambda} \underline{R}^T & \underline{0} \end{pmatrix} \begin{pmatrix} \underline{Y}_1 \\ \underline{Y}_2 \end{pmatrix} \approx \begin{pmatrix} \underline{b} \\ \sqrt{\lambda} \underline{V}_2^T \underline{d} \end{pmatrix}. \quad (3.21)$$

Whereas (3.18) is now

$$\|\underline{R}^T \underline{V}_1 \underline{Y}_1 - \underline{d}\|^2 = \alpha^2. \quad (3.22)$$

Now we perform another QR decomposition

$$\underline{A} \underline{V}_2 = (\underline{Q}_1, \underline{Q}_2) \begin{pmatrix} \underline{U} \\ \underline{0} \end{pmatrix} \quad \begin{matrix} \text{p} \\ \text{n} \end{matrix} \quad (3.23)$$

Since we assumed that A has rank m and $\underline{V}_2$ is orthogonal, $\underline{U}$ is a non-singular upper triangular matrix. Now (3.21) is the same problem

$$\begin{pmatrix} \underline{Q}_1^T \underline{A} \underline{V}_1 & \underline{U} \\ \underline{Q}_2^T \underline{A} \underline{V}_1 & \underline{0} \\ \sqrt{\lambda} \underline{R}^T & \underline{0} \end{pmatrix} \begin{pmatrix} \underline{Y}_1 \\ \underline{Y}_2 \end{pmatrix} \approx \begin{pmatrix} \underline{Q}_1^T \underline{b} \\ \underline{Q}_2^T \underline{b} \\ \sqrt{\lambda} \underline{V}_2^T \underline{d} \end{pmatrix}. \quad (3.25)$$

as:

Now for every $\underline{Y}_1$ we can determine $\underline{Y}_2$ such that the first p -n equations are exactly satisfied. Therefore the problem splits as follows

$$\begin{pmatrix} \underline{Q}_2^T \underline{A} \underline{V}_1 \\ \sqrt{\lambda} \underline{R}^T \end{pmatrix} \underline{Y}_1 \approx \begin{pmatrix} \underline{Q}_2^T \underline{b} \\ \sqrt{\lambda} \underline{V}_2^T \underline{d} \end{pmatrix} \quad (3.25)$$

with (3.22) as constraint. The final change of variables

$$\underline{Y}_1 := \underline{R}^T \underline{Y}_1$$

gives the desired problem (P2).

3.4 Rutishauser's Relaxed and Doubly Relaxed Least Squares Problem.

In [35] Rutishauser remarked that if we minimize

$$Q(\underline{x}) = \|\underline{A}\underline{x} - \underline{b}\|^2 \quad (3.26)$$

for a ill-conditioned matrix $\underline{A}$ (i.e., $\kappa = \|\underline{A}\| \|\underline{A}^+\| = \sigma_1 / \sigma_p \gg 1$, where $\sigma_1 > \dots \geq \sigma_p > \sigma_{p+1} = \dots = \sigma_n = 0$ are the singular values of A), that then the exact solution $\hat{\underline{x}} = \underline{A}^+ \underline{b}$ may be not satisfactory because $\|\hat{\underline{x}}\| \gg 1$ and evaluation of $\hat{\underline{A}} \hat{\underline{x}}$ will be affected by cancellation. Replacing (3.26) by

$$Q_\epsilon(\underline{x}) = \|\underline{A}\underline{x} - \underline{b}\|^2 + \epsilon \|\underline{x}\| \quad (3.27)$$

will give a shorter solution $\underline{x}_\epsilon$ which however does not minimize (3.27) any more. The process of relaxing can be described as follows: The solution of (3.26) is the solution of the least squares problem

$$\underline{A}\underline{x} \approx \underline{b}. \quad (3.28)$$

The relaxed solution $\underline{\mathbf{x}}_{-\epsilon}$ of (3.27) is obtained by choosing some $\epsilon > 0$ and solving

$$\begin{pmatrix} \underline{\mathbf{A}} \\ \epsilon \underline{\mathbf{I}} \end{pmatrix} \underline{\mathbf{x}} \approx \begin{pmatrix} \underline{\mathbf{b}} \\ 0 \end{pmatrix} . \quad (3.29)$$

The connection to problem (P2) is as follows: Instead of presenting a bound for the length of $\mathbf{x}$: $\|\underline{\mathbf{x}}\|^2 \leq \alpha^2$ and solve the equation

$$\mathbf{r}(\lambda) = \alpha^2$$

to determine $\lambda > 0$, we simply set $\lambda = \epsilon^2 > 0$ and solve for $\mathbf{x}$.

Rutishauser [35] defines the double relaxed solution $\underline{\mathbf{x}}_{\mathbf{d}\epsilon}$ to be the solution of

$$(\underline{\mathbf{A}}^T \underline{\mathbf{A}} + \epsilon^2 \underline{\mathbf{I}} + \epsilon(\underline{\mathbf{A}}^T \underline{\mathbf{A}} + \epsilon^2 \underline{\mathbf{I}})^{-1}) \underline{\mathbf{x}} = \underline{\mathbf{A}}^T \underline{\mathbf{b}} . \quad (3.30)$$

Observe that $\underline{\mathbf{x}}_{\mathbf{d}\epsilon}$ is obtained by relaxing the **normal** equations of the relaxed solution:

$$\begin{pmatrix} \underline{\mathbf{A}}^T \underline{\mathbf{A}} + \epsilon^2 \underline{\mathbf{I}} \\ \epsilon \underline{\mathbf{I}} \end{pmatrix} \underline{\mathbf{x}} \approx \begin{pmatrix} \underline{\mathbf{A}}^T \underline{\mathbf{b}} \\ 0 \end{pmatrix} . \quad (3.31)$$

If we put $\underline{\mathbf{B}} := \underline{\mathbf{A}}^T \underline{\mathbf{A}} + \epsilon^2 \underline{\mathbf{I}}$ then the normal equations of (3.31) are

$$(\underline{\mathbf{B}}^T \underline{\mathbf{B}} + \epsilon^2 \underline{\mathbf{I}}) \underline{\mathbf{x}} = \underline{\mathbf{B}}^T \underline{\mathbf{A}}^T \underline{\mathbf{b}} . \quad (3.32)$$

But $\underline{\mathbf{B}}$ is symmetric and non-singular, therefore we may multiply 6.32)

from left by $\underline{\mathbf{B}}^{-1}$ and we get

$$(\underline{\mathbf{B}} + \epsilon^2 \underline{\mathbf{B}}^{-1}) \underline{\mathbf{x}} = \underline{\mathbf{A}}^T \underline{\mathbf{b}}$$

which is (3.30). From (3.30) we have

Lemma 3.1. The double relaxed solution $\underline{\mathbf{x}}_{-\mathbf{d}\epsilon}$ minimizes

$$\mathbf{Q}_{\mathbf{d}}(\underline{\mathbf{x}}) = \|(\underline{\mathbf{A}}^T \underline{\mathbf{A}} + \epsilon \underline{\mathbf{I}}) \underline{\mathbf{x}} - \underline{\mathbf{A}}^T \underline{\mathbf{b}}\|^2 + \epsilon \|\underline{\mathbf{x}}\|^2 .$$

The aim of relaxing is to approximate the "ideal solution" [13]

without **explicitely** computing the singular value decomposition. By the "ideal" solution we mean the following: Let

$$\underline{\mathbf{A}} = \underline{\mathbf{u}} \underline{\Sigma} \underline{\mathbf{v}}^T , \quad \underline{\mathbf{u}} , \underline{\mathbf{v}} \text{ orthogonal}$$

$$\underline{\Sigma} = \text{diag}(\sigma_1, \dots, \sigma_n) , \quad \sigma_1 \geq \dots \geq \sigma_r > \sigma_{r+1} = \dots = \sigma_n = 0$$

be the singular value decomposition of $\underline{\mathbf{A}}$. If the **datas** and the computation were exact then the shortest solution of

$$\|\underline{\mathbf{A}} \underline{\mathbf{x}} - \underline{\mathbf{b}}\|^2 = \min$$

would be $\underline{\mathbf{x}} = \underline{\mathbf{A}}^+ \underline{\mathbf{b}} = \underline{\mathbf{u}} \underline{\mathbf{y}}$ where $\underline{\mathbf{y}}$ is defined by

$$\underline{\mathbf{y}}_i = \begin{cases} \frac{c_i}{\sigma_i} & \text{if } \sigma_i \neq 0 \\ 0 & \text{if } \sigma_i = 0 . \end{cases} \quad (c := \underline{\mathbf{v}}^T \underline{\mathbf{b}})$$

In practical **computation** however it is not clear which σ_i are to be interpreted as zero [11,26]. One is forced to make a rank decision,

to prescribe a tolerance τ and to **compute**

$$\tilde{\underline{\mathbf{y}}}_i = \begin{cases} \frac{c_i}{\sigma_i} & \text{if } \sigma_i \geq \tau \\ 0 & \text{if } \sigma_i < \tau \end{cases} . \quad (3.33)$$

We will call $\tilde{\underline{\mathbf{y}}}_i$ the "ideal solution".

Now consider the function defined on $[0, \infty)$

$$k(u) := \begin{cases} 0 & \text{if } \sigma < \tau \\ \frac{1}{\sigma} & \text{if } \sigma > \tau \end{cases} \quad (3.34)$$

The coefficient of $\tilde{y}_i$ multiplying c_i is given by $k(\sigma_i)$. The length of y depends on the choice of τ and is

$$\|\tilde{y}\|^2 = \sum_{\sigma_i \geq \tau} \left(\frac{c_i}{\sigma_i} \right)^2.$$

Let $\gamma := \|\tilde{y}\| / \|b\|$. We can ask the question: how to choose ϵ such that the relaxed and the doubly relaxed solution have the same or at least approximately the same length as the ideal solution. The answer is given in [13]. We have

$$\begin{aligned} \text{if } \epsilon \geq \frac{1}{2\gamma} &\Rightarrow \|x_\epsilon\| < \gamma \|b\| = \|\tilde{y}\| \\ \text{if } \epsilon \geq \frac{1}{2\gamma^2} &\Rightarrow \|x_{d\epsilon}\| < \gamma \|b\| = \|\tilde{y}\|. \end{aligned}$$

and

If we transform the relaxed problem (3.29) and the doubly relaxed problem (3.31) using SVD we get

$$(y_\epsilon)_i = \frac{\sigma_i}{\sigma_i^2 + \epsilon^2} c_i \quad (3.35)$$

$$(y_{d\epsilon})_i = \frac{\sigma_i}{\sigma_i^2 + \epsilon^2 + \frac{\epsilon^2}{\sigma_i^2}} c_i. \quad (3.36)$$

We see that we can think of (3.35), (3.36) as approximations of (3.33). More precisely, we want to choose ϵ so that the two functions

$$k_\epsilon(\sigma) := \frac{\sigma}{\sigma^2 + \epsilon^2} \quad (3.37)$$

and

$$k_{d\epsilon}(\sigma) := \frac{\sigma}{\sigma^2 + \epsilon^2 + \frac{\epsilon^2}{\sigma^2}} \quad (3.38)$$

approximate the function $k(\sigma)$ (3.34). Rutishauser [35] and Molinari [13] show that indeed $k_{d\epsilon}$ is a better approximation than k_ϵ . More research could be done in this direction, e.g. relaxing the normal equations of the original problem

$$\begin{pmatrix} \bar{A}^T \bar{A} \\ \epsilon \bar{I} \end{pmatrix} \bar{x} \approx \begin{pmatrix} \bar{A}^T \bar{b} \\ 0 \end{pmatrix}$$

would yield the function

$$k_{1\epsilon} := \frac{\sigma}{\sigma^2 + \epsilon^2 \frac{2}{\sigma^2}}$$

which is somehow between both discussed above. It is clear that those different relaxed solutions have to be **computed** without forming the normal equations. A program for $\bar{x}_\epsilon$ and $\bar{x}_{d\epsilon}$ is given in [13].

Theorem 3.2. Let $\bar{A} = \begin{pmatrix} \bar{A} \\ \epsilon \bar{I} \end{pmatrix}$ and $\bar{A}^+ = \begin{pmatrix} \bar{B}^+ & C^+ \\ -\bar{\epsilon} & -\bar{\epsilon} \end{pmatrix}$. ($\bar{B}^+$ is an $n \times n$

matrix.) Then for ϵ sufficiently small we have

$$\bar{B}_\epsilon^+ = \bar{A}^+ + \sum_{j=1}^{\infty} (-1)^j \bar{A}^+ (\bar{A}^+)^{2j} \bar{\epsilon}^{2j},$$

Proof. Let $A = U \Sigma V^T$ with $\Sigma = \begin{pmatrix} \sigma_1 & & 0 \\ & \ddots & \\ 0 & & \sigma_r & \\ & & & 0 \end{pmatrix}$ be the singular

value decomposition, with $\sigma_1 \geq \dots \geq \sigma_r > 0$. Then

$$\begin{aligned} \bar{A}_{-\epsilon}^+ &= (\bar{A}^T \bar{A} + \epsilon^2 I)^{-1} \bar{A}_{-\epsilon}^T = (\bar{V} \bar{\Sigma}^T \bar{U}^T \bar{U} \bar{\Sigma} \bar{V}^T + \epsilon^2 I)^{-1} (\bar{V} \bar{\Sigma}^T \bar{U}^T, I_{-\epsilon}) \\ &\quad \underbrace{\phantom{(\bar{V} \bar{\Sigma}^T \bar{U}^T \bar{U} \bar{\Sigma} \bar{V}^T + \epsilon^2 I)^{-1}}}_{I_m} \\ &= \bar{V} (\bar{\Sigma}^T \bar{\Sigma} + \epsilon^2 I)^{-1} (\bar{\Sigma}^T \bar{U}^T, \epsilon \bar{V}^T) \\ &= \bar{V} ([(\bar{\Sigma}^T \bar{\Sigma} + \epsilon^2 I)^{-1} \bar{\Sigma}^T] \bar{U}^T, \epsilon (\bar{\Sigma}^T \bar{\Sigma} + \epsilon^2 I)^{-1} \bar{V}^T) \\ &= [\bar{B}_{-\epsilon}^+, \bar{C}_{-\epsilon}^+] \end{aligned}$$

$$\bar{B}_{-\epsilon}^+ = \bar{V} (\bar{\Sigma}^T \bar{\Sigma} + \epsilon^2 I)^{-1} \bar{\Sigma}^T \bar{U}^T$$

$$\bar{B}_{-\epsilon}^+ = \bar{V} \left| \begin{array}{c|c} \frac{\sigma_1^2}{a_1^2 + \epsilon^2} & 0 \\ \hline 0 & \frac{\sigma_k^2}{\sigma_r^2 + \epsilon^2} \end{array} \right| \bar{U}^T$$

$$\text{now } \frac{a_{-1}^2}{a_1^2 + \epsilon^2} = \frac{1}{\sigma_1^2 + \epsilon^2} \frac{1}{1 + (\epsilon/\sigma_1)^2} = \frac{1}{\sigma_1^2} \sum_{j=0}^{\infty} (\epsilon/\sigma_1)^{2j} \quad \text{for } |\epsilon| < \sigma_r,$$

$$\bar{B}_{-\epsilon}^+ = \sum_{j=0}^{\infty} \bar{V} \underbrace{\begin{bmatrix} 1/\sigma_1^{2j+1} & & 0 \\ & \ddots & \\ 0 & & 1/\sigma_r^{2j+1} \end{bmatrix}}_{\bar{\Sigma}^+ (\bar{\Sigma}^+)^T \bar{\Sigma}^+ 2^j} \bar{U}^T$$

$$\text{where } \bar{\Sigma}_{-}^+ = \begin{bmatrix} 1/\sigma_1 & & 0 \\ & \ddots & \\ 0 & & 1/\sigma_r \end{bmatrix}_n,$$

we observe

$$\bar{V} \bar{\Sigma}_{-}^+ (\bar{\Sigma}_{-}^+)^T \bar{\Sigma}_{-}^+ 2^j \bar{U}^T = \bar{A}^+ (\bar{A}^+)^T \bar{A}^+ 2^j \quad . \quad 0$$

Remark. Theorem 3.2 shows that we can extrapolate $\bar{x}_M = \bar{A}^+ \bar{b}$ from the relaxed solution $\bar{x}_{\epsilon} = \bar{B}_{-\epsilon}^+ \bar{b}_{-}$. Using $\epsilon_1 = \epsilon/2^1$ we can apply the Romberg extrapolation since only even powers of ϵ occur. It has the advantage that the rank of $\bar{A}$ must not be determined.

...

4. Minimum Norm Solution with Given Norm of the Residual.

In this section we discuss the special case of a least squares problem with a quadratic constraint where $\underline{A} = \underline{I}$:

$$\left. \begin{aligned} \|\underline{x}\| &= \min \\ \|\underline{C}\underline{x} - \underline{d}\| &= \alpha \end{aligned} \right\} \quad \text{subject to} \quad (P3E)$$

The problem with inequality constraint $\|\underline{C}\underline{x} - \underline{d}\| \leq \alpha$ will be denoted by (P3). Geometrically (P3E) means that we are looking for a point on the ellipsoid $\|\underline{C}\underline{x} - \underline{d}\|^2 = \alpha^2$ that is nearest to the origin. From the theory in Section 2 we have that the solution of (P3E) is a solution $(\underline{x}, \lambda)$ of the normal equations

$$(\underline{I} + \lambda \underline{C}^T \underline{C}) \underline{x} = \lambda \underline{C}^T \underline{d} \quad (4.1)$$

$$\|\underline{C}\underline{x} - \underline{d}\|^2 = \alpha^2 \quad ; \quad (4.2)$$

with largest λ . Since $NS(\underline{I}) \cap NS(\underline{C}) = \{0\}$ the solution is unique if $\lambda \neq -\gamma_1^2$ where γ_1 is a singular value of $\underline{C}$. The solution exists only if

$$\alpha \geq \|(\underline{C} \underline{C}^T - \underline{I}) \underline{d}\| := \alpha_{\min}.$$

For problem (P3) the solution is $\underline{x} = \underline{0}$ if $\alpha \geq \|\underline{d}\|$. The length function f is in this case

$$f(h) = \|(\underline{\lambda} - \underline{C}(\underline{I} + \underline{\lambda} \underline{C}^T \underline{C})^{-1} \underline{C}^T - \underline{I}) \underline{d}\|^2. \quad (4.3)$$

By Corollary 3.1, (3.13), this is equal to

$$f(h) = \|(\underline{I} + \lambda \underline{C} \underline{C}^T)^{-1} \underline{d}\|^2 \quad (4.5)$$

which gives again the connection to the dual equations.

4.1 The Dual Normal Equations.

We transform the normal equations (4.1), (4.2) as follows: We multiply (4.1) from left by $\underline{C}$ giving

$$(\underline{I} + \lambda \underline{C} \underline{C}^T) \underline{C} \underline{x} = \lambda \underline{C} \underline{C}^T \underline{d}. \quad (4.6)$$

If we subtract from (4.6) on both sides $\lambda \underline{C} \underline{C}^T \underline{d} + \underline{d}$ we get

$$(\underline{I} + h \underline{C} \underline{C}^T)(\underline{C} \underline{x} - \underline{d}) = -\underline{d}.$$

Finally introducing the new variable $\underline{z} = \underline{C} \underline{x} - \underline{d}$ we get the theorem:

Theorem 4.1.

(i) Let $(\underline{x}, \lambda)$ be a solution of the primal normal equations

$$\begin{aligned} (\underline{I} + \lambda \underline{C}^T \underline{C}) \underline{x} &= \lambda \underline{C}^T \underline{d} \\ \|\underline{C} \underline{x} - \underline{d}\|^2 &= \alpha^2 \end{aligned} \quad (4.7)$$

then $(\underline{z}, \lambda)$ with $\underline{z} := \underline{C} \underline{x} - \underline{d}$ is a solution of the dual normal equations

$$\begin{aligned} (\underline{I} + \lambda \underline{C} \underline{C}^T) \underline{z} &= -\underline{d} \\ \|\underline{z}\|^2 &= \alpha^2 \end{aligned}, \quad (4.8)$$

(ii) Let $(\underline{z}, \lambda)$ be a solution of (4.8), then $(\underline{x}, \lambda)$ with $\underline{x} = -h \underline{C}^T \underline{z}$ is a solution of (4.7).

Proof.

$$\begin{aligned} (i) \quad (\underline{I} + \lambda \underline{C} \underline{C}^T)(\underline{C} \underline{x} - \underline{d}) &= \underline{C}(\underline{I} + \lambda \underline{C}^T \underline{C}) \underline{x} - \underline{d} - \lambda \underline{C} \underline{C}^T \underline{d} \\ &= \underline{C}(\lambda \underline{C}^T \underline{d}) - \underline{d} - \lambda \underline{C} \underline{C}^T \underline{d} = -\underline{d}. \end{aligned}$$

and

5. Computational Aspects.

To compute the solution of a least squares problem with a quadratic constraint we have to determine iteratively the largest solution of the secular equation. For the problem (PIE), (P2E) and (P3E) we can expect numerical difficulties since at every step of the iteration we have to solve a system of equations with the matrix

$$\underline{\underline{A}}^T \underline{\underline{A}} + \lambda \underline{\underline{C}}^T \underline{\underline{C}} \quad \text{where } h < 0 \quad (5.1)$$

which is in general not positive definite. A good way to solve this system is to transform it to diagonal form using BSVD of von Loan [45]. If $\underline{\underline{A}} = \underline{\underline{I}}$ or $\underline{\underline{C}} = \underline{\underline{I}}$ we can use SVD to **diagonalize** the system stably [48]. To avoid the forming of $\underline{\underline{A}}^T \underline{\underline{A}}$ and $\underline{\underline{C}}^T \underline{\underline{C}}$ we could consider the augmented $(m+n+p) \times (m+n+p)$ system

$$\begin{pmatrix} -1 & \underline{\underline{A}} & \underline{\underline{O}} \\ \underline{\underline{A}}^T & \underline{\underline{O}} & \underline{\underline{C}}^T \\ 0 & \underline{\underline{C}} & \frac{1}{\lambda} \underline{\underline{I}} \end{pmatrix} \begin{pmatrix} \underline{\underline{y}} \\ \underline{\underline{x}} \\ \underline{\underline{z}} \end{pmatrix} = \begin{pmatrix} \underline{\underline{b}} \\ \underline{\underline{O}} \\ \underline{\underline{d}} \end{pmatrix} \quad (5.2)$$

Björk has shown [1] that for the **augmented** system

$$\begin{pmatrix} \underline{\underline{I}} & \underline{\underline{A}} \\ \underline{\underline{A}}^T & \underline{\underline{O}} \end{pmatrix} \begin{pmatrix} \underline{\underline{r}} \\ \underline{\underline{x}} \end{pmatrix} = \begin{pmatrix} \underline{\underline{b}} \\ \underline{\underline{O}} \end{pmatrix} \quad (5.3)$$

We can obtain the condition number $2 \cdot \kappa(\underline{\underline{A}})$ with appropriate scaling (instead of $\kappa^2(\underline{\underline{A}})$ for the normal equations $\underline{\underline{A}}^T \underline{\underline{A}} \underline{\underline{x}} = \underline{\underline{A}}^T \underline{\underline{b}}$). However in (5.2) the condition number depends on the value of λ and may be large if h is near $-\mu_1$ where μ_1 is an eigenvalue of $\det(\underline{\underline{A}}^T \underline{\underline{A}} - \mu \underline{\underline{C}}^T \underline{\underline{C}}) = 0$.

$$\alpha^2 = \|\underline{\underline{z}}\|^2 = \|\underline{\underline{C}}\underline{\underline{x}} - \underline{\underline{d}}\|^2.$$

$$(ii) \quad (\underline{\underline{I}} + h \underline{\underline{C}}^T \underline{\underline{C}})(-\lambda \underline{\underline{C}}^T \underline{\underline{z}}) = -h \underline{\underline{C}}^T (\underline{\underline{I}} + h \underline{\underline{C}} \underline{\underline{C}}^T) \underline{\underline{z}} = h \underline{\underline{C}}^T \underline{\underline{b}}.$$

$$\|\underline{\underline{C}}\underline{\underline{x}} - \underline{\underline{d}}\|^2 = \|\underline{\underline{C}}(-\lambda \underline{\underline{C}}^T \underline{\underline{z}}) - \underline{\underline{d}}\|^2 = \|\underline{\underline{z}}\|^2 = \alpha^2. \quad \square$$

Equating both expressions for x and z again gives us the identities of Corollary 3.1.

contrast to the relaxed least squares problem, the dual equations exists for every λ .

4.2 Representation as Least Squares Problem.

If we want to solve (P3) then the **solution** is $x = 0$ if $\alpha > \|\underline{\underline{d}}\|$ and exists only if $\alpha > \alpha_{\min} = \|(\underline{\underline{C}} \underline{\underline{C}}^+ - \underline{\underline{I}})\underline{\underline{d}}\|$. For

$$\alpha_{\min} < \alpha < \|\underline{\underline{d}}\|.$$

We have to **compute** the solution iteratively solving the secular equation $f(h) = \alpha^2$ where

$$f(\lambda) = \|\underline{\underline{C}}\underline{\underline{x}} - \underline{\underline{d}}\|^2 = \|\underline{\underline{z}}\|^2.$$

$\underline{\underline{x}}$ and $\underline{\underline{z}}$ are obtained solving either

$$\begin{pmatrix} \underline{\underline{I}} \\ \sqrt{\lambda} \underline{\underline{C}} \end{pmatrix} \underline{\underline{z}} \approx \begin{pmatrix} \underline{\underline{O}} \\ \sqrt{\lambda} \underline{\underline{d}} \end{pmatrix}$$

or the dual problem

$$\begin{pmatrix} \underline{\underline{I}} \\ \sqrt{\lambda} \underline{\underline{C}}^T \end{pmatrix} \underline{\underline{z}} \approx \begin{pmatrix} -\underline{\underline{d}} \\ \underline{\underline{O}} \end{pmatrix}$$

The solution is numerically much better defined for problem (P1), (P2) and (P3) where $\lambda > 0$. We can formulate the normal equations as least squares problems and besides (BSVD)- and SVD-diagonalization, there are several different ways to compute stably and cheaply the solution, especially if we can use a structure of the matrix. Since we are solving the secular equation with some iterative method we may have to compute derivatives of the length function f (see Section 5.4). If the iterative method converges fast it may not be necessary to transform the problem at all (especially if $\underline{A}$ and $\underline{C}$ are sparse). In general, however, it appears to be best [8], [12] to bidiagonalize the matrix for problems (P2) and (P3) before iterating.

5.1 Solution of a Relaxed Least Squares Problem with Band Matrix.

We consider (P2) with $\underline{A}$ respectively (P3) with $\underline{C}$ being a band matrix. For a given $\lambda > 0$ we have to solve the least squares problems

$$\begin{pmatrix} \underline{A} \\ \sqrt{\lambda} \underline{I} \end{pmatrix} \underline{x} \approx \begin{pmatrix} \underline{b} \\ \underline{0} \end{pmatrix} \quad (5.4)$$

respectively

$$\begin{pmatrix} \underline{I} \\ \sqrt{\lambda} \underline{C} \end{pmatrix} \underline{x} \approx \begin{pmatrix} \underline{0} \\ \sqrt{\lambda} \underline{d} \end{pmatrix} . \quad (5.5)$$

If we interchange the equations in (5.5) we have in either case a least squares problem of the form

$$\begin{pmatrix} \underline{B} \\ \underline{D} \end{pmatrix} \underline{x} \approx \begin{pmatrix} \underline{c}_1 \\ \underline{c}_2 \end{pmatrix} \quad (5.6)$$

where $\underline{D}$ is a diagonal matrix ($\underline{I}$ or $\sqrt{\lambda} \underline{I}$) and $\underline{B}$ is a band matrix ($\underline{A}$ or $\sqrt{\lambda} \underline{C}$).

Using Givens-rotations we shall show how an orthogonal matrix G can be found such that

$$G \begin{pmatrix} \underline{B} \\ \underline{D} \end{pmatrix} = \begin{pmatrix} \underline{0} \\ \underline{F} \end{pmatrix} \quad (5.7)$$

where $\underline{F}$ is a band upper triangular matrix. If the same rotations are applied to the right hand side the solution can be found by backsubstitution using $\underline{F}$.

To describe the transformation we assume that B is an $(m \times n)$ matrix with

$$b_{ij} = 0 \quad \text{if } i-j > m_1 \text{ or } j-i > m_2 .$$

This means that $\underline{B}$ has m_1 diagonals below the main diagonal and m_2 diagonals above, that contains the possibly non-zero elements.

Let

$$k_{\min}(k) := \max\{k-m_2, 1\}$$

$$k_{\max}(k) := \min\{k+m_1, m\} .$$

Then the k -th column of B has the only (possibly) non-zero elements

$$b_{ik} , \quad i = k_{\min}(k), \dots, k_{\max}(k) .$$

In n -steps we now perform the transformation (5.7) annihilating in the k -th step the k -th column of B and producing the k -th row of F . Let $\underline{F}$ be at the beginning the diagonal matrix D and defined by

$$\text{rot}(i, k, x)$$

the multiplication of $\begin{pmatrix} \underline{\mathbf{B}} \\ \underline{\mathbf{F}} \end{pmatrix}$ from left by a Givens matrix that changes the i -th row of $\underline{\mathbf{B}}$ and the k -th row of $\underline{\mathbf{F}}$ annihilating the element x . Then the transformation (5.7) is described by

```

for k := 1 step 1 until n do
  for i := k_min(k) step 1 until k_max(k) do
    rot(i, k, b_ik);

```

In the k -th step of this algorithm we annihilate $m_1 + m_2 + 1$ elements (or less at the border of $\underline{\mathbf{B}}$) using $\sim 2(m_1 + m_2 + 1)^2$ multiplications. For $i = k_{\min} = k - m_2$ the Givens rotation changes only 2 elements: b_{ik} and f_{kk} . Then for $i = k_{\min} + 1$, 4 elements are affected, etc., until for $i = k_{\max} = k + m_1$ when $b_{k+m_1, k}$ is rotated to zero, we have to change $m_1 + m_2 + 1$ elements. The whole transformation requires therefore $\sim 2n(m_1 + m_2 + 1)^2$ multiplications which is comparable to Gaussian elimination of the normal equations.

Especially if $\underline{\mathbf{B}}$ is upper bidiagonal the first step to annihilate the first and second column is done as follows:

$$\begin{array}{c}
 \left[\begin{array}{cccc|cccc}
 q_1 & e_1 & & & 0 & e'_1 & & \\
 & q_2 & e_2 & & & q_2 & e_2 & \\
 & & q_3 & e_3 & & & q_3 & e_3 \\
 & & & q_4 & & & & q_4 \\
 & d_1 & & & & u_1 & v_1 & \\
 & & d_2 & & & & d_2 & \\
 & & & d_3 & & & & d_3 \\
 & & & & & & & d_4
 \end{array} \right] \xrightarrow{\text{rot}(1,1,q_1)} \left[\begin{array}{cccc|cccc}
 0 & 0 & & & 0 & 0 & & \\
 & q_2 & e_2 & & & 0 & e_2 & \\
 & & q_3 & e_3 & & & q_3 & e_3 \\
 & & & q_4 & & & & q_4 \\
 & u_1 & v_1 & & & u_1 & v_1 & \\
 & & d'_2 & & & & u_2 & v_2 \\
 & & & 4 & & & & d_3 \\
 & & & & & & & d_4
 \end{array} \right] \xrightarrow{\text{rot}(2,2,q_2)} \left[\begin{array}{cccc|cccc}
 0 & 0 & & & 0 & 0 & & \\
 & q_2 & e_2 & & & 0 & e_2 & \\
 & & q_3 & e_3 & & & q_3 & e_3 \\
 & & & q_4 & & & & q_4 \\
 & u_1 & v_1 & & & u_1 & v_1 & \\
 & & d'_2 & & & & u_2 & v_2 \\
 & & & 4 & & & & d_3 \\
 & & & & & & & d_4
 \end{array} \right]
 \end{array}$$

The following procedure performs the transformation

$$\underline{\mathbf{G}} \begin{pmatrix} \underline{\mathbf{B}} \\ \underline{\mathbf{D}} \end{pmatrix} = \begin{pmatrix} \underline{\mathbf{O}} \\ \underline{\mathbf{F}} \end{pmatrix}$$

where

$$\bar{B} = \begin{pmatrix} q_1 & e_1 & 0 \\ & \ddots & \vdots \\ & & e_{n-1} \\ 0 & & & q_n \end{pmatrix} \quad \bar{F} = \begin{pmatrix} u_1 & v_1 & 0 \\ & \ddots & \vdots \\ & & v_{n-1} \\ & 0 & & u_n \end{pmatrix}$$

and. stores the Givens rotations in two arrays $co[i]$, $si[i]$

```

procedure updec (n,q,e,d,u,v,c0,si);
  value n; integer n;
  array q,e,d,u,v,c0,si;
  begin integer i; real t,c,s,h;
    procedure rot(a,b);
      value a,b; real a,b;
      begin real t;
        if b = 0 then begin c := 0; s := 1 end
        else
          begin t := -a/b; c := 1/sqrt(1+t*t2); s := t*c end
        end;
      for i := 1 step 1 until n do u[i] := d[i];
      for i := 1 step 1 until n do
        begin
          rot(q[i],u[i]);
          c0[2*i-1] := c; si[2*i-1] := s;
          u[i] := -q[i]*s + u[i]*c;
          if i < n then
            begin
              v[i] := -e[i]*s; h := e[i]*c;
              rot(h,u[i+1]);
              c0[2*i] := c; si[2*i] := s;
              u[i+1] := -h*s + u[i+1]*c
            end if;
          end i;
        end i;
      end updec;

```

The following procedure performs the transformation

$$G \begin{pmatrix} c_1 \\ \vdots \\ c_l \end{pmatrix}$$

and solves the bidiagonal system by backsubstitution.

```

procedure solve (n,u,v,c0,si,cl,c2,y);
  value n; integer n;
  array c0, si, cl, c2, y;
  begin
    integer i; real t,c,s,h;
    for i := 1 step 1 until n do
      begin
        c := c0[2*i-1]; s := si[2*i-1];
        h := cl[i]*c + c2[i]*s;
        c2[i] := -cl[i]*s + c2[i]*c;
        cl[i] := h;
        if i < n then
          begin
            c := c0[2*i]; s := si[2*i];
            h := cl[i]*c + c2[i]*s;
            c2[i+1] := -cl[i]*s + c2[i+1]*c;
            cl[i] := h;
          end if;
        end i;
      end i;
    y[n1] := c2[n] / u[n];
    for i := n-1 step -1 until 1 do
      y[i] := (c2[i] - v[i]*y[i+1]) / u[i];
    end solve;

```


In step 1 we annihilate the elements of the first column of A using the first row of C . This produces the new non-zero elements $\odot$. In a second "cleaning" step we zero the elements $\odot$ in the top part using the third row:

$$\begin{bmatrix} \textcircled{X} & \textcircled{X} & \textcircled{X} \\ \textcircled{X} & \textcircled{X} & \textcircled{X} \\ \textcircled{X} & \textcircled{X} & \textcircled{X} \end{bmatrix} \rightarrow \begin{bmatrix} \textcircled{X} & \textcircled{X} & \textcircled{X} \\ \textcircled{X} & \textcircled{X} & \textcircled{X} \\ \textcircled{X} & \textcircled{X} & \textcircled{X} \end{bmatrix} \rightarrow \begin{bmatrix} \textcircled{X} & \textcircled{X} & \textcircled{X} \\ \textcircled{X} & \textcircled{X} & \textcircled{X} \\ \textcircled{X} & \textcircled{X} & \textcircled{X} \end{bmatrix}$$

After these two steps the remaining matrix is:

[illegible]

Now we can use the second **row** of $\underline{C}$ to zero the second column of **etc.** A special treatment is needed at the border of the matrix. using the last row of $\underline{C}$ to zero out the n-5 column gives

Diagram illustrating the construction of the 2×2 block matrix A_1 from the 1×1 block matrix A .

The top part shows the 1×1 block matrix A (a single row of blocks) and the 2×2 block matrix A_1 (a 2×2 grid of blocks). The blocks in A_1 are labeled with 0 or x .

The bottom part shows the construction of A_1 from A using the Kronecker product. The matrix A_1 is formed by taking the Kronecker product of the 2×2 identity matrix I_2 (represented by a 2×2 grid of 1 's and 0 's) and the 1×1 block matrix A .

The resulting 2×2 block matrix A_1 is:

$$A_1 = \begin{pmatrix} 1 \otimes A & 0 \otimes A \\ 0 \otimes A & 1 \otimes A \end{pmatrix}$$

where 1 and 0 are the 1×1 identity and zero matrices, respectively, and $\otimes$ denotes the Kronecker product.

In the cleaning step 2 we can zero **all** but **3** elements in the upper part of the **matrix** using the last row of **A** . Now the solution can be computed by backsubstitution: first **x₆** and **x₇** using **0'** and the other unknown using **A'** . In Reinsch's "Smoothering with Spline Functions" [33] the problem

$$(Q_D^T)^2 Q + \lambda E)z = Q^T y \quad (5.10)$$

subject to

$$\frac{\|D_Q\|_2}{\|Z\|_2} = \frac{\sigma_2}{\sigma_1} \quad (5.11)$$

occurs, where $\lambda > 0$, $\underline{D}$ is an $(n+1) \times (n+1)$ diagonal matrix, $\underline{Q}$ an $(n-1) \times (n-1)$ tridiagonal matrix and $\underline{T}$ a positive definite tridiagonal $(n-1) \times (n-1)$ matrix. S is a given constant and y a given $(n+1)$ vector. Reinsch solves the normal equation (6.0) using the Cholesky decomposition of the **penta-diagonal** coefficient matrix. Using the above described algorithm we can solve (5.10) as a least squares problem without sacrificing the sparseness of the **matrices**. As preparation we compute the Cholesky decomposition of T :

$$T = \underline{\underline{P}}^T \underline{\underline{B}}, \quad \underline{\underline{B}} \text{ upper bidiagonal}$$

and (5.10) then becomes

$$\begin{pmatrix} \underline{\underline{D}} & \underline{\underline{Q}} \\ \sqrt{\lambda} & \underline{\underline{B}} \end{pmatrix} \underline{\underline{z}} \approx \begin{pmatrix} \underline{\underline{D}}^{-1} \underline{\underline{y}} \\ \underline{\underline{0}} \end{pmatrix}. \quad (5.12)$$

The matrix of (5.12) has the form (n = 5)

$$\left[\begin{array}{cccccc|cccccc} \text{X} & & & & & & 0 & \textcircled{\text{X}} & & & \\ \text{x} & \text{x} & & & & & 0 & \text{X} & & & \\ \text{x} & \text{x} & \text{x} & & & & 0 & \text{X} & \text{X} & & \\ & \text{x} & \text{x} & \text{x} & & & \text{X} & \text{X} & \text{X} & \text{X} & \\ & & \text{x} & \text{x} & \text{x} & & & \text{X} & \text{X} & \text{X} & \\ & & & \text{x} & \text{x} & \text{X} & & & & & \\ & & & & \text{x} & & \text{X} & \text{X} & \textcircled{\text{X}} & & \\ & & & & & & & \text{x} & \text{x} & \text{x} & \text{X} \end{array} \right] \xrightarrow{\text{1}} \left[\begin{array}{cccccc|cccccc} \text{x} & \text{x} & & & & & \text{X} & \text{X} & \textcircled{\text{X}} & & \\ & \text{x} & \text{x} & & & & & \text{x} & \text{x} & \text{x} & \\ & & \text{x} & \text{x} & & & & & & & \\ & & & \text{x} & \text{x} & & & & & & \\ & & & & \text{x} & & & & & & \\ & & & & & & & & & & \end{array} \right];$$

Zeroing the first column in step 1 leaves one element $\textcircled{\text{X}}$ that has to be removed in the cleaning step.

5.3 Bidiagonalization.

If the matrices $\underline{\underline{A}}$ and $\underline{\underline{C}}$ are dense then for problems (P2), (P3) and (P1) (after performing Eldén's transformation, see Section 3.3) we have to solve a least square problem of the form

$$\begin{pmatrix} \underline{\underline{A}} \\ \sqrt{\lambda} & \underline{\underline{I}} \end{pmatrix} \underline{\underline{x}} \approx \begin{pmatrix} \underline{\underline{b}} \\ \underline{\underline{0}} \end{pmatrix} \quad (5.13)$$

for every new value of λ . If $\underline{\underline{P}}$ and $\underline{\underline{Q}}$ are orthogonal then an equivalent system to (5.13) is

$$\begin{pmatrix} \underline{\underline{P}}^T & \underline{\underline{0}} \\ \underline{\underline{0}} & \underline{\underline{Q}}^T \end{pmatrix} \begin{pmatrix} \underline{\underline{A}} \\ \sqrt{\lambda} & \underline{\underline{I}} \end{pmatrix} \underline{\underline{Q}}^T \underline{\underline{x}} \approx \begin{pmatrix} \underline{\underline{P}}^T \underline{\underline{b}} \\ \underline{\underline{0}} \end{pmatrix} \\ \Rightarrow \begin{pmatrix} \underline{\underline{P}}^T \underline{\underline{A}} \underline{\underline{Q}} \\ \sqrt{\lambda} & \underline{\underline{I}} \end{pmatrix} \underline{\underline{y}} \approx \begin{pmatrix} \underline{\underline{c}} \\ \underline{\underline{0}} \end{pmatrix} \quad (5.14)$$

with $\underline{\underline{Y}} := \underline{\underline{Q}}^T \underline{\underline{x}}$ and $\underline{\underline{c}} := \underline{\underline{P}}^T \underline{\underline{b}}$.

Now we can choose $\underline{\underline{P}}$ and $\underline{\underline{Q}}$ so that (5.14) is simpler to solve than (5.13). Moré [29] proposed to choose $\underline{\underline{P}}$ and $\underline{\underline{Q}}$ so that

$$\underline{\underline{B}} := \underline{\underline{P}}^T \underline{\underline{A}} \underline{\underline{Q}} = \begin{pmatrix} \boxed{\underline{\underline{B}}} & \\ & 0 & 0 \end{pmatrix}, \quad (5.15)$$

i.e., to perform the QR decomposition of $\underline{\underline{A}}$ with column pivoting. However the solution of (5.14) in this case is still an n^3 process. It has been pointed out by More/ (private communication) that in his application [29] only $l \leq 2$ iterations are needed to solve the secular equation. Therefore it does not matter in this case if the solution of (5.14) is an n^3 process. In general however we will have to perform much more iterations (see Section 6). With about double the amount of work to perform the decomposition (5.15) we can bidiagonalize $\underline{\underline{A}}$:

$$\underline{\underline{B}} := \underline{\underline{P}}^T \underline{\underline{A}} \underline{\underline{Q}} = \begin{pmatrix} \text{---} & \\ \text{---} & 0 \end{pmatrix} \quad (5.16)$$

as proposed by [8] and [12]. Then to solve (5.14) we can use the algorithm of Section 5.1 that requires only $\sim n$ multiplications.

Bidiagonalizing a matrix A is the first step in computing the singular value decomposition [48]. But the diagonalizing of B requires that P and Q be formed explicitly and updated after every QR-step which makes the process fairly expensive. To transform our problem to bidiagonal form we do not need P and Q explicitly. It is sufficient to store the vectors of the Householder matrices that perform the transformation.

Let A be an $m \times n$ matrix. Then we can bidiagonalize A in n steps [48]:

$$\underline{B} = \underbrace{\underline{P}_n \dots \underline{P}_1}_{\underline{P}^T} \underline{A} \underbrace{\underline{Q}_1 \dots \underline{Q}_{n-1}}_{\underline{Q}}$$

where

$$\underline{P}_{-j} = \begin{pmatrix} \underline{I}_{j-1} & \underline{0} \\ \underline{0} & \underline{\tilde{P}}_{-j} \end{pmatrix} \begin{matrix} \} j-1 \\ \} m-j+1 \end{matrix}$$

$$\underline{\tilde{P}}_{-j} = \underline{I} - \underline{w}_{-j} \underline{w}_{-j}^T \quad \text{with} \quad \|\underline{w}_{-j}\| = \sqrt{2}$$

and

$$\underline{Q}_{-j} = \begin{pmatrix} \underline{I}_j & \underline{0} \\ \underline{0} & \underline{\tilde{Q}}_{-j} \end{pmatrix} \begin{matrix} \} j \\ \} n-j \end{matrix}$$

$$\underline{\tilde{Q}}_{-j} = \underline{I} - \underline{v}_{-j} \underline{v}_{-j}^T \quad \text{with} \quad \|\underline{v}_{-j}\| = \sqrt{2} .$$

The transformation vectors $\underline{w}_{-j}$ and $\underline{v}_{-j}$ are stored in the matrix A as follows.

$$A = \begin{array}{|c|c|c|} \hline & & \underline{v}_{-1}^T \\ \hline & \underline{w}_2 & \underline{w}_2^T \\ \hline \underline{w}_1 & & \ddots \\ \hline \end{array} \quad (5.17)$$

To compute the vectors

$$\underline{P} \underline{x} \quad , \quad \underline{P}^T \underline{x} \quad , \quad \underline{Q} \underline{y} \quad , \quad \underline{Q}^T \underline{y}$$

we do not need P and Q explicitly. In (5.17) we have all the information to compute the operator P , P^T , Q or Q^T that act on a certain vector.

The computation of $\underline{y} = \underline{P}^T \underline{x}$, i.e., is done by the following procedure:

```

procedure ptx(m,n,a,x,y);
value m,n; integer n,m;
array x, y;
begin integer i,k; real s;

  for i := 1 step 1 until m do y[i] := x[i];
  t: for i := 1 step 1 until n do
    begin
      s := 0;
      for k := 1 step 1 until m do s := s + a[k,i] x[k];
      for k := 1 step 1 until m do y[k] := y[k] - a[k,i] s;
    end i;
  end

```

If we wish to compute $\underline{y} = \underline{P} \underline{x}$, the only change we have to make in the above procedure is

```
t : for i := n step -1 until 1 do
```

The following procedure `bidiā` bidiagonalizes $\underline{A}$, i.e., computes the **diagonal** $\underline{q}$ and the superdiagonal $\underline{e}$ of $\underline{B}$ and stores the transformation vectors $\underline{w}_j$ and $\underline{v}_{-j}$ in $\underline{A}$ (5.17):

```
procedure bidiā(m,n,a,q,e);
  integer m,n; _____ m,n;
  array a,q,e;
  begin
    integer i,j,k; real s, fak;
    for i := 1 step 1 until n do
      begin
        comment transform (ai1,...,am1) to (q1,0,...,0);
        s := 0;
        for j := i step 1 until m do s := s + a[j,i]2;
        if s = 0 then q[i] := 0 else
          begin
            s := sqrt(s);
            q[i] := if a[i,i] > 0 then -s else s;
            fak := sqrt(s x (s + abs(a[i,i])));
            a[i,i] := a[i,i] - q[i];
            for k := i step 1 until m do a[k,i] := a[k,i]/fak;
            for j := i+1 step 1 until n do
```

```
      begin
        s := 0;
        for k := i step 1 until m do s := s + a[k,i] x a[k,j];
        fork := i step 1 until m do a[k,j] := a[k,j] - a[k,i] x s;
      end j;
    end s;
    comment transform (ai,i+1,...,ain) to (e1,0,...,0);
    if i = n then goto ende;
    s := 0;
    for j := i+1 step 1 until n do s := s + a[i,j]2;
    if s = 0 then e[i] := 0 else
      begin
        s := sqrt(s);
        e[i] := if a[i,i+1] > 0 then -s else s;
        fak := sqrt(s x (s + abs(a[i,i+1])));
        a[i,i+1] := a[i,i+1] - e[i];
        for k := i+1 step 1 until n do a[i,k] := a[i,k]/fak;
        for j := i+1 step 1 until m do
          begin
            s := 0;
            for k := i+1 step 1 until n do s := s + a[j,k] x a[i,k];
            for k := i+1 step 1 until n do a[j,k] := a[j,k] - a[i,k] x s;
          end j;
        end s;
      end i;
    ende;
  end bidiā;
```

The decomposition using `_bigdia` requires $\sim 2(m n^2 - n^3/3)$ multiplications.

If $m > n$ it is possible to bidiagonalize A even cheaper [26], [5 I, [8].

The idea is to transform $\underline{A}$ first to an upper triangular matrix R and

then bidiagonalize R . In this case the operation count is $\sim m n^2 + n^3$.

An alternative approach to **bidiagonalizing** $\underline{A}$ has been suggested by Golub and Kahan [14]. This algorithm uses the matrix $\underline{A}$ not explicitly. Only

the operator $\underline{A} \underline{x}$ is needed. Unfortunately it is not numerically stable but nevertheless it seems to be useful for sparse matrices [31].

5.4 Computation of the Derivatives of the Length Function.

If we want to solve the secular equation

$$f(\lambda) = \alpha^2 \quad (5.18)$$

using some high order iteration method we have to compute the derivatives

of f . For the general problem (P1) we have

$$(\underline{A}^T \underline{A} + \lambda \underline{C}^T \underline{C}) \underline{x}^{(k)} = \underline{A}^T \underline{b} + \lambda \underline{C}^T \underline{d} \quad (5.19)$$

$$f(\lambda) = \|\underline{C} \underline{x}(\lambda) - \underline{d}\|^2.$$

By differentiating (5.19) we get

$$(\underline{A}^T \underline{A} + \lambda \underline{C}^T \underline{C}) \underline{x}' = -\underline{C}^T (\underline{C} \underline{x} - \underline{d}) \quad (5.20)$$

$$(\underline{A}^T \underline{A} + \lambda \underline{C}^T \underline{C}) \underline{x}'' = -2 \underline{C}^T \underline{C} \underline{x}'.$$

In general we have for $k \geq 2$

$$(\underline{A}^T \underline{A} + \lambda \underline{C}^T \underline{C}) \underline{x}^{(k)} = -k \underline{C}^T \underline{C} \underline{x}^{(k-1)}. \quad (5.21)$$

Lemma 5.1. If $\underline{B}_\lambda := \underline{A}^T \underline{A} + \lambda \underline{C}^T \underline{C}$ and $\underline{x}^{(k)}$ is the solution of (5.21)

then

$$\underline{x}^{(k)T} \underline{B}_\lambda \underline{x}^{(k)} = \begin{cases} \underline{x}^T (\underline{A}^T \underline{b} + \lambda \underline{C}^T \underline{d}), & k = 0 \\ -\underline{x}'^T \underline{C}^T (\underline{C} \underline{x} - \underline{d}), & k = 1 \\ -k \underline{x}^{(k)T} \underline{C}^T \underline{C} \underline{x}^{(k-1)}, & k \geq 2 \end{cases}$$

$$\underline{x}^{(k)T} \underline{B}_\lambda \underline{x}^{(k)} = \begin{cases} \frac{1}{2} \underline{x}^T \underline{B}_\lambda \underline{x}'' + \underline{x}^T \underline{C}^T \underline{d}, & k = 1 \\ \frac{k}{k+1} \underline{x}^{(k-1)T} \underline{B}_\lambda \underline{x}^{(k+1)}, & k \geq 2 \end{cases}$$

$$\|\underline{C} \underline{x}^{(k)}\|^2 = \begin{cases} -\underline{x}^T \underline{B}_\lambda \underline{x}' + \underline{x}^T \underline{C}^T \underline{d}, & k = 0 \\ -\frac{1}{k+1} \underline{x}^{(k)T} \underline{B}_\lambda \underline{x}^{(k+1)}, & k > 1 \end{cases}$$

$$\|\underline{C} \underline{x}^{(k)}\| = \begin{cases} \frac{1}{2} \underline{x}^{(2)T} \underline{C}^T (\underline{C} \underline{x} - \underline{d}), & k = 1 \\ \frac{k}{k+1} \underline{x}^{(k+1)T} \underline{C}^T \underline{C} \underline{x}^{(k-1)}, & k > 2 \end{cases}$$

$$\frac{d}{d\lambda} (\underline{x}^{(k)T} \underline{B}_\lambda \underline{x}^{(k)}) = \begin{cases} -\|\underline{C} \underline{x} - \underline{d}\|^2 + \|\underline{d}\|^2, & k = 0 \\ -(2k+1) \|\underline{C} \underline{x}^{(k)}\|^2, & k \geq 1 \end{cases}$$

$$\frac{d}{d\lambda} \|\underline{C} \underline{x}^{(k)}\|^2 = \begin{cases} -2 \underline{x}'^T \underline{B}_\lambda \underline{x}' + 2 \underline{x}'^T \underline{C}^T \underline{d}, & k = 0 \\ -\frac{2}{k+1} \underline{x}^{(k+1)T} \underline{B}_\lambda \underline{x}^{(k+1)}, & k > 1 \end{cases}$$

Proof. Equations (i) follow, by multiplying (5.19), (5.20), (5.21) by $\underline{x}^{(k)T}$ from the left.

Looking at two consecutive equations of (5.21):

$$\underline{B}_\lambda \underline{x}^{(k)} = -k \underline{C}^T \underline{C} \underline{x}^{(k-1)} \quad (5.22)$$

$$\underline{B}_\lambda \underline{x}^{(k+1)} = -(k+1) \underline{C}^T \underline{C} \underline{x}^{(k)} \quad (5.23)$$

We obtain (ii) by multiplying the first (5.22) by $\underline{x}^{(k)T}$ and the second by $\underline{x}^{(k-1)T}$ and by eliminating the expression $\underline{x}^{(k)T} \underline{C}^T \underline{C} \underline{x}^{(k-1)}$. To prove (iii) we multiply the equation for $\underline{x}^{(k+1)}$ of (5.21) by $\underline{x}^{(k)}$ and (5.20) by $\underline{x}^T$. Equation (iv) follows by multiplying (5.22) by $\underline{x}^{(k+1)T}$ and (5.23) by $\underline{x}^{(k)}$ and subtracting both equations.

Finally observe that $\underline{B}'_\lambda = \underline{C}^T \underline{C}$, therefore

$$\begin{aligned} \frac{d}{d\lambda} (\underline{x}^{(k)T} \underline{B}'_\lambda \underline{x}^{(k)}) \\ = \underline{x}^{(k+1)T} \underline{B}_\lambda \underline{x}^{(k)} + \underline{x}^{(k)T} \{ \underline{B}'_\lambda \underline{x}^{(k)} + \underline{B}_\lambda \underline{x}^{(k+1)} \} \\ = 2 \underline{x}^{(k+1)T} \underline{B}_\lambda \underline{x}^{(k)} + \|\underline{C}\underline{x}^{(k)}\|^2 \end{aligned}$$

Now for $k = 0$ using (iii) we have

$$\frac{d}{d\lambda} (\underline{x}^T \underline{B}_\lambda \underline{x}) = -\|\underline{C}\underline{x}\|^2 + 2 \underline{x}^T \underline{C}^T \underline{d} = -\|\underline{C}\underline{x} - \underline{d}\|^2 + \|\underline{d}\|^2.$$

For $k \geq 1$ using (iii) we get

$$\frac{d}{d\lambda} (\underline{x}^{(k)T} \underline{B}_\lambda \underline{x}^{(k)}) = -(2k+1) \|\underline{C}\underline{x}^{(k)}\|^2.$$

Using (i) we can write

$$\frac{d}{d\lambda} \|\underline{C}\underline{x}^{(k)}\|^2 = -\frac{2}{k+1} \underline{x}^{(k+1)T} \underline{B}_\lambda \underline{x}^{(k+1)}, \quad \square$$

Corollary 5.1. If $\underline{x}$ is the solution of (5.19) then

$$\left. \begin{aligned} \int f(\lambda) d\lambda &= \lambda \|\underline{d}\|^2 - \underline{x}^T (\underline{A}^T \underline{A} + \lambda \underline{C}^T \underline{C}) \underline{x} + \text{const.} \\ &= -\underline{A} \underline{d}^T (\underline{C}\underline{x} - \underline{d}) - \underline{b}^T \underline{A}\underline{x} + \text{const.} \end{aligned} \right\} \quad (5.24)$$

Proof. From equation (v) of Lemma 5.1 we have

$$f(\lambda) = \|\underline{C}\underline{x} - \underline{d}\|^2 = \|\underline{d}\|^2 - \frac{d}{d\lambda} (\underline{x}^T \underline{B}_\lambda \underline{x}).$$

By integrating and using (5.19) the result follows immediately. $\square$

Corollary 5.2. If $\underline{B}_\lambda = \underline{A}^T \underline{A} + \lambda \underline{C}^T \underline{C}$ then for $k \geq 1$,

$$(i) \quad \frac{d}{d\lambda} (\underline{x}^{(k)T} \underline{B}_\lambda \underline{x}^{(k)}) = \frac{2k+1}{k+1} \underline{x}^{(k)T} \underline{B}_\lambda \underline{x}^{(k+1)}$$

$$(ii) \quad \frac{d}{d\lambda} (\underline{x}^{(k)T} \underline{B}_\lambda \underline{x}^{(k+1)}) = 2 \underline{x}^{(k+1)T} \underline{B}_\lambda \underline{x}^{(k+1)}$$

Corollary 5.3. If $\underline{x}^{(k)}$ is a solution of (5.21), then for $k \geq 1$

$$(i) \quad \frac{d}{d\lambda} \|\underline{C}\underline{x}^{(k)}\|^2 = 2 \underline{x}^{(k+1)T} \underline{C}^T \underline{C} \underline{x}^{(k)}$$

$$(ii) \quad \frac{d}{d\lambda} (\underline{x}^{(k+1)T} \underline{C}^T \underline{C} \underline{x}^{(k)}) = \frac{2k+3}{k+1} \|\underline{C}\underline{x}^{(k+1)}\|^2.$$

Theorem 5.1. Let $\underline{x}$, $\underline{x}'$, and $\underline{x}^{(k)}$ be solutions of (5.19), (5.20), and (5.21) and let $\underline{B}_\lambda = \underline{A}^T \underline{A} + \lambda \underline{C}^T \underline{C}$. Then

$$f(A) = \|\underline{C}\underline{x} - \underline{d}\|^2 = -\underline{x}^T \underline{B}_\lambda \underline{x}' - \underline{d}^T (\underline{C}\underline{x} - \underline{d}) \quad (5.25)$$

$$f'(A) = 2 \underline{x}^T \underline{C}^T (\underline{C}\underline{x} - \underline{d}) = -2 \underline{x}^T \underline{B}_\lambda \underline{x}' \quad (5.26)$$

and for $k \geq 1$

$$f^{(2k)}(\lambda) = (k+1)\gamma_{2k} \|\bar{C}\bar{x}^{(k)}\|^2 = -\gamma_{2k} \bar{x}^{(k)T} \bar{B}_\lambda \bar{x}^{(k+1)} \quad (5.27)$$

$$\begin{aligned} f^{(2k+1)}(\lambda) &= (k+1)\gamma_{2k+1} \bar{x}^{(k)T} \bar{C}^T \bar{C} \bar{x}^{(k+1)} \\ &= -\gamma_{2k+1} \bar{x}^{(k+1)T} \bar{B}_\lambda \bar{x}^{(k+1)} \end{aligned} \quad (5.28)$$

where

$$\left. \begin{aligned} \gamma_{2k} &= \frac{1 \cdot 3 \cdot 5 \cdots (2k+1)}{(k+1)!} 2^k \\ \gamma_{2k+1} &= \frac{1 \cdot 3 \cdot 5 \cdots (2k+1)}{(k+1)!} 2^{k+1} \end{aligned} \right\} \quad (5.29)$$

Proof. The proof is by induction. If

$$f(\lambda) = \|\bar{C}\bar{x} - \bar{d}\|^2$$

then by differentiating and by Lemma 5.1 (i) we have

$$f'(A) = 2 \bar{x}'^T \bar{C}^T (\bar{C}\bar{x} - \bar{d}) = -2 \bar{x}'^T \bar{B}_\lambda \bar{x}'$$

Differentiating again using (v) and (iii) of Lemma 5.1 we get

$$f''(A) = 2 \cdot 3 \|\bar{C}\bar{x}'\|^2 = -3 \bar{x}'^T \bar{B}_\lambda \bar{x}''$$

which is (5.27) for $k = 1$. Now we can use Corollary 5.2 and 5.3 to **compute** the higher derivatives.

Assume

$$f^{(2k)}(\lambda) = (k+1)\gamma_{2k} \|\bar{C}\bar{x}^{(k)}\|^2 = -\gamma_{2k} \bar{x}^{(k)T} \bar{B}_\lambda \bar{x}^{(k+1)}.$$

Then

$$\begin{aligned} f^{(2k+1)}(\lambda) &= (k+1)\gamma_{2k} \cdot 2 \bar{x}^{(k)T} \bar{C}^T \bar{C} \bar{x}^{(k+1)} \\ &= -\gamma_{2k} \cdot 2 \bar{x}^{(k+1)T} \bar{B}_\lambda \bar{x}^{(k+1)} \end{aligned}$$

which is (5.28) for $\gamma_{2k+1} = 2\gamma_{2k}$. Again differentiating using Corollary 5.2 and 5.3 yields

$$\begin{aligned} f^{(2k+2)}(\lambda) &= \gamma_{2k+1} (2k+3) \|\bar{C}\bar{x}^{(k+1)}\|^2 \\ &= -\gamma_{2k+1} \frac{2k+3}{k+2} \bar{x}^{(k+1)T} \bar{B}_\lambda \bar{x}^{(k+2)} \end{aligned}$$

which is (5.27) for $k+1$ and $\gamma_{2k+2} = \gamma_{2k+1} \frac{2k+3}{k+2}$. Prom this recursion for γ_i it is easy to verify (5.29). $\square$

Theorem 5.1 shows that we can compute cheaply derivatives of f . To compute $x^{(k)}$ we have to solve a linear system with the **same** matrix $\bar{B}_\lambda$ as for x . Therefore we can use a factorization of $\bar{B}_\lambda$, If $A > 0$ then $\bar{x}, \bar{x}', \dots, \bar{x}^{(k)}$ is a solution of the least squares problem

$$\begin{pmatrix} \bar{A} \\ \sqrt{\lambda} \bar{C} \end{pmatrix} \bar{x} \approx \begin{pmatrix} \bar{b} \\ \sqrt{\lambda} \bar{d} \end{pmatrix}$$

$$\begin{pmatrix} \bar{A} \\ \sqrt{\lambda} \bar{C} \end{pmatrix} \bar{x}' \approx -\frac{1}{\sqrt{\lambda}} \begin{pmatrix} \bar{0} \\ \bar{C}\bar{x} - \bar{d} \end{pmatrix}$$

$$\begin{pmatrix} \bar{A} \\ \sqrt{\lambda} \bar{C} \end{pmatrix} \bar{x}^{(k)} \approx -\frac{k}{\sqrt{\lambda}} \begin{pmatrix} \bar{0} \\ \bar{C}\bar{x}^{(k-1)} \end{pmatrix}, \quad k > 2.$$

If we transform by an orthogonal matrix $\underline{\mathbf{G}}$

$$\underline{\mathbf{G}} \begin{pmatrix} \underline{\mathbf{A}} \\ \sqrt{\lambda} \underline{\mathbf{C}} \end{pmatrix} = \begin{pmatrix} \underline{\mathbf{R}}_{\lambda} \\ \underline{\mathbf{0}} \end{pmatrix}, \quad \underline{\mathbf{R}}_{\lambda} \text{ upper triangular,}$$

then $\underline{\mathbf{R}}_{\lambda}^T \underline{\mathbf{R}}_{\lambda}$ is the **Cholesky decomposition** of $\underline{\mathbf{R}}_{\lambda} = \underline{\mathbf{A}}^T \underline{\mathbf{A}} + \lambda \underline{\mathbf{C}}^T \underline{\mathbf{C}}$. To solve

$$\underline{\mathbf{R}}_{\lambda} \underline{\mathbf{x}}^{(\mathbf{k})} = -\mathbf{k} \underline{\mathbf{C}}^T \underline{\mathbf{C}} \underline{\mathbf{x}}^{(\mathbf{k}-1)}$$

we have two possibilities

(1) **compute $\underline{\mathbf{y}}^{(\mathbf{k})}$** by forward substitution **from**

$$\underline{\mathbf{R}}_{\lambda}^T \underline{\mathbf{y}}^{(\mathbf{k})} = -\mathbf{k} \underline{\mathbf{C}}^T (\underline{\mathbf{C}} \underline{\mathbf{x}}^{(\mathbf{k}-1)});$$

(2) compute $\underline{\mathbf{y}}^{(\mathbf{k})}$ using $\underline{\mathbf{G}}$

$$\begin{pmatrix} \underline{\mathbf{y}}^{(\mathbf{k})} \\ \underline{\mathbf{k}} \end{pmatrix} = \underline{\mathbf{G}} \begin{pmatrix} \underline{\mathbf{0}} \\ -\frac{\mathbf{k}}{\sqrt{\lambda}} \underline{\mathbf{C}} \underline{\mathbf{x}}^{(\mathbf{k}-1)} \end{pmatrix}.$$

Then we obtain $\underline{\mathbf{x}}^{(\mathbf{k})}$ by back substitution in

$$\underline{\mathbf{R}}_{\lambda} \underline{\mathbf{x}}^{(\mathbf{k})} = \underline{\mathbf{y}}^{(\mathbf{k})}.$$

If we define $\underline{\mathbf{z}}^{(\mathbf{k})} := \underline{\mathbf{C}} \underline{\mathbf{x}}^{(\mathbf{k})}$ then

$$\underline{\mathbf{f}}^{(2\mathbf{k}-1)}(\lambda) = \mathbf{k} \gamma_{2\mathbf{k}-1} \underline{\mathbf{z}}^{(\mathbf{k}-1)T} \underline{\mathbf{z}}^{(\mathbf{k})}$$

or equivalently

$$\underline{\mathbf{f}}^{(2\mathbf{k}-1)}(\lambda) = -\gamma_{2\mathbf{k}-1} \|\underline{\mathbf{y}}^{\mathbf{k}}\|^2.$$

Similarly

...

$$\underline{\mathbf{f}}^{(2\mathbf{k})}(\lambda) = (\mathbf{k}+1) \gamma_{2\mathbf{k}} \|\underline{\mathbf{z}}^{(\mathbf{k})}\|^2$$

or if $\underline{\mathbf{y}}^{(\mathbf{k}+1)}$ has been computed

$$\underline{\mathbf{f}}^{(2\mathbf{k})}(\lambda) = -\gamma_{2\mathbf{k}} \underline{\mathbf{y}}^{(\mathbf{k})T} \underline{\mathbf{y}}^{(\mathbf{k}+1)}.$$

In Section 6 we shall consider third order iteration methods to solve the secular equation for $\lambda > 0$. We need therefore the values of $\underline{\mathbf{f}}$, $\underline{\mathbf{f}}'$, and $\underline{\mathbf{f}}''$, which can be computed as follows:

1. **Compute $\underline{\mathbf{G}}$** and $\underline{\mathbf{R}}_{\lambda}$ so that

$$\underline{\mathbf{G}} \begin{pmatrix} \underline{\mathbf{A}} \\ \sqrt{\lambda} \underline{\mathbf{C}} \end{pmatrix} = \begin{pmatrix} \underline{\mathbf{R}}_{\lambda} \\ \underline{\mathbf{0}} \end{pmatrix}.$$

2. **Compute $\underline{\mathbf{y}}$** from

$$\underline{\mathbf{R}}_{\lambda}^T \underline{\mathbf{y}} = \underline{\mathbf{A}}^T \underline{\mathbf{b}} + \lambda \underline{\mathbf{C}}^T \underline{\mathbf{d}}$$

or using $\underline{\mathbf{G}}$ by

$$\begin{pmatrix} \underline{\mathbf{y}} \\ \underline{\mathbf{h}} \end{pmatrix} = \underline{\mathbf{G}} \begin{pmatrix} \underline{\mathbf{b}} \\ \sqrt{\lambda} \underline{\mathbf{d}} \end{pmatrix}.$$

3. **Compute $\underline{\mathbf{x}}$** from $\underline{\mathbf{R}}_{\lambda} \underline{\mathbf{x}} = \underline{\mathbf{y}}$ and form $\underline{\mathbf{z}} := \underline{\mathbf{C}} \underline{\mathbf{x}} - \underline{\mathbf{d}}$.

4. $\underline{\mathbf{f}}(\lambda) = \|\underline{\mathbf{z}}\|^2$.

5. Compute $\underline{\mathbf{y}}'$ by solving

$$\underline{\mathbf{R}}_{\lambda}^T \underline{\mathbf{y}}' = -\underline{\mathbf{C}}^T \underline{\mathbf{z}}$$

or by using $\underline{\mathbf{G}}$

$$\begin{pmatrix} \underline{y}' \\ \underline{h}' \end{pmatrix} := \underline{G} \begin{pmatrix} 0 \\ -\frac{1}{\sqrt{\lambda}} \underline{z} \end{pmatrix}$$

$$6. \quad f'(A) := -2 \|\underline{y}'\|^2 \quad .$$

$$7. \quad \text{Compute } \underline{x}' \text{ from } \underline{R}_{\lambda} \underline{x}' = \underline{y}' \quad \text{and form } \underline{z}' := \underline{C} \underline{x}' \quad .$$

$$8. \quad f''(A) = 6 \|\underline{z}'\|^2 \quad .$$

If we do not want to store $\underline{g}$ then we have to compute $\underline{y}'$ in step 5 by forward substitution. However, $\underline{y}$ in step 2 can be computed together with the decomposition in step 1 without forming $\underline{g}$ explicitly. Eldén gives in [8] similar recursions to compute the derivatives.

6. One-point Iteration Methods to Solve the Secular Equation.

In this chapter we discuss how to find the solution $A^* > 0$ of the secular equation

$$f(\lambda) = \alpha^2 \tag{6.1}$$

that we need to solve problem (P1), (P2) or (P3). The length function f is a positive rational function for all A and decreasing in $(0,\infty)$.

If we start with $A = 0$ Newton's method will produce a strictly

increasing sequence $\{\lambda_n\}$ which converges globally to A^* . However

as it was observed in [34] if α is small, convergence is slow and as a remedy Reinsch suggested solving the equation

$$\frac{1}{\sqrt{f}} - \frac{1}{\alpha} = 0 \tag{6.2}$$

instead of $\sqrt{f} - \alpha = 0$ using Newton's method. Indeed convergence is much better in this case. There are several possible interpretations and explanations for this fact. For our purpose we compare the Newton step resulting for the three equations

$$g_1(\lambda) := f(A) - \alpha^2 = 0$$

$$g_2(\lambda) := \sqrt{f(\lambda)} - \alpha = 0$$

$$g_3(\lambda) := \frac{1}{\sqrt{f(\lambda)}} - \frac{1}{\alpha} = 0 \quad .$$

A short calculation yields the following Newton iteration functions for the three equations:

$$\lambda = \frac{f - \alpha^2}{f'} \quad \text{for } \mathcal{E}_1, \quad (6.3)$$

$$\lambda = \frac{f - \alpha^2}{f'} \frac{2}{1 + \frac{\alpha}{\sqrt{f}}} \quad \text{for } \mathcal{E}_2, \quad (6.4)$$

$$\lambda = \frac{f - \alpha^2}{f'} \frac{2}{1 + \frac{\alpha}{\sqrt{f}}} \frac{\sqrt{f}}{\alpha} \quad \text{for } \mathcal{E}_3. \quad (6.5)$$

For $\lambda = 0$ we typically have $f(0) \gg \alpha^2$. Since $f' < 0$ for $A > 0$ the step in (6.4) will be about twice as big as in (6.3). But (6.5) will produce an even larger step, proportional to $\frac{\sqrt{f}}{\alpha}$ the discrepancy of $\sqrt{f}$ and α . If $f \approx \alpha^2$ then all the three steps are of about the same size.

We can look at the two functions

$$h_1(\lambda) = 2 / \left(1 + \frac{\alpha}{\sqrt{f(\lambda)}} \right)$$

$$h_2(\lambda) = h_1(\lambda) \frac{\sqrt{f}}{\alpha}$$

as functions that help to accelerate the global convergence of (6.3) by preserving the order (all iterations are Newton sequences and of second order).

6.1 Convergence Factors.

For the following discussion we change notation. Let f be a given function. We are looking for a number s such that $f(s) = 0$. We assume

that f has sufficient continuous derivatives in a neighborhood of s and furthermore we assume that s is a simple zero of f . We consider one point iteration methods without memory [43]

$$\left\{ \begin{array}{l} x_0 \text{ arbitrary} \\ x_{n+1} = F(x_n) \end{array} \right\} \quad n = 0, 1, \dots \quad (6.6)$$

where

$$F(x) = x - \frac{f(x)}{f'(x)}. \quad (6.7)$$

and $G(x)$ is an appropriate chosen function which we will call the "convergence factor". The idea is to choose G so that the global convergence using (6.7) is better than Newton's iteration ($G(x) \equiv 1$). As has been pointed out by Kahan [24], every sequence generated by an iteration (6.6) can be interpreted as a sequence obtained by applying Newton's method to a certain equation

$$g(x) = 0. \quad (6.8)$$

Indeed if we put

$$x - \frac{g(x)}{g'(x)} = F(x)$$

a short calculation yields (with some $c \neq 0$)

$$g(x) = c \cdot \exp \left(\int \frac{dx}{x - F(x)} \right). \quad (6.9)$$

Solving (6.8) with (6.9) using Newton's method yields the sequence (6.6). Some examples may illustrate the point.

(1) Let $x_{n+1} = 1 + \frac{x}{2}$, $x_0 = 0$. This sequence converges

linearly to $s = 2$. Indeed using (6.9) we get

$$g(x) = \exp\left(\int dx / \left(\frac{x}{2} - 1\right)\right) = \left(\frac{x}{2} - 1\right)^2.$$

$g(x)$ has a double zero $s = 2$ and therefore Newton's iteration converges linearly.

(2) Consider the Halley iteration formula $x_{n+1} = F(x_n)$ with

$$F(x) = x - \frac{2 f(x) f'(x)}{2 (f'(x))^2 - f''(x) f(x)}.$$

using (6.9) we obtain

$$g(x) = \exp\left(\int \left(\frac{f(x)}{f'(x)} - \frac{f''(x)}{2f'(x)}\right) dx\right) = \frac{f(x)}{\sqrt{f'(x)}}.$$

Therefore using Halley's method for $f(x) = 0$ is applying Newton's method to

$$g(x) = \frac{f(x)}{\sqrt{f'(x)}} = 0, \quad [3].$$

(3) Consider the iteration

$$x_{n+1} = x_n - \frac{f(x_n) - \alpha}{f'(x_n)} \cdot \frac{f'(x_n)}{\alpha} \quad (6.10)$$

to solve $f(x) = \alpha$. Using (6.9) we get

$$g(x) = 1 - \frac{\alpha}{f(x)}$$

or since g is only determined to a constant we may also divide by $-a$ and get

$$g(x) = \frac{1}{f'(x)} - \frac{1}{\alpha}.$$

Therefore (6.10) is obtained by applying Newton's method to

$$\frac{1}{f'(x)} - \frac{1}{\alpha} = 0$$

instead of $f(x) - \alpha = 0$.

For the class of iteration functions $F(x)$ (6.7) we would like to consider however if it will not be possible in general to evaluate the integral for g in (6.9) explicitly and thus provide a nice explanation of the iteration function.

The order of convergence [22] of an iteration formula

$$x_{n+1} = F(x_n)$$

is

one (or linear convergence), if $|F'(s)| < 1$
 two (or quadratic convergence), if $F'(s) = 0$
 three (or cubic convergence), if $F'(s) = F''(s) = 0$
 m , if $F'(s) = F''(s) = \dots = F^{(m-1)}(s) = 0$.

For $G(x) \equiv 1$ the iteration with

$$F(x) = x - \frac{f(x)}{f'(x)} \cdot G(x) \quad (6.11)$$

is Newton's method and it is of second order for a zero of f with multiplicity one. Differentiation (6.11) gives

$$F'(x) = 1 - \left(\frac{f(x)}{f'(x)}\right)' G(x) - \frac{f(x)}{f'(x)} G'(x).$$

Since $f(s) = 0$ we have $F'(s) = 0$ if $G(s) = 1$ and $f(s) \cdot G'(s) = 0$.

Therefore we have

Lemma 6.1. Let G be differentiable, $G(s) = 1$, $f(s) \cdot G'(s) = 0$.

Then the iteration

$$x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)} G(x_n)$$

is of second order for simple zeros s of f . Since s is unknown we have to choose

$$G(x) = H(f(x), f'(x), \dots).$$

Examples:

$$(4) \quad G(x) = H(f) = 1 + f(x)$$

$$\Rightarrow F(x) = x - \frac{f(x)}{f'(x)} (1 + f(x)).$$

(6.13) yields a second order iteration formula. Indeed using (6.9)

we see that the iteration is obtained solving $g(x) = 0$ with Newton

where

$$g(x) = \exp\left(\int \frac{f'(x)dx}{f(x)(1+f(x))}\right) = \frac{1}{1+f(x)} - 1.$$

(6.13) is a special case ($a = 1$) of

$$(5) \quad H(f) = \frac{\alpha + f}{\alpha}$$

for some $\alpha \neq 0$. Therefore the iteration

$$x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)} \frac{\alpha + f(x_n)}{\alpha}$$

is the same as applying Newton's method to

$$g(x) = \frac{1}{f(x)} + \alpha - \frac{1}{\alpha} = 0, \quad \alpha \neq 0.$$

$$(6) \quad G(x) = 1 / \left(1 - \frac{1}{2} \frac{f(x)f''(x)}{(f'(x))^2} \right).$$

This choice is Halley's method. Again we have $G(s) = 1$.

6.2 Third Order Iterative Methods.

We consider again the iteration

$$x_{n+1} = F(x_n) = x_n \frac{f(x_n)}{f'(x_n)} \cdot G(x_n). \quad (6.15)$$

In Section 6.1 we saw that if $G(s) = 1$ then (6.15) is quadratic convergent. Let $u(x) := \frac{f(x)}{f'(x)}$. Then

$$F(x) = x - u(x), \dots \quad (6.16)$$

and we can ask for the conditions that the iteration (6.15) be cubically convergent. Differentiating (6.16) we get (dropping the argument)

$$\begin{aligned} F' &= 1 - u'G - uG' \\ F'' &= -u''G - 2u'G' - uG'' \\ u &= f/f' \\ u' &= 1 - \frac{ff''}{f'^2} \\ u'' &= -\frac{f'''}{f'} + 2 \frac{ff'''}{f'^3} - \frac{f'f'''}{f'^2} \end{aligned}$$

Now we want to have

$$F'(s) = F''(s) = 0. \quad (6.17)$$

Since $u(s) = 0$, $u'(s) = 1$ and $u''(s) = -\frac{f''(s)}{f'(s)}$, this is the case if

$$G(s) = 1 \text{ and } G'(s) = \frac{1}{2} \frac{f''(s)}{f'(s)}. \quad (6.18)$$

Theorem 6.1. Let $H \in C^2[-a, a]$ for some $a > 0$. The iteration $x_{n+1} = F(x_n)$ with

$$F(x) = x - \frac{f(x)}{f'(x)} \cdot H\left(\frac{f(x)f''(x)}{(f'(x))^2}\right) \quad (6.19)$$

converges **cubically** to a simple zero of f if and only if

$$H(0) = 1 \quad \text{and} \quad H'(0) = \frac{1}{2}.$$

Proof. The function F in (6.19) is the special case of (6.16) with

$$G(x) = H\left(\frac{f(x)f''(x)}{(f'(x))^2}\right).$$

Let $t(x) = \frac{f(x)f''(x)}{(f'(x))^2}$. Then $t(x) = 1 - u'(x)$ and therefore

$$t'(s) = -u''(s) = \frac{f''(s)^2}{f'(s)^3}.$$

Clearly $G(s) = 1 \Leftrightarrow H(0) = 1$ and

$$\begin{aligned} G'(s) &= H'(t(s))t'(s) \\ &= H'(0) \cdot \frac{f''(s)}{f'(s)} = \frac{1}{2} \frac{f''(s)}{f'(s)} \end{aligned}$$

$$\Leftrightarrow H'(0) = \frac{1}{2}. \quad \square$$

Many well known third order iterative methods are special cases of Theorem 6.1. Let

$$t := \frac{f(x)f''(x)}{(f'(x))^2} \quad (6.20)$$

(1) Euler's formula.

$$H(t) = \frac{2}{1 + \sqrt{1-2t}} = 1 + \frac{1}{2}t + \frac{1}{2}t^2 + \frac{5}{8}t^3 + \dots$$

clearly $H(0) = 1$ and $H'(0) = 1/2$.

(2) Halley's formula.

$$H(t) = \frac{1}{1 - \frac{1}{2}t} = 1 + \frac{1}{2}t + \frac{1}{4}t^2 + \frac{1}{8}t^3 + \dots$$

(3) Quadratic inverse interpolation [43].

$$H(t) = 1 + \frac{1}{2}t.$$

(4) Ostrowski's square root iteration [30].

$$H(t) = \frac{1}{\sqrt{1-t}} = 1 + \frac{1}{2}t + \frac{3}{8}t^2 + \frac{5}{16}t^3 + \dots$$

(5) Hansen-Patrick family [19].

$$H(t) = \frac{\alpha+1}{\alpha + \sqrt{1-(\alpha+1)t}} = 1 + \frac{1}{2}t + \frac{\alpha+3}{8}t^2 + \dots$$

(6) To solve the secular equation we will use

$$H(t) = e^{\frac{1}{2}t} = 1 + \frac{1}{2}t + \frac{1}{8}t^2 + \frac{1}{48}t^3 + \dots$$

This formula has the advantage that $H(t) > 0$ even if the starting point is far away from the solution s . For large t , the other formulas may

not work (wrong sign in H , argument of the square root is negative).

The following lemma is useful for construction of third order iteration methods.

Lemma 6.2. Let $H_1(t)$ and $H_2(t)$ be two functions with

$$H_1(0) = a \quad \text{and} \quad H_2(0) = b, \quad i = 1, 2.$$

Then the three mean functions

$$A = (H_1 + H_2)/2$$

$$B = \sqrt{H_1 H_2}$$

$$C = 2 / \left(\frac{1}{H_1} + \frac{1}{H_2} \right)$$

have the same property.

Proof. Obviously $A(0) = B(0) = C(0) = a$.

$$A' = (H_1' + H_2')/2 = b.$$

$$B' = (H_1' H_2 + H_1 H_2') / (2 \sqrt{H_1 H_2})$$

$$B'(0) = (b \cdot a + a \cdot b) / (2a) = b.$$

$$C = \frac{2 H_1 H_2}{H_1 + H_2}$$

$$C' = 2 \frac{(H_1 + H_2)(H_1' H_2 + H_1 H_2') - (H_1' + H_2') H_1 H_2}{(H_1 + H_2)^2}$$

$$C'(0) = 2 \frac{2a(2ab) - 2ba^2}{4a^2} = b. \quad \square$$

Any of the three means of the examples (1) to (6) yields a new third order iterative method.

Notice that not every third order iterative method must have the form (6.19). For example, consider

$$G(x) = H(t(x)) + x^2(x),$$

with $H(0) = 1$, $H'(0) = 1/2$, and $t(x)$ defined by (6.20). Clearly this choice of the convergence factor G leads to a third order formula which has not the form of Theorem 6.1.

However it is possible to describe all the third order iteration methods. Let

$$F_k := x + \sum_{i=1}^{k-1} \frac{(-1)^i}{i!} x^i \frac{1}{f^i} \frac{d^{i-1}}{dx} \frac{1}{f^i}$$

denote the Schröder iteration functions (see [7] or [21]). The iteration

$$x_{n+1} = F_k(x_n)$$

is of k -th order. Now every k -th order iteration function can be written as

$$F(x) = F_k(x) + x^k(x) \cdot \varphi(x)$$

where φ is an arbitrary function [7], [43]. Therefore for $k = 2$ we obtain the most general G for a third order method

$$G(x) = 1 + \frac{1}{2} t(x) + x^2(x) \cdot \varphi(x)$$

which we can write

$$G(x) = H(t(x)) + x^2(x) \cdot \psi(x)$$

with arbitrary ψ .

6.3 The Convergence Factor for a Third Order Method.

Assume that $x_{n+1} = x_n - K(x_n)$ is a third order iterative method for solving $f(x) = 0$. We clearly have

$$K(s) = K''(s) = 0, \quad K'(s) = 1. \quad (6.21)$$

Now consider the iteration $x_{n+1} = F(x_n)$ with

$$F(x) = x - K(x) \cdot G(x). \quad (6.22)$$

We have

$$\left. \begin{aligned} F' &= 1 - K'G - KG' \\ F'' &= -K''G - 2K'G' - KG'' \end{aligned} \right\} \quad (6.23)$$

Lemma 6.3. The iteration $x_{n+1} = F(x_n)$ with F defined in (6.22) is also of third order if $G \in C^2[a, b]$ with $SE(a, b)$ and

$$G(s) = 1, \quad G'(s) = 0. \quad (6.24)$$

Proof. If we use (6.21) and (6.24) in (6.23) then we have

$$F'(s) = F''(s) = 0. \quad \square$$

If we choose G so that

$$G''(s) = -\frac{1}{2} K''(s)$$

then we would have a fourth order iteration. However we think that a third order method with good global convergence is more **useful** than a local convergent fourth order formula.

The following functions are examples of possible convergence factors for a third order iteration:

$$(1) \quad G(x) = (t^\beta(x) + t^{-\beta}(x))/2$$

$$\text{where } t(x) = \frac{f(x) + \alpha}{\alpha} > 0$$

and $\alpha, \beta \in \mathbb{R}$.

$$(2) \quad G(x) = \cosh(f(x) \cdot r(x))$$

where $r \in C^2[a, b]$ with $s \in (a, b)$.

$$(3) \quad G(x) = L(f^2(x) \cdot r(x))$$

where $L \in C^2[-d, d]$ $d > 0$ and $L(0) = 1$,

$r \in C^2[a, b]$ with $s \in (a, b)$.

We know now how to choose convergence factors to preserve or increase the order of an iteration formula. Our aim is however to improve global convergence. Let

$$F(x) = x - K(x)G(x)$$

be the iteration function. Then G has to be chosen so that

$$g(x) = \exp \left(\int \frac{dx}{K(x) \cdot G(x)} \right) \quad (6.25)$$

is nearly linear if x is far away from s . Since the iteration

$x_{n+1} = F(x_n)$ is the same as applying Newton's method to $g(x) = 0$,

global convergence will be good if g is linear far away from the

solution s . However, since (6.25) may be impossible to compute

explicitly, this requirement seems not to be practical to determine G .

We have to estimate G by other means. In Section 6.4 we shall interpret

some G geometrically. Those functions can serve as models for others.

6.4 Geometrical Interpretation of the Convergence Factors.

It is possible to derive iterative methods as follows:

(i) choose a "simple" function h so that

$$f^{(i)}(x) = h^{(i)}(x), \quad i = 0, 1, \dots, k;$$

(ii) solve analytically $h(z) = \alpha$ obtaining $z = z(x)$;

(iii) use the iteration $x_{n+1} = z(x_n)$ to solve the equation $f(x) = a$.

The function h should be simple so that $h(z) = \alpha$ can be solved analytically. h approximates f and some derivatives locally at one point x and we thus obtain a one point iteration formula without memory.

Instead of approximating f one can also find iteration methods by approximating the inverse function locally.

(i) Choose a function $h(y)$ so that

$$(f^{[-1]}(y))^{(i)} = h^{(i)}(y), \quad i = 0, 1, \dots, k.$$

(ii) Put $y = f(x_n)$ and use the iteration formula

$$x_{n+1} = h(\alpha)$$

to solve $f(x) = \alpha$.

Since we do not know $f^{[-1]}$, the derivatives must be replaced using derivatives of f [21]:

$$(f^{[-1]}(y))' = \frac{1}{f'(f^{[-1]}(y))}$$

$$(f^{[-1]}(y))'' = -\frac{f''(f^{[-1]}(y))}{f'^2(f^{[-1]}(y))}.$$

If h approximates f resp $f^{[-1]}$ well we can expect a good global convergence. We give in the following some examples of methods derived by interpolation. The convergence factors of these examples may help to choose a method analytically.

(1) Newton's method $x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)}$ is obtained approximating f or $f^{[-1]}$ locally by a linear function h ($h(x) = ax+b$ resp. $h(y) = ay+b$).

(2) We choose $h(z) = \frac{a}{(z-b)^2}$ and determine a, b so that

$h^{(i)}(x) = f^{(i)}(x)$, $i = 1, 2$, giving (dropping the argument x):

$$\left. \begin{aligned} a &= 4 f^3 / f'^2 \\ b &= -x - 2 f / f' \end{aligned} \right\} \quad (6.26)$$

Now $h(z) = \alpha$ gives

$$z = -b \pm \sqrt{\frac{a}{\alpha}}$$

using (6.26) we get

$$z = x - \frac{f}{f'} \pm 2 \left(\pm \sqrt{\frac{f}{\alpha}} - 1 \right).$$

If we use only the + sign we have the iteration

$$\left. \begin{aligned} x_{n+1} &= x_n - \frac{f(x_n) - \alpha}{f'(x_n)} G(x_n) \\ \text{with } G(x) &= 2 \sqrt{\frac{f}{\alpha}} / \left(1 + \sqrt{\frac{\alpha}{f}} \right) \end{aligned} \right\} \quad (6.27)$$

But (6.27) is Reinsch's proposal to solve $\frac{1}{\sqrt{f}} - \frac{1}{\sqrt{\alpha}} = 0$ with Newton

instead of $f - a = 0$. We know from Section 3 that the length function f

has the form

$$f(x) = \sum_{i=1}^n c_i^2 (x - x_i)$$

Therefore it is reasonable to approximate f by $h = a/(x-b)^2$.

(3) Instead of solving $f(x) - \alpha = 0$ with Newton's method we can also solve

$$g(x) = (f(x))^{-1/\beta} - \alpha^{-1/\beta} \quad (6.28)$$

if $f(x) > 0$ and $a > 0$. Newton's iteration for (6.28) yields the same iteration function as if we approximate f and f' locally by

$$h(x) = \frac{a}{(x+b)^{-1/\beta}}.$$

The resulting iteration formula is

$$x_{n+1} = x_n - \frac{f(x_n) - \alpha}{f'(x_n)} \cdot G(x_n)$$

with

$$G(x) = \frac{\beta f}{f - \alpha} \left(\beta \sqrt{\frac{f}{\alpha}} - 1 \right). \quad (6.29)$$

If we let $\beta \rightarrow \infty$ in (6.29) we get

$$G(x) = \frac{f}{f - \alpha} \ln \left(\frac{f}{\alpha} \right). \quad (6.30)$$

This convergence factor is also obtained by solving $\ln(f) - \ln(\alpha) = \ln(f/\alpha) = 0$ with Newton or by approximating locally f and f' by

$$h(x) = ae^{bx}.$$

The next examples use functions that approximate f , f' and f'' locally.

(4) We choose $h(z) = a/(z+b) + c$ and determine a , b , c so that $h^{(i)}(x) = f^{(i)}(x)$, $i = 0, 1, 2$.

We obtain (dropping the argument x):

$$a = -4 f'^3 / f''^2$$

$$b = -x - 2 f' / f''$$

$$c = f - 2 f'^2 / f''.$$

Now solving $h(z) = \alpha$ gives the iteration (with $f_n^{(i)} = f^{(i)}(x_n)$)

$$x_{n+1} = x_n - \frac{f_n - \alpha}{f'_n} \frac{1}{1 - \frac{1}{2} \frac{(f_n - \alpha) f''_n}{f'^2_n}} \quad (6.31)$$

which is Halley's formula to solve $g(x) = f(x) - \alpha = 0$. Recall (see Section 6.1) that the same iteration is obtained if we solve with Newton's method

$$g(x) = \frac{f(x) - \alpha}{\sqrt{f'(x)}} = 0.$$

If we approximate the inverse function $f^{[-1]}(y)$ locally by

$h(y) = a/(y+b) + c$ we get

$$a = -4 f'^3 / f''^2$$

$$b = -f + 2 f'^2 / f''$$

$$c = x + 2 f' / f''.$$

Finally putting $x_n = x$ and $x_{n+1} = h(a)$ yields again (6.31). Therefore Halley's formula is obtained by locally approximating f or $f^{[-1]}$ by a hyperbola.

This rational approximation of f has the following property. Suppose we want to solve $f(x) = \alpha$ using Halley's method. We obtain the same iteration (6.31) for

$$g_1(x) = f(x) - \alpha = 0$$

or

$$g_2(x) = \frac{1}{f(x)} - \frac{1}{\alpha} = 0, \quad (6.32)$$

i.e., for Halley's method the transformation (6.32) has no effect. But as we saw (6.32) yield another convergence factor for Newton's method.

(5) **Euler's** method is obtained by approximating f , f' , and f'' locally with

$$h(z) = az^2 + bz + c.$$

We get

$$\begin{aligned} a &= f'/2, \quad b = f' - f''x \\ c &= f - f'x + \frac{f''}{2}x^2 \end{aligned}$$

and

$$\begin{aligned} x_{n+1} &= x_n - \frac{(f_n - \alpha)}{f'_n} \cdot G(x_n) \\ G(x) &= 2 / \left(1 + \sqrt{1 - 2 \frac{(f - \alpha)f''}{f'^2}} \right) \end{aligned}$$

with

(6) Approximating $f^{[-1]}$ by a parabola

$$h(y) = ay^2 + by + c,$$

yields

$$\begin{aligned} a &= -\frac{1}{2} \frac{f''}{f'f} \\ b &= \frac{1}{f'} + \frac{f''}{f'f} f \\ c &= x - \frac{f}{f'} - \frac{1}{2} \frac{f''}{f'} f^2. \end{aligned}$$

We obtain the iteration formula for $f(x)$ $-a = 0$

$$x_{n+1} = x_n - \frac{(f_n - \alpha)}{f'_n} \left(1 + \frac{1}{2} \frac{(f_n - \alpha)f''_n}{f'^2_n} \right)$$

(7) Approximating f locally by

$$h(z) = ae^{bz} + c$$

yields

$$\begin{aligned} a &= \frac{f'^2}{f''} \exp - \frac{f''}{f'} x \\ b &= f''/f' \\ c &= f - f'^2/f'' \end{aligned}$$

Solving $h(z) = \alpha$ gives the iteration

$$x_{n+1} = x_n + \frac{f'_n}{f''_n} \ln \left(1 - \frac{(f_n - \alpha)f''_n}{f'^2_n} \right)$$

which we can write

$$x_{n+1} = x_n + \frac{f_n - \alpha}{f'_n} G(x_n)$$

with

$$\begin{aligned} G(x) &= -\ln(1 - t(x))/t(x) \\ t(x) &= (f(x) - a) f''(x)/f'^2(x). \end{aligned}$$

From Theorem 6.1 we know that all the methods (4) - (7) are cubically convergent.

6.5 Solving the Secular Equation.

A very good and simple method to solve $f(h) = \alpha^2$ is Reinsch's

proposal

$$\lambda_{n+1} = \lambda_n - \frac{f(\lambda_n) - \alpha^2}{f'(\lambda_n)} g(\lambda_n) \quad (6.33)$$

where

$$g(\lambda) = 2 \frac{\sqrt{f(\lambda)}}{\alpha} / \left(1 + \frac{\alpha}{\sqrt{f(\lambda)}} \right).$$

Iteration (6.33) can also be written as

$$\lambda_{n+1} = \lambda_n - \frac{f(\lambda_n)}{f'(\lambda_n)} \cdot 2 \left(\frac{\sqrt{f(\lambda_n)}}{\alpha} - 1 \right)$$

Reinsch proved in [34] that the sequence λ_n obtained with this iteration converges starting with $\lambda_1 = 0$ monotonically **increasing** to the solution. **This** property together with good global convergence makes this iteration very attractive.

If we use another method with starting value $\lambda_1 = 0$, e.g. Halley's iteration, global convergence may not be guaranteed. As an example, **consider** problem (P3)

$$\min \|x\|$$

subject to

$$\|Ax - b\| \leq \alpha$$

where $n = m = 6$, A Hilbert matrix, $a_{i,j} = 1/(i+j-1)$, $b = (1, 0, 0, 0, 0, 0)^T$.

If we start with $\lambda_0 = 0.1$ we get $\lambda_1 = -0.9975$ instead of a value bigger than λ_0 . The other third order methods involving a term

$$1 - \frac{1}{\gamma} t \quad \text{or} \quad 1 - t, \quad t = \frac{f'''}{f'f'}$$

may also fail if $t > 1$. The quadratic inverse interpolation

$$H(t) = 1 + \frac{1}{\gamma} t \quad (6.34)$$

yields a correction with the right sign but numerical experiments show that if Halley converges, it converges globally much faster.

An improvement of (6.34) is to use

$$H(t) = \exp(t/\gamma). \quad (6.35)$$

However global convergence has to be accelerated by a convergence factor H .

Therefore we propose to use

$$\lambda_{n+1} = \lambda_n - \frac{f' - \alpha^2}{f''} \exp\left(\frac{1}{\gamma} \frac{f''''}{f''} \frac{1}{\lambda_n}\right) \cdot H(\lambda_n) \quad (6.36)$$

with

$$H(\lambda) = \frac{1}{\gamma} \left(\frac{\sqrt{f(\lambda)}}{\alpha} + \frac{1}{\sqrt{f(\lambda)}} \right).$$

Since we cannot prove monotonicity or global convergence properties, it is necessary to take measures against too big steps. If the step is too big (usually at the beginning) we use Reinsch's iteration (6.33). For the

above mentioned example with the **Hilbert** matrix we obtain using (6.36) and $\lambda_0 = 0.1$ the solution of the secular equation $\lambda \approx 7.7 \times 10^{11}$ in four iterations.

To illustrate the behavior of the different methods we **computed** the only positive solution of the equation

$$f(\lambda) = 1$$

$$f(\lambda) = 0.6 + \sum_{i=1}^{20} \frac{2 + 0.8^i}{(\lambda + 0.8^i)^2} \quad .$$

where

We used ALGOL W and double precision (ca. 16 decimal digits) on the **IBM 360 of SIAC**. The results are displayed in Table 1. Newton's method uses 21 steps to converge to machine precision. **Reinsch's** variant converges after seven **steps**. The improved Halley's **method** needs only four steps.

Table 1

Example for the behaviour of different methods to solve $g(\lambda) = f(\lambda) - \alpha = 0$.

Second order methods $\lambda_{n-1} = F(\lambda_n)$:

Newton	Reinsch	Newton for $\ln(f/\alpha) = 0$
$F(\lambda) = \lambda - \frac{g}{g'}$	$\lambda - \frac{f}{f'}, 2(\frac{f}{f'} - 1)$	$\lambda - \frac{f}{f'}, \ln(\frac{f}{\alpha})$
0.00783209855661240	3.20297243199116	0
0.0209101056943383	7.50002221633766	0.0634187403647334
0.0424438763172333	9.6945598729093	0.53074605218083
0.077467170580476	10.2422618393796	1.97352054359389
0.13329409609385	10.2699361392028	4.58739426667985
0.223664677428177	10.2700022239607	7.4796424219632
0.365485460269776	10.2700022243362	9.52785436035909
0.567031564547462	10.2700022243362	10.2135824785941
0.926977169376511	10.2700022243362	10.2696683866710
1.44940702417750	10.2700022243362	10.2700022126256
2.22759607120317	10.2700022243362	10.2700022243362
3.36604538155582	10.2700022243362	10.2700022243362
4.92591522457686	10.2700022243362	10.2700022243362
6.87134065977040	10.2700022243362	10.2700022243362
8.79350315070676	10.2700022243362	10.2700022243362
9.97230829940186	10.2700022243362	10.2700022243362
10.257425942621	10.2700022243362	10.2700022243362
10.2669795745767	10.2700022243362	10.2700022243362
10.2700022242627	10.2700022243362	10.2700022243362

Third order Methods $\lambda_{n+1} = \lambda_n - \frac{g}{g'}, H(t)$ with $t := \frac{g}{g'} \frac{g''}{g'^2}$

The following two methods fail for starting values < 6.8 :

Ostrowski	Euler
$H(t) = 1/(1-t)^{0.5}$	$2/(1 + (1-2t)^{0.5})$
6.87134069977040	6.87134069977040
11.3578564654989	8.79350313070076
10.2615901313624	10.4062426841626
10.2700022241005	10.269936114743
10.2700022243362	10.2700022243362
10.2700022243362	10.2700022243362
.	.

Table 1 (cont.)

Halley	inverse interp.	$\frac{1}{1 - t/2}$	$\exp(t/2)$
0.0494307764343190	0.0144232340488059	0	0.01617023288293991
0.253229265175140	0.0460967990981273	0	0.3030975341263060
0.986381372252434	0.113492541266652	0	0.169596942032953
3.21756490475624	0.253299532203121	0	0.413519282769038
7.85449586417966	0.536339290697594	0	0.951636742732932
10.2238730991803	1.05310528695540	2	2.06316964524166
10.2700019998135	2.14957041757370	4	4.26001192403909
10.270022243362	4.03642343012226	7	7.65236939221572
10.270022243362	6.91272434349204	10	10.0374630402832
10.270022243362	9.60092062566452	10	10.2698435997489
10.270022243362	10.2634961753614	10	10.2700022243362
10.270022243362	10.2700022160612	10	10.2700022243362
10.270022243362	10.2700022243362	10	10.2700022243362
10.270022243362	10.2700022243362	10	10.2700022243362

Third order methods with convergence factor $G = \left\{ \frac{\sqrt{F}}{\alpha} + \frac{\alpha}{\sqrt{F}} \right\} / 2 :$

Halley	$G / (1 - t/2)$	$G \exp(t/2)$
5.07846031436509	0	0
9.89047170050903	1.86673871095732	1.86673871095732
10.2699142055339	5.5130340570731	5.5130340570731
10.2700022543362	9.03602564637327	9.03602564637327
10.2700022243362	10.2593265907900	10.2593265907900
10.2700022243362	10.2700021410736	10.2700021410736
10.2700022243362	10.2700022243362	10.2700022243362

7. Generalizations.

Let A be an $(n \times n)$ matrix, b a given n vector, $\underline{c}$ an $m \times n$ matrix, d a given m vector and γ, α real numbers. We consider the problem

$$F(\underline{x}) = \underline{x}^T A \underline{x} + \underline{b}^T \underline{x} + \gamma = \min \quad (7.1)$$

subject to

$$\|\underline{C}\underline{x} - \underline{d}\| = a. \quad (7.2)$$

Using a Lagrange multiplier $\lambda/2$ the solution of (7.1), (7.2) is a stationary point $(\underline{x}, \lambda)$ of

$$L(\underline{x}, \lambda) = F(\underline{x}) + \frac{\lambda}{2} \{\|\underline{C}\underline{x} - \underline{d}\|^2 - \alpha^2\}. \quad (7.3)$$

$\frac{\partial L}{\partial \underline{x}} = 0$ and $\frac{\partial L}{\partial \lambda} = 0$ gives

$$\left. \begin{aligned} (\underline{A} + \underline{A}^T + \lambda \underline{C}^T \underline{C}) \underline{x} &= -\underline{b} + \lambda \underline{C}^T \underline{d} \\ \|\underline{C}\underline{x} - \underline{d}\|^2 &= \alpha^2 \end{aligned} \right\} \quad (7.4)$$

Theorem 7.1. Let $(\underline{x}_i, \lambda_i)$, $i = 1, 2$, be two solutions of (7.4), then

$$F(\underline{x}_1) - F(\underline{x}_2) = \frac{\lambda_2 - \lambda_1}{4} \|\underline{C}(\underline{x}_1 - \underline{x}_2)\|^2. \quad (7.5)$$

Proof. This theorem is a generalization of Theorem 2.1 and the proof is very similar. We have

$$(\underline{A} + \underline{A}^T + \lambda_1 \underline{C}^T \underline{C}) \underline{x}_1 = -\underline{b} + \lambda_1 \underline{C}^T \underline{d} \quad (7.6)$$

$$(\underline{A} + \underline{A}^T + \lambda_2 \underline{C}^T \underline{C}) \underline{x}_2 = -\underline{b} + \lambda_2 \underline{C}^T \underline{d} \quad (7.7)$$

$\underline{x}_1^T$ (7.6) gives

$$2 \mathbf{x}_1^T \mathbf{A} \mathbf{x}_1 + \lambda_1 \|\mathbf{C} \mathbf{x}_1\|^2 = -\mathbf{x}_1^T \mathbf{b} + \lambda_1 \mathbf{x}_1^T \mathbf{C}^T \mathbf{d} \quad ,$$

or rearranged

$$2(\mathbf{x}_1^T \mathbf{A} \mathbf{x}_1 + \mathbf{x}_1^T \mathbf{b}) = \mathbf{x}_1^T \mathbf{b} - \lambda_1 (\|\mathbf{C} \mathbf{x}_1\|^2 - \mathbf{x}_1^T \mathbf{C}^T \mathbf{d}) . \quad (7.8)$$

Similarly we have

$$2(\mathbf{x}_2^T \mathbf{A} \mathbf{x}_2 + \mathbf{x}_2^T \mathbf{b}) = \mathbf{x}_2^T \mathbf{b} - \lambda_2 (\|\mathbf{C} \mathbf{x}_2\|^2 - \mathbf{x}_2^T \mathbf{C}^T \mathbf{d}) . \quad (7.9)$$

Now $\mathbf{x}_2^T$ (7.6) - $\mathbf{x}_1^T$ (7.7) gives

$$\begin{aligned} (\lambda_1 - \lambda_2) \mathbf{x}_2^T \mathbf{C}^T \mathbf{C} \mathbf{x}_1 &= -\mathbf{x}_2^T \mathbf{b} + \mathbf{x}_1^T \mathbf{b} + \lambda_1 \mathbf{x}_2^T \mathbf{C}^T \mathbf{d} - \lambda_2 \mathbf{x}_1^T \mathbf{C}^T \mathbf{d} \\ \Rightarrow (\mathbf{x}_1 - \mathbf{x}_2)^T \mathbf{b} &= (\lambda_1 - \lambda_2) \mathbf{x}_2^T \mathbf{C}^T \mathbf{C} \mathbf{x}_1 - \mathbf{d}^T (\lambda_1 \mathbf{C} \mathbf{x}_2 - \lambda_2 \mathbf{C} \mathbf{x}_1) . \end{aligned} \quad (7.10)$$

Subtracting (7.8) - (7.9) and replacing $(\mathbf{x}_1 - \mathbf{x}_2)^T \mathbf{b}$ by (7.10) we get

$$\begin{aligned} 2(\mathbf{F}(\mathbf{x}_1) - \mathbf{F}(\mathbf{x}_2)) &= \\ &- \lambda_1 \{\|\mathbf{C} \mathbf{x}_1\|^2 - \mathbf{x}_1^T \mathbf{C}^T \mathbf{d} - \mathbf{x}_2^T \mathbf{C}^T \mathbf{C} \mathbf{x}_1 + \mathbf{x}_2^T \mathbf{C}^T \mathbf{d}\} \\ &+ \lambda_2 \{\|\mathbf{C} \mathbf{x}_2\|^2 - \mathbf{x}_2^T \mathbf{C}^T \mathbf{d} - \mathbf{x}_2^T \mathbf{C}^T \mathbf{C} \mathbf{x}_1 + \mathbf{x}_1^T \mathbf{C}^T \mathbf{d}\} . \end{aligned}$$

Now like in the proof of Theorem 2.1 we know that both $\{ \}$ are equal and therefore we replace them by their arithmetic mean which gives the desired result. C1

Corollary 7.1. The solution of (7.1), (7.2) is the solution $(\mathbf{x}, \lambda)$ of (7.4) with largest λ .

Proof. From (7.5) we have

$$\lambda_2 < \lambda_1 \Rightarrow \mathbf{F}(\mathbf{x}_2) > \mathbf{F}(\mathbf{x}_1) . \quad \square$$

Theorem 7.2. Let $(\mathbf{x}_i, \lambda_i)$, $i = 1, 2$, be two solutions of (7.4), then

$$(\lambda_1 + \lambda_2) (\mathbf{F}(\mathbf{x}_1) - \mathbf{F}(\mathbf{x}_2)) = (\lambda_2 - \lambda_1) (\mathbf{x}_1 - \mathbf{x}_2)^T \mathbf{A} (\mathbf{x}_1 - \mathbf{x}_2) . \quad (7.11)$$

Proof.

$$\lambda_1 \mathbf{C}^T \mathbf{C} \mathbf{x}_1 - \lambda_1 \mathbf{C}^T \mathbf{d} = -(\mathbf{A} + \mathbf{A}^T) \mathbf{x}_1 - \mathbf{b} \quad (7.12)$$

$$\lambda_2 \mathbf{C}^T \mathbf{C} \mathbf{x}_2 - \lambda_2 \mathbf{C}^T \mathbf{d} = -(\mathbf{A} + \mathbf{A}^T) \mathbf{x}_2 - \mathbf{b} . \quad (7.13)$$

$\lambda_1 \mathbf{x}_1^T$ (7.13) - $\lambda_2 \mathbf{x}_2^T$ (7.12) gives

$$\lambda_1 \lambda_2 \{(\mathbf{x}_2 - \mathbf{x}_1)^T \mathbf{C}^T \mathbf{d}\} = (\lambda_2 - \lambda_1) \mathbf{x}_1^T (\mathbf{A} + \mathbf{A}^T) \mathbf{x}_2 - (\lambda_1 \mathbf{x}_1 - \lambda_2 \mathbf{x}_2)^T \mathbf{b} . \quad (7.14)$$

$\lambda_1 \mathbf{x}_2^T$ (7.13) - $\lambda_2 \mathbf{x}_1^T$ (7.12) gives

$$\begin{aligned} \lambda_1 \lambda_2 \{\|\mathbf{C} \mathbf{x}_2\|^2 - \|\mathbf{C} \mathbf{x}_1\|^2 + (\mathbf{x}_1 - \mathbf{x}_2)^T \mathbf{C}^T \mathbf{d}\} \\ = -2\lambda_1 \mathbf{x}_2^T \mathbf{A} \mathbf{x}_2 + 2\lambda_2 \mathbf{x}_1^T \mathbf{A} \mathbf{x}_1 - (\lambda_1 \mathbf{x}_2 - \lambda_2 \mathbf{x}_1)^T \mathbf{b} . \end{aligned} \quad (7.15)$$

Observe that

$$0 = \|\mathbf{C} \mathbf{x}_2 - \mathbf{d}\|^2 - \|\mathbf{C} \mathbf{x}_1 - \mathbf{d}\|^2 = \|\mathbf{C} \mathbf{x}_2\|^2 - \|\mathbf{C} \mathbf{x}_1\|^2 + 2(\mathbf{x}_1 - \mathbf{x}_2)^T \mathbf{C}^T \mathbf{d} .$$

So that if we subtract (7.15) - (7.14) we get

$$\begin{aligned} 0 &= 2\lambda_2 \mathbf{x}_1^T \mathbf{A} \mathbf{x}_1 - 2\lambda_1 \mathbf{x}_2^T \mathbf{A} \mathbf{x}_2 - (\lambda_1 \mathbf{x}_2 - \lambda_2 \mathbf{x}_1 - \lambda_1 \mathbf{x}_1 + \lambda_2 \mathbf{x}_2)^T \mathbf{b} \\ &\quad + (\lambda_1 - \lambda_2) \mathbf{x}_1^T (\mathbf{A} + \mathbf{A}^T) \mathbf{x}_2 . \end{aligned}$$

Rearranged we have

$$\begin{aligned} \lambda_1 \{ 2 \bar{x}_2^T A \bar{x}_2 + \bar{x}_2^T \bar{b} - \bar{x}_1^T \bar{b} - \bar{x}_1^T (A + A^T) \bar{x}_2 \} \\ = \lambda_2 \{ 2 \bar{x}_1^T A \bar{x}_1 + \bar{x}_1^T \bar{b} - \bar{x}_2^T \bar{b} - \bar{x}_1^T (A + A^T) \bar{x}_2 \} \end{aligned} \quad (7.16)$$

Now observe that the left hand side of (7.16) is

$$\{ \} = F(\bar{x}_2) - F(\bar{x}_1) + (\bar{x}_1 - \bar{x}_2)^T A (\bar{x}_1 - \bar{x}_2) \cdot$$

Similarly the right hand side simplifies. Therefore

$$\begin{aligned} \lambda_1 \{ F(\bar{x}_2) - F(\bar{x}_1) + (\bar{x}_1 - \bar{x}_2)^T A (\bar{x}_1 - \bar{x}_2) \} \\ = \lambda_2 \{ F(\bar{x}_1) - F(\bar{x}_2) + (\bar{x}_1 - \bar{x}_2)^T A (\bar{x}_1 - \bar{x}_2) \} \end{aligned}$$

or rearranged

$$(\lambda_1 + \lambda_2) (F(\bar{x}_1) - F(\bar{x}_2)) = (\lambda_1 - \lambda_2) (\bar{x}_1 - \bar{x}_2)^T A (\bar{x}_1 - \bar{x}_2) \cdot \square$$

Corollary 7.2. Let $(\bar{x}_i, \lambda_i)$, $i = 1, 2$, be two solutions of (7.4) with $\lambda_1 \neq \lambda_2$. Then

$$(\bar{x}_1 - \bar{x}_2)^T A (\bar{x}_1 - \bar{x}_2) = -\frac{\lambda_1 + \lambda_2}{4} \|C(\bar{x}_1 - \bar{x}_2)\|^2 \quad (7.17)$$

Proof. We combine the results of Theorems 7.1 and 7.2. $\square$ $\quad ;$

8. Smoothing of Datas.

Rutishauser proposed in [31] a method to smooth datas. We present here a modified version. Let

$$d_i, \quad i = 1, \dots, n$$

be given datas of a smooth function. However let's assume the d_i are perturbed by measurement errors. We look for a new set of datas

$$x_i, \quad i = 1, \dots, n$$

that does not deviate too much from d_i and that is smoother. If we

assume that the d_i are equidistant then we might want to solve

$$\sum_{i=2}^{n-1} (x_{i+1} - 2x_i + x_{i-1})^2 = \min \quad (8.1)$$

subject to

$$\sum_{i=1}^n (x_i - d_i)^2 \leq n \varepsilon^2 \quad (8.2)$$

Here ε^2 is a measure for the variance, the mean deviation we want to allow the new data x_i to differ from d_i :

$$\varepsilon \approx \sqrt{\frac{\sum_{i=1}^n (x_i - d_i)^2}{n}}$$

The larger ε the more the x_i will be smoothed.

Introducing the $(n-2) \times n$ tridiagonal matrix

$$\bar{A} = \begin{pmatrix} 1 & -2 & 1 & & & \\ & 1 & -2 & 1 & & \\ & & \ddots & \ddots & \ddots & \\ & & & 1 & -2 & 1 \end{pmatrix}$$

and the vectors d and $\bar{x}$, equations (8.1) and (8.2) become

$$\left. \begin{aligned} \|\underline{Ax}\| &= \min \\ \text{subject to} \\ \|\underline{x} - \underline{d}\| &\leq \alpha := \sqrt{n} \cdot \delta \end{aligned} \right\} \quad (8.3)$$

and we have problem (P1) for $\underline{b} = \underline{0}$ and $\underline{C} = \underline{I}_n$. The normal equations are

$$(\underline{A}^T \underline{A} + \lambda \underline{I}) \underline{x} = \lambda \underline{d} \quad (8.4)$$

$$\|\underline{x} - \underline{d}\|^2 = \alpha^2 \quad (8.5)$$

Instead of solving (8.4), Rutishauser proposed to choose some "smoothing" parameter γ and to solve

$$(\underline{I} + \gamma \underline{A}^T \underline{A}) \underline{x} = \underline{d} \quad (8.6)$$

which is (8.4) for $\lambda = 1/\gamma$. The condition number of the matrix in

(8.6) is 16γ . For large n one may have to choose γ large which

leads to numerical problems [36]. However large γ or small λ may

be meaningful: for $\gamma \rightarrow \infty$ ($\lambda \rightarrow 0$) the new values $\underline{x}$ lie on a straight

line, i.e., are values of a linear regression. Of course one would like

to have as limit the linear regression to the datas $\underline{d}$.

If we make the change of variables in (8.3)

$$\begin{aligned} \underline{w} &:= \underline{x} - \underline{d} \\ \underline{b} &:= -\underline{Ad} \end{aligned} \quad (8.7)$$

then our problem becomes

$$\left. \begin{aligned} \|\underline{Aw} - \underline{b}\| &= \min \\ \text{subject to} \\ \|\underline{w}\| &\leq \alpha \end{aligned} \right\} \quad (8.8)$$

Problem (8.8) leads to a relaxed least squares problem

$$\left. \begin{aligned} (\underline{A}^T \underline{A} + \lambda \underline{I}) \underline{w} &= \underline{A}^T \underline{b} \\ \|\underline{w}\| &= \alpha \end{aligned} \right\} \quad (8.9)$$

We can use the dual equations for (8.9) (see Section 3):

$$(\underline{AA}^T + \underline{A} \underline{I}) \underline{z} = -\underline{b} = -\underline{Ad} \quad (8.10)$$

$$\|\underline{A}^T \underline{z}\| = \alpha \quad (8.11)$$

with

$$\underline{w} = -\underline{A}^T \underline{z} \quad (8.12)$$

Observe that now $\lambda \rightarrow 0$ causes no problems since $\underline{AA}^T$ is not singular.

Furthermore since $\underline{b} = -\underline{Ad}$ we can write (8.10) as

$$\begin{pmatrix} \underline{A}^T \\ \sqrt{\lambda} \underline{I} \end{pmatrix} \underline{z} \approx \begin{pmatrix} \underline{d} \\ \underline{0} \end{pmatrix} \quad (8.13)$$

$\underline{A}$ is tridiagonal, therefore we shall use the algorithm described in

Section 5.1 to solve for a given $\lambda \geq 0$ the least squares problem (8.13).

For $n = 8$, e.g. the remaining matrix before and after the third step is:

We cannot apply all the results of Section 5.4 since f is not of the same form. We have

$$(\bar{A}\bar{A}^T + \bar{A} \bar{I})\bar{z}^{(k)} = -\bar{z}^{(k-1)}, \quad k \geq 1.$$

We shall use Reinsch's proposal and therefore we need only f' :

$$\begin{aligned} f'(A) &= 2 \bar{z}'^T \bar{A} \bar{A}^T \bar{z} \\ &= -2 \bar{z}'^T (\bar{z} + \lambda \bar{z}') . \end{aligned}$$

An ALGOL W procedure **smooth**(n, δ, d, x) is given at the end of this section. The following example has been **computed** with this procedure using single precision (ca. 6 decimal digits) on the SLAC computer (IBM 360).

Example.

we choose $n = 30$ and

$$d_i = \sqrt{i} + 0.2 \sin(i), \quad i = 1, \dots, 30.$$

For this **example** we have $f(0) = 1.825814$, therefore if

$$\delta > \sqrt{\frac{f(0)}{n}} = 0.246699$$

the smoothed value x_i are the same as obtained by linear regression.

On the other hand if $\delta = 10^{-4}$ we have $\|\bar{x} - \bar{d}\| = 5.4 \times 10^{-4}$ and the x_i differ only in the fourth decimal from the d_i . If we interpret the d_i as perturbed value of $\sqrt{i}$ then because the mean of $|0.2 \sin(i)|$ is approximately

...

comment column $n-3$;

$$\begin{pmatrix} 0 & c_{n-3} \\ d_{n-3} & y_{n-3} \end{pmatrix} := G_1 \begin{pmatrix} r & c_{n-3} \\ \sqrt{\lambda} & 0 \end{pmatrix};$$

$$\begin{pmatrix} 0 & r & c_{n-2} \\ d_{n-3} & d_{n-3} & y_{n-3} \end{pmatrix} := G_2 \begin{pmatrix} s & t & \\ d_{n-3} & 0 & y_{n-3} \end{pmatrix};$$

$$\begin{pmatrix} 0 & s & c_{n-1} \\ d_{n-3} & d_{n-3} & y_{n-3} \end{pmatrix} := G_3 \begin{pmatrix} 1 & -2 & d_{n-1} \\ d_{n-3} & d_{n-3} & y_{n-3} \end{pmatrix};$$

comment column $n-2$;

$$\begin{pmatrix} 0 & c_{n-2} \\ d_{n-2} & y_{n-2} \end{pmatrix} := G_1 \begin{pmatrix} r & c_{n-2} \\ \sqrt{\lambda} & 0 \end{pmatrix};$$

$$\begin{pmatrix} 0 & c_{n-1} \\ d_{n-2} & y_{n-2} \end{pmatrix} := G_2 \begin{pmatrix} s & c_{n-1} \\ d_{n-2} & y_{n-2} \end{pmatrix};$$

$$\begin{pmatrix} 0 & c_n \\ d_{n-2} & y_{n-2} \end{pmatrix} := G_3 \begin{pmatrix} 1 & d_n \\ d_{n-2} & y_{n-2} \end{pmatrix};$$

end.

To solve the problem we need the derivatives of $f(\lambda)$ defined by

$$(\bar{A}\bar{A}^T + \lambda \bar{I})\bar{z}(\lambda) = \bar{A} \bar{d}$$

$$f(\lambda) = \|\bar{A}^T \bar{z}(\lambda)\|^2.$$

$$0.2 \frac{1}{\pi} \int_0^{\pi} \sin(x) dx \quad \frac{0.4}{\pi} \approx 1.3$$

we expect the best smoothing for $\delta \approx 1.3$ which indeed can be seen clearly in Table 2.

Table 2 : Smoothing of $d_i = \sqrt{i} + 0.2 \sin(i)$, $i = 1, \dots, 30$.

$\delta =$	≥ 0.2467	0.2466	0.2	0.17	0.15	0.13	0.12
	1.722253	1.722027	1.603430	1.507798	1.414468	1.261987	1.243047
	1.861272	1.861084	1.763119	1.684705	1.609509	1.496164	1.483359
	2.000287	2.000137	1.922686	1.861354	1.804054	1.727389	1.717025
	2.139301	2.139186	2.081965	2.031419	1.997584	1.955655	1.947423
	2.270316	2.278234	2.240736	2.212497	2.189486	2.182792	2.181764
	2.417336	2.417287	2.398729	2.386040	2.378851	2.406053	2.418491
	2.556347	2.556335	2.555611	2.557375	2.564482	2.619121	2.643821
	2.695338	2.695371	2.711038	2.725835	2.745208	2.815103	2.841743
	2.834349	2.834425	2.864746	2.890921	2.920292	2.992062	3.003115
	2.973354	2.973438	3.016541	3.052362	3.089548	3.154735	3.155210
	3.112376	3.112476	3.166302	3.210031	3.253120	3.310709	3.301294
	3.251370	3.251490	3.313931	3.363794	3.411095	3.464372	3.457992
	3.390381	3.390501	3.459267	3.513468	3.563265	3.613997	3.616866
	3.529382	3.529527	3.602212	3.658839	3.703311	3.754458	3.762860
	3.668403	3.668556	3.742742	3.799843	3.849200	3.883019	3.866796
	3.807439	3.807585	3.880861	3.936624	3.983325	4.002785	3.995223
	3.946456	3.946573	4.016701	4.069475	4.112409	4.120646	4.105030
	4.085469	4.085618	4.150337	4.196711	4.237081	4.241594	4.228400
	4.224488	4.224601	4.282032	4.324522	4.357614	4.364625	4.362045
	4.363507	4.363605	4.411738	4.447031	4.473988	4.404037	4.490513
	4.502520	4.502613	4.539614	4.566405	4.586239	4.594892	4.600982
	4.641560	4.641610	4.665841	4.682961	4.694780	4.697631	4.694865
	4.780583	4.730605	4.790689	4.797174	4.800344	4.797594	4.787060
	4.919610	4.919595	4.914412	4.400512	4.903647	4.899828	4.891946
	5.058637	5.058589	5.037205	5.020307	5.005067	5.003985	5.001230
	5.137662	5.197586	5.159234	5.129760	5.104605	5.103989	5.123779
	5.336688	5.336576	5.280628	5.238020	5.202195	5.192799	5.216838
	5.475712	5.475564	5.401550	5.345339	5.298068	5.268023	5.281142
	5.614742	5.614557	5.522194	5.452085	5.392826	5.333412	5.324587
	5.753768	5.753551	5.642737	5.558625	5.487171	5.395162	5.3601306

$\ x - d\ $	1.351226	1.350682	1.095428	0.9311236	0.8215933	0.7120381	0.6572663
$\left[\sum (\Delta^2 x_i)^2 \right]^{1/2}$	$6.67 \cdot 10^{-5}$	$1.11 \cdot 10^{-4}$	0.008449	0.0150982	0.021434	0.0445737	0.0796041
$\left[\sum (x_i - \sqrt{i})^2 \right]^{1/2}$	1.205420	1.204810	0.9014646	0.6828826	0.5084146	0.306330	0.3031789
# of iteration	0	1	5	5	7		
$\lambda = f^{-1}(nd^2)$	0	$3.83 \cdot 10^{-7}$	$2.79 \cdot 10^{-4}$	$7.60 \cdot 10^{-4}$	$2.01 \cdot 10^{-3}$	$3.15 \cdot 10^{-2}$	$8.89 \cdot 10^{-2}$

References

- [1] Bjork, A., "Iterative Refinement of Linear Least Squares Solutions I," BIT 7 (1967), 257-278.
- [2] Bjork, A., "Solving Linear Least Squares Problems by Gram-Schmidt Orthogonalization," BIT 7 (1967), 1-21.
- [3] Brown, G. H. Jr., "On Halley's Variation of Newton's Method," American Mathematical Monthly 84, 9 (1977).
- [4] Bunb, J. R. and Rose, D. J., Sparse Matrix Computations, Academic Press, (1976).
- [5] Chan, T. F. C., "On Computing the Singular Value Decomposition," Stanford Computer Science Department Report STAN-CS-77-588, (1977).
- [6] Davies, M. and Dawson, B., "On Global Convergence of Halley's Iteration Formula," Numer. Math. 24 (1975).
- [7] Ehrmann, H., "Konstruktion und Durchfuehrung von Iterationsverfahren hoeherer Ordnung," Arch. Rational Mech. Anal. 4 (1959).
- [8] Elden, L., "Algorithms for the Regularization of Ill-Conditioned Least Squares Problems," BIT 17 (1977), 134-145.
- [9] Elden, L., "Numerical Analysis of Regularization and Constrained Least Squares Methods," Thesis No. 21, Linköping University (1977).
- [10] Forsythe, G. E. and Golub, G. H., "On the Stat. Values of a Second Degree Polynomial on the Unit Sphere," J. Soc. Indust. Appl. Math. 13 (1965).
- [11] Forsythe, G. E., Malcolm, M. A. and Moler, C. B., Computer Methods for Mathematical Computations, Prentice Hall, 147.
- [12] Gander, W., "How to Apply the Dual Problem to Solve a Certain Constrained Minimum Norm Problem," SIAM Spring Meeting Madison 1978.
- [13] Gander, W., Molinari, L. and Svecowa, H., Numerische Prozeduren ISBN 33, Birkhaeuser Verlag, 1977.
- [14] Golub, G. H. and Kahan, W., "Calculating the Singular Values and Pseudo-Inverse of a Matrix," SIAM J. Numer. Anal. 2, 2 (1965).
- [15] Golub, G. H., Kliema, V. and Stewart, G. W., "Rank Degeneracy and Least Squares Problems," Stanford Computer Science Department Report STAN-CS-76-559 (1976).
- [16] Golub, G. H., "Some Modified Matrix Eigenvalue Problems," SIAM Review 15, 2 (1973).
- [17] Golub, G. H., Heath, M. and Wahba, G., "Generalized Cross-Validation," Stanford Computer Science Department Report STAN-CS-77-622 (1977).
- [18] Golub, G. H. and Luk, F. T., "Singular Value Decomposition: Applications and Computations," Internal Report Serra House Stanford University Computer Science Department (1978).
- [19] Hansen, E. and Patrick, M., "A Family of Root Finding Methods," Numer. Math. 27 (1977), 257-269.
- [20] Hanson, R. J. and Phillips, J. L., "An Adaptive Numerical Method for Fredholm Integral Equations," Numer. Math. 24 (1975), 291-307.
- [21] Henrici, P., Applied and Computational Complex Analysis, Wiley, 1974.
- [22] Henrici, P., Elements of Numerical Analysis, Wiley, 1964.
- [23] Kahan, W., Manuscript in Box 5 G of the G. Forsythe archive, Stanford Main Library.
- [24] Kahan, W., "Why Use Tangents When Secants Will Do," Gatlinburg Conference 1977, Asilomar.
- [25] Kalman, R. E., "Algebraic Aspects of the Generalized Inverse . . ." Generalized Inverses and Applications, Academic Press, (1976).
- [26] Lawson, Ch. L. and Hanson, R. J., Solving Least Squares Problems, Prentice Hall, 1974.
- [27] Marquart, D. W., "An Algorithm for Least-Squares Estimation of Nonlinear Parameters," SIAM 11, 2 (1963).
- [28] Marti, J. T., "Minimum Norm Solutions of Fredholm Integral Equations of the First Kind," Report 78-01 Sem. f. angew. Math. ETHZ, (1978).
- [29] Moré, J. J., "The Levenberg-Marquart Algorithm: Implementation and Theory," Dundee Conference on Numerical Analysis (1977).
- [30] Ostrowski, A. M., Solution of Equations and Systems of Equations, Academic Press, (1973).
- [31] Paige, C. C., "Bidiagonalization of Matrices and Solution of Linear Equations," SIAM J. Numer. Anal. 11, 1 (1974).
- [32] Peters, G. and Wilkinson, J. H., "Ax = λBx and the Generalized Eigenproblem," SIAM J. Numer. Anal. 7, 4 (1970).
- [33] Reinsch, Chr. H., "Smoothing by Spline Functions," Numer. Math. 10 (1967), 177-183.
- [34] Reinsch, Chr. H., "Smoothing by Spline Functions. II," Numer. Math. 16 (1971), 451-454.

- [35] Rutishauser, H., "Once Again: The Least Square Problem," Linear Algebra and Its Applications 1 (1968), 479-488.
- [36] Rutishauser, H., "Vorlesungen ueber numerische Mathematik, hrsg. von M. Gutknecht," Band I, II, Birkhaeuser Verlag, 1976.
- [37] Schek, H.-J. and Eggenberger, R., "Least Squares - Loesungen und Daempfung bei unterbestimmten Gleichungssystemen," Computing 19 (1977).
- [38] Schwarz, H. R., Rutishauser, H. and Stiefel, E., Matrizen-Numerik, Teubner Verlag, 1968.
- [39] Spjøtvoll, E., "A Note on a Theorem of Forsythe and Golub," SIAM J. Appl. Math. 23, 3 (142).
- [40] Stewart, G. W., "On the Continuity of the Generalized Inverse," SIAM J. Appl. Math. 17, 1 (1969).
- [41] Stewart, G. W., "On the Numerical Properties of an Iteration for Computing the Generalized Inverse," CNA 12, Austin, Texas (1971).
- [42] Tikhonov, A. N., Solution of Ill Posed Problems, Wiley, 1977.
- [43] Traub, J. F., Iterative Methods for Solution of Equations, Prentice Hall, 1964.
- [44] Van Loan, Ch. "Lectures in Least Squares," Technical Report TR 76-279, Cornell University (1976).
- [45] Van Loan, Ch. F., "Generalizing the Singular Value Decomposition," SIAM J. Numer. Anal. 13, 1 (1976).
- [46] Varah, J. M., "On the Numerical Solution of Ill-Conditioned Linear Systems," SIAM J. Num. Anal. 10 (1973).
- [47] Varah, J. M., "A Practical Examination of Num. Methods for Ill Posed Problems," Technical Report 76-08, University of British Columbia, 1976.
- [48] Wilkinson, J. H. and Reinsch, C., Linear Algebra, Springer, 1971.