

Andreas Meister

Numerik linearer Gleichungssysteme

Aus dem Programm
Numerische Mathematik

Numerische Mathematik
von M. Bollhöfer und V. Mehrmann

Finanzderivate mit MATLAB®
von M. Günther und A. Jüngel

Numerik linearer Gleichungssysteme
von A. Meister

Numerische Mathematik für Anfänger
von G. Opfer

Numerische Mathematik kompakt
von R. Plato

Übungsbuch zur Numerischen Mathematik
von R. Plato

Keine Probleme mit Inversen Problemen
von A. Rieder

Elementare Numerische Mathematik
von B. Schuppar

Andreas Meister

Numerik linearer Gleichungssysteme

Eine Einführung in moderne Verfahren

**Mit MATLAB-Implementierungen
von C. Vömel**

3., überarbeitete Auflage



Bibliografische Information Der Deutschen Nationalbibliothek
Die Deutsche Nationalbibliothek verzeichnet diese Publikation in der
Deutschen Nationalbibliografie; detaillierte bibliografische Daten sind im Internet über
<<http://dnb.d-nb.de>> abrufbar.

Prof. Dr. Andreas Meister

Universität Kassel

Fachbereich Mathematik

Heinrich-Plett-Str. 40

34132 Kassel

meister@mathematik.uni-kassel.de

1. Auflage 1999

2., überarbeitete Auflage 2005

3., überarbeitete Auflage 2008

Alle Rechte vorbehalten

© Friedr. Vieweg & Sohn Verlag | GWV Fachverlage GmbH, Wiesbaden 2008

Lektorat: Ulrike Schmickler-Hirzebruch | Susanne Jahnel

Der Vieweg Verlag ist ein Unternehmen von Springer Science+Business Media.

www.vieweg.de



Das Werk einschließlich aller seiner Teile ist urheberrechtlich geschützt. Jede Verwertung außerhalb der engen Grenzen des Urheberrechtsgesetzes ist ohne Zustimmung des Verlags unzulässig und strafbar. Das gilt insbesondere für Vervielfältigungen, Übersetzungen, Mikroverfilmungen und die Einspeicherung und Verarbeitung in elektronischen Systemen.

Umschlaggestaltung: Ulrike Weigel, www.CorporateDesignGroup.de

Druck und buchbinderische Verarbeitung: MercedesDruck, Berlin

Gedruckt auf säurefreiem und chlorfrei gebleichtem Papier.

Printed in Germany

ISBN 978-3-8348-0431-0

Für Jonas, Anika, Lara und Jannik

Vorwort

Im Rahmen der Numerik linearer Gleichungssysteme befassen wir uns mit der effizienten Lösung großer, linearer Systeme, wodurch ein wichtiges Teilgebiet der Numerischen Linearen Algebra betrachtet wird, das in den letzten Jahren immer größere Bedeutung gewonnen hat.

Der in den vergangenen zwei Dekaden vollzogene drastische Anstieg der Leistungsfähigkeit von Personal Computern, Workstations und Großrechneranlagen hat zu einer weitverbreiteten Entwicklung numerischer Verfahren zur Simulation praxisrelevanter Problemstellungen in der Medizin, der Physik, den Ingenieurwissenschaften und vielen weiteren Bereichen geführt. Neben der Methode der Finiten Elemente, die inhärent auf ein lineares Gleichungssystem führt, benötigen auch die häufig verwendeten Finite-Differenzen- und Finite-Volumen-Verfahren in Kombination mit einem impliziten Zeitschrittverfahren einen Algorithmus zur Lösung linearer Gleichungssysteme. So ist es nicht verwunderlich, daß die Forschungsaktivitäten auf dem Gebiet der Gleichungssystemlöser einen deutlichen Aufschwung erfahren und zur Entwicklung einer Vielzahl effizienter Verfahren geführt haben.

Das vorliegende Buch basiert auf den Inhalten einer vom Autor am Fachbereich Mathematik der Universität Hamburg gehaltenen vierstündigen Vorlesung, die sich im Kontext der erwähnten Entwicklungen mit der Herleitung und Analyse klassischer sowie moderner Methoden zur Lösung linearer Gleichungssysteme befaßte.

Aufgrund der in diesem Gebiet vorliegenden Vielzahl unterschiedlicher direkter und iterativer Verfahren, ist es innerhalb einer vierstündigen Vorlesung sicherlich nicht möglich alle existierenden Algorithmen vorzustellen. Ziel dieses Manuskriptes ist es daher, dem interessierten Leser einen Überblick über weite Bereiche dieses Gebietes zu vermitteln, die für praktische Anwendungen wichtigen Methoden zu diskutieren, verwandte Algorithmen durch Bemerkungen und Literaturhinweise zu integrieren und die Erarbeitung weiterer Verfahren zu erleichtern.

Die vorausgesetzten Grundkenntnisse beschränken sich hierbei gezielt auf übliche Inhalte der Analysis und linearen Algebra mathematischer Vorlesungen in den ersten zwei Semestern eines Hochschulfaches, da die beschriebenen Methoden weit über die Grenzen eines Mathematikstudiums von großem Interesse sind. Zur Unterstützung eines Selbststudiums werden zudem alle benötigten Grundlagen in einem eigenständigen Kapitel bereitgestellt.

Nachdem das Auftreten linearer Gleichungssysteme im ersten Kapitel anhand einiger Modellbeispiele beschrieben wird, stellen wir im zweiten Kapitel die für die folgenden Methoden benötigten Grundlagen der linearen Algebra zur Verfügung. Das dritte Kapitel widmet sich den direkten Verfahren, die häufig in modernen Gleichungssystemlösern involviert sind oder teilweise in unvollständigen Versionen als Vorkonditionierer verwendet werden. Der Schwerpunkt liegt auf der Beschreibung iterativer Verfahren, die im anschließenden vierten Kapitel vorgestellt werden. Hierbei wird stets besonderen Wert auf eine Motivation sowie eine übersichtliche, einheitliche und mathematisch abgesicherte Herleitung der iterativen Gleichungssystemlöser gelegt. Neben den Splitting-Methoden, wie zum Beispiel dem Jacobi- und Gauß-Seidel-Verfahren, der Richardson-Iteration und

den Relaxationsverfahren, beschreiben wir zunächst die Zweigittermethode und anschließend das Mehrgitterverfahren sowie dessen vollständige Variante. Die Herleitung des CG-Verfahrens nehmen wir durch eine Kombination der zuvor beschriebenen Verfahren des steilsten Abstiegs und der konjugierten Richtungen vor. Desweiteren betrachten wir vom GMRES-Verfahren über die BiCG-Methode bis zum QMRGStAB-Verfahren eine große Bandbreite moderner Krylov-Unterraum-Methoden, die zur Lösung von Gleichungssystemen mit einer unsymmetrischen und indefiniten Matrix geeignet sind. Das abschließende fünfte Kapitel ist einer ausführlichen Beschreibung und Untersuchung möglicher Präkonditionierungstechniken gewidmet, da die dargestellten Verfahren bei praxisrelevanten Problemstellungen in der Regel erst in Kombination mit einem geeigneten Vorkonditionierer eine stabile und effiziente Gesamtmethode liefern.

Die präsentierten Methoden finden ihre Anwendung in unterschiedlichsten numerischen Verfahren, die in verschiedensten Programmiersprachen entwickelt wurden. Daher wurde gezielt eine Darstellung der Algorithmen gewählt, die eine Umsetzung in ein beliebiges Computer Programm ermöglichen, weshalb die Verwendung einer expliziten Programmiersprache vermieden wurde. Viele der präsentierten Verfahren sind bereits in Softwarepaketen wie zum Beispiel LAPACK [3], LINSOL [75], MATLAB [1] und Templates [8] verfügbar.

Bedanken möchte ich mich an dieser Stelle bei Frau Monika Jampert, die weite Teile meines handschriftlichen Manuskriptes in ein $\text{\LaTeX}$ -Skript verwandelt hat und bei Dipl. Math. Dirk Nitschke für seine Unterstützung bei allen $\text{\LaTeX}$ -Fragen. Zudem gilt mein Dank Dr. Christoph Bäsler, Dipl. Math. Michael Breuss, Martin Ludwig, Christian Nagel, Dipl. Math. Stefanie Schmidt und Dipl. Math. Christof Vömel für das intensive Korrekturlesen und die vielen konstruktiven Hinweise und Bemerkungen, die sich in vielen Bereichen positiv ausgewirkt haben. Besonders bedanken möchte ich mich an dieser Stelle bei Prof. Dr. Thomas Sonar, der durch seine ermutigende Unterstützung und jahrelange fachliche Begleitung wesentlich zur Entstehung dieses Buches beigetragen hat.

Hamburg, im September 1999

ANDREAS MEISTER

Vorwort zur zweiten Auflage

Innerhalb der zweiten Auflage wurden neben der Korrektur entdeckter Fehler auch textliche Erweiterungen und ergänzende Bemerkungen eingefügt, die hoffentlich eine bessere Lesbarkeit zur Folge haben. Mein Dank gilt hierbei vielen Lesern, die durch zahlreiche Hinweise und konstruktive Kritik wesentlich zur Verbesserung beigetragen haben.

Eine zentrale Erweiterung stellen die im ergänzten Anhang aufgeführten MATLAB-Implementierung zu gängigen Krylov-Unterraum-Methoden dar, die von Dr. C. Vömel entwickelt wurden. Für diese hilfreiche Unterstützung möchte ich mich an dieser Stelle recht herzlich bedanken. Durch diese Ergänzung ergibt sich für den interessierten Anwender eine Möglichkeit zur unmittelbaren Nutzung der diskutierten Verfahren, die einen tieferen Einblick in die jeweiligen Eigenschaften der betrachteten Methode im Kontext individueller Problemstellungen eröffnet.

Über Kommentare, Verbesserungsvorschläge und Korrekturhinweise, die mir beispielsweise über meine Email-Adresse `meister@mathematik.uni-kassel.de` zugesendet werden können, würde ich mich sehr freuen. Desweiteren besteht unter

`http://www.mathematik.uni-kassel.de/~meister/buch\_online`

ein Online-Service zum Buch. Neben aktuellen Informationen und Korrekturhinweisen können dieser Seite auch die angesprochenen MATLAB-Implementierungen entnommen werden.

Kassel, im November 2004

ANDREAS MEISTER

Vorwort zur dritten Auflage

Neben geringfügigen Korrekturen von Tippfehlern wurden im Rahmen dieser Auflage zahlreiche Übungsaufgaben ergänzt, die sich jeweils am Ende der Kapitel 2 bis 5 befinden. Bei einigen Problemstellungen konnte ich dabei auf Übungen aus meinem eigenen Studium zurückgreifen. Den damaligen Aufgabenstellern gilt daher rückwirkend mein herzlicher Dank. Die zugehörigen Lösungen können dem unter

`http://www.vieweg.de/tu/8y`

bereitgestellten Online-Service entnommen werden. Hier findet der interessierte Leser zudem Hinweise zu themenbegleitenden Kompaktkursen sowie die im Buch aufgeführten MATLAB-Implementierungen.

Kassel, im Oktober 2007

ANDREAS MEISTER

Inhaltsverzeichnis

1	Beispiele für das Auftreten linearer Gleichungssysteme	1
2	Grundlagen der linearen Algebra	7
2.1	Vektornormen und Skalarprodukt	7
2.2	Lineare Operatoren, Matrizen und Matrixnormen	13
2.3	Konditionszahl und singuläre Werte	25
2.4	Der Banachsche Fixpunktsatz	29
2.5	Übungsaufgaben	32
3	Direkte Verfahren	36
3.1	Gauß-Elimination	36
3.2	Cholesky-Zerlegung	46
3.3	QR-Zerlegung	49
3.3.1	Das Gram-Schmidt-Verfahren	49
3.3.2	Die QR-Zerlegung nach Givens	55
3.4	Übungsaufgaben	59
4	Iterative Verfahren	62
4.1	Splitting-Methoden	65
4.1.1	Jacobi-Verfahren	69
4.1.2	Gauß-Seidel-Verfahren	74
4.1.3	Relaxationsverfahren	77
4.1.3.1	Jacobi-Relaxationsverfahren	78
4.1.3.2	Gauß-Seidel-Relaxationsverfahren	80
4.1.4	Richardson-Verfahren	89
4.1.5	Symmetrische Splitting-Methoden	92
4.2	Mehrgitterverfahren	97
4.2.1	Zweigitterverfahren	105
4.2.2	Der Mehrgitteralgorithmus	109
4.2.3	Das vollständige Mehrgitterverfahren	111

4.3	Projektionsmethoden und Krylov-Unterraum-Verfahren	112
4.3.1	Verfahren für symmetrische, positiv definite Matrizen	117
4.3.1.1	Die Methode des steilsten Abstiegs	118
4.3.1.2	Das Verfahren der konjugierten Richtungen	124
4.3.1.3	Das Verfahren der konjugierten Gradienten	126
4.3.2	Verfahren für reguläre Matrizen	135
4.3.2.1	Der Arnoldi-Algorithmus und die FOM	135
4.3.2.2	Der Lanczos-Algorithmus und die D-Lanczos-Methode	140
4.3.2.3	Der Bi-Lanczos-Algorithmus	144
4.3.2.4	Das GMRES-Verfahren	149
4.3.2.5	Das BiCG-Verfahren	164
4.3.2.6	Das CGS-Verfahren	171
4.3.2.7	Das BiCGSTAB-Verfahren	174
4.3.2.8	Das TFQMR-Verfahren	180
4.3.2.9	Das QMRCGSTAB-Verfahren	189
4.3.2.10	Konvergenzanalysen	192
4.4	Übungsaufgaben	194
5	Präkonditionierer	200
5.1	Skalierungen	201
5.2	Polynomiale Präkonditioner	204
5.3	Splitting-assoziierte Präkonditionierer	207
5.4	Die unvollständige LU-Zerlegung	208
5.5	Die unvollständige Cholesky-Zerlegung	211
5.6	Die unvollständige QR-Zerlegung	212
5.7	Die unvollständige Frobenius-Inverse	214
5.8	Das präkonditionierte CG-Verfahren	216
5.9	Das präkonditionierte BiCGSTAB-Verfahren	219
5.10	Vergleich der Präkonditionierer	221
5.11	Übungsaufgaben	225
A	Implementierungen in MATLAB	226
	Literaturverzeichnis	239
	Index	244

1 Beispiele für das Auftreten linearer Gleichungssysteme

In diesem Kapitel beschäftigen wir uns mit der Entstehung linearer Gleichungssysteme auf der Basis physikalisch relevanter Problemstellungen. Die ausgewählten Beispiele verdeutlichen einerseits das Auftreten schwachbesetzter (Beispiele 1.1 und 1.4) sowie vollbesetzter (Beispiele 1.2 und 1.3) Matrizen und dienen andererseits als Modellprobleme (Beispiele 1.1 und 1.4) für die Analyse der in den folgenden Abschnitten beschriebenen Verfahren.

Beispiel 1.1 Die Poisson-Gleichung

Die Poisson-Gleichung stellt eine elliptische partielle Differentialgleichung zweiter Ordnung dar, die zu einem Standardproblem bei der Untersuchung partieller Differentialgleichungen und deren numerischer Behandlung geworden ist. Sie tritt unter anderem bei der numerischen Simulation inkompressibler reibungsbehafteter Strömungsfelder, das heißt bei der Diskretisierung der inkompressiblen Navier-Stokes-Gleichungen auf [17, 22].

Unter Verwendung des Laplace-Operators

$$\Delta = \frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2}$$

läßt sich die Poisson-Gleichung als Randwertproblem auf dem Gebiet $\Omega \subset \mathbb{R}^2$ in der Form

$$\begin{aligned} -\Delta u(x,y) &= f(x,y) \text{ für } (x,y) \in \Omega \subset \mathbb{R}^2, \\ u(x,y) &= \varphi(x,y) \text{ für } (x,y) \in \partial\Omega \end{aligned} \tag{1.0.1}$$

schreiben, wobei $\partial\Omega$ den Rand von Ω beschreibt.

Bei gegebener Funktion $f \in C(\Omega; \mathbb{R})$ und gegebenen Randwerten $\varphi \in C(\partial\Omega; \mathbb{R})$ wird folglich eine Funktion $u \in C^2(\Omega; \mathbb{R}) \cap C(\overline{\Omega}; \mathbb{R})$ gesucht, die dem Randwertproblem (1.0.1) genügt.

Zur Diskretisierung der Poisson-Gleichung auf dem Einheitsquadrat $\Omega = (0,1) \times (0,1)$ mittels einer zentralen Finite-Differenzen-Methode wird $\overline{\Omega} = \Omega \cup \partial\Omega$ mit einem Gitter Ω_h der Schrittweite $h = 1/(N+1)$ mit $N \in \mathbb{N}$ versehen.

Wir schreiben

$$(x_i, y_j) = (ih, jh) \text{ für } i, j = 0, \dots, N+1$$

sowie

$$u_{ij} = u(x_i, y_j) \text{ und } f_{ij} = f(x_i, y_j) \text{ für } i, j = 0, \dots, N+1.$$

Desweiteren approximieren wir

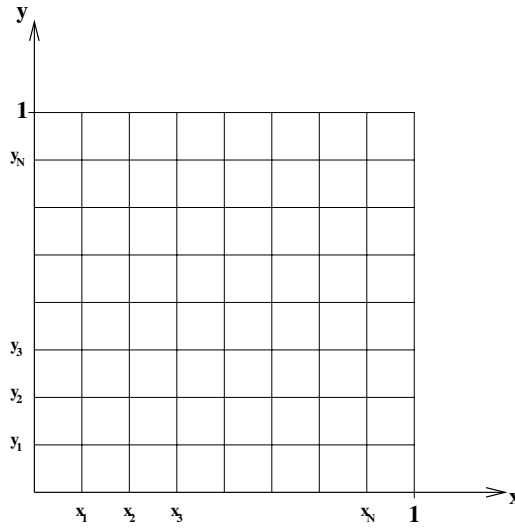


Bild 1.1 Diskretisierung des Einheitsquadrates

$$\begin{aligned}
 \frac{\partial^2 u}{\partial x^2}(x_i, y_j) &\approx \frac{1}{h} \left(\underbrace{\frac{u_{i+1,j} - u_{i,j}}{h}}_{\approx \frac{\partial u}{\partial x}(x_{i+1/2}, y_j)} - \frac{u_{i,j} - u_{i-1,j}}{h} \right) \\
 &= \frac{1}{h^2} (u_{i+1,j} - 2u_{i,j} + u_{i-1,j}).
 \end{aligned} \tag{1.0.2}$$

Mit einer analogen Vorgehensweise für $\frac{\partial^2 u}{\partial y^2}$ erhalten wir für den Laplace-Operator

$$-\Delta u(x_i, y_j) \approx \frac{1}{h^2} (4u_{i,j} - u_{i-1,j} - u_{i+1,j} - u_{i,j-1} - u_{i,j+1})$$

und somit die diskrete Form der Gleichung (1.0.1)

$$\begin{aligned}
 4u_{i,j} - u_{i-1,j} - u_{i+1,j} - u_{i,j-1} - u_{i,j+1} &= h^2 f_{ij} && \text{für } 1 \leq i, j \leq N, \\
 u_{0,j} = \varphi_{0,j}, \quad u_{N+1,j} = \varphi_{N+1,j} &&& \text{für } j = 0, \dots, N+1, \\
 u_{i,0} = \varphi_{i,0}, \quad u_{i,N+1} = \varphi_{i,N+1} &&& \text{für } i = 0, \dots, N+1.
 \end{aligned}$$

Eine zeilenweise Neunummerierung, die einer lexikographischen Anordnung der inneren Gitterpunkte entspricht, ergibt

$$u_1 = u_{1,1}, \quad u_2 = u_{2,1}, \quad u_3 = u_{3,1}, \dots, \quad u_N = u_{N,1}, \quad u_{N+1} = u_{1,2}, \dots, u_{N^2} = u_{N,N}.$$

Die Neunummerierung wird ebenfalls für f durchgeführt, und wir erhalten somit ein Gleichungssystem $\mathbf{A}\mathbf{u} = \mathbf{g}$ für die Bestimmung des Lösungsvektors $\mathbf{u} := (u_1, \dots, u_{N^2})^T$. Hierbei gilt

$$\mathbf{A} = \begin{pmatrix} \mathbf{B} & -\mathbf{I} & & \\ -\mathbf{I} & \ddots & \ddots & \\ & \ddots & \ddots & -\mathbf{I} \\ & & -\mathbf{I} & \mathbf{B} \end{pmatrix} \in \mathbb{R}^{N^2 \times N^2} \quad (1.0.3)$$

mit

$$\mathbf{B} = \begin{pmatrix} 4 & -1 & & \\ -1 & \ddots & \ddots & \\ & \ddots & \ddots & -1 \\ & & -1 & 4 \end{pmatrix} \in \mathbb{R}^{N \times N} \quad \text{und} \quad \mathbf{I} = \begin{pmatrix} 1 & & \\ & \ddots & \\ & & 1 \end{pmatrix} \in \mathbb{R}^{N \times N}.$$

Die rechte Seite weist für den Spezialfall $\varphi \equiv 0$ die Gestalt

$$\mathbf{g} = h^2 \begin{pmatrix} f_1 \\ \vdots \\ f_{N^2} \end{pmatrix} \in \mathbb{R}^{N^2} \quad (1.0.4)$$

auf. Im Fall $\varphi \not\equiv 0$ müssen die entsprechenden Randwerte im Vektor der rechten Seite berücksichtigt werden. Für die erste Komponente erhalten wir in diesem Fall

$$g_1 = h^2 f_1 + \varphi_{0,1} + \varphi_{1,0}.$$

Bemerkung:

Die Matrix $\mathbf{A}$ ist symmetrisch, positiv definit und schwachbesetzt. Ihre Struktur ist zudem abhängig von der gewählten Anordnung der inneren Gitterpunkte.

Beispiel 1.2 Lineare Integralgleichung zweiter Art

Lineare Integralgleichungen erster und zweiter Art treten zur Festlegung von Dichtefunktionen auf, mittels derer die Lösung einer partiellen Differentialgleichung (zum Beispiel der Helmholtz-Gleichung) in Form eines Einfach- oder Doppelschichtpotentials dargestellt werden kann [43, 35].

Wir betrachten die Integralgleichung zweiter Art

$$u(x) = v(x) + \int_0^1 k(x,y)u(y)dy \quad \text{für } x \in [0,1]. \quad (1.0.5)$$

Bei gegebenem Integralkern $k : [0,1] \times [0,1] \rightarrow \mathbb{R}$ mit $k(x,0) = 0$ und gegebener Funktion $v : [0,1] \rightarrow \mathbb{R}$ werden wir im folgenden ein lineares Gleichungssystem zur Berechnung einer Näherungslösung der gesuchten Funktion $u : [0,1] \rightarrow \mathbb{R}$ herleiten.

Zur Diskretisierung der Integralgleichung verwenden wir eine Nyström-Methode. Hierzu unterteilen wir das Intervall $[0,1]$ in N Teilintervalle der Länge $h = 1/N$ mit $N \in \mathbb{N}$. Sei

$$x_j = \frac{j}{N} = jh \quad \text{für } j = 0, \dots, N,$$

dann erhalten wir unter Verwendung der numerischen Integration

$$\int_0^1 k(x,y)u(y)dy \approx h \sum_{j=1}^N k(x,x_j)u(x_j) = h \sum_{j=0}^N k(x,x_j)u(x_j)$$

aus (1.0.5) die approximative Darstellung

$$u(x) = v(x) + \frac{1}{N} \sum_{j=0}^N k(x,x_j)u(x_j) \text{ für } x \in [0,1]. \quad (1.0.6)$$

Betrachten wir die Gleichung (1.0.6) an den Stützstellen x_i , $i = 0, \dots, N$ und definieren $u_i = u(x_i)$, $v_i = v(x_i)$, dann erhalten wir mit $\mathbf{v} = (v_0, \dots, v_N)^T$ das lineare Gleichungssystem $\mathbf{A}\mathbf{u} = \mathbf{v}$ zur Bestimmung des Vektors $\mathbf{u} = (u_0, \dots, u_N)^T$. Hierbei gilt

$$\mathbf{A} = (a_{ij})_{i,j=0,\dots,N} \in \mathbb{R}^{(N+1) \times (N+1)}$$

mit

$$a_{ij} = \delta_{ij} - \frac{1}{N} k(x_i, x_j) \quad (1.0.7)$$

unter Verwendung des Kronecker-Symbols

$$\delta_{ij} = \begin{cases} 0, & i \neq j \\ 1, & i = j. \end{cases}$$

Bemerkung:

Aus der Darstellung der Matrixkoeffizienten (1.0.7) ist unmittelbar ersichtlich, daß die Besetzungsstruktur sowie alle weiteren Eigenschaften der Matrix (Symmetrie, Definitheit) von dem zugrundeliegenden Integralkern abhängen. Im Normalfall ist die Matrix $\mathbf{A}$ hierbei vollbesetzt.

Beispiel 1.3 Funktionsdefinition aus Meßwerten

Es sei bekannt, daß sich eine physikalische Größe z als Polynom n -ten Grades der Zeit t schreiben läßt:

$$z(t) = \sum_{i=0}^n \alpha_i t^i \text{ für } t \in \mathbb{R}_0^+.$$

Aus Messungen ist die physikalische Größe zu den Zeitpunkten $0 \leq t_0 < \dots < t_N$, $N \geq n$, bekannt. Die zugehörigen Meßwerte lauten $z_0, \dots, z_N$. Die Aufgabe besteht in der Ermittlung der Koeffizienten $\alpha_0, \dots, \alpha_n$, so daß die Summe der Fehlerquadrate minimal wird (Methode der kleinsten Quadrate). Das bedeutet

$$f(\boldsymbol{\alpha}) = \sum_{j=0}^N \left[\sum_{i=0}^n \alpha_i t_j^i - z_j \right]^2 = \sum_{j=0}^N [z(t_j) - z_j]^2$$

soll über alle $\boldsymbol{\alpha} = (\alpha_0, \dots, \alpha_n)^T \in \mathbb{R}^{n+1}$ minimiert werden. Der Lösungsvektor $\boldsymbol{\alpha}$ erfüllt somit

$$\frac{\partial f(\boldsymbol{\alpha})}{\partial \alpha_k} = 0 \text{ für } k = 0, \dots, n.$$

Aufgrund der Gestalt der Funktion f folgt hieraus

$$\begin{aligned}
 \sum_{j=0}^N 2[z(t_j) - z_j] t_j^k &= 0 && \text{für } k = 0, \dots, n \\
 \Leftrightarrow \sum_{j=0}^N z(t_j) t_j^k &= \sum_{j=0}^N z_j t_j^k && \text{für } k = 0, \dots, n \\
 \Leftrightarrow \sum_{j=0}^N \left(\sum_{i=0}^n \alpha_i t_j^i \right) t_j^k &= \sum_{j=0}^N z_j t_j^k && \text{für } k = 0, \dots, n \\
 \Leftrightarrow \sum_{i=0}^n \left(\sum_{j=0}^N t_j^{i+k} \right) \alpha_i &= \sum_{j=0}^N z_j t_j^k && \text{für } k = 0, \dots, n.
 \end{aligned}$$

Zu lösen bleibt somit das lineare Gleichungssystem $\mathbf{A}\boldsymbol{\alpha} = \mathbf{b}$ mit

$$\mathbf{A} = (a_{ki})_{k,i=0,\dots,n} \in \mathbb{R}^{(n+1) \times (n+1)} \text{ mit } a_{ki} = \sum_{j=0}^N t_j^{k+i} \text{ für } k, i = 0, \dots, n$$

und

$$\mathbf{b} = (b_0, \dots, b_n)^T \in \mathbb{R}^{n+1} \text{ mit } b_k = \sum_{j=0}^N z_j t_j^k \text{ für } k = 0, \dots, n.$$

Bemerkung:

Die resultierende Matrix $\mathbf{A}$ ist stets vollbesetzt und symmetrisch. Das vorliegende Gleichungssystem wird üblicherweise als die zum linearen Ausgleichsproblem $\min_{\boldsymbol{\alpha} \in \mathbb{R}^{n+1}} f(\boldsymbol{\alpha})$ gehörende Normalengleichung bezeichnet.

Beispiel 1.4 Die Konvektions-Diffusions-Gleichung

Die Konvektions-Diffusions-Gleichung wird häufig als Modellproblem bei der Entwicklung neuartiger Verfahren innerhalb der Strömungsmechanik genutzt, da sie eine strukturelle Ähnlichkeit zu den Euler- und Navier-Stokes-Gleichungen aufweist. In Anlehnung an die Arbeiten [67] und [50] betrachten wir die stationäre Konvektions-Diffusions-Gleichung auf dem Einheitsquadrat $\Omega = (0,1) \times (0,1)$ in der Form

$$\begin{aligned}
 \boldsymbol{\beta} \cdot \nabla u(x,y) - \varepsilon \Delta u(x,y) &= 0 && \text{für } (x,y) \in \Omega \subset \mathbb{R}^2, \\
 u(x,y) &= x^2 + y^2 && \text{für } (x,y) \in \partial\Omega
 \end{aligned} \tag{1.0.8}$$

mit $\boldsymbol{\beta} = (\cos \alpha, \sin \alpha)^T$, $\alpha = 45^\circ$ und $\varepsilon \in \mathbb{R}_0^+$.

Wir nutzen die in Beispiel 1.1 vorgestellte äquidistante Diskretisierung des Einheitsquadrates. Um im Grenzfall eines verschwindenden Diffusionsparameters ε keine Entkopplung des diskreten Systems in der Form eines Schachbrettmusters zu erhalten, wird der Gradient ∇u innerhalb des konvektiven Anteils mittels einer einseitigen Differenz gemäß

$$\frac{\partial u}{\partial x}(x_{i+1}, y_j) \approx \frac{u_{i+1,j} - u_{i,j}}{h} \quad \text{und} \quad \frac{\partial u}{\partial y}(x_i, y_{j+1}) \approx \frac{u_{i,j+1} - u_{i,j}}{h},$$

diskretisiert. Diese Strategie wird innerhalb der Strömungsmechanik als Upwind-Methode bezeichnet. Die Verwendung der Approximation (1.0.2) für den Laplace-Operator ergibt in Kombination mit einer lexikographischen Anordnung ein Gleichungssystem der Form $\mathbf{A}\mathbf{u} = \mathbf{g}$ zur Berechnung des Vektors $\mathbf{u} := (u_1, \dots, u_{N^2})^T$. Die Besetzungsstruktur der Matrix stimmt hierbei mit der Struktur der Matrix innerhalb des Beispiels der Poisson-Gleichung überein. Die explizite Darstellung der Matrix lautet

$$\mathbf{A} = \begin{pmatrix} \mathbf{B} & -\varepsilon \mathbf{I} & & \\ \mathbf{D} & \ddots & \ddots & \\ & \ddots & \ddots & -\varepsilon \mathbf{I} \\ & & \mathbf{D} & \mathbf{B} \end{pmatrix} \in \mathbb{R}^{N^2 \times N^2} \quad (1.0.9)$$

mit

$$\mathbf{B} = \begin{pmatrix} 4\varepsilon + h(\cos \alpha + \sin \alpha) & -\varepsilon & & \\ -\varepsilon - h \cos \alpha & \ddots & \ddots & \\ & \ddots & \ddots & -\varepsilon \\ & & -\varepsilon - h \cos \alpha & 4\varepsilon + h(\cos \alpha + \sin \alpha) \end{pmatrix} \in \mathbb{R}^{N \times N}$$

und

$$\mathbf{D} = \begin{pmatrix} -\varepsilon - h \sin \alpha & & \\ & \ddots & \\ & & -\varepsilon - h \sin \alpha \end{pmatrix} \in \mathbb{R}^{N \times N},$$

wobei $\mathbf{I} \in \mathbb{R}^{N \times N}$ wiederum die Einheitsmatrix repräsentiert. Die rechte Seite $\mathbf{g}$ hängt in diesem Fall ausschließlich von den über (1.0.8)₂ gegebenen Randbedingungen ab.

Bemerkung:

Die Matrix $\mathbf{A}$ ist stets unsymmetrisch und schwachbesetzt.

2 Grundlagen der linearen Algebra

Die Herleitung und Analyse der im folgenden betrachteten Verfahren zur Lösung linearer Gleichungssysteme basiert auf den in diesem Kapitel dargestellten Grundlagen. Wir betrachten hierbei stets Abbildungen zwischen reellen beziehungsweise komplexen linearen Räumen, die oftmals auch als Vektorräume über $\mathbb{R}$ respektive $\mathbb{C}$ bezeichnet werden. Derartige lineare Räume, wie zum Beispiel der $\mathbb{R}^n$ oder $\mathbb{C}^n$ stellen nichtleere Mengen dar, deren Elemente durch eine Addition verknüpft und mit Skalaren $\lambda \in \mathbb{R}$ oder $\mathbb{C}$ multipliziert werden können, die den Vektorraumaxiomen genügen. Diese Axiome garantieren lediglich, daß mit Addition, Subtraktion und Multiplikation wie gewohnt gerechnet werden kann.

2.1 Vektornormen und Skalarprodukt

Notwendige Begriffe wie Orthogonalität, Länge, Abstand und Konvergenz basieren auf Skalarprodukten beziehungsweise Normen. Diese Abbildungen werden wir daher in diesem Abschnitt einführen und einige wesentliche Aussagen präsentieren.

Definition 2.1 Sei X ein komplexer beziehungsweise reeller linearer Raum. Eine Abbildung

$$\|\cdot\| : X \longrightarrow \mathbb{R}$$

mit den Eigenschaften

- | | | |
|------|--|-----------------------|
| (N1) | $\ \mathbf{x}\ \geq 0$ | (Positivität) |
| (N2) | $\ \mathbf{x}\ = 0 \Leftrightarrow \mathbf{x} = \mathbf{0}$ | (Definitheit) |
| (N3) | $\ \alpha \cdot \mathbf{x}\ = \alpha \cdot \ \mathbf{x}\ \quad \forall \mathbf{x} \in X, \forall \alpha \in \mathbb{C} \text{ (bzw. } \mathbb{R})$ | (Homogenität) |
| (N4) | $\ \mathbf{x} + \mathbf{y}\ \leq \ \mathbf{x}\ + \ \mathbf{y}\ \quad \forall \mathbf{x}, \mathbf{y} \in X$ | (Dreiecksungleichung) |

nennt man eine Norm auf X . Ein linearer Raum X mit einer Norm heißt normierter Raum. Falls $X = \mathbb{C}^n$ beziehungsweise $X = \mathbb{R}^n$ gilt, so wird die Norm auch als Vektornorm bezeichnet.

Die Positivität einer Norm kann hierbei auch aus den Axiomen (N3) und (N4) hergeleitet werden und bräuchte daher aus formalen Gründen nicht innerhalb der Definition aufgeführt werden. Wir verwenden dennoch die obige Formulierung, um diese Eigenschaft der Norm explizit vorliegen zu haben.

Auf $\mathbb{C}^n$ beziehungsweise $\mathbb{R}^n$ sind Normen zum Beispiel wie folgt gegeben:

$$(a) \quad \|\mathbf{x}\|_1 := \sum_{i=1}^n |x_i| \quad (\text{Betragssummennorm})$$

$$(b) \quad \|\mathbf{x}\|_2 := \left(\sum_{i=1}^n |x_i|^2 \right)^{\frac{1}{2}} \quad (\text{Euklidische Norm})$$

$$(c) \quad \|\mathbf{x}\|_\infty := \max_{i=1, \dots, n} |x_i| \quad (\text{Maximumnorm})$$

Wir kommen nun zur Einführung des benötigten Konvergenzbegriffs.

Definition 2.2 Eine Folge $\{\mathbf{x}_n\}_{n \in \mathbb{N}}$ von Elementen aus einem normierten Raum X heißt konvergent mit dem Grenzelement $\mathbf{x} \in X$, wenn zu jedem $\varepsilon > 0$ eine natürliche Zahl $N = N(\varepsilon)$ existiert, so daß

$$\|\mathbf{x}_n - \mathbf{x}\| < \varepsilon \quad \forall n \geq N$$

gilt. Eine Folge, die nicht konvergiert, heißt divergent.

Definition 2.3 Sei X ein normierter Raum. Eine Folge $\{\mathbf{x}_n\}_{n \in \mathbb{N}}$ aus X heißt Cauchy-Folge, wenn es zu jedem $\varepsilon > 0$ ein $N = N(\varepsilon) \in \mathbb{N}$ derart gibt, daß

$$\|\mathbf{x}_n - \mathbf{x}_m\| < \varepsilon \quad \forall n, m \geq N$$

gilt.

Der anschließende Satz liefert einen generellen Zusammenhang zwischen konvergenten Folgen und Cauchy-Folgen.

Satz 2.4 Sei X ein normierter Raum. Jede konvergente Folge $\{\mathbf{x}_n\}_{n \in \mathbb{N}}$ aus X ist eine Cauchy-Folge.

Beweis:

Sei $\varepsilon > 0$ gegeben und $\mathbf{x}$ das Grenzelement der Folge $\{\mathbf{x}_n\}_{n \in \mathbb{N}}$

$$\Rightarrow \exists N = N\left(\frac{\varepsilon}{2}\right) \in \mathbb{N} \text{ mit } \|\mathbf{x}_n - \mathbf{x}\| < \frac{\varepsilon}{2} \quad \forall n \geq N$$

$$\Rightarrow \|\mathbf{x}_n - \mathbf{x}_m\| \leq \|\mathbf{x}_n - \mathbf{x}\| + \|\mathbf{x} - \mathbf{x}_m\| < \varepsilon \quad \forall n, m \geq N. \quad \square$$

Definition 2.5 Sei X ein komplexer oder reeller linearer Raum. Eine Abbildung

$$(\cdot, \cdot) : X \times X \longrightarrow \mathbb{C}$$

mit den Eigenschaften

$$(H1) \quad (\mathbf{x}, \mathbf{x}) \in \mathbb{R}_0^+ \quad \forall \mathbf{x} \in X \quad (\text{Positivität})$$

$$(H2) \quad (\mathbf{x}, \mathbf{x}) = 0 \Leftrightarrow \mathbf{x} = \mathbf{0} \quad (\text{Definitheit})$$

$$(H3) \quad (\mathbf{x}, \mathbf{y}) = \overline{(\mathbf{y}, \mathbf{x})} \quad \forall \mathbf{x}, \mathbf{y} \in X \quad (\text{Symmetrie})$$

$$(H4) \quad (\alpha \mathbf{x} + \beta \mathbf{y}, \mathbf{z}) = \alpha(\mathbf{x}, \mathbf{z}) + \beta(\mathbf{y}, \mathbf{z}) \quad \forall \mathbf{x}, \mathbf{y}, \mathbf{z} \in X \quad \alpha, \beta \in \mathbb{C} \quad (\text{Linearität})$$

heißt Skalarprodukt oder inneres Produkt auf X . Ein linearer Raum X versehen mit einem Skalarprodukt heißt Prä-Hilbert-Raum.

Bemerkung:

Es gilt zudem

$$\begin{aligned}
 (\text{H4}') \quad (\mathbf{x}, \alpha \mathbf{y} + \beta \mathbf{z}) &\stackrel{(\text{H3})}{=} \overline{(\alpha \mathbf{y} + \beta \mathbf{z}, \mathbf{x})} \stackrel{(\text{H4})}{=} \overline{\alpha (\mathbf{y}, \mathbf{x}) + \beta (\mathbf{z}, \mathbf{x})} \\
 &\stackrel{(\text{H3})}{=} \overline{\alpha} (\mathbf{x}, \mathbf{y}) + \overline{\beta} (\mathbf{x}, \mathbf{z}).
 \end{aligned}$$

Diese Eigenschaft wird als Antilinearität bezeichnet.

Satz 2.6 Auf jedem Prä-Hilbert-Raum X ist durch

$$\|\mathbf{x}\| := \sqrt{(\mathbf{x}, \mathbf{x})}, \quad \mathbf{x} \in X$$

eine Norm erklärt.

Beweis:

Die ersten drei Normeigenschaften ergeben sich durch direktes Nachrechnen. Die Eigenschaft (N4) folgt durch Anwendung der Cauchy-Schwarz'schen Ungleichung

$$|(\mathbf{x}, \mathbf{y})| \leq \|\mathbf{x}\| \|\mathbf{y}\|.$$

□

Üblicherweise wird auf dem $\mathbb{C}^n$ das euklidische Skalarprodukt gemäß

$$(\mathbf{x}, \mathbf{y})_2 := \sum_{i=1}^n x_i \overline{y_i} = \mathbf{y}^* \mathbf{x}$$

genutzt. Hierbei gilt der Zusammenhang

$$\|\mathbf{x}\|_2 := \sqrt{(\mathbf{x}, \mathbf{x})_2}.$$

Satz 2.7 Sei X ein normierter Raum und $\{\mathbf{x}_n\}_{n \in \mathbb{N}}$ eine konvergente Folge in X , dann ist das Grenzelement eindeutig bestimmt.

Beweis:

Seien $\mathbf{x}$ und $\mathbf{y}$ zwei Grenzelemente, d. h. gelte $\mathbf{x}_n \rightarrow \mathbf{x}$ für $n \rightarrow \infty$ und $\mathbf{x}_n \rightarrow \mathbf{y}$ für $n \rightarrow \infty$, dann folgt

$$\begin{aligned}
 0 \leq \|\mathbf{x} - \mathbf{y}\| &= \|\mathbf{x} - \mathbf{x}_n + \mathbf{x}_n - \mathbf{y}\| \\
 &\leq \underbrace{\|\mathbf{x} - \mathbf{x}_n\|}_{\xrightarrow{n \rightarrow \infty} 0} + \underbrace{\|\mathbf{x}_n - \mathbf{y}\|}_{\xrightarrow{n \rightarrow \infty} 0} \xrightarrow{n \rightarrow \infty} 0 \\
 \Rightarrow 0 \leq \|\mathbf{x} - \mathbf{y}\| &\leq 0 \\
 \Rightarrow \|\mathbf{x} - \mathbf{y}\| = 0 &\stackrel{(\text{N2})}{\Rightarrow} \mathbf{x} - \mathbf{y} = 0 \Rightarrow \mathbf{x} = \mathbf{y}.
 \end{aligned}$$

□

Definition 2.8 Zwei Normen $\|\cdot\|_a$ und $\|\cdot\|_b$ auf einem linearen Raum X heißen äquivalent, wenn jede Folge genau dann bezüglich $\|\cdot\|_a$ konvergiert, wenn sie bezüglich $\|\cdot\|_b$ konvergiert.

Satz und Definition 2.9 Zwei Normen $\|\cdot\|_a$ und $\|\cdot\|_b$ auf einem linearen Raum X sind genau dann äquivalent, wenn es reelle Zahlen $\alpha, \beta > 0$ gibt, so daß

$$\alpha\|\mathbf{x}\|_b \leq \|\mathbf{x}\|_a \leq \beta\|\mathbf{x}\|_b \quad \forall \mathbf{x} \in X \quad (2.1.1)$$

gilt. Die größte derartige Zahl α und die kleinste derartige Zahl β werden als Äquivalenzkonstanten bezeichnet.

Beweis:

„ $\Leftarrow$ “

Sei $\{\mathbf{x}_n\}_{n \in \mathbb{N}}$ eine beliebige Folge mit $\mathbf{x}_n \xrightarrow{\|\cdot\|_b} \mathbf{x}$ für $n \rightarrow \infty$, dann folgt mit (2.1.1)

$$\|\mathbf{x}_n - \mathbf{x}\|_a \leq \beta\|\mathbf{x}_n - \mathbf{x}\|_b \rightarrow 0 \text{ für } n \rightarrow \infty$$

und somit die geforderte Konvergenz der Folge in der Norm $\|\cdot\|_a$. Die zweite Ungleichung in (2.1.1) liefert auf analoge Weise die Konvergenz einer Folge in der Norm $\|\cdot\|_b$ aus der Konvergenz der Folge in der Norm $\|\cdot\|_a$.

„ $\Rightarrow$ “

Annahme: Es existiert kein $\beta \in \mathbb{R}^+$ mit $\|\mathbf{x}\|_a \leq \beta$ für alle $\mathbf{x} \in X$ mit $\|\mathbf{x}\|_b = 1$. Dann existiert eine Folge $\{\mathbf{x}_n\}_{n \in \mathbb{N}}$ mit

$$\|\mathbf{x}_n\|_a \geq n^2 \quad \text{und} \quad \|\mathbf{x}_n\|_b = 1.$$

Sei

$$\mathbf{y}_n := \frac{\mathbf{x}_n}{n},$$

dann erhalten wir

$$\|\mathbf{y}_n\|_a = \frac{1}{n}\|\mathbf{x}_n\|_a \geq n \quad \text{und} \quad \|\mathbf{y}_n\|_b = \frac{1}{n}. \quad (2.1.2)$$

Somit folgt $\mathbf{y}_n \xrightarrow{\|\cdot\|_b} \mathbf{0}$, für $n \rightarrow \infty$, wodurch die Folge aufgrund der vorausgesetzten Äquivalenz der Normen $\|\cdot\|_a$ und $\|\cdot\|_b$ auch in der Norm $\|\cdot\|_a$ konvergiert. Hiermit ergibt sich direkt der angestrebte Widerspruch zur Eigenschaft (2.1.2).

Folglich existiert ein $\beta \in \mathbb{R}^+$ mit $\|\mathbf{x}\|_a \leq \beta$ für alle $\mathbf{x} \in X$ mit $\|\mathbf{x}\|_b = 1$. Sei nun $\mathbf{y} \in X \setminus \{\mathbf{0}\}$, dann erhalten wir durch

$$\|\mathbf{y}\|_a = \left\| \|\mathbf{y}\|_b \frac{\mathbf{y}}{\|\mathbf{y}\|_b} \right\|_a = \|\mathbf{y}\|_b \left\| \frac{\mathbf{y}}{\|\mathbf{y}\|_b} \right\|_a \leq \beta\|\mathbf{y}\|_b \quad (2.1.3)$$

die Übertragung der Eigenschaft auf alle Vektoren aus $X \setminus \{\mathbf{0}\}$. Da das Nullelement stets der Gleichung (2.1.1) genügt, folgt

$$\|\mathbf{y}\|_a \leq \beta\|\mathbf{y}\|_b \quad \text{für alle } \mathbf{y} \in X.$$

Die zweite Abschätzung ergibt sich analog. □

Aufgrund des Satzes 2.9 wird die Definition äquivalenter Normen oftmals auf der Basis der Ungleichungen (2.1.1) vorgenommen. Als direkte Folgerung dieser Ungleichungen erhalten wir die folgende Aussage.

Korollar 2.10 Die Grenzelemente einer Folge stimmen bezüglich zweier äquivalenter Normen überein.

Satz 2.11 Auf einem endlichdimensionalen reellen oder komplexen linearen Raum sind alle Normen äquivalent.

Beweis:

Für den Nachweis der obigen Behauptung genügt es zu zeigen, daß eine bestimmte Norm äquivalent zu jeder Norm ist. Sei $X := \text{span}\{\mathbf{u}_1, \dots, \mathbf{u}_m\}$ ein m -dimensionaler linearer Raum, dann läßt sich jedes $\mathbf{x} \in X$ in der Form

$$\mathbf{x} = \sum_{i=1}^m \alpha_i \mathbf{u}_i, \quad \alpha_i \in \mathbb{C}$$

darstellen. Definieren wir auf X die Norm

$$\|\mathbf{x}\|_{\max} := \max_{i=1, \dots, m} |\alpha_i|,$$

dann folgt für jede weitere Norm $\|\cdot\|$ auf X die Abschätzung

$$\begin{aligned} \|\mathbf{x}\| &= \left\| \sum_{i=1}^m \alpha_i \mathbf{u}_i \right\| \stackrel{(\text{N4})}{\leq} \sum_{i=1}^m \|\alpha_i \mathbf{u}_i\| \stackrel{(\text{N3})}{=} \sum_{i=1}^m |\alpha_i| \|\mathbf{u}_i\| \\ &\leq \max_{i=1, \dots, m} |\alpha_i| \underbrace{\sum_{i=1}^m \|\mathbf{u}_i\|}_{\beta :=} = \beta \|\mathbf{x}\|_{\max}. \end{aligned} \quad (2.1.4)$$

Die zweite Ungleichung werden wir mittels des folgenden Widerspruchsbeweises nachweisen.

Annahme: Es existiert kein $\alpha \in \mathbb{R}^+$ mit $\|\mathbf{x}\|_{\max} \leq \alpha \|\mathbf{x}\|$ für alle $\mathbf{x} \in X$.

Wie wir bereits in Gleichung (2.1.3) gesehen haben, ist diese Eigenschaft äquivalent zur Aussage:

Es existiert kein $\alpha \in \mathbb{R}^+$ mit $\|\mathbf{x}\|_{\max} \leq \alpha$ für alle $\mathbf{x} \in X$ mit $\|\mathbf{x}\| = 1$.

Somit existiert eine Folge $\{\mathbf{x}_n\}_{n \in \mathbb{N}}$ mit

$$\|\mathbf{x}_n\|_{\max} \geq n \text{ und } \|\mathbf{x}_n\| = 1.$$

Sei

$$\mathbf{y}_n := \frac{\mathbf{x}_n}{\|\mathbf{x}_n\|_{\max}}, \quad \text{dann gilt } \|\mathbf{y}_n\|_{\max} = 1 \quad \text{für alle } n \in \mathbb{N}. \quad (2.1.5)$$

Wir schreiben

$$\mathbf{y}_n = \sum_{i=1}^m \alpha_{i,n} \mathbf{u}_i$$

und erhalten mit (2.1.5) $\max_{i=1,\dots,m} |\alpha_{i,n}| = 1$ für alle $n \in \mathbb{N}$. Mit $\{\alpha_{1,n}\}_{n \in \mathbb{N}}, \dots, \{\alpha_{m,n}\}_{n \in \mathbb{N}}$ liegen daher m beschränkte Folgen komplexer Zahlen vor, so daß für $i = 1, \dots, m$ aufgrund des Satzes von Bolzano-Weierstraß je eine konvergente Teilfolge

$$\alpha_{i,n(j)} \xrightarrow{j \rightarrow \infty} \tilde{\alpha}_i \text{ mit } n(j+1) > n(j)$$

existiert. Definieren wir

$$\mathbf{y} := \sum_{i=1}^m \tilde{\alpha}_i \mathbf{u}_i,$$

dann folgt für $\mathbf{y}_{n(j)} := \sum_{i=1}^m \alpha_{i,n(j)} \mathbf{u}_i$

$$\|\mathbf{y}_{n(j)} - \mathbf{y}\|_{\max} = \max_{i=1,\dots,m} |\alpha_{i,n(j)} - \tilde{\alpha}_i| \xrightarrow{j \rightarrow \infty} 0. \quad (2.1.6)$$

Unter Verwendung der Ungleichung (2.1.4) erhalten wir

$$\mathbf{y}_{n(j)} \xrightarrow{\|\cdot\|} \mathbf{y} \text{ für } j \rightarrow \infty,$$

so daß sich mit (2.1.5)

$$\|\mathbf{y}_{n(j)}\| = \left\| \frac{\mathbf{x}_{n(j)}}{\|\mathbf{x}_{n(j)}\|_{\max}} \right\| = \frac{1}{\|\mathbf{x}_{n(j)}\|_{\max}} \leq \frac{1}{n(j)} \xrightarrow{j \rightarrow \infty} 0$$

und folglich $\mathbf{y} = \mathbf{0}$ ergibt. Hierdurch erhalten wir mittels $\|\mathbf{y}_{n(j)}\|_{\max} \xrightarrow{j \rightarrow \infty} 0$ einen Widerspruch zu $\|\mathbf{y}_n\|_{\max} = 1$ für alle $n \in \mathbb{N}$. $\square$

Für die aufgeführten Vektornormen erhalten wir die Äquivalenzkonstanten für $\mathbf{x} \in \mathbb{R}^n$ gemäß

$$\|\mathbf{x}\|_2 \leq \|\mathbf{x}\|_1 \leq \sqrt{n} \|\mathbf{x}\|_2, \quad (2.1.7)$$

$$\|\mathbf{x}\|_{\infty} \leq \|\mathbf{x}\|_2 \leq \sqrt{n} \|\mathbf{x}\|_{\infty}, \quad (2.1.8)$$

und

$$\frac{1}{n} \|\mathbf{x}\|_1 \leq \|\mathbf{x}\|_{\infty} \leq \|\mathbf{x}\|_1. \quad (2.1.9)$$

Definition 2.12 Eine Teilmenge V eines normierten Raumes X heißt vollständig, wenn jede Cauchy-Folge aus V gegen ein Grenzelement aus V konvergiert. Ein vollständiger normierter Raum heißt Banach-Raum.

Bei den betrachteten linearen Gleichungssystemen betrachten wir stets endlichdimensionale normierte Räume. Für diese Räume ergibt sich die folgende hilfreiche Eigenschaft.

Satz 2.13 *Jeder endlichdimensionale normierte Raum ist ein Banach-Raum.*

Beweis:

Sei X ein endlichdimensionaler normierter Raum. Laut Satz 2.11 sind auf X alle Normen äquivalent, so daß der Nachweis für eine beliebige Norm erbracht werden kann.

Sei $X := \text{span}\{\mathbf{u}_1, \dots, \mathbf{u}_k\}$, dann läßt sich jedes $\mathbf{x} \in X$ in der Form

$$\mathbf{x} = \sum_{i=1}^k \alpha_i \mathbf{u}_i \text{ mit } \alpha_i \in \mathbb{C}$$

schreiben. Wir betrachten nun wiederum die Norm

$$\|\mathbf{x}\|_{\max} := \max_{i=1, \dots, k} |\alpha_i|.$$

Sei $\{\mathbf{x}_n\}_{n \in \mathbb{N}}$ eine Cauchy-Folge in X , dann schreiben wir

$$\mathbf{x}_n = \sum_{i=1}^k \alpha_{i,n} \mathbf{u}_i \text{ mit } \alpha_{i,n} \in \mathbb{C}$$

und erhalten

$$\max_{i=1, \dots, k} |\alpha_{i,n} - \alpha_{i,m}| = \|\mathbf{x}_n - \mathbf{x}_m\|_{\max} \rightarrow 0 \text{ für } n, m \rightarrow \infty.$$

Folglich stellen die k Koeffizientenfolgen jeweils Cauchy-Folgen in $\mathbb{C}$ dar, so daß aufgrund der Vollständigkeit des Raumes der komplexen Zahlen

$$\alpha_{i,n} \xrightarrow{n \rightarrow \infty} \tilde{\alpha}_i \in \mathbb{C} \text{ für } i = 1, \dots, k$$

gilt. Mit $\tilde{\mathbf{x}} = \sum_{i=1}^k \tilde{\alpha}_i \mathbf{u}_i$ liegt wegen

$$\|\mathbf{x}_n - \tilde{\mathbf{x}}\|_{\max} = \max_{i=1, \dots, k} |\alpha_{i,n} - \tilde{\alpha}_i| \xrightarrow{n \rightarrow \infty} 0$$

der gesuchte Vektor vor. □

2.2 Lineare Operatoren, Matrizen und Matrixnormen

Die Eigenschaften der Matrix eines linearen Gleichungssystem sind wesentlich für die Auswahl eines geeigneten iterativen Verfahrens und bestimmen in der Regel zudem das Konvergenzverhalten der gewählten Methode. Neben der Einführung spezieller Klassen von Matrizen und der Festlegung unterschiedlicher Matrixnormen werden wir in diesem Abschnitt den Spektralradius erläutern und dessen Zusammenhang zu den Matrixnormen diskutieren, der sich als fundamental für Konvergenzaussagen von Splitting-Methoden erweisen wird.

Definition 2.14 Seien X, Y normierte Räume mit den Normen $\|\cdot\|_X$ beziehungsweise $\|\cdot\|_Y$. Ein Operator $\mathbf{A} : X \rightarrow Y$ heißt

- (a) stetig an der Stelle $\mathbf{x} \in X$, falls für alle Folgen $\{\mathbf{x}_n\}_{n \in \mathbb{N}}$ aus X mit $\mathbf{x}_n \xrightarrow{\|\cdot\|_X} \mathbf{x}$ für $n \rightarrow \infty$

$$\mathbf{A}\mathbf{x}_n \xrightarrow{\|\cdot\|_Y} \mathbf{A}\mathbf{x} \text{ für } n \rightarrow \infty$$

folgt.

(b) stetig, falls $\mathbf{A}$ an allen Stellen $\mathbf{x} \in X$ stetig ist.

(c) linear, falls

$$\mathbf{A}(\alpha \mathbf{x} + \beta \mathbf{y}) = \alpha \mathbf{A}\mathbf{x} + \beta \mathbf{A}\mathbf{y} \quad \forall \mathbf{x}, \mathbf{y} \in X \quad \forall \alpha, \beta \in \mathbb{C}$$

gilt.

(d) beschränkt, wenn $\mathbf{A}$ linear ist und ein $C \geq 0$ existiert, so daß

$$\|\mathbf{A}\mathbf{x}\|_Y \leq C \|\mathbf{x}\|_X \quad \forall \mathbf{x} \in X$$

gilt. Jede Zahl C mit dieser Eigenschaft heißt Schranke von $\mathbf{A}$.

Wir beschäftigen uns im weiteren mit sogenannten induzierten Matrixnormen, die jeweils auf der Grundlage einer Vektornorm durch den folgenden Satz festgelegt werden.

Satz und Definition 2.15 *Ein linearer Operator $\mathbf{A} : X \rightarrow Y$ ist genau dann beschränkt, wenn*

$$\|\mathbf{A}\| := \sup_{\|\mathbf{x}\|_X=1} \|\mathbf{A}\mathbf{x}\|_Y < \infty$$

gilt. $\|\mathbf{A}\|$ ist die kleinste Schranke von $\mathbf{A}$ und heißt Norm des Operators.

Beweis:

„ $\Rightarrow$ “

Sei $\mathbf{A}$ beschränkt, dann existiert insbesondere ein $C \geq 0$ mit

$$\|\mathbf{A}\mathbf{x}\|_Y \leq C \text{ für alle } \mathbf{x} \in X \text{ mit } \|\mathbf{x}\|_X = 1$$

und wir erhalten

$$\|\mathbf{A}\| = \sup_{\|\mathbf{x}\|_X=1} \|\mathbf{A}\mathbf{x}\|_Y \leq C < \infty.$$

Insbesondere stellt $\|\mathbf{A}\|$ demzufolge die kleinste Schranke von $\mathbf{A}$ dar.

„ $\Leftarrow$ “

Sei $\mathbf{A}$ linear und gelte $\|\mathbf{A}\| < \infty$

Betrachten wir ein beliebiges $\mathbf{x} \in X \setminus \{\mathbf{0}\}$, dann folgt

$$\begin{aligned} \|\mathbf{A}\mathbf{x}\|_Y &= \left\| \|\mathbf{x}\|_X \mathbf{A} \left(\frac{\mathbf{x}}{\|\mathbf{x}\|_X} \right) \right\|_Y \\ &= \|\mathbf{x}\|_X \left\| \mathbf{A} \left(\frac{\mathbf{x}}{\|\mathbf{x}\|_X} \right) \right\|_Y \\ &\leq \|\mathbf{x}\|_X \sup_{\|\mathbf{z}\|_X=1} \|\mathbf{A}\mathbf{z}\|_Y \\ &= \|\mathbf{x}\|_X \|\mathbf{A}\|. \end{aligned}$$

Somit ist $\mathbf{A}$ beschränkt mit Schranke $\|\mathbf{A}\|$.

□

Die Eigenschaft

$$\|\mathbf{A}\mathbf{x}\| \leq \|\mathbf{A}\|\|\mathbf{x}\|$$

wird als Verträglichkeitsbedingung bezeichnet.

Stetigkeit und Beschränktheit stellen im Kontext linearer Operatoren äquivalente Begriffe dar. Diese Tatsache wird durch den folgenden Satz nachgewiesen.

Satz 2.16 *Für einen linearen Operator $\mathbf{A} : X \rightarrow Y$ sind die folgenden drei Eigenschaften äquivalent:*

(a) $\mathbf{A}$ ist stetig an der Stelle $\mathbf{x} = \mathbf{0}$.

(b) $\mathbf{A}$ ist stetig.

(c) $\mathbf{A}$ ist beschränkt.

Beweis:

„(a) $\Rightarrow$ (b)“

Sei $\mathbf{x} \in X$ gegeben und $\{\mathbf{x}_n\}_{n \in \mathbb{N}}$ eine beliebige Folge aus X mit $\mathbf{x}_n \rightarrow \mathbf{x}$, $n \rightarrow \infty$, dann folgt mit der Stetigkeit an der Stelle $\mathbf{x} = \mathbf{0}$ und der Linearität des Operators die Gleichung

$$\mathbf{A}\mathbf{x}_n = \underbrace{\mathbf{A}(\mathbf{x}_n - \mathbf{x})}_{\xrightarrow{\|\cdot\|_Y} 0, n \rightarrow \infty} + \mathbf{A}\mathbf{x} \xrightarrow{\|\cdot\|_Y} \mathbf{A}\mathbf{x} \text{ für } n \rightarrow \infty.$$

„(b) $\Rightarrow$ (c)“

Annahme: $\mathbf{A}$ ist stetig und nicht beschränkt.

Dann existiert eine Folge $\{\mathbf{x}_n\}_{n \in \mathbb{N}}$ aus X mit $\|\mathbf{x}_n\|_X = 1$ und $\|\mathbf{A}\mathbf{x}_n\|_Y \geq n$. Wir definieren

$$\mathbf{y}_n := \frac{\mathbf{x}_n}{\|\mathbf{A}\mathbf{x}_n\|_Y}$$

und erhalten hiermit

$$\|\mathbf{y}_n\|_X = \frac{\|\mathbf{x}_n\|_X}{\|\mathbf{A}\mathbf{x}_n\|_Y} = \frac{1}{\|\mathbf{A}\mathbf{x}_n\|_Y} \leq \frac{1}{n}.$$

Folglich konvergiert die Folge $\mathbf{y}_n$ in der Norm auf X gegen das Nullelement, und es folgt aus der Stetigkeit des Operators

$$\mathbf{A}\mathbf{y}_n \xrightarrow{\|\cdot\|_Y} \mathbf{A}(\mathbf{0}) = \mathbf{0} \text{ für } n \rightarrow \infty,$$

so daß ein Widerspruch zu

$$\|\mathbf{A}\mathbf{y}_n\|_Y = \frac{\|\mathbf{A}\mathbf{x}_n\|_Y}{\|\mathbf{A}\mathbf{x}_n\|_Y} = 1 \text{ für alle } n \in \mathbb{N}$$

vorliegt.

„(c) $\Rightarrow$ (a)“

Sei $\mathbf{A}$ beschränkt und $\{\mathbf{x}_n\}_{n \in \mathbb{N}}$ eine Folge aus X mit $\mathbf{x}_n \xrightarrow{\|\cdot\|_X} \mathbf{0}$, für $n \rightarrow \infty$. Dann folgt

$$\|\mathbf{A}\mathbf{x}_n\|_Y \leq \|\mathbf{A}\| \|\mathbf{x}_n\|_X \rightarrow 0 \text{ für } n \rightarrow \infty$$

und hieraus mit

$$\mathbf{A}\mathbf{x}_n \xrightarrow{\|\cdot\|_Y} \mathbf{0} = \mathbf{A}(\mathbf{0}) \text{ für } n \rightarrow \infty$$

die Stetigkeit des Operators an der Stelle $\mathbf{x} = \mathbf{0}$. □

Bei der Betrachtung von Kompositionen linearer Abbildungen erweist sich der folgende Satz zur Abschätzung der Norm der Komposition als hilfreich.

Satz 2.17 *Seien $\mathbf{A} : X \rightarrow Y$ und $\mathbf{B} : Y \rightarrow Z$ beschränkte lineare Operatoren, dann ist $\mathbf{BA} : X \rightarrow Z$ beschränkt mit*

$$\|\mathbf{BA}\| \leq \|\mathbf{B}\| \|\mathbf{A}\|. \quad (2.2.1)$$

Beweis:

Mit

$$\|\mathbf{BA}\mathbf{x}\|_Z \leq \|\mathbf{B}\| \|\mathbf{A}\mathbf{x}\|_Y \leq \|\mathbf{B}\| \|\mathbf{A}\| \|\mathbf{x}\|_X$$

folgt der Nachweis direkt aus

$$\|\mathbf{BA}\| = \sup_{\|\mathbf{x}\|_X=1} \|\mathbf{BA}\mathbf{x}\|_Z \leq \sup_{\|\mathbf{x}\|_X=1} \|\mathbf{B}\| \|\mathbf{A}\| \|\mathbf{x}\|_X = \|\mathbf{B}\| \|\mathbf{A}\|.$$

□

Die Eigenschaft (2.2.1) wird auch als Submultiplikativität der zugrundeliegenden Norm bezeichnet. Wir werden uns im folgenden, solange nicht ausdrücklich erwähnt, stets auf Normen beschränken, die auf der Grundlage der Definition 2.15 festgelegt sind. Im Kontext der betrachteten Matrizen sind jedoch auch hiervon abweichende Normen denkbar, die nicht der Ungleichung (2.2.1) genügen.

Jede lineare Abbildung zwischen endlichdimensionalen Vektorräumen kann durch eine Matrix

$$\mathbf{A} = \begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & \ddots & \vdots \\ a_{m1} & \dots & a_{mn} \end{pmatrix} \in \mathbb{C}^{m \times n}$$

repräsentiert werden. Die Menge aller Matrizen wollen wir nun vorab in unterschiedliche Klassen unterteilen.

Definition 2.18 Zu einer gegebenen Matrix

$$\mathbf{A} = \begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & \ddots & \vdots \\ a_{n1} & \dots & a_{nn} \end{pmatrix} \in \mathbb{R}^{n \times n}$$

heißt

$$\mathbf{A}^T = \begin{pmatrix} a_{11} & \dots & a_{n1} \\ \vdots & \ddots & \vdots \\ a_{1n} & \dots & a_{nn} \end{pmatrix} \in \mathbb{R}^{n \times n}$$

die zu $\mathbf{A}$ transponierte Matrix.

Definition 2.19 Eine Matrix $A \in \mathbb{R}^{n \times n}$ heißt

- (a) symmetrisch, falls $A^T = A$ gilt,
- (b) orthogonal, falls $A^T A = I$ gilt.

Definition 2.20 Zu gegebener Matrix

$$A = \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & \ddots & \vdots \\ a_{n1} & \cdots & a_{nn} \end{pmatrix} \in \mathbb{C}^{n \times n}$$

heißt

$$A^* = \begin{pmatrix} \overline{a_{11}} & \cdots & \overline{a_{n1}} \\ \vdots & \ddots & \vdots \\ \overline{a_{1n}} & \cdots & \overline{a_{nn}} \end{pmatrix} \in \mathbb{C}^{n \times n}$$

die zu A adjungierte Matrix.

Definition 2.21 Eine Matrix $A \in \mathbb{C}^{n \times n}$ heißt

- (a) hermitesch, falls $A^* = A$ gilt,
- (b) unitär, falls $A^* A = I$ gilt,
- (c) normal, falls $A^* A = A A^*$ gilt,
- (d) ähnlich zur Matrix $B \in \mathbb{C}^{n \times n}$, falls eine reguläre Matrix $C \in \mathbb{C}^{n \times n}$ mit $B = C^{-1} A C$ existiert,
- (e) linke untere Dreiecksmatrix, falls $a_{ij} = 0 \quad \forall j > i$ gilt,
- (f) rechte obere Dreiecksmatrix, falls $a_{ij} = 0 \quad \forall j < i$ gilt,
- (g) Diagonalmatrix, falls $a_{ij} = 0 \quad \forall j \neq i$ gilt.

Definition 2.22 Sei $X = \mathbb{R}^n$ beziehungsweise $\mathbb{C}^n$. Eine Matrix $A : X \rightarrow X$ heißt

- (a) positiv semidefinit, falls $(Ax, x)_2 \geq 0$ für alle $x \in X$ gilt,
- (b) positiv definit, falls $(Ax, x)_2 > 0$ für alle $x \in X \setminus \{0\}$ gilt,
- (c) negativ semidefinit, falls $-A$ positiv semidefinit ist,
- (d) negativ definit, falls $-A$ positiv definit ist.

Definition 2.23 Zwei Gleichungssysteme heißen äquivalent, wenn ihre Lösungsmengen identisch sind.

Lemma 2.24 Sei $P \in \mathbb{C}^{n \times n}$ regulär und $A \in \mathbb{C}^{n \times n}$, dann sind die Gleichungssysteme $Ax = y$ und $PAx = Py$ äquivalent.

Beweis:

Aufgrund der Regularität der Matrix P folgt

$$Px = 0 \Leftrightarrow x = 0.$$

Damit erhalten wir die Behauptung direkt aus

$$P(Ax - y) = 0 \Leftrightarrow Ax - y = 0.$$

□

Lemma und Definition 2.25 Das lineare Gleichungssystem $Ax = b$ mit $A \in \mathbb{C}^{m \times n}$ ($\mathbb{R}^{m \times n}$) ist genau dann lösbar, wenn $\text{rang}(A) = \text{rang}(A, b)$ gilt, wobei $\text{rang}(A)$ die Dimension des Bildes von A darstellt, das heißt

$$\text{rang}(A) = \dim \text{bild}(A)$$

mit

$$\text{bild}(A) = \{y \in \mathbb{C}^m (\mathbb{R}^m) \mid \exists x \in \mathbb{C}^n (\mathbb{R}^n) \text{ mit } y = Ax\}.$$

Beweis:

„ $\Rightarrow$ “

Sei $Ax = b$, dann gilt $b \in \text{bild}(A)$. Hiermit erhalten wir $\text{rang}(A) = \text{rang}(A, b)$.

„ $\Leftarrow$ “

Sei $\text{rang}(A) = \text{rang}(A, b)$, dann läßt sich b in der Form $b = \sum_{i=1}^n x_i a_i$ mit $x_i \in \mathbb{C}$ schreiben, wobei a_i für $i = 1, \dots, n$ die i -te Spalte der Matrix A darstellt. Hieraus ergibt sich direkt

$$Ax = b \text{ mit } x = (x_1, \dots, x_n)^T.$$

□

Lemma 2.26 Seien $L, \tilde{L} \in \mathbb{C}^{n \times n}$ linke untere und $R, \tilde{R} \in \mathbb{C}^{n \times n}$ rechte obere Dreiecksmatrizen, dann sind

$$L\tilde{L} \text{ und } R\tilde{R}$$

ebenfalls linke untere beziehungsweise rechte obere Dreiecksmatrizen.

Beweis:

Sei $\bar{L} = (\bar{l}_{ij})_{i,j=1,\dots,n} = L\tilde{L}$, dann folgt für $j > i$

$$\bar{l}_{ij} = \sum_{m=1}^n l_{im} \tilde{l}_{mj} = \sum_{m=1}^{j-1} l_{im} \underbrace{\tilde{l}_{mj}}_{=0} + \sum_{m=j}^n \underbrace{l_{im}}_{=0} \tilde{l}_{mj} = 0.$$

Analog ergibt sich die Behauptung für die rechten oberen Dreiecksmatrizen.

□

Lemma 2.27 Seien $Q, \tilde{Q} \in \mathbb{R}^{n \times n} (\mathbb{C}^{n \times n})$ orthogonale (unitäre) Matrizen, dann ist auch

$$Q\tilde{Q}$$

orthogonal (unitär).

Beweis:

Bei dem Beweis beschränken wir uns auf orthogonale Matrizen. Der Nachweis der Behauptung für unitäre Matrizen verläuft analog. Aufgrund der Orthogonalität der Matrizen $\tilde{Q}$ und Q folgt

$$(\tilde{Q}Q)^T Q\tilde{Q} = \tilde{Q}^T Q^T Q\tilde{Q} = \tilde{Q}^T I\tilde{Q} = \tilde{Q}^T \tilde{Q} = I.$$

□

Wie bereits zuvor erwähnt, kann eine Matrixnorm mittels einer vorliegenden Vektornorm definiert werden. Ist $A \in \mathbb{C}^{n \times n}$ und $\|\cdot\|_a : \mathbb{C}^n \rightarrow \mathbb{R}$ eine Norm, dann bezeichnet man

$$\|A\|_a := \sup_{\|x\|_a=1} \|Ax\|_a \quad (2.2.2)$$

als die von der Vektornorm induzierte Matrixnorm. In diesem Sinne gilt für $A \in \mathbb{C}^{n \times n}$:

$$\begin{aligned} \|A\|_1 &= \max_{k=1,\dots,n} \sum_{i=1}^n |a_{ik}| \quad (\text{Spaltensummennorm}), \\ \|A\|_\infty &= \max_{i=1,\dots,n} \sum_{k=1}^n |a_{ik}| \quad (\text{Zeilensummennorm}), \\ \|A\|_2 &\leq \left(\sum_{i,k=1}^n |a_{ik}|^2 \right)^{\frac{1}{2}} = \|A\|_F. \end{aligned}$$

Hierbei wird $\|\cdot\|_F : \mathbb{C}^{n \times n} \rightarrow \mathbb{R}$ als Frobeniusnorm bezeichnet. Wegen $\|I\|_F = \sqrt{n}$ ist aus (2.2.2) sofort ersichtlich, daß die Frobeniusnorm keine induzierte Matrixnorm darstellt. Auch die durch $\|A\| := \max_{i,k=1,\dots,n} |a_{ik}|$ gegebene Norm besitzt keine zugehörige Vektornorm. Einfache Beispiele zeigen, daß diese Norm nicht submultiplikativ ist. Es sei an dieser Stelle nochmals betont, daß wir uns bei den folgenden Betrachtungen, solange nicht ausdrücklich erwähnt, stets auf induzierte Matrixnormen beschränken werden, obwohl einige Aussagen (zum Beispiel der Satz 2.28) allgemeine Gültigkeit besitzen. Die einzige Ausnahme bildet die im Abschnitt 5.7 betrachtete Frobenius-Norm bei der Herleitung der unvollständigen Frobenius-Inversen.

Satz 2.28 *Eine Matrix $A \in \mathbb{C}^{n \times n}$ ist in jeder Norm beschränkt.*

Beweis:

Sei $x \in X$ beliebig, dann folgt

$$\begin{aligned} \|Ax\|_\infty &= \max_{i=1,\dots,n} |(Ax)_i| = \max_{i=1,\dots,n} \left| \sum_{j=1}^n a_{ij}x_j \right| \\ &\leq \max_{i=1,\dots,n} \sum_{j=1}^n |a_{ij}| |x_j| \leq S \|x\|_\infty \end{aligned}$$

mit $S = \max_{i=1,\dots,n} \sum_{j=1}^n |a_{ij}|$. Diese Aussage gilt durch die Äquivalenz der Normen (Satz 2.11) für jede beliebige Norm. □

Folglich erhalten wir mit Satz 2.28 auch die Stetigkeit jeder Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$.

Definition 2.29 Eine komplexe Zahl $\lambda \in \mathbb{C}$ heißt Eigenwert der Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$, falls ein Vektor $\mathbf{x} \in \mathbb{C}^n \setminus \{\mathbf{0}\}$ mit

$$\mathbf{A}\mathbf{x} = \lambda\mathbf{x}$$

existiert. Der Vektor $\mathbf{x}$ heißt Eigenvektor zum Eigenwert λ .

Die Menge

$$\sigma(\mathbf{A}) = \{\lambda \mid \lambda \text{ ist Eigenwert von } \mathbf{A}\}$$

wird als Spektrum von $\mathbf{A}$ bezeichnet.

Die Zahl

$$\rho(\mathbf{A}) = \max \{|\lambda| \mid \lambda \in \sigma(\mathbf{A})\}$$

heißt Spektralradius von $\mathbf{A}$.

Die Eigenwerte $\lambda \in \sigma(\mathbf{A})$ stellen wegen

$$(\mathbf{A} - \lambda\mathbf{I})\mathbf{x} = \mathbf{0} \text{ mit } \mathbf{x} \in \mathbb{C}^n \setminus \{\mathbf{0}\}$$

die Nullstellen des charakteristischen Polynoms

$$p(\lambda) = \det(\mathbf{A} - \lambda\mathbf{I})$$

dar. Da jedes Polynom über $\mathbb{C}$ nach dem Fundamentalsatz der Algebra in Linearfaktoren zerfällt, das heißt $p(\lambda) = \prod_{i=1}^n (\lambda - \lambda_i)$ gilt, erhalten wir direkt $\sigma(\mathbf{A}) \neq \emptyset$ für alle $\mathbf{A} \in \mathbb{C}^{n \times n}$.

Der folgende Satz von Schur stellt das wesentliche Hilfsmittel zum Nachweis des Satzes 2.35 dar, wodurch in Kombination mit Satz 2.34 ein direkter Zusammenhang zwischen einer Norm und dem Spektralradius einer Matrix vorliegt. Diese Aussagen werden sich später bei der Konvergenzanalyse von Splitting-Methoden als entscheidend erweisen.

Satz 2.30 Zu jeder Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$ ($\mathbb{R}^{n \times n}$ mit $\sigma(\mathbf{A}) \subset \mathbb{R}$) existiert eine unitäre (orthogonale) Matrix $\mathbf{U} \in \mathbb{C}^{n \times n}$ ($\mathbb{R}^{n \times n}$) derart, daß

$$\mathbf{U}^* \mathbf{A} \mathbf{U}$$

eine rechte obere Dreiecksmatrix darstellt.

Beweis:

Wir beschränken uns beim Beweis der Behauptung zunächst auf den komplexen Fall.

Der Beweis wird mittels einer vollständigen Induktion geführt.

Für $n = 1$ erfüllt $\mathbf{U} = \mathbf{I}$ die Behauptung.

Sei die Behauptung für $j = 1, \dots, n$ erfüllt, dann wähle ein $\lambda \in \sigma(\mathbf{A})$ mit zugehörigem Eigenvektor $\tilde{\mathbf{v}}_1 \in \mathbb{C}^{n+1} \setminus \{\mathbf{0}\}$. Durch Erweiterung von $\mathbf{v}_1 = \tilde{\mathbf{v}}_1 / \|\tilde{\mathbf{v}}_1\|_2$ durch $\mathbf{v}_2, \dots, \mathbf{v}_{n+1}$ zu einer Orthonormalbasis des $\mathbb{C}^{n+1}$ ergibt sich mit

$$\mathbb{C}^{(n+1) \times (n+1)} \ni \mathbf{V} = (\mathbf{v}_1 \dots \mathbf{v}_{n+1})$$

die Gleichung

$$\mathbf{V}^* \mathbf{A} \mathbf{V} \mathbf{e}_1 = \mathbf{V}^* \mathbf{A} \mathbf{v}_1 = \mathbf{V}^* \lambda \mathbf{v}_1 = \lambda \mathbf{e}_1,$$

wobei $\mathbf{e}_1 = (1, 0, \dots, 0)^T \in \mathbb{C}^{n+1}$ gilt. Hiermit folgt

$$\mathbf{V}^* \mathbf{A} \mathbf{V} = \begin{pmatrix} \lambda & \tilde{\mathbf{a}}^T \\ 0 & \tilde{\mathbf{A}} \\ \vdots & \\ 0 & \end{pmatrix} \quad \text{mit } \tilde{\mathbf{A}} \in \mathbb{C}^{n \times n} \text{ und } \tilde{\mathbf{a}} \in \mathbb{C}^n.$$

Zu $\tilde{\mathbf{A}}$ existiert laut Induktionsvoraussetzung eine unitäre Matrix $\tilde{\mathbf{W}} \in \mathbb{C}^{n \times n}$ derart, daß $\tilde{\mathbf{W}}^* \tilde{\mathbf{A}} \tilde{\mathbf{W}}$ eine rechte obere Dreiecksmatrix darstellt. Mit $\tilde{\mathbf{W}}$ ist auch

$$\mathbf{W} = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & & & \\ \vdots & & \tilde{\mathbf{W}} & \\ 0 & & & \end{pmatrix}$$

unitär und wir erhalten mit $\mathbf{U} := \mathbf{V} \mathbf{W} \in \mathbb{C}^{(n+1) \times (n+1)}$ laut Lemma 2.27 eine unitäre Matrix, für die einfaches Nachrechnen zeigt, daß $\mathbf{U}^* \mathbf{A} \mathbf{U}$ eine rechte obere Dreiecksmatrix darstellt.

Im Fall einer regulären Matrix $\mathbf{A} \in \mathbb{R}^{(n+1) \times (n+1)}$ erhalten wir wegen $\sigma(\mathbf{A}) \subset \mathbb{R}$ zu jedem Eigenvektor $\tilde{\mathbf{v}}_1 \in \mathbb{C}^{n+1} \setminus \{\mathbf{0}\}$ zum Eigenwert $\lambda \in \mathbb{R}$ wegen $\tilde{\mathbf{v}}_1 = \tilde{\mathbf{x}}_1 + i\tilde{\mathbf{y}}_1$, $\tilde{\mathbf{x}}_1, \tilde{\mathbf{y}}_1 \in \mathbb{R}^{n+1}$ aus

$$\mathbf{A} \tilde{\mathbf{x}}_1 + i \mathbf{A} \tilde{\mathbf{y}}_1 = \mathbf{A} \tilde{\mathbf{v}}_1 = \lambda \tilde{\mathbf{v}}_1 = \lambda \tilde{\mathbf{x}}_1 + i \lambda \tilde{\mathbf{y}}_1$$

die Eigenschaft

$$\mathbf{A} \tilde{\mathbf{x}}_1 = \lambda \tilde{\mathbf{x}}_1 \text{ sowie } \mathbf{A} \tilde{\mathbf{y}}_1 = \lambda \tilde{\mathbf{y}}_1,$$

so daß mit $\tilde{\mathbf{x}}_1$ oder $\tilde{\mathbf{y}}_1$ ein Eigenvektor aus $\mathbb{R}^{n+1} \setminus \{\mathbf{0}\}$ vorliegt. Unter Berücksichtigung dieser Eigenschaft ergibt sich der Nachweis im Fall einer reellen Matrix analog zum obigen Vorgehen. $\square$

Lemma 2.31 Für jede hermitesche Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$ gilt $\sigma(\mathbf{A}) \subset \mathbb{R}$.

Beweis:

Sei $\mathbf{x} \in \mathbb{C}^n \setminus \{\mathbf{0}\}$ Eigenvektor zum Eigenwert $\lambda \in \sigma(\mathbf{A})$. Dann folgt mit

$$\lambda \underbrace{\|\mathbf{x}\|_2}_{\in \mathbb{R}^+} = \lambda(\mathbf{x}, \mathbf{x})_2 = (\lambda \mathbf{x}, \mathbf{x})_2 = (\mathbf{A} \mathbf{x}, \mathbf{x})_2 = (\mathbf{x}, \mathbf{A} \mathbf{x})_2 = (\mathbf{x}, \lambda \mathbf{x})_2 = \overline{\lambda}(\mathbf{x}, \mathbf{x})_2 = \overline{\lambda} \underbrace{\|\mathbf{x}\|_2}_{\in \mathbb{R}^+}$$

direkt $\lambda = \overline{\lambda}$ und somit $\lambda \in \mathbb{R}$. $\square$

Als direkte Folgerung des letzten Satzes erhalten wir die anschließende Aussage für hermitesche und wegen Lemma 2.31 somit auch für symmetrische Matrizen.

Korollar 2.32 Sei $\mathbf{A} \in \mathbb{C}^{n \times n}$ ($\mathbb{R}^{n \times n}$) hermitesch (symmetrisch), dann existiert eine unitäre (orthogonale) Matrix $\mathbf{U} \in \mathbb{C}^{n \times n}$ ($\mathbb{R}^{n \times n}$), derart, daß

$$\mathbf{U}^* \mathbf{A} \mathbf{U} = \text{diag}\{\lambda_1, \dots, \lambda_n\} \in \mathbb{R}^{n \times n}$$

gilt. Hierbei stellt für $i = 1, \dots, n$ jeweils $\lambda_i \in \mathbb{R}$ den Eigenwert der Matrix $\mathbf{A}$ mit der i -ten Spalte von $\mathbf{U}$ als zugehörigen Eigenvektor dar.

Satz 2.33 Sei $\mathbf{A} \in \mathbb{C}^{n \times n}$, dann gilt

$$\|\mathbf{A}\|_2 = \sqrt{\rho(\mathbf{A}^* \mathbf{A})}. \quad (2.2.3)$$

Beweis:

Da $\mathbf{A}^* \mathbf{A}$ hermitesch ist, existiert laut Korollar 2.32 eine unitäre Matrix $\mathbf{U} \in \mathbb{C}^{n \times n}$, so daß $\mathbf{U}^* \mathbf{A}^* \mathbf{A} \mathbf{U} = \text{diag}\{\lambda_1, \dots, \lambda_n\} \in \mathbb{R}^{n \times n}$ gilt. Jedes $\mathbf{x} \in \mathbb{C}^n$ läßt sich daher unter Verwendung der Spalten $\mathbf{u}_1, \dots, \mathbf{u}_n$ der unitären Matrix $\mathbf{U}$ in der Form $\mathbf{x} = \sum_{i=1}^n \alpha_i \mathbf{u}_i$ mit $\alpha_i \in \mathbb{C}$ darstellen und es gilt

$$\mathbf{A}^* \mathbf{A} \mathbf{x} = \sum_{i=1}^n \lambda_i \alpha_i \mathbf{u}_i.$$

Hiermit erhalten wir

$$\begin{aligned} \|\mathbf{A} \mathbf{x}\|_2^2 &= (\mathbf{A} \mathbf{x}, \mathbf{A} \mathbf{x})_2 = (\mathbf{x}, \mathbf{A}^* \mathbf{A} \mathbf{x})_2 \\ &= \left(\sum_{i=1}^n \alpha_i \mathbf{u}_i, \sum_{i=1}^n \lambda_i \alpha_i \mathbf{u}_i \right)_2 = \sum_{i=1}^n (\alpha_i \mathbf{u}_i, \lambda_i \alpha_i \mathbf{u}_i)_2 \\ &= \sum_{i=1}^n \lambda_i |\alpha_i|^2 \\ &\leq \rho(\mathbf{A}^* \mathbf{A}) \sum_{i=1}^n |\alpha_i|^2 = \rho(\mathbf{A}^* \mathbf{A}) \|\mathbf{x}\|_2^2. \end{aligned} \quad (2.2.4)$$

Hiermit erhalten wir

$$\frac{\|\mathbf{A} \mathbf{x}\|_2^2}{\|\mathbf{x}\|_2^2} \leq \rho(\mathbf{A}^* \mathbf{A}).$$

Die Gleichheit ergibt sich durch Betrachtung des Eigenvektors $\mathbf{u}_j$ zum betragsgrößten Eigenwert λ_j . Mit (2.2.4) folgt

$$0 \leq \|\mathbf{A} \mathbf{u}_i\|_2^2 = \lambda_i, \quad i = 1, \dots, n$$

und somit erhalten wir

$$\frac{\|\mathbf{A} \mathbf{u}_j\|_2^2}{\|\mathbf{u}_j\|_2^2} = \frac{\lambda_j \|\mathbf{u}_j\|_2^2}{\|\mathbf{u}_j\|_2^2} = \lambda_j = \rho(\mathbf{A}^* \mathbf{A}).$$

□

Aufgrund der Eigenschaft (2.2.3) wird $\|\mathbf{A}\|_2$ auch als Spektralnorm der Matrix $\mathbf{A}$ bezeichnet. Die folgenden zwei Sätze sind von entscheidender Bedeutung für die Konvergenzaussagen iterativer Verfahren, die auf einer Aufteilung der Matrix des Gleichungssystems beruhen.

Satz 2.34 Für jede Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$ gilt

- (a) $\rho(\mathbf{A}) \leq \|\mathbf{A}\|$ für jede induzierte Matrixnorm,
- (b) $\|\mathbf{A}\|_2 = \rho(\mathbf{A})$, falls $\mathbf{A}$ hermitesch ist.

Beweis:

Sei $\lambda \in \mathbb{C}$ der betragsgrößte Eigenwert von $\mathbf{A}$ zum Eigenvektor $\mathbf{u} \in \mathbb{C}^n \setminus \{\mathbf{0}\}$ für den o.B.d.A. $\|\mathbf{u}\| = 1$ gilt. Dann folgt

$$\|\mathbf{A}\| = \sup_{\|\mathbf{x}\|=1} \|\mathbf{Ax}\| \geq \|\mathbf{Au}\| = |\lambda| \|\mathbf{u}\| = |\lambda|,$$

womit sich direkt $\rho(\mathbf{A}) \leq \|\mathbf{A}\|$ ergibt.

Mit Korollar 2.32 erhalten wir für hermitesche Matrizen die Eigenschaft $\rho(\mathbf{A}^2) = \rho(\mathbf{A})^2$, so daß sich für hermitesche Matrizen unter Verwendung des Satzes 2.33

$$\|\mathbf{A}\|_2 = \sqrt{\rho(\mathbf{A}^* \mathbf{A})} = \sqrt{\rho(\mathbf{A}^2)} = \sqrt{\rho(\mathbf{A})^2} = \rho(\mathbf{A})$$

ergibt. □

Satz 2.35 *Zu jeder Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$ und zu jedem $\varepsilon > 0$ existiert eine Norm auf $\mathbb{C}^{n \times n}$, so daß*

$$\|\mathbf{A}\| \leq \rho(\mathbf{A}) + \varepsilon$$

gilt.

Beweis:

Für $n = 1$ ist die Aussage ebenso trivial wie für jede Nullmatrix $\mathbf{A} \in \mathbb{C}^{n \times n}$. Seien daher $n \geq 2$ und $\mathbf{A}$ ungleich der Nullmatrix. Der Satz von Schur liefert die Existenz einer unitären Matrix $\mathbf{U} \in \mathbb{C}^{n \times n}$ mit

$$\mathbf{R} = \mathbf{U}^* \mathbf{A} \mathbf{U} = \begin{pmatrix} r_{11} & \cdots & r_{1n} \\ & \ddots & \vdots \\ & & r_{nn} \end{pmatrix},$$

wobei $\lambda_i = r_{ii}$, $i = 1, \dots, n$ die Eigenwerte von $\mathbf{A}$ sind. Sei $\alpha := \max_{i,k=1,\dots,n} |r_{ik}| > 0$, dann definieren wir zu gegebenem $\varepsilon > 0$

$$\delta := \min \left(1, \frac{\varepsilon}{(n-1)\alpha} \right) > 0. \quad (2.2.5)$$

Mit

$$\mathbf{D} = \begin{pmatrix} 1 & & & \\ & \delta & & \\ & & \ddots & \\ & & & \delta^{n-1} \end{pmatrix}$$

erhalten wir

$$\mathbf{C} := \mathbf{D}^{-1} \mathbf{R} \mathbf{D} = \begin{pmatrix} r_{11} & \delta r_{12} & \cdots & \cdots & \delta^{n-1} r_{1n} \\ & \ddots & \ddots & & \vdots \\ & & \ddots & \ddots & \vdots \\ & & & \ddots & \delta r_{n-1,n} \\ & & & & r_{n,n} \end{pmatrix}.$$

Unter Verwendung der Definition (2.2.5) folgt

$$\begin{aligned}
 \|C\|_\infty &\leq \max_{i=1,\dots,n} |r_{ii}| + (n-1)\delta\alpha \\
 &= \rho(A) + (n-1)\delta\alpha \\
 &\leq \rho(A) + \varepsilon.
 \end{aligned} \tag{2.2.6}$$

Da D und U reguläre Matrizen darstellen, ist durch

$$\|x\| := \|D^{-1}U^{-1}x\|_\infty$$

eine Norm auf $\mathbb{C}^n$ gegeben. Sei $y := D^{-1}U^{-1}x$, dann folgt mit (2.2.6)

$$\begin{aligned}
 \|A\| &= \sup_{\|x\|=1} \|Ax\| = \sup_{\|D^{-1}U^{-1}x\|_\infty=1} \|D^{-1}U^{-1}Ax\|_\infty \\
 &= \sup_{\|y\|_\infty=1} \|D^{-1}U^{-1}AUDy\|_\infty \\
 &= \sup_{\|y\|_\infty=1} \|Cy\|_\infty \\
 &= \|C\|_\infty \leq \rho(A) + \varepsilon.
 \end{aligned}$$

□

Mit den obigen beiden Sätzen ergibt sich somit stets für alle Matrizen A und alle $\varepsilon > 0$ die Existenz einer Norm derart, daß

$$\rho(A) \leq \|A\| \leq \rho(A) + \varepsilon$$

gilt. Hierdurch können Konvergenzaussagen, die auf der Norm einer Matrix basieren und dabei unabhängig von der speziellen Wahl der Norm sind, direkt auf den Spektralradius der Matrix übertragen werden.

Betrachten wir nun normale Matrizen, so werden diese teilweise auch über ihre Diagonalisierbarkeit mittels einer unitären Matrix definiert [23]. Die Möglichkeit einer solchen Definition normaler Matrizen wird durch den folgenden Satz nachgewiesen, der uns im weiteren auch eine spezielle Darstellung der Konditionszahl solcher Matrizen ermöglicht.

Satz 2.36 *Eine Matrix $A \in \mathbb{C}^{n \times n}$ ist genau dann normal, wenn eine unitäre Matrix $U \in \mathbb{C}^{n \times n}$ existiert, so daß $D = U^*AU$ eine Diagonalmatrix darstellt.*

Beweis:

„ $\Rightarrow$ “

Mit Satz 2.30 existiert eine unitäre Matrix U derart, daß $R = (r_{ij})_{i,j=1,\dots,n} = U^*AU$ eine rechte obere Dreiecksmatrix darstellt. Mit

$$R^*R = U^*A^*UU^*AU = U^*A^*AU = U^*AA^*U = U^*AUU^*A^*U = RR^*$$

ist R ebenfalls normal. Für die Diagonalelemente der Matrix $B = (b_{ij})_{i,j=1,\dots,n} = R^*R = RR^*$ gilt

$$\sum_{j=1}^i |r_{ji}|^2 = b_{ii} = \sum_{j=i}^n |r_{ij}|^2.$$

Eine sukzessive Auswertung dieser Gleichung für $i = 1, \dots, n$ liefert $r_{ij} = 0$ für alle $i \neq j$. Somit stellt $\mathbf{D} = \mathbf{R}$ die gesuchte Diagonalmatrix dar.

„ $\Leftarrow$ “

Sei $\mathbf{U} \in \mathbb{C}^{n \times n}$ unitär mit $\mathbf{D} = \text{diag}\{d_{11}, \dots, d_{nn}\} = \mathbf{U}^* \mathbf{A} \mathbf{U}$, dann folgt

$$\mathbf{A}^* \mathbf{A} = \mathbf{U} \mathbf{D}^* \mathbf{D} \mathbf{U}^* = \mathbf{U} \mathbf{D} \mathbf{D}^* \mathbf{U}^* = \mathbf{A} \mathbf{A}^*.$$

□

2.3 Konditionszahl und singuläre Werte

Neben dem Spektralradius stellt auch die Konditionszahl einer Matrix eine interessante Größe dar, mit der Fehlereinflüsse wie auch Konvergenzgeschwindigkeiten iterativer Methoden abgeschätzt werden können. Wir werden daher in diesem Abschnitt den Begriff der Konditionszahl einführen und seine Bedeutung hinsichtlich der Lösung linearer Gleichungssysteme studieren. Es wird sich dabei zeigen, daß sich die bezüglich der Spektralnorm gebildete Konditionszahl einer normalen Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ durch deren Eigenwerte darstellen läßt. Im allgemeinen Fall einer regulären Matrix muß hierzu ihre Singulärwertverteilung betrachtet werden.

Definition 2.37 Sei $\mathbf{A} \in \mathbb{C}^{n \times n}$ regulär, dann heißt

$$\text{cond}_a(\mathbf{A}) := \|\mathbf{A}\|_a \|\mathbf{A}^{-1}\|_a$$

die Konditionszahl der Matrix $\mathbf{A}$ bezüglich der induzierten Matrixnorm $\|\cdot\|_a$.

Es ist leicht ersichtlich, daß die Konditionszahl einer regulären Matrix unabhängig von der zugrundeliegenden induzierten Matrixnorm nach unten beschränkt ist. Den mathematischen Nachweis hierzu liefert das folgende Lemma.

Lemma 2.38 Sei $\mathbf{A} \in \mathbb{C}^{n \times n}$ regulär, dann gilt

$$\text{cond}(\mathbf{A}) \geq \text{cond}(\mathbf{I}) = 1$$

für $\text{cond}(\mathbf{A}) = \|\mathbf{A}\| \|\mathbf{A}^{-1}\|$ mit einer induzierten Matrixnorm $\|\cdot\|$.

Beweis:

Der Nachweis ergibt sich direkt aus der folgenden Ungleichung

$$\text{cond}(\mathbf{I}) = \|\mathbf{I}\| \|\mathbf{I}^{-1}\| = 1 = \|\mathbf{I}\| = \|\mathbf{A} \mathbf{A}^{-1}\| \stackrel{\text{Satz 2.17}}{\leq} \|\mathbf{A}\| \|\mathbf{A}^{-1}\| = \text{cond}(\mathbf{A}).$$

□

Die Relevanz der Konditionszahl hinsichtlich der iterativen Lösung linearer Gleichungssysteme werden wir nun anhand zweier Sätze verdeutlichen. In praktischen Anwendungen ist es üblich, die Norm des Residuenvektors $\mathbf{r}_m = \mathbf{b} - \mathbf{A}\mathbf{x}_m$ als Maß für die Güte der vorliegenden Näherungslösung innerhalb eines Iterationsverfahrens zu verwenden, da der Fehlervektor $\mathbf{e}_m = \mathbf{A}^{-1}\mathbf{b} - \mathbf{x}_m$ aufgrund der Unkenntnis über die wahre Lösung nicht zur Verfügung steht. Ein Grund hierfür liegt in der Eigenschaft, daß das Residuum analog zum Fehler genau dann identisch verschwindet, wenn die Iterierte mit der exakten Lösung übereinstimmt. Es wird sich zeigen, daß für eine kleine Konditionszahl der Matrix des linearen Gleichungssystems eine Konvergenzabschätzung auf der Basis des Residuums sinnvoll ist. Liegt jedoch eine große Konditionszahl vor, so kann trotz Verringerung des Residuums eine deutlich anwachsende Fehlernorm vorliegen.

Satz 2.39 *Gegeben sei ein Iterationsverfahren zur Lösung von $\mathbf{A}\mathbf{x} = \mathbf{b}$ mit einer regulären Matrix $\mathbf{A}$. Es bezeichne $\mathbf{e}_k = \mathbf{A}^{-1}\mathbf{b} - \mathbf{x}_k$ den Fehlervektor und $\mathbf{r}_k = \mathbf{b} - \mathbf{A}\mathbf{x}_k$ den Residuenvektor des k -ten Iterationschritts, dann gilt*

$$\frac{1}{\text{cond}(\mathbf{A})} \frac{\|\mathbf{r}_k\|}{\|\mathbf{r}_0\|} \leq \frac{\|\mathbf{e}_k\|}{\|\mathbf{e}_0\|} \leq \text{cond}(\mathbf{A}) \frac{\|\mathbf{r}_k\|}{\|\mathbf{r}_0\|} \leq \text{cond}(\mathbf{A})^2 \frac{\|\mathbf{e}_k\|}{\|\mathbf{e}_0\|}. \quad (2.3.1)$$

Beweis:

Mit

$$\|\mathbf{r}_k\| = \|\mathbf{b} - \mathbf{A}\mathbf{x}_k\| = \|\mathbf{A}\mathbf{e}_k\| \leq \|\mathbf{A}\| \|\mathbf{e}_k\|$$

und

$$\|\mathbf{e}_k\| = \|\mathbf{A}^{-1}\mathbf{b} - \mathbf{x}_k\| \leq \|\mathbf{A}^{-1}\| \|\mathbf{b} - \mathbf{A}\mathbf{x}_k\| = \|\mathbf{A}^{-1}\| \|\mathbf{r}_k\|$$

ergibt sich der erste Teil der Behauptung aus

$$\frac{1}{\text{cond}(\mathbf{A})} \frac{\|\mathbf{r}_k\|}{\|\mathbf{r}_0\|} = \frac{1}{\|\mathbf{A}\| \|\mathbf{A}^{-1}\|} \frac{\|\mathbf{r}_k\|}{\|\mathbf{r}_0\|} \leq \frac{1}{\|\mathbf{A}\|} \frac{\|\mathbf{r}_k\|}{\|\mathbf{A}^{-1}\mathbf{r}_0\|} = \frac{1}{\|\mathbf{A}\|} \frac{\|\mathbf{A}\mathbf{e}_k\|}{\|\mathbf{e}_0\|} \leq \frac{\|\mathbf{e}_k\|}{\|\mathbf{e}_0\|}.$$

Die weiteren Ungleichungen folgen analog. $\square$

Dem obigen Satz können wir entnehmen, daß im Fall einer orthogonalen Matrix $\mathbf{A}$ stets

$$\frac{\|\mathbf{e}_k\|_2}{\|\mathbf{e}_0\|_2} = \frac{\|\mathbf{r}_k\|_2}{\|\mathbf{r}_0\|_2}$$

gilt, wodurch eine direkte Konvergenzanalyse auf der Grundlage des Residuums durchgeführt werden kann.

Meßdaten aus physikalischen Experimenten weisen in natürlicher Weise Ungenauigkeiten auf. Die Auswirkung solcher oder anderer fehlerhafter Eingangsdaten auf die Lösung können ebenfalls unter Verwendung der Konditionszahl der Matrix abgeschätzt werden. Auch hierbei erweist sich eine kleine Konditionszahl als vorteilhaft, da sich hier kleine relative Fehler in den Daten $\mathbf{b}$ auch nur als kleine relative Fehler in den Lösungen $\mathbf{x}$ bemerkbar machen.

Satz 2.40 *Seien $\mathbf{A}$ regulär, $\mathbf{x}$ die Lösung des Gleichungssystems $\mathbf{A}\mathbf{x} = \mathbf{b}$ und $\mathbf{x} + \Delta\mathbf{x}$ die Lösung von $\mathbf{A}(\mathbf{x} + \Delta\mathbf{x}) = \mathbf{b} + \Delta\mathbf{b}$, dann gilt*

$$\frac{\|\Delta\mathbf{x}\|}{\|\mathbf{x}\|} \leq \text{cond}(\mathbf{A}) \frac{\|\Delta\mathbf{b}\|}{\|\mathbf{b}\|}.$$

Beweis:

Aufgrund der Linearität von $\mathbf{A}$ folgt

$$\Delta \mathbf{b} = (\mathbf{b} + \Delta \mathbf{b}) - \mathbf{b} = \mathbf{A}(\mathbf{x} + \Delta \mathbf{x}) - \mathbf{A}\mathbf{x} = \mathbf{A}(\Delta \mathbf{x}).$$

Hiermit erhalten wir

$$\|\Delta \mathbf{x}\| = \|\mathbf{A}^{-1}(\Delta \mathbf{b})\| \leq \|\mathbf{A}^{-1}\| \|\Delta \mathbf{b}\|,$$

so daß mit

$$\|\mathbf{b}\| = \|\mathbf{A}\mathbf{x}\| \leq \|\mathbf{A}\| \|\mathbf{x}\|$$

die behauptete Ungleichung mit

$$\frac{\|\Delta \mathbf{x}\|}{\|\mathbf{x}\|} \leq \frac{\|\mathbf{A}^{-1}\| \|\Delta \mathbf{b}\|}{\|\mathbf{A}\|^{-1} \|\mathbf{b}\|} = \text{cond}(\mathbf{A}) \frac{\|\Delta \mathbf{b}\|}{\|\mathbf{b}\|}$$

folgt. □

Liegt ein Gleichungssystem vor, bei dem die Matrix eine sehr große Konditionszahl besitzt, dann erweist es sich aufgrund der Sätze 2.39 und 2.40 als sinnvoll, zunächst eine Umformulierung auf ein äquivalentes System $\tilde{\mathbf{A}}\mathbf{x} = \tilde{\mathbf{b}}$ mit $\text{cond}(\tilde{\mathbf{A}}) \ll \text{cond}(\mathbf{A})$ vorzunehmen und anschließend ein Iterationsverfahren auf das transformierte System anzuwenden. Eine äquivalente Umformulierung kann zum Beispiel durch Multiplikation mit einer regulären Matrix $\mathbf{P}$ gemäß Lemma 2.24 durchgeführt werden. Hierbei stellt sich die Frage nach der Wahl der Matrix $\mathbf{P}$. Die Matrix sollte aus Rechenzeitgründen einfach berechenbar sein und aus Effektivitätsgründen eine möglichst gute Approximation der Inversen der Matrix $\mathbf{A}$ repräsentieren, so daß $\text{cond}(\mathbf{P}\mathbf{A}) \ll \text{cond}(\mathbf{A})$ gilt. Solche äquivalenten Transformationen werden als Präkonditionierungen bezeichnet und ausführlich im Kapitel 5 untersucht. Es sei an dieser Stelle bereits erwähnt, daß die Stabilisierung numerischer Verfahren zwar einen Grund zur Nutzung einer Präkonditionierung darstellt, die wesentliche Zielsetzung derartiger Techniken allerdings in der Beschleunigung iterativer Verfahren liegt.

Der folgende Abschnitt liefert eine Darstellung der durch die Spektralnorm gegebenen Konditionszahl einer regulären Matrix $\mathbf{A}$ mittels ihrer singulären Werte beziehungsweise ihrer Eigenwerte. Die Konditionszahlen hinsichtlich der Zeilen- respektive Spaltensummennorm kann anschließend über $\text{cond}_2(\mathbf{A})$ abgeschätzt werden.

Satz und Definition 2.41 *Für jede reguläre Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ existieren orthogonale Matrizen $\mathbf{U}, \mathbf{V} \in \mathbb{R}^{n \times n}$ derart, daß*

$$\mathbf{U}^T \mathbf{A} \mathbf{V} = \text{diag}\{\sigma_1, \dots, \sigma_n\} \quad (2.3.2)$$

mit $0 < \sigma_1 \leq \dots \leq \sigma_n$ gilt. Die aufgeführten reellen Zahlen σ_i ($i = 1, \dots, n$) heißen singuläre Werte der Matrix $\mathbf{A}$ und genügen jeweils der Gleichung

$$\det(\mathbf{A}^T \mathbf{A} - \sigma_i^2 \mathbf{I}) = 0. \quad (2.3.3)$$

Die i -te Spalte von $\mathbf{U}$ bzw. $\mathbf{V}$ heißt i -ter Links- bzw. Rechtssingulärvektor der Matrix $\mathbf{A}$.

Beweis:

Da $\mathbf{A}$ regulär ist, liegt mit $\mathbf{A}^T \mathbf{A}$ eine symmetrische, positiv definite Matrix vor. Somit existiert laut Korollar 2.32 eine orthogonale Matrix $\mathbf{V} \in \mathbb{R}^{n \times n}$ derart, daß

$$\mathbf{V}^T \mathbf{A}^T \mathbf{A} \mathbf{V} = \text{diag}\{\lambda_1, \dots, \lambda_n\}$$

mit $0 < \lambda_1 \leq \dots \leq \lambda_n$ gilt. Definieren wir $\sigma_i = \sqrt{\lambda_i}$, so erhalten wir $0 < \sigma_1 \leq \dots \leq \sigma_n$ mit $\det(\mathbf{A}^T \mathbf{A} - \sigma_i^2 \mathbf{I}) = 0$. Zum Nachweis der verbleibenden Aussage (2.3.2) definieren wir $\mathbf{U} = \mathbf{A} \mathbf{V} \mathbf{D}^{-1}$ mit $\mathbf{D} = \text{diag}\{\sigma_1, \dots, \sigma_n\}$. Wegen

$$\mathbf{U}^T \mathbf{U} = \mathbf{D}^{-T} \mathbf{V}^T \mathbf{A}^T \mathbf{A} \mathbf{V} \mathbf{D}^{-1} = \mathbf{D}^{-1} \text{diag}\{\lambda_1, \dots, \lambda_n\} \mathbf{D}^{-1} = \mathbf{I}$$

repräsentiert $\mathbf{U}$ eine orthogonale Matrix, und wir erhalten die behauptete Darstellung durch

$$\mathbf{U}^T \mathbf{A} \mathbf{V} = \mathbf{D}^{-T} \mathbf{V}^T \mathbf{A}^T \mathbf{A} \mathbf{V} = \mathbf{D}^{-T} \text{diag}\{\lambda_1, \dots, \lambda_n\} = \mathbf{D}.$$

□

Satz 2.42 Sei $\mathbf{A} \in \mathbb{R}^{n \times n}$, dann gelten mit den Bezeichnungen aus Satz 2.41 die Aussagen

$$\sup_{\mathbf{x} \neq 0} \frac{\|\mathbf{A}\mathbf{x}\|_2}{\|\mathbf{x}\|_2} = \sigma_n$$

und

$$\inf_{\mathbf{x} \neq 0} \frac{\|\mathbf{A}\mathbf{x}\|_2}{\|\mathbf{x}\|_2} = \sigma_1.$$

Beweis:

Aus der Definition der euklidischen Norm und der Regularität von $\mathbf{V}$ folgt

$$\sup_{\mathbf{x} \neq 0} \frac{\|\mathbf{A}\mathbf{x}\|_2^2}{\|\mathbf{x}\|_2^2} = \sup_{\mathbf{V}\mathbf{x} \neq 0} \frac{\mathbf{x}^T \mathbf{V}^T \mathbf{A}^T \mathbf{A} \mathbf{V} \mathbf{x}}{\mathbf{x}^T \mathbf{V}^T \mathbf{V} \mathbf{x}} = \sup_{\mathbf{V}\mathbf{x} \neq 0} \frac{\mathbf{x}^T \mathbf{V}^T \mathbf{A}^T \mathbf{U} \mathbf{U}^T \mathbf{A} \mathbf{V} \mathbf{x}}{\mathbf{x}^T \mathbf{V}^T \mathbf{V} \mathbf{x}}$$

wegen $\mathbf{U} \mathbf{U}^T = \mathbf{I}$. Hiermit ergibt sich wegen $\mathbf{V}^T \mathbf{V} = \mathbf{I}$ die erste Aussage. Analog schließt man auf die Gültigkeit der zweiten Behauptung. □

Nach diesen Überlegungen können wir die Spektralnorm einer regulären Matrix $\mathbf{A}$ und ihrer Inversen durch die singulären Werte der Matrix gemäß

$$\|\mathbf{A}\|_2 = \sigma_n \quad \text{und} \quad \|\mathbf{A}^{-1}\|_2 = \frac{1}{\sigma_1}$$

ausdrücken. Da normale Matrizen laut Satz 2.36 unitär diagonalisierbar sind, erhalten wir als unmittelbare Folgerung der obigen Gleichung den folgenden Zusammenhang zwischen der Konditionszahl und den singulären Werten, respektive Eigenwerten, einer regulären Matrix in der folgenden Form:

Korollar 2.43 Sei $\mathbf{A} \in \mathbb{R}^{n \times n}$ regulär, dann gilt

$$\text{cond}_2(\mathbf{A}) = \frac{\sigma_n}{\sigma_1}, \tag{2.3.4}$$

wenn σ_n der größte und σ_1 der kleinste singuläre Wert der Matrix ist.
Ist $\mathbf{A}$ zudem normal, so gilt

$$\text{cond}_2(\mathbf{A}) = \frac{|\lambda_n|}{|\lambda_1|}, \quad (2.3.5)$$

falls λ_n den betragsgrößten und λ_1 den betragskleinsten Eigenwert der Matrix bezeichnet.

Obwohl eine zum Korollar 2.43 äquivalente Aussage bezogen auf eine andere Matrixnorm im allgemeinen nicht existiert, können Fehlerabschätzungen und Konvergenzaussagen auf die entsprechenden Konditionszahlen übertragen werden. Analog zu den Vektor- und Matrixnormen legen wir hierzu den Begriff der äquivalenten Konditionszahl fest.

Definition 2.44 Seien $\|\cdot\|_a$ und $\|\cdot\|_b$ zwei Vektornormen auf $\mathbb{R}^n$. Dann heißen die Konditionszahlen cond_a und cond_b äquivalent, wenn es reelle Zahlen $\alpha, \beta > 0$ gibt, so daß

$$\alpha \text{cond}_b(\mathbf{A}) \leq \text{cond}_a(\mathbf{A}) \leq \beta \text{cond}_b(\mathbf{A}) \quad (2.3.6)$$

für alle regulären $\mathbf{A} \in \mathbb{R}^{n \times n}$ gilt.

Als direkte Folgerung aus den Matrix- respektive Vektornormäquivalenzen erhalten wir für jede reguläre Matrix $\mathbf{A} \in \mathbb{R}^n$ die Ungleichungen

$$\frac{1}{n} \text{cond}_2(\mathbf{A}) \leq \text{cond}_1(\mathbf{A}) \leq n \text{cond}_2(\mathbf{A}), \quad (2.3.7)$$

$$\frac{1}{n} \text{cond}_\infty(\mathbf{A}) \leq \text{cond}_2(\mathbf{A}) \leq n \text{cond}_\infty(\mathbf{A}) \quad (2.3.8)$$

und

$$\frac{1}{n^2} \text{cond}_1(\mathbf{A}) \leq \text{cond}_\infty(\mathbf{A}) \leq n^2 \text{cond}_1(\mathbf{A}). \quad (2.3.9)$$

2.4 Der Banachsche Fixpunktsatz

Viele Iterationsverfahren zur Lösung eines linearen Gleichungssystems $\mathbf{Ax} = \mathbf{b}$ lassen sich in der Form

$$\mathbf{x}_{n+1} = F(\mathbf{x}_n) \text{ für } n = 0, 1, 2, \dots \quad (2.4.1)$$

schreiben, wobei F eine Abbildung der betrachteten Grundmenge in sich darstellt. Die durch die Vorgehensweise (2.4.1) gesuchte Lösung muß hierbei einen Fixpunkt der Abbildung F darstellen. Die Iterationsvorschrift (2.4.1) wird daher auch als Fixpunktiteration bezeichnet. Für diese Verfahrensklasse liefert der Banachsche Fixpunktsatz Aussagen zur Existenz und Eindeutigkeit eines Fixpunktes sowie eine *a priori* und eine *a posteriori* Fehlerabschätzung. Zunächst führen wir die hierzu notwendigen Begriffe des Fixpunktes und der Kontraktionszahl ein.

Definition 2.45 Ein Element x einer Menge $D \subset X$ heißt Fixpunkt eines Operators $F : D \subset X \rightarrow X$, falls

$$F(x) = x$$

gilt.

Definition 2.46 Sei X ein normierter Raum. Ein Operator

$$F : D \subset X \rightarrow X$$

heißt kontrahierend, wenn eine Zahl $0 \leq q < 1$ mit

$$\|F(x) - F(y)\| \leq q \|x - y\| \quad \forall x, y \in D$$

existiert. Die Zahl q heißt Kontraktionszahl des Operators F .

Satz 2.47 Kontrahierende Operatoren sind stetig und besitzen höchstens einen Fixpunkt.

Beweis:

Wir betrachten zunächst die Stetigkeit des Operators. Seien X ein normierter Raum, $F : D \subset X \rightarrow X$ ein kontrahierender Operator und $\{x_n\}_{n \in \mathbb{N}}$ eine Folge aus D mit $x_n \rightarrow x \in D$ für $n \rightarrow \infty$, dann folgt mit

$$0 \leq \|F(x_n) - F(x)\| \leq q \|x_n - x\| \rightarrow 0 \text{ für } n \rightarrow \infty$$

die Stetigkeit des Operators F .

Seien $x, y \in D$ Fixpunkte von F , dann erhalten wir

$$\|x - y\| = \|F(x) - F(y)\| \leq q \|x - y\|,$$

so daß mit

$$\underbrace{(1 - q)}_{>0} \underbrace{\|x - y\|}_{\geq 0} \leq 0$$

die Gleichung $\|x - y\| = 0$ und somit $x = y$ folgt. □

Satz 2.48 (Banachscher Fixpunktsatz)

Sei D eine vollständige Teilmenge eines normierten Raumes X und

$$F : D \rightarrow D$$

ein kontrahierender Operator, dann existiert genau ein Fixpunkt $x \in D$ von F , und die durch

$$x_{n+1} = F(x_n) \text{ für } n = 0, 1, 2, \dots$$

gegebene Folge konvergiert für jeden Startwert $x_0 \in D$ gegen x . Es gelten zudem die a priori Fehlerabschätzung

$$\|x_n - x\| \leq \frac{q^n}{1 - q} \|x_1 - x_0\|$$

und die a posteriori Fehlerabschätzung

$$\|x_n - x\| \leq \frac{q}{1 - q} \|x_n - x_{n-1}\|,$$

wobei q die Kontraktionszahl des Operators repräsentiert.

Beweis:

Mit $\mathbf{x}_0 \in D$ ist $\mathbf{x}_{n+1} = F(\mathbf{x}_n)$, $n = 0, 1, \dots$ wegen $F : D \rightarrow D$ wohldefiniert. Es gilt

$$\|\mathbf{x}_{n+1} - \mathbf{x}_n\| \leq q \|\mathbf{x}_n - \mathbf{x}_{n-1}\| \leq \dots \leq q^n \|\mathbf{x}_1 - \mathbf{x}_0\|. \quad (2.4.2)$$

Sei $m > n$, dann folgt

$$\begin{aligned} \|\mathbf{x}_n - \mathbf{x}_m\| &\leq \|\mathbf{x}_n - \mathbf{x}_{n+1}\| + \|\mathbf{x}_{n+1} - \mathbf{x}_{n+2}\| + \dots + \|\mathbf{x}_{m-1} - \mathbf{x}_m\| \\ &\stackrel{(2.4.2)}{\leq} (q^n + \dots + q^{m-1}) \|\mathbf{x}_1 - \mathbf{x}_0\| \\ &\leq q^n \sum_{i=0}^{\infty} q^i \|\mathbf{x}_1 - \mathbf{x}_0\| \\ &= \frac{q^n}{1 - q} \|\mathbf{x}_1 - \mathbf{x}_0\|. \end{aligned} \quad (2.4.3)$$

Wegen $|q| < 1$ erhalten wir

$$\|\mathbf{x}_n - \mathbf{x}_m\| \rightarrow 0 \text{ für } n \rightarrow \infty,$$

so daß mit $\{\mathbf{x}_n\}_{n \in \mathbb{N}}$ eine Cauchy-Folge vorliegt. Aufgrund der Vollständigkeit der Teilmenge D existiert ein $\mathbf{x} \in D$ mit $\mathbf{x}_n \rightarrow \mathbf{x}$, für $n \rightarrow \infty$. Mit

$$\mathbf{x} = \lim_{n \rightarrow \infty} \mathbf{x}_n = \lim_{n \rightarrow \infty} F(\mathbf{x}_{n-1}) \stackrel{\text{Satz 2.47}}{=} F\left(\lim_{n \rightarrow \infty} \mathbf{x}_{n-1}\right) = F(\mathbf{x})$$

ist $\mathbf{x}$ ein Fixpunkt von F , der nach Satz 2.47 eindeutig bestimmt ist.

Aus der Gleichung (2.4.3) erhalten wir die *a priori* Fehlerabschätzung

$$\|\mathbf{x}_n - \mathbf{x}\| = \lim_{m \rightarrow \infty} \|\mathbf{x}_n - \mathbf{x}_m\| \stackrel{(2.4.3)}{\leq} \frac{q^n}{1 - q} \|\mathbf{x}_1 - \mathbf{x}_0\|.$$

Aus

$$\begin{aligned} \|\mathbf{x}_n - \mathbf{x}\| &= \|F(\mathbf{x}_{n-1}) - F(\mathbf{x}_n) + F(\mathbf{x}_n) - F(\mathbf{x})\| \\ &\leq \|F(\mathbf{x}_{n-1}) - F(\mathbf{x}_n)\| + \|F(\mathbf{x}_n) - F(\mathbf{x})\| \\ &\leq q \|\mathbf{x}_{n-1} - \mathbf{x}_n\| + q \|\mathbf{x}_n - \mathbf{x}\| \end{aligned}$$

ergibt sich durch einfache Umformulierung die behauptete *a posteriori* Fehlerabschätzung

$$\|\mathbf{x}_n - \mathbf{x}\| \leq \frac{q}{1 - q} \|\mathbf{x}_{n-1} - \mathbf{x}_n\|.$$

□

Beispiel 2.49 Gesucht sei ein $x \in \mathbb{R}$ mit $x = 1 - \sin x$. Wir betrachten mit $D = [\varepsilon, 1]$, $0 < \varepsilon \leq 1 - \sin 1 < 1$, eine bezüglich der Norm $\|x\| = |x|$ vollständige Teilmenge des $\mathbb{R}$ und haben mit $f(x) = 1 - \sin x$ eine beliebig oft stetig differenzierbare Funktion vorliegen, die zudem $f : D \rightarrow D$ und $|f(x) - f(y)| \leq q|x - y|$ mit der Kontraktionszahl $q = \max_{x \in D} |f'(x)| = |\cos \varepsilon| < 1$ erfüllt. Somit besitzt $f(x)$ genau einen Fixpunkt $x \in D$ und die Iteration

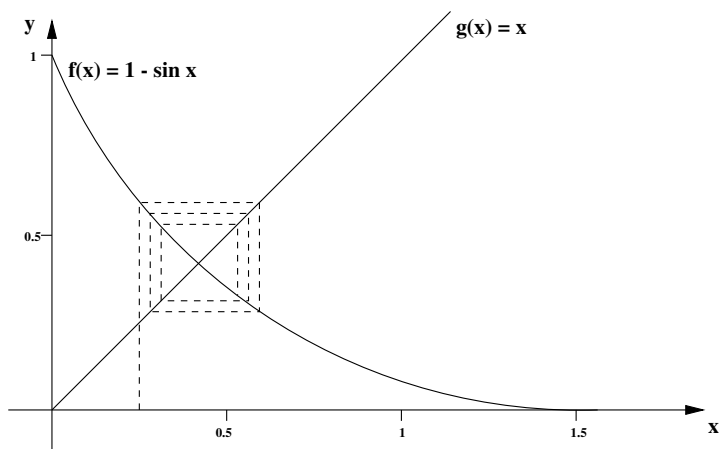


Bild 2.1 Iterationsverlauf für $f(x) = 1 - \sin x$ mit $x_0 = 0.25$.

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

konvergiert für alle $x_0 \in D$ gegen x . Die folgende Skizze verdeutlicht geometrisch den Verlauf der Iteration für $x_0 = 0.25$.

2.5 Übungsaufgaben

Aufgabe 1:

Gegeben sei ein lineares Gleichungssystem $\mathbf{Ax} = \mathbf{b}$ (1). Zeigen Sie, dass die Operationen

- (a) Multiplikation einer Gleichung in (1) mit einer komplexen Zahl ($\neq 0$),
- (b) Vertauschen zweier Gleichungen in (1)
- (c) Addition der j -ten Gleichung aus (1) zur i -ten Gleichung in (1),

stets ein zu (1) äquivalentes Gleichungssystem liefern.

Hinweis: Drücken Sie alle obigen Operationen durch eine Matrixmultiplikation aus und nutzen Sie Lemma 2.24.

Aufgabe 2:

Sei $\mathbf{A} = (a_{ij})_{i,j=1,\dots,n} \in \mathbb{R}^{n \times n}$ eine symmetrische und positiv definite Matrix. Zeigen Sie:

- (a) $a_{ii} > 0$, $i = 1, \dots, n$,
- (b) $\max_{i,j=1,\dots,n} |a_{ij}| = \max_{i=1,\dots,n} |a_{ii}|$.

Aufgabe 3:

Beweisen Sie die Aussage: Ist $\mathbf{A} \in \mathbb{R}^{n \times n}$ eine symmetrische, streng diagonal dominante Matrix mit positiven Diagonalelementen, dann ist $\mathbf{A}$ positiv definit.

Aufgabe 4:

Es sei $\mathbf{A} \in \mathbb{C}^{n \times n}$. Zeigen Sie:

$$\mathbf{A}^j \rightarrow \mathbf{0}, \quad j \rightarrow \infty \quad \Longleftrightarrow \quad \rho(\mathbf{A}) < 1.$$

Aufgabe 5:

Sei

$$\mathbf{A} = \begin{pmatrix} 0.78 & 0.563 \\ 0.913 & 0.659 \end{pmatrix}.$$

Berechnen Sie $\text{cond}_\infty(\mathbf{A}) = \|\mathbf{A}\|_\infty \|\mathbf{A}^{-1}\|_\infty$.

Aufgabe 6:

Sei $\mathbf{A} \in \mathbb{C}^{n \times n}$ regulär und $\mathbf{U} \in \mathbb{C}^{n \times n}$ unitär. Zeigen Sie:

$$\text{cond}_2(\mathbf{A}\mathbf{U}) = \text{cond}_2(\mathbf{A}) = \text{cond}_2(\mathbf{U}\mathbf{A}).$$

Aufgabe 7:

(a) Es sei

$$\mathbf{A} = \begin{pmatrix} \frac{1}{2} & 100 \\ 0 & \frac{1}{4} \end{pmatrix}$$

gegeben. Bestimmen Sie $\rho(\mathbf{A})$, $\|\mathbf{A}\|_\infty$, $\|\mathbf{A}\|_1$, $\|\mathbf{A}\|_F$.

(b) Skizzieren Sie die Einheitskugeln im $\mathbb{R}^2$ für die Vektornormen $\|\cdot\|_\infty$, $\|\cdot\|_1$, $\|\cdot\|_2$ und bestimmen Sie die besten Äquivalenzkonstanten für diese drei Normen im $\mathbb{R}^n$.

Aufgabe 8:

Es sei X ein Banachraum und $\mathbf{A} : X \rightarrow X$ ein Operator, dessen Potenz $\mathbf{A}^m$ für ein $m \in \mathbb{N}$ ein kontrahierender Operator ist. Beweisen Sie, dass genau ein Fixpunkt $\mathbf{x}$ von $\mathbf{A}$ existiert und dass die durch die Iterationsvorschrift

$$\begin{aligned} \mathbf{x}_n &= \mathbf{A}\mathbf{x}_{n-1}, \quad n = 1, 2, \dots \\ \mathbf{x}_0 &\in X \quad \text{beliebig} \end{aligned}$$

erzeugte Folge $\{\mathbf{x}_n\}_{n \in \mathbb{N}}$ gegen $\mathbf{x}$ konvergiert.

Aufgabe 9:

Sei $\|\cdot\|$ eine beliebige induzierte Matrixnorm auf $\mathbb{C}^{n \times n}$. Zeigen Sie: Es gilt

$$\rho(\mathbf{A}) = \lim_{k \rightarrow \infty} \|\mathbf{A}^k\|^{1/k}$$

für alle $\mathbf{A} \in \mathbb{C}^{n \times n}$.

Aufgabe 10:

Zeigen Sie, dass der Spektralradius keine Norm auf dem $\mathbb{R}^{n \times n}$ definiert. Wie sieht es mit dem Raum $\mathbb{R}_{\text{Symm}}^{n \times n}$ der symmetrischen, reellen $n \times n$ -Matrizen aus?

Aufgabe 11:

Gegeben sei die Funktion f mit

$$f(x) = 0.5 + \sin x - 2x.$$

- (a) Man zeige, dass f genau eine Nullstellen besitzt.
- (b) Man berechne die Nullstelle von f mit Hilfe des Fixpunktverfahrens, wobei die Voraussetzungen des Banachschen Fixpunktsatzes jeweils zu überprüfen sind.
- (c) Für die berechnete Näherung x_5 führe man eine a priori und eine a posteriori Fehlerabschätzung durch.

Hinweis: Eine einfache Funktionsabtastung zeigt, dass sich die Nullstelle im Intervall $[0.4, 0.5]$ befindet.

Aufgabe 12:

Gegeben sei die Funktion f mit

$$f(x) = \frac{x}{4} - 2 - 3 \ln x.$$

- (a) Man zeige, dass f genau zwei Nullstellen besitzt.
- (b) Man berechne die beiden Nullstellen von f mit Hilfe des Fixpunktverfahrens, wobei die Voraussetzungen des Banachschen Fixpunktsatzes jeweils zu überprüfen sind.
- (c) Für die berechnete Näherung x_5 führe man jeweils eine a priori und eine a posteriori Fehlerabschätzung durch.

Hinweis: Eine einfache Funktionsabtastung zeigt, dass sich die Nullstellen im Intervall $[0.5, 0.6]$ sowie im Intervall $[56, 57]$ befinden.

Aufgabe 13:

Susanne und Frank machen ein Physik-Experiment. Sie versuchen die Parameter x_1 und x_2 zu bestimmen und kennen die funktionale Beziehung

$$\begin{pmatrix} 6 & 3 \\ 1 & 2 \end{pmatrix} \mathbf{x} = \begin{pmatrix} a \\ 1 \end{pmatrix}.$$

Die beiden führen einen Versuch durch und messen $a = 1$. Als erfahrene Experimentatoren wissen Sie, dass in Experimenten immer Fehler auftauchen, haben jedoch keine Zeit, den Versuch zu wiederholen.

Schätzen Sie für Susanne und Frank in Abhängigkeit vom unbekannten Messfehler ϵ den relativen Fehler $\|\Delta \mathbf{x}\|/\|\mathbf{x}\|$ in der ∞ -Norm nach oben ab.

Aufgabe 14:

Berechnen Sie die Kondition der 3×3 Hilbert-Matrix

$$\mathbf{A} = \begin{pmatrix} 1 & 1/2 & 1/3 \\ 1/2 & 1/3 & 1/4 \\ 1/3 & 1/4 & 1/5 \end{pmatrix}$$

in der ∞ - und der 1-Norm. Wenn Sie wollen, überprüfen Sie ihr Ergebnis mit MATLAB. Nutzen Sie hierzu die Befehle `cond(A,1)` bzw. `cond(A,inf)` für die Konditionszahlen. Eine $n \times n$ Hilbert-Matrix wird mittels `hilb(n)` erzeugt.

Aufgabe 15:

Gesucht sind die Lösungen des nichtlinearen Gleichungssystems

$$\begin{aligned}x^2 + y^2 &= 4, \\ \frac{1}{16}x^2 + y^2 &= 1.\end{aligned}$$

- (a) Fertigen Sie eine Skizze an, die die Lage der Lösungen verdeutlicht. Bestimmen Sie für den 1. Quadranten einen *guten* ganzzahligen Startwert (x_0, y_0) .
- (b) Gesucht ist die Lösung im ersten Quadranten. Geben Sie eine geeignete Fixpunktgleichung an und weisen Sie hierfür die Voraussetzungen des Fixpunktsatzes von Banach nach.

3 Direkte Verfahren

Unter einem direkten Verfahren zur Lösung eines linearen Gleichungssystems versteht man eine Rechenvorschrift, die unter Vernachlässigung von Rundungsfehlern die exakte Lösung in endlich vielen Schritten ermittelt. Direkte Verfahren werden heutzutage nur selten zur unmittelbaren Lösung großer linearer Gleichungssysteme verwendet. Sie werden jedoch häufig in einer unvollständigen Form als Vorkonditionierer innerhalb iterativer Methoden genutzt und zur Lösung von Subproblemen eingesetzt.

3.1 Gauß-Elimination

Die Grundidee des Gaußschen Eliminationsverfahrens liegt in einer sukzessiven Transformation des Systems $\mathbf{Ax} = \mathbf{b}$ in ein äquivalentes System der Form

$$\mathbf{LRx} = \mathbf{b}$$

mit einer rechten oberen Dreiecksmatrix $\mathbf{R}$ und einer linken unteren Dreiecksmatrix $\mathbf{L}$, das durch einfaches Vorwärts- und anschließendes Rückwärtseinsetzen gelöst werden kann. Hierbei wird in der üblichen Formulierung des Verfahrens die Multiplikation $\tilde{\mathbf{b}} = \mathbf{L}^{-1}\mathbf{b}$ simultan mit der Berechnung der Matrix $\mathbf{R}$ durchgeführt. Wir bezeichnen im folgenden mit

$$\mathbf{P}_{kj} = \begin{pmatrix} 1 & & & & & & \\ & \ddots & & & & & \\ & & 1 & & & & \\ & & & 0 & \dots & \dots & 1 \\ & & & \vdots & 1 & & \vdots \\ & & & \vdots & & \ddots & \vdots \\ & & & \vdots & & & 1 & \vdots \\ & & 1 & \dots & \dots & \dots & 0 & \\ & & & & & & & 1 & \ddots & \\ & & & & & & & & & 1 \end{pmatrix} \quad \begin{array}{l} \leftarrow k\text{-te Zeile} \\ \\ \\ \leftarrow j\text{-te Zeile} \end{array} \quad (3.1.1)$$

stets eine Permutationsmatrix, die aus der Einheitsmatrix $\mathbf{I}$ durch Vertauschung der j -ten mit der k -ten Zeile ($j \geq k$) hervorgegangen ist. Für $k = j$ gilt hierbei $\mathbf{P}_{kj} = \mathbf{I}$. Wir werden im weiteren die Matrix $\mathbf{L}$ durch eine multiplikative Verknüpfung von Matrizen der Form

$$\mathbf{L}_k = \begin{pmatrix} 1 & & & & & \\ & \ddots & & & & \\ & & 1 & & & \\ & & -\ell_{k+1,k} & \ddots & & \\ & & \vdots & & \ddots & \\ & & -\ell_{n,k} & & & 1 \end{pmatrix} \quad (3.1.2)$$

darstellen. Derartige Matrizen, die sich höchstens in einer Spalte von der Einheitsmatrix unterscheiden, werden als Frobeniusmatrizen bezeichnet.

Definition 3.1 Die Zerlegung einer Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ in ein Produkt

$$\mathbf{A} = \mathbf{L}\mathbf{R}$$

aus einer linken unteren Dreiecksmatrix $\mathbf{L} \in \mathbb{R}^{n \times n}$ und einer rechten oberen Dreiecksmatrix $\mathbf{R} \in \mathbb{R}^{n \times n}$ heißt

LR-Zerlegung.

In der Literatur wird die LR-Zerlegung auch oft als LU-Zerlegung (lower, upper) bezeichnet. Speziell bei unvollständigen Zerlegungen, wie sie im Kapitel 5 beschrieben werden, ist der Begriff der unvollständigen LU-Zerlegung anstelle der unvollständigen LR-Zerlegung gängig.

Mit Hilfe der Matrizen (3.1.1) und (3.1.2) läßt sich der zentrale Teil des Gaußschen Eliminationsverfahrens wie folgt formulieren:

Algorithmus Gauß-Elimination I —

$\mathbf{A}^{(1)} := \mathbf{A}$
Für $k = 1, \dots, n-1$
Wähle aus der k -ten Spalte von $\mathbf{A}^{(k)}$ ein beliebiges Element $a_{jk}^{(k)} \neq 0$ mit $j \geq k$.
Definiere $\mathbf{P}_{kj}$ mit obigem j und k gemäß (3.1.1).
$\tilde{\mathbf{A}}^{(k)} := \mathbf{P}_{kj} \mathbf{A}^{(k)}$
Definiere $\mathbf{L}_k$ gemäß (3.1.2) mit $l_{ik} = \tilde{a}_{ik}^{(k)} / \tilde{a}_{kk}^{(k)}$, $i = k+1, \dots, n$.
$\mathbf{A}^{(k+1)} := \mathbf{L}_k \tilde{\mathbf{A}}^{(k)}$

Mit $\mathbf{A}^{(n)}$ liegt hierdurch die rechte obere Dreiecksmatrix $\mathbf{R}$ vor, die wie bereits beschrieben zur einfachen Lösung des linearen Gleichungssystems verwendet wird. Allgemein enthält $\mathbf{A}^{(i)}$ durch die spezielle Konstruktion der Matrizen $\mathbf{L}_k$, $k = 1, \dots, i-1$ bis einschließlich zur $(i-1)$ -ten Spalte unterhalb der Diagonalen ausschließlich Nullelemente.

Wir wenden uns jetzt der Frage nach der Existenz und Eindeutigkeit von LR-Zerlegungen zu. Einfache Beispiele wie die Matrix

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

zeigen, daß nicht jede reguläre Matrix notwendigerweise eine LR-Zerlegung besitzt. Analog ist auch der Nachweis der Eindeutigkeit der LR-Zerlegungen nur unter einer zusätzlichen Bedingung an eine der beiden Dreiecksmatrizen möglich. Bevor wir uns mit den eigentlichen Existenz- und Eindeutigkeitsaussagen beschäftigen, ist es sinnvoll, zunächst die Begriffe der Hauptabschnittsmatrix und -determinante einzuführen und einige hilfreiche Lemmata zu beweisen.

Definition 3.2 Sei $\mathbf{A} \in \mathbb{R}^{n \times n}$ gegeben, dann heißt

$$\mathbf{A}[k] := \begin{pmatrix} a_{11} & \dots & a_{1k} \\ \vdots & \ddots & \vdots \\ a_{k1} & \dots & a_{kk} \end{pmatrix} \in \mathbb{R}^{k \times k} \text{ für } k \in \{1, \dots, n\}$$

die führende $k \times k$ -Hauptabschnittsmatrix von $\mathbf{A}$ und $\det \mathbf{A}[k]$ die führende $k \times k$ -Hauptabschnittsdeterminante von $\mathbf{A}$.

Lemma 3.3 Seien $\ell_i = (0, \dots, 0, \ell_{i+1,i}, \dots, \ell_{n,i})^T \in \mathbb{R}^n$ und $\mathbf{e}_i \in \mathbb{R}^n$ der i -te Einheitsvektor, dann gilt für $\mathbf{L}_i = \mathbf{I} - \ell_i \mathbf{e}_i^T \in \mathbb{R}^{n \times n}$

$$(a) \quad \mathbf{L}_i^{-1} = \mathbf{I} + \ell_i \mathbf{e}_i^T.$$

$$(b) \quad \mathbf{L}_1^{-1} \mathbf{L}_2^{-1} \dots \mathbf{L}_k^{-1} = \mathbf{I} + \sum_{i=1}^k \ell_i \mathbf{e}_i^T \text{ für } k = 1, \dots, n-1.$$

Beweis:

Zu (a):

Da $\mathbf{L}_i$ eine untere Dreiecksmatrix mit Einheitsdiagonale darstellt, existiert genau eine Matrix $\mathbf{L}_i^{-1}$ mit $\mathbf{L}_i^{-1} \mathbf{L}_i = \mathbf{L}_i \mathbf{L}_i^{-1} = \mathbf{I}$. Hieraus folgt die Behauptung (a) durch

$$(\mathbf{I} - \ell_i \mathbf{e}_i^T)(\mathbf{I} + \ell_i \mathbf{e}_i^T) = \mathbf{I} - \ell_i \mathbf{e}_i^T + \underbrace{\ell_i \mathbf{e}_i^T - \ell_i \mathbf{e}_i^T \ell_i \mathbf{e}_i^T}_{=0} = \mathbf{I}.$$

Zu (b):

Wir führen den Beweis durch Induktion über k . Für $k = 1$ liefert (a) die Behauptung. Gelte die Aussage für $j = 1, \dots, k < n-1$, dann folgt

$$\begin{aligned}
L_1^{-1} \dots L_k^{-1} L_{k+1}^{-1} &= \left(I + \sum_{i=1}^k \ell_i e_i^T \right) \left(I + \ell_{k+1} e_{k+1}^T \right) \\
&= I + \ell_{k+1} e_{k+1}^T + \sum_{i=1}^k \ell_i e_i^T + \sum_{i=1}^k \ell_i \underbrace{e_i^T \ell_{k+1}}_{=0} e_{k+1}^T \\
&= I + \sum_{i=1}^{k+1} \ell_i e_i^T.
\end{aligned}$$

□

Lemma 3.4 Sei $L \in \mathbb{R}^{n \times n}$ eine reguläre linke untere Dreiecksmatrix, dann stellt auch $L^{-1} \in \mathbb{R}^{n \times n}$ eine linke untere Dreiecksmatrix dar. Für reguläre rechte obere Dreiecksmatrizen $R \in \mathbb{R}^{n \times n}$ gilt die analoge Aussage.

Beweis:

Definieren wir $D = \text{diag} \{\ell_{11}, \dots, \ell_{nn}\}$ mittels der Diagonaleinträge der Matrix L , dann gilt $\det D \neq 0$ und

$$\tilde{L} := D^{-1} L$$

stellt laut Lemma 2.26 ebenfalls eine untere Dreiecksmatrix dar, die zudem eine Einheitsdiagonale besitzt. Somit hat $\tilde{L}$ die Form $\tilde{L} = I + \sum_{i=1}^{n-1} \tilde{\ell}_i e_i^T$ mit

$$\tilde{\ell}_i = (0, \dots, 0, \tilde{\ell}_{i+1,i}, \dots, \tilde{\ell}_{n,i})^T$$

und kann unter Verwendung der Matrizen $\tilde{L}_i = I + \tilde{\ell}_i e_i^T$ ($i = 1, \dots, n-1$) als ein Produkt

$$\tilde{L} = \tilde{L}_1 \cdot \dots \cdot \tilde{L}_{n-1}$$

dargestellt werden. Für die Inverse ergibt sich $\tilde{L}^{-1} = \tilde{L}_{n-1}^{-1} \cdot \dots \cdot \tilde{L}_1^{-1}$ mit $\tilde{L}_i^{-1} = I - \tilde{\ell}_i e_i^T$ ($i = 1, \dots, n-1$) laut Lemma 3.3. Die Matrix L^{-1} läßt sich folglich als Produkt unterer Dreiecksmatrizen in der Form $L^{-1} = \tilde{L}_{n-1}^{-1} \cdot \dots \cdot \tilde{L}_1^{-1} D^{-1}$ schreiben und stellt somit nach Lemma 2.26 ebenfalls eine linke untere Dreiecksmatrix dar.

Mit $R^T (R^{-1})^T = (R^{-1} R)^T = I = R^T (R^T)^{-1}$ folgt $(R^T)^{-1} = (R^{-1})^T$. Unter Verwendung des obigen Beweisteils stellt mit $L = R^T$ auch $L^{-1} = (R^T)^{-1}$ eine linke untere Dreiecksmatrix dar, wodurch $R^{-1} = \left((R^T)^{-1} \right)^T$ eine rechte obere Dreiecksmatrix repräsentiert. □

Lemma 3.5 Sei $A = (a_{ij})_{i,j=1,\dots,n} \in \mathbb{R}^{n \times n}$ und sei $L = (\ell_{ij})_{i,j=1,\dots,n} \in \mathbb{R}^{n \times n}$ eine untere Dreiecksmatrix, dann gilt

$$(LA)[k] = L[k]A[k] \quad \text{für } k = 1, \dots, n.$$

Beweis:

Sei $k \in \{1, \dots, n\}$. Für $i, j \in \{1, \dots, k\}$ folgt mit $\ell_{im} = 0$ für $m > k \geq i$

$$((LA)[k])_{ij} = \sum_{m=1}^n \ell_{im} a_{mj} = \sum_{m=1}^k \ell_{im} a_{mj} + \sum_{m=k+1}^n \ell_{im} a_{mj} = \sum_{m=1}^k \ell_{im} a_{mj} = (L[k]A[k])_{ij}.$$

□

Satz 3.6 (Existenz einer LR-Zerlegung I)

Sei $A \in \mathbb{R}^{n \times n}$ eine reguläre Matrix, dann existiert eine Permutationsmatrix $P \in \mathbb{R}^{n \times n}$ derart, daß PA eine LR-Zerlegung besitzt.

Beweis:

Für alle im Algorithmus I berechneten Matrizen $A^{(k)}$ gilt

$$a_{i\ell}^{(k)} = 0 \quad \text{für } \ell = 1, \dots, k-1, \quad i > \ell. \quad (3.1.3)$$

Da $\det L_\ell \neq 0 \neq \det P_{\ell, j_\ell}$ für alle $\ell = 1, \dots, k-1$, $j_\ell \geq \ell$ gilt, folgt

$$\det A^{(k)} = \det (L_{k-1} P_{k-1, j_{k-1}} \dots L_1 P_{1, j_1} A) \neq 0.$$

Somit existiert ein $i \in \{k, \dots, n\}$ mit $a_{ik}^{(k)} \neq 0$ und der Algorithmus I bricht nicht vor der Berechnung von $A^{(n)}$ ab. Zudem stellt

$$R := A^{(n)} = L_{n-1} P_{n-1, j_{n-1}} \dots L_1 P_{1, j_1} A \quad (3.1.4)$$

mit (3.1.3) eine obere Dreiecksmatrix dar. Alle L_i lassen sich hierbei in der Form

$$L_i = I - \ell_i e_i^T, \quad \ell_i = (0, \dots, 0, \ell_{i+1, i}, \dots, \ell_{n, i})^T$$

schreiben, und es gilt

$$P_{k, j_k} e_i = e_i \text{ sowie } P_{k, j_k} \ell_i = \hat{\ell}_i = (0, \dots, 0, \hat{\ell}_{i+1, i}, \dots, \hat{\ell}_{n, i})^T \quad \text{für alle } i < k \leq j_k.$$

Für $i < k \leq j_k$ folgt hiermit unter Berücksichtigung von $P_{k, j_k} = P_{k, j_k}^T$ die Gleichung

$$\begin{aligned} P_{k, j_k} L_i &= P_{k, j_k} (I - \ell_i e_i^T) = P_{k, j_k} - \hat{\ell}_i e_i^T \\ &= P_{k, j_k} - \hat{\ell}_i (P_{k, j_k} e_i)^T = P_{k, j_k} - \hat{\ell}_i e_i^T P_{k, j_k} = \hat{L}_i P_{k, j_k} \end{aligned}$$

mit einer unteren Dreiecksmatrix $\hat{L}_i = I - \hat{\ell}_i e_i^T$. Die Verwendung der Gleichung (3.1.4) liefert

$$R = \underbrace{L_{n-1} \hat{L}_{n-2} \dots \hat{L}_1}_{\tilde{L} :=} \underbrace{P_{n-1, j_{n-1}} \dots P_{1, j_1}}_{P :=} A$$

mit einer Permutationsmatrix P . Da $\tilde{L}$ eine untere Dreiecksmatrix darstellt, folgt mit Lemma 3.4, daß

$$L := \tilde{L}^{-1}$$

eine untere Dreiecksmatrix repräsentiert und es ergibt sich die behauptete Darstellung

$$LR = PA.$$

□

Satz 3.7 (Existenz einer LR-Zerlegung II)

Sei $A \in \mathbb{R}^{n \times n}$ regulär, dann besitzt A genau dann eine LR-Zerlegung, wenn

$$\det A[k] \neq 0 \quad \forall k = 1, \dots, n$$

gilt.

Beweis:

„ $\Rightarrow$ “: Gelte $\mathbf{A} = \mathbf{LR}$.

Aufgrund der Regularität der Matrix $\mathbf{A}$ liefert der Determinantenmultiplikationssatz

$$\det \mathbf{L}[n] \cdot \det \mathbf{R}[n] = \det \mathbf{A}[n] \neq 0$$

und folglich

$$\det \mathbf{L}[n] \neq 0 \neq \det \mathbf{R}[n].$$

Da $\mathbf{L}$ und $\mathbf{R}$ Dreiecksmatrizen repräsentieren, folgt hierdurch

$$\det \mathbf{L}[k] \neq 0 \neq \det \mathbf{R}[k]$$

für $k = 1, \dots, n$ und es ergibt sich mit Lemma 3.5

$$\det \mathbf{A}[k] = \det(\mathbf{LR})[k] = \det \mathbf{L}[k] \cdot \det \mathbf{R}[k] \neq 0.$$

„ $\Leftarrow$ “: Gelte $\det \mathbf{A}[k] \neq 0$ für alle $k = 1, \dots, n$.

$\mathbf{A}$ besitzt eine LR-Zerlegung, falls Algorithmus I mit $\mathbf{P}_{kk} = \mathbf{I}$, $k = 1, \dots, n-1$ durchgeführt werden kann, das heißt, wenn $a_{kk}^{(k)} \neq 0$ für $k = 1, \dots, n-1$ gilt.

Für $k = 1$ gilt $a_{11}^{(1)} = \det \mathbf{A}[1] \neq 0$, wodurch $\mathbf{P}_{11} = \mathbf{I}$ wählbar ist.

Sei $a_{kk}^{(k)} \neq 0$ für $k < n-1$, dann folgt

$$\mathbf{A}^{(k+1)} = \mathbf{L}_k \dots \mathbf{L}_1 \mathbf{A},$$

und wir erhalten wiederum mit Lemma 3.5

$$\det \mathbf{A}^{(k+1)}[k+1] = \det \mathbf{L}_k[k+1] \cdot \dots \cdot \det \mathbf{L}_1[k+1] \cdot \det \mathbf{A}[k+1] \neq 0.$$

Da $\mathbf{A}^{(k+1)}[k+1]$ eine obere Dreiecksmatrix darstellt, folgt

$$a_{k+1,k+1}^{(k+1)} \neq 0.$$

□

Satz 3.8 (Eindeutigkeit der LR-Zerlegung)

Sei $\mathbf{A} \in \mathbb{R}^{n \times n}$ regulär mit $\det \mathbf{A}[k] \neq 0$ für $k = 1, \dots, n$, dann existiert genau eine LR-Zerlegung von $\mathbf{A}$ derart, daß $\mathbf{L}$ eine Einheitsdiagonale besitzt.

Beweis:

Mit Satz 3.7 existiert mindestens eine LR-Zerlegung der Matrix $\mathbf{A}$. Seien zwei LR-Zerlegungen der Matrix $\mathbf{A}$ durch $\mathbf{L}_1 \mathbf{R}_1 = \mathbf{A} = \mathbf{L}_2 \mathbf{R}_2$ gegeben, wobei $\mathbf{L}_1$ und $\mathbf{L}_2$ Einheitsdiagonalen besitzen, dann folgt

$$\mathbf{R}_2 \mathbf{R}_1^{-1} = \mathbf{L}_2^{-1} \mathbf{L}_1.$$

Mit Lemma 2.26 und Lemma 3.4 ist somit $\mathbf{L}_2^{-1} \mathbf{L}_1$ zugleich eine linke untere und rechte obere Dreiecksmatrix, die eine Einheitsdiagonale besitzt. Folglich gilt

$$\mathbf{L}_2^{-1} \mathbf{L}_1 = \mathbf{I}$$

und wir erhalten $L_1 = L_2$ und $R_1 = R_2$.

□

Bemerkung:

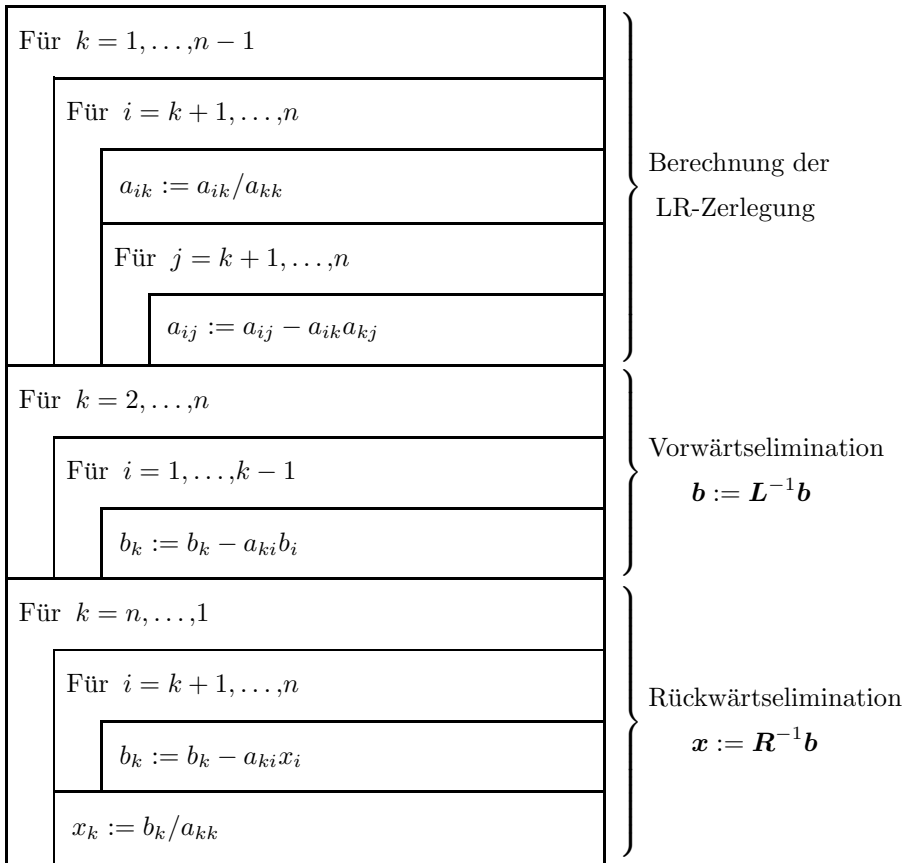
Ohne die Forderung, daß L eine Einheitsdiagonale besitzt, folgt

$$L_1 = L_2 D, R_1 = D^{-1} R_2$$

mit einer regulären Diagonalmatrix D . Somit sind LR-Zerlegungen durch Multiplikation mit einer Diagonalmatrix ineinander überführbar.

Zur Lösung des linearen Gleichungssystems $Ax = b$ ergibt sich die folgende explizite Form des Gauß-Algorithmus.

Algorithmus Gauß-Elimination ohne Pivotisierung —



Zur Analyse des Rechenaufwandes betrachten wir lediglich die zeitaufwendigen Multiplikationen und Divisionen:

$$\begin{aligned}
\# \text{ Divisionen} &= \sum_{k=1}^{n-1} (n-k) + n = \sum_{i=1}^{n-1} i + n = \sum_{i=1}^n i = \frac{n(n+1)}{2} \\
\# \text{ Multiplikationen} &= \sum_{k=1}^{n-1} (n-k)^2 + \frac{n(n-1)}{2} + \frac{n(n-1)}{2} \\
&= \sum_{i=1}^{n-1} i^2 + n(n-1) = \frac{(n-1)n(2n-1)}{6} + n(n-1) \\
&= \frac{n^3}{3} + \frac{n^2}{2} - \frac{5n}{6}.
\end{aligned}$$

Hiermit erhalten wir $\# \text{ Division} + \# \text{ Multiplikation} = \frac{n^3}{3} + n^2 - \frac{n}{3}$.

Beträgt die Rechenzeit für eine Multiplikation bzw. Division $1\mu\text{sec} = 10^{-6}\text{sec}$, dann benötigt der Gauß-Algorithmus für eine reelle $n \times n$ Matrix $\mathbf{A}$ die in der folgenden Tabelle aufgeführten Rechenzeiten:

n	Zeit
10^3	$\approx 5 \text{ min } 30 \text{ sec}$
10^4	$\approx 4 \text{ Tage}$
10^5	$\approx 10 \text{ Jahre } 7 \text{ Monate}$

Tabelle 3.1 Rechenzeiten des Gaußschen Eliminationsverfahrens

Gleichungssysteme mit einer Anzahl von 10^5 Unbekannten treten bereits bei impliziten Finite-Volumen- und Finite-Differenzen-Verfahren für die zweidimensionalen Euler-Gleichungen der Gasdynamik auf, wenn mit 25000 Kontrollvolumina respektive Gitterpunkten eine durchaus gängige Diskretisierung des Strömungsgebietes vorliegt. Bei solchen Anwendungsfällen erweist sich der Gauß-Algorithmus auch bei Vernachlässigung der Rundungsfehler und des immensen Speicherplatzbedarfs bereits aus Rechenzeitgründen als unpraktikabel. Betrachten wir hingegen Gleichungssysteme mit einer kleinen Anzahl von Unbekannte, so liegt mit dem Gaußschen Eliminationsverfahren eine direkte Methode vor, die häufig den iterativen Algorithmen überlegen ist. Mit dem folgenden Beispiel sollen die Vorteile einer Pivotisierung im Hinblick auf die Genauigkeit des ermittelten Lösungsvektors verdeutlicht werden.

Beispiel 3.9 Sei $\varepsilon \ll 1$ derart, daß bei Maschinengenauigkeit

$$1 \pm \varepsilon = 1 \quad \text{respektive} \quad 1 \pm \frac{1}{\varepsilon} = \frac{1}{\varepsilon}$$

gilt. Wir betrachten das Gleichungssystem $\mathbf{Ax} = \mathbf{b}$ in der Form

$$\begin{array}{rcl}
\varepsilon x_1 & + & 2x_2 = 1 \\
x_1 & + & x_2 = 1,
\end{array}$$

das die exakten Lösung

$$x_1 = \frac{1}{2 - \varepsilon} \approx 0.5, \quad x_2 = \frac{1 - \varepsilon}{2 - \varepsilon} \approx 0.5$$

besitzt. Mit dem Gaußschen Eliminationsverfahren erhalten wir

$$\begin{aligned} \varepsilon x_1 + 2x_2 &= 1 \\ \left(1 - \frac{2}{\varepsilon}\right)x_2 &= 1 - \frac{1}{\varepsilon} \end{aligned}$$

und damit aufgrund der vorliegenden Rechengenauigkeit $x_2 = \frac{1}{\varepsilon} \cdot \frac{\varepsilon}{2} = 0.5$. Einsetzen in die erste Gleichung liefert $\varepsilon x_1 + 1 = 1$, wodurch $x_1 = 0$ folgt.

Eine vorherige Zeilenvertauschung führt dagegen zum Gleichungssystem

$$\begin{aligned} x_1 + x_2 &= 1 \\ \varepsilon x_1 + 2x_2 &= 1. \end{aligned}$$

Hieraus folgt mit dem Gaußschen Eliminationsverfahren:

$$\begin{aligned} x_1 + x_2 &= 1 \\ (2 - \varepsilon)x_2 &= 1 - \varepsilon, \end{aligned}$$

so daß wir $x_2 = 0.5$ und $x_1 = 1 - x_2 = 0.5$ als Lösung erhalten. Anhand des vorliegenden Gleichungssystems wollen wir die Vorteile der Pivotisierung näher untersuchen. Betrachten wir zunächst das Gaußsche Eliminationsverfahren ohne Pivotisierung. Das Element $a_{11}^{(1)} = \varepsilon$ der Matrix

$$\mathbf{A} = \begin{pmatrix} a_{11}^{(1)} & a_{12}^{(1)} \\ a_{21}^{(1)} & a_{22}^{(1)} \end{pmatrix} = \begin{pmatrix} \varepsilon & 2 \\ 1 & 1 \end{pmatrix}$$

ist deutlich kleiner als $a_{21}^{(1)} = 1$. Hierdurch wird im Algorithmus ein sehr großer Quotient $\frac{a_{21}^{(1)}}{a_{11}^{(1)}} = \frac{1}{\varepsilon}$ gebildet. Da auch die weiteren Matrixelemente in der Größenordnung von $a_{21}^{(1)}$ liegen, erhalten wir

$$a_{22}^{(2)} = \underbrace{a_{22}^{(1)}}_{=\mathcal{O}(1) \text{ für } \varepsilon \rightarrow 0} - \frac{a_{12}^{(1)} a_{21}^{(1)}}{\underbrace{a_{11}^{(1)}}_{=\mathcal{O}(\frac{1}{\varepsilon}) \text{ für } \varepsilon \rightarrow 0}} = 1 - \frac{2}{\varepsilon} = \mathcal{O}\left(\frac{1}{\varepsilon}\right) \quad \text{für } \varepsilon \rightarrow 0.$$

Innerhalb der Matrix $\mathbf{A}^{(2)}$ weisen somit alle verbleibenden Nichtnullelemente eine deutlich unterschiedliche Quantität auf. Analog ergibt sich

$$\mathbf{b} = \begin{pmatrix} 1 \\ 1 - \frac{1}{\varepsilon} \end{pmatrix} = \begin{pmatrix} \mathcal{O}(1) \\ \mathcal{O}(\frac{1}{\varepsilon}) \end{pmatrix} \quad \text{für } \varepsilon \rightarrow 0.$$

Liegt ein sehr kleiner Matrixkoeffizient $a_{11}^{(1)} = \varepsilon$ vor, so führt die Maschinengenauigkeit zu $a_{22}^{(2)} = -\frac{2}{\varepsilon}$ und $b_2^{(2)} = -\frac{1}{\varepsilon}$. Folglich haben im vorliegenden Fall die Werte $a_{22}^{(1)}$ und $b_2^{(1)}$ keinen Einfluß auf die Lösung des Gleichungssystems. Die zweite Gleichung ist deshalb für hinreichend kleines $\varepsilon \neq 0$ numerisch äquivalent zur ersten Gleichung für $\varepsilon = 0$. Daher erhalten wir stets das Endergebnis

$$x_2 = \frac{b_1^{(1)}}{a_{12}^{(1)}} \text{ und } x_1 = 0.$$

Durch die vorgestellte Zeilenvertauschung liegt ein kleiner Quotient $\frac{\tilde{a}_{21}^{(1)}}{\tilde{a}_{11}^{(1)}} = \varepsilon$ vor, wodurch sich alle relevanten Elemente der Matrix $\mathbf{A}^{(2)}$ in der gleichen Größenordnung befinden. Folglich ergeben sich keine Probleme aufgrund der vorliegenden Rechengenauigkeit. Bei Verwendung einer Pivotisierung auf der Basis des größten Zeilen- oder Spaltenelementes liegt der Quotient stets im Intervall $[-1, 1]$, wogegen bei einer direkten Nutzung des Gaußschen Eliminationsverfahrens keine allgemeingültige Schranke für den Quotienten angegeben werden kann.

Die Permutation von Zeilen oder Spalten erweist sich somit nicht nur im Fall $a_{kk}^{(k)} = 0$ als sinnvoll.

Wir unterscheiden drei Pivotisierungsarten:

(a) Spaltenpivotisierung:

Definiere $\mathbf{P}_{kj}$ gemäß (3.1.1) mit

$$j = \text{index} \max_{j=k, \dots, n} |a_{jk}^{(k)}|,$$

und betrachte das zu

$$\mathbf{A}^{(k)} \mathbf{x} = \mathbf{b}$$

äquivalente System

$$\mathbf{P}_{kj} \mathbf{A}^{(k)} \mathbf{x} = \mathbf{P}_{kj} \mathbf{b}.$$

(b) Zeilenpivotisierung:

Definiere $\mathbf{P}_{kj}$ gemäß (3.1.1) mit

$$j = \text{index} \max_{j=k, \dots, n} |a_{kj}^{(k)}|,$$

und betrachte das System

$$\begin{aligned} \mathbf{A}^{(k)} \mathbf{P}_{kj} \mathbf{y} &= \mathbf{b} \\ \mathbf{x} &= \mathbf{P}_{kj} \mathbf{y}. \end{aligned}$$

(c) Vollständige Pivotisierung:

Definiere $\mathbf{P}_{k,j_1}$ und $\mathbf{P}_{k,j_2}$ gemäß (3.1.1) mit

$$\begin{aligned} j_1 &= \text{index} \max_{j=k, \dots, n} \left(\max_{i=k, \dots, n} |a_{ji}^{(k)}| \right) \\ j_2 &= \text{index} \max_{j=k, \dots, n} \left(\max_{i=k, \dots, n} |a_{ij}^{(k)}| \right) \end{aligned}$$

und betrachte das System

$$\begin{aligned} \mathbf{P}_{k,j_1} \mathbf{A}^{(k)} \mathbf{P}_{k,j_2} \mathbf{y} &= \mathbf{P}_{k,j_1} \mathbf{b} \\ \mathbf{x} &= \mathbf{P}_{k,j_2} \mathbf{y}. \end{aligned}$$

Bemerkung:

Wird das Gaußsche Eliminationsverfahren auf ein Gleichungssystem

$$\mathbf{A}\mathbf{x} = \mathbf{b}$$

mit einer singulären Matrix bei vollständiger Pivotisierung angewendet, so existiert ein $k \leq n$ mit

$$\max_{i,j=k,\dots,n} |a_{ij}^{(k)}| = 0.$$

In diesem Fall ist das System mit Lemma 2.25 genau dann lösbar, wenn sich bei Durchführung aller Multiplikationen mit $\mathbf{L}_i$ und $\mathbf{P}_{i,j_i}$, $i = 1, \dots, k-1$ angewandt auf die erweiterte Koeffizientenmatrix

$$(\mathbf{A}^{(1)}, \mathbf{b}^{(1)}) = (\mathbf{A}, \mathbf{b})$$

die Form

$$(\mathbf{A}^{(k)}, \mathbf{b}^{(k)}) \text{ mit } \mathbf{b}^{(k)} = (b_1^{(k)}, \dots, b_{k-1}^{(k)}, 0, \dots, 0)^T$$

ergeben hat.

3.2 Cholesky-Zerlegung

Für symmetrische, positiv definite Matrizen kann der beim Gaußschen Eliminationsverfahren benötigte Aufwand zur Berechnung einer LR-Zerlegung verringert werden.

Definition 3.10 Die Zerlegung einer Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ in ein Produkt

$$\mathbf{A} = \mathbf{L}\mathbf{L}^T$$

mit einer linken unteren Dreiecksmatrix $\mathbf{L} \in \mathbb{R}^{n \times n}$ heißt Cholesky-Zerlegung.

Satz 3.11 (Existenz und Eindeutigkeit der Cholesky-Zerlegung)

Zu jeder symmetrischen, positiv definiten Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ existiert genau eine linke untere Dreiecksmatrix $\mathbf{L} \in \mathbb{R}^{n \times n}$ mit $\ell_{ii} > 0, i = 1, \dots, n$ derart, daß

$$\mathbf{A} = \mathbf{L}\mathbf{L}^T$$

gilt.

Beweis:

Sei $\mathbf{x} \in \mathbb{R}^k \setminus \{\mathbf{0}\}$, $k < n$, dann definieren wir $\mathbf{y} = (\mathbf{x}, 0, \dots, 0)^T \in \mathbb{R}^n \setminus \{\mathbf{0}\}$. Somit folgt

$$\mathbf{x}^T \mathbf{A}[k] \mathbf{x} = \mathbf{y}^T \mathbf{A} \mathbf{y} > 0,$$

so daß alle $\mathbf{A}[k]$ für $k = 1, \dots, n$ positiv definit sind.

Eine Induktion über n liefert nun die Behauptung: Für $n = 1$ gilt $\mathbf{A} = (a_{11}) > 0$, wodurch $\ell_{11} := \sqrt{a_{11}} > 0$ die Darstellung $\mathbf{A} = \mathbf{L}\mathbf{L}^T$ mit $\mathbf{L} = (\ell_{11})$ liefert. Sei die Behauptung für $n = 1, \dots, j$ erfüllt, dann folgt für $n = j+1$

$$\mathbf{A} = \begin{pmatrix} \mathbf{A}[j] & \mathbf{c} \\ \mathbf{c}^T & a_{nn} \end{pmatrix},$$

wobei $\mathbf{A}[j]$ positiv definit ist. Somit existiert genau eine linke untere Dreiecksmatrix $\mathbf{L}_j \in \mathbb{R}^{j \times j}$ mit $\ell_{ii} > 0$, $i = 1, \dots, j$ und

$$\mathbf{A}[j] = \mathbf{L}_j \mathbf{L}_j^T.$$

Wir machen nun den Ansatz

$$\mathbf{L}_n := \begin{pmatrix} \mathbf{L}_j & \mathbf{0} \\ \mathbf{d}^T & \alpha \end{pmatrix} \quad (3.2.1)$$

mit $\mathbf{d} \in \mathbb{R}^j$, $\alpha \in \mathbb{R}$ derart, daß

$$\begin{pmatrix} \mathbf{A}[j] & \mathbf{c} \\ \mathbf{c}^T & a_{nn} \end{pmatrix} = \mathbf{A}[n] = \mathbf{L}_n \mathbf{L}_n^T = \begin{pmatrix} \mathbf{A}[j] & \mathbf{L}_j \mathbf{d} \\ \mathbf{d}^T \mathbf{L}_j^T & \mathbf{d}^T \mathbf{d} + \alpha^2 \end{pmatrix}$$

gelten soll. Wegen $0 \neq \det \mathbf{A}[j] = (\det \mathbf{L}_j)^2$ ist $\mathbf{L}_j$ invertierbar und damit $\mathbf{d} = \mathbf{L}_j^{-1} \mathbf{c}$ eindeutig bestimmt. Weiter gilt $\alpha^2 = a_{nn} - \mathbf{d}^T \mathbf{d}$. Aufgrund der positiven Definitheit von $\mathbf{A} = \mathbf{A}[n]$ gilt $\det \mathbf{A}[n] > 0$, so daß sich

$$0 < \frac{\det(\mathbf{A}[n])}{(\det(\mathbf{L}_j))^2} = \alpha^2$$

ergibt. Damit erhalten wir

$$\alpha = \sqrt{a_{nn} - \mathbf{d}^T \mathbf{d}} \in \mathbb{R}^+,$$

wodurch mit (3.2.1) die gesuchte Matrix vorliegt. $\square$

Wir nehmen eine spaltenweise Berechnung der Matrixkoeffizienten vor. Bei der Herleitung des Algorithmus gehen wir somit davon aus, daß alle ℓ_{ij} für $i = 1, \dots, n$ und $j \leq k-1$ bekannt sind. Dann folgt aus

$$\mathbf{A} = \begin{pmatrix} \ell_{11} & & \\ \vdots & \ddots & \\ \ell_{n1} & \dots & \ell_{nn} \end{pmatrix} \begin{pmatrix} \ell_{11} & \dots & \ell_{n1} \\ & \ddots & \vdots \\ & & \ell_{nn} \end{pmatrix}$$

die Beziehung

$$a_{kk} = \ell_{k1}^2 + \dots + \ell_{kk}^2 = \sum_{j=1}^k \ell_{kj}^2,$$

wodurch sich ℓ_{kk} gemäß

$$\ell_{kk} = \sqrt{a_{kk} - \sum_{j=1}^{k-1} \ell_{kj}^2} \quad (3.2.2)$$

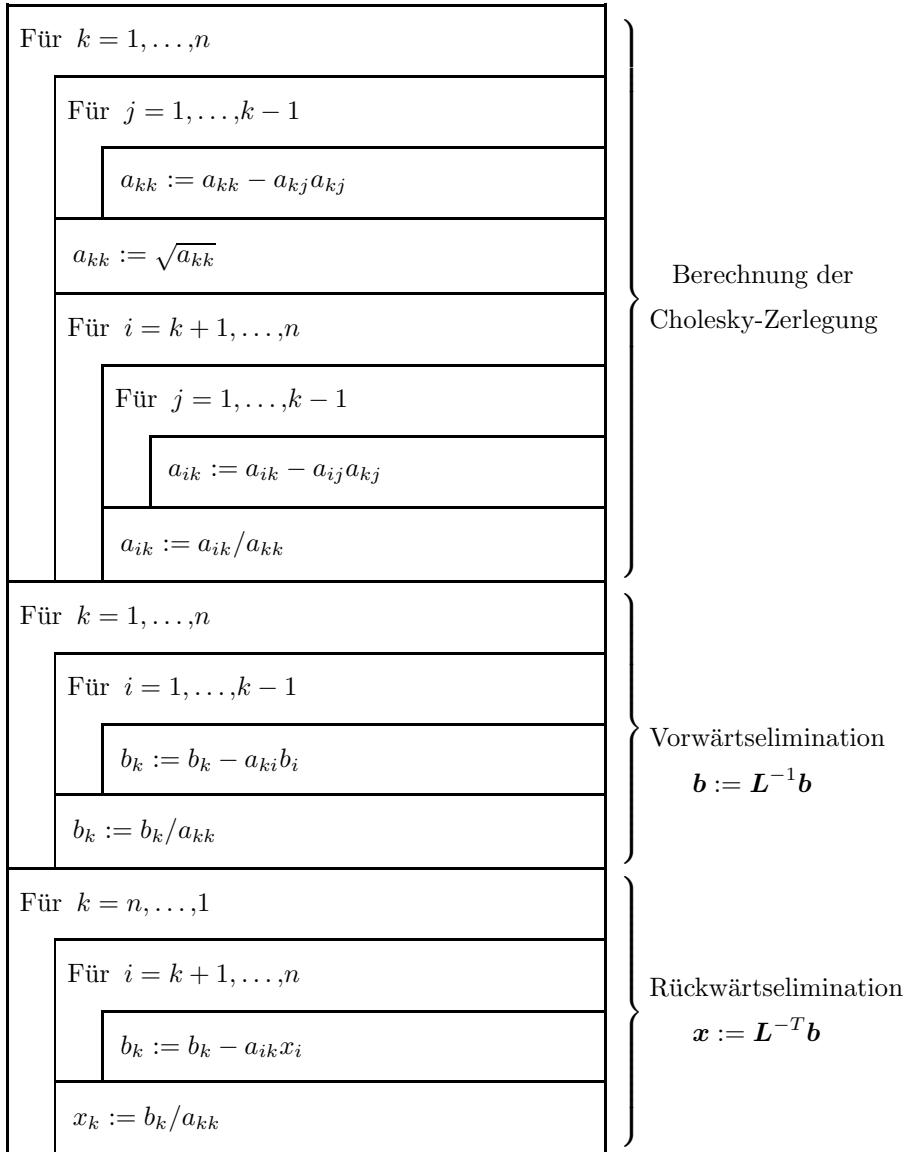
berechnen läßt. Aus

$$a_{ik} = \ell_{i1} \ell_{k1} + \dots + \ell_{ik} \ell_{kk} = \sum_{j=1}^k \ell_{ij} \ell_{kj} \quad \text{für } i = k+1, \dots, n$$

erhalten wir die Berechnungsvorschrift für die unterhalb der Diagonale befindlichen Elemente der k -ten Spalte in der Form

$$\ell_{ik} = \frac{1}{\ell_{kk}} \left(a_{ik} - \sum_{j=1}^{k-1} \ell_{ij} \ell_{kj} \right) \quad \text{für } i = k+1, \dots, n. \quad (3.2.3)$$

Algorithmus Cholesky-Zerlegung —



Betrachten wir für die Analyse der Rechenzeit lediglich den für große n entscheidenden ersten Teil, so erhalten wir

Multiplikationen + # Divisionen + # Wurzeln

$$\begin{aligned}
 &= \underbrace{\sum_{k=1}^n (k-1)}_{\text{Multiplikationen}} + \underbrace{\sum_{k=1}^n 1}_{\text{Wurzeln}} + \underbrace{\sum_{k=1}^n (n-k)(k-1)}_{\text{Multiplikationen}} + \underbrace{\sum_{k=1}^n (n-k)}_{\text{Divisionen}} \\
 &= \sum_{k=1}^n k + n \sum_{k=1}^n k - \sum_{k=1}^n k^2 \\
 &= \frac{n(n+1)}{2} + n \frac{n(n+1)}{2} - \frac{n(n+1)(2n+1)}{6} \\
 &= \frac{n^3}{6} + \frac{n^2}{2} + \frac{n}{3}.
 \end{aligned}$$

Für große n benötigt die Cholesky-Zerlegung damit nur etwa die Hälfte der aufwendigen Operationen des Gauß-Algorithmus.

3.3 QR-Zerlegung

Eine wesentliche Grundlage zur Definition des im Abschnitt 4.3.2.4 betrachteten GMRES-Verfahrens wie auch vieler Methoden zur Lösung von Eigenwertproblemen und linearen Ausgleichsproblemen stellt die QR-Zerlegung einer Matrix dar. Der Vorteil der QR-Zerlegung im Vergleich zur LR-Zerlegung liegt in der Normerhaltung der unitären Transformation bezüglich der euklidischen Norm. Zudem kann das Gleichungssystem $\mathbf{Ax} = \mathbf{b}$ wegen $\mathbf{Q}^* = \mathbf{Q}^{-1}$ mittels

$$\mathbf{Ax} = \mathbf{b} \Leftrightarrow \mathbf{QRx} = \mathbf{b} \Leftrightarrow \mathbf{Rx} = \mathbf{Q}^* \mathbf{b}$$

leicht gelöst werden. Die drei bekanntesten Methoden zur Berechnung dieser Zerlegung sind das Gram-Schmidt-Verfahren und die Algorithmen nach Givens und Householder. In den folgenden Kapiteln werden wir ausschließlich die beiden erstgenannten Methoden verwenden. Eine ausführliche Herleitung der QR-Zerlegung nach Householder findet man zum Beispiel in [15, 31, 53, 19].

Definition 3.12 Die Zerlegung einer Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$ in ein Produkt

$$\mathbf{A} = \mathbf{QR}$$

aus einer unitären Matrix $\mathbf{Q}$ und einer rechten oberen Dreiecksmatrix $\mathbf{R}$ heißt

QR-Zerlegung.

3.3.1 Das Gram-Schmidt-Verfahren

Die Herleitung des Gram-Schmidt-Verfahrens ergibt sich unmittelbar aus dem folgenden konstruktiven Existenznachweis der QR-Zerlegung.

Satz 3.13 (Existenz der QR-Zerlegung)

Sei $\mathbf{A} \in \mathbb{C}^{n \times n}$ eine reguläre Matrix, dann existieren eine unitäre Matrix $\mathbf{Q} \in \mathbb{C}^{n \times n}$ und eine rechte obere Dreiecksmatrix $\mathbf{R} \in \mathbb{C}^{n \times n}$ derart, daß

$$\mathbf{A} = \mathbf{Q}\mathbf{R}$$

gilt.

Beweis:

Wir führen den Beweis, indem wir sukzessive für $k = 1, \dots, n$ die Existenz von Vektoren $\mathbf{q}_1, \dots, \mathbf{q}_k \in \mathbb{C}^n$ mit

$$(\mathbf{q}_i, \mathbf{q}_j)_2 = \delta_{ij} \text{ für } i, j = 1, \dots, k \quad (3.3.1)$$

und

$$\text{span}\{\mathbf{q}_1, \dots, \mathbf{q}_k\} = \text{span}\{\mathbf{a}_1, \dots, \mathbf{a}_k\} \quad (3.3.2)$$

nachweisen, wobei $\mathbf{a}_j$ ($1 \leq j \leq k$) den j -ten Spaltenvektor der Matrix $\mathbf{A}$ darstellt.

Für $k = 1$ sind die beiden Bedingungen wegen $\mathbf{a}_1 \in \mathbb{C}^n \setminus \{\mathbf{0}\}$ mit

$$\mathbf{q}_1 = \frac{\mathbf{a}_1}{\|\mathbf{a}_1\|_2} \quad (3.3.3)$$

erfüllt.

Seien nun $\mathbf{q}_1, \dots, \mathbf{q}_k$ mit $k \in \{1, \dots, n-1\}$ gegeben, die die Bedingungen (3.3.1) und (3.3.2) erfüllen, dann läßt sich jeder Vektor

$$\mathbf{q}_{k+1} \in \text{span}\{\mathbf{a}_1, \dots, \mathbf{a}_{k+1}\} \setminus \text{span}\{\mathbf{q}_1, \dots, \mathbf{q}_k\} \quad (3.3.4)$$

in der Form

$$\mathbf{q}_{k+1} = c_{k+1} \left(\mathbf{a}_{k+1} - \sum_{i=1}^k c_i \mathbf{q}_i \right) \quad (3.3.5)$$

mit $c_{k+1} \neq 0$ schreiben. Motiviert durch

$$(\mathbf{q}_{k+1}, \mathbf{q}_j)_2 = c_{k+1} [(\mathbf{a}_{k+1}, \mathbf{q}_j)_2 - c_j] \text{ für } j = 1, \dots, k$$

und der Zielsetzung der Orthogonalität setzen wir in (3.3.5) $c_i := (\mathbf{a}_{k+1}, \mathbf{q}_i)_2$ für $i = 1, \dots, k$ und erhalten hierdurch die Gleichung

$$\begin{aligned} (\mathbf{q}_{k+1}, \mathbf{q}_j)_2 &= c_{k+1} \left[(\mathbf{a}_{k+1}, \mathbf{q}_j)_2 - \sum_{i=1}^k (\mathbf{a}_{k+1}, \mathbf{q}_i)_2 (\mathbf{q}_i, \mathbf{q}_j)_2 \right] \\ &= c_{k+1} [(\mathbf{a}_{k+1}, \mathbf{q}_j)_2 - (\mathbf{a}_{k+1}, \mathbf{q}_j)_2] = 0 \text{ für } j = 1, \dots, k. \end{aligned} \quad (3.3.6)$$

Da $\mathbf{A}$ regulär ist, gilt $\mathbf{a}_{k+1} \notin \text{span}\{\mathbf{q}_1, \dots, \mathbf{q}_k\}$, so daß

$$\tilde{\mathbf{q}}_{k+1} := \mathbf{a}_{k+1} - \sum_{i=1}^k (\mathbf{a}_{k+1}, \mathbf{q}_i)_2 \mathbf{q}_i \neq \mathbf{0}$$

folgt. Mit

$$c_{k+1} := \frac{1}{\|\tilde{\mathbf{q}}_{k+1}\|_2}$$

ergibt sich $\|\mathbf{q}_{k+1}\|_2 = 1$, so daß durch

$$\mathbf{q}_{k+1} = \frac{1}{\|\tilde{\mathbf{q}}_{k+1}\|_2} \tilde{\mathbf{q}}_{k+1} = \frac{1}{\|\tilde{\mathbf{q}}_{k+1}\|_2} \left(\mathbf{a}_{k+1} - \sum_{i=1}^k (\mathbf{a}_{k+1}, \mathbf{q}_i)_2 \mathbf{q}_i \right) \quad (3.3.7)$$

wegen (3.3.6) der gesuchte Vektor vorliegt. Die Definition

$$\mathbf{Q} = (\mathbf{q}_1 \dots \mathbf{q}_n)$$

mit $\mathbf{q}_1, \dots, \mathbf{q}_n$ gemäß (3.3.3) respektive (3.3.7) und

$$\mathbf{R} := \begin{pmatrix} r_{11} & \cdots & r_{1n} \\ & \ddots & \vdots \\ & & r_{nn} \end{pmatrix} \quad \text{mit} \quad r_{ik} = \begin{cases} \|\tilde{\mathbf{q}}_i\|_2 & \text{für } k = i, \\ (\mathbf{a}_k, \mathbf{q}_i)_2 & \text{für } k > i \end{cases}$$

liefert somit

$$\mathbf{A} = \mathbf{Q}\mathbf{R}.$$

□

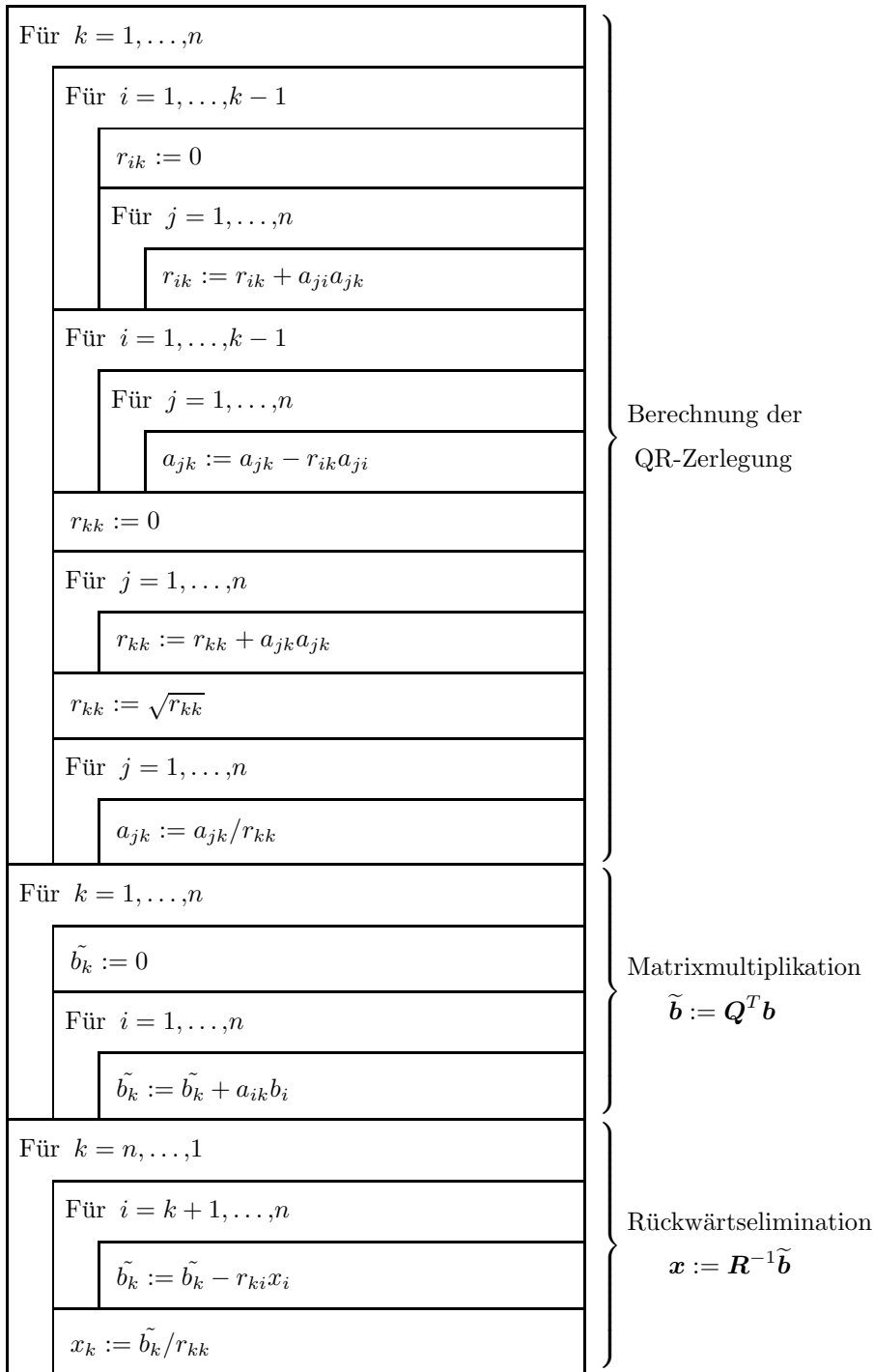
Die Lösung des Gleichungssystems $\mathbf{A}\mathbf{x} = \mathbf{b}$ erhalten wir durch Rückwärtsauflösen des Gleichungssystems $\mathbf{R}\mathbf{x} = \tilde{\mathbf{b}} = \mathbf{Q}^*\mathbf{b}$.

Seien die Spalten der Matrix $\mathbf{A} \in \mathbb{C}^{m \times n}$, $m \geq n$ linear unabhängig, dann folgt aus dem obigen Beweis die Existenz einer Matrix $\tilde{\mathbf{Q}} \in \mathbb{C}^{m \times n}$, deren Spalten paarweise orthonormal sind, und einer regulären rechten oberen Dreiecksmatrix $\mathbf{R} \in \mathbb{C}^{n \times n}$ mit $\mathbf{A} = \tilde{\mathbf{Q}}\mathbf{R}$. Erweitern wir $\tilde{\mathbf{Q}}$ zu einer unitären Matrix $\mathbf{Q} \in \mathbb{C}^{m \times m}$, so folgt die Darstellung

$$\mathbf{A} = \mathbf{Q} \begin{pmatrix} \mathbf{R} \\ \mathbf{0} \end{pmatrix}.$$

Zusammenfassend erhalten wir aus dem Beweis des obigen Satzes das Gram-Schmidt-Verfahren für eine reguläre Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ in der folgenden Form:

Algorithmus Gram-Schmidt —



Bei der Aufwandsanalyse der QR-Zerlegung nach Gram-Schmidt betrachten wir lediglich den aufwendigen ersten Teil und vernachlässigen auch hier das explizite Lösen des Gleichungssystems. Somit erhalten wir

Multiplikationen + # Divisionen + # Wurzeln

$$\begin{aligned}
 &= \underbrace{2n \sum_{k=1}^n (k-1)}_{\text{Multiplikationen}} + \underbrace{n \sum_{k=1}^n 1}_{\text{Divisionen}} + \underbrace{n \sum_{k=1}^n 1}_{\text{Wurzeln}} + \sum_{k=1}^n 1 \\
 &= 2n \sum_{k=1}^n k + n = 2 \frac{n^2(n+1)}{2} + n = n^3 + n^2 + n.
 \end{aligned}$$

Für große n ist der Aufwand somit ungefähr dreimal so hoch wie beim Gauß-Algorithmus.

Die numerische Problematik des Gram-Schmidt-Verfahrens liegt im Auftreten von Rundungsfehlern, die zu Auslöschungseffekten und folglich zu einer ungenügenden Orthogonalität der Vektoren $\mathbf{q}_1, \dots, \mathbf{q}_n$ führen können. Zur Verbesserung betrachten wir zwei Varianten des Verfahrens.

Variante I: Das modifizierte Gram-Schmidt-Verfahren

Grundidee: Nach der Berechnung des k -ten Spaltenvektors $\mathbf{q}_k$ der unitären Matrix $\mathbf{Q}$ werden die Vektoren $\mathbf{a}_{k+1}, \dots, \mathbf{a}_n$ der Matrix $\mathbf{A}$ derart modifiziert, daß sie senkrecht auf allen Spaltenvektoren $\mathbf{q}_1, \dots, \mathbf{q}_k$ stehen.

Durchführung: Sei $\mathbf{A} = (\mathbf{a}_1 \dots \mathbf{a}_n) \in \mathbb{R}^{n \times n}$

1. Schritt: Setze $\mathbf{a}_k^{(1)} = \mathbf{a}_k$ für $k = 1, \dots, n$.

2. Schritt: Für $k = 1, \dots, n$

$$\begin{aligned}
 \mathbf{q}_k &:= \frac{\mathbf{a}_k^{(k)}}{\|\mathbf{a}_k^{(k)}\|_2} \\
 \mathbf{a}_j^{(k+1)} &= \mathbf{a}_j^{(k)} - (\mathbf{a}_j^{(k)}, \mathbf{q}_k)_2 \mathbf{q}_k \text{ für } j = k+1, \dots, n.
 \end{aligned}$$

Der Rechenaufwand des modifizierten Gram-Schmidt-Verfahrens ist identisch zum ursprünglichen Algorithmus. Vergleichen wir beide Methoden, so kann sich eine Veränderung frühestens für $n = 3$ ergeben.

Für $\mathbf{A} = (\mathbf{a}_1, \mathbf{a}_2, \mathbf{a}_3) = (\mathbf{a}_1^{(1)}, \mathbf{a}_2^{(1)}, \mathbf{a}_3^{(1)}) \in \mathbb{R}^{3 \times 3}$ folgt:

Gram-Schmidt-Verfahren	Modifiziertes Gram-Schmidt-Verfahren
$\mathbf{q}_1 = \frac{\mathbf{a}_1}{\ \mathbf{a}_1\ _2}$	$\mathbf{q}_1 = \frac{\mathbf{a}_1^{(1)}}{\ \mathbf{a}_1^{(1)}\ _2}$ $\mathbf{a}_2^{(2)} = \mathbf{a}_2^{(1)} - \left(\mathbf{a}_2^{(1)}, \mathbf{q}_1\right)_2 \mathbf{q}_1$ $\mathbf{a}_3^{(2)} = \mathbf{a}_3^{(1)} - \left(\mathbf{a}_3^{(1)}, \mathbf{q}_1\right)_2 \mathbf{q}_1$
$\tilde{\mathbf{a}}_2 = \mathbf{a}_2 - (\mathbf{a}_2, \mathbf{q}_1)_2 \mathbf{q}_1$ $\mathbf{q}_2 = \frac{\tilde{\mathbf{a}}_2}{\ \tilde{\mathbf{a}}_2\ _2}$	$\mathbf{q}_2 = \frac{\mathbf{a}_2^{(2)}}{\ \mathbf{a}_2^{(2)}\ _2}$ $\mathbf{a}_3^{(3)} = \mathbf{a}_3^{(2)} - \left(\mathbf{a}_3^{(2)}, \mathbf{q}_2\right)_2 \mathbf{q}_2$
$\tilde{\mathbf{a}}_3 = \mathbf{a}_3 - (\mathbf{a}_3, \mathbf{q}_1)_2 \mathbf{q}_1$ $\quad - (\mathbf{a}_3, \mathbf{q}_2)_2 \mathbf{q}_2$ $\mathbf{q}_3 = \frac{\tilde{\mathbf{a}}_3}{\ \tilde{\mathbf{a}}_3\ _2}$	$\mathbf{q}_3 = \frac{\mathbf{a}_3^{(3)}}{\ \mathbf{a}_3^{(3)}\ _2}$

Hieraus ergibt sich der Zusammenhang

$$\begin{aligned}
 \mathbf{a}_3^{(3)} &= \mathbf{a}_3^{(1)} - \left(\mathbf{a}_3^{(1)}, \mathbf{q}_1\right)_2 \mathbf{q}_1 - \left(\mathbf{a}_3^{(1)}, \mathbf{q}_2\right)_2 \mathbf{q}_2 + \left(\left(\mathbf{a}_3^{(1)}, \mathbf{q}_1\right)_2 \mathbf{q}_1, \mathbf{q}_2\right)_2 \mathbf{q}_2 \\
 &= \tilde{\mathbf{a}}_3 + ((\mathbf{a}_3, \mathbf{q}_1)_2 \mathbf{q}_1, \mathbf{q}_2)_2 \mathbf{q}_2,
 \end{aligned}$$

so daß der Unterschied im Vektor $((\mathbf{a}_3, \mathbf{q}_1)_2 \mathbf{q}_1, \mathbf{q}_2)_2 \mathbf{q}_2$ liegt, der bei exakter Arithmetik verschwindet.

Variante II: Gram-Schmidt-Verfahren mit Nachorthogonalisierung

Grundidee: Nach der Berechnung des unnormierten k -ten Spaltenvektors $\mathbf{q}_k$ wird dieser bezüglich $\mathbf{q}_1, \dots, \mathbf{q}_{k-1}$ orthogonalisiert und anschließend normiert.

Durchführung: Sei $\tilde{\mathbf{q}}_k$ der unnormierte Vektor, dann setze

$$\begin{aligned}\bar{\mathbf{q}}_k &= \tilde{\mathbf{q}}_k - \sum_{i=1}^{k-1} (\tilde{\mathbf{q}}_k, \mathbf{q}_i)_2 \mathbf{q}_i \\ \mathbf{q}_k &= \frac{\bar{\mathbf{q}}_k}{\|\bar{\mathbf{q}}_k\|_2}.\end{aligned}$$

Zusätzlich müssen die Koeffizienten der Matrix $\mathbf{R}$ durch

$$r_{i,k}^{\text{neu}} = r_{i,k}^{\text{alt}} + (\tilde{\mathbf{q}}_k, \mathbf{q}_i)_2 \text{ für } i = 1, \dots, k-1$$

modifiziert werden.

Der Rechenaufwand ist im Vergleich zum ursprünglichen Gram-Schmidt-Verfahren etwa doppelt so hoch.

Satz 3.14 (Eindeutigkeit der QR-Zerlegung)

Sei $\mathbf{A} \in \mathbb{C}^{n \times n}$ regulär, dann existiert zu je zwei QR-Zerlegungen

$$\mathbf{Q}_1 \mathbf{R}_1 = \mathbf{A} = \mathbf{Q}_2 \mathbf{R}_2 \quad (3.3.8)$$

eine unitäre Diagonalmatrix $\mathbf{D} \in \mathbb{C}^{n \times n}$ mit

$$\mathbf{Q}_1 = \mathbf{Q}_2 \mathbf{D} \quad \text{und} \quad \mathbf{R}_2 = \mathbf{D} \mathbf{R}_1.$$

Beweis:

Mit $\mathbf{D} = \mathbf{Q}_2^* \mathbf{Q}_1$ liegt wegen Lemma 2.27 eine unitäre Matrix vor, und es gilt

$$\mathbf{Q}_1 = \mathbf{Q}_2 \mathbf{D}.$$

Da $\mathbf{A}$ regulär ist, sind auch $\mathbf{R}_1$ und $\mathbf{R}_2$ regulär, und wir erhalten mit (3.3.8)

$$\mathbf{D} = \mathbf{Q}_2^* \mathbf{Q}_1 = \mathbf{R}_2 \mathbf{R}_1^{-1}.$$

Lemma 3.4 in Kombination mit Lemma 2.26 und Lemma 2.27 besagt, daß $\mathbf{D}$ eine rechte obere Dreiecksmatrix darstellt. Da $\mathbf{D}$ zudem unitär ist, stellt $\mathbf{D}$ sogar eine Diagonalmatrix dar. $\square$

Korollar 3.15 Sei $\mathbf{A} \in \mathbb{C}^{n \times n}$ regulär, dann existiert genau eine QR-Zerlegung der Matrix $\mathbf{A}$ derart, daß die Diagonalelemente der Matrix $\mathbf{R}$ reell und positiv sind.

3.3.2 Die QR-Zerlegung nach Givens

Wir beschränken uns auf den Fall einer Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$. Die Idee der Givens-Methode liegt in einer sukzessiven Elimination der Unterdiagonalelemente. Beginnend mit der ersten Spalte werden hierzu die Subdiagonalelemente jeder Spalte in aufsteigender Reihenfolge mittels orthogonaler Drehmatrizen annulliert.

Gehen wir zum Beispiel von der Matrix

$$\mathbf{A} = \begin{pmatrix} * & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & * \\ 0 & \ddots & & & & & & & & & \vdots \\ \vdots & \ddots & \ddots & & & & & & & & \vdots \\ \vdots & & 0 & * & & & & & & & \vdots \\ \vdots & & 0 & 0 & * & & & & & & \vdots \\ \vdots & & \vdots & \vdots & * & \ddots & & & & & \vdots \\ \vdots & & \vdots & \vdots & \vdots & \ddots & \ddots & & & & \vdots \\ \vdots & & \vdots & 0 & \vdots & & \ddots & \ddots & & & \vdots \\ \vdots & & \vdots & * & \vdots & & & \ddots & \ddots & & \vdots \\ \vdots & & \vdots & \vdots & \vdots & & & & \ddots & & \vdots \\ 0 & \dots & 0 & * & * & \dots & \dots & \dots & \dots & \dots & * \end{pmatrix} \in \mathbb{R}^{n \times n},$$

$\uparrow$
 i -te Spalte

$\leftarrow j$ -te Zeile

aus, das heißt, es gilt

$$a_{k\ell} = 0 \quad \forall \ell \in \{1, \dots, i-1\} \text{ mit } \ell < k \in \{1, \dots, n\}, \quad (3.3.9)$$

$$a_{i+1,i} = \dots = a_{j-1,i} = 0 \quad (3.3.10)$$

und $a_{ji} \neq 0$. Dann suchen wir eine orthogonale Matrix

$$\mathbf{G}_{ji} = \begin{pmatrix} 1 & & & & & & & & & & \\ & \ddots & & & & & & & & & \\ & & 1 & & & & & & & & \\ & & & g_{ii} & 0 & \dots & 0 & g_{ij} & & & \\ & & & 0 & 1 & & & 0 & & & \\ & & & \vdots & & \ddots & & \vdots & & & \\ & & & 0 & & & 1 & 0 & & & \\ & & & g_{ji} & 0 & \dots & 0 & g_{jj} & & & \\ & & & & & & & & 1 & & \\ & & & & & & & & & \ddots & \\ & & & & & & & & & & 1 \end{pmatrix} \in \mathbb{R}^{n \times n}$$

derart, daß für

$$\tilde{\mathbf{A}} = \mathbf{G}_{ji} \mathbf{A}$$

neben

$$\tilde{a}_{k\ell} = 0 \quad \forall \ell \in \{1, \dots, i-1\} \text{ mit } \ell < k \in \{1, \dots, n\}, \quad (3.3.11)$$

und

$$\tilde{a}_{i+1,i} = \dots = \tilde{a}_{j-1,i} = 0 \quad (3.3.12)$$

auch

$$\tilde{a}_{ji} = 0 \quad (3.3.13)$$

gilt.

Zunächst unterscheidet sich $\tilde{\mathbf{A}}$ von $\mathbf{A}$ lediglich in der i -ten und j -ten Zeile, und es gilt für $\ell = 1, \dots, n$

$$\tilde{a}_{i\ell} = g_{ii}a_{i\ell} + g_{ij}a_{j\ell}$$

$$\tilde{a}_{j\ell} = g_{ji}a_{i\ell} + g_{jj}a_{j\ell}.$$

Mit (3.3.9) folgt $a_{i\ell} = a_{j\ell} = 0$ für $\ell < i < j$, so daß

$$\tilde{a}_{i\ell} = \tilde{a}_{j\ell} = 0 \quad \text{für} \quad \ell = 1, \dots, i-1$$

gilt und folglich die Forderungen (3.3.11) und (3.3.12) erfüllt sind. Wohldefiniert durch $a_{ji} \neq 0$ setzen wir

$$g_{ii} = g_{jj} = \frac{a_{ii}}{\sqrt{a_{ii}^2 + a_{ji}^2}}$$

und

$$g_{ij} = -g_{ji} = \frac{a_{ji}}{\sqrt{a_{ii}^2 + a_{ji}^2}}.$$

Somit stellt $\mathbf{G}_{ji}$ eine orthogonale Drehmatrix um den Winkel $\alpha = \arccos g_{ii}$ dar, und es gilt

$$\tilde{a}_{ji} = -\frac{a_{ji}}{\sqrt{a_{ii}^2 + a_{ji}^2}}a_{ii} + \frac{a_{ii}}{\sqrt{a_{ii}^2 + a_{ji}^2}}a_{ji} = 0.$$

Definieren wir $\mathbf{G}_{ji} = \mathbf{I}$ im Fall einer Matrix $\mathbf{A}$, die (3.3.9) und (3.3.10) genügt und zudem $a_{ji} = 0$ beinhaltet, dann haben wir mit

$$\tilde{\mathbf{Q}} := \prod_{i=n-1}^1 \prod_{j=n}^{i+1} \mathbf{G}_{ji} := \mathbf{G}_{n,n-1} \cdot \dots \cdot \mathbf{G}_{3,2} \cdot \mathbf{G}_{n,1} \cdot \dots \cdot \mathbf{G}_{3,1} \cdot \mathbf{G}_{2,1}$$

eine orthogonale Matrix, für die

$$\mathbf{R} = \tilde{\mathbf{Q}}\mathbf{A}$$

eine obere Dreiecksmatrix ist. Mit $\mathbf{Q} = \tilde{\mathbf{Q}}^T$ folgt

$$\mathbf{A} = \mathbf{Q}\mathbf{R}.$$

Bei der Givens-Methode haben wir die Möglichkeit, das Gleichungssystem $\mathbf{A}\mathbf{x} = \mathbf{b}$ mittels eines QR-Verfahrens ohne explizite Abspeicherung der orthogonalen Matrix zu lösen.

Wir ergänzen die Matrix $\mathbf{A}$ um die rechte Seite $\mathbf{b}$ gemäß $\mathbf{a}_{n+1} = \mathbf{b}$ und erhalten:

Algorithmus Givens-Methode —

Für $i = 1, \dots, n-1$		
Für $j = i+1, \dots, n$		
Y	$a_{ji} = 0$	N
$t := 1/\sqrt{a_{ii}^2 + a_{ji}^2}$		
$s := ta_{ji}$		
$c := ta_{ii}$		
Für $k = i, \dots, n+1$		
$t := ca_{ik} + sa_{jk}$		
Y	$k = i$	N
		$a_{jk} := -sa_{ik} + ca_{jk}$
$a_{ik} := t$		
$a_{ji} := 0$		
Für $i = n, \dots, 1$		
Für $j = i+1, \dots, n$		
$a_{i,n+1} := a_{i,n+1} - a_{ij}x_j$		
$x_i := a_{i,n+1}/a_{ii}$		

Bei der Givens-Methode ergibt sich ohne Berücksichtigung der rechten Seite $\mathbf{b}$ und des expliziten Auflörens des Gleichungssystems durch Rückwärtseinsetzen der folgende Rechenaufwand für eine vollbesetzte Matrix:

Multiplikationen + # Divisionen + # Wurzeln

$$\begin{aligned}
 &\leq \underbrace{\sum_{i=1}^{n-1} \{2(n-i) + 4(n-i)(n-i+1)\}}_{\text{Multiplikationen}} + \underbrace{\sum_{i=1}^{n-1} (n-i)}_{\text{Divisionen}} + \underbrace{\sum_{i=1}^{n-1} (n-i)}_{\text{Wurzeln}} \\
 &= \frac{4}{3}n^3 + 2n^2 - \frac{10}{3}n.
 \end{aligned}$$

3.4 Übungsaufgaben

Aufgabe 1:

Man berechne die LR-Zerlegung der Matrix

$$\mathbf{A} = \begin{pmatrix} 1 & 4 & 5 \\ 1 & 6 & 11 \\ 2 & 14 & 31 \end{pmatrix}$$

und löse hiermit das lineare Gleichungssystem $\mathbf{Ax} = (17, 31, 82)^T$.

Aufgabe 2:

Gegeben ist das lineare Gleichungssystem $\mathbf{Ax} = \mathbf{b}$ mit

$$\mathbf{A} = \begin{pmatrix} 4 & \alpha \\ 6 & 15 \end{pmatrix}, \quad \mathbf{b} = \begin{pmatrix} 8 \\ \beta \end{pmatrix}.$$

- (a) Für welche Werte von α und β besitzt dieses Gleichungssystem
 - (1) eine eindeutige Lösung,
 - (2) keine Lösung,
 - (3) unendlich viele Lösungen?
- (b) Für den Fall, dass $\mathbf{Ax} = \mathbf{b}$ eindeutig lösbar ist, gebe man die Lösung in Abhängigkeit von den Parametern α, β an.

Aufgabe 3:

Beweisen Sie, dass zu jeder symmetrischen, positiv definiten Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ genau eine untere Dreiecksmatrix $\mathbf{L} \in \mathbb{R}^{n \times n}$ mit positiven Diagonalelementen existiert, für welche $\mathbf{A} = \mathbf{LL}^T$ gilt.

Aufgabe 4:

Man berechne die QR-Zerlegung der Matrix

$$\mathbf{A} = \begin{pmatrix} 1 & 3 \\ -1 & 1 \end{pmatrix}$$

und löse hiermit das lineare Gleichungssystem $\mathbf{Ax} = (16, 0)^T$.

Aufgabe 5:

Gegeben sei die Matrix

$$\mathbf{A} = \begin{pmatrix} 1 & 3 & 4 & 1 \\ 2 & 7 & a & 4 \\ 1 & 4 & 6 & 1 \\ 3 & 4 & 9 & 0 \end{pmatrix}.$$

- (a) Für welchen Wert von a besitzt $\mathbf{A}$ keine LR-Zerlegung.
- (b) Berechnen Sie eine LR-Zerlegung von $\mathbf{A}$ im Existenzfall.
- (c) Für den Fall, dass $\mathbf{A}$ keine LR-Zerlegung besitzt, geben Sie eine Permutationsmatrix $\mathbf{P}$ derart an, dass $\mathbf{PA}$ eine LR-Zerlegung besitzt.

Aufgabe 6:

Gegeben sei eine reguläre Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ mit Bandbreite d , das heißt es gilt $\mathbf{A} = (a_{ij})_{i,j=1,\dots,n}$ mit $a_{ij} = 0$ für $|i - j| \geq d$. Zudem seien alle Hauptabschnittsdeterminanten von $\mathbf{A}$ ungleich Null.

Zeigen Sie: Bei der LR-Zerlegung von $\mathbf{A}$ bleibt die Bandstruktur erhalten, d.h. $\mathbf{L}$ und $\mathbf{R}$ besitzen die Bandbreite d .

Aufgabe 7:

Gegeben sei die Tridiagonalmatrix

$$\mathbf{A}_n = \begin{pmatrix} a_1 & c_2 & & & \\ b_1 & a_2 & \ddots & & \\ & \ddots & \ddots & \ddots & \\ & & b_{n-1} & a_n & c_n \end{pmatrix}, \quad c_1 := 0, \quad b_n := 0.$$

Zeigen Sie, dass unter den Voraussetzungen

$$|b_j| + |c_j| \leq |a_j| \quad \text{und} \quad |b_j| < |a_j|, \quad j = 1, \dots, n$$

eine LR-Zerlegung von $\mathbf{A}_n$ in der Form

$$\mathbf{L}_n = \begin{pmatrix} 1 & & & & \\ \alpha_1 & \ddots & & & \\ & \ddots & \ddots & & \\ & & \ddots & \ddots & \\ & & & \alpha_{n-1} & 1 \end{pmatrix} \quad \text{und} \quad \mathbf{R}_n = \begin{pmatrix} \beta_1 & c_2 & & & \\ & \ddots & \ddots & & \\ & & \ddots & \ddots & \\ & & & \ddots & c_n \\ & & & & \beta_n \end{pmatrix}$$

existiert. Geben Sie eine Rekursionsformel zur Berechnung der α_j und β_j an.

Hinweis: Zeigen Sie induktiv: $|\alpha_j| < 1$, $j = 1, \dots, n-1$ und $|\beta_j| > 0$, $j = 1, \dots, n$.

Aufgabe 8:

Es sei $\mathbf{A} \in \mathbb{C}^{n \times n}$ regulär und es gelte $\mathbf{A} = \mathbf{Q}\mathbf{R}$, wobei $\mathbf{Q} \in \mathbb{C}^{n \times n}$ unitär ist und $\mathbf{R} \in \mathbb{C}^{n \times n}$ eine rechte obere Dreiecksmatrix darstellt. Es bezeichne $\mathbf{a}_j \in \mathbb{C}^n$ bzw. $\mathbf{q}_j \in \mathbb{C}^n$ den j -ten Spaltenvektor von $\mathbf{A}$ bzw. $\mathbf{Q}$, $j = 1, \dots, n$. Zeigen Sie:

$$\text{span}\{\mathbf{a}_1, \dots, \mathbf{a}_j\} = \text{span}\{\mathbf{q}_1, \dots, \mathbf{q}_j\} \quad \text{für alle } j = 1, \dots, n.$$

Aufgabe 9:

Berechnen Sie eine Cholesky-Zerlegung der Matrix

$$\mathbf{A} = \begin{pmatrix} 9 & 3 & 9 \\ 3 & 9 & 11 \\ 9 & 11 & 17 \end{pmatrix}$$

und lösen Sie hiermit das lineare Gleichungssystem

$$\mathbf{A}\mathbf{x} = \begin{pmatrix} 24 \\ 16 \\ 32 \end{pmatrix}.$$

Aufgabe 10:

(a) Gegeben sei die Matrix

$$\mathbf{A} = \begin{pmatrix} \mathbf{A}_{11} & \mathbf{A}_{12} \\ \mathbf{0} & \mathbf{A}_{22} \end{pmatrix} \in \mathbb{R}^{n \times n}$$

mit $\mathbf{A}_{ii} \in \mathbb{R}^{m_i \times m_i}$, $i = 1, 2$, wobei $m_1 + m_2 = n$ und $\det \mathbf{A}_{ii} \neq 0$, $i = 1, 2$ gilt. Man zeige, dass $\mathbf{A}$ invertierbar ist und bestimme eine Blockdarstellung der Inversen $\mathbf{A}^{-1}$.

(b) Unter Benutzung geeigneter Blockbildungen berechne man die Inverse von

$$\mathbf{B} = \begin{pmatrix} 1 & 2 & 3 & 1 & 0 \\ 2 & 5 & 7 & 1 & 1 \\ 0 & 1 & 2 & 1 & 1 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}.$$

Aufgabe 11:

Gegeben seien eine symmetrische Matrix $\mathbf{A}_{n-1} \in \mathbb{R}^{(n-1) \times (n-1)}$ und eine untere Dreiecksmatrix $\mathbf{L}_{n-1} \in \mathbb{R}^{(n-1) \times (n-1)}$ mit positiven Diagonalelementen, die $\mathbf{L}_{n-1} \mathbf{L}_{n-1}^T = \mathbf{A}_{n-1}$ erfülle. Weiter seien $\mathbf{b} \in \mathbb{R}^{n-1}$ ein Spaltenvektor, $\alpha \in \mathbb{R}$ und

$$\mathbf{A}_n := \begin{pmatrix} \mathbf{A}_{n-1} & \mathbf{b} \\ \mathbf{b}^T & \alpha \end{pmatrix} \in \mathbb{R}^{n \times n}$$

symmetrisch und positiv definit. Zeigen Sie, dass ein Vektor $\mathbf{c} \in \mathbb{R}^{n-1}$ und eine positive reelle Zahl β existieren, so dass

$$\mathbf{A}_n = \begin{pmatrix} \mathbf{L}_{n-1} & 0 \\ \mathbf{c}^T & \beta \end{pmatrix} \begin{pmatrix} \mathbf{L}_{n-1}^T & \mathbf{c} \\ 0 & \beta \end{pmatrix}$$

gilt.

Aufgabe 12:

Gegeben sei die Matrix

$$\mathbf{A} = \begin{pmatrix} 1 & 4 & 7 \\ 2 & \alpha & \beta \\ 0 & 1 & 1 \end{pmatrix}.$$

Unter welchen Voraussetzungen an die Werte $\alpha, \beta \in \mathbb{R}$ ist die Matrix regulär und besitzt zudem eine LR-Zerlegung. Geben Sie zudem ein Parameterpaar (α, β) derart an, dass die Matrix $\mathbf{A}$ regulär ist und keine LR-Zerlegung besitzt.

4 Iterative Verfahren

Die präsentierten direkten Verfahren stellen bei einer kleinen Anzahl von Unbekannten oftmals eine effiziente Vorgehensweise dar. Praxisrelevante Problemstellungen (siehe Beispiele 1.1 und 1.4) führen jedoch häufig auf große schwachbesetzte Gleichungssysteme. Die Speicherung derartiger Gleichungssysteme wird gewöhnlich erst durch die Vernachlässigung der Nullelemente der Matrix, die teilweise über 99% der Matrixkoeffizienten darstellen, ermöglicht. Betrachten wir die erläuterte Diskretisierung der Konvektions-Diffusions-Gleichung mit $N = 500$, so benötigen wir bei einem vorausgesetzten Speicherplatzbedarf von 8 Byte für jede reelle Zahl etwa 10 Megabyte zur Speicherung der nichtverschwindenden Matrixkoeffizienten. Dagegen würde das Abspeichern der gesamten Matrix 500 Gigabyte beanspruchen.

Direkte Verfahren können in der Regel die besondere Gestalt der Gleichungssysteme nicht ausnutzen, wodurch vollbesetzte Zwischenmatrizen generiert werden, die einerseits den verfügbaren Speicherplatz überschreiten und andererseits zu unakzeptablen Rechenzeiten führen (siehe Tabelle 3.1). Desweiteren entstehen solche Gleichungssysteme zumeist durch eine Diskretisierung der zugrundeliegenden Aufgabenstellung, wodurch auch die exakte Lösung des Gleichungssystems nur eine Approximation an die gesuchte Lösung darstellt. Folglich erweist sich eine Näherungslösung für das Gleichungssystem mit einem Fehler in der Größenordnung des Diskretisierungsfehlers als ausreichend. Hierzu eignen sich iterative Methoden hervorragend.

Wir betrachten ein lineares Gleichungssystem der Form

$$\mathbf{A}\mathbf{x} = \mathbf{b} \quad (4.0.1)$$

mit gegebener rechter Seite $\mathbf{b} \in \mathbb{C}^n$ und regulärer Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$. Iterative Verfahren ermitteln sukzessive Näherungen $\mathbf{x}_m$ an die exakte Lösung $\mathbf{A}^{-1}\mathbf{b}$ durch wiederholtes Ausführen einer festgelegten Rechenvorschrift

$$\mathbf{x}_{m+1} = \phi(\mathbf{x}_m, \mathbf{b}) \text{ für } m = 0, 1, \dots$$

bei gewähltem Startvektor $\mathbf{x}_0 \in \mathbb{C}^n$.

Bevor wir uns mit speziellen numerischen Methoden beschäftigen werden, wollen wir in diesem Abschnitt zunächst einige zweckdienliche Eigenschaften iterativer Verfahren beschreiben.

Definition 4.1 Eine Iterationsverfahren ist gegeben durch eine Abbildung

$$\phi : \mathbb{C}^n \times \mathbb{C}^n \rightarrow \mathbb{C}^n$$

und heißt linear, falls Matrizen $\mathbf{M}, \mathbf{N} \in \mathbb{C}^{n \times n}$ derart existieren, daß

$$\phi(\mathbf{x}, \mathbf{b}) = \mathbf{M}\mathbf{x} + \mathbf{N}\mathbf{b}$$

gilt. Die Matrix $\mathbf{M}$ wird als Iterationsmatrix der Iteration ϕ bezeichnet.

Die Matrizen $\mathbf{M}$ und $\mathbf{N}$ werden hierbei eindeutig durch die Iterationsvorschrift festgelegt.

Definition 4.2 Einen Vektor $\tilde{\mathbf{x}} \in \mathbb{C}^n$ bezeichnen wir als Fixpunkt des Iterationsverfahrens $\phi : \mathbb{C}^n \times \mathbb{C}^n \rightarrow \mathbb{C}^n$ zu $\mathbf{b} \in \mathbb{C}^n$, falls

$$\tilde{\mathbf{x}} = \phi(\tilde{\mathbf{x}}, \mathbf{b})$$

gilt.

Definition 4.3 Ein Iterationsverfahren ϕ heißt konsistent zur Matrix $\mathbf{A}$, wenn für alle $\mathbf{b} \in \mathbb{C}^n$ die Lösung $\mathbf{A}^{-1}\mathbf{b}$ ein Fixpunkt von ϕ zu $\mathbf{b}$ ist. Ein Iterationsverfahren ϕ heißt konvergent, wenn für alle $\mathbf{b} \in \mathbb{C}^n$ und alle Startwerte $\mathbf{x}_0 \in \mathbb{C}^n$ ein vom Startwert unabhängiger Grenzwert

$$\hat{\mathbf{x}} = \lim_{m \rightarrow \infty} \mathbf{x}_m = \lim_{m \rightarrow \infty} \phi(\mathbf{x}_{m-1}, \mathbf{b})$$

existiert.

Die Konsistenz stellt eine notwendige Bedingung an jedes Iterationsverfahren dar, da mit ihr ein sinnvoller Zusammenhang zwischen der numerischen Methode und dem Gleichungssystem sichergestellt wird. Bei einem inkonsistenten Algorithmus müßte der Anwender die Iteration in einem geeigneten Moment abbrechen, da die exakte Lösung keinen stationären Punkt der Iteration darstellt und sich die Folge der Näherungslösungen nach Erreichen des Lösungsvektors notwendigerweise wieder von diesem entfernen wird. Bei einem linearen Iterationsverfahren kann die Konsistenz der Methode unmittelbar anhand der verwendeten Matrizen $\mathbf{M}$ und $\mathbf{N}$ bestimmt werden. Diese wesentliche Eigenschaft wird durch den folgenden Satz belegt.

Satz 4.4 Ein lineares Iterationsverfahren ist genau dann konsistent zur Matrix $\mathbf{A}$, wenn

$$\mathbf{M} = \mathbf{I} - \mathbf{N}\mathbf{A}$$

gilt.

Beweis:

Sei $\tilde{\mathbf{x}} = \mathbf{A}^{-1}\mathbf{b}$.

„ $\Rightarrow$ “ ϕ sei konsistent zur Matrix $\mathbf{A}$.

Damit erhalten wir

$$\tilde{\mathbf{x}} = \phi(\tilde{\mathbf{x}}, \mathbf{b}) = \mathbf{M}\tilde{\mathbf{x}} + \mathbf{N}\mathbf{b} = \mathbf{M}\tilde{\mathbf{x}} + \mathbf{N}\mathbf{A}\tilde{\mathbf{x}}.$$

Da die Konsistenz für alle $\mathbf{b} \in \mathbb{C}^n$ gilt, ergibt sich unter Berücksichtigung der Regularität der Matrix $\mathbf{A}$ die Gültigkeit der obigen Gleichung für alle $\tilde{\mathbf{x}} \in \mathbb{C}^n$, wodurch

$$\mathbf{M} = \mathbf{I} - \mathbf{N}\mathbf{A}$$

folgt.

„ $\Leftarrow$ “ Es gelte $\mathbf{M} = \mathbf{I} - \mathbf{N}\mathbf{A}$.

Dann ergibt sich

$$\tilde{\mathbf{x}} = \mathbf{M}\tilde{\mathbf{x}} + \mathbf{N}\mathbf{A}\tilde{\mathbf{x}} = \mathbf{M}\tilde{\mathbf{x}} + \mathbf{N}\mathbf{b} = \phi(\tilde{\mathbf{x}}, \mathbf{b}),$$

wodurch die Konsistenz des Iterationsverfahrens ϕ zur Matrix $\mathbf{A}$ folgt. $\square$

Die Wahl $\mathbf{M} = \mathbf{I}$ und $\mathbf{N} = \mathbf{0}$ zeigt bereits deutlich, daß die Forderung nach der Konsistenz des Verfahrens zwar eine notwendige, jedoch keine hinreichende Bedingung zur Festlegung praktikabler linearer Iterationsverfahren repräsentiert. Wir benötigen folglich eine zusätzliche Forderung, um die Konvergenz der Methode gegen die gesuchte Lösung sicherzustellen.

Satz 4.5 *Ein lineares Iterationsverfahren ϕ ist genau dann konvergent, wenn der Spektralradius der Iterationsmatrix $\mathbf{M}$ die Bedingung*

$$\rho(\mathbf{M}) < 1$$

erfüllt.

Beweis:

„ $\Rightarrow$ “ ϕ sei konvergent.

Sei λ Eigenwert von $\mathbf{M}$ mit $|\lambda| = \rho(\mathbf{M})$ und $\mathbf{x} \in \mathbb{C}^n \setminus \{0\}$ der zugehörige Eigenvektor. Wählen wir $\mathbf{b} = \mathbf{0} \in \mathbb{C}^n$, dann folgt für $\mathbf{x}_0 = c\mathbf{x}$ mit beliebigem $c \in \mathbb{R} \setminus \{0\}$ die Iterationsfolge

$$\mathbf{x}_m = \phi(\mathbf{x}_{m-1}, \mathbf{b}) = \mathbf{M}\mathbf{x}_{m-1} = \dots = \mathbf{M}^m\mathbf{x}_0 = \lambda^m\mathbf{x}_0.$$

Im Fall $|\lambda| > 1$ folgt aus $\|\mathbf{x}_m\| = |\lambda|^m\|\mathbf{x}_0\|$ die Divergenz der Folge $\{\mathbf{x}_m\}_{m \in \mathbb{N}}$.

Für $|\lambda| = 1$ stellt $\mathbf{M}$ für den Eigenvektor eine Drehung dar. Die Konvergenz der Folge $\{\mathbf{x}_m\}_{m \in \mathbb{N}}$ liegt daher nur im Fall $\lambda = 1$ vor. Hierbei erhalten wir $\mathbf{x}_m = \mathbf{x}_0$ für alle $m \in \mathbb{N}$ unabhängig vom gewählten Skalierungsparameter c , so daß sich mit

$$\hat{\mathbf{x}} = \lim_{m \rightarrow \infty} \mathbf{x}_m = \mathbf{x}_0$$

ein vom Startvektor abhängiger Grenzwert ergibt und daher das Iterationsverfahren nicht konvergent ist.

Die Bedingung $|\lambda| < 1$ und damit $\rho(\mathbf{M}) < 1$ stellt demzufolge ein notwendiges Kriterium für die Konvergenz des Iterationsverfahrens dar.

„ $\Leftarrow$ “ Gelte $\rho(\mathbf{M}) < 1$.

Da laut Satz 2.11 alle Normen auf dem $\mathbb{C}^n$ äquivalent sind, kann die Konvergenz in einer beliebigen Norm nachgewiesen werden.

Sei $\varepsilon := \frac{1}{2}(1 - \rho(\mathbf{M})) > 0$, dann existiert mit Satz 2.35 eine Norm auf $\mathbb{C}^{n \times n}$ derart, daß

$$q := \|\mathbf{M}\| \leq \rho(\mathbf{M}) + \varepsilon < 1$$

gilt. Bei gegebenem $\mathbf{b} \in \mathbb{C}^n$ definieren wir

$$\begin{aligned} F: \mathbb{C}^n &\rightarrow \mathbb{C}^n \\ \mathbf{x} &\xrightarrow{F} F(\mathbf{x}) = \mathbf{M}\mathbf{x} + \mathbf{N}\mathbf{b}. \end{aligned}$$

Hiermit erhalten wir

$$\|F(\mathbf{x}) - F(\mathbf{y})\| = \|\mathbf{M}\mathbf{x} - \mathbf{M}\mathbf{y}\| \leq \|\mathbf{M}\| \|\mathbf{x} - \mathbf{y}\| = q \|\mathbf{x} - \mathbf{y}\|,$$

so daß aufgrund des Banachschen Fixpunktsatzes die durch

$$\mathbf{x}_{m+1} = F(\mathbf{x}_m)$$

definierte Folge $\{\mathbf{x}_m\}_{m \in \mathbb{N}}$ für ein beliebiges Startelement $\mathbf{x}_0 \in \mathbb{C}^n$ gegen den eindeutig bestimmten Fixpunkt

$$\hat{\mathbf{x}} = \lim_{m \rightarrow \infty} \mathbf{x}_{m+1} = \lim_{m \rightarrow \infty} F(\mathbf{x}_m) = \lim_{m \rightarrow \infty} \phi(\mathbf{x}_m, \mathbf{b})$$

konvergiert und folglich mit ϕ ein konvergentes Iterationsverfahren vorliegt. $\square$

Natürlich können analog zur Konsistenz auch Matrizen $\mathbf{M}$ und $\mathbf{N}$ angegeben werden, die ein konvergentes lineares Iterationsverfahren generieren, ohne in einem zweckmäßigen Verhältnis zum Gleichungssystem zu stehen (z. B. $\mathbf{M} = \mathbf{0}$, $\mathbf{N} = \mathbf{I}$). Erst das Zusammenwirken von Konsistenz und Konvergenz liefert eine geeignete Iterationsvorschrift.

Satz 4.6 *Sei ϕ ein konvergentes und zur Matrix $\mathbf{A}$ konsistentes lineares Iterationsverfahren, dann erfüllt das Grenzelement $\tilde{\mathbf{x}}$ der Folge*

$$\mathbf{x}_m = \phi(\mathbf{x}_{m-1}, \mathbf{b}) \text{ für } m = 1, 2, \dots$$

für jedes $\mathbf{x}_0 \in \mathbb{C}^n$ das Gleichungssystem (4.0.1).

Beweis:

Mit Satz 4.5 konvergiert die Folge $\mathbf{x}_m = \phi(\mathbf{x}_{m-1}, \mathbf{b})$ gegen den eindeutig bestimmten Fixpunkt

$$\tilde{\mathbf{x}} = \phi(\tilde{\mathbf{x}}, \mathbf{b}),$$

der wegen der Konsistenz des Iterationsverfahrens die Lösung der Gleichung $\mathbf{A}\mathbf{x} = \mathbf{b}$ darstellt. $\square$

4.1 Splitting-Methoden

Splitting-Methoden zur Lösung des Gleichungssystems (4.0.1) basieren auf einer Aufteilung der Matrix $\mathbf{A}$ in der Form

$$\mathbf{A} = \mathbf{B} + (\mathbf{A} - \mathbf{B}), \quad \mathbf{B} \in \mathbb{C}^{n \times n}, \quad (4.1.1)$$

so daß sich aus

$$\mathbf{A}\mathbf{x} = \mathbf{b}$$

das äquivalente System

$$Bx = (B - A)x + b$$

ergibt. Ist B zudem regulär, dann erhalten wir

$$x = B^{-1}(B - A)x + B^{-1}b$$

und definieren hierdurch das lineare Iterationsverfahren

$$x_{m+1} = \phi(x_m, b) = Mx_m + Nb \text{ für } m = 0, 1, \dots$$

mit

$$M := B^{-1}(B - A)$$

und

$$N := B^{-1}.$$

Bevor wir Splitting-Verfahren hinsichtlich ihrer Konsistenz und ihres Konvergenzverhaltens untersuchen, werden wir mit der folgenden Definition den Begriff der symmetrischen Splitting-Methode einführen. Solche Methoden erweisen sich bei der im Kapitel 5 beschriebenen Präkonditionierung symmetrischer, positiv definiter Matrizen als vorteilhaft, da mit Hilfe der Iterationsmatrix symmetrischer Splitting-Methoden diese Eigenschaften der Matrix auch über die äquivalente Umformulierung hinaus erhalten werden können. Folglich können Methoden für positiv definite, symmetrische Matrizen, wie zum Beispiel das im Abschnitt 4.3.1.3 hergeleitete Verfahren der konjugierten Gradienten, auch nach der Präkonditionierung angewendet werden.

Definition 4.7 Die Splitting-Methode

$$x_{m+1} = B^{-1}(B - A)x_m + B^{-1}b$$

zur Lösung der Gleichung (4.0.1) heißt symmetrisch, falls für jede positiv definite und symmetrische Matrix A die Matrix B ebenfalls positiv definit und symmetrisch ist.

Satz 4.8 Sei $B \in \mathbb{C}^{n \times n}$ regulär, dann ist das lineare Iterationsverfahren

$$x_{m+1} = \phi(x_m, b) = B^{-1}(B - A)x_m + B^{-1}b$$

zur Matrix A konsistent.

Beweis:

Mit $M = B^{-1}(B - A) = I - B^{-1}A = I - NA$ folgt die Behauptung durch Anwendung des Satzes 4.4. $\square$

Gilt für eine Splitting-Methode $\rho(M) < 1$, dann stellt das eindeutig bestimmte Grenzelement der Folge

$$x_m = \phi(x_{m-1}, b) \text{ für } m = 1, 2, \dots$$

für beliebigen Startvektor $x_0 \in \mathbb{C}^n$ die Lösung der zugehörigen Gleichung $Ax = b$ dar. Betrachten wir ein lineares Iterationsverfahren

$$x_m = Mx_{m-1} + Nb \text{ für } m = 1, 2, \dots$$

mit $\rho(M) < 1$, dann existiert laut Satz 2.35 zu jedem ε mit $0 < \varepsilon < 1 - \rho(M)$ eine Norm derart, daß

$$\rho(\mathbf{M}) \leq \underbrace{\|\mathbf{M}\|}_{q:=} \leq \rho(\mathbf{M}) + \varepsilon < 1$$

gilt. Aus dem Banachschen Fixpunktsatz folgt die *a priori* Fehlerabschätzung

$$\|\mathbf{x}_m - \mathbf{A}^{-1}\mathbf{b}\| \leq \frac{q^m}{1-q} \|\mathbf{x}_1 - \mathbf{x}_0\| \text{ für } m = 1, 2, \dots$$

Für jede weitere Norm $\|\cdot\|_a$ gilt

$$\|\mathbf{x}_m - \mathbf{A}^{-1}\mathbf{b}\|_a \leq \frac{q^m}{1-q} C_a \|\mathbf{x}_1 - \mathbf{x}_0\|_a$$

mit einer Konstanten $C_a > 0$, die nur von den Normen abhängt. Somit stellt der Spektralradius in jeder Norm ein Maß für die Konvergenzgeschwindigkeit dar.

Satz 4.9 Sei ϕ ein zur Matrix $\mathbf{A}$ konsistentes lineares Iterationsverfahren, für dessen zugehörige Iterationsmatrix $\mathbf{M}$ eine Norm derart existiert, daß $q := \|\mathbf{M}\| < 1$ gilt, dann folgt für gegebenes $\varepsilon > 0$

$$\|\mathbf{x}_m - \mathbf{A}^{-1}\mathbf{b}\| \leq \varepsilon$$

für alle $m \in \mathbb{N}$ mit

$$m \geq \frac{\ln \frac{\varepsilon(1-q)}{\|\mathbf{x}_1 - \mathbf{x}_0\|}}{\ln q}$$

und $\mathbf{x}_1 = \phi(\mathbf{x}_0, \mathbf{b}) \neq \mathbf{x}_0$.

Beweis:

Mit $\|\phi(\mathbf{x}, \mathbf{b}) - \phi(\mathbf{y}, \mathbf{b})\| \leq q \|\mathbf{x} - \mathbf{y}\|$ folgt mit der *a priori* Fehlerabschätzung des Banachschen Fixpunktsatzes die Ungleichung

$$\|\mathbf{x}_m - \mathbf{A}^{-1}\mathbf{b}\| \leq \frac{q^m}{1-q} \|\mathbf{x}_1 - \mathbf{x}_0\|.$$

Zu gegebenem $\varepsilon > 0$ erhalten wir unter Ausnutzung von $\mathbf{x}_1 \neq \mathbf{x}_0$ für

$$m \geq \frac{\ln \frac{\varepsilon(1-q)}{\|\mathbf{x}_1 - \mathbf{x}_0\|}}{\ln q}.$$

somit die Abschätzung

$$\|\mathbf{x}_m - \mathbf{A}^{-1}\mathbf{b}\| \leq \frac{q^m}{1-q} \|\mathbf{x}_1 - \mathbf{x}_0\| \leq \frac{\frac{\varepsilon(1-q)}{\|\mathbf{x}_1 - \mathbf{x}_0\|}}{1-q} \|\mathbf{x}_1 - \mathbf{x}_0\| = \varepsilon.$$

□

Gilt $0 < q < 1$ und $\tilde{q} = q^2$, dann folgt

$$\frac{\tilde{q}^m}{1-\tilde{q}} = \frac{q^{2m}}{1-q^2} < \frac{q^{2m}}{1-q}.$$

Betrachtet man folglich zwei konvergente lineare Iterationsverfahren ϕ_1 und ϕ_2 deren zugeordnete Iterationsmatrizen $\mathbf{M}_1$ und $\mathbf{M}_2$ die Eigenschaft

$$\rho(\mathbf{M}_1) = \rho(\mathbf{M}_2)^2$$

erfüllen, dann liefert Satz 4.9 eine gesicherte Genauigkeitsaussage für die Methode ϕ_1 in der Regel nach der Hälfte der für das Verfahren ϕ_2 benötigten Iterationszahl. Innerhalb eines iterativen Verfahrens dieser Klasse darf daher mit einer Halbierung der benötigten Iterationen gerechnet werden, wenn der Spektralradius beispielsweise von 0.9 auf 0.81 gesenkt wird.

Beispiel 4.10 Triviales Verfahren

Wir betrachten das Modellproblem

$$\underbrace{\begin{pmatrix} 0.7 & -0.4 \\ -0.2 & 0.5 \end{pmatrix}}_{\mathbf{A}:=} \underbrace{\begin{pmatrix} x_1 \\ x_2 \end{pmatrix}}_{\mathbf{x}:=} = \underbrace{\begin{pmatrix} 0.3 \\ 0.3 \end{pmatrix}}_{\mathbf{b}:=}. \quad (4.1.2)$$

Die exakte Lösung lautet $\mathbf{A}^{-1}\mathbf{b} = (1,1)^T$. Natürlich besteht für dieses Gleichungssystem keine Notwendigkeit zur Nutzung eines iterativen Verfahrens. Das Beispiel eignet sich jedoch sehr gut zur Verdeutlichung der Effizienz der einzelnen Splitting-Methoden. Mit $\mathbf{A} = \mathbf{I} - (\mathbf{I} - \mathbf{A})$ folgt die Äquivalenz zwischen $\mathbf{A}\mathbf{x} = \mathbf{b}$ und

$$\mathbf{x} = (\mathbf{I} - \mathbf{A})\mathbf{x} + \mathbf{b}.$$

Hierdurch ergibt sich das einfache konsistente und lineare Iterationsverfahren

$$\mathbf{x}_{m+1} = \phi(\mathbf{x}_m, \mathbf{b}) = \underbrace{(\mathbf{I} - \mathbf{A})}_{\mathbf{M}:=} \mathbf{x}_m + \underbrace{\mathbf{I}}_{\mathbf{N}:=} \mathbf{b}. \quad (4.1.3)$$

Es gilt $\det(\mathbf{M} - \lambda\mathbf{I}) = (\lambda - 0.4)^2 - 0.09$, so daß die Eigenwerte der Iterationsmatrix $\lambda_1 = 0.1$ und $\lambda_2 = 0.7$ sind und damit das Verfahren (4.1.3) mit $\rho(\mathbf{M}) = 0.7 < 1$ konvergiert. Sei $\mathbf{x}_0 = (21, -19)^T$, dann erhalten wir den in der folgenden Tabelle aufgeführten Konvergenzverlauf.

Triviales Verfahren				
m	$x_{m,1}$	$x_{m,2}$	$\varepsilon_m := \ \mathbf{x}_m - \mathbf{A}^{-1}\mathbf{b}\ _\infty$	$\varepsilon_m/\varepsilon_{m-1}$
0	2.100000e+01	-1.900000e+01	2.000000e+01	
10	8.116832e-01	8.116832e-01	1.883168e-01	7.000000e-01
40	9.999958e-01	9.999958e-01	4.244537e-06	7.000000e-01
70	1.000000e-00	1.000000e-00	9.566903e-11	7.000002e-01
96	1.000000e-00	1.000000e-00	8.881784e-15	6.956522e-01

Der Konvergenzverlauf zeigt, daß eine derartig primitive Wahl der Iterationsmatrix in der Regel zu keinem zufriedenstellenden Algorithmus führen wird. Die Entwicklung effektiver Splitting-Methoden korreliert mit der gezielten Festlegung der Matrix $\mathbf{B}$, die einerseits leicht invertierbar sein muß und andererseits eine gute Approximation der Matrix $\mathbf{A}$ mit dem Ziel $\rho(\mathbf{M}) \ll 1$ darstellen soll.

4.1.1 Jacobi-Verfahren

Das Jacobi-Verfahren zur iterativen Lösung des linearen Gleichungssystems $\mathbf{Ax} = \mathbf{b}$ mit regulärer Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$ setzt nichtverschwindende Diagonalelemente $a_{ii} \neq 0$, $i = 1, \dots, n$ der Matrix $\mathbf{A}$ voraus, so daß mit $\mathbf{D} = \text{diag}\{a_{11}, \dots, a_{nn}\}$ eine reguläre Diagonalmatrix vorliegt. Der Grundidee der Splitting-Methoden folgend, schreiben wir das lineare Gleichungssystem $\mathbf{Ax} = \mathbf{b}$ in der äquivalenten Form

$$\mathbf{x} = \underbrace{\mathbf{D}^{-1}(\mathbf{D} - \mathbf{A})}_{\mathbf{M}_J :=} \mathbf{x} + \underbrace{\mathbf{D}^{-1}\mathbf{b}}_{\mathbf{N}_J :=}.$$

Mit Satz 4.8 ist das hiermit definierte lineare Iterationsverfahren

$$\mathbf{x}_{m+1} = \mathbf{D}^{-1}(\mathbf{D} - \mathbf{A})\mathbf{x}_m + \mathbf{D}^{-1}\mathbf{b} \text{ für } m = 0, 1, 2, \dots \quad (4.1.4)$$

konsistent zur Matrix $\mathbf{A}$, und wir erhalten die Komponentenschreibweise

$$x_{m+1,i} = \frac{1}{a_{ii}} \left(b_i - \sum_{\substack{j=1 \\ j \neq i}}^n a_{ij} x_{m,j} \right) \text{ für } i = 1, \dots, n \text{ und } m = 0, 1, 2, \dots$$

Beim Jacobi-Verfahren wird die neue Iterierte $\mathbf{x}_{m+1}$ somit ausschließlich mittels der alten Iterierten $\mathbf{x}_m$ ermittelt. Die Methode wird aus diesem Grund auch als Gesamtschrittverfahren bezeichnet und ist folglich unabhängig von der gewählten Nummerierung der Unbekannten $\mathbf{x} = (x_1, \dots, x_n)^T$. Da es sich um eine Splitting-Methode handelt, ist die Konvergenz des Jacobi-Verfahrens einzig vom Spektralradius der Iterationsmatrix $\mathbf{M}_J = \mathbf{D}^{-1}(\mathbf{D} - \mathbf{A})$ abhängig. Dieser Wert kann laut Satz 2.34 durch jede beliebige Matrixnorm abgeschätzt werden, wodurch sich die Konvergenz des Verfahrens ohne Berechnung der Eigenwerte anhand der Größe der Matrixkoeffizienten überprüfen läßt.

Satz 4.11 *Erfüllt die reguläre Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$ mit $a_{ii} \neq 0$, $i = 1, \dots, n$ das starke Zeilensummenkriterium*

$$q_\infty := \max_{i=1, \dots, n} \sum_{\substack{k=1 \\ k \neq i}}^n \frac{|a_{ik}|}{|a_{ii}|} < 1$$

oder das starke Spaltensummenkriterium

$$q_1 := \max_{k=1, \dots, n} \sum_{\substack{i=1 \\ i \neq k}}^n \frac{|a_{ik}|}{|a_{ii}|} < 1$$

oder das Quadratsummenkriterium

$$q_2 := \sum_{\substack{i,k=1 \\ i \neq k}}^n \left(\frac{|a_{ik}|}{|a_{ii}|} \right)^2 < 1,$$

dann konvergiert das Jacobi-Verfahren bei beliebigem Startvektor $\mathbf{x}_0 \in \mathbb{C}^n$ und für beliebige rechte Seite $\mathbf{b} \in \mathbb{C}^n$ gegen $\mathbf{A}^{-1}\mathbf{b}$.

Beweis:

Wegen

$$\mathbf{M}_J = \mathbf{D}^{-1}(\mathbf{D} - \mathbf{A})$$

folgt

$$q_\infty = \|\mathbf{M}_J\|_\infty, \quad q_1 = \|\mathbf{M}_J\|_1 \quad \text{und} \quad 1 > \sqrt{q_2} \geq \|\mathbf{M}_J\|_2,$$

wodurch sich die Behauptung durch Anwendung des Satzes 4.5 auf der Grundlage

$$\rho(\mathbf{M}_J) \leq \min\{\|\mathbf{M}_J\|_\infty, \|\mathbf{M}_J\|_1, \|\mathbf{M}_J\|_2\} < 1$$

ergibt. □

Bemerkung:

Eine Matrix, die das starke Zeilensummenkriterium erfüllt, wird als strikt diagonaldominant bezeichnet.

Häufig treten in praktischen Anwendungen Matrizen auf, die keiner der drei in Satz 4.11 aufgeführten Konvergenzkriterien genügen. Einen bekannten Repräsentanten stellt die innerhalb des Beispiels 1.1 bei der Diskretisierung der Poisson-Gleichung hergeleitete Matrix dar. Dennoch kann für derartige Gleichungssysteme die Konvergenz des Jacobi-Verfahrens nachgewiesen werden. Wie wir sehen werden, liegt die entscheidende Eigenschaft dabei in der Irreduzibilität der Matrix begründet.

Definition 4.12 Eine Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$ heißt *reduzibel* oder *zerlegbar*, falls eine Permutationsmatrix $\mathbf{P} \in \mathbb{R}^{n \times n}$ derart existiert, daß

$$\mathbf{PAP}^T = \begin{pmatrix} \tilde{\mathbf{A}}_{11} & \tilde{\mathbf{A}}_{12} \\ \mathbf{0} & \tilde{\mathbf{A}}_{22} \end{pmatrix}$$

mit $\tilde{\mathbf{A}}_{ii} \in \mathbb{C}^{n_i \times n_i}$, $n_i > 0$, $i = 1, 2$, $n_1 + n_2 = n$ gilt. Andernfalls heißt $\mathbf{A}$ *irreduzibel* oder *unzerlegbar*.

Die Reduzibilität ist somit ausschließlich von der Besetzungsstruktur der Matrix und nicht von den absoluten Werten der Matrixkoeffizienten abhängig. Zudem haben die Diagonalelemente keinen Einfluß auf die Reduzibilität.

Satz 4.13 Sei die reguläre Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$ irreduzibel und diagonaldominant, das heißt, es gilt

$$\max_{i=1, \dots, n} \sum_{\substack{j=1 \\ j \neq i}}^n \frac{|a_{ij}|}{|a_{ii}|} \leq 1, \quad (4.1.5)$$

und es existiere ein $k \in \{1, \dots, n\}$ mit

$$\sum_{\substack{j=1 \\ j \neq k}}^n \frac{|a_{kj}|}{|a_{kk}|} < 1, \quad (4.1.6)$$

dann konvergiert das Jacobi-Verfahren bei beliebigem Startvektor $\mathbf{x}_0 \in \mathbb{C}^n$ und für jede beliebige rechte Seite $\mathbf{b} \in \mathbb{C}^n$ gegen $\mathbf{A}^{-1}\mathbf{b}$.

Die Grundidee des Satzes liegt in einer graphentheoretischen Überlegung: Bezeichnen wir mit $G^{\mathbf{A}} := \{V^{\mathbf{A}}, E^{\mathbf{A}}\}$ den gerichteten Graphen der Matrix $\mathbf{A}$, der aus den Knoten $V^{\mathbf{A}} := \{1, \dots, n\}$ und der Kantenmenge geordneter Paare $E^{\mathbf{A}} := \{(i, j) \in V^{\mathbf{A}} \times V^{\mathbf{A}} \mid a_{ij} \neq 0\}$ besteht. Eine Matrix $\mathbf{A}$ ist genau dann irreduzibel, wenn der zugehörige gerichtete Graph $G^{\mathbf{A}}$ zusammenhängend ist, das heißt, wenn es zu je zwei Indizes $i_0 = i, i_\ell = j \in V^{\mathbf{A}}$ einen gerichteten Weg der Länge $\ell \in \mathbb{N}$

$$(i_0, i_1)(i_1, i_2) \dots (i_{\ell-1}, i_\ell)$$

mit $(i_k, i_{k+1}) \in E^{\mathbf{A}}$ für $k = 0, \dots, \ell - 1$ gibt.

Ist

$$\mathbf{e}_m := \mathbf{x}_m - \mathbf{A}^{-1}\mathbf{b}$$

der Fehlervektor der m -ten Iterierten des Jacobi-Verfahrens, dann folgt mit (4.1.5) $|e_{m+1,i}| \leq \|\mathbf{e}_m\|_\infty$ für $i = 1, \dots, n$, und mit (4.1.6) existiert ein $k \in \{1, \dots, n\}$ mit

$$|e_{m+1,k}| \leq \gamma \|\mathbf{e}_m\|_\infty \text{ mit } \gamma < 1.$$

Hiermit folgt

$$|e_{m+2,j}| \leq \gamma \|\mathbf{e}_m\|_\infty \text{ mit } \gamma < 1$$

für alle $j \in \{1, \dots, n\}$ mit $(j, k) \in E^{\mathbf{A}}$, und aufgrund des zusammenhängenden Graphen gilt

$$\|\mathbf{e}_{m+n}\|_\infty \leq \gamma \|\mathbf{e}_m\|_\infty \quad (4.1.7)$$

mit $\gamma < 1$, wodurch sich die Konvergenz des Jacobi-Verfahrens ergibt. Die Irreduzibilität der Matrix stellt somit sicher, daß die Eigenschaft (4.1.6) bei diagonaldominanten Matrizen spätestens nach n Iterationen zu einer betragsmäßigen Verringerung aller Komponenten innerhalb des Fehlervektors führt. Die Anzahl der benötigten Iterationen wird hierbei durch den maximalen Wert der Länge des jeweils kürzesten gerichteten Weges zwischen Zeilen, die der Ungleichung (4.1.6) genügen, und Zeilen $i \in \{1, \dots, n\}$, die die Gleichung

$$\sum_{\substack{j=1 \\ j \neq i}}^n \frac{|a_{ij}|}{|a_{ii}|} = 1$$

erfüllen, bestimmt. Liegt eine vollbesetzte Matrix vor, die den Bedingungen des Satzes 4.13 genügt, so ergibt sich die Ungleichung (4.1.7) spätestens nach zwei Iterationen.

Wir kommen nun zum Beweis des Satzes 4.13.

Beweis:

Betrachten wir die Iterationsmatrix $\mathbf{M}_J = (m_{ij})_{i,j=1,\dots,n}$ des Jacobi-Verfahrens und definieren $|\mathbf{M}_J| := (|m_{ij}|)_{i,j=1,\dots,n}$ sowie $\mathbf{y} := (1, \dots, 1)^T \in \mathbb{R}^n$, so gilt mit (4.1.5)

$$0 \leq (|\mathbf{M}_J|\mathbf{y})_i \leq y_i \text{ für } i = 1, \dots, n, \quad (4.1.8)$$

und mit (4.1.6) existiert ein $k \in \{1, \dots, n\}$ mit

$$0 \leq (|\mathbf{M}_J|\mathbf{y})_k < y_k. \quad (4.1.9)$$

Aus (4.1.8) folgt $(|\mathbf{M}_J|^m \mathbf{y})_i \leq y_i$ für $i = 1, \dots, n$ und alle $m \in \mathbb{N}$. Existiert ein $\tilde{m} \in \mathbb{N}$ mit $(|\mathbf{M}_J|^{\tilde{m}} \mathbf{y})_j < y_j$, so erhalten wir für alle $m \geq \tilde{m}$ mit (4.1.5) die Ungleichung

$$(|\mathbf{M}_J|^m \mathbf{y})_j < y_j. \quad (4.1.10)$$

Mit

$$\mathbf{t}^m := \mathbf{y} - |\mathbf{M}_J|^m \mathbf{y}$$

und

$$\tau^m := \#\{t_i^m \mid t_i^m \neq 0, i = 1, \dots, n\}$$

ergibt sich unter Verwendung von (4.1.9) und (4.1.10)

$$0 < \tau^1 \leq \tau^2 \leq \dots \leq \tau^m \leq \tau^{m+1} \leq \dots$$

Nun nehmen wir an, es gäbe ein $m \in \{1, \dots, n-1\}$ mit $\tau^m = \tau^{m+1} < n$ und führen dieses zum Widerspruch.

O.B.d.A. weise der Vektor $\mathbf{t}^m$ die Form

$$\mathbf{t}^m = \begin{pmatrix} \mathbf{u} \\ \mathbf{0} \end{pmatrix}, \quad \mathbf{u} \in \mathbb{R}^p \text{ für } 1 \leq p < n, \text{ und } u_i > 0 \text{ für } i = 1, \dots, p$$

auf, dann folgt mit (4.1.10) und $\tau^m = \tau^{m+1}$

$$\mathbf{t}^{m+1} = \begin{pmatrix} \mathbf{v} \\ \mathbf{0} \end{pmatrix}, \quad \mathbf{v} \in \mathbb{R}^p, \text{ und } v_i > 0 \text{ für } i = 1, \dots, p.$$

Sei

$$|\mathbf{M}_J| = \begin{pmatrix} |\mathbf{M}_{11}| & |\mathbf{M}_{12}| \\ |\mathbf{M}_{21}| & |\mathbf{M}_{22}| \end{pmatrix}$$

mit $|\mathbf{M}_{11}| \in \mathbb{R}^{p \times p}$, dann ergibt sich unter Verwendung der Ungleichung (4.1.8)

$$\begin{pmatrix} \mathbf{v} \\ \mathbf{0} \end{pmatrix} = \mathbf{t}^{m+1} = \mathbf{y} - |\mathbf{M}_J|^{m+1} \mathbf{y} \geq |\mathbf{M}_J| \mathbf{y} - |\mathbf{M}_J|^{m+1} \mathbf{y} = |\mathbf{M}_J| \mathbf{t}^m = \begin{pmatrix} |\mathbf{M}_{11}| \mathbf{u} \\ |\mathbf{M}_{21}| \mathbf{u} \end{pmatrix}.$$

Mit $u_i > 0$, $i = 1, \dots, p$ erhalten wir aufgrund der Nichtnegativität der Elemente der Matrix $|\mathbf{M}_{21}|$ die Gleichung

$$|\mathbf{M}_{21}| = \mathbf{0},$$

wodurch $\mathbf{M}_j$ reduzibel ist. Da sich die Besetzungsstrukturen von $\mathbf{M}_j$ und $\mathbf{A}$ bis auf die Diagonalelemente gleichen und diese keinen Einfluß auf die Reduzibilität haben, liegt demzufolge ein Widerspruch zur Irreduzibilität der Matrix $\mathbf{A}$ vor. Somit folgt im Fall $\tau^m < n$ direkt

$$0 < \tau^1 < \tau^2 < \dots < \tau^m < \tau^{m+1}$$

für $m \in \{1, \dots, n-1\}$, wodurch sich die Existenz eines $m \in \{1, \dots, n\}$ mit

$$0 \leq (|\mathbf{M}_J|^m \mathbf{y})_i < y_i$$

für $i = 1, \dots, n$ ergibt. Hiermit erhalten wir

$$\rho(\mathbf{M}_J)^m \leq \rho(\mathbf{M}_J^m) \leq \|\mathbf{M}_J^m\|_\infty \leq \| |\mathbf{M}_J|^m \|_\infty < 1$$

und folglich $\rho(\mathbf{M}_J) < 1$. □

Beispiel 4.14 Für das Gleichungssystem $\mathbf{A}\mathbf{u} = \mathbf{g}$ mit der aus der Diskretisierung des Poisson-Problems resultierenden Matrix $\mathbf{A}$ (siehe (1.0.3)) konvergiert das Jacobi-Verfahren bei beliebigem $\mathbf{g} \in \mathbb{R}^N$ und $\mathbf{u}_0 \in \mathbb{R}^N$. Diese Behauptung belegen wir durch den Nachweis der Voraussetzungen des Satzes 4.13. Zunächst erfüllt $\mathbf{A}$ das schwache Zeilensummenkriterium, und es gilt

$$\sum_{i=2}^n \frac{|a_{1i}|}{|a_{11}|} = \frac{1}{4} + \frac{1}{4} < 1,$$

so daß die Irreduzibilität der Matrix zu zeigen bleibt. Als ersten Schritt weisen wir die Irreduzibilität der in der Matrix $\mathbf{A} \in \mathbb{R}^{N^2 \times N^2}$ enthaltenen Submatrizen $\mathbf{B} \in \mathbb{R}^{N \times N}$ nach. Seien $i, j \in V^{\mathbf{B}}$ gegeben. Aufgrund der vorliegenden Symmetrie der Matrix können wir o. B. d. A. $i < j$ voraussetzen. Wegen $(\ell, \ell+1) \in E^{\mathbf{B}}$, $\ell = 1, \dots, N-1$ erhalten wir durch

$$(i, i+1)(i+1, i+2) \dots (i+k-1, j)$$

mit $k = j - i$ einen gerichteten Weg von i nach j , womit die Irreduzibilität der Matrix $\mathbf{B}$ folgt.

Wir kommen nun zur Matrix $\mathbf{A}$: Seien $i, j \in V^{\mathbf{A}}$ gegeben, wobei analog zur Matrix $\mathbf{B}$ aufgrund der Symmetrie der Matrix $\mathbf{A}$ wiederum $i < j$ o. B. d. A. angenommen wird, dann existiert stets eine eindeutige Darstellung der Indizes in der Form

$$i = i_1 N + i_2 \quad \text{und} \quad j = j_1 N + j_2$$

mit $0 \leq i_1 \leq j_1 \leq N-1$ und $i_2, j_2 \in \{1, \dots, N\}$. Gilt $i_1 = j_1$, so existiert aufgrund der Irreduzibilität der Matrix $\mathbf{B}$ ein gerichteter Weg von i nach j . Im Fall $i_1 < j_1$ existiert zunächst ein gerichteter Weg von i nach $\tilde{j} = i_1 N + j_2 \leq N^2 - N$. Wegen $(\ell, \ell+N) \in E^{\mathbf{A}}$, für $\ell = 1, \dots, N^2 - N$ erhalten wir durch

$$(\tilde{j}, \tilde{j}+N)(\tilde{j}+N, \tilde{j}+2N) \dots (\tilde{j}+(k-1)N, j)$$

mit $k = j_1 - i_1$ die Vervollständigung des gerichteten Weges von i nach j .

Eine reduzierbare Matrix ergibt sich beispielsweise bei der Diskretisierung der Konvektions-Diffusions-Gleichung im Grenzfall eines verschwindenden Diffusionsparameters. In diesem Fall kann das entsprechende Gleichungssystem jedoch durch einfache Vorwärtselimination gelöst werden.

Beispiel 4.15 Für das Modellproblem (4.1.2)

$$\mathbf{A} = \begin{pmatrix} 0.7 & -0.4 \\ -0.2 & 0.5 \end{pmatrix}, \quad \mathbf{b} = \begin{pmatrix} 0.3 \\ 0.3 \end{pmatrix}$$

liegt mit

$$q_\infty := \max_{i=1,2} \sum_{\substack{j=1 \\ j \neq i}}^2 \frac{|a_{ij}|}{|a_{ii}|} = \max \left\{ \frac{4}{7}, \frac{2}{5} \right\} < 1$$

der Nachweis der Konvergenz des Jacobi-Verfahrens vor. Die Eigenwerte der Iterationsmatrix

$$\mathbf{M}_J = \mathbf{D}^{-1}(\mathbf{D} - \mathbf{A}) = \begin{pmatrix} 0 & \frac{4}{7} \\ \frac{2}{5} & 0 \end{pmatrix}$$

lauten

$$\lambda_{1,2} = \pm \sqrt{\frac{8}{35}},$$

so daß

$$\rho(\mathbf{M}_J) = \sqrt{\frac{8}{35}} \approx 0.4781$$

folgt. Unter Verwendung des Startvektors $\mathbf{x}_0 = (21, -19)^T$ erhalten wir den folgenden Iterationsverlauf, der auch die im Vergleich zur Konvergenz des trivialen Verfahrens (siehe Beispiel 4.10) wegen $\rho(\mathbf{M}_J) \approx (\rho(\mathbf{I} - \mathbf{A}))^2$ zu erwartende Halbierung der Iterationsanzahl belegt.

Jacobi Verfahren				
m	$x_{m,1}$	$x_{m,2}$	$\varepsilon_m := \ \mathbf{x}_m - \mathbf{A}^{-1}\mathbf{b}\ _\infty$	$\varepsilon_m/\varepsilon_{m-1}$
0	2.100000e+01	-1.900000e+01	2.000000e+01	
15	9.996275e-01	1.000261e+00	3.725165e-04	5.714286e-01
30	1.000000e+00	1.000000e-00	4.856900e-09	4.000000e-01
45	1.000000e-00	1.000000e+00	9.037215e-14	5.700280e-01
48	1.000000e+00	1.000000e-00	8.437695e-15	4.086022e-01

4.1.2 Gauß-Seidel-Verfahren

Analog zum vorgestellten Jacobi-Verfahren wird auch beim Gauß-Seidel-Verfahren ein lineares Gleichungssystem $\mathbf{Ax} = \mathbf{b}$ mit einer regulären Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$ betrachtet, die einen ebenfalls regulären Diagonalanteil $\mathbf{D} = \text{diag}\{a_{11}, \dots, a_{nn}\}$ aufweist. Wir definieren die strikte linke untere Dreiecksmatrix

$$\mathbf{L} = (\ell_{ij})_{i,j=1,\dots,n} \quad \text{mit} \quad \ell_{ij} = \begin{cases} a_{ij}, & i > j \\ 0, & \text{sonst} \end{cases} \quad (4.1.11)$$

und die strikte rechte obere Dreiecksmatrix

$$\mathbf{R} = (r_{ij})_{i,j=1,\dots,n} \quad \text{mit} \quad r_{ij} = \begin{cases} a_{ij}, & i < j \\ 0, & \text{sonst.} \end{cases} \quad (4.1.12)$$

Hierdurch erhalten wir das zu $\mathbf{Ax} = \mathbf{b}$ äquivalente Gleichungssystem

$$(\mathbf{D} + \mathbf{L})\mathbf{x} = -\mathbf{Rx} + \mathbf{b} \quad (4.1.13)$$

und

$$\mathbf{x} = \underbrace{-(\mathbf{D} + \mathbf{L})^{-1}\mathbf{R}}_{\mathbf{M}_{GS}:=} \mathbf{x} + \underbrace{(\mathbf{D} + \mathbf{L})^{-1}\mathbf{b}}_{\mathbf{N}_{GS}:=}.$$

Somit gilt

$$\mathbf{M}_{GS} = (\mathbf{D} + \mathbf{L})^{-1}(\mathbf{D} + \mathbf{L} - \mathbf{A}) = \mathbf{I} - \mathbf{N}_{GS}\mathbf{A},$$

und das lineare Iterationsverfahren

$$\mathbf{x}_{m+1} = -(\mathbf{D} + \mathbf{L})^{-1} \mathbf{R} \mathbf{x}_m + (\mathbf{D} + \mathbf{L})^{-1} \mathbf{b} \text{ für } m = 0, 1, \dots \quad (4.1.14)$$

ist konsistent zur Matrix $\mathbf{A}$.

Zur Herleitung der Komponentenschreibweise betrachten wir die i -te Zeile des Iterationsverfahrens (4.1.14) in der Form gemäß Gleichungssystem (4.1.13)

$$\sum_{j=1}^i a_{ij} x_{m+1,j} = - \sum_{j=i+1}^n a_{ij} x_{m,j} + b_i.$$

Seien $x_{m+1,j}$ für $j = 1, \dots, i-1$ bekannt, dann kann $x_{m+1,i}$ durch

$$x_{m+1,i} = \frac{1}{a_{ii}} \left(b_i - \sum_{j=1}^{i-1} a_{ij} x_{m+1,j} - \sum_{j=i+1}^n a_{ij} x_{m,j} \right) \text{ für } i = 1, \dots, n \text{ und } m = 0, 1, 2, \dots \quad (4.1.15)$$

ermittelt werden. Aus dieser Darstellung wird deutlich, daß beim Gauß-Seidel-Verfahren zur Berechnung der i -ten Komponente der $(m+1)$ -ten Iterierten neben den Komponenten der alten m -ten Iterierten $\mathbf{x}_m$ die bereits bekannten ersten $i-1$ Komponenten der $(m+1)$ -ten Iterierten $\mathbf{x}_{m+1}$ verwendet werden. Das Verfahren wird daher auch als Einzelschrittverfahren bezeichnet und ist abhängig von der gewählten Nummerierung der Unbekannten $(x_1, \dots, x_n)^T$. Verglichen mit dem Jacobi-Verfahren wird beim Gauß-Seidel-Algorithmus mit $\mathbf{D} + \mathbf{L}$ eine bessere Approximation der Matrix $\mathbf{A}$ verwendet, wodurch ein kleiner Spektralradius der Iterationsmatrix und folglich eine schnellere Konvergenz erwartet werden darf.

Satz 4.16 *Sei die reguläre Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$ mit $a_{ii} \neq 0$ für $i = 1, \dots, n$ gegeben. Erfüllen die durch*

$$p_i = \sum_{j=1}^{i-1} \frac{|a_{ij}|}{|a_{ii}|} p_j + \sum_{j=i+1}^n \frac{|a_{ij}|}{|a_{ii}|} \text{ für } i = 1, 2, \dots, n$$

rekursiv definierten Zahlen $p_1, \dots, p_n$ die Bedingung

$$p := \max_{i=1, \dots, n} p_i < 1,$$

dann konvergiert das Gauß-Seidel-Verfahren bei beliebigem Startvektor $\mathbf{x}_0$ und für jede beliebige rechte Seite $\mathbf{b}$ gegen $\mathbf{A}^{-1} \mathbf{b}$.

Beweis:

Unser Ziel ist der Nachweis $\|\mathbf{M}_{GS}\|_\infty < 1$. Sei $\mathbf{x} \in \mathbb{C}^n$ mit $\|\mathbf{x}\|_\infty = 1$. Für

$$\mathbf{z} := \mathbf{M}_{GS} \mathbf{x} = -(\mathbf{D} + \mathbf{L})^{-1} \mathbf{R} \mathbf{x}$$

gilt

$$z_i = - \sum_{j=1}^{i-1} \frac{a_{ij}}{a_{ii}} z_j - \sum_{j=i+1}^n \frac{a_{ij}}{a_{ii}} x_j. \quad (4.1.16)$$

Somit folgt unter Verwendung von $\|\mathbf{x}\|_\infty = 1$ die Abschätzung

$$|z_1| \leq \sum_{j=2}^n \frac{|a_{1j}|}{|a_{11}|} = p_1 < 1.$$

Seien $z_1, \dots, z_{i-1}$ mit $|z_j| \leq p_j$, $j = 1, \dots, i-1 < n$ gegeben, dann folgt für die i -te Komponente des Vektors $\mathbf{z}$ mit (4.1.16)

$$|z_i| \leq \sum_{j=1}^{i-1} \frac{|a_{ij}|}{|a_{ii}|} p_j + \sum_{j=i+1}^n \frac{|a_{ij}|}{|a_{ii}|} = p_i < 1.$$

Hieraus ergibt sich $\|\mathbf{z}\|_\infty < 1$ und damit aufgrund der Kompaktheit des Einheitskreises die Abschätzung

$$\|\mathbf{M}_{GS}\|_\infty = \sup_{\substack{\mathbf{x} \in \mathbb{C}^n \\ \|\mathbf{x}\|_\infty = 1}} \|\mathbf{M}_{GS}\mathbf{x}\|_\infty < 1,$$

wodurch $\rho(\mathbf{M}_{GS}) < 1$ gilt und die Konvergenz des Gauß-Seidel-Verfahrens vorliegt. $\square$

Beispiel 4.17 Wie bereits im Beispiel 4.14 betrachten wir die Matrix $\mathbf{A}$ gemäß Beispiel 1.1. Mit $p_j = 0$ für alle $j \leq 0$ erhalten wir

$$p_1 = \frac{1}{4} + \frac{1}{4} < 1,$$

$$p_i \leq \frac{1}{4}p_{i-1} + \frac{1}{4} + \frac{1}{4} < 1 \text{ für } i = 2, \dots, N,$$

$$p_i \leq \frac{1}{4}p_{i-N} + \frac{1}{4}p_{i-1} + \frac{1}{4} + \frac{1}{4} < 1 \text{ für } i = N+1, \dots, N^2,$$

und folglich die Konvergenz des Gauß-Seidel-Verfahrens.

Strikt diagonaldominante Matrizen erfüllen wegen $\max_{i=1, \dots, n} \sum_{\substack{j=1 \\ j \neq i}}^n \frac{|a_{ij}|}{|a_{ii}|} < 1$ inhärent die Bedingungen des Satzes 4.16, wodurch sich das folgende Korollar ergibt:

Korollar 4.18 *Sei die reguläre Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$ strikt diagonaldominant, dann konvergiert das Gauß-Seidel-Verfahren bei beliebigem Startvektor $\mathbf{x}_0$ und für jede beliebige rechte Seite $\mathbf{b}$ gegen $\mathbf{A}^{-1}\mathbf{b}$.*

Beispiel 4.19 Die Matrix

$$\mathbf{A} = \begin{pmatrix} 0.7 & -0.4 \\ -0.2 & 0.5 \end{pmatrix}$$

ist strikt diagonaldominant, wodurch die Konvergenz des Gauß-Seidel-Verfahrens sichergestellt ist. Die zugehörige Iterationsmatrix

$$\mathbf{M}_{GS} = -(\mathbf{D} + \mathbf{L})^{-1}\mathbf{R} = \begin{pmatrix} 0 & \frac{4}{7} \\ 0 & \frac{8}{35} \end{pmatrix}$$

weist die Eigenwerte $\lambda_1 = 0$ und $\lambda_2 = \frac{8}{35}$ auf, so daß

$$\rho(\mathbf{M}_{GS}) = \rho(\mathbf{M}_J)^2 = \frac{8}{35} \approx 0.22857$$

gilt und mit etwa doppelt so schneller Konvergenz wie beim Jacobi-Verfahren gerechnet werden darf. Für den Startvektor $\mathbf{x}_0 = (21, -19)^T$ und die rechte Seite $\mathbf{b} = (0.3, 0.3)^T$ erhalten wir diese Erwartung mit dem in der folgenden Tabelle aufgelisteten Konvergenzverlauf bestätigt.

Gauß-Seidel-Verfahren				
m	$x_{m,1}$	$x_{m,2}$	$\varepsilon_m := \ \mathbf{x}_m - \mathbf{A}^{-1}\mathbf{b}\ _\infty$	$\varepsilon_m/\varepsilon_{m-1}$
0	2.100000e+01	-1.900000e+01	2.000000e+01	
5	9.688054e-01	9.875222e-01	3.119462e-02	2.285714e-01
10	9.999805e-01	9.999922e-01	1.946209e-05	2.285714e-01
15	1.000000e-00	1.000000e-00	1.214225e-08	2.285714e-01
20	1.000000e-00	1.000000e-00	7.575385e-12	2.285702e-01
25	1.000000e-00	1.000000e-00	4.551914e-15	2.204301e-01

4.1.3 Relaxationsverfahren

Wir schreiben das lineare Iterationsverfahren

$$\mathbf{x}_{m+1} = \mathbf{B}^{-1}(\mathbf{B} - \mathbf{A})\mathbf{x}_m + \mathbf{B}^{-1}\mathbf{b}$$

in der Form

$$\mathbf{x}_{m+1} = \mathbf{x}_m + \underbrace{\mathbf{B}^{-1}(\mathbf{b} - \mathbf{A}\mathbf{x}_m)}_{\mathbf{r}_m :=}. \quad (4.1.17)$$

Somit kann $\mathbf{x}_{m+1}$ als Korrektur von $\mathbf{x}_m$ unter Verwendung des Vektors $\mathbf{r}_m$ interpretiert werden.

Beschränken wir unsere Betrachtungen zunächst auf Verfahren, die einen Gesamtschrittcharakter aufweisen, so liegt das Ziel der Relaxationsverfahren in einer Verbesserung der Konvergenzgeschwindigkeit der Methode (4.1.17) durch Gewichtung des Korrekturvektors $\mathbf{r}_m$. Wir modifizieren (4.1.17) zu

$$\mathbf{x}_{m+1} = \mathbf{x}_m + \omega \mathbf{B}^{-1}(\mathbf{b} - \mathbf{A}\mathbf{x}_m)$$

mit $\omega \in \mathbb{R}^+$. Ausgehend von $\mathbf{x}_m$ suchen wir das optimale $\mathbf{x}_{m+1}$ in Richtung $\mathbf{r}_m$.

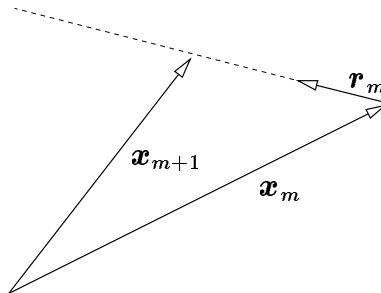


Bild 4.1 Berechnung der Iterierten beim Relaxationsverfahren

Optimal bedeutet im obigen Sinne, daß der Spektralradius der Iterationsmatrix minimal wird. Mit

$$\begin{aligned}
\mathbf{x}_{m+1} &= \mathbf{x}_m + \omega \mathbf{B}^{-1}(\mathbf{b} - \mathbf{A}\mathbf{x}_m) \\
&= \underbrace{(\mathbf{I} - \omega \mathbf{B}^{-1} \mathbf{A})}_{\mathbf{M}(\omega) :=} \mathbf{x}_m + \underbrace{\omega \mathbf{B}^{-1}}_{\mathbf{N}(\omega) :=} \mathbf{b}
\end{aligned} \tag{4.1.18}$$

muß $\omega \in \mathbb{R}^+$ folglich derart bestimmt werden, daß $\rho(\mathbf{M}(\omega))$ minimal wird, das heißt

$$\omega = \arg \min_{\alpha \in \mathbb{R}^+} \rho(\mathbf{M}(\alpha)).$$

Der Gewichtungsfaktor ω heißt Relaxationsparameter, und die Methode (4.1.18) wird für $\omega < 1$ als Unterrelaxationsverfahren und für $\omega > 1$ als Überrelaxationsverfahren bezeichnet.

Betrachten wir nun als zugrundeliegende Verfahren diejenigen, die auf einem Einzelschrittansatz beruhen. Für diese Algorithmen ist die erwähnte Vorgehensweise zwar ebenfalls durchführbar, jedoch unüblich. Die Relaxation wird bei diesen Methoden bei der Iteration jeder Einzelkomponente berücksichtigt. Die genaue Vorgehensweise wird am Beispiel des Gauß-Seidel-Verfahrens im Abschnitt 4.1.3.2 erläutert.

4.1.3.1 Jacobi-Relaxationsverfahren

Gemäß (4.1.17) schreiben wir das Jacobi-Verfahren in der Form

$$\mathbf{x}_{m+1} = \mathbf{x}_m + \mathbf{D}^{-1}(\mathbf{b} - \mathbf{A}\mathbf{x}_m) \text{ für } m = 0, 1, \dots$$

Somit besitzt das Jacobi-Relaxationsverfahren die Darstellung

$$\begin{aligned}
\mathbf{x}_{m+1} &= \mathbf{x}_m + \omega \mathbf{D}^{-1}(\mathbf{b} - \mathbf{A}\mathbf{x}_m) \\
&= \underbrace{(\mathbf{I} - \omega \mathbf{D}^{-1} \mathbf{A})}_{\mathbf{M}_J(\omega) :=} \mathbf{x}_m + \underbrace{\omega \mathbf{D}^{-1}}_{\mathbf{N}_J(\omega) :=} \mathbf{b} \text{ für } m = 0, 1, \dots,
\end{aligned}$$

und wir erhalten die Komponentenschreibweise

$$\begin{aligned}
x_{m+1,i} &= x_{m,i} + \frac{\omega}{a_{ii}} \left(b_i - \sum_{j=1}^n a_{ij} x_{m,j} \right) \\
&= (1 - \omega) x_{m,i} + \frac{\omega}{a_{ii}} \left(b_i - \sum_{\substack{j=1 \\ j \neq i}}^n a_{ij} x_{m,j} \right) \text{ für } i = 1, \dots, n \text{ und } m = 0, 1, \dots
\end{aligned}$$

Satz 4.20 Die Iterationsmatrix des Jacobi-Verfahrens $\mathbf{M}_J$ habe nur reelle Eigenwerte $\lambda_1 \leq \dots \leq \lambda_n$ mit den zugehörigen linear unabhängigen Eigenvektoren $\mathbf{u}_1, \dots, \mathbf{u}_n$, und es gelte $\rho(\mathbf{M}_J) < 1$. Dann besitzt die Iterationsmatrix $\mathbf{M}_J(\omega)$ des Jacobi-Relaxationsverfahrens die Eigenwerte

$$\mu_i = 1 - \omega + \omega \lambda_i \text{ für } i = 1, \dots, n,$$

und es gilt

$$\omega_{opt} = \arg \min_{\omega \in \mathbb{R}^+} \rho(\mathbf{M}_J(\omega)) = \frac{2}{2 - \lambda_1 - \lambda_n}.$$

Beweis:

Mit

$$-D^{-1}(L + R)u_i = M_J u_i = \lambda_i u_i \text{ für } i = 1, \dots, n$$

erhalten wir

$$\begin{aligned} M_J(\omega)u_i &= (I - \omega D^{-1}A)u_i \\ &= ((1 - \omega)I - \omega D^{-1}(L + R))u_i \\ &= (1 - \omega + \omega \lambda_i)u_i \text{ für } i = 1, \dots, n. \end{aligned}$$

Da die Eigenvektoren $u_1, \dots, u_n$ linear unabhängig sind, existieren keine weiteren Eigenwerte. Für die Eigenwerte $\mu_i(\omega) = 1 - \omega + \omega \lambda_i$, ($i = 1, \dots, n$) der Iterationsmatrix des Jacobi-Relaxationsverfahrens gilt mit $\omega \geq 0$

$$\mu_1(\omega) \leq \dots \leq \mu_n(\omega).$$

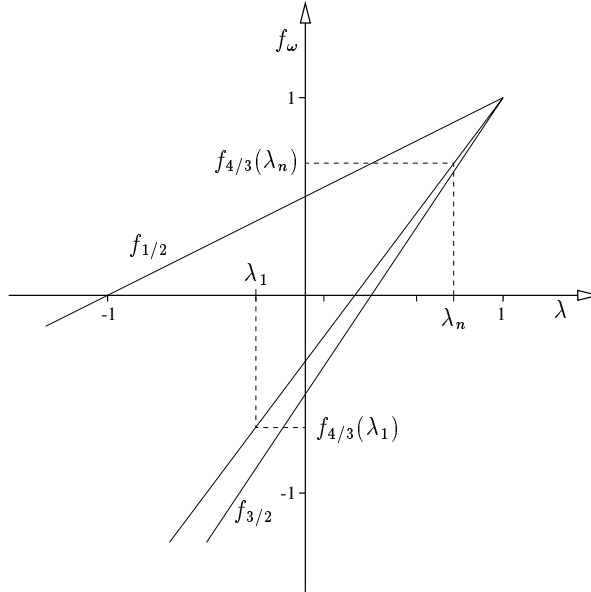


Bild 4.2 Verlauf der Relaxationsfunktionen f_ω in Abhängigkeit vom gewählten Relaxationsparameter ω

Betrachten wir die in der Abbildung 4.2 für verschiedene Relaxationsparameter ω dargestellte Relaxationsfunktion

$$\begin{aligned} f_\omega : \mathbb{R} &\rightarrow \mathbb{R} \\ \lambda &\xrightarrow{f_\omega} f_\omega(\lambda) = 1 - \omega + \omega \lambda, \end{aligned}$$

so liegt die Idee nahe, den optimalen Relaxationsparameter durch die Bedingung

$$\mu_n(\omega^*) = f_{\omega^*}(\lambda_n) = -f_{\omega^*}(\lambda_1) = -\mu_1(\omega^*)$$

zu bestimmen. Hierdurch erhalten wir aus

$$\mu_n(\omega^*) = 1 - \omega^* + \omega^* \lambda_n = -(1 - \omega^* + \omega^* \lambda_1) = -\mu_1(\omega^*)$$

unter Verwendung von $\rho(\mathbf{M}_j) < 1$ die Darstellung

$$\omega^* = \frac{2}{2 - \lambda_1 - \lambda_n} > 0.$$

Berücksichtigen wir $f_\omega(1) = 1$ für alle $\omega \in \mathbb{R}^+$, so gilt stets $\mu_n(\omega^*) = -\mu_1(\omega^*) \geq 0$. Zum Nachweis der Optimalität sei $\omega > \omega^* > 0$. Somit folgt

$$\mu_1(\omega) = 1 - \omega \underbrace{(1 - \lambda_1)}_{>0} < 1 - \omega^*(1 - \lambda_1) = \mu_1(\omega^*) \leq 0,$$

wodurch sich

$$\rho(\mathbf{M}_J(\omega)) \geq |\mu_1(\omega)| > |\mu_1(\omega^*)| = \rho(\mathbf{M}_J(\omega^*))$$

ergibt. Eine analoge Aussage erhalten wir im Fall $\omega < \omega^*$ durch Betrachtung der Eigenwerte $\mu_n(\omega)$ und $\mu_n(\omega^*)$, so daß

$$\omega^* = \omega_{\text{opt}} = \arg \min_{\omega \in \mathbb{R}^+} \rho(\mathbf{M}_J(\omega))$$

folgt. □

Gleichungssysteme, bei denen die Iterationsmatrix des Jacobi-Verfahrens eine um den Nullpunkt symmetrische Eigenwertverteilung besitzt, liefern aufgrund des Satzes 4.20 den optimalen Relaxationsparameter $\omega_{\text{opt}} = 1$. Folglich kann mit dem Relaxationsansatz in solchen Fällen keine Beschleunigung des zugrundeliegenden Gesamtschrittverfahrens erzielt werden. Die betrachtete Modellgleichung (4.1.2) weist diese Eigenschaft auf, weshalb auf eine Diskussion der Methode in Bezug auf diesen Sachverhalt verzichtet wird.

4.1.3.2 Gauß-Seidel-Relaxationsverfahren

Wir betrachten die Komponentenschreibweise des Gauß-Seidel-Verfahrens (4.1.15) mit gewichteter Korrekturvektorkomponente $r_{m,i}$ in der Form

$$\begin{aligned} x_{m+1,i} &= x_{m,i} + \frac{\omega}{a_{ii}} \left(b_i - \sum_{j=1}^{i-1} a_{ij} x_{m+1,j} - \sum_{j=i}^n a_{ij} x_{m,j} \right) \\ &= (1 - \omega) x_{m,i} + \frac{\omega}{a_{ii}} \left(b_i - \sum_{j=1}^{i-1} a_{ij} x_{m+1,j} - \sum_{j=i+1}^n a_{ij} x_{m,j} \right) \end{aligned}$$

für $i = 1, \dots, n$ und $m = 0, 1, \dots$. Hieraus erhalten wir

$$(\mathbf{I} + \omega \mathbf{D}^{-1} \mathbf{L}) \mathbf{x}_{m+1} = [(1 - \omega) \mathbf{I} - \omega \mathbf{D}^{-1} \mathbf{R}] \mathbf{x}_m + \omega \mathbf{D}^{-1} \mathbf{b},$$

wodurch

$$\mathbf{D}^{-1}(\mathbf{D} + \omega \mathbf{L}) \mathbf{x}_{m+1} = \mathbf{D}^{-1}[(1 - \omega) \mathbf{D} - \omega \mathbf{R}] \mathbf{x}_m + \omega \mathbf{D}^{-1} \mathbf{b}$$

und somit das Gauß-Seidel-Relaxationsverfahren in der Darstellung

$$\mathbf{x}_{m+1} = \underbrace{(\mathbf{D} + \omega \mathbf{L})^{-1} [(1 - \omega) \mathbf{D} - \omega \mathbf{R}]}_{\mathbf{M}_{GS}(\omega) :=} \mathbf{x}_m + \underbrace{\omega (\mathbf{D} + \omega \mathbf{L})^{-1} \mathbf{b}}_{\mathbf{N}_{GS}(\omega) :=}$$

folgt.

Satz 4.21 Sei $\mathbf{A} \in \mathbb{C}^{n \times n}$ mit $a_{ii} \neq 0$ für $i = 1, \dots, n$, dann gilt für $\omega \in \mathbb{R}$

$$\rho(\mathbf{M}_{GS}(\omega)) \geq |\omega - 1|.$$

Beweis:

Seien $\lambda_1, \dots, \lambda_n$ die Eigenwerte von $\mathbf{M}_{GS}(\omega)$, dann folgt

$$\begin{aligned} \prod_{i=1}^n \lambda_i &= \det \mathbf{M}_{GS}(\omega) = \det((\mathbf{D} + \omega \mathbf{L})^{-1}) \det((1 - \omega) \mathbf{D} - \omega \mathbf{R}) \\ &= \det(\mathbf{D}^{-1}) \det((1 - \omega) \mathbf{D}) \\ &= \det(\mathbf{D})^{-1} (1 - \omega)^n \det \mathbf{D} \\ &= (1 - \omega)^n. \end{aligned}$$

Hiermit ergibt sich

$$\rho(\mathbf{M}_{GS}(\omega)) = \max_{i=1, \dots, n} |\lambda_i| \geq |1 - \omega|.$$

□

Der obige Satz besagt, daß eine Wahl des Relaxationsparameters $\omega \leq 0$ stets zu einem divergenten Verfahren führt, so daß die beim Relaxationsverfahren zunächst willkürlich geforderte Positivität des Relaxationsparameters an dieser Stelle im Kontext des Gauß-Seidel-Relaxationsverfahrens seine Begründung findet. Analog beinhaltet der Satz die Forderung $\omega < 2$. Beide Bedingungen fassen wir in dem folgenden Korollar zusammen.

Korollar 4.22 Das Gauß-Seidel-Relaxationsverfahren konvergiert höchstens für einen Relaxationsparameter $\omega \in (0, 2)$.

Satz 4.23 Sei $\mathbf{A}$ hermitesch und positiv definit, dann konvergiert das Gauß-Seidel-Relaxationsverfahren genau dann, wenn $\omega \in (0, 2)$ ist.

Beweis:

Aus der positiven Definitheit der Matrix $\mathbf{A}$ folgt $a_{ii} \in \mathbb{R}^+$ für $i = 1, \dots, n$, wodurch sich die Wohldefiniertheit des Gauß-Seidel-Relaxationsverfahrens ergibt.

„ $\Rightarrow$ “ Das Gauß-Seidel-Relaxationsverfahren sei konvergent.

In diesem Fall ergibt sich $\omega \in (0, 2)$ unmittelbar aus Korollar 4.22.

„ $\Leftarrow$ “ Gelte $\omega \in (0, 2)$.

Sei λ Eigenwert von $\mathbf{M}_{GS}(\omega)$ zum Eigenvektor $\mathbf{x} \in \mathbb{C}^n$. Da $\mathbf{A}$ hermitesch ist folgt $\mathbf{L}^* = \mathbf{R}$ und somit

$$((1 - \omega) \mathbf{D} - \omega \mathbf{L}^*) \mathbf{x} = \lambda (\mathbf{D} + \omega \mathbf{L}) \mathbf{x}.$$

Mit

$$\begin{aligned} 2[(1 - \omega) \mathbf{D} - \omega \mathbf{L}^*] &= (2 - \omega) \mathbf{D} + \omega(-\mathbf{D} - 2\mathbf{L}^*) \\ &= (2 - \omega) \mathbf{D} - \omega \mathbf{A} + \omega(\mathbf{L} - \mathbf{L}^*) \end{aligned}$$

und

$$\begin{aligned} 2(\mathbf{D} + \omega \mathbf{L}) &= (2 - \omega) \mathbf{D} + \omega(\mathbf{D} + 2\mathbf{L}) \\ &= (2 - \omega) \mathbf{D} + \omega \mathbf{A} + \omega(\mathbf{L} - \mathbf{L}^*) \end{aligned}$$

ergibt sich für den Eigenvektor $\mathbf{x} \in \mathbb{C}^n$

$$\begin{aligned}
 \lambda((2-\omega)\mathbf{x}^* \mathbf{D} \mathbf{x} + \omega \mathbf{x}^* \mathbf{A} \mathbf{x} + \omega \mathbf{x}^* (\mathbf{L} - \mathbf{L}^*) \mathbf{x}) \\
 &= 2\lambda \mathbf{x}^* (\mathbf{D} + \omega \mathbf{L}) \mathbf{x} \\
 &= 2\mathbf{x}^* [(1-\omega)\mathbf{D} - \omega \mathbf{L}^*] \mathbf{x} \\
 &= (2-\omega)\mathbf{x}^* \mathbf{D} \mathbf{x} - \omega \mathbf{x}^* \mathbf{A} \mathbf{x} + \omega \mathbf{x}^* (\mathbf{L} - \mathbf{L}^*) \mathbf{x}.
 \end{aligned}$$

Unter Verwendung der imaginären Einheit i schreiben wir

$$\mathbf{x}^* (\mathbf{L} - \mathbf{L}^*) \mathbf{x} = \mathbf{x}^* \mathbf{L} \mathbf{x} - \mathbf{x}^* \mathbf{L}^* \mathbf{x} = \mathbf{x}^* \mathbf{L} \mathbf{x} - \overline{\mathbf{x}^* \mathbf{L} \mathbf{x}} = i \cdot s, \quad s \in \mathbb{R}.$$

Zudem gilt

$$\begin{aligned}
 d &:= \mathbf{x}^* \mathbf{D} \mathbf{x} \in \mathbb{R}^+, \\
 a &:= \mathbf{x}^* \mathbf{A} \mathbf{x} \in \mathbb{R}^+,
 \end{aligned}$$

so daß

$$\lambda((2-\omega)d + \omega a + i\omega s) = (2-\omega)d - \omega a + i\omega s$$

folgt. Division durch ω und Einsetzen von $\mu = \frac{2-\omega}{\omega}$ liefert

$$\lambda(\mu d + a + is) = \mu d - a + is.$$

Aus der Voraussetzung $\omega \in (0,2)$ erhalten wir $\mu \in \mathbb{R}^+$, so daß mit $a, d \in \mathbb{R}^+$ die Ungleichung $|\mu d + is - a| < |\mu d + is - (-a)|$ und daher die Abschätzung

$$|\lambda| = \frac{|\mu d + is - a|}{|\mu d + is + a|} < 1$$

folgt. Da der Eigenwert beliebig aus dem Spektrum der Iterationsmatrix gewählt wurde, erhalten wir die für die Konvergenz des Verfahrens notwendige und hinreichende Bedingung

$$\rho(\mathbf{M}_{GS}(\omega)) < 1.$$

□

Zur Bestimmung des optimalen Relaxationsparameters erweisen sich konsistent geordnete Matrizen als vorteilhaft, die daher mit der folgenden Definition eingeführt werden.

Definition 4.24 Seien $\mathbf{L} \in \mathbb{C}^{n \times n}$ eine strikte linke untere und $\mathbf{R} \in \mathbb{C}^{n \times n}$ eine strikte rechte obere Dreiecksmatrix, dann heißt die Matrix $\mathbf{A} = \mathbf{D} + \mathbf{L} + \mathbf{R} \in \mathbb{C}^{n \times n}$ mit regulärem Diagonalanteil $\mathbf{D}$ konsistent geordnet, falls die Eigenwerte von

$$\mathbf{C}(\alpha) = -(\alpha \mathbf{D}^{-1} \mathbf{L} + \alpha^{-1} \mathbf{D}^{-1} \mathbf{R}) \text{ mit } \alpha \in \mathbb{C} \setminus \{0\}$$

unabhängig von α sind.

Beispiel 4.25 Jede Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$ der Form

$$\mathbf{A} = \begin{pmatrix} \mathbf{I} & \mathbf{A}_{12} \\ \mathbf{A}_{21} & \mathbf{I} \end{pmatrix}$$

ist konsistent geordnet. Es gilt

$$L = \begin{pmatrix} \mathbf{0} & \mathbf{0} \\ \mathbf{A}_{21} & \mathbf{0} \end{pmatrix} \quad \text{und} \quad R = \begin{pmatrix} \mathbf{0} & \mathbf{A}_{12} \\ \mathbf{0} & \mathbf{0} \end{pmatrix},$$

so daß $C(\alpha)$ die Darstellung

$$C(\alpha) = \begin{pmatrix} \mathbf{0} & -\alpha^{-1}\mathbf{A}_{12} \\ -\alpha\mathbf{A}_{21} & \mathbf{0} \end{pmatrix} = \begin{pmatrix} \mathbf{I} & \mathbf{0} \\ \mathbf{0} & -\alpha\mathbf{I} \end{pmatrix} \begin{pmatrix} \mathbf{0} & \mathbf{A}_{12} \\ \mathbf{A}_{21} & \mathbf{0} \end{pmatrix} \begin{pmatrix} \mathbf{I} & \mathbf{0} \\ \mathbf{0} & -\alpha^{-1}\mathbf{I} \end{pmatrix}$$

besitzt. Folglich ist λ genau dann Eigenwert von $C(\alpha)$, wenn λ Eigenwert von $C(1) = \mathbf{A} - \mathbf{I}$ ist, wodurch sich die geforderte Unabhängigkeit der Eigenwerte von der gewählten Größe $\alpha \in \mathbb{C} \setminus \{0\}$ ergibt.

Beispiel 4.26 Jede Tridiagonalmatrix

$$A = \begin{pmatrix} a_1 & b_1 & & & \\ c_2 & a_2 & b_2 & & \\ & \ddots & \ddots & \ddots & \\ & & \ddots & \ddots & b_{n-1} \\ & & & c_n & a_n \end{pmatrix} \in \mathbb{C}^{n \times n}$$

mit $a_i \neq 0$ für $i = 1, \dots, n$ ist konsistent geordnet, denn es gilt

$$C(1) = -D^{-1}(L + R) = \begin{pmatrix} 0 & d_1 & & & \\ \ell_2 & \ddots & \ddots & & \\ & \ddots & \ddots & d_{n-1} & \\ & & \ell_n & 0 & \end{pmatrix}$$

mit $d_i = -\frac{b_i}{a_i}$ für $i = 1, \dots, n-1$ und $\ell_i = -\frac{c_i}{a_i}$ für $i = 2, \dots, n$, so daß mit der regulären Matrix

$$S(\alpha) = \begin{pmatrix} 1 & & & & \\ & \alpha & & & \\ & & \alpha^2 & & \\ & & & \ddots & \\ & & & & \alpha^{n-1} \end{pmatrix}, \quad \alpha \in \mathbb{C} \setminus \{0\}$$

die Gleichung $S(\alpha)C(1)S^{-1}(\alpha) = C(\alpha)$ folgt.

Satz 4.27 Seien $A \in \mathbb{C}^{n \times n}$ konsistent geordnet und $\omega \in (0, 2)$, dann ist $\mu \in \mathbb{C} \setminus \{0\}$ genau dann Eigenwert von $M_{GS}(\omega)$, wenn

$$\lambda = \frac{\mu + \omega - 1}{\omega\mu^{1/2}}$$

Eigenwert von M_J ist.

Beweis:

Sei $\mu \in \mathbb{C} \setminus \{0\}$, dann gilt

$$\begin{aligned}
 & (\mathbf{I} + \omega \mathbf{D}^{-1} \mathbf{L})(\mu \mathbf{I} - \mathbf{M}_{GS}(\omega)) \\
 &= \mu(\mathbf{I} + \omega \mathbf{D}^{-1} \mathbf{L}) - \mathbf{D}^{-1}(\mathbf{D} + \omega \mathbf{L}) \underbrace{(\mathbf{D} + \omega \mathbf{L})^{-1}((1 - \omega)\mathbf{D} - \omega \mathbf{R})}_{\mathbf{M}_{GS}(\omega)=} \\
 &= (\mu - (1 - \omega))\mathbf{I} + \omega \mathbf{D}^{-1}(\mu \mathbf{L} + \mathbf{R}) \\
 &= (\mu - (1 - \omega))\mathbf{I} + \omega \mu^{1/2} \mathbf{D}^{-1}(\mu^{1/2} \mathbf{L} + \mu^{-1/2} \mathbf{R}).
 \end{aligned}$$

Mit $\det(\mathbf{I} + \omega \mathbf{D}^{-1} \mathbf{L}) = 1$ ist aufgrund der obigen Gleichung $\mu \in \mathbb{C} \setminus \{0\}$ genau dann Eigenwert von $\mathbf{M}_{GS}(\omega)$, wenn

$$\det\left((\mu - (1 - \omega))\mathbf{I} + \omega \mu^{1/2} \mathbf{D}^{-1}(\mu^{1/2} \mathbf{L} + \mu^{-1/2} \mathbf{R})\right) = 0$$

gilt, das heißt

$$\frac{\mu - (1 - \omega)}{\omega \mu^{1/2}}$$

Eigenwert von $-\mathbf{D}^{-1}(\mu^{1/2} \mathbf{L} + \mu^{-1/2} \mathbf{R})$ ist. Mit der Voraussetzung, daß die Matrix $\mathbf{A}$ konsistent geordnet ist, stimmen die Eigenwerte der beiden Matrizen $-\mathbf{D}^{-1}(\mu^{1/2} \mathbf{L} + \mu^{-1/2} \mathbf{R})$ und $-\mathbf{D}^{-1}(\mathbf{L} + \mathbf{R}) = \mathbf{M}_J$ überein, wodurch die Behauptung des Satzes vorliegt. $\square$

Für konsistent geordnete Matrizen $\mathbf{A} \in \mathbb{C}^{n \times n}$ gilt somit

$$\rho(\mathbf{M}_{GS}) = \rho(\mathbf{M}_J)^2,$$

wodurch das Gauß-Seidel-Verfahren verglichen mit der Jacobi-Methode in der Regel die Hälfte an Iterationen benötigt, um eine vorgegebene Genauigkeitschranke zu erreichen. Desweiteren erhalten wir im Fall einer konsistent geordneten Matrix $\mathbf{A}$ stets

$$\sigma(\mathbf{M}_J) = -\sigma(\mathbf{M}_J),$$

wodurch laut Satz 4.20 für derartige Matrizen keine Konvergenzbeschleunigung mittels einer Relaxation des Jacobi-Verfahrens erzielt werden kann.

Satz 4.28 Sei $\mathbf{A} \in \mathbb{C}^{n \times n}$ konsistent geordnet. Die Eigenwerte von $\mathbf{M}_J$ seien reell, und es gelte

$$\rho := \rho(\mathbf{M}_J) < 1.$$

Dann gilt

- (a) das Gauß-Seidel-Relaxationsverfahren konvergiert für alle $\omega \in (0, 2)$.
- (b) der Spektralradius der Iterationsmatrix $\mathbf{M}_{GS}(\omega)$ wird minimal für

$$\omega_{opt} = \frac{2}{1 + \sqrt{1 - \rho^2}},$$

womit

$$\rho(\mathbf{M}_{GS}(\omega_{opt})) = \omega_{opt} - 1 = \frac{1 - \sqrt{1 - \rho^2}}{1 + \sqrt{1 - \rho^2}}$$

vorliegt.

Beweis:

Seien $\lambda_1, \dots, \lambda_n \in \mathbb{R}$ Eigenwerte von $\mathbf{M}_J$, dann ist mit Satz 4.27 μ genau dann Eigenwert von $\mathbf{M}_{GS}(\omega)$, wenn

$$\lambda = \frac{\mu + \omega - 1}{\omega \mu^{1/2}} \in \{\lambda_1, \dots, \lambda_n\} \quad (4.1.19)$$

gilt. Da die Matrix $\mathbf{A}$ konsistent geordnet ist, ist mit $\lambda \in \mathbb{R}$ auch $-\lambda$ Eigenwert von $\mathbf{M}_J$, wodurch das Vorzeichen in (4.1.19) keine Bedeutung besitzt. Wir können daher die Gleichung

$$\lambda^2 \omega^2 \mu = (\mu + \omega - 1)^2 \quad (4.1.20)$$

betrachten und o.B.d.A. $\lambda \geq 0$ voraussetzen.

Aus $\rho(\mathbf{M}_J) < 1$ folgt somit $\lambda \in [0, 1)$. Desweiteren betrachten wir aufgrund des Korollars 4.22 stets $\omega \in (0, 2)$. Aus (4.1.20) erhalten wir die zwei Eigenwerte in der Form

$$\mu^\pm = \mu^\pm(\omega, \lambda) = \frac{1}{2} \lambda^2 \omega^2 - (\omega - 1) \pm \lambda \omega \sqrt{\frac{1}{4} \lambda^2 \omega^2 - (\omega - 1)}. \quad (4.1.21)$$

Wir definieren

$$g(\omega, \lambda) := \frac{1}{4} \lambda^2 \omega^2 - (\omega - 1).$$

Für gegebenes $\lambda \in [0, 1)$ lauten die Nullstellen dieser Funktion

$$\omega^\pm = \omega^\pm(\lambda) = \frac{2}{1 \pm \sqrt{1 - \lambda^2}}. \quad (4.1.22)$$

Mit $\omega \in (0, 2)$ können wir ω^- vernachlässigen und es ergibt sich $\omega^+(\lambda) > 1$ für alle $\lambda \in [0, 1)$. Zudem gilt für alle $\lambda \in [0, 1)$ und $\omega \in (0, 2)$

$$\frac{\partial g}{\partial \omega}(\omega, \lambda) = \frac{1}{2} \lambda^2 \omega - 1 < 0.$$

Wir erhalten die folgenden drei Fälle:

(1) $2 > \omega > \omega^+(\lambda)$:

Die beiden Eigenwerte $\mu^+(\omega, \lambda)$ und $\mu^-(\omega, \lambda)$ sind komplex, und es gilt

$$|\mu^+(\omega, \lambda)| = |\mu^-(\omega, \lambda)| = |\omega - 1| = \omega - 1$$

(2) $\omega = \omega^+(\lambda)$:

Aus (4.1.22) folgt $\lambda^2 = \frac{4}{\omega} - \frac{4}{\omega^2}$, wodurch sich

$$|\mu^+(\omega, \lambda)| = |\mu^-(\omega, \lambda)| = \frac{1}{2} \lambda^2 \omega^2 - (\omega - 1) = 2\omega - 2 - (\omega - 1) = \omega - 1$$

ergibt.

(3) $0 < \omega < \omega^+(\lambda)$:

Die Gleichung (4.1.21) liefert zwei reelle Eigenwerte

$$\mu^\pm(\omega, \lambda) = \underbrace{\frac{1}{2}\lambda^2\omega^2 - (\omega - 1)}_{>0} \pm \underbrace{\lambda\omega\sqrt{\frac{1}{4}\lambda^2\omega^2 - (\omega - 1)}}_{\geq 0}$$

mit

$$\max\{|\mu^+(\omega, \lambda)|, |\mu^-(\omega, \lambda)|\} = \mu^+(\omega, \lambda).$$

Zur Bestimmung von $\rho(\mathbf{M}_{GS}(\omega))$ sind wir in allen drei Fällen nur an $\mu^+(\omega, \lambda)$ interessiert. Damit betrachten wir für $\lambda \in [0, 1)$

$$\mu(\omega, \lambda) = \begin{cases} \mu^+(\omega, \lambda) & \text{für } 0 < \omega < \omega^+(\lambda) \\ \omega - 1 & \text{für } \omega^+(\lambda) \leq \omega < 2. \end{cases} \quad (4.1.23)$$

Hiermit gilt für $0 < \omega < \omega^+(\lambda)$ und $\lambda \in [0, 1)$

$$\frac{\partial \mu}{\partial \lambda}(\omega, \lambda) = \underbrace{\lambda\omega^2}_{\geq 0} + \underbrace{\omega\sqrt{\frac{1}{4}\lambda^2\omega^2 - (\omega - 1)}}_{>0} + \underbrace{\lambda\omega\frac{1}{2}\frac{\frac{1}{2}\lambda^2\omega^2}{\sqrt{\frac{1}{4}\lambda^2\omega^2 - (\omega - 1)}}}_{\geq 0} > 0, \quad (4.1.24)$$

und wegen

$$\mu(\omega, \lambda) = \left(\frac{\omega\lambda}{2} + \sqrt{\frac{1}{4}\lambda^2\omega^2 - (\omega - 1)} \right)^2$$

folgt

$$\frac{\partial \mu}{\partial \omega}(\omega, \lambda) = 2 \underbrace{\left(\frac{\omega\lambda}{2} + \sqrt{\frac{1}{4}\lambda^2\omega^2 - (\omega - 1)} \right)}_{>0} \underbrace{\left[\frac{\lambda}{2} + \frac{1}{2} \frac{\frac{1}{2}\lambda^2\omega - 1}{\sqrt{\frac{1}{4}\lambda^2\omega^2 - (\omega - 1)}} \right]}_{q(\omega, \lambda) :=}.$$

Wir schreiben

$$q(\omega, \lambda) = \underbrace{\frac{1}{2\sqrt{\frac{1}{4}\lambda^2\omega^2 - (\omega - 1)}}}_{>0} \left(\underbrace{\lambda\sqrt{\frac{1}{4}\lambda^2\omega^2 - (\omega - 1)}}_{q_1(\omega, \lambda) :=} + \underbrace{\frac{1}{2}\lambda^2\omega - 1}_{q_2(\omega, \lambda) :=} \right).$$

Für die Funktionen q_1 und q_2 gilt hierbei für alle $\lambda \in [0, 1)$ und $\omega \in (0, \omega^+(\lambda))$

$$q_1(\omega, \lambda) \geq 0 \quad \text{und} \quad q_2(\omega, \lambda) < 0.$$

Desweiteren liefert

$$[q_1(\omega, \lambda)]^2 = \frac{\omega^2\lambda^4}{4} + \lambda^2 - \omega\lambda^2 < \frac{\omega^2\lambda^4}{4} + 1 - \omega\lambda^2 = [q_2(\omega, \lambda)]^2$$

die Ungleichung

$$\frac{\partial \mu}{\partial \omega}(\omega, \lambda) < 0 \quad \text{für alle } \lambda \in [0, 1) \quad \text{und} \quad \omega \in (0, \omega^+(\lambda)).$$

Aus (4.1.23) erhalten wir zudem $\mu(0, \lambda) = 1 = \mu(2, \lambda)$, so daß $|\mu(\omega, \lambda)| < 1$ für alle $\lambda \in [0, 1)$ und $\omega \in (0, 2)$ folgt, wodurch sich direkt $\rho(\mathbf{M}_{GS}(\omega)) < 1$ ergibt. Für jeden Eigenwert λ wird $|\mu(\omega, \lambda)|$ minimal für $\omega_{\text{opt}} = \omega^+(\lambda)$.

Gleichung (4.1.23) liefert somit

$$\begin{aligned} \rho(\mathbf{M}_{GS}(\omega_{\text{opt}})) &= |\mu(\omega_{\text{opt}}, \rho(\mathbf{M}_J))| \\ &= \omega_{\text{opt}}(\rho(\mathbf{M}_J)) - 1 \\ &\stackrel{(4.1.22)}{=} \frac{2}{1 + \sqrt{1 - \rho^2}} - 1 \\ &= \frac{1 - \sqrt{1 - \rho^2}}{1 + \sqrt{1 - \rho^2}}. \end{aligned}$$

□

Bemerkung:

Da für alle $\lambda \in (0, 1)$

$$\lim_{\omega \searrow \omega_{\text{opt}}} \frac{\partial \mu}{\partial \omega}(\omega, \lambda) = 1 \quad \text{und} \quad \lim_{\omega \nearrow \omega_{\text{opt}}} \frac{\partial \mu}{\partial \omega}(\omega, \lambda) = -\infty \quad (4.1.25)$$

gilt, sollte ω im Zweifelsfall eher größer ω_{opt} als kleiner ω_{opt} gewählt werden.

Zudem liegt der optimale Relaxationsparameter für die den Voraussetzungen des Satzes 4.28 genügenden Matrizen im Intervall $[1, 2)$, weshalb dieses Relaxationsverfahren auch als SOR-Methode (*successive overrelaxation method*) bezeichnet wird. Bei einem gedämpften Gauß-Seidel-Verfahren ($\omega < 1$) spricht man hingegen auch von der *successive underrelaxation method*.

Beispiel 4.29 Wir betrachten wiederum das Modellproblem $\mathbf{Ax} = \mathbf{b}$ mit

$$\mathbf{A} = \begin{pmatrix} 0.7 & -0.4 \\ -0.2 & 0.5 \end{pmatrix}, \quad \mathbf{b} = \begin{pmatrix} 0.3 \\ 0.3 \end{pmatrix}.$$

Die Matrix $\mathbf{A}$ ist als Tridiagonalmatrix konsistent geordnet, und die Eigenwerte von

$$\mathbf{M}_J = -\mathbf{D}^{-1}(\mathbf{L} + \mathbf{R})$$

sind laut Beispiel 4.15

$$\lambda_{1,2} = \pm \sqrt{\frac{8}{35}} \in \mathbb{R},$$

so daß $\rho(\mathbf{M}_J) < 1$ gilt. Unter Verwendung dieser Eigenschaften liefert der Satz 4.28 die Konvergenz des Gauß-Seidel-Relaxationsverfahrens

$$\mathbf{x}_{m+1} = \mathbf{M}_{GS}(\omega)\mathbf{x}_m + \mathbf{N}_{GS}(\omega)\mathbf{b}$$

für alle $\omega \in (0,2)$. Der optimale Relaxationsparameter lautet

$$\omega_{\text{opt}} = \frac{2}{1 + \sqrt{1 - \frac{8}{35}}} \approx 1.0648$$

und liefert

$$\rho(\mathbf{M}_{GS}(\omega^*)) = \frac{1 - \sqrt{1 - \frac{8}{35}}}{1 + \sqrt{1 - \frac{8}{35}}} \approx 0.0648.$$

Die Abbildung 4.3 zeigt den für $\lambda = \sqrt{\frac{8}{35}}$ aufgetragenen Verlauf des Spektralradius in Abhängigkeit vom Relaxationsparameter. Desweiteren präsentiert die Tabelle den Konvergenzverlauf des Gauß-Seidel-Relaxationsverfahrens mit optimalem Relaxationsparameter beim Modellproblem unter Verwendung des Startvektors $\mathbf{x}_0 = (21, -19)^T$.

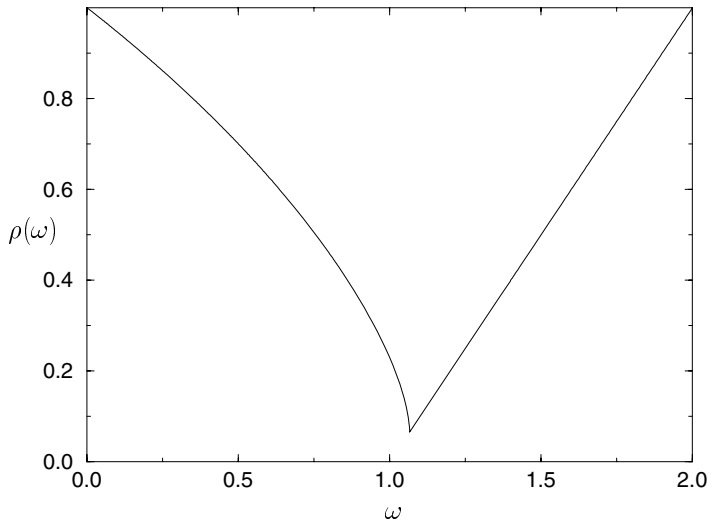


Bild 4.3 Spektralradius in Abhängigkeit vom Relaxationsparameter

SOR-Gauß-Seidel-Verfahren				
m	$x_{m,1}$	$x_{m,2}$	$\varepsilon_m := \ \mathbf{x}_m - \mathbf{A}^{-1}\mathbf{b}\ _\infty$	$\varepsilon_m/\varepsilon_{m-1}$
0	2.100000e+01	-1.900000e+01	2.000000e+01	
5	9.987226e-01	9.997003e-01	1.277401e-03	8.134709e-02
10	1.000000e-00	1.000000e-00	2.942099e-09	7.205638e-02
15	1.000000e-00	1.000000e-00	4.884981e-15	6.727829e-02

Die Abbildung 4.3 verdeutlicht nachdrücklich den Verlauf des Spektralradius in der Nähe des optimalen Relaxationsparameters, wodurch das analytisch ermittelte Verhalten unterstrichen wird.

4.1.4 Richardson-Verfahren

Das Richardson-Verfahren basiert auf dem Grundalgorithmus

$$\mathbf{x}_{m+1} = (\mathbf{I} - \mathbf{A})\mathbf{x}_m + \mathbf{b}$$

(siehe Beispiel 4.10) und einer Gewichtung des Korrekturvektors

$$\mathbf{r}_m = \mathbf{b} - \mathbf{A}\mathbf{x}_m$$

mit einer Zahl $\Theta \in \mathbb{C}$. Wir erhalten

$$\mathbf{x}_{m+1} = \underbrace{(\mathbf{I} - \Theta\mathbf{A})}_{\mathbf{M}_R(\Theta)} \mathbf{x}_m + \underbrace{\Theta\mathbf{I}}_{\mathbf{N}_R(\Theta)} \mathbf{b} \quad (4.1.26)$$

oder in Komponentenschreibweise

$$x_{m+1,i} = x_{m,i} + \Theta(b_i - \sum_{j=1}^n a_{ij}x_{m,j}).$$

Lemma 4.30 Sei $\mathbf{A} \in \mathbb{C}^{n \times n}$ mit $\sigma(\mathbf{A}) \subset \mathbb{R}$ gegeben, und seien $\lambda_{\max} = \max_{\lambda \in \sigma(\mathbf{A})} \lambda$ sowie $\lambda_{\min} = \min_{\lambda \in \sigma(\mathbf{A})} \lambda$, dann gilt für alle $\Theta \in \mathbb{R}$

$$\sigma(\mathbf{M}_R(\Theta)) \subset \mathbb{R}$$

und für alle $\Theta \in \mathbb{C}$

$$\rho(\mathbf{M}_R(\Theta)) = \max\{|1 - \Theta\lambda_{\max}|, |1 - \Theta\lambda_{\min}|\}.$$

Beweis:

Mit $\mathbf{M}_R(\Theta) = \mathbf{I} - \Theta\mathbf{A}$ weisen alle Eigenwerte der Matrix $\mathbf{M}_R(\Theta)$ die Form $\mu = 1 - \Theta\lambda$ mit $\lambda \in \sigma(\mathbf{A})$ auf. Mit $\Theta \in \mathbb{R}$ und $\sigma(\mathbf{A}) \subset \mathbb{R}$ sind alle Eigenwerte von $\mathbf{M}_R(\Theta)$ reellwertig, und es gilt $\sigma(\mathbf{M}_R(\Theta)) \subset \mathbb{R}$.

Für gegebenes $\Theta \in \mathbb{C}$ betrachten wir die durch

$$\begin{aligned} p &: [a, b] \rightarrow \mathbb{R} \\ \lambda &\mapsto p(\lambda) = |1 - \Theta\lambda| \end{aligned}$$

definierte Funktion, die auf $[a, b]$ kein lokales Maximum besitzt. Hiermit gilt

$$\max_{\lambda \in [a, b]} p(\lambda) = \max\{|1 - \Theta a|, |1 - \Theta b|\},$$

und wir erhalten

$$\rho(\mathbf{M}_R(\Theta)) = \max_{\lambda \in \sigma(\mathbf{A})} p(\lambda) = \max\{|1 - \Theta\lambda_{\min}|, |1 - \Theta\lambda_{\max}|\}.$$

□

In Anlehnung an das Gauß-Seidel-Relaxationsverfahren werden wir mit den folgenden zwei Sätzen den Konvergenzbereich der Richardson-Iteration in Abhängigkeit vom Parameter Θ sowie dessen optimalen Wert Θ_{opt} bestimmen.

Satz 4.31 Seien $\mathbf{A} \in \mathbb{C}^{n \times n}$ mit $\sigma(\mathbf{A}) \subset \mathbb{R}^+$ und $\lambda_{\max} = \max_{\lambda \in \sigma(\mathbf{A})} \lambda$, $\lambda_{\min} = \min_{\lambda \in \sigma(\mathbf{A})} \lambda$, dann konvergiert das Richardson-Verfahren (4.1.26) für $\Theta \in \mathbb{R}$ genau dann, wenn

$$0 < \Theta < \frac{2}{\lambda_{\max}}$$

gilt.

Beweis:

„ $\Rightarrow$ “ Das Richardson-Verfahren sei konvergent.

In diesem Fall gilt mit Lemma 4.30

$$1 > \rho(\mathbf{M}_R(\Theta)) \geq |1 - \Theta\lambda_{\max}| \geq 1 - \Theta\lambda_{\max},$$

womit $\Theta\lambda_{\max} > 0$ und damit aufgrund des positiven reellen Spektrums von $\mathbf{A}$ die Ungleichung $\Theta > 0$ folgt. Weiter gilt

$$-1 < -\rho(\mathbf{M}_R(\Theta)) \leq -|1 - \Theta\lambda_{\max}| \leq 1 - \Theta\lambda_{\max},$$

wodurch sich $\Theta\lambda_{\max} < 2$ und mit $\lambda_{\max} > 0$ die Abschätzung $\Theta < \frac{2}{\lambda_{\max}}$ ergibt.

„ $\Leftarrow$ “ Sei $0 < \Theta < \frac{2}{\lambda_{\max}}$.

Unter Verwendung von $\sigma(\mathbf{A}) \subset \mathbb{R}^+$ erhalten wir

$$-1 < 1 - \Theta\lambda_{\max} \leq 1 - \Theta\lambda_{\min} < 1,$$

so daß mit Lemma 4.30

$$\rho(\mathbf{M}_R(\Theta)) = \max\{|1 - \Theta\lambda_{\min}|, |1 - \Theta\lambda_{\max}|\} < 1$$

folgt. □

Satz 4.32 Sei $\mathbf{A} \in \mathbb{C}^{n \times n}$ mit $\sigma(\mathbf{A}) \subset \mathbb{R}^+$, und gelten die Festlegungen $\lambda_{\min} = \min_{\lambda \in \sigma(\mathbf{A})} \lambda$ und $\lambda_{\max} = \max_{\lambda \in \sigma(\mathbf{A})} \lambda$, dann wird der Spektralradius von $\mathbf{M}_R(\Theta)$ minimal für

$$\Theta_{opt} = \frac{2}{\lambda_{\min} + \lambda_{\max}},$$

und es gilt

$$\rho(\mathbf{M}_R(\Theta_{opt})) = \frac{\lambda_{\max} - \lambda_{\min}}{\lambda_{\max} + \lambda_{\min}}.$$

Beweis:

Für $\Theta \in \mathbb{C}$ und $\lambda \in \mathbb{R}^+$ folgt

$$|1 - \Theta\lambda| = (1 - \lambda\operatorname{Re}(\Theta))^2 + (\lambda\operatorname{Im}(\Theta))^2 \geq (1 - \lambda\operatorname{Re}(\Theta))^2,$$

wodurch sich $\Theta_{opt} \in \mathbb{R}$ ergibt. Seien desweiteren die Funktionen $g_{\max}, g_{\min} : \mathbb{R} \rightarrow \mathbb{R}$ durch $g_{\max}(\Theta) = |1 - \Theta\lambda_{\max}|$, $g_{\min}(\Theta) = |1 - \Theta\lambda_{\min}|$ gegeben, so können wir unter Berücksichtigung von

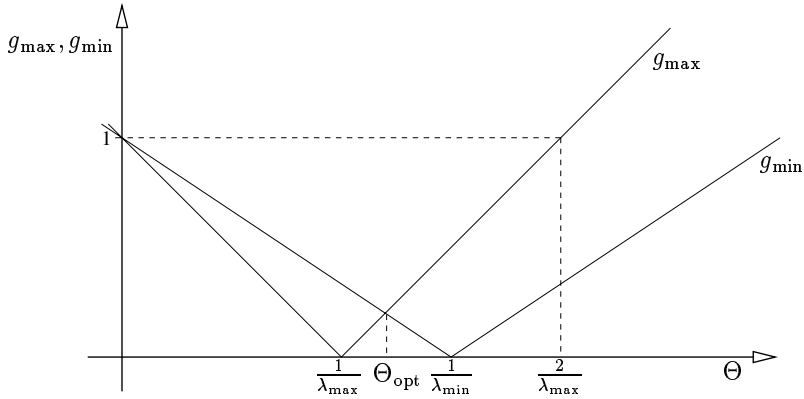


Bild 4.4 Bestimmung des Parameters Θ_{opt} .

$$\Theta_{\text{opt}} = \arg \min_{\Theta \in \mathbb{C}} \rho(\mathbf{M}_R(\Theta)) = \arg \min_{\Theta \in \mathbb{R}} \rho(\mathbf{M}_R(\Theta)) = \arg \min_{\Theta \in \mathbb{R}} \max \{g_{\max}(\Theta), g_{\min}(\Theta)\}$$

der Abbildung 4.4 die Bedingung

$$\Theta_{\text{opt}} \lambda_{\max} - 1 = 1 - \Theta_{\text{opt}} \lambda_{\min}$$

für das optimale Θ entnehmen, wodurch sich

$$\Theta_{\text{opt}} = \frac{2}{\lambda_{\max} + \lambda_{\min}}$$

und

$$\rho(\mathbf{M}_R(\Theta_{\text{opt}})) = \frac{\lambda_{\max} - \lambda_{\min}}{\lambda_{\max} + \lambda_{\min}}$$

ergeben.

□

Beispiel 4.33 Wir betrachten das bereits bekannte Modellproblem $\mathbf{Ax} = \mathbf{b}$ mit

$$\mathbf{A} = \begin{pmatrix} 0.7 & -0.4 \\ -0.2 & 0.5 \end{pmatrix} \quad \text{und} \quad \mathbf{b} = \begin{pmatrix} 0.3 \\ 0.3 \end{pmatrix}.$$

Mit

$$0 = \det(\mathbf{A} - \lambda \mathbf{I}) = (0.7 - \lambda)(0.5 - \lambda) - 0.08 = (\lambda - 0.6)^2 - 0.09$$

erhalten wir die Eigenwerte $\lambda_{1,2} = 0.6 \pm 0.3$ und somit ein positives reelles Spektrum $\sigma(\mathbf{A})$. Satz 4.31 liefert die Konvergenz des Richardson-Verfahrens für alle $\Theta \in \mathbb{R}$ mit $0 < \Theta < \frac{2}{0.9} = 2.\bar{2}$. Mit Satz 4.32 ergeben sich

$$\Theta_{\text{opt}} = \frac{2}{0.9 + 0.3} = \frac{5}{3}$$

und

$$\rho(\mathbf{M}_R(\Theta_{\text{opt}})) = \frac{0.9 - 0.3}{0.9 + 0.3} = 0.5.$$

Für den üblichen Startvektor $\mathbf{x}_0 = (21, -19)^T$ ergibt sich der in der folgenden Tabelle aufgelistete Konvergenzverlauf:

Richardson-Verfahren				
m	$x_{m,1}$	$x_{m,2}$	$\varepsilon_m := \ \mathbf{x}_m - \mathbf{A}^{-1}\mathbf{b}\ _\infty$	$\varepsilon_m/\varepsilon_{m-1}$
0	2.100000e+01	-1.900000e+01	2.000000e+01	
15	9.989827e-01	1.000203e+00	1.017253e-03	8.333333e-01
30	1.000000e+00	1.000000e-00	1.862645e-08	3.000000e-01
45	1.000000e-00	1.000000e+00	9.473533e-13	8.331381e-01
52	1.000000e+00	1.000000e-00	4.662937e-15	3.157895e-01

4.1.5 Symmetrische Splitting-Methoden

Die bei den betrachteten Splitting-Methoden

$$\mathbf{x}_{m+1} = \mathbf{B}^{-1}(\mathbf{B} - \mathbf{A})\mathbf{x}_m + \mathbf{B}^{-1}\mathbf{b}$$

vorliegende Matrix $\mathbf{B}$ kann als leicht invertierbare Approximation der Matrix $\mathbf{A}$ verstanden werden und stellt daher eine mögliche Prädiktionierungsmatrix für weitere iterative Verfahren dar. Setzt ein iteratives Verfahren die Symmetrie und positive Definitheit der Matrix $\mathbf{A}$ voraus, so sollte auch das prädiktionierte System, zum Beispiel

$$\underbrace{\mathbf{B}^{-1/2}\mathbf{A}\mathbf{B}^{-1/2}}_{\tilde{\mathbf{A}}:=} \mathbf{y} = \mathbf{B}^{-1/2}\mathbf{b},$$

mit $\tilde{\mathbf{A}}$ eine symmetrische und positiv definite Matrix vorliegen haben. Um dieses Ziel zu erreichen, sind symmetrische Splitting-Methoden interessant.

Satz 4.34

- (1) Das Jacobi-Relaxationsverfahren ist symmetrisch.
- (2) Das Richardson-Verfahren ist genau symmetrisch, wenn $\Theta \in \mathbb{R}^+$ gilt.
- (3) Das Gauß-Seidel-Relaxationsverfahren ist nicht symmetrisch.

Beweis:

Sei $\mathbf{A} = \mathbf{D} + \mathbf{L} + \mathbf{R}$ eine positiv definite und symmetrische Matrix, dann gilt $a_{ii} > 0$ für $i = 1, \dots, n$, so daß neben der Richardson-Methode auch die Jacobi- und Gauß-Seidel-Verfahren wohldefiniert sind.

- (1) $\mathbf{B}_J(\omega) = \frac{1}{\omega}\mathbf{D} = \frac{1}{\omega}\text{diag}\{a_{11}, \dots, a_{nn}\}$ ist symmetrisch und mit $\omega, a_{ii} \in \mathbb{R}^+$ für $i = 1, \dots, n$ zudem positiv definit.
- (2) Mit $\mathbf{B}_R(\Theta) = \frac{1}{\Theta}\mathbf{I}$ ist das Richardson-Verfahren genau dann eine symmetrische Splitting-Methode, wenn $\Theta \in \mathbb{R}^+$ gilt.
- (3) Das Gauß-Seidel-Relaxationsverfahren ist unabhängig von einer speziellen Wahl des Relaxationsparameters $\omega \in (0, 2)$ keine symmetrische Splitting-Methode, da die Matrix $\mathbf{B}_{GS}(\omega) = \frac{1}{\omega}(\mathbf{D} + \omega\mathbf{L})$ für alle $\mathbf{L} \neq 0$ unsymmetrisch ist.

□

Wir wollen im folgenden am Beispiel des Gauß-Seidel-Relaxationsverfahrens eine Möglichkeit der Symmetrisierung unsymmetrischer Splitting-Methoden vorstellen. Entspreche $\mathbf{A} = \mathbf{D} + \mathbf{L} + \mathbf{R}$ der üblichen Zerlegung der Matrix $\mathbf{A}$, dann definieren wir mit

$$\mathbf{x}_{m+1} = \underbrace{-(\mathbf{D} + \mathbf{R})^{-1}\mathbf{L}}_{\mathbf{M}_{RGS} :=} \mathbf{x}_m + \underbrace{(\mathbf{D} + \mathbf{R})^{-1}\mathbf{b}}_{\mathbf{N}_{RGS} :=} \text{ für } m = 0, 1, \dots$$

das rückwärts durchgeführte Gauß-Seidel-Verfahren. Die Komposition der beiden Gauß-Seidel-Verfahren entspricht der Durchführung zweier aufeinanderfolgender Iterationsvorschriften. Durch eine Zusammenfassung der beiden Einzelschritte zu einem Iterationsschritt erhalten wir mit

$$\mathbf{x}_{m+1/2} = \mathbf{M}_{GS}\mathbf{x}_m + \mathbf{N}_{GS}\mathbf{b}$$

die Darstellung

$$\mathbf{x}_{m+1} = \mathbf{M}_{RGS}\mathbf{x}_{m+1/2} + \mathbf{N}_{RGS}\mathbf{b} = \underbrace{\mathbf{M}_{RGS}\mathbf{M}_{GS}}_{\mathbf{M}_{SGS} :=} \mathbf{x}_m + \underbrace{(\mathbf{M}_{RGS}\mathbf{N}_{GS} + \mathbf{N}_{RGS})}_{\mathbf{N}_{SGS} :=} \mathbf{b}.$$

Für das hieraus entstandene symmetrische Gauß-Seidel-Verfahren gelten folglich

$$\begin{aligned} \mathbf{M}_{SGS} &= (-(\mathbf{D} + \mathbf{R})^{-1}\mathbf{L}) (-(\mathbf{D} + \mathbf{L})^{-1}\mathbf{R}) \\ &= (\mathbf{D} + \mathbf{R})^{-1}\mathbf{L}(\mathbf{D} + \mathbf{L})^{-1}\mathbf{R} \end{aligned}$$

und

$$\begin{aligned} \mathbf{N}_{SGS} &= (-(\mathbf{D} + \mathbf{R})^{-1}\mathbf{L}) (\mathbf{D} + \mathbf{L})^{-1} + (\mathbf{D} + \mathbf{R})^{-1} \\ &= (\mathbf{D} + \mathbf{R})^{-1} (\mathbf{I} - \mathbf{L}(\mathbf{D} + \mathbf{L})^{-1}) \\ &= (\mathbf{D} + \mathbf{R})^{-1}\mathbf{D}(\mathbf{D} + \mathbf{L})^{-1}. \end{aligned}$$

Analog erhalten wir das symmetrische Gauß-Seidel-Relaxationsverfahren (symmetrisches SOR-Verfahren = SSOR-Verfahren) als Kombination der relaxierten Einzelverfahren in der Form

$$\mathbf{x}_{m+1} = \mathbf{M}_{SGS}(\omega)\mathbf{x}_m + \mathbf{N}_{SGS}(\omega)\mathbf{b}$$

mit

$$\mathbf{M}_{SGS}(\omega) = (\mathbf{D} + \omega\mathbf{R})^{-1} ((1 - \omega)\mathbf{D} - \omega\mathbf{L}) (\mathbf{D} + \omega\mathbf{L})^{-1} ((1 - \omega)\mathbf{D} - \omega\mathbf{R})$$

und

$$\mathbf{N}_{SGS}(\omega) = \omega(2 - \omega) (\mathbf{D} + \omega\mathbf{R})^{-1}\mathbf{D}(\mathbf{D} + \omega\mathbf{L})^{-1}. \quad (4.1.27)$$

Für jede positiv definite und hermitesche Matrix $\mathbf{A} \in \mathbb{C}^{n \times n}$ existiert eine unitäre Matrix $\mathbf{Q} \in \mathbb{C}^{n \times n}$ mit

$$\mathbf{D} = \mathbf{Q}^* \mathbf{A} \mathbf{Q} = \text{diag} \{d_{11}, \dots, d_{nn}\} \text{ mit } d_{ii} > 0 \text{ für } i = 1, \dots, n.$$

Dann sei im folgenden stets

$$\mathbf{A}^{1/\alpha} := \mathbf{Q}^* \mathbf{D}^{1/\alpha} \mathbf{Q} = \mathbf{Q}^* \text{diag} \{d_{11}^{1/\alpha}, \dots, d_{nn}^{1/\alpha}\} \mathbf{Q}.$$

Im Fall einer positiv definiten, symmetrischen Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ ist $\mathbf{Q} \in \mathbb{R}^{n \times n}$ orthogonal.

Satz 4.35 *Das symmetrische Gauß-Seidel-Relaxationsverfahren ist für $\omega \in (0,2)$ eine konsistente symmetrische Splitting-Methode.*

Beweis:

Mit

$$\mathbf{M}_{GS}(\omega) = \mathbf{I} - \mathbf{N}_{GS}(\omega)\mathbf{A} \quad \text{und} \quad \mathbf{M}_{RGS}(\omega) = \mathbf{I} - \mathbf{N}_{RGS}(\omega)\mathbf{A}$$

erhalten wir

$$\begin{aligned} & \mathbf{I} - \mathbf{N}_{SGS}(\omega)\mathbf{A} \\ &= \mathbf{I} - (\mathbf{M}_{RGS}(\omega)\mathbf{N}_{GS}(\omega) + \mathbf{N}_{RGS}(\omega))\mathbf{A} \\ &= \underbrace{\mathbf{I} - \mathbf{N}_{RGS}(\omega)\mathbf{A}}_{=\mathbf{M}_{RGS}(\omega)} - \mathbf{M}_{RGS}(\omega) \underbrace{\mathbf{N}_{GS}(\omega)\mathbf{A}}_{=\mathbf{I} - \mathbf{M}_{GS}(\omega)} \\ &= \mathbf{M}_{RGS}(\omega)\mathbf{M}_{GS}(\omega) \\ &= \mathbf{M}_{SGS}(\omega), \end{aligned} \tag{4.1.28}$$

wodurch das SSOR eine konsistente Splitting-Methode repräsentiert.

Sei $\mathbf{A}$ positiv definit und symmetrisch, dann gelten $a_{ii} > 0$ ($i = 1, \dots, n$) und $\mathbf{L} = \mathbf{R}^T$.

Sei $\mathbf{D}^{1/2} = \text{diag} \{ \sqrt{a_{11}}, \dots, \sqrt{a_{nn}} \}$, dann folgt mit $\mathbf{B}_{SGS}(\omega) = \mathbf{N}_{SGS}^{-1}(\omega)$ die Gleichung

$$\begin{aligned} \mathbf{B}_{SGS}(\omega) &= \frac{1}{\omega(2-\omega)} (\mathbf{D} + \omega\mathbf{L})\mathbf{D}^{-1}(\mathbf{D} + \omega\mathbf{R}) \\ &= \frac{1}{\sqrt{\omega(2-\omega)}} \left(\mathbf{D} + \omega\mathbf{R}^T \right) \mathbf{D}^{-1/2} \underbrace{\frac{1}{\sqrt{\omega(2-\omega)}} \mathbf{D}^{-1/2}(\mathbf{D} + \omega\mathbf{R})}_{\mathbf{C} :=} \\ &= \mathbf{C}^T \mathbf{C}. \end{aligned}$$

Die Voraussetzungen $\omega \in (0,2)$ und $a_{ii} > 0$ ($i = 1, \dots, n$) liefern die Regularität der Matrix $\mathbf{C}$, wodurch sich neben der Symmetrie auch die positive Definitheit der Matrix $\mathbf{B}_{SGS} = \mathbf{C}^T \mathbf{C}$ ergibt. $\square$

Satz und Definition 4.36 *Sei $\mathbf{A} \in \mathbb{C}^{n \times n}$ eine positiv definite und hermitesche Matrix, dann stellt*

$$\|\cdot\|_{\mathbf{A}} : \mathbb{C}^n \rightarrow \mathbb{R}$$

$$\mathbf{x} \mapsto \|\mathbf{x}\|_{\mathbf{A}} := \left\| \mathbf{A}^{1/2} \mathbf{x} \right\|_2$$

eine Norm dar. Diese Norm heißt *Energienorm* (bezüglich $\mathbf{A}$). Die zugehörige Matrixnorm ist durch

$$\|\cdot\|_{\mathbf{A}} : \mathbb{C}^{n \times n} \rightarrow \mathbb{R}$$

$$\mathbf{B} \mapsto \|\mathbf{B}\|_{\mathbf{A}} = \left\| \mathbf{A}^{1/2} \mathbf{B} \mathbf{A}^{-1/2} \right\|_2$$

gegeben.

Beweis:

Der Nachweis der ersten Aussage beruht auf einfachem Nachrechnen der Vektornormaxiome und verbleibt als Übungsaufgabe. Für die Matrixnorm erhalten wir

$$\begin{aligned} \|\mathbf{B}\|_{\mathbf{A}} &= \sup_{\|\mathbf{x}\|_{\mathbf{A}}=1} \|\mathbf{B}\mathbf{x}\|_{\mathbf{A}} = \sup_{\|\mathbf{A}^{1/2}\mathbf{x}\|_2=1} \|\mathbf{A}^{1/2}\mathbf{B}\mathbf{x}\|_2 \\ &= \sup_{\|\mathbf{y}\|_2=1} \|\mathbf{A}^{1/2}\mathbf{B}\mathbf{A}^{-1/2}\mathbf{y}\|_2 = \left\| \mathbf{A}^{1/2}\mathbf{B}\mathbf{A}^{-1/2} \right\|_2. \end{aligned}$$

□

Satz 4.37 Sei $\mathbf{A} \in \mathbb{C}^{n \times n}$ hermitesch und positiv definit, dann konvergiert das symmetrische Gauß-Seidel-Relaxationsverfahren für alle $\omega \in (0,2)$, und es gilt

$$\rho(\mathbf{M}_{SGS}(\omega)) = \|\mathbf{M}_{SGS}(\omega)\|_{\mathbf{A}} = \|\mathbf{M}_{GS}(\omega)\|_{\mathbf{A}}^2,$$

sowie

$$\sigma(\mathbf{M}_{SGS}(\omega)) \subset [0, \rho(\mathbf{M}_{SGS}(\omega))].$$

Beweis:

Da $\mathbf{A}^{1/2} \in \mathbb{C}^{n \times n}$ regulär ist, ist $\mathbf{M}_{SGS}(\omega)$ zu

$$\mathbf{A}^{1/2} \mathbf{M}_{SGS}(\omega) \mathbf{A}^{-1/2} = \left(\mathbf{A}^{1/2} \mathbf{M}_{RGS}(\omega) \mathbf{A}^{-1/2} \right) \left(\mathbf{A}^{1/2} \mathbf{M}_{GS}(\omega) \mathbf{A}^{-1/2} \right) \quad (4.1.29)$$

ähnlich. Mit

$$\mathbf{M}_{RGS}(\omega) = (\mathbf{D} + \omega \mathbf{R})^{-1} ((1 - \omega) \mathbf{D} - \omega \mathbf{L}) = \mathbf{I} - \omega (\mathbf{D} + \omega \mathbf{R})^{-1} \mathbf{A}$$

und

$$\mathbf{M}_{GS}(\omega) = (\mathbf{D} + \omega \mathbf{L})^{-1} ((1 - \omega) \mathbf{D} - \omega \mathbf{R}) = \mathbf{I} - \omega (\mathbf{D} + \omega \mathbf{L})^{-1} \mathbf{A}$$

erhalten wir

$$\begin{aligned} \mathbf{A}^{1/2} \mathbf{M}_{RGS}(\omega) \mathbf{A}^{-1/2} &= \mathbf{I} - \omega \mathbf{A}^{1/2} (\mathbf{D} + \omega \mathbf{R})^{-1} \mathbf{A}^{1/2} \\ &= \left(\mathbf{I} - \omega \mathbf{A}^{1/2} ((\mathbf{D} + \omega \mathbf{R})^{-1})^* \mathbf{A}^{1/2} \right)^* \\ &= \left(\mathbf{I} - \omega \mathbf{A}^{1/2} (\mathbf{D} + \omega \mathbf{L})^{-1} \mathbf{A}^{1/2} \right)^* \\ &= \left(\mathbf{A}^{1/2} \mathbf{M}_{GS}(\omega) \mathbf{A}^{-1/2} \right)^*. \end{aligned} \quad (4.1.30)$$

Aus $\mathbf{A}^{1/2} \mathbf{M}_{SGS}(\omega) \mathbf{A}^{-1/2} = \left(\mathbf{A}^{1/2} \mathbf{M}_{GS}(\omega) \mathbf{A}^{-1/2} \right)^* \left(\mathbf{A}^{1/2} \mathbf{M}_{GS}(\omega) \mathbf{A}^{-1/2} \right)$ folgt

$$\sigma(\mathbf{M}_{SGS}(\omega)) = \sigma\left(\mathbf{A}^{1/2} \mathbf{M}_{SGS}(\omega) \mathbf{A}^{-1/2}\right) \subset \mathbb{R}_0^+.$$

Somit gilt

$$\sigma(\mathbf{M}_{SGS}(\omega)) \subset [0, \rho(\mathbf{M}_{SGS}(\omega))]. \quad (4.1.31)$$

Aus (4.1.29) und (4.1.30) wird direkt ersichtlich, daß die Matrix $\mathbf{A}^{1/2} \mathbf{M}_{SGS}(\omega) \mathbf{A}^{-1/2}$ hermitesch ist und somit

$$\begin{aligned} \rho(\mathbf{M}_{SGS}(\omega)) &= \rho\left(\mathbf{A}^{1/2} \mathbf{M}_{SGS}(\omega) \mathbf{A}^{-1/2}\right) \\ &\stackrel{\text{Satz 2.34}}{=} \|\mathbf{A}^{1/2} \mathbf{M}_{SGS}(\omega) \mathbf{A}^{-1/2}\|_2 = \|\mathbf{M}_{SGS}(\omega)\|_{\mathbf{A}} \\ &\stackrel{(4.1.29), (4.1.30)}{=} \left\| \left(\mathbf{A}^{1/2} \mathbf{M}_{GS}(\omega) \mathbf{A}^{-1/2} \right)^* \left(\mathbf{A}^{1/2} \mathbf{M}_{GS}(\omega) \mathbf{A}^{-1/2} \right) \right\|_2 \\ &= \left\| \mathbf{A}^{1/2} \mathbf{M}_{GS}(\omega) \mathbf{A}^{-1/2} \right\|_2^2 = \|\mathbf{M}_{GS}(\omega)\|_{\mathbf{A}}^2 \end{aligned}$$

folgt.

Für hermitesche Matrizen führen wir die folgende Ordnungsrelation ein: Seien die Matrizen $\mathbf{F}, \mathbf{G} \in \mathbb{C}^{n \times n}$ hermitesch, dann schreiben wir $\mathbf{G} > \mathbf{F}$, wenn $\mathbf{G} - \mathbf{F}$ positiv definit ist. Seien $\mathbf{F}, \mathbf{G} \in \mathbb{C}^{n \times n}$ mit $\mathbf{G} > \mathbf{F}$ und $\mathbf{F}$ positiv definit gegeben, dann sei $\lambda \in \mathbb{R}^+$ Eigenwert von $\mathbf{F}$ mit $\lambda = \rho(\mathbf{F})$ und zugehörigem Eigenvektor $\mathbf{x} \in \mathbb{C}^n$ mit $\|\mathbf{x}\|_2 = 1$. Somit erhalten wir

$$\begin{aligned} 0 &< \rho(\mathbf{F}) = (\mathbf{F}\mathbf{x}, \mathbf{x})_2 \stackrel{(\mathbf{G} > \mathbf{F})}{<} (\mathbf{G}\mathbf{x}, \mathbf{x})_2 \leq \|\mathbf{G}\mathbf{x}\|_2 \|\mathbf{x}\|_2 \\ &\leq \|\mathbf{G}\|_2 \|\mathbf{x}\|_2^2 = \|\mathbf{G}\|_2 = \rho(\mathbf{G}). \end{aligned} \quad (4.1.32)$$

Laut Voraussetzung ist $\mathbf{A}$ positiv definit und hermitesch, wodurch sich für $\omega \in (0, 2)$ die Gleichung

$$\begin{aligned} \mathbf{N}_{GS}^{-1}(\omega) + \underbrace{\mathbf{N}_{GS}^{-*}(\omega)}_{:= (\mathbf{N}_{GS}^{-1}(\omega))^*} &= \frac{1}{\omega} (\mathbf{D} + \omega \mathbf{L}) + \frac{1}{\omega} (\mathbf{D} + \omega \mathbf{R}) \\ &= \mathbf{L} + \mathbf{R} + \mathbf{D} + \left(\frac{2}{\omega} - 1 \right) \mathbf{D} \\ &= \mathbf{A} + \left(\frac{2}{\omega} - 1 \right) \mathbf{D} > \mathbf{A} \end{aligned} \quad (4.1.33)$$

ergibt. Betrachten wir die Matrix

$$\begin{aligned} \widehat{\mathbf{M}}_{GS}(\omega) &:= \mathbf{A}^{1/2} \mathbf{M}_{GS}(\omega) \mathbf{A}^{-1/2} \\ &= \mathbf{A}^{1/2} (\mathbf{I} - \mathbf{N}_{GS}(\omega) \mathbf{A}) \mathbf{A}^{-1/2} \\ &= \mathbf{I} - \mathbf{A}^{1/2} \mathbf{N}_{GS}(\omega) \mathbf{A}^{1/2}, \end{aligned}$$

dann folgt

$$\begin{aligned}
& \widehat{\mathbf{M}}_{GS}(\omega) \widehat{\mathbf{M}}_{GS}^*(\omega) \\
&= \left(\mathbf{I} - \mathbf{A}^{1/2} \mathbf{N}_{GS}(\omega) \mathbf{A}^{1/2} \right) \left(\mathbf{I} - \mathbf{A}^{1/2} \mathbf{N}_{GS}^*(\omega) \mathbf{A}^{1/2} \right) \\
&= \mathbf{I} - \mathbf{A}^{1/2} (\mathbf{N}_{GS}(\omega) + \mathbf{N}_{GS}^*(\omega)) \mathbf{A}^{1/2} \\
&\quad + \mathbf{A}^{1/2} \mathbf{N}_{GS}(\omega) \mathbf{A} \mathbf{N}_{GS}^*(\omega) \mathbf{A}^{1/2} \\
&= \mathbf{I} - \mathbf{A}^{1/2} \mathbf{N}_{GS}(\omega) (\mathbf{N}_{GS}^{-1}(\omega) + \mathbf{N}_{GS}^{-*}(\omega)) \mathbf{N}_{GS}^*(\omega) \mathbf{A}^{1/2} \\
&\quad + \mathbf{A}^{1/2} \mathbf{N}_{GS}(\omega) \mathbf{A} \mathbf{N}_{GS}^*(\omega) \mathbf{A}^{1/2} \\
&= \mathbf{I} - \mathbf{A}^{1/2} \mathbf{N}_{GS}(\omega) (\mathbf{N}_{GS}^{-1}(\omega) + \mathbf{N}_{GS}^{-*}(\omega) - \mathbf{A}) \mathbf{N}_{GS}^*(\omega) \mathbf{A}^{1/2} \\
&\stackrel{(4.1.33)}{<} \mathbf{I}. \tag{4.1.34}
\end{aligned}$$

Mit (4.1.32) erhalten wir $\rho(\mathbf{M}_{SGS}(\omega)) = \rho(\widehat{\mathbf{M}}_{GS}(\omega) \widehat{\mathbf{M}}_{GS}^*(\omega)) < \rho(\mathbf{I}) = 1$. $\square$

Der obige Beweis beinhaltet zudem die Aussage, daß für eine positiv definite Matrix $\mathbf{A}$ das Gauß-Seidel-Relaxationsverfahren für $\omega \in (0, 2)$ monoton in der Energienorm konvergiert.

4.2 Mehrgitterverfahren

Mehrgitterverfahren werden neben der Lösung linearer Gleichungssysteme, die zum Beispiel bei der Diskretisierung elliptischer Differentialgleichungen (Poisson-Gleichung) entstehen, auch zur Konvergenzbeschleunigung innerhalb expliziter numerischer Verfahren zur Lösung hyperbolischer (Euler-Gleichungen) und hyperbolisch-parabolischer (Navier-Stokes-Gleichungen) Differentialgleichungen angewendet [11, 4, 29]. Obwohl im zweiten Fall kein lineares Gleichungssystem betrachtet wird, bleibt die prinzipielle Vorgehensweise identisch.

Wir werden in diesem Abschnitt die Grundidee der Mehrgitterverfahren vorstellen und ihre Wirkung bei der Lösung eines speziellen Modellproblems studieren. Konvergenzaussagen werden zum Beispiel in [36, 34] hergeleitet.

Wir vereinfachen die im Beispiel 1.1 vorgestellte Poisson-Gleichung zum eindimensionalen Dirichlet-Randwertproblem.

Gegeben: $\Omega = (0, 1)$ und $f \in C(\Omega; \mathbb{R})$

Gesucht: $u \in C^2(\Omega; \mathbb{R}) \cap C(\overline{\Omega}; \mathbb{R})$ mit

$$\begin{aligned}
-u''(x) &= f(x) & \text{für } x \in \Omega, \\
u(x) &= 0 & \text{für } x \in \partial\Omega = \{0, 1\}.
\end{aligned} \tag{4.2.1}$$

Wir definieren die Schrittweitenfolge $\{h_\ell\}_{\ell=0}^\infty$ mit $h_0 = \frac{1}{2}$ und $h_\ell = h_0/2^\ell = 2^{-(\ell+1)}$ sowie die Gitterfolge

$$\Omega_\ell := \Omega_{h_\ell} = \{jh_\ell \mid j = 1, \dots, 2^{\ell+1} - 1\} \text{ für } \ell = 0, 1, \dots,$$

wobei ℓ als Stufenzahl oder Stufenindex bezeichnet wird.

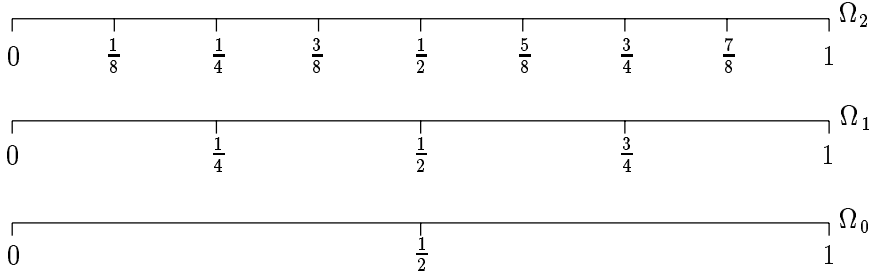


Bild 4.5 Gitterhierarchie

Approximieren wir wiederum die Ableitung der gesuchten Funktion u mittels eines zentralen Differenzenquotienten, so erhalten wir mit

$$u_j^\ell := u(jh_\ell) \text{ für } j = 1, \dots, N_\ell := 2^{\ell+1} - 1$$

und

$$f_j^\ell := f(jh_\ell) \text{ für } j = 1, \dots, N_\ell$$

unter Berücksichtigung der Randbedingungen das lineare Gleichungssystem der ℓ -ten Stufe in der Form

$$\mathbf{A}_\ell \mathbf{u}^\ell = \mathbf{f}^\ell$$

mit

$$\mathbf{u}^\ell = (u_1^\ell, \dots, u_{N_\ell}^\ell)^T, \quad \mathbf{f}^\ell = (f_1^\ell, \dots, f_{N_\ell}^\ell)^T$$

und

$$\mathbf{A}_\ell = \frac{1}{h_\ell^2} \begin{pmatrix} 2 & -1 & & & \\ -1 & 2 & -1 & & \\ & \ddots & \ddots & \ddots & \\ & & \ddots & \ddots & -1 \\ & & & -1 & 2 \end{pmatrix}.$$

Die Matrix $\mathbf{A}_\ell$ ist irreduzibel und diagonaldominant, so daß sowohl das Jacobi- als auch das Gauß-Seidel-Verfahren konvergiert.

Bemerkung:

Das vorliegende Modellproblem eignet sich sehr gut zur Beschreibung des Mehrgitteralgorithmus. Es sei an dieser Stelle jedoch erwähnt, daß derartige Gleichungssysteme in der Praxis direkt gelöst werden sollten. Der Grund hierfür liegt in der speziellen Gestalt der Matrix $\mathbf{A}_\ell$. Zunächst gilt $\det \mathbf{A}_\ell[k] \neq 0$ für $k = 1, \dots, n$, wodurch laut Satz 3.7

eine LR-Zerlegung der Matrix $\mathbf{A}_\ell$ ohne Pivotisierung ermittelt werden kann. Berücksichtigt man innerhalb der Zerlegung die Tridiagonalgestalt der Matrix, so läßt sich das Gleichungssystem mittels des Gaußschen Eliminationsverfahrens derart lösen, daß der Rechenaufwand proportional zur Anzahl der Unbekannten ist. Eine explizite Herleitung dieser Methode wird in [63] beschrieben.

Wir betrachten das Jacobi-Relaxationsverfahren

$$\mathbf{u}_{m+1}^\ell = \mathbf{u}_m^\ell + \tilde{\omega} \mathbf{D}_\ell^{-1} (\mathbf{b} - \mathbf{A}_\ell \mathbf{u}_m^\ell) \text{ für } m = 0, 1, \dots,$$

das aufgrund der speziellen Gestalt der Diagonalmatrix $\mathbf{D}_\ell$ mit $\omega = \frac{1}{2} \tilde{\omega}$ die Form

$$\begin{aligned} \mathbf{u}_{m+1}^\ell &= \mathbf{u}_m^\ell + \omega h_\ell^2 (\mathbf{b} - \mathbf{A}_\ell \mathbf{u}_m^\ell) \\ &= \underbrace{(\mathbf{I} - \omega h_\ell^2 \mathbf{A}_\ell)}_{\mathbf{M}_\ell(\omega) :=} \mathbf{u}_m^\ell + \underbrace{\omega h_\ell^2 \mathbf{I}}_{\mathbf{N}_\ell(\omega) :=} \mathbf{b} \end{aligned} \quad (4.2.2)$$

besitzt und dem Richardson-Verfahren mit $\Theta = \omega h_\ell^2$ entspricht. Wir werden vorab die Konvergenzeigenschaften der Methode untersuchen.

Da die Eigenfunktionen des homogenen Randwertproblems (4.2.1) durch

$$u(x) = c \sin(j\pi x) \text{ mit } j \in \mathbb{N} \text{ und } c \in \mathbb{R}$$

gegeben sind, erweist es sich als nicht besonders überraschend, daß die Eigenvektoren $\mathbf{e}^{\ell,j}$ der Matrix $\mathbf{A}_\ell$ durch

$$\mathbf{e}^{\ell,j} = \sqrt{2h_\ell} \begin{pmatrix} \sin j\pi h_\ell \\ \vdots \\ \sin j\pi N_\ell h_\ell \end{pmatrix} \text{ für } j = 1, \dots, N_\ell \quad (4.2.3)$$

deren diskrete Formulierung repräsentieren. Die zugehörigen Eigenwerte $\lambda^{\ell,j}$ erhalten wir durch

$$\mathbf{A}_\ell \mathbf{e}^{\ell,j} = \underbrace{4h_\ell^{-2} \sin^2 \left(\frac{j\pi h_\ell}{2} \right)}_{\lambda^{\ell,j} =} \mathbf{e}^{\ell,j} \text{ für } j = 1, \dots, N_\ell.$$

Mit $\mathbf{M}_\ell(\omega) = \mathbf{I} - \omega h_\ell^2 \mathbf{A}_\ell$ stimmen die Eigenvektoren von $\mathbf{A}_\ell$ und $\mathbf{M}_\ell(\omega)$ überein, und die Eigenwerte der Iterationsmatrix lauten

$$\lambda^{\ell,j}(\omega) = 1 - 4\omega \sin^2 \left(\frac{j\pi h_\ell}{2} \right) \text{ für } j = 1, \dots, N_\ell. \quad (4.2.4)$$

Lemma 4.38 Die durch (4.2.3) gegebenen Vektoren $\{\mathbf{e}^{\ell,1}, \dots, \mathbf{e}^{\ell,N_\ell}\}$ stellen eine Orthonormalbasis des $\mathbb{R}^{N_\ell}$ dar.

Beweis:

Sei i die komplexe Einheit und $z = e^{i \frac{2\pi j}{N_\ell+1}} \in \mathbb{C}$ mit $j \in \mathbb{Z}$. Für $\frac{j}{N_\ell+1} \in \mathbb{Z}$ folgt direkt $z = 1$ und somit

$$\sum_{k=1}^{N_\ell+1} z^k = N_\ell + 1. \quad (4.2.5)$$

Gilt $\frac{j}{N_\ell+1} \notin \mathbb{Z}$, so ergibt sich $z \neq 1 = z^{N_\ell+1}$ und folglich

$$\sum_{k=1}^{N_\ell+1} z^k = z \frac{z^{N_\ell+1} - 1}{z - 1} = 0. \quad (4.2.6)$$

Zusammenfassend erhalten wir aus den Gleichungen (4.2.5) und (4.2.6) für den Realteil der betrachteten Summen die Darstellung

$$\sum_{k=1}^{N_\ell+1} \cos\left(k \frac{2\pi j}{N_\ell+1}\right) = \begin{cases} 0 & , \text{ falls } \frac{j}{N_\ell+1} \notin \mathbb{Z} \\ N_\ell + 1 & \text{ sonst.} \end{cases}$$

Die Orthogonalität der Vektoren erhalten wir mit $j, m \in \{1, \dots, N_\ell\}$ unter Verwendung von $\cos((j-m)\pi) - \cos((j+m)\pi) = 0$ mittels

$$\begin{aligned} (\mathbf{e}^{\ell,j}, \mathbf{e}^{\ell,m})_2 &= \frac{2}{N_\ell+1} \sum_{k=1}^{N_\ell} \sin\left(j \frac{\pi k}{N_\ell+1}\right) \sin\left(m \frac{\pi k}{N_\ell+1}\right) \\ &= \frac{1}{N_\ell+1} \sum_{k=1}^{N_\ell} \left\{ \cos\left((j-m) \frac{\pi k}{N_\ell+1}\right) - \cos\left((j+m) \frac{\pi k}{N_\ell+1}\right) \right\} \\ &= \begin{cases} 0 & \text{für } j \neq m, \\ 1 & \text{für } j = m. \end{cases} \end{aligned}$$

Die Basiseigenschaft folgt abschließend direkt aus der Orthonormalität der Vektoren. $\square$

Da die Matrix $\mathbf{A}$ als Tridiagonalmatrix laut Beispiel 4.26 konsistent geordnet ist, liegt eine um Null symmetrische Verteilung der Eigenwerte der Iterationsmatrix $\mathbf{M}_J$ des Jacobi-Verfahrens vor. Unter Berücksichtigung dieser Eigenschaft kann durch eine Relaxation des Verfahrens laut Satz 4.20 keine Verringerung des Spektralradius $\rho(\mathbf{M}_J)$ erzielt werden. Der Vorteil einer Relaxation liegt vielmehr in den Dämpfungseigenschaften der Methode hinsichtlich unterschiedlicher Fehlerfrequenzen, die wir im folgenden untersuchen werden.

Mit dem obigen Lemma läßt sich der Fehler zwischen dem Startvektor $\mathbf{u}_0^\ell$ und der exakten Lösung $\mathbf{u}^{\ell,*} = \mathbf{A}_\ell^{-1} \mathbf{f}^\ell$ in der Form

$$\mathbf{u}_0^\ell - \mathbf{u}^{\ell,*} = \sum_{j=1}^{N_\ell} \alpha_j \mathbf{e}^{\ell,j}, \quad \alpha_j \in \mathbb{R}$$

schreiben, und wir erhalten

$$\begin{aligned} \mathbf{u}_1^\ell - \mathbf{u}^{\ell,*} &= \mathbf{M}_\ell(\omega) \mathbf{u}_0^\ell + \mathbf{N}_\ell(\omega) \mathbf{f}^\ell - \left(\mathbf{M}_\ell(\omega) \mathbf{u}^{\ell,*} + \mathbf{N}_\ell(\omega) \mathbf{f}^\ell \right) \\ &= \mathbf{M}_\ell(\omega) (\mathbf{u}_0^\ell - \mathbf{u}^{\ell,*}) = \sum_{j=1}^{N_\ell} \alpha_j \lambda^{\ell,j}(\omega) \mathbf{e}^{\ell,j} \end{aligned}$$

und entsprechend

$$\mathbf{u}_m^\ell - \mathbf{u}^{\ell,*} = \sum_{j=1}^{N_\ell} \alpha_j [\lambda^{\ell,j}(\omega)]^m \mathbf{e}^{\ell,j} \text{ für } m = 0, 1, \dots$$

Für die Eigenwerte $\lambda^{\ell,j}(\omega)$ erhalten wir beim Jacobi-Verfahren ($\tilde{\omega} = 1$, d.h. $\omega = \frac{\tilde{\omega}}{2} = \frac{1}{2}$) die in Abbildung 4.6 dargestellte graphische Verteilung, wobei die diskreten Werte für $\ell = 3$ durch * gekennzeichnet sind.

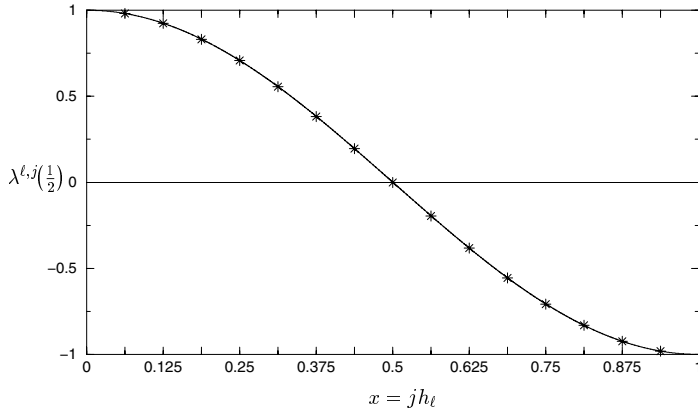


Bild 4.6 Eigenwertverteilung der Iterationsmatrix des Jacobi-Verfahrens

Wir können der Abbildung 4.6 drei wesentliche Aussagen entnehmen:

- (1) Das Jacobi-Verfahren liefert einen schnellen Abfall der im mittleren Frequenzbereich befindlichen Fehlerkomponenten, während die hoch- und niedrigfrequenten Anteile deutlich geringer gedämpft werden.
- (2) Je feiner die Diskretisierung gewählt wird, d. h. je höher die Stufenzahl ℓ ist, desto größer wird der Spektralradius der Iterationsmatrix $\rho(\mathbf{M}_\ell(\frac{1}{2}))$ des Jacobi-Verfahrens.
- (3) Es gilt stets

$$\lambda^{\ell,1}\left(\frac{1}{2}\right) = -\lambda^{\ell,N_\ell}\left(\frac{1}{2}\right) = \rho\left(\mathbf{M}_\ell\left(\frac{1}{2}\right)\right),$$

so daß der Spektralradius der Iterationsmatrix durch eine einfache Relaxation des Jacobi-Verfahrens nicht verkleinert werden kann.

Betrachten wir die Taylorentwicklung des Spektralradius nach der Schrittweite, dann erhalten wir

$$\begin{aligned} \rho(\mathbf{M}_\ell(\omega)) &= 1 - 4\omega \left(\sum_{n=0}^{\infty} (-1)^n \frac{\left(\pi \frac{h_\ell}{2}\right)^{2n+1}}{(2n+1)!} \right)^2 \\ &= 1 - 4\omega \frac{\pi^2 h_\ell^2}{4} + O(h_\ell^4), \end{aligned}$$

so daß der Spektralradius quadratisch mit der Schrittweite gegen 1 konvergiert. Eine Verfeinerung des Gitters zur exakten Berechnung der Lösung bringt somit neben einem erhöhten Aufwand pro Iteration eine drastische Reduktion der Konvergenzgeschwindigkeit mit sich.

Die Wahl $\omega = \frac{1}{2}$ liefert zwar die optimale Konvergenzgeschwindigkeit für das beschriebene Verfahren, beinhaltet jedoch das Problem, daß sowohl langwellige wie auch kurzwellige Fehleranteile sehr schlecht gedämpft werden.

Grundlegend für das Mehrgitterverfahren ist jedoch eine Dämpfung der hochfrequenten Fehleranteile auf dem feinsten Gitter. Wir wählen daher im folgenden $\omega = \frac{1}{4}$, wodurch die Eigenwertverteilung gemäß Abbildung 4.7 folgt, der wir entnehmen, daß die langwelligsten Fehlerterme nur sehr langsam, die hochfrequenten Fehleranteile allerdings sehr schnell gedämpft werden.

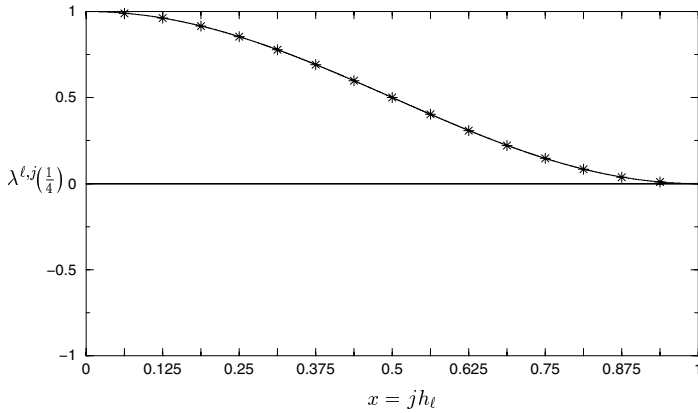


Bild 4.7 Eigenwertverteilung der Iterationsmatrix des gedämpften Jacobi-Verfahrens

Sei beispielsweise der Fehler $\mathbf{e}_0^\ell$ in der Form

$$\mathbf{e}_0^2 := \mathbf{u}_0^2 - \mathbf{u}^{2,*} = (0.75, 0.2, 0.9, 0.8, 0.55, 0.9, 0.3)^T$$

auf der 2-ten Stufe ($N_\ell = N_2 = 7$) gegeben, so zeigt die Abbildung 4.8 die Entwicklung des Fehlers.

Ideal wäre somit eine Kopplung des Jacobi-Relaxationsverfahrens mit einer Iterationsvorschrift, die die komplementäre Eigenschaft der schnellen Dämpfung langwelliger Fehlerterme aufweist. Alle bisherigen im Abschnitt 4.1 vorgestellten Verfahren weisen eine solche Eigenschaft jedoch nicht auf. Beim Mehrgitterverfahren nutzen wir die Kenntnis der Glattheit des Fehlers, indem wir diesen auf einem größeren Gitter approximieren und anschließend, mittels zum Beispiel einer linearen Interpolation, auf das feine Gitter abbilden. Zunächst benötigen wir hierzu eine Abbildung vom feinen Gitter Ω_ℓ auf das gröbere Gitter $\Omega_{\ell-1}$, die wir *Restriktion* nennen, und eine Abbildung von $\Omega_{\ell-1}$ auf Ω_ℓ , die wir als *Prolongation* bezeichnen.

Als Restriktion von Ω_ℓ auf $\Omega_{\ell-1}$ bezeichnen wir eine lineare surjektive Abbildung

$$\mathbf{R}_\ell^{\ell-1} : \mathbb{R}^{N_\ell} \rightarrow \mathbb{R}^{N_{\ell-1}}.$$

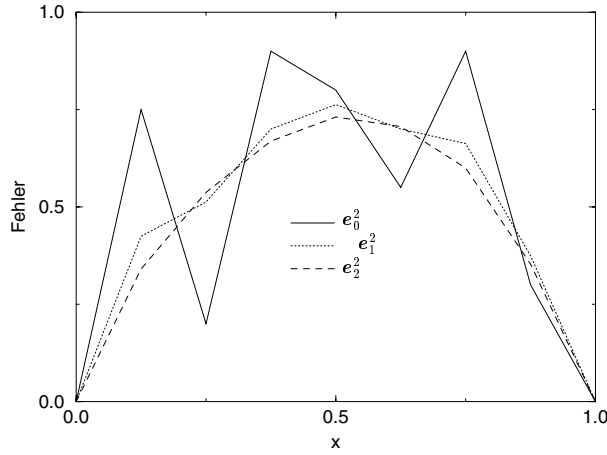


Bild 4.8 Entwicklung des Fehlers beim gedämpften Jacobi-Verfahren

Bei der speziellen Schachtelung der Gitterfolge kann zum Beispiel die triviale Restriktion gemäß Abbildung 4.9 verwendet werden, die durch

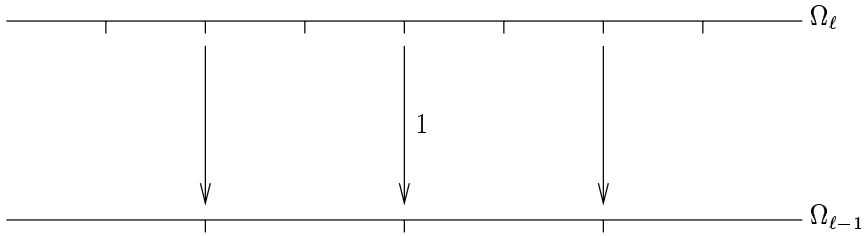
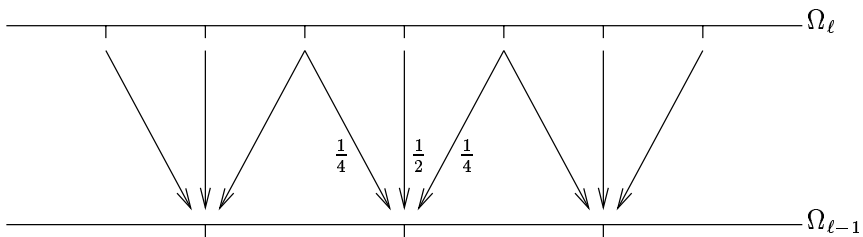
$$\mathbf{u}^{\ell-1} = \begin{pmatrix} u_1^{\ell-1} \\ \vdots \\ u_{N_{\ell-1}}^{\ell-1} \end{pmatrix} = \mathbf{R}_{\ell}^{\ell-1} \mathbf{u}^{\ell} = \begin{pmatrix} u_2^{\ell} \\ u_4^{\ell} \\ \vdots \\ u_{N_{\ell}-1}^{\ell} \end{pmatrix}$$

gegeben ist und durch die Matrix

$$\mathbf{R}_{\ell}^{\ell-1} = \begin{pmatrix} 0 & 1 & 0 & & & \\ & 0 & 1 & 0 & & \\ & & 0 & 1 & 0 & \\ & & & \ddots & \ddots & \ddots \\ & & & & \ddots & \ddots \\ & & & & & 0 & 1 & 0 \end{pmatrix} \in \mathbb{R}^{N_{\ell-1} \times N_{\ell}} \quad (4.2.7)$$

repräsentiert wird.

Diese einfache Restriktion erweist sich oftmals als ungünstig, da die Werte an den Gitterpunkten $\Omega_{\ell} \setminus \Omega_{\ell-1}$ und hierdurch deren Informationen verlorengehen. Daher wird häufig die Restriktion gemäß Abbildung 4.10 mit der entsprechenden Matrixdarstellung

**Bild 4.9** Triviale Restriktion**Bild 4.10** Lineare Restriktion

$$\mathbf{R}_\ell^{\ell-1} = \frac{1}{4} \begin{pmatrix} 1 & 2 & 1 & & & & \\ & 1 & 2 & 1 & & & \\ & & 1 & 2 & 1 & & \\ & & & \ddots & \ddots & \ddots & \\ & & & & \ddots & \ddots & \ddots \\ & & & & & 1 & 2 & 1 \end{pmatrix}$$

genutzt wird.

Als Prolongation von $\Omega_{\ell-1}$ auf Ω_ℓ bezeichnen wir eine lineare injektive Abbildung

$$\mathbf{P}_{\ell-1}^\ell : \mathbb{R}^{N_{\ell-1}} \rightarrow \mathbb{R}^{N_\ell}.$$

Hierzu kann zum Beispiel eine lineare Interpolation zur Definition der Werte an den Zwischenstellen genutzt werden. In unserem Modellfall ergibt sich die graphische Darstellung gemäß Abbildung 4.11 und damit die Matrix

$$P_{\ell-1}^\ell = \frac{1}{2} \begin{pmatrix} 1 & & & & & & & \\ 2 & & & & & & & \\ & 1 & & & & & & \\ & & 2 & & & & & \\ & & & 1 & & & & \\ & & & & 2 & & & \\ & & & & & 1 & & \\ & & & & & & \ddots & \\ & & & & & & & 1 \\ & & & & & & & 2 \\ & & & & & & & 1 \end{pmatrix} \in \mathbb{R}^{N_\ell \times N_{\ell-1}}. \quad (4.2.8)$$

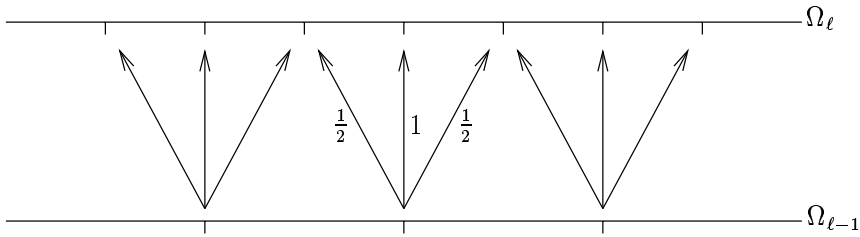


Bild 4.11 Lineare Prolongation

4.2.1 Zweigitterverfahren

Nach j Schritten des Iterationsverfahrens (4.2.2) mit $\omega = 1/4$ erhalten wir einen weitgehend glatten Fehler

$$\mathbf{e}_j^\ell = \mathbf{u}_j^\ell - \mathbf{u}^{\ell,*}. \quad (4.2.9)$$

Dieser Vektor lässt sich somit gut auf dem nächstgrößeren Gitter $\Omega_{\ell-1}$ mit deutlich weniger Rechenaufwand approximieren. Mit dem Defekt

$$\mathbf{d}_j^\ell := \mathbf{A}_\ell \mathbf{u}_j^\ell - \mathbf{f}^\ell \quad (4.2.10)$$

erhalten wir

$$\mathbf{A}_\ell \mathbf{e}_j^\ell \stackrel{(4.2.9)}{=} \mathbf{A}_\ell (\mathbf{u}_j^\ell - \mathbf{u}^{\ell,*}) = \mathbf{A}_\ell \mathbf{u}_j^\ell - \mathbf{A}_\ell \mathbf{u}^{\ell,*} = \mathbf{d}_j^\ell. \quad (4.2.11)$$

Wir ermitteln nun eine Näherung an $\mathbf{e}_j^\ell$, indem wir die Gleichung (4.2.11) auf dem Gitter $\Omega_{\ell-1}$ betrachten, dort exakt lösen und den somit berechneten Vektor auf das ursprüngliche Gitter Ω_ℓ prolongieren.

Wir betrachten daher die Gleichung

$$\mathbf{A}_{\ell-1} \mathbf{e}^{\ell-1} = \mathbf{d}^{\ell-1} \quad (4.2.12)$$

mit dem restringierten Defekt

$$\mathbf{d}^{\ell-1} = \mathbf{R}_{\ell-1}^{\ell-1} \mathbf{d}_j^\ell. \quad (4.2.13)$$

Wird (4.2.12) exakt gelöst, dann liefert

$$\mathbf{P}_{\ell-1}^\ell \mathbf{e}^{\ell-1} = \mathbf{P}_{\ell-1}^\ell \mathbf{A}_{\ell-1}^{-1} \mathbf{d}^{\ell-1} \quad (4.2.14)$$

eine Approximation an den gesuchten Fehlervektor $\mathbf{e}_j^\ell$.

Die unter Verwendung des groben Gitters ermittelte Korrektur der vorliegenden Näherungslösung $\mathbf{u}_j^\ell$ läßt sich somit durch (4.2.9) bis (4.2.14) in der Form

$$\mathbf{u}_j^{\ell,\text{neu}} = \mathbf{u}_j^\ell - \mathbf{P}_{\ell-1}^\ell \mathbf{A}_{\ell-1}^{-1} \mathbf{R}_\ell^{\ell-1} \left(\mathbf{A}_\ell \mathbf{u}_j^\ell - \mathbf{f}^\ell \right)$$

zusammenfassen.

Definition 4.39 Sei $\mathbf{u}_j^\ell$ eine Näherungslösung der Gleichung $\mathbf{A}_\ell \mathbf{u}^\ell = \mathbf{f}^\ell$, dann heißt die Methode

$$\mathbf{u}_j^{\ell,\text{neu}} = \phi_\ell^{GGK} \left(\mathbf{u}_j^\ell, \mathbf{f}^\ell \right)$$

mit

$$\phi_\ell^{GGK} \left(\mathbf{u}_j^\ell, \mathbf{f}^\ell \right) = \mathbf{u}_j^\ell - \mathbf{P}_{\ell-1}^\ell \mathbf{A}_{\ell-1}^{-1} \mathbf{R}_\ell^{\ell-1} \left(\mathbf{A}_\ell \mathbf{u}_j^\ell - \mathbf{f}^\ell \right) \quad (4.2.15)$$

Grobgitterkorrekturverfahren.

Lemma 4.40 Die Grobgitterkorrekturmethode ϕ_ℓ^{GGK} stellt ein lineares konsistentes Iterationsverfahren mit

$$\mathbf{M}_\ell^{GGK} = \mathbf{I} - \mathbf{P}_{\ell-1}^\ell \mathbf{A}_{\ell-1}^{-1} \mathbf{R}_\ell^{\ell-1} \mathbf{A}_\ell \quad (4.2.16)$$

und

$$\mathbf{N}_\ell^{GGK} = \mathbf{P}_{\ell-1}^\ell \mathbf{A}_{\ell-1}^{-1} \mathbf{R}_\ell^{\ell-1} \quad (4.2.17)$$

dar.

Beweis:

Mit (4.2.15) erhalten wir

$$\phi_\ell^{GGK} (\mathbf{u}, \mathbf{f}) = \mathbf{M}_\ell^{GGK} \mathbf{u} + \mathbf{N}_\ell^{GGK} \mathbf{f}.$$

Hierbei gilt $\mathbf{M}_\ell^{GGK}, \mathbf{N}_\ell^{GGK} \in \mathbb{R}^{N_\ell \times N_\ell}$ mit $\mathbf{M}_\ell^{GGK} = \mathbf{I} - \mathbf{N}_\ell^{GGK} \mathbf{A}_\ell$, wodurch sich die Behauptung mit Satz 4.4 ergibt. $\square$

Die Grobgitterkorrekturmethode kann somit als eigenständiges Iterationsverfahren genutzt werden. Es erweist sich hierfür jedoch als ungeeignet, wie uns das folgende Lemma belegen wird.

Lemma 4.41 Das Grobgitterkorrekturverfahren ϕ_ℓ^{GGK} ist nicht konvergent.

Beweis:

Wegen $N_\ell > N_{\ell-1}$ ist der Kern von $\mathbf{R}_\ell^{\ell-1}$, kurz $\ker(\mathbf{R}_\ell^{\ell-1})$, nicht trivial. Sei $\mathbf{0} \neq \mathbf{v} \in \ker(\mathbf{R}_\ell^{\ell-1})$, dann folgt mit $\mathbf{w} := \mathbf{A}_\ell^{-1} \mathbf{v} \neq \mathbf{0}$ die Gleichung

$$\mathbf{M}_\ell^{GGK} \mathbf{w} = \mathbf{w} - \mathbf{P}_{\ell-1}^\ell \mathbf{A}_{\ell-1}^{-1} \mathbf{R}_\ell^{\ell-1} \underbrace{\mathbf{A}_\ell \mathbf{w}}_{=\mathbf{v}} = \mathbf{w},$$

$\underbrace{\hspace{10em}}_{=\mathbf{0}}$

wodurch sich $\rho \left(\mathbf{M}_\ell^{GGK} \right) \geq 1$ ergibt. $\square$

Betrachten wir wiederum als Beispiel den Fehler $\mathbf{e}_0^\ell = (0.75, 0.2, 0.9, 0.8, 0.55, 0.9, 0.3)^T$, so erhalten wir mit zwei Iterationsschritten auf Ω_ℓ und einer anschließenden Grobgitterkorrektur die in Abbildung 4.12 aufgeführte Darstellung.

Definition 4.42 Sind $\phi, \psi : \mathbb{C}^n \times \mathbb{C}^n \rightarrow \mathbb{C}^n$ zwei Iterationsverfahren, dann heißt

$$\phi \circ \psi : \mathbb{C}^n \times \mathbb{C}^n \rightarrow \mathbb{C}^n$$

mit

$$\mathbf{x}_{m+1} = (\phi \circ \psi)(\mathbf{x}_m, \mathbf{b}) := \phi(\psi(\mathbf{x}_m, \mathbf{b}), \mathbf{b})$$

Produktiteration.

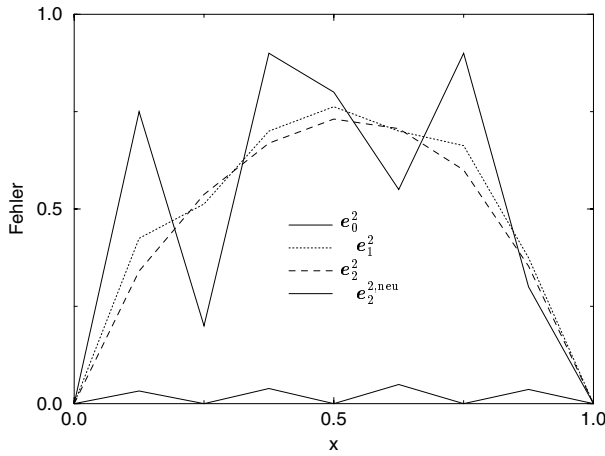


Bild 4.12 Entwicklung des Fehlers beim Zweigitterverfahren

Lemma 4.43 Sind ϕ, ψ zwei lineare Iterationsverfahren mit den Iterationsmatrizen $\mathbf{M}_\phi$ und $\mathbf{M}_\psi$, dann gilt:

- (a) Sind ϕ und ψ konsistent, dann ist auch die Produktiteration $\phi \circ \psi$ konsistent.
- (b) Die Iterationsmatrix der Produktiteration $\phi \circ \psi$ hat die Form

$$\mathbf{M}_{\phi \circ \psi} = \mathbf{M}_\phi \mathbf{M}_\psi.$$

- (c) Die beiden Produktiterationen $\phi \circ \psi$ und $\psi \circ \phi$ besitzen die gleichen Konvergenzeigenschaften im Sinne von

$$\rho(\mathbf{M}_{\phi \circ \psi}) = \rho(\mathbf{M}_{\psi \circ \phi}).$$

Beweis:

zu (a): Sei $\hat{\mathbf{x}} = \mathbf{A}^{-1}\mathbf{b}$, dann folgt die Konsistenz aus

$$(\phi \circ \psi)(\hat{\mathbf{x}}, \mathbf{b}) = \phi(\psi(\hat{\mathbf{x}}, \mathbf{b}), \mathbf{b}) = \phi(\hat{\mathbf{x}}, \mathbf{b}) = \hat{\mathbf{x}}.$$

zu (b): Mit $\phi(\mathbf{x}_m, \mathbf{b}) = \mathbf{M}_\phi \mathbf{x}_m + \mathbf{N}_\phi \mathbf{b}$ und $\psi(\mathbf{x}_m, \mathbf{b}) = \mathbf{M}_\psi \mathbf{x}_m + \mathbf{N}_\psi \mathbf{b}$ folgt die Behauptung gemäß

$$(\phi \circ \psi)(\mathbf{x}_m, \mathbf{b}) = \mathbf{M}_\phi (\mathbf{M}_\psi \mathbf{x}_m + \mathbf{N}_\psi \mathbf{b}) + \mathbf{N}_\phi \mathbf{b} = \underbrace{\mathbf{M}_\phi \mathbf{M}_\psi}_{=\mathbf{M}_{\phi \circ \psi}} \mathbf{x}_m + \underbrace{(\mathbf{M}_\phi \mathbf{N}_\psi + \mathbf{N}_\phi)}_{=\mathbf{N}_{\phi \circ \psi}} \mathbf{b}.$$

zu (c): Sei $\mathbf{x}$ Eigenvektor der Matrix $\mathbf{M}_\phi \mathbf{M}_\psi$ zum Eigenwert $\lambda \neq 0$, so ergibt sich aus $\mathbf{M}_\phi \mathbf{M}_\psi \mathbf{x} = \lambda \mathbf{x} \neq \mathbf{0}$ die Eigenschaft $\mathbf{M}_\psi \mathbf{x} \neq \mathbf{0}$, wodurch aufgrund der Gleichung $\mathbf{M}_\psi \mathbf{M}_\phi \mathbf{M}_\psi \mathbf{x} = \lambda \mathbf{M}_\psi \mathbf{x}$ stets $\rho(\mathbf{M}_\phi \mathbf{M}_\psi) \leq \rho(\mathbf{M}_\psi \mathbf{M}_\phi)$ gilt. Analog erhalten wir $\rho(\mathbf{M}_\psi \mathbf{M}_\phi) \leq \rho(\mathbf{M}_\phi \mathbf{M}_\psi)$. Damit stimmen die Spektralradien der Iterationsmatrizen und somit auch das Konvergenzverhalten der Produktiterationen $\phi \circ \psi$ und $\psi \circ \phi$ überein.

□

Lineare Iterationsverfahren wie zum Beispiel die Jacobi-Methode wirken sich glättend auf den Fehlerverlauf aus und werden daher im folgenden als *Glätter* bezeichnet. Seien $\nu \in \mathbb{N}$ die Anzahl der Iterationsschritte auf dem feinen Gitter Ω_ℓ und ϕ_ℓ das Glättungsverfahren, dann erhalten wir das Zweigitterverfahren als Produktiteration in der Form

$$\phi_\ell^{ZGM} = \phi_\ell^{GK} \circ \phi_\ell^\nu. \quad (4.2.18)$$

Bezeichnet R die Restriktion, P die Prolongation und E das exakte Lösen des Gleichungssystems, dann läßt sich (4.2.18) mit dem Glätter G gemäß der in der Abbildung 4.13 dargestellten Form visualisieren.

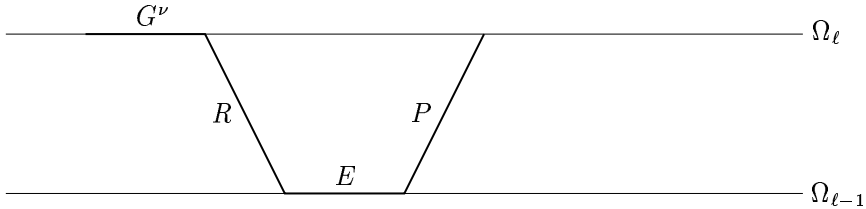


Bild 4.13 Zweigitterverfahren ohne Nachglättung

Seien $\nu_1, \nu_2 \in \mathbb{N}$ mit $\nu = \nu_1 + \nu_2$, dann liegt laut Lemma 4.43 mit

$$\phi_\ell^{ZGM(\nu_1, \nu_2)} := \phi_\ell^{\nu_2} \circ \phi_\ell^{GK} \circ \phi_\ell^{\nu_1} \quad (4.2.19)$$

ein Verfahren vor, das die gleichen Konvergenzeigenschaften wie (4.2.18) aufweist. Man spricht hierbei von ν_1 Vor- und ν_2 Nachglättungen. Wir erhalten somit die in Abbildung 4.14 präsentierte graphische Darstellung.

Lemma 4.44 Sei ϕ_ℓ ein konsistentes Iterationsverfahren mit Iterationsmatrix $\mathbf{M}_\ell$, dann ist das Zweigitteriterationsverfahren $\phi_\ell^{ZGM(\nu_1, \nu_2)}$ konsistent mit der Iterationsmatrix

$$\mathbf{M}_\ell^{ZGM(\nu_1, \nu_2)} = \mathbf{M}_\ell^{\nu_2} \left(\mathbf{I} - \mathbf{P}_{\ell-1}^\ell \mathbf{A}_{\ell-1}^{-1} \mathbf{R}_\ell^{\ell-1} \mathbf{A}_\ell \right) \mathbf{M}_\ell^{\nu_1}.$$

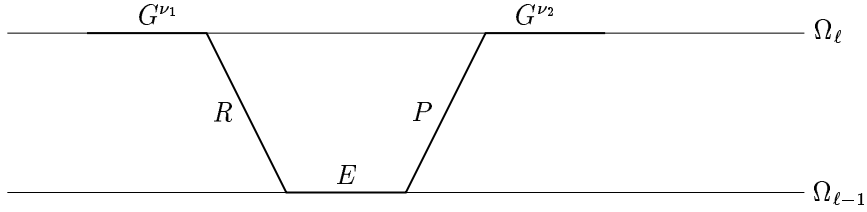


Bild 4.14 Zweigitterverfahren mit Nachglättung

Beweis:

Laut Lemma 4.40 ist die Grobgitterkorrektur ϕ_ℓ^{GK} konsistent mit Iterationsmatrix $M_\ell^{GK} = I - P_{\ell-1}^\ell A_{\ell-1}^{-1} R_\ell^{\ell-1} A_\ell$. Lemma 4.43 (a) liefert somit die behauptete Konsistenz, und mit Lemma 4.43 (b) folgt

$$M_\ell^{ZGM(\nu_1, \nu_2)} = M_{\phi_\ell^{\nu_2} \circ \phi_\ell^{GK} \circ \phi_\ell^{\nu_1}} = M_\ell^{\nu_2} \left(I - P_{\ell-1}^\ell A_{\ell-1}^{-1} R_\ell^{\ell-1} A_\ell \right) M_\ell^{\nu_1}.$$

□

Zur Vor- und Nachiteration können hierbei natürlich auch unterschiedliche Glättungsalgorithmen genutzt werden.

In der Form eines Diagramms läßt sich das Zweigitterverfahren wie folgt schreiben:

Algorithmus - Zweigitterverfahren —

Für $i = 1, \dots, \nu_1$
$u^\ell := \phi_\ell(u^\ell, f^\ell)$
$d^{\ell-1} := R_\ell^{\ell-1} (A_\ell u^\ell - f^\ell)$
$e^{\ell-1} := A_{\ell-1}^{-1} d^{\ell-1}$
$u^\ell := u^\ell - P_{\ell-1}^\ell e^{\ell-1}$
Für $i = 1, \dots, \nu_2$
$u^\ell := \phi_\ell(u^\ell, f^\ell)$

4.2.2 Der Mehrgitteralgorithmus

Das Zweigitterverfahren hat sich als effizient herausgestellt. Es ist jedoch für große Systeme unpraktikabel, da es die exakte Lösung der Korrekturgleichung

$$\mathbf{A}_{\ell-1}\mathbf{e}^{\ell-1} = \mathbf{d}^{\ell-1} \quad (4.2.20)$$

auf $\Omega_{\ell-1}$ benötigt. Da aber die Prolongation der exakten Lösung $\mathbf{e}^{\ell-1}$ auf das feine Gitter Ω_ℓ nur eine Näherung an den gesuchten Fehlervektor

$$\mathbf{e}^\ell = \mathbf{u}^\ell - \mathbf{u}^{\ell,*}$$

liefert, erweist sich auch eine approximative Lösung der Korrekturgleichung als ausreichend.

Die Gleichung (4.2.20) weist die gleiche Form wie die Ausgangsgleichung $\mathbf{A}_\ell \mathbf{u}^\ell = \mathbf{f}^\ell$ auf. Die Idee liegt daher in der Nutzung einer Zweigittermethode auf $\Omega_{\ell-1}$ und $\Omega_{\ell-2}$ zur Approximation von $\mathbf{e}^{\ell-1} = \mathbf{A}_{\ell-1}^{-1} \mathbf{d}^{\ell-1}$. Damit erhalten wir eine Dreigittermethode, bei der $\mathbf{A}_{\ell-2} \mathbf{e}^{\ell-2} = \mathbf{d}^{\ell-2}$ exakt gelöst werden muß. Sukzessives Fortsetzen dieser Idee liefert ein Verfahren auf $\ell + 1$ Gittern $\Omega_\ell, \dots, \Omega_0$ bei dem lediglich $\mathbf{A}_0 \mathbf{e}^0 = \mathbf{d}^0$ exakt gelöst werden muß. Bei der von uns gewählten Gitterverfeinerung mit $h_{\ell-1} = 2h_\ell$ gilt $\mathbf{A}_0 \in \mathbb{R}^{1 \times 1}$. In der Praxis wird in der Regel mit Ω_0 ein Gitter genutzt, das eine approximative Lösung des Gleichungssystems mit der Matrix $\mathbf{A}_0$ auf effiziente und einfache Weise ermöglicht.

Der Mehrgitteralgorithmus läßt sich folglich als rekursives Verfahren in der anschließenden Form darstellen:

Algorithmus - Mehrgitterverfahren $\phi_\ell^{\text{MGM}(\nu_1, \nu_2)}(\mathbf{u}^\ell, \mathbf{f}^\ell) \text{ —}$

Y $\ell = 0$	N
$\mathbf{u}^0 := \mathbf{A}_0^{-1} \mathbf{f}^0$	Für $i = 1, \dots, \nu_1$
Rückgabe von $\mathbf{u}^0$	$\mathbf{u}^\ell := \phi_\ell(\mathbf{u}^\ell, \mathbf{f}^\ell)$
	$\mathbf{d}^{\ell-1} := \mathbf{R}_\ell^{\ell-1}(\mathbf{A}_\ell \mathbf{u}^\ell - \mathbf{f}^\ell)$
	$\mathbf{e}_0^{\ell-1} := \mathbf{0}$
	Für $i = 1, \dots, \gamma$
	$\mathbf{e}_i^{\ell-1} := \phi_{\ell-1}^{\text{MGM}(\nu_1, \nu_2)}(\mathbf{e}_{i-1}^{\ell-1}, \mathbf{d}^{\ell-1})$
	$\mathbf{u}^\ell := \mathbf{u}^\ell - \mathbf{P}_{\ell-1}^\ell \mathbf{e}_\gamma^{\ell-1}$
	Für $i = 1, \dots, \nu_2$
	$\mathbf{u}^\ell := \phi_\ell(\mathbf{u}^\ell, \mathbf{f}^\ell)$
	Rückgabe von $\mathbf{u}^\ell$

In der hier gewählten Darstellung werden zur iterativen Lösung der Grobgittergleichung γ Schritte verwendet. In der Praxis erweisen sich in der Regel $\gamma = 1$ respektive $\gamma = 2$ als geeignet.

Der Fall $\gamma = 1$ liefert den sogenannten V-Zyklus, der sich für $\ell = 3$ graphisch gemäß Abbildung 4.15 darstellen läßt, während für $\gamma = 2$ die als W-Zyklus bezeichnete Iterationsfolge vorliegt. Der für diese Vorgehensweise entstehende algorithmische Ablauf des Verfahrens ist für den Fall von vier genutzten Gittern in Abbildung 4.16 dargestellt.

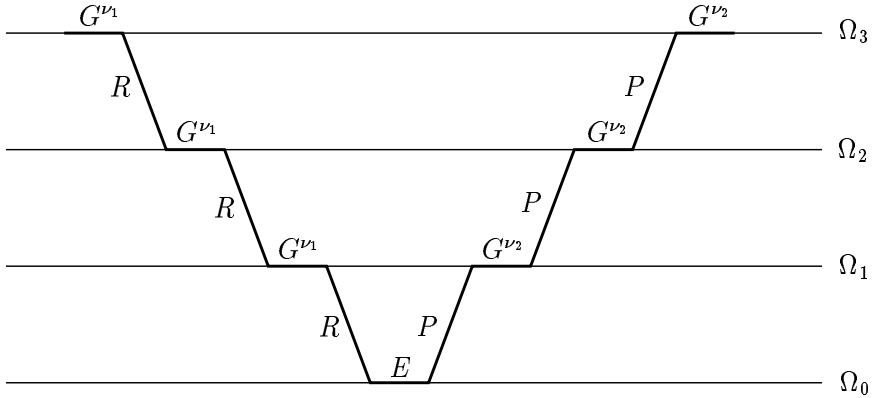


Bild 4.15 Mehrgitterverfahren mit V-Zyklus

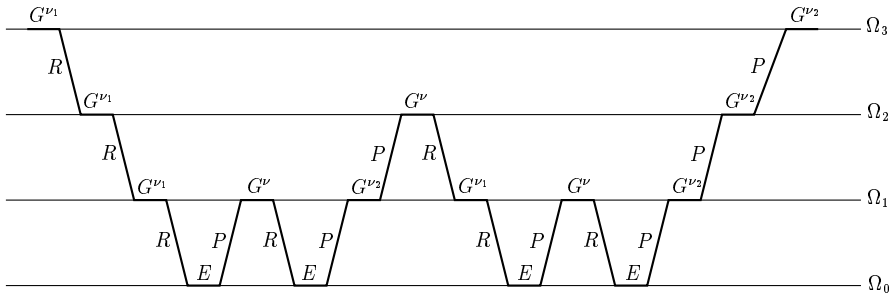


Bild 4.16 Mehrgitterverfahren mit W-Zyklus mit $\nu = \nu^1 + \nu^2$

4.2.3 Das vollständige Mehrgitterverfahren

Die Idee des vollständigen Mehrgitterverfahrens (nested iterations, full multigrid method (FMGM)) liegt in der Nutzung der gröberen Gitter zur Verbesserung der Startnäherung $\mathbf{u}_0^\ell$ auf dem feinsten Gitter Ω_ℓ .

Die Durchführung läßt sich wie folgt beschreiben:

- Löse auf Ω_0 die Gleichung $\mathbf{A}_0 \mathbf{u}^0 = \mathbf{f}^0$ exakt.
- Prolongiere das Ergebnis auf das nächstfeinere Gitter und führe einige Iterationsschritte mit dem Glättungsverfahren durch.

- (c) Wiederhole (b), bis eine Näherungslösung $\mathbf{u}_0^\ell$ auf Ω_ℓ ermittelt wurde.
- (d) Führe das Mehrgitterverfahren mit der Startnäherung $\mathbf{u}_0^\ell$ durch.

Graphisch erhalten wir für den V-Zyklus auf 3 Gittern die Darstellung des vollständigen Mehrgitterverfahrens in der in Abbildung 4.17 präsentierten Form.

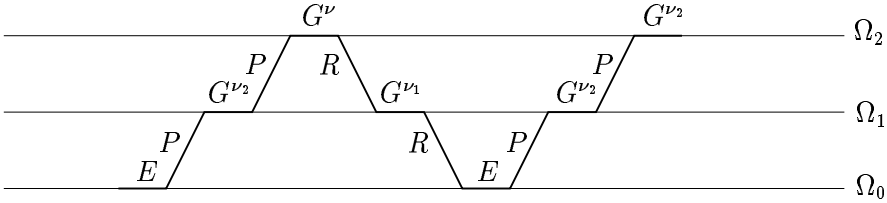


Bild 4.17 Vollständiges Mehrgitterverfahren mit $\nu = \nu^1 + \nu^2$

Varianten des Verfahrens erhält man zum Beispiel durch

- (a) eine approximative anstelle einer exakten Lösung der Gleichung $\mathbf{A}_0 \mathbf{u}^0 = \mathbf{f}^0$.
- (b) die Nutzung unterschiedlicher Anzahlen von Glättungsschritten.
- (c) die Verwendung verschiedener Prolongationen und Restriktionen.

4.3 Projektionsmethoden und Krylov-Unterraum-Verfahren

Wir betrachten in diesem Abschnitt stets lineare Gleichungssysteme der Form

$$\mathbf{A}\mathbf{x} = \mathbf{b} \quad (4.3.1)$$

mit einer regulären Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ und einer rechten Seite $\mathbf{b} \in \mathbb{R}^n$.

Definition 4.45 Eine Projektionsmethode zur Lösung der Gleichung (4.3.1) ist ein Verfahren zur Berechnung von Näherungslösungen $\mathbf{x}_m \in \mathbf{x}_0 + K_m$ unter Berücksichtigung der Bedingung

$$(\mathbf{b} - \mathbf{A}\mathbf{x}_m) \perp L_m, \quad (4.3.2)$$

wobei $\mathbf{x}_0 \in \mathbb{R}^n$ beliebig ist und K_m sowie L_m m -dimensionale Unterräume des $\mathbb{R}^n$ repräsentieren. Die Orthogonalitätsbedingung ist hierbei durch das euklidische Skalarprodukt mittels

$$\mathbf{x} \perp \mathbf{y} \Leftrightarrow (\mathbf{x}, \mathbf{y})_2 = 0$$

definiert.

Gilt $K_m = L_m$, so besagt (4.3.2), daß der Residuenvektor $\mathbf{r}_m = \mathbf{b} - \mathbf{A}\mathbf{x}_m$ senkrecht auf K_m steht. In diesem Fall liegt daher eine orthogonale Projektionsmethode vor, und (4.3.2) heißt Galerkin-Bedingung.

Für $K_m \neq L_m$ liegt eine schiefe Projektionsmethode vor, und (4.3.2) wird als Petrov-Galerkin-Bedingung bezeichnet.

Beispiel 4.46 Jeder Schritt des Gauß-Seidel-Verfahrens kann als orthogonale Projektionsmethode mit $\mathbf{x}_0 = (x_{m+1,1}, \dots, x_{m+1,i-1}, 0, x_{m,i+1}, \dots, x_{m,n})^T$ und den eindimensionalen Räumen $K = L = \text{span}\{\mathbf{e}_i\}$ interpretiert werden, denn es gilt

$$\begin{aligned} x_i &= -\frac{1}{a_{ii}} \left(\sum_{j=1}^{i-1} a_{ij} x_{m+1,j} + \sum_{j=i+1}^n a_{ij} x_{m,j} - b_i \right) \\ \Leftrightarrow b_i - (\mathbf{A}\mathbf{x})_i &= 0 \quad \text{mit} \quad \mathbf{x} \in \mathbf{x}_0 + K \\ \Leftrightarrow \mathbf{b} - \mathbf{A}\mathbf{x} &\perp L \quad \text{mit} \quad \mathbf{x} \in \mathbf{x}_0 + K. \end{aligned}$$

Diese Eigenschaft läßt sich auch für weitere iterative Verfahren dieser Art nachweisen.

Generell unterscheiden sich Splitting-Methoden jedoch wesentlich von den Projektionsmethoden. Eine Gegenüberstellung dieser beiden Verfahrensklassen gibt die Tabelle 4.1.

Splitting-Methoden	Projektionsmethoden
Berechnung von Näherungslösungen $\mathbf{x}_m \in \mathbb{R}^n$	Berechnung von Näherungslösungen $\mathbf{x}_m \in \mathbf{x}_0 + K_m \subset \mathbb{R}^n$ $\dim K_m = m \leq n$
Berechnungsvorschrift $\mathbf{x}_m = \mathbf{M}\mathbf{x}_{m-1} + \mathbf{N}\mathbf{b}$	Berechnungsvorschrift (Orthogonalitätsbed.) $\mathbf{b} - \mathbf{A}\mathbf{x}_m \perp L_m \subset \mathbb{R}^n$ $\dim L_m = m \leq n$

Tabelle 4.1 Gegenüberstellung von Splitting-Methoden und Projektionsverfahren

Definition 4.47 Eine Krylov-Unterraum-Methode ist eine Projektionsmethode zur Lösung der Gleichung (4.3.1), bei der K_m den *Krylov-Unterraum*

$$K_m = K_m(\mathbf{A}, \mathbf{r}_0) = \text{span}\{\mathbf{r}_0, \mathbf{A}\mathbf{r}_0, \dots, \mathbf{A}^{m-1}\mathbf{r}_0\}$$

mit $\mathbf{r}_0 = \mathbf{b} - \mathbf{A}\mathbf{x}_0$ darstellt.

Krylov-Unterraum-Methoden werden oftmals durch eine Umformulierung des linearen Gleichungssystems in eine Minimierungsaufgabe beschrieben. Zwei der bekanntesten Vertreter dieser Algorithmengruppe sind das von Hestenes und Stiefel [37] entwickelte Verfahren der konjugierten Gradienten und die von Saad und Schulz [62] hergeleitete GMRES-Methode. Beide Verfahren ermitteln die optimale Approximation $\mathbf{x}_m \in \mathbf{x}_0 + K_m$ an die gesuchte Lösung $\mathbf{A}^{-1}\mathbf{b}$ im Sinne der Orthogonalitätsbedingung (4.3.2), wobei bei jeder Iteration die Dimension des Unterraums um eins inkrementiert wird. Vernachlässigt man die auftretenden Rundungsfehler, so würden beide Methoden spätestens nach n Iterationen die exakte Lösung liefern.

Bevor wir uns mit der Herleitung der einzelnen Verfahren befassen, werden wir zunächst eine allgemeingültige Konvergenzaussage für Krylov-Unterraum-Methoden formulieren. Hierzu erweist sich das folgende Lemma als hilfreich.

Lemma 4.48 *Gegeben sei eine Projektionsmethode zur Lösung der Gleichung $\mathbf{Ax} = \mathbf{b}$ mit einer regulären Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$. Sei $m \in \mathbb{N}$ fest, und bilden die Spaltenvektoren der Matrix $\mathbf{V}_m \in \mathbb{R}^{n \times m}$ beziehungsweise $\mathbf{W}_m \in \mathbb{R}^{n \times m}$ eine Basis der Räume K_m respektive L_m derart, daß $\mathbf{W}_m^T \mathbf{AV}_m \in \mathbb{R}^{m \times m}$ regulär ist, dann besitzt die Lösung der Projektionsmethode die Darstellung*

$$\mathbf{x}_m = \mathbf{x}_0 + \mathbf{V}_m \left(\mathbf{W}_m^T \mathbf{AV}_m \right)^{-1} \mathbf{W}_m^T \mathbf{r}_0.$$

Beweis:

Aufgrund der Basiseigenschaft der Spaltenvektoren der Matrix $\mathbf{V}_m$ läßt sich der Lösungsvektor in der Form $\mathbf{x}_m = \mathbf{x}_0 + \mathbf{V}_m \boldsymbol{\alpha}_m$ mit $\boldsymbol{\alpha}_m \in \mathbb{R}^m$ schreiben. Aus der Orthogonalitätsbedingung (4.3.2) erhalten wir

$$\mathbf{W}_m^T (\mathbf{b} - \mathbf{A}(\mathbf{x}_0 + \mathbf{V}_m \boldsymbol{\alpha}_m)) = \mathbf{0},$$

wodurch

$$\mathbf{W}_m^T \mathbf{AV}_m \boldsymbol{\alpha}_m = \mathbf{W}_m^T (\mathbf{b} - \mathbf{Ax}_0)$$

folgt. Aufgrund der vorausgesetzten Regularität der Matrix $\mathbf{W}_m^T \mathbf{AV}_m$ gilt daher

$$\boldsymbol{\alpha}_m = \left(\mathbf{W}_m^T \mathbf{AV}_m \right)^{-1} \mathbf{W}_m^T \mathbf{r}_0,$$

wodurch sich die behauptete Gestalt der Iterierten ergibt. $\square$

Aus dem obigen Lemma erhalten wir unmittelbar die Darstellung des zugehörigen Residuenvektors in der Form

$$\mathbf{r}_m = \mathbf{b} - \mathbf{Ax}_m = \mathbf{r}_0 - \mathbf{AV}_m \left(\mathbf{W}_m^T \mathbf{AV}_m \right)^{-1} \mathbf{W}_m^T \mathbf{r}_0. \quad (4.3.3)$$

Wir kommen nun zum angekündigten Konvergenzsatz, der in [38] vorgestellt wurde.

Satz 4.49 *Sei die Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ regulär. Desweiteren bezeichnen $\mathbf{v}_1, \dots, \mathbf{v}_m \in \mathbb{R}^n$ und $\mathbf{w}_1, \dots, \mathbf{w}_m \in \mathbb{R}^n$ die durch ein beliebiges Krylov-Unterraum-Verfahren erzeugten Basisvektoren des K_m und L_m . Liegt mit den Matrizen $\mathbf{V}_m = (\mathbf{v}_1 \dots \mathbf{v}_m) \in \mathbb{R}^{n \times m}$ und $\mathbf{W}_m = (\mathbf{w}_1 \dots \mathbf{w}_m) \in \mathbb{R}^{n \times m}$ eine reguläre Matrix $\mathbf{W}_m^T \mathbf{AV}_m \in \mathbb{R}^{m \times m}$ vor, dann folgen mit der Projektion*

$$\mathbf{P}_m = \mathbf{I} - \mathbf{AV}_m \left(\mathbf{W}_m^T \mathbf{AV}_m \right)^{-1} \mathbf{W}_m^T \quad (4.3.4)$$

die Abschätzungen für den Fehlervektor $\mathbf{e}_m = \mathbf{A}^{-1} \mathbf{b} - \mathbf{x}_m$ und den Residuenvektor $\mathbf{r}_m = \mathbf{A} \mathbf{e}_m$ der Krylov-Unterraum-Methode in der Form

$$\|\mathbf{e}_m\| \leq \|\mathbf{A}^{-1} \mathbf{P}_m\| \min_{p \in \mathcal{P}_m^1} \|p(\mathbf{A}) \mathbf{r}_0\| \quad (4.3.5)$$

und

$$\|\mathbf{r}_m\| \leq \|\mathbf{P}_m\| \min_{p \in \mathcal{P}_m^1} \|p(\mathbf{A}) \mathbf{r}_0\|, \quad (4.3.6)$$

wobei $\mathcal{P}_m^1$ die Menge aller Polynome p vom Höchstgrad m bezeichnet, die zudem die Nebenbedingung $p(\mathbf{0}) = \mathbf{I}$ erfüllen.

Beweis:

Aus der Definition der Projektion $\mathbf{P}_m$ folgt unter Berücksichtigung der Regularität der Matrix $\mathbf{W}_m^T \mathbf{A} \mathbf{V}_m$ die Gleichung

$$\mathbf{P}_m \mathbf{A} \mathbf{V}_m = \mathbf{A} \mathbf{V}_m - \mathbf{A} \mathbf{V}_m \left(\mathbf{W}_m^T \mathbf{A} \mathbf{V}_m \right)^{-1} \mathbf{W}_m^T \mathbf{A} \mathbf{V}_m = \mathbf{0}. \quad (4.3.7)$$

Unter Verwendung der Gleichung (4.3.3) ergibt sich $\mathbf{r}_m = \mathbf{P}_m \mathbf{r}_0$, wodurch mit (4.3.7) zudem

$$\mathbf{r}_m = \mathbf{P}_m (\mathbf{r}_0 + \mathbf{A} \mathbf{V}_m \boldsymbol{\alpha})$$

für jeden Vektor $\boldsymbol{\alpha} \in \mathbb{R}^m$ gilt. Wegen $\mathbf{A} \mathbf{V}_m \boldsymbol{\alpha} \in \mathbf{A} K_m$ erhalten wir

$$\mathbf{r}_m = \mathbf{P}_m p(\mathbf{A}) \mathbf{r}_0$$

für jedes beliebige Polynom $p \in \mathcal{P}_m^1$. Hierdurch folgt

$$\|\mathbf{r}_m\| = \min_{p \in \mathcal{P}_m^1} \|\mathbf{P}_m p(\mathbf{A}) \mathbf{r}_0\| \leq \|\mathbf{P}_m\| \min_{p \in \mathcal{P}_m^1} \|p(\mathbf{A}) \mathbf{r}_0\|.$$

Analog liefert $\mathbf{e}_m = \mathbf{A}^{-1} \mathbf{r}_m$ die Ungleichung (4.3.5). $\square$

Die im obigen Satz geforderte Regularität der Matrix $\mathbf{W}_m^T \mathbf{A} \mathbf{V}_m$ läßt sich für spezielle Krylov-Unterraum-Verfahren direkt nachweisen. Für eine symmetrische, positiv definite Matrix $\mathbf{A}$ existiert laut Korollar 2.32 eine orthogonale Matrix $\mathbf{U}$ mit $\mathbf{U}^T \mathbf{A} \mathbf{U} = \mathbf{D} = \text{diag}\{\lambda_1, \dots, \lambda_n\}$. Aus der positiven Definitheit folgt $\lambda_i > 0$ für $i = 1, \dots, n$ und wir erhalten $\mathbf{A} = \mathbf{U} \mathbf{D}^{1/2} \mathbf{D}^{1/2} \mathbf{U}^T$. Betrachten wir hierbei eine orthogonale Krylov-Unterraum-Methode, so können wegen $L_m = K_m$ die Matrizen $\mathbf{W}_m = \mathbf{V}_m$ genutzt werden, wodurch sich

$$\mathbf{W}_m^T \mathbf{A} \mathbf{V}_m = \left(\mathbf{D}^{1/2} \mathbf{U} \mathbf{V}_m \right)^T \left(\mathbf{D}^{1/2} \mathbf{U} \mathbf{V}_m \right) \in \mathbb{R}^{m \times m}$$

ergibt. Da die Vektoren der Matrix $\mathbf{V}_m$ eine Basis von K_m darstellen, erhalten wir $\text{rang}(\mathbf{V}_m) = m$. Hierdurch ist $\mathbf{D}^{1/2} \mathbf{U} \mathbf{V}_m$ injektiv, so daß

$$\left(\mathbf{x}, \left(\mathbf{D}^{1/2} \mathbf{U} \mathbf{V}_m \right)^T \left(\mathbf{D}^{1/2} \mathbf{U} \mathbf{V}_m \right) \mathbf{x} \right)_2 = \left\| \left(\mathbf{D}^{1/2} \mathbf{U} \mathbf{V}_m \right) \mathbf{x} \right\|_2^2 \neq 0 \quad \forall \mathbf{x} \in \mathbb{R}^m \setminus \{\mathbf{0}\}$$

die Regularität der Matrix $\mathbf{W}_m^T \mathbf{A} \mathbf{V}_m$ liefert. Eine derartige Methode stellt zum Beispiel das im folgenden betrachtete Verfahren der konjugierten Gradienten dar.

Für reguläre Matrizen $\mathbf{A} \in \mathbb{R}^{n \times n}$ betrachten wir $L_m = \mathbf{A} K_m$, wodurch $\mathbf{W}_m = \mathbf{A} \mathbf{V}_m$ gewählt werden kann. Analog zu der obigen Betrachtung ergibt sich die Regularität der Matrix $\mathbf{W}_m^T \mathbf{A} \mathbf{V}_m$ aus

$$\mathbf{W}_m^T \mathbf{A} \mathbf{V}_m = (\mathbf{A} \mathbf{V}_m)^T \mathbf{A} \mathbf{V}_m.$$

Diese Bedingungen erfüllt das im Abschnitt 4.3.2.4 hergeleitete GMRES-Verfahren.

Generell ergibt sich für $m = n$ die Regularität der Matrizen $\mathbf{V}_m$ und $\mathbf{W}_m$, so daß $\mathbf{P}_m = \mathbf{0}$ gilt, und wir aus den Abschätzungen (4.3.5) und (4.3.6) jeweils die Aussage erhalten, daß derartige Krylov-Unterraum-Verfahren spätestens nach n Schritten die exakte Lösung ermitteln.

Die Entwicklung modernster numerischer Algorithmen zur Simulation praxisrelevanter Problemstellungen hat zu einer großen Nachfrage hinsichtlich effizienter, schneller und robuster iterativer Gleichungssystemlöser geführt und einen wesentlichen Impuls zum in den letzten zwanzig Jahren erzielten Fortschritt im Bereich der Krylov-Unterraum-Verfahren beigetragen. Aufgrund der Vielzahl dieser Verfahren werden wir uns im vorliegenden Buch neben dem Verfahren der konjugierten Gradienten im wesentlichen auf die sehr häufig verwendeten Methoden GMRES, CGS, BiCGSTAB, TFQMR und QMRCGSTAB beschränken und die Einordnung weiterer Algorithmen in den vorliegenden Rahmen soweit möglich durch Anmerkungen vornehmen. Der Zusammenhang zwischen den genannten Algorithmen wird in der Abbildung 4.18 schematisch verdeutlicht. Hinsichtlich weiterer Literaturstellen zu dieser Verfahrensklasse sei auf die Bücher von Axelsson [7], Demmel [18], Fischer [24], Greenbaum [32], Kelley [42], Meurant [51], Saad [61], Trefethen und Bau [69], van der Vorst [72] sowie Weiss [74] verwiesen.

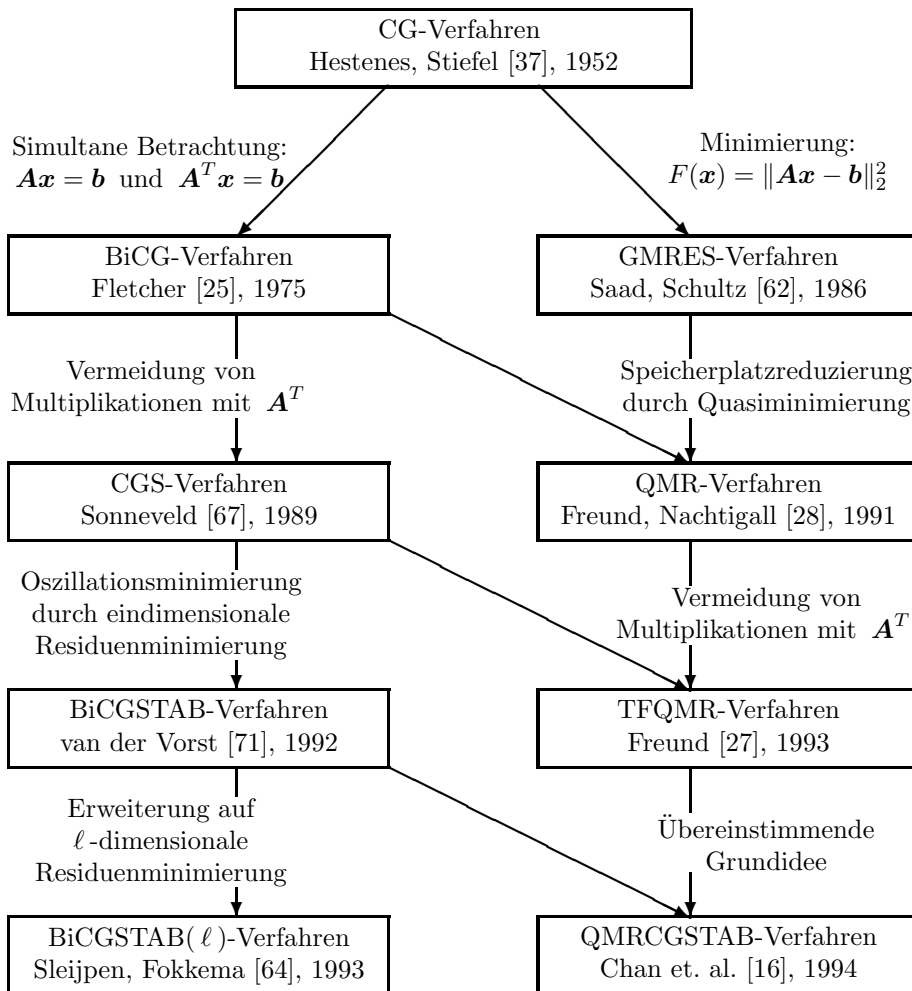


Bild 4.18 Zusammenhänge zwischen Krylov-Unterraum-Verfahren

4.3.1 Verfahren für symmetrische, positiv definite Matrizen

Innerhalb dieses Abschnitts setzen wir voraus, daß das betrachtete Gleichungssystem (4.3.1) eine symmetrische und positiv definite Matrix aufweist. Zur Herleitung der Verfahren betrachten wir die Funktion

$$\begin{aligned} F : \mathbb{R}^n &\rightarrow \mathbb{R} \\ \mathbf{x} &\mapsto \frac{1}{2}(\mathbf{A}\mathbf{x}, \mathbf{x})_2 - (\mathbf{b}, \mathbf{x})_2 \end{aligned} \quad (4.3.8)$$

und werden zunächst einige ihrer grundlegenden Eigenschaften im Zusammenhang mit dem betrachteten Gleichungssystem studieren.

Lemma 4.50 *Seien $\mathbf{A} \in \mathbb{R}^{n \times n}$ symmetrisch, positiv definit und $\mathbf{b} \in \mathbb{R}^n$ gegeben, dann gilt mit der durch (4.3.8) gegebenen Funktion F*

$$\hat{\mathbf{x}} = \arg \min_{\mathbf{x} \in \mathbb{R}^n} F(\mathbf{x})$$

genau dann, wenn

$$\mathbf{A}\hat{\mathbf{x}} = \mathbf{b}$$

gilt.

Beweis:

Da $\mathbf{A}$ positiv definit ist, existiert die positiv definite Inverse $\mathbf{A}^{-1} \in \mathbb{R}^{n \times n}$, und wir erhalten

$$\arg \min_{\mathbf{x} \in \mathbb{R}^n} F(\mathbf{x}) = \arg \min_{\mathbf{x} \in \mathbb{R}^n} G(\mathbf{x}) \quad (4.3.9)$$

mit

$$\begin{aligned} G(\mathbf{x}) &= F(\mathbf{x}) + \frac{1}{2}\mathbf{b}^T \mathbf{A}^{-1} \mathbf{b} \\ &= \frac{1}{2}(\mathbf{A}\mathbf{x}, \mathbf{x})_2 - (\mathbf{b}, \mathbf{x})_2 + \frac{1}{2}\mathbf{b}^T \mathbf{A}^{-1} \mathbf{b} \\ &= \frac{1}{2}(\mathbf{A}\mathbf{x} - \mathbf{b})^T \mathbf{x} - \frac{1}{2}\mathbf{b}^T (\mathbf{x} - \mathbf{A}^{-1} \mathbf{b}) \\ &= \frac{1}{2}(\mathbf{A}\mathbf{x} - \mathbf{b})^T \mathbf{x} - \frac{1}{2}(\mathbf{A}^{-1} \mathbf{b})^T (\mathbf{A}\mathbf{x} - \mathbf{b}) \\ &= \frac{1}{2}(\mathbf{A}\mathbf{x} - \mathbf{b})^T \mathbf{A}^{-1} (\mathbf{A}\mathbf{x} - \mathbf{b}). \end{aligned} \quad (4.3.10)$$

Somit gelten $G(\mathbf{x}) \geq 0$ und

$$G(\mathbf{x}) = 0 \quad \Leftrightarrow \quad \mathbf{x} = \mathbf{A}^{-1} \mathbf{b}.$$

□

Bei allen in diesem Abschnitt betrachteten Verfahren nehmen wir eine sukzessive Minimierung der Funktion F ausgehend vom Punkt $\mathbf{x} \in \mathbb{R}^n$ entlang spezieller Richtungen $\mathbf{p} \in \mathbb{R}^n$ vor. Wir definieren daher für $\mathbf{x}, \mathbf{p} \in \mathbb{R}^n$ die Funktion

$$\begin{aligned} f_{\mathbf{x}, \mathbf{p}} : \mathbb{R} &\rightarrow \mathbb{R} \\ \lambda &\mapsto f_{\mathbf{x}, \mathbf{p}}(\lambda) := F(\mathbf{x} + \lambda \mathbf{p}). \end{aligned} \quad (4.3.11)$$

Lemma und Definition 4.51 Seien die Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ symmetrisch, positiv definit und die Vektoren $\mathbf{x}, \mathbf{p} \in \mathbb{R}^n$ mit $\mathbf{p} \neq \mathbf{0}$ gegeben, dann gilt

$$\lambda_{\text{opt}} = \lambda_{\text{opt}}(\mathbf{x}, \mathbf{p}) := \arg \min_{\lambda \in \mathbb{R}} f_{\mathbf{x}, \mathbf{p}}(\lambda) = \frac{(\mathbf{r}, \mathbf{p})_2}{(\mathbf{A}\mathbf{p}, \mathbf{p})_2}$$

mit $\mathbf{r} := \mathbf{b} - \mathbf{A}\mathbf{x}$. Der Vektor $\mathbf{r}$ wird als Residuenvektor und seine euklidische Norm $\|\mathbf{r}\|_2$ als Residuum bezeichnet.

Beweis:

Es gilt

$$\begin{aligned} f_{\mathbf{x}, \mathbf{p}}(\lambda) &= \frac{1}{2} (\mathbf{A}(\mathbf{x} + \lambda\mathbf{p}), \mathbf{x} + \lambda\mathbf{p})_2 - (\mathbf{b}, \mathbf{x} + \lambda\mathbf{p})_2 \\ &= F(\mathbf{x}) + \lambda(\mathbf{A}\mathbf{x} - \mathbf{b}, \mathbf{p})_2 + \frac{1}{2}\lambda^2(\mathbf{A}\mathbf{p}, \mathbf{p})_2. \end{aligned}$$

Somit folgt

$$f'_{\mathbf{x}, \mathbf{p}}(\lambda) = (\mathbf{A}\mathbf{x} - \mathbf{b}, \mathbf{p})_2 + \lambda \underbrace{(\mathbf{A}\mathbf{p}, \mathbf{p})_2}_{>0}$$

und

$$f'_{\mathbf{x}, \mathbf{p}}(\lambda_{\text{opt}}) = (\mathbf{A}\mathbf{x} - \mathbf{b}, \mathbf{p})_2 + \frac{(\mathbf{b} - \mathbf{A}\mathbf{x}, \mathbf{p})_2}{(\mathbf{A}\mathbf{p}, \mathbf{p})_2} (\mathbf{A}\mathbf{p}, \mathbf{p})_2 = 0.$$

Mit

$$f''_{\mathbf{x}, \mathbf{p}}(\lambda) = (\mathbf{A}\mathbf{p}, \mathbf{p})_2 \quad \begin{matrix} \mathbf{A} \text{ pos. def.} \\ \mathbf{p} \neq \mathbf{0} \\ > \end{matrix} \quad 0 \quad (4.3.12)$$

ist λ_{opt} globales Minimum von $f_{\mathbf{x}, \mathbf{p}}$. $\square$

Ist nun mit $\{\mathbf{p}_m\}_{m \in \mathbb{N}_0}$ eine Folge von Suchrichtungen aus $\mathbb{R}^n \setminus \{\mathbf{0}\}$ gegeben, dann können wir ein erstes Verfahren erstellen:

Algorithmus - Basisl ser —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$
F�r $m = 0, 1, \dots$
$\mathbf{r}_m = \mathbf{b} - \mathbf{A}\mathbf{x}_m$
$\lambda_m = \frac{(\mathbf{r}_m, \mathbf{p}_m)_2}{(\mathbf{A}\mathbf{p}_m, \mathbf{p}_m)_2}$
$\mathbf{x}_{m+1} = \mathbf{x}_m + \lambda_m \mathbf{p}_m$

4.3.1.1 Die Methode des steilsten Abstiegs

Zur Vervollst ndigung des Basisl sers ben tigen wir eine Berechnungsvorschrift zur Ermittlung der Suchrichtungen $\mathbf{p}_m \in \mathbb{R}^n$. Wir fordern zudem o.B.d.A. $\|\mathbf{p}_m\|_2 = 1$. F r $\mathbf{x} \neq \mathbf{A}^{-1}\mathbf{b}$ erhalten wir eine global optimale Wahl durch

$$\mathbf{p} = \frac{\hat{\mathbf{x}} - \mathbf{x}}{\|\hat{\mathbf{x}} - \mathbf{x}\|_2} \quad \text{mit} \quad \hat{\mathbf{x}} = \mathbf{A}^{-1}\mathbf{b},$$

denn hiermit folgt bei Definition von λ_{opt} gemäß Lemma 4.51

$$\begin{aligned} \tilde{\mathbf{x}} &= \mathbf{x} + \lambda_{\text{opt}}\mathbf{p} = \\ &= \mathbf{x} + \|\hat{\mathbf{x}} - \mathbf{x}\|_2 \frac{(\mathbf{b} - \mathbf{A}\mathbf{x}, \hat{\mathbf{x}} - \mathbf{x})_2}{(\mathbf{b} - \mathbf{A}\mathbf{x}, \hat{\mathbf{x}} - \mathbf{x})_2} \frac{\hat{\mathbf{x}} - \mathbf{x}}{\|\hat{\mathbf{x}} - \mathbf{x}\|_2} \\ &= \hat{\mathbf{x}}. \end{aligned}$$

Jedoch benötigen wir hierbei bereits zur Definition der Suchrichtung die exakte Lösung. Beschränken wir uns auf lokale Optimalität, so erhalten wir diese mit dem negativen Gradienten der Funktion F . Hier gilt

$$\nabla F(\mathbf{x}) = \frac{1}{2}(\mathbf{A} + \mathbf{A}^T)\mathbf{x} - \mathbf{b} \stackrel{\mathbf{A} \text{ sym.}}{=} \mathbf{A}\mathbf{x} - \mathbf{b} = -\mathbf{r}.$$

Somit liefert

$$\mathbf{p} := \begin{cases} \frac{\mathbf{r}}{\|\mathbf{r}\|_2} & \text{für } \mathbf{r} \neq 0, \\ \mathbf{0} & \text{für } \mathbf{r} = 0 \end{cases} \quad (4.3.13)$$

die Richtung des steilsten Abstiegs.

Die Funktion F ist wegen

$$\nabla^2 F(\mathbf{x}) = \mathbf{A}$$

und positiv definitem $\mathbf{A}$ strikt konvex. Auch hiermit erkennen wir, daß

$$\hat{\mathbf{x}} = \mathbf{A}^{-1}\mathbf{b}$$

wegen

$$\nabla F(\hat{\mathbf{x}}) = \mathbf{0}$$

das einzige und globale Minimum von F darstellt.

Mit (4.3.13) erhalten wir das auch als Gradientenverfahren bezeichnete Verfahren des steilsten Abstiegs in der folgenden Form:

Algorithmus - Verfahren des steilsten Abstiegs —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$	
Für $m = 0, 1, \dots$	
$\mathbf{r}_m = \mathbf{b} - \mathbf{A}\mathbf{x}_m$	
Y $\mathbf{r}_m \neq \mathbf{0}$ N	
$\lambda_m = \frac{\ \mathbf{r}_m\ _2^2}{(\mathbf{A}\mathbf{r}_m, \mathbf{r}_m)_2}$	$\lambda_m = 0$
$\mathbf{x}_{m+1} = \mathbf{x}_m + \lambda_m \mathbf{r}_m$	

Bemerkung:

In der praktischen Anwendung wird $\mathbf{r}_0$ außerhalb der Schleife ermittelt und innerhalb der Schleife

$$\begin{aligned}\mathbf{r}_{m+1} &= \mathbf{b} - \mathbf{A}\mathbf{x}_{m+1} = \mathbf{b} - \mathbf{A}\mathbf{x}_m - \lambda_m \mathbf{A}\mathbf{r}_m \\ &= \mathbf{r}_m - \lambda_m \mathbf{A}\mathbf{r}_m\end{aligned}$$

verwendet, wodurch pro Iteration eine Matrix-Vektor-Multiplikation vermieden wird. Eine MATLAB-Implementierung ist im Anhang A aufgelistet.

Satz 4.52 *Das Verfahren des steilsten Abstiegs ist konsistent und nicht linear.*

Beweis:

Das Iterationsverfahren $\phi : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}^n$ ist gegeben durch

$$\phi(\mathbf{x}, \mathbf{b}) = (\mathbf{I} - \lambda(\mathbf{x}, \mathbf{b})\mathbf{A})\mathbf{x} + \lambda(\mathbf{x}, \mathbf{b})\mathbf{b}, \quad (4.3.14)$$

wobei $\lambda : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$ in der Form

$$\lambda(\mathbf{x}, \mathbf{b}) = \begin{cases} \frac{\|\mathbf{b} - \mathbf{A}\mathbf{x}\|_2^2}{(\mathbf{A}(\mathbf{b} - \mathbf{A}\mathbf{x}), \mathbf{b} - \mathbf{A}\mathbf{x})_2} & \text{für } \mathbf{b} - \mathbf{A}\mathbf{x} \neq \mathbf{0} \\ 0 & \text{sonst.} \end{cases}$$

festgelegt ist. Einfaches Nachrechnen an kleinen Beispielen zeigt, daß λ nicht konstant ist. Für $\hat{\mathbf{x}} = \mathbf{A}^{-1}\mathbf{b}$ ergibt sich $\lambda(\hat{\mathbf{x}}, \mathbf{b}) = 0$. Folglich erhalten wir $\hat{\mathbf{x}} = \phi(\hat{\mathbf{x}}, \mathbf{b})$ und $\hat{\mathbf{x}}$ stellt einen Fixpunkt der Iteration ϕ dar. $\square$

Satz 4.53 *Sei $\mathbf{A}$ positiv definit und symmetrisch, dann konvergiert die durch das Verfahren des steilsten Abstiegs definierte Folge $\{\mathbf{x}_m\}_{m \in \mathbb{N}_0}$ für jeden Startvektor $\mathbf{x}_0 \in \mathbb{R}^n$ gegen die Lösung $\hat{\mathbf{x}} = \mathbf{A}^{-1}\mathbf{b}$, und es gilt für den Fehlervektor $\mathbf{e}_m = \mathbf{x}_m - \hat{\mathbf{x}}$ die Abschätzung*

$$\|\mathbf{e}_m\|_{\mathbf{A}} \leq \left(\frac{\text{cond}_2(\mathbf{A}) - 1}{\text{cond}_2(\mathbf{A}) + 1} \right)^m \|\mathbf{e}_0\|_{\mathbf{A}}. \quad (4.3.15)$$

Beweis:

Betrachten wir die Gleichung (4.3.14), so kann das Verfahren als spezielle Richardson-Iteration mit variablem $\Theta = \lambda(\mathbf{x}, \mathbf{b})$ interpretiert werden. Da der Startvektor $\mathbf{x}_0 \in \mathbb{R}^n$ beliebig ist, reicht der Nachweis der Abschätzung (4.3.15) für $m = 1$.

Seien $\lambda_{\max} = \max_{\lambda \in \sigma(\mathbf{A})} \lambda$ und $\lambda_{\min} = \min_{\lambda \in \sigma(\mathbf{A})} \lambda$, dann ergibt sich beim herkömmlichen Richardson-Verfahren mit optimalem Gewichtungparameter

$$\Theta_{\text{opt}} \stackrel{\text{Satz 4.32}}{=} \frac{2}{\lambda_{\max} + \lambda_{\min}}$$

für den Fehlervektor $\mathbf{e}_1^R$ die Darstellung

$$\mathbf{e}_1^R = \mathbf{M}_R(\Theta_{\text{opt}})\mathbf{e}_0$$

mit

$$\mathbf{M}_R(\Theta_{\text{opt}}) = \mathbf{I} - \Theta_{\text{opt}}\mathbf{A}.$$

Aus der Symmetrie der Matrix $\mathbf{A}$ folgt, daß $\mathbf{M}_R(\Theta_{\text{opt}})$ ebenfalls symmetrisch ist und sich somit

$$\begin{aligned} \|\mathbf{M}_R(\Theta_{\text{opt}})\|_2 &\stackrel{\text{Satz 2.34}}{=} \rho(\mathbf{M}_R(\Theta_{\text{opt}})) \\ &\stackrel{\text{Satz 4.32}}{=} \frac{\lambda_{\max} - \lambda_{\min}}{\lambda_{\max} + \lambda_{\min}} \end{aligned}$$

ergibt. Zudem gilt für beliebiges $p \in \mathbb{R}$

$$\mathbf{M}_R(\Theta_{\text{opt}})\mathbf{A}^p = \mathbf{A}^p\mathbf{M}_R(\Theta_{\text{opt}}),$$

und es folgt mit

$$\tilde{\mathbf{e}}_0 = \mathbf{A}^{1/2}\mathbf{e}_0 \text{ und } \tilde{\mathbf{e}}_1^R = \mathbf{A}^{1/2}\mathbf{e}_1^R$$

die Darstellung

$$\tilde{\mathbf{e}}_1^R = \mathbf{A}^{1/2}\mathbf{M}_R(\Theta_{\text{opt}})\mathbf{e}_0 = \mathbf{M}_R(\Theta_{\text{opt}})\tilde{\mathbf{e}}_0.$$

Hiermit erhalten wir

$$\|\mathbf{e}_1^R\|_{\mathbf{A}} = \|\tilde{\mathbf{e}}_1^R\|_2 \leq \|\mathbf{M}_R(\Theta_{\text{opt}})\|_2 \|\tilde{\mathbf{e}}_0\|_2 = \|\mathbf{M}_R(\Theta_{\text{opt}})\|_2 \|\mathbf{e}_0\|_{\mathbf{A}}.$$

Unter Verwendung der Gleichung (4.3.10) folgern wir

$$G(\mathbf{x}) = \frac{1}{2}(\mathbf{A}\mathbf{x} - \mathbf{b})^T \mathbf{A}^{-1}(\mathbf{A}\mathbf{x} - \mathbf{b}) = \frac{1}{2}\|\mathbf{x} - \hat{\mathbf{x}}\|_{\mathbf{A}}^2$$

so daß sich mit

$$\mathbf{x}_1^R = \mathbf{x}_0 + \Theta_{\text{opt}}\mathbf{r}_0$$

und

$$\mathbf{x}_1 = \mathbf{x}_0 + \lambda_0\mathbf{r}_0$$

durch

$$\mathbf{x}_1 = \arg \min_{\mathbf{x} \in \mathbf{x}_0 + \text{span}\{\mathbf{r}_0\}} F(\mathbf{x}) = \arg \min_{\mathbf{x} \in \mathbf{x}_0 + \text{span}\{\mathbf{r}_0\}} G(\mathbf{x}) = \arg \min_{\mathbf{x} \in \mathbf{x}_0 + \text{span}\{\mathbf{r}_0\}} \frac{1}{2}\|\mathbf{x} - \hat{\mathbf{x}}\|_{\mathbf{A}}^2$$

die Abschätzung

$$\|\mathbf{e}_1\|_{\mathbf{A}} \leq \|\mathbf{e}_1^R\|_{\mathbf{A}} \leq \underbrace{\frac{\lambda_{\max} - \lambda_{\min}}{\lambda_{\max} + \lambda_{\min}}}_{\xi :=} \|\mathbf{e}_0\|_{\mathbf{A}} \quad (4.3.16)$$

ergibt. Für positiv definite, symmetrische Matrizen $\mathbf{A} \in \mathbb{R}^{n \times n}$ gilt $\|\mathbf{A}\|_2 = \rho(\mathbf{A}) = \lambda_{\max} > 0$ und $\|\mathbf{A}^{-1}\|_2 = \rho(\mathbf{A}^{-1}) = \lambda_{\min}^{-1} > 0$, so daß einerseits mit $|\xi| < 1$ die behauptete Konvergenz des Verfahrens vorliegt und andererseits aus (4.3.16) die Abschätzung

$$\|\mathbf{e}_1\|_{\mathbf{A}} \leq \left(\frac{\frac{\lambda_{\max}}{\lambda_{\min}} - 1}{\frac{\lambda_{\max}}{\lambda_{\min}} + 1} \right) \|\mathbf{e}_0\|_{\mathbf{A}} = \left(\frac{\|\mathbf{A}\|_2 \|\mathbf{A}^{-1}\|_2 - 1}{\|\mathbf{A}\|_2 \|\mathbf{A}^{-1}\|_2 + 1} \right) \|\mathbf{e}_0\|_{\mathbf{A}} = \left(\frac{\text{cond}_2(\mathbf{A}) - 1}{\text{cond}_2(\mathbf{A}) + 1} \right) \|\mathbf{e}_0\|_{\mathbf{A}}$$

folgt. □

Wegen $\text{cond}_2(\mathbf{A}) \geq 1$ ist es aufgrund des obigen Satzes somit stets vorteilhaft, eine möglichst kleine Konditionzahl vorliegen zu haben.

Beispiel 4.54 Wir betrachten

$$\mathbf{A}\mathbf{x} = \mathbf{b}$$

mit

$$\mathbf{A} = \begin{pmatrix} 2 & 0 \\ 0 & 10 \end{pmatrix}, \quad \mathbf{b} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}, \quad \mathbf{x}_0 = \begin{pmatrix} 4 \\ \sqrt{1.8} \end{pmatrix}.$$

Hiermit erhalten wir den folgenden Konvergenzverlauf:

Verfahren des steilsten Abstiegs (Gradientenverfahren)				
m	$x_{m,1}$	$x_{m,2}$	$\varepsilon_m := \ \mathbf{x}_m - \mathbf{A}^{-1}\mathbf{b}\ _{\mathbf{A}}$	$\varepsilon_m/\varepsilon_{m-1}$
0	4.000000e+00	1.341641e+00	7.071068e+00	
10	3.271049e-02	1.097143e-02	5.782453e-02	6.183904e-01
40	1.788827e-08	5.999910e-09	3.162230e-08	6.183904e-01
70	9.782499e-15	3.281150e-15	1.729318e-14	6.183904e-01
72	3.740893e-15	1.254734e-15	6.613026e-15	6.183904e-01

Wir nutzen die in Abbildung 4.19 qualitativ dargestellten Höhenlinien der Funktion F zur Verdeutlichung des Konvergenzverlaufs. Liegen bei der betrachteten Diagonalmatrix gleiche Diagonaleinträge vor, dann beschreiben die Höhenlinien Kreise, und das Verfahren konvergiert bei beliebigem Startvektor bereits bei der ersten Iteration, da der Residuenvektor stets in die Richtung des Koordinatenursprungs zeigt. Weist die Diagonalmatrix positive, jedoch sehr unterschiedlich große Diagonaleinträge auf, dann stellen die Höhenlinien der Funktion F zunehmend gestreckte Ellipsen dar, wodurch die Näherungslösung bei jeder Iteration das Vorzeichen wechseln kann und nur sehr langsam gegen das Minimum der Funktion F konvergiert. Die zunehmende Verringerung der Konvergenzgeschwindigkeit lässt sich hierbei auch deutlich durch Betrachten der Konditionzahl der Matrix $\mathbf{A}$ erklären, die bei der zugrundeliegenden Diagonalmatrix mit dem Streckungsverhältnis der Ellipse übereinstimmt.

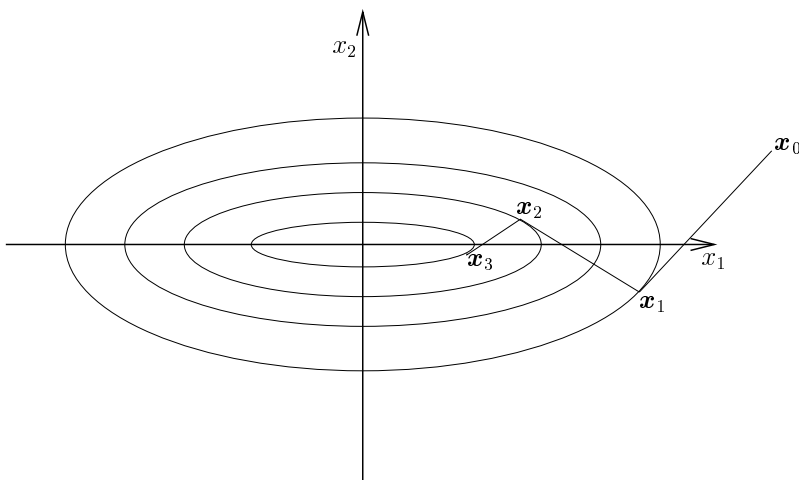


Bild 4.19 Höhenlinien der Funktion F mit qualitativem Konvergenzverlauf

Als Motivation zur Verbesserung führen wir die Begriffe Optimalität bezüglich einer Richtung und eines Unterraums ein.

Definition 4.55 Sei $F : \mathbb{R}^n \rightarrow \mathbb{R}$ gegeben, dann heißt $\mathbf{x} \in \mathbb{R}^n$

- (a) optimal bezüglich der Richtung $\mathbf{p} \in \mathbb{R}^n$, falls

$$F(\mathbf{x}) \leq F(\mathbf{x} + \lambda \mathbf{p}) \quad \forall \lambda \in \mathbb{R}$$

gilt.

- (b) optimal bezüglich eines Unterraums $U \subset \mathbb{R}^n$, falls

$$F(\mathbf{x}) \leq F(\mathbf{x} + \boldsymbol{\xi}) \quad \forall \boldsymbol{\xi} \in U$$

gilt.

Lemma 4.56 Sei F durch (4.3.8) gegeben, dann ist $\mathbf{x} \in \mathbb{R}^n$ genau dann bezüglich $U \subset \mathbb{R}^n$ optimal, wenn

$$\mathbf{r} = \mathbf{b} - \mathbf{A}\mathbf{x} \perp U$$

gilt.

Beweis:

Für beliebiges $\boldsymbol{\xi} \in U \setminus \{\mathbf{0}\}$ betrachten wir für $\lambda \in \mathbb{R}$

$$f_{\mathbf{x}, \boldsymbol{\xi}}(\lambda) = F(\mathbf{x} + \lambda \boldsymbol{\xi}).$$

Die Funktion $f_{\mathbf{x}, \boldsymbol{\xi}}$ ist laut (4.3.12) strikt konvex und aus

$$f'_{\mathbf{x}, \boldsymbol{\xi}}(\lambda) = (\mathbf{A}\mathbf{x} - \mathbf{b}, \boldsymbol{\xi})_2 + \lambda(\mathbf{A}\boldsymbol{\xi}, \boldsymbol{\xi})_2$$

folgt

$$f'_{\mathbf{x}, \boldsymbol{\xi}}(0) = 0 \quad \Leftrightarrow \quad \mathbf{A}\mathbf{x} - \mathbf{b} \perp \boldsymbol{\xi}.$$

□

Satz 4.57 Die Iterierten $\mathbf{x}_m$, $m \in \mathbb{N}$ des Gradientenverfahrens sind optimal bezüglich der Richtung $\mathbf{r}_{m-1} = \mathbf{b} - \mathbf{A}\mathbf{x}_{m-1}$.

Beweis:

Im Fall $\mathbf{r}_{m-1} = \mathbf{0}$ folgt direkt $\mathbf{r}_m \perp \mathbf{r}_{m-1}$. Sei $\mathbf{r}_{m-1} \neq \mathbf{0}$, so erhalten wir mit

$$\lambda_{m-1} = \frac{\|\mathbf{r}_{m-1}\|_2^2}{(\mathbf{A}\mathbf{r}_{m-1}, \mathbf{r}_{m-1})_2}$$

die Orthogonalität der Residuenvektoren durch

$$\begin{aligned} (\mathbf{r}_m, \mathbf{r}_{m-1})_2 &= (\mathbf{r}_{m-1} - \lambda_{m-1} \mathbf{A}\mathbf{r}_{m-1}, \mathbf{r}_{m-1})_2 \\ &= (\mathbf{r}_{m-1}, \mathbf{r}_{m-1})_2 - \frac{\|\mathbf{r}_{m-1}\|_2^2}{(\mathbf{A}\mathbf{r}_{m-1}, \mathbf{r}_{m-1})_2} (\mathbf{A}\mathbf{r}_{m-1}, \mathbf{r}_{m-1})_2 \\ &= 0. \end{aligned}$$

Lemma 4.56 liefert damit die Optimalität.

□

Korollar 4.58 *Das Gradientenverfahren stellt in jedem Schritt eine orthogonale Projektionsmethode mit $K = L = \text{span}\{\mathbf{r}_{m-1}\}$ dar.*

Wünschenswert wäre eine Optimalität der Iterierten bezüglich des gesamten Unterraums $U = \text{span}\{\mathbf{r}_0, \dots, \mathbf{r}_{m-1}\}$, da für linear unabhängige Residuenvektoren hierdurch spätestens

$$\mathbf{x}_n = \mathbf{A}^{-1}\mathbf{b}$$

folgt. Im Verfahren des steilsten Abstiegs liegt jedoch das Problem vor, daß die ermittelte Näherungslösung $\mathbf{x}_m$ stets nur eine Optimalität bezüglich $\mathbf{r}_{m-1}$ aufweist und die Bedingung $\mathbf{r} \perp \mathbf{p}$ nicht transitiv ist, so daß aus $\mathbf{r}_{m-2} \perp \mathbf{r}_{m-1}$ und $\mathbf{r}_{m-1} \perp \mathbf{r}_m$ nicht notwendigerweise $\mathbf{r}_{m-2} \perp \mathbf{r}_m$ folgt.

4.3.1.2 Das Verfahren der konjugierten Richtungen

Die Idee dieses Verfahrens ist die Erweiterung der Optimalität der ermittelten Näherungen $\mathbf{x}_m$ auf den gesamten Unterraum $U_m = \text{span}\{\mathbf{p}_0, \dots, \mathbf{p}_{m-1}\}$ mit linear unabhängigen Suchrichtungen $\mathbf{p}_0, \dots, \mathbf{p}_{m-1}$. Hierzu werden wir mit dem folgenden Satz eine Bedingung an die zu verwendenden Suchrichtungen formulieren, die den Erhalt der Optimalität bezüglich U_m im $(m+1)$ -ten Schritt garantiert.

Satz 4.59 *Sei F gemäß (4.3.8) gegeben und $\mathbf{x} \in \mathbb{R}^n$ optimal bezüglich des Unterraums $U = \text{span}\{\mathbf{p}_0, \dots, \mathbf{p}_{m-1}\} \subset \mathbb{R}^n$, dann ist $\tilde{\mathbf{x}} = \mathbf{x} + \boldsymbol{\xi}$ genau dann optimal bezüglich U , wenn*

$$\mathbf{A}\boldsymbol{\xi} \perp U$$

gilt.

Beweis:

Sei $\boldsymbol{\eta} \in U$ beliebig, dann folgt die Behauptung unmittelbar aus

$$(\mathbf{b} - \mathbf{A}\tilde{\mathbf{x}}, \boldsymbol{\eta})_2 = (\mathbf{b} - \mathbf{A}\mathbf{x}, \boldsymbol{\eta})_2 - \underbrace{(\mathbf{A}\boldsymbol{\xi}, \boldsymbol{\eta})_2}_{=0}.$$

□

Wird mit $\mathbf{p}_m$ eine Suchrichtung gewählt, für die

$$\mathbf{A}\mathbf{p}_m \perp U_m = \text{span}\{\mathbf{p}_0, \dots, \mathbf{p}_{m-1}\}$$

oder äquivalent

$$\mathbf{A}\mathbf{p}_m \perp \mathbf{p}_j, j = 0, \dots, m-1$$

gilt, so erbt die Näherungslösung

$$\mathbf{x}_{m+1} = \mathbf{x}_m + \lambda_m \mathbf{p}_m$$

laut Satz 4.59 die Optimalität von $\mathbf{x}_m$ bezüglich U_m unabhängig von der Wahl des skalaren Gewichtungsparameters λ_m . Dieser Freiheitsgrad wird im weiteren zur Erweiterung der Optimalität auf

$$U_{m+1} = \text{span} \{\mathbf{p}_0, \dots, \mathbf{p}_m\}$$

genutzt.

Definition 4.60 Sei $\mathbf{A} \in \mathbb{R}^{n \times n}$, dann heißen die Vektoren $\mathbf{p}_0, \dots, \mathbf{p}_m \in \mathbb{R}^n$ paarweise konjugiert oder $\mathbf{A}$ -orthogonal, falls

$$(\mathbf{p}_i, \mathbf{p}_j)_{\mathbf{A}} := (\mathbf{A}\mathbf{p}_i, \mathbf{p}_j)_2 = 0 \quad \forall i, j \in \{0, \dots, m\} \text{ und } i \neq j$$

gilt.

Eine für das Verfahren wichtige Eigenschaft liegt in der sukzessiven Dimensionserhöhung innerhalb der betrachteten Folge von Untervektorräumen $\{U_m\}_{m=1,2,\dots}$, die durch das anschließende Lemma nachgewiesen wird.

Lemma 4.61 Seien $\mathbf{A} \in \mathbb{R}^{n \times n}$ symmetrisch, positiv definit und $\mathbf{p}_0, \dots, \mathbf{p}_{m-1} \in \mathbb{R}^n \setminus \{\mathbf{0}\}$ paarweise $\mathbf{A}$ -orthogonal, dann gilt

$$\dim \text{span} \{\mathbf{p}_0, \dots, \mathbf{p}_{m-1}\} = m$$

für $m = 1, \dots, n$.

Beweis:

Gelte $\sum_{j=0}^{m-1} \alpha_j \mathbf{p}_j = \mathbf{0}$ mit $\alpha_j \in \mathbb{R}$, dann erhalten wir für $i = 0, \dots, m-1$

$$0 = (\mathbf{0}, \mathbf{A}\mathbf{p}_i)_2 = \left(\sum_{j=0}^{m-1} \alpha_j \mathbf{p}_j, \mathbf{A}\mathbf{p}_i \right)_2 = \sum_{j=0}^{m-1} \alpha_j (\mathbf{p}_j, \mathbf{A}\mathbf{p}_i)_2 = \alpha_i \underbrace{(\mathbf{p}_i, \mathbf{A}\mathbf{p}_i)_2}_{\neq 0, \text{ da } \mathbf{A} \text{ pos.def.}} \quad .$$

Folglich ergibt sich $\alpha_i = 0$ für $i = 0, \dots, m-1$, wodurch die lineare Unabhängigkeit der Vektoren nachgewiesen ist. $\square$

Seien mit $\mathbf{p}_0, \dots, \mathbf{p}_m \in \mathbb{R}^n \setminus \{\mathbf{0}\}$ paarweise konjugierte Suchrichtungen gegeben und $\mathbf{x}_m$ optimal bezüglich $U_m = \text{span} \{\mathbf{p}_0, \dots, \mathbf{p}_{m-1}\}$, dann erhalten wir die Optimalität von

$$\mathbf{x}_{m+1} = \mathbf{x}_m + \lambda \mathbf{p}_m$$

bezüglich U_{m+1} , wenn

$$0 = (\mathbf{b} - \mathbf{A}\mathbf{x}_{m+1}, \mathbf{p}_j)_2 = \underbrace{(\mathbf{b} - \mathbf{A}\mathbf{x}_m, \mathbf{p}_j)_2}_{= 0 \text{ für } j \neq m} - \lambda \underbrace{(\mathbf{A}\mathbf{p}_m, \mathbf{p}_j)_2}_{= 0 \text{ für } j \neq m}$$

für $j = 0, \dots, m$ gilt. Hieraus ergibt sich für λ die Darstellung

$$\lambda = \frac{(\mathbf{r}_m, \mathbf{p}_m)_2}{(\mathbf{A}\mathbf{p}_m, \mathbf{p}_m)_2},$$

und wir können das Verfahren der konjugierten Richtungen in der folgenden Form schreiben:

Algorithmus - Verfahren der konjugierten Richtungen —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$
$\mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0$
Für $m = 0, \dots, n-1$
$\lambda_m := \frac{(\mathbf{r}_m, \mathbf{p}_m)_2}{(\mathbf{A}\mathbf{p}_m, \mathbf{p}_m)_2}$
$\mathbf{x}_{m+1} := \mathbf{x}_m + \lambda_m \mathbf{p}_m$
$\mathbf{r}_{m+1} := \mathbf{r}_m - \lambda_m \mathbf{A}\mathbf{p}_m$

Sind die gegebenen Suchrichtungen ungünstig gewählt, so kann mit $\mathbf{x}_n$ die exakte Lösung vorliegen, obwohl $\mathbf{x}_{n-1}$ noch einen sehr großen Fehler aufweist. Im Extremfall kann hierbei $\mathbf{x}_m = \mathbf{x}_0 \neq \mathbf{A}^{-1}\mathbf{b}$ für $m = 1, 2, \dots, n-1$ und $\mathbf{x}_n = \mathbf{A}^{-1}\mathbf{b}$ gelten. Das Verfahren kann daher bei fest vorgegebenen Suchrichtungen in der Regel nur als direktes Verfahren genutzt werden. Diese Eigenschaft führt bei großen n zu einem hohen Rechenaufwand. Eine problemangepaßte Auswahl der Suchrichtungen ist also gefordert.

4.3.1.3 Das Verfahren der konjugierten Gradienten

Die von Hestenes und Stiefel [37] vorgestellte Methode der konjugierten Gradienten (CG-Verfahren) kombiniert das Gradientenverfahren mit dem Verfahren der konjugierten Richtungen. Das Verfahren nutzt die Residuenvektoren zur Definition der konjugierten Suchrichtungen, wodurch ein problemangepaßtes Vorgehen in Bezug auf die Auswahl der Suchrichtungen vorliegt, und zudem die Optimalität des Verfahrens der konjugierten Richtungen erhalten wird (siehe Schaubild 4.20). Analog zu der bereits vorgestellten Methode der konjugierten Richtungen kann das Verfahren der konjugierten Gradienten als direktes und iteratives Verfahren interpretiert werden.

Mit den Residuenvektoren $\mathbf{r}_0, \dots, \mathbf{r}_m$ ermitteln wir die Suchrichtungen sukzessive für $m = 1, \dots, n-1$ gemäß

$$\begin{aligned} \mathbf{p}_0 &= \mathbf{r}_0, \\ \mathbf{p}_m &= \mathbf{r}_m + \sum_{j=0}^{m-1} \alpha_j \mathbf{p}_j. \end{aligned} \quad (4.3.17)$$

Für $\alpha_j = 0$ ($j = 0, \dots, m-1$) erhalten wir damit eine zum Verfahren des steilsten Abstiegs analoge Auswahl der Suchrichtungen. Mit der Berücksichtigung der bereits genutzten Suchrichtungen $\mathbf{p}_0, \dots, \mathbf{p}_{m-1} \in \mathbb{R}^n \setminus \{\mathbf{0}\}$ in der obigen Form liegen m Freiheitsgrade in der Wahl der Koeffizienten α_j vor, die zur Gewährleistung der Konjugiertheit der Suchrichtungen genutzt werden. Aus der geforderten $\mathbf{A}$ -Orthogonalitätsbedingung folgt

$$0 = (\mathbf{A}\mathbf{p}_m, \mathbf{p}_i)_2 = (\mathbf{A}\mathbf{r}_m, \mathbf{p}_i)_2 + \sum_{j=0}^{m-1} \alpha_j (\mathbf{A}\mathbf{p}_j, \mathbf{p}_i)_2$$

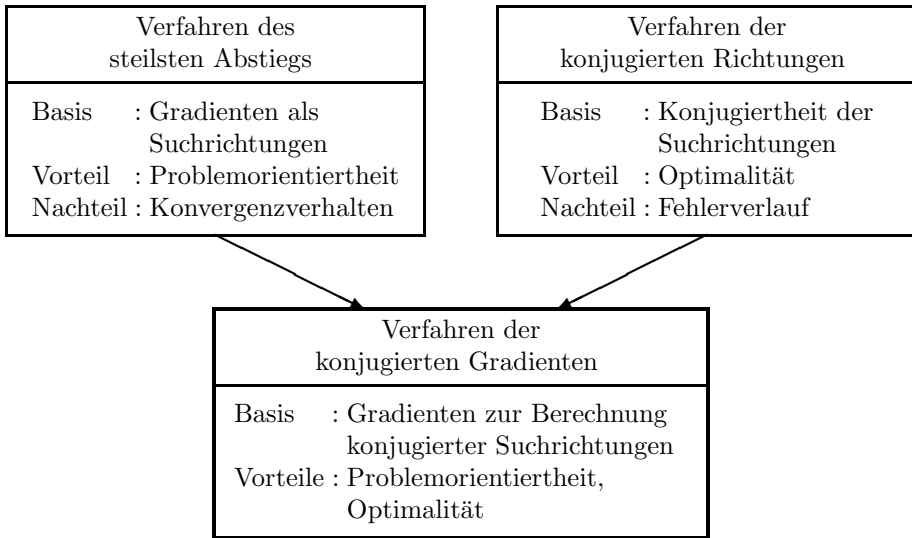


Bild 4.20 Herleitung des CG-Verfahrens

für $i = 0, \dots, m-1$. Mit

$$(\mathbf{A}\mathbf{p}_j, \mathbf{p}_i)_2 = 0 \quad \text{für } i, j \in \{0, \dots, m-1\} \quad \text{und } i \neq j$$

erhalten wir die benötigte Vorschrift zur Berechnung der Koeffizienten in der Form

$$\alpha_i = -\frac{(\mathbf{A}\mathbf{r}_m, \mathbf{p}_i)_2}{(\mathbf{A}\mathbf{p}_i, \mathbf{p}_i)_2}. \quad (4.3.18)$$

Somit ergibt sich die vorläufige Version des Verfahrens der konjugierten Gradienten in der folgenden Darstellung:

Algorithmus - Vorläufiges Verfahren der konjugierten Gradienten —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$
$\mathbf{p}_0 := \mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0$
Für $m = 0, \dots, n-1$
$\lambda_m := \frac{(\mathbf{r}_m, \mathbf{p}_m)_2}{(\mathbf{A}\mathbf{p}_m, \mathbf{p}_m)_2}$
$\mathbf{x}_{m+1} := \mathbf{x}_m + \lambda_m \mathbf{p}_m$
$\mathbf{r}_{m+1} := \mathbf{r}_m - \lambda_m \mathbf{A}\mathbf{p}_m$
$\mathbf{p}_{m+1} := \mathbf{r}_{m+1} - \sum_{j=0}^m \frac{(\mathbf{A}\mathbf{r}_{m+1}, \mathbf{p}_j)_2}{(\mathbf{A}\mathbf{p}_j, \mathbf{p}_j)_2} \mathbf{p}_j \quad (4.3.19)$

Das vorläufige CG-Verfahren weist den entscheidenden Nachteil auf, daß zur Berechnung von $\mathbf{p}_{m+1}$ scheinbar alle $\mathbf{p}_j$ ($j = 0, \dots, m$) benötigt werden. Im ungünstigsten Fall benötigen wir daher den Speicherplatz einer vollbesetzten $n \times n$ -Matrix für die Suchrichtungen. Bei großen schwachbesetzten Matrizen ist das Verfahren somit ineffizient und eventuell unpraktikabel. Durch die folgende Analyse zeigt sich jedoch, daß das Verfahren im Hinblick auf den Speicherplatzbedarf und die Rechenzeit entscheidend verbessert werden kann.

Satz 4.62 *Vorausgesetzt, das vorläufige CG-Verfahren bricht nicht vor der Berechnung von $\mathbf{p}_k$ für $k > 0$ ab, dann gilt*

- (a) $\mathbf{p}_m$ ist konjugiert zu allen $\mathbf{p}_j$ mit $0 \leq j < m \leq k$,
- (b) $U_{m+1} := \text{span} \{\mathbf{p}_0, \dots, \mathbf{p}_m\} = \text{span} \{\mathbf{r}_0, \dots, \mathbf{r}_m\}$ mit $\dim U_{m+1} = m + 1$ für $m = 0, \dots, k - 1$,
- (c) $\mathbf{r}_m \perp U_m$ für $m = 1, \dots, k$,
- (d) $\mathbf{x}_k = \mathbf{A}^{-1}\mathbf{b} \iff \mathbf{r}_k = \mathbf{0} \iff \mathbf{p}_k = \mathbf{0}$,
- (e) $U_{m+1} = \text{span} \{\mathbf{r}_0, \dots, \mathbf{A}^m \mathbf{r}_0\}$ für $m = 0, \dots, k - 1$,
- (f) $\mathbf{r}_m$ ist konjugiert zu allen $\mathbf{p}_j$ mit $0 \leq j < m - 1 < k - 1$.

Beweis:

zu (a): Da $\mathbf{p}_k$ berechnet wurde, gilt $\mathbf{p}_0, \dots, \mathbf{p}_{k-1} \in \mathbb{R}^n \setminus \{\mathbf{0}\}$. Die Behauptung folgt aus den Berechnungsvorschriften (4.3.17) und (4.3.18).

zu (b): Induktion über m .

Für $m = 0$ folgt $\mathbf{p}_0 = \mathbf{r}_0 \in \mathbb{R}^n \setminus \{\mathbf{0}\}$ und damit die Behauptung. Sei (b) für $m < k - 1$ erfüllt. Wegen $\mathbf{p}_{m+1} \in \mathbb{R}^n \setminus \{\mathbf{0}\}$ folgt mit (a) die Konjugiertheit von $\mathbf{p}_{m+1}$ zu allen $\mathbf{p}_0, \dots, \mathbf{p}_m$. Aus Lemma 4.61 erhalten wir somit $\dim U_{m+2} = m + 2$, und

$$\mathbf{p}_{m+1} - \mathbf{r}_{m+1} \stackrel{(4.3.17)}{=} \sum_{j=0}^m \alpha_j \mathbf{p}_j \in U_{m+1}$$

liefert

$$U_{m+2} = \text{span} \{U_{m+1}, \mathbf{p}_{m+1}\} = \text{span} \{U_{m+1}, \mathbf{r}_{m+1}\}.$$

zu (c): Induktion über m :

Für $m = 1$ erhalten wir mit $\mathbf{p}_0 \neq \mathbf{0}$ die Gleichung

$$(\mathbf{r}_1, \mathbf{r}_0)_2 = (\mathbf{r}_0, \mathbf{r}_0)_2 - \frac{(\mathbf{r}_0, \mathbf{p}_0)_2}{(\mathbf{A}\mathbf{p}_0, \mathbf{p}_0)_2} (\mathbf{A}\mathbf{p}_0, \mathbf{r}_0)_2 \stackrel{\mathbf{r}_0 = \mathbf{p}_0}{=} 0.$$

Sei (c) für $m < k$ erfüllt und $\boldsymbol{\eta} \in U_m$, dann folgt mit Teil (a)

$$(\mathbf{r}_{m+1}, \boldsymbol{\eta})_2 = \underbrace{(\mathbf{r}_m, \boldsymbol{\eta})_2}_{=0} - \lambda_m \underbrace{(\mathbf{A}\mathbf{p}_m, \boldsymbol{\eta})_2}_{=0} = 0.$$

Unter Berücksichtigung von $\mathbf{p}_m \neq \mathbf{0}$ erhalten wir die Behauptung durch

$$(\mathbf{r}_{m+1}, \mathbf{p}_m)_2 = (\mathbf{r}_m, \mathbf{p}_m)_2 - \frac{(\mathbf{r}_m, \mathbf{p}_m)_2}{(\mathbf{A}\mathbf{p}_m, \mathbf{p}_m)_2} (\mathbf{A}\mathbf{p}_m, \mathbf{p}_m)_2 = 0.$$

zu (d): Aus $\mathbf{r}_k = \mathbf{b} - \mathbf{A}\mathbf{x}_k$ folgt direkt die Äquivalenz zwischen $\mathbf{x}_k = \mathbf{A}^{-1}\mathbf{b}$ und $\mathbf{r}_k = \mathbf{0}$. Sei $\mathbf{r}_k = \mathbf{0}$, dann liefert (4.3.19) die Gleichung $\mathbf{p}_k = \mathbf{0}$. Gelte $\mathbf{p}_k = \mathbf{0}$, dann gilt wiederum mit (4.3.19) $\mathbf{r}_k \in U_k$. Laut Teil (c) gilt aber $\mathbf{r}_k \perp U_k$, womit $\mathbf{r}_k = \mathbf{0}$ folgt.

zu (e): Induktion über m :

Für $m = 0$ ist die Aussage trivial.

Sei (e) für $m < k - 1$ erfüllt, dann folgt mit (b)

$$\mathbf{r}_m \in U_{m+1} = \text{span} \{\mathbf{r}_0, \dots, \mathbf{r}_m\} = \text{span} \{\mathbf{r}_0, \dots, \mathbf{A}^m \mathbf{r}_0\}$$

sowie

$$\mathbf{A}\mathbf{p}_m \in \mathbf{A}U_{m+1} = \text{span} \{\mathbf{A}\mathbf{r}_0, \dots, \mathbf{A}^{m+1}\mathbf{r}_0\}.$$

Folglich erhalten wir $\mathbf{r}_{m+1} = \mathbf{r}_m - \lambda_m \mathbf{A}\mathbf{p}_m \in \text{span} \{\mathbf{r}_0, \dots, \mathbf{A}^{m+1}\mathbf{r}_0\}$, so daß $U_{m+2} = \text{span} \{\mathbf{r}_0, \dots, \mathbf{r}_{m+1}\} \subset \text{span} \{\mathbf{r}_0, \dots, \mathbf{A}^{m+1}\mathbf{r}_0\}$ gilt. Teil (b) liefert $\dim U_{m+2} = m + 2$, wodurch $U_{m+2} = \text{span} \{\mathbf{r}_0, \dots, \mathbf{A}^{m+1}\mathbf{r}_0\}$ folgt.

zu (f): Für $0 \leq j < m - 1 \leq k - 1$ gilt $\mathbf{p}_j \in U_{m-1}$. Somit gilt $\mathbf{A}\mathbf{p}_j \in U_m$, und wir erhalten

$$(\mathbf{A}\mathbf{r}_m, \mathbf{p}_j)_2 \stackrel{\mathbf{A} \text{ symm.}}{=} (\mathbf{r}_m, \mathbf{A}\mathbf{p}_j)_2 \stackrel{(c)}{=} 0.$$

□

Dem obigen Satz können wir drei wesentliche Aussagen zur Verbesserung des Verfahrens entnehmen:

- Mit Aussage (f) folgt

$$\mathbf{p}_m = \mathbf{r}_m - \sum_{j=0}^{m-1} \frac{(\mathbf{A}\mathbf{r}_m, \mathbf{p}_j)_2}{(\mathbf{A}\mathbf{p}_j, \mathbf{p}_j)_2} \mathbf{p}_j = \mathbf{r}_m - \frac{(\mathbf{A}\mathbf{r}_m, \mathbf{p}_{m-1})_2}{(\mathbf{A}\mathbf{p}_{m-1}, \mathbf{p}_{m-1})_2} \mathbf{p}_{m-1}. \quad (4.3.20)$$

Damit ist der Speicheraufwand unabhängig von der Anzahl der Iterationen.

- Das Verfahren bricht in der $k + 1$ -ten Iteration genau dann ab, wenn $\mathbf{p}_k = \mathbf{0}$ gilt. Mit (d) liefert $\mathbf{x}_k$ in diesem Fall bereits die exakte Lösung, so daß $\mathbf{p}_k = \mathbf{0}$ als Abbruchkriterium genutzt werden kann. Im folgenden Diagramm stellen wir zudem eine Variante dar, die ohne zusätzlichen Rechenaufwand als Abbruchkriterium das Residuum verwendet.
- Aus $\mathbf{r}_{m+1} = \mathbf{r}_m - \lambda_m \mathbf{A}\mathbf{p}_m$ erhalten wir aufgrund der Eigenschaft (c) die Gleichung $(\mathbf{r}_m - \lambda_m \mathbf{A}\mathbf{p}_m, \mathbf{r}_m)_2 = 0$, wodurch sich die skalare Größe in der Form

$$\lambda_m = \frac{(\mathbf{r}_m, \mathbf{p}_m)_2}{(\mathbf{A}\mathbf{p}_m, \mathbf{p}_m)_2} = \frac{(\mathbf{r}_m, \mathbf{r}_m)_2}{(\mathbf{A}\mathbf{p}_m, \mathbf{r}_m)_2} \quad (4.3.21)$$

schreiben läßt. Verwendung der Gleichung (4.3.20) offenbart

$$(\mathbf{A}\mathbf{p}_m, \mathbf{r}_m)_2 = \left(\mathbf{A}\mathbf{p}_m, \mathbf{p}_m + \frac{(\mathbf{A}\mathbf{r}_m, \mathbf{p}_{m-1})_2}{(\mathbf{A}\mathbf{p}_{m-1}, \mathbf{p}_{m-1})_2} \mathbf{p}_{m-1} \right)_2 = (\mathbf{A}\mathbf{p}_m, \mathbf{p}_m)_2,$$

so daß (4.3.21) die Eigenschaft

$$(\mathbf{r}_m, \mathbf{r}_m)_2 = (\mathbf{r}_m, \mathbf{p}_m)_2 \quad (4.3.22)$$

liefert. Desweiteren ergibt sich im Fall $\lambda_m \neq 0$ aus dem vorläufigen CG-Verfahren stets

$$\mathbf{A}\mathbf{p}_m = -\frac{1}{\lambda_m}(\mathbf{r}_{m+1} - \mathbf{r}_m), \quad (4.3.23)$$

so daß

$$\begin{aligned} \frac{(\mathbf{A}\mathbf{r}_{m+1}, \mathbf{p}_m)_2}{(\mathbf{A}\mathbf{p}_m, \mathbf{p}_m)_2} &\stackrel{\text{A symm.}}{=} \frac{(\mathbf{A}\mathbf{p}_m, \mathbf{r}_{m+1})_2}{(\mathbf{A}\mathbf{p}_m, \mathbf{p}_m)_2} \stackrel{(4.3.23)}{=} \frac{(\mathbf{r}_{m+1} - \mathbf{r}_m, \mathbf{r}_{m+1})_2}{(\mathbf{r}_{m+1} - \mathbf{r}_m, \mathbf{p}_m)_2} \\ &\stackrel{(b),(c)}{=} -\frac{(\mathbf{r}_{m+1}, \mathbf{r}_{m+1})_2}{(\mathbf{r}_m, \mathbf{p}_m)_2} \stackrel{(4.3.22)}{=} -\frac{(\mathbf{r}_{m+1}, \mathbf{r}_{m+1})_2}{(\mathbf{r}_m, \mathbf{r}_m)_2} \end{aligned}$$

folgt und hierdurch innerhalb jeder Iteration eine Matrix-Vektor-Multiplikation entfällt.

Hiermit ergibt sich das CG-Verfahren in der folgenden Form. Der interessierte Leser findet eine Implementierung der Methode im Anhang A.

Algorithmus - Verfahren der konjugierten Gradienten —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$															
$\mathbf{p}_0 := \mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0$, $\alpha_0 := \ \mathbf{r}_0\ _2^2$															
Für $m = 0, \dots, n-1$															
<table border="1"> <tr> <td colspan="2">$\alpha_m \neq 0$</td><td></td></tr> <tr> <td>Y</td><td></td><td>N</td></tr> <tr> <td>$\mathbf{v}_m := \mathbf{A}\mathbf{p}_m$, $\lambda_m := \frac{\alpha_m}{(\mathbf{v}_m, \mathbf{p}_m)_2}$</td><td colspan="2" rowspan="5">STOP</td></tr> <tr> <td>$\mathbf{x}_{m+1} := \mathbf{x}_m + \lambda_m \mathbf{p}_m$</td></tr> <tr> <td>$\mathbf{r}_{m+1} := \mathbf{r}_m - \lambda_m \mathbf{v}_m$</td></tr> <tr> <td>$\alpha_{m+1} := \ \mathbf{r}_{m+1}\ _2^2$</td></tr> <tr> <td>$\mathbf{p}_{m+1} := \mathbf{r}_{m+1} + \frac{\alpha_{m+1}}{\alpha_m} \mathbf{p}_m$</td></tr> </table>			$\alpha_m \neq 0$			Y		N	$\mathbf{v}_m := \mathbf{A}\mathbf{p}_m$, $\lambda_m := \frac{\alpha_m}{(\mathbf{v}_m, \mathbf{p}_m)_2}$	STOP		$\mathbf{x}_{m+1} := \mathbf{x}_m + \lambda_m \mathbf{p}_m$	$\mathbf{r}_{m+1} := \mathbf{r}_m - \lambda_m \mathbf{v}_m$	$\alpha_{m+1} := \ \mathbf{r}_{m+1}\ _2^2$	$\mathbf{p}_{m+1} := \mathbf{r}_{m+1} + \frac{\alpha_{m+1}}{\alpha_m} \mathbf{p}_m$
$\alpha_m \neq 0$															
Y		N													
$\mathbf{v}_m := \mathbf{A}\mathbf{p}_m$, $\lambda_m := \frac{\alpha_m}{(\mathbf{v}_m, \mathbf{p}_m)_2}$	STOP														
$\mathbf{x}_{m+1} := \mathbf{x}_m + \lambda_m \mathbf{p}_m$															
$\mathbf{r}_{m+1} := \mathbf{r}_m - \lambda_m \mathbf{v}_m$															
$\alpha_{m+1} := \ \mathbf{r}_{m+1}\ _2^2$															
$\mathbf{p}_{m+1} := \mathbf{r}_{m+1} + \frac{\alpha_{m+1}}{\alpha_m} \mathbf{p}_m$															

Beispiel 4.63 Wir betrachten das eindimensionale Randwertproblem (4.2.1) mit $h_\ell = \frac{1}{8}$ (das heißt $N_\ell = 7$), wodurch sich $\mathbf{A}_\ell \mathbf{u}^\ell = \mathbf{f}^\ell$ mit

$$\mathbb{R}^{N_\ell \times N_\ell} \ni \mathbf{A}_\ell = \text{tridiag}\{-64, 128, -64\}$$

ergibt. Zudem sei die rechte Seite gemäß

$$\mathbf{f}^\ell = (128, -448, 704, -832, 512, 128, 320)^T$$

gegeben. Mit dem Startvektor $\mathbf{u}_0 = \mathbf{0}$ erhalten wir den in der folgenden Tabelle präsentierten Konvergenzverlauf:

Verfahren der konjugierten Gradienten (CG-Verfahren)								
m	$u_{m,1}$	$u_{m,2}$	$u_{m,3}$	$u_{m,4}$	$u_{m,5}$	$u_{m,6}$	$u_{m,7}$	$\ \mathbf{r}_m\ _2$
0	0.00	0.00	0.00	0.00	0.00	0.00	0.00	1336.36
1	0.58	-2.04	3.21	-3.79	2.33	0.58	1.46	363.57
2	-0.39	-1.72	2.81	-4.57	3.00	4.99	4.26	252.76
3	-0.01	-2.38	2.06	-3.53	4.87	6.07	6.25	153.30
4	-0.14	-2.88	2.57	-2.13	6.50	7.48	5.93	117.64
5	-0.70	-2.18	3.53	-1.12	7.65	7.81	6.27	103.52
6	0.13	-1.14	5.40	0.54	8.23	8.54	6.98	89.70
7	1.00	0.00	6.00	1.00	9.00	9.00	7.00	0.00

Dieses einfache Beispiel bestätigt die theoretische Eigenschaft des Verfahrens, nach spätestens N_ℓ Iterationen die exakte Lösung zu ermitteln. Im allgemeinen weist der numerische Algorithmus jedoch Ungenauigkeiten aufgrund von Rundungsfehlern auf, wodurch bei einem Gleichungssystem (4.3.1) das nach n Iterationen vorliegende Ergebnis in der Regel von der gesuchten Lösung abweicht. Daher führt man in der Praxis bei kleinem n zusätzliche Iterationen durch, bis eine hinreichend genaue Lösung berechnet wurde. Ein Algol-Programm für eine Variante dieses Algorithmus findet man in [76], einen umfangreichen Bericht über numerische Erfahrungen in [57].

Viele aus technischen Anwendungen resultierende Gleichungssysteme besitzen eine große, schwachbesetzte Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$. An einer exakten Lösung dieses Gleichungssystems ist man dabei aus folgenden zwei Gründen nicht interessiert: Zunächst würde auch bei exakter Arithmetik die eventuell benötigten n Iterationen auf extrem hohe Rechenzeiten und folglich zu einem unbrauchbaren Verfahren führen. Desweiteren stellt die exakte Lösung des Gleichungssystems gewöhnlich nur eine Approximation an die Lösung der zugrundeliegenden Problemstellung dar. Man ist deshalb nur an einer Fehlerordnung interessiert, die im Bereich der Approximationsgüte des Diskretisierungsverfahrens liegt. Diese Argumente zeigen, daß auch ein theoretisch direkter Algorithmus wie das CG-Verfahren üblicherweise nur als iterative Methode verwendet wird.

Die im CG-Algorithmus auftretenden Quotienten verdeutlichen die grundlegende Abhängigkeit des Verfahrens von der positiven Definitheit der Matrix. Gleichungssystemslöser, die die Symmetrie der Matrix ausnutzen, jedoch keine positive Definitheit fordern, werden ausführlich von Fischer [24] diskutiert.

Beispiel 4.64 Als weiteres Beispiel betrachten wir die Poisson-Gleichung (siehe Beispiel 1.1), wobei $N = 200$ gewählt wird und folglich 40.000 Unbekannte vorliegen. Zudem verwenden wir $\varphi \equiv 0$ und $f(x,y) = 2x + 2y$. Aus der folgenden Tabelle wird ersichtlich, daß der CG-Algorithmus bereits nach 641 Iterationen eine übliche Grenze der Maschinengenauigkeit erreicht hat. Auch hierdurch zeigt sich nachdrücklich, daß das Verfahren bei großen Gleichungssystemen als iterative Methode genutzt werden sollte. Innerhalb des Beispiels 4.14 haben wir bereits nachgewiesen, daß die Matrix des betrachteten Gleichungssystems den Voraussetzungen des Satzes 4.13 genügt und folglich das Jacobi-Verfahren konvergiert. Der Residuenverlauf dieser Splitting-Methode ist ebenfalls in der Tabelle enthalten und unterstreicht die durch das CG-Verfahren gewonnene immense Effizienzsteigerung.

	CG-Verfahren	Jacobi-Verfahren
Iterationen	Residuenverlauf	
0	140.348	140.348
150	1.83245	134.735
300	2.40822e-05	131.221
450	1.77391e-12	128.135
600	1.88161e-15	125.292
641	8.91038e-17	124.547

Bemerkung:

Es gilt beim CG-Verfahren stets

$$\mathbf{x}_m \in \mathbf{x}_0 + \underbrace{\text{span} \{ \mathbf{p}_0, \dots, \mathbf{p}_{m-1} \}}_{=\text{span} \{ \mathbf{r}_0, \dots, \mathbf{A}^{m-1} \mathbf{r}_0 \} = K_m},$$

und $\mathbf{x}_m$ ist optimal bezüglich K_m , da

$$\mathbf{r}_m \perp K_m$$

gilt. Das CG-Verfahren stellt somit eine orthogonale Krylov-Unterraum-Methode dar. Zudem gilt

$$\mathbf{x}_m = \arg \min_{\mathbf{x} \in \mathbf{x}_0 + K_m} F(\mathbf{x}). \quad (4.3.24)$$

Wir haben bereits beim Verfahren des steilsten Abstiegs eine Konvergenzaussage auf der Basis der Konditionszahl der Matrix des linearen Gleichungssystems herleiten können. Eine analoge Aussage werden wir nun für das CG-Verfahren formulieren.

Satz 4.65 Seien $\mathbf{A} \in \mathbb{R}^{n \times n}$ symmetrisch, positiv definit und $\{\mathbf{x}_m\}_{m \in \mathbb{N}_0}$ die durch das Verfahren der konjugierten Gradienten erzeugte Folge von Näherungslösungen. Dann erfüllt der zu $\mathbf{x}_m$ korrespondierende Fehlervektor $\mathbf{e}_m = \mathbf{x}_m - \mathbf{A}^{-1} \mathbf{b}$ die Ungleichung

$$\|\mathbf{e}_m\|_{\mathbf{A}} \leq 2 \left(\frac{\sqrt{\text{cond}_2(\mathbf{A})} - 1}{\sqrt{\text{cond}_2(\mathbf{A})} + 1} \right)^m \|\mathbf{e}_0\|_{\mathbf{A}}.$$

Beweis:

Wir nutzen wiederum die Menge $\mathcal{P}_m^1$ aller Polynome p vom Höchstgrad m , die der Nebenbedingung $p(\mathbf{0}) = \mathbf{I}$ genügen. Folglich existiert aufgrund der Gleichung

$$\mathbf{e}_m = \mathbf{x}_m - \mathbf{A}^{-1}\mathbf{b} = \underbrace{\mathbf{x}_0 - \mathbf{A}^{-1}\mathbf{b}}_{=\mathbf{e}_0} - \sum_{i=1}^m c_i \underbrace{\mathbf{A}^{i-1}\mathbf{r}_0}_{=-\mathbf{A}^i\mathbf{e}_0}$$

ein Polynom $p \in \mathcal{P}_m^1$ mit

$$\mathbf{e}_m = p(\mathbf{A})\mathbf{e}_0. \quad (4.3.25)$$

Die Gleichung (4.3.24) liefert $\|\mathbf{x}_m - \mathbf{A}^{-1}\mathbf{b}\|_{\mathbf{A}} = \min_{\mathbf{x} \in \mathbf{x}_0 + K_m} \|\mathbf{x} - \mathbf{A}^{-1}\mathbf{b}\|_{\mathbf{A}}$, wodurch

$$\|\mathbf{e}_m\|_{\mathbf{A}} = \min_{p \in \mathcal{P}_m^1} \|p(\mathbf{A})\mathbf{e}_0\|_{\mathbf{A}} \quad (4.3.26)$$

folgt. Als symmetrische und positiv definite Matrix besitzt $\mathbf{A}$ reelle und positive Eigenwerte $\lambda_n \geq \dots \geq \lambda_1 > 0$ und die zugehörigen Eigenvektoren $\mathbf{v}_1, \dots, \mathbf{v}_n$ können als Orthonormalbasis des $\mathbb{R}^n$ gewählt werden. Somit existiert eine Darstellung des Fehlervektors $\mathbf{e}_0$ in der Form

$$\mathbf{e}_0 = \sum_{i=1}^n \alpha_i \mathbf{v}_i$$

mit $\alpha_i \in \mathbb{R}$ für $i = 1, \dots, n$. Hieraus erhalten wir

$$\|\mathbf{e}_0\|_{\mathbf{A}}^2 = \left(\sum_{i=1}^n \alpha_i \mathbf{v}_i, \sum_{i=1}^n \alpha_i \lambda_i \mathbf{v}_i \right)_2 = \sum_{i=1}^n \alpha_i^2 \lambda_i$$

und analog

$$\|p(\mathbf{A})\mathbf{e}_0\|_{\mathbf{A}}^2 = \sum_{i=1}^n p(\lambda_i)^2 \alpha_i^2 \lambda_i.$$

Folglich ergibt sich unter Ausnutzung der Gleichung (4.3.26) die Ungleichung

$$\begin{aligned} \|\mathbf{e}_m\|_{\mathbf{A}} &= \min_{p \in \mathcal{P}_m^1} \|p(\mathbf{A})\mathbf{e}_0\|_{\mathbf{A}} \\ &= \min_{p \in \mathcal{P}_m^1} \left(\sum_{i=1}^n p(\lambda_i)^2 \alpha_i^2 \lambda_i \right)^{1/2} \\ &\leq \min_{p \in \mathcal{P}_m^1} \max_{j=1, \dots, n} |p(\lambda_j)| \left(\sum_{i=1}^n \alpha_i^2 \lambda_i \right)^{1/2} \\ &\leq \min_{p \in \mathcal{P}_m^1} \max_{\lambda \in [\lambda_1, \lambda_n]} |p(\lambda)| \|\mathbf{e}_0\|_{\mathbf{A}}. \end{aligned} \quad (4.3.27)$$

Die weitere Aufgabe zum Nachweis der Behauptung besteht in der Angabe eines speziellen Polynoms. Für die Tschebyscheff-Polynome $T_m : [-1, 1] \mapsto [-1, 1]$ mit

$$T_m(x) = \cos(m \arccos x), \quad m \in \mathbb{N}_0$$

läßt sich mittels einer Induktion unter Verwendung von $T_0(x) = 1$ und $T_1(x) = x$ die Rekursionsformel

$$T_{m+1}(x) = 2xT_m(x) - T_{m-1}(x) \text{ für } m \in \mathbb{N}$$

aus dem Kosinus-Additionstheorem herleiten. Diese Darstellung belegt $T_m \in \mathcal{P}_m$ und wird zudem zur Fortsetzung der Tschebyscheff-Polynome auf $\mathbb{R}$ genutzt. Ebenfalls ergibt sich durch eine vollständige Induktion die Gleichung

$$T_m\left(\frac{1}{2}\left(x + \frac{1}{x}\right)\right) = \frac{1}{2}\left(x^m + \frac{1}{x^m}\right) \text{ für } m \in \mathbb{N}_0. \quad (4.3.28)$$

Wir betrachten zunächst den Fall $\lambda_1 \neq \lambda_n$. Unter dieser Voraussetzung ist das Polynom

$$p_m(\lambda) = \frac{T_m\left(\frac{2\lambda - (\lambda_n + \lambda_1)}{\lambda_1 - \lambda_n}\right)}{T_m\left(\frac{\lambda_n + \lambda_1}{\lambda_n - \lambda_1}\right)}$$

wohldefiniert und es gilt $p_m \in \mathcal{P}_m^1$. Ausgehend von der Ungleichung (4.3.27) erhalten wir mit

$$\left|T_m\left(\frac{2\lambda - (\lambda_n + \lambda_1)}{\lambda_1 - \lambda_n}\right)\right| \leq 1 \quad \forall \lambda \in [\lambda_1, \lambda_n]$$

und

$$\frac{\lambda_n + \lambda_1}{\lambda_n - \lambda_1} = \frac{\text{cond}_2(\mathbf{A}) + 1}{\text{cond}_2(\mathbf{A}) - 1} = \frac{1}{2} \left(\frac{\sqrt{\text{cond}_2(\mathbf{A})} + 1}{\sqrt{\text{cond}_2(\mathbf{A})} - 1} + \frac{\sqrt{\text{cond}_2(\mathbf{A})} - 1}{\sqrt{\text{cond}_2(\mathbf{A})} + 1} \right) \quad (4.3.29)$$

die Abschätzung

$$\begin{aligned} \|e_m\|_{\mathbf{A}} &\leq \max_{\lambda \in [\lambda_1, \lambda_n]} |p_m(\lambda)| \|e_0\|_{\mathbf{A}} \\ &\leq \frac{1}{\left|T_m\left(\frac{\lambda_n + \lambda_1}{\lambda_n - \lambda_1}\right)\right|} \|e_0\|_{\mathbf{A}} \\ &\stackrel{(4.3.29), (4.3.28)}{=} \left| 2 \left(\underbrace{\left(\frac{\sqrt{\text{cond}_2(\mathbf{A})} + 1}{\sqrt{\text{cond}_2(\mathbf{A})} - 1} \right)^m}_{>0} + \underbrace{\left(\frac{\sqrt{\text{cond}_2(\mathbf{A})} - 1}{\sqrt{\text{cond}_2(\mathbf{A})} + 1} \right)^m}_{>0} \right)^{-1} \right| \|e_0\|_{\mathbf{A}} \\ &\leq 2 \left(\frac{\sqrt{\text{cond}_2(\mathbf{A})} - 1}{\sqrt{\text{cond}_2(\mathbf{A})} + 1} \right)^m \|e_0\|_{\mathbf{A}}. \end{aligned}$$

Für den Spezialfall $\lambda_1 = \lambda_n > 0$ nutzen wir $p_m \in \mathcal{P}_m^1$ mit $p_m(\lambda) = 1 - \frac{\lambda}{\lambda_n}$ und erhalten mit $\text{cond}_2(\mathbf{A}) = \frac{\lambda_n}{\lambda_1} = 1$ die behauptete Ungleichung

$$\|e_m\|_{\mathbf{A}} \leq \underbrace{\max_{\lambda \in [\lambda_1, \lambda_n]} |p_m(\lambda)|}_{=0} \|e_0\|_{\mathbf{A}} = 0 = 2 \left(\frac{\sqrt{\text{cond}_2(\mathbf{A})} - 1}{\sqrt{\text{cond}_2(\mathbf{A})} + 1} \right)^m \|e_0\|_{\mathbf{A}}.$$

□

4.3.2 Verfahren für reguläre Matrizen

Im Kontext komplexer Anwendungsfälle können häufig keine Aussagen über die Eigenschaften der auftretenden Gleichungssysteme getroffen werden. Bereits die Diskretisierung der Konvektions-Diffusions-Gleichung (Beispiel 1.4) zeigt, daß selbst bei einfachen Grundgleichungen die Symmetrie der Matrix nicht gewährleistet werden kann. Wir werden uns in diesem Abschnitt daher mit Verfahren beschäftigen, die neben der Regularität der Matrix zunächst keine weiteren Forderungen an das betrachtete Gleichungssystem stellen. Diese Bedingung erweist sich beim GMRES-Verfahren als theoretisch hinreichend für die Konvergenz und Stabilität des Verfahrens. Die weiteren iterativen Gleichungssystemlöser basieren direkt oder indirekt auf dem Bi-Lanczos-Algorithmus und erben demzufolge im Fall einer indefiniten Matrix auch dessen mögliche vorzeitige Verfahrensabbrüche. Trotz dieser Problematik sind die angesprochenen Algorithmen im Bereich des wissenschaftlichen Rechnens sehr verbreitet und haben sich als robust, stabil und effizient erwiesen.

4.3.2.1 Der Arnoldi-Algorithmus und die FOM

Arnoldi entwickelte 1951 ein Verfahren zur sukzessiven Transformation dichtbesetzter Matrizen auf die obere Hessenbergform. Dieser Algorithmus wurde später auch zur Ermittlung von Eigenwerten verwendet und dient in der folgenden Form zur Berechnung einer Orthonormalbasis $\mathcal{V}_m = \{\mathbf{v}_1, \dots, \mathbf{v}_m\}$ des Krylov-Unterraums K_m , die für weitere Verfahren von entscheidender Bedeutung ist.

Zur Herleitung des Arnoldi-Algorithmus setzen wir voraus, daß mit $\{\mathbf{v}_1, \dots, \mathbf{v}_j\}$ eine Orthogonalbasis des $K_j = \text{span}\{\mathbf{r}_0, \dots, \mathbf{A}^{j-1}\mathbf{r}_0\}$ für $j = 1, \dots, m$ vorliegt. Wegen

$$\mathbf{A}K_m = \text{span}\{\mathbf{A}\mathbf{r}_0, \dots, \mathbf{A}^m\mathbf{r}_0\} \subset K_{m+1}$$

liegt die Idee nahe, $\mathbf{v}_{m+1}$ in der Form

$$\mathbf{v}_{m+1} = \mathbf{A}\mathbf{v}_m + \boldsymbol{\xi} \text{ mit } \boldsymbol{\xi} \in \text{span}\{\mathbf{v}_1, \dots, \mathbf{v}_m\} = K_m$$

zu definieren. Mit $\boldsymbol{\xi} = -\sum_{j=1}^m \alpha_j \mathbf{v}_j$ folgt

$$(\mathbf{v}_{m+1}, \mathbf{v}_j)_2 = (\mathbf{A}\mathbf{v}_m, \mathbf{v}_j)_2 - \alpha_j (\mathbf{v}_j, \mathbf{v}_j)_2,$$

wodurch aufgrund der Orthogonalitätsbedingung

$$\alpha_j = \frac{(\mathbf{A}\mathbf{v}_m, \mathbf{v}_j)_2}{(\mathbf{v}_j, \mathbf{v}_j)_2}$$

für $j = 1, \dots, m$ gilt. Betrachten wir zudem ausschließlich normierte Basisvektoren, dann vereinfacht sich die Berechnung der Koeffizienten zu

$$\alpha_j = (\mathbf{v}_j, \mathbf{A}\mathbf{v}_m)_2$$

und wir erhalten unter der Voraussetzung $\mathbf{r}_0 \neq 0$ das folgende Verfahren:

Algorithmus - Arnoldi —

$\mathbf{v}_1 := \frac{\mathbf{r}_0}{\ \mathbf{r}_0\ _2}$	
Für $j = 1, \dots, m$	
Für $i = 1, \dots, j$	
$h_{ij} := (\mathbf{v}_i, \mathbf{A}\mathbf{v}_j)_2$	(4.3.30)
$\mathbf{w}_j := \mathbf{A}\mathbf{v}_j - \sum_{i=1}^j h_{ij}\mathbf{v}_i$	(4.3.31)
$h_{j+1,j} := \ \mathbf{w}_j\ _2$	(4.3.32)
Y $h_{j+1,j} \neq 0$ N	
$\mathbf{v}_{j+1} := \frac{\mathbf{w}_j}{h_{j+1,j}}$	(4.3.33) $\mathbf{v}_{j+1} := \mathbf{0}$
	STOP

Eine Implementierung dieses Verfahrens in Form eines MATLAB-Files befindet sich im Anhang A.

Satz 4.66 *Vorausgesetzt, der Arnoldi-Algorithmus bricht nicht vor der Berechnung von $\mathbf{v}_m \neq \mathbf{0}$ ab, dann stellt $\mathcal{V}_j = \{\mathbf{v}_1, \dots, \mathbf{v}_j\}$ eine Orthonormalbasis des j -ten Krylov-Unterraums K_j für $j = 1, \dots, m$ dar.*

Beweis:

Wir weisen zunächst nach, daß die Vektoren $\mathbf{v}_1, \dots, \mathbf{v}_j$ für alle $j = 1, \dots, m$ ein Orthonormalsystem (ONS) repräsentieren. Hierzu führen wir eine Induktion über j durch.

Für $j = 1$ ist die Aussage wegen $\mathbf{r}_0 \neq \mathbf{0}$ trivial. Sei $\mathcal{V}_k$ für $k = 1, \dots, j < m$ ein ONS, dann folgt unter Ausnutzung der Voraussetzung $\mathbf{v}_{j+1} \neq \mathbf{0}$ die Gleichung

$$\begin{aligned}
 (\mathbf{v}_{j+1}, \mathbf{v}_k)_2 &\stackrel{(4.3.31)}{=} \frac{1}{h_{j+1,j}} \left(\mathbf{A}\mathbf{v}_j - \sum_{i=1}^j h_{ij}\mathbf{v}_i, \mathbf{v}_k \right)_2 \\
 &= \frac{1}{h_{j+1,j}} ((\mathbf{A}\mathbf{v}_j, \mathbf{v}_k)_2 - h_{kj}) \\
 &\stackrel{(4.3.30)}{=} \frac{1}{h_{j+1,j}} ((\mathbf{A}\mathbf{v}_j, \mathbf{v}_k)_2 - (\mathbf{v}_k, \mathbf{A}\mathbf{v}_j)_2) \\
 &= 0.
 \end{aligned}$$

Die Normierung (4.3.33) liefert die Behauptung.

Wir kommen nun zum Nachweis der Basiseigenschaft und führen wiederum eine Induktion über j durch.

Für $j = 1$ ist die Aussage trivial. Sei $\mathcal{V}_k$ für $k = 1, \dots, j < m$ eine Basis von K_k , dann folgt

$$\mathbf{w}_j = \underbrace{\mathbf{A} \mathbf{v}_j}_{\in K_j} - \underbrace{\sum_{i=1}^j h_{ij} \mathbf{v}_i}_{\in K_j} \in K_{j+1}.$$

Somit gilt $\text{span}\{\mathbf{v}_1, \dots, \mathbf{v}_{j+1}\} \subset K_{j+1}$. Aufgrund der nachgewiesenen Orthonormalität der Vektoren $\mathbf{v}_1, \dots, \mathbf{v}_{j+1} \in \mathbb{R}^n \setminus \{\mathbf{0}\}$ gilt $\dim \text{span}\{\mathbf{v}_1, \dots, \mathbf{v}_{j+1}\} = j + 1$, wodurch $\text{span}\{\mathbf{v}_1, \dots, \mathbf{v}_{j+1}\} = K_{j+1}$ folgt. $\square$

Satz 4.67 *Vorausgesetzt, der Arnoldi-Algorithmus bricht nicht vor der Berechnung von $\mathbf{v}_m \neq \mathbf{0}$ ab, dann erhalten wir unter Verwendung von $\mathbf{V}_m = (\mathbf{v}_1 \dots \mathbf{v}_m) \in \mathbb{R}^{n \times m}$ mit*

$$\mathbf{H}_m := \mathbf{V}_m^T \mathbf{A} \mathbf{V}_m \in \mathbb{R}^{m \times m} \quad (4.3.34)$$

eine obere Hessenbergmatrix, für die

$$(\mathbf{H}_m)_{ij} = \begin{cases} h_{ij} & \text{aus dem Arnoldi-Algorithmus für } i \leq j + 1, \\ 0 & \text{für } i > j + 1 \end{cases}$$

gilt.

Beweis:

Bezeichnen wir die Elemente der Matrix $\mathbf{H}_m$ mit $\tilde{h}_{ij}$, dann folgt mit (4.3.34)

$$\tilde{h}_{ij} = (\mathbf{v}_i, \mathbf{A} \mathbf{v}_j)_2 \stackrel{(4.3.30)}{=} h_{ij} \quad \text{für } i \leq j.$$

Seien $j \in \{1, \dots, m-1\}$ beliebig, aber fest, dann erhalten wir für $k \in \{1, \dots, m-j\}$ die Darstellung

$$\begin{aligned} \tilde{h}_{j+k,j} &= (\mathbf{v}_{j+k}, \mathbf{A} \mathbf{v}_j)_2 \\ &\stackrel{(4.3.31)}{=} (\mathbf{v}_{j+k}, \mathbf{w}_j)_2 + \sum_{i=1}^j h_{ij} \underbrace{(\mathbf{v}_{j+k}, \mathbf{v}_i)_2}_{=0} \\ &\stackrel{(4.3.33)}{=} h_{j+1,j} (\mathbf{v}_{j+k}, \mathbf{v}_{j+1})_2 \\ &= \begin{cases} h_{j+1,j} & \text{für } k = 1 \\ 0 & \text{für } k = 2, \dots, m-j. \end{cases} \end{aligned}$$

$\square$

Satz 4.68 *Vorausgesetzt, der Arnoldi-Algorithmus bricht nicht vor der Berechnung von $\mathbf{v}_{m+1}$ ab, dann gilt*

$$\mathbf{A} \mathbf{V}_m = \mathbf{V}_{m+1} \overline{\mathbf{H}}_m, \quad (4.3.35)$$

wobei $\overline{\mathbf{H}}_m \in \mathbb{R}^{(m+1) \times m}$ durch

$$\overline{\mathbf{H}}_m = \begin{pmatrix} \mathbf{H}_m \\ 0 \dots 0 \ h_{m+1,m} \end{pmatrix} \quad (4.3.36)$$

gegeben ist.

Beweis:

Aus (4.3.31) und (4.3.33) folgt

$$\mathbf{A}\mathbf{v}_j = h_{j+1,j}\mathbf{v}_{j+1} + \sum_{i=1}^j h_{ij}\mathbf{v}_i \text{ für } j = 1, \dots, m$$

und somit (4.3.35). $\square$

Mit dem Arnoldi-Algorithmus kann direkt ein iteratives Verfahren zur Lösung der Gleichung $\mathbf{A}\mathbf{x} = \mathbf{b}$ definiert werden. Zunächst läßt sich jeder Vektor $\mathbf{x}_m \in \mathbf{x}_0 + K_m$ in der Form

$$\mathbf{x}_m = \mathbf{x}_0 + \mathbf{V}_m \boldsymbol{\alpha}_m \quad \text{mit} \quad \boldsymbol{\alpha}_m \in \mathbb{R}^m \quad (4.3.37)$$

schreiben, da die Spalten der Matrix $\mathbf{V}_m$ eine Basis des Krylov-Unterraums K_m darstellen.

Machen wir den Ansatz einer orthogonalen Krylov-Unterraum-Methode, dann ergibt sich unter Verwendung des ersten Einheitsvektors $\mathbf{e}_1 = (1, 0, \dots, 0)^T \in \mathbb{R}^m$ die Bedingung

$$\begin{aligned} \mathbf{r}_m &= \mathbf{b} - \mathbf{A}\mathbf{x}_m \perp K_m \\ \iff (\mathbf{b} - \mathbf{A}\mathbf{x}_m, \mathbf{v}_j)_2 &= 0 \text{ für } j = 1, \dots, m \\ \iff \mathbb{R}^m \ni \mathbf{0} &= \mathbf{V}_m^T (\mathbf{b} - \mathbf{A}\mathbf{x}_m) \\ &\stackrel{(4.3.37)}{=} \mathbf{V}_m^T (\mathbf{r}_0 - \mathbf{A}\mathbf{V}_m \boldsymbol{\alpha}_m) \\ &= \|\mathbf{r}_0\|_2 \mathbf{e}_1 - \underbrace{\mathbf{V}_m^T \mathbf{A} \mathbf{V}_m}_{=\mathbf{H}_m} \boldsymbol{\alpha}_m. \end{aligned} \quad (4.3.38)$$

Hiermit erhalten wir die als Full Orthogonalization Method (FOM) bezeichnete orthogonale Krylov-Unterraum-Methode in der folgenden Form:

Algorithmus - Full Orthogonalization Method (FOM) —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$ und $m \in \mathbb{N}$.	
$\mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0$	
Y $\mathbf{r}_0 \neq \mathbf{0}$ N	
Berechne $\mathbf{V}_m$ und $\mathbf{H}_m$ gemäß dem Arnoldi-Algorithmus unter Verwendung von $\mathbf{r}_0$	STOP
$\boldsymbol{\alpha}_m := \ \mathbf{r}_0\ _2 \mathbf{H}_m^{-1} \mathbf{e}_1$	
$\mathbf{x}_m := \mathbf{x}_0 + \mathbf{V}_m \boldsymbol{\alpha}_m$	

Die Lösung der Gleichung

$$\mathbf{H}_m \boldsymbol{\alpha}_m = \|\mathbf{r}_0\|_2 \mathbf{e}_1$$

kann zum Beispiel durch $m - 1$ Givens-Rotationen mit $\mathbf{G} = \mathbf{G}_{m,m-1} \cdot \dots \cdot \mathbf{G}_{2,1}$ und anschließendem Rückwärtslösen des verbleibenden Systems

$$\mathbf{R} \boldsymbol{\alpha}_m = \|\mathbf{r}_0\|_2 \mathbf{G} \mathbf{e}_1$$

mit einer rechten oberen Dreiecksmatrix $\mathbf{R}$ erfolgen. Für das Residuum ergibt sich hierbei

$$\begin{aligned} \|\mathbf{r}_m\|_2 &= \|\mathbf{b} - \mathbf{A}(\mathbf{x}_0 + \mathbf{V}_m \boldsymbol{\alpha}_m)\|_2 \\ &\stackrel{\text{Satz 4.68}}{=} \|\mathbf{V}_{m+1} (\|\mathbf{r}_0\|_2 \mathbf{e}_1 - \overline{\mathbf{H}}_m \boldsymbol{\alpha}_m)\|_2 \\ &= \left\| \mathbf{G} \|\mathbf{r}_0\|_2 \mathbf{e}_1 - \begin{pmatrix} \mathbf{R} \\ 0 \dots 0 \ h_{m+1,m} \end{pmatrix} \boldsymbol{\alpha}_m \right\|_2 \\ &= h_{m+1,m} (\boldsymbol{\alpha}_m)_m. \end{aligned}$$

Bei der FOM muß durchaus m sehr groß gewählt werden, um eine akzeptable Lösung erwarten zu dürfen. Eine Überprüfung des Residuums als Indikator für einen Verfahrensabbruch wäre daher wünschenswert. Eine Möglichkeit besteht darin, mit $\varepsilon > 0$ eine Genauigkeitsschranke vorzugeben. Erfüllt $\mathbf{x}_m$ die Bedingung $\|\mathbf{b} - \mathbf{A}\mathbf{x}_m\|_2 \leq \varepsilon$ dann wird das Verfahren beendet, ansonsten wird ein „Restart“ mit $\mathbf{x}_0 = \mathbf{x}_m$ durchgeführt. Hierdurch wird der Speicherplatzbedarf für die Orthonormalbasis auf $n \cdot m$ reelle Zahlen beschränkt.

Zur Aufwandsminimierung in Bezug auf die Rechenzeit und den benötigten Speicherplatz erhalten wir somit die folgende Variante der Full Orthogonalization Method, zu der eine MATLAB-Implementierung im Anhang A aufgeführt ist.

Algorithmus - Restarted FOM —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$, $\varepsilon > 0$ und $m \in \mathbb{N}$.					
$\mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0$					
Solange $\ \mathbf{r}_0\ _2 > \varepsilon$					
<table> <tr> <td>Berechne $\mathbf{V}_m$ und $\mathbf{H}_m$ gemäß Arnoldi-Algorithmus unter Verwendung von $\mathbf{r}_0$</td></tr> <tr> <td>$\boldsymbol{\alpha}_m := \ \mathbf{r}_0\ _2 \mathbf{H}_m^{-1} \mathbf{e}_1$</td></tr> <tr> <td>$\mathbf{x}_m := \mathbf{x}_0 + \mathbf{V}_m \boldsymbol{\alpha}_m$</td></tr> <tr> <td>$\mathbf{x}_0 := \mathbf{x}_m$</td></tr> <tr> <td>$\mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0$</td></tr> </table>	Berechne $\mathbf{V}_m$ und $\mathbf{H}_m$ gemäß Arnoldi-Algorithmus unter Verwendung von $\mathbf{r}_0$	$\boldsymbol{\alpha}_m := \ \mathbf{r}_0\ _2 \mathbf{H}_m^{-1} \mathbf{e}_1$	$\mathbf{x}_m := \mathbf{x}_0 + \mathbf{V}_m \boldsymbol{\alpha}_m$	$\mathbf{x}_0 := \mathbf{x}_m$	$\mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0$
Berechne $\mathbf{V}_m$ und $\mathbf{H}_m$ gemäß Arnoldi-Algorithmus unter Verwendung von $\mathbf{r}_0$					
$\boldsymbol{\alpha}_m := \ \mathbf{r}_0\ _2 \mathbf{H}_m^{-1} \mathbf{e}_1$					
$\mathbf{x}_m := \mathbf{x}_0 + \mathbf{V}_m \boldsymbol{\alpha}_m$					
$\mathbf{x}_0 := \mathbf{x}_m$					
$\mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0$					

Ein Problem des vorgestellten Algorithmus liegt darin, daß die Orthogonalität der Vektoren bereits nach wenigen Iterationen verloren gehen kann. Eine Verbesserung liefert der sogenannte Arnoldi-modifizierte Gram-Schmidt-Algorithmus. Dieser ergibt sich aus dem Arnoldi-Verfahren, indem (4.3.30) und (4.3.31)

Aktuelle Schleife (Arnoldi) —

Für $i = 1, \dots, j$
$h_{ij} := (\mathbf{v}_i, \mathbf{A}\mathbf{v}_j)_2$
$\mathbf{w}_j := \mathbf{A}\mathbf{v}_j - \sum_{i=1}^j h_{ij}\mathbf{v}_i$

durch

Modifizierte Schleife (Arnoldi) —

$\mathbf{w}_j := \mathbf{A}\mathbf{v}_j$
Für $i = 1, \dots, j$
$h_{ij} := (\mathbf{v}_i, \mathbf{w}_j)_2$
$\mathbf{w}_j := \mathbf{w}_j - h_{ij}\mathbf{v}_i$

(4.3.39)

ersetzt werden.

Legt man eine maximale Bandbreite der Matrix $\mathbf{H}_m$ fest und vernachlässigt anschließend alle außerhalb liegenden Matrixeinträge, so ergibt sich die IOM (Incomplete Orthogonalization Method), die analog zu DIOM (Direct IOM) eine schiefe Projektionsmethode darstellt. DIOM und IOM unterscheiden sich einzig darin, daß bei der direkten Version eine LU-Zerlegung der Matrix $\mathbf{H}_m$ vorgenommen wird, wodurch sich ein einfach programmierbarer Algorithmus ergibt. Eine spezielle Implementierung der IOM wird in [58] vorgestellt.

4.3.2.2 Der Lanczos-Algorithmus und die D-Lanczos-Methode

Mit dem Ziel der Eigenvektor- und Eigenwertbestimmung stellte Lanczos [44] bereits 1950 ein Verfahren zur Umformung symmetrischer Matrizen auf Tridiagonalgestalt vor. Wie wir sehen werden, kann die Methode als Vereinfachung des Arnoldi-Algorithmus für symmetrische Matrizen aufgefaßt werden. Betrachten wir eine symmetrische Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$, dann folgt mit (4.3.34)

$$\mathbf{H}_m = \mathbf{V}_m^T \mathbf{A} \mathbf{V}_m = \mathbf{V}_m^T \mathbf{A}^T \mathbf{V}_m = (\mathbf{V}_m^T \mathbf{A} \mathbf{V}_m)^T = \mathbf{H}_m^T.$$

Hiermit erhalten wir den Lanczos-Algorithmus aus dem Arnoldi-Algorithmus, da

- (a) das Matricelement $h_{j-1,j}$ bereits durch $h_{j,j-1}$ im vorhergehenden Schritt berechnet wurde und somit die Schleife (4.3.30) zu $h_{jj} = (\mathbf{v}_j, \mathbf{A}\mathbf{v}_j)_2$ degeneriert und
- (b) (4.3.31) sich auf $\mathbf{w}_j = \mathbf{A}\mathbf{v}_j - h_{j-1,j}\mathbf{v}_{j-1} - h_{jj}\mathbf{v}_j$ verkürzt.

Wählen wir die Notationen $a_j = h_{jj}$ und $c_j = h_{j-1,j}$, dann hat $\mathbf{H}_m$ die Gestalt

$$\mathbf{H}_m = \begin{pmatrix} a_1 & c_2 & & & \\ c_2 & a_2 & \ddots & & \\ & \ddots & \ddots & \ddots & \\ & & \ddots & a_{m-1} & c_m \\ & & & c_m & a_m \end{pmatrix}.$$

Aus der modifizierten Variante des Arnoldi-Algorithmus (siehe (4.3.39)) ergibt sich unter Ausnutzung von $h_{j-1,j} = h_{j,j-1}$ der Lanczos-Algorithmus in der Form

Algorithmus - Lanczos —

$\mathbf{v}_1 := \frac{\mathbf{r}_0}{\ \mathbf{r}_0\ _2}, c_1 := 0, \mathbf{v}_0 := \mathbf{0}$	
Für $j = 1, \dots, m$	
$\mathbf{w}_j := \mathbf{A}\mathbf{v}_j - c_j\mathbf{v}_{j-1}$	
$a_j := (\mathbf{w}_j, \mathbf{v}_j)_2$	
$\mathbf{w}_j := \mathbf{w}_j - a_j\mathbf{v}_j$	
$c_{j+1} := \ \mathbf{w}_j\ _2$	
Y $\swarrow$ $c_{j+1} \neq 0$ $\searrow$ N	
$\mathbf{v}_{j+1} := \frac{\mathbf{w}_j}{c_{j+1}}$	$\mathbf{v}_{j+1} := \mathbf{0}$
	STOP

Analog zum Arnoldi-Algorithmus lässt sich auch die Lanczos-Methode zur Lösung des linearen Gleichungssystems $\mathbf{A}\mathbf{x} = \mathbf{b}$ verwenden. Vorausgesetzt, die Matrix $\mathbf{H}_m$ lässt sich in der Form

$$\mathbf{H}_m = \underbrace{\begin{pmatrix} 1 & & & & \\ \gamma_2 & \ddots & & & 0 \\ & \ddots & \ddots & & \\ 0 & & \ddots & \ddots & \\ & & & \gamma_m & 1 \end{pmatrix}}_{=: \mathbf{L}_m} \underbrace{\begin{pmatrix} \beta_1 & \delta_2 & & & \\ & \ddots & \ddots & & 0 \\ & & \ddots & \ddots & \\ 0 & & & \ddots & \delta_m \\ & & & & \beta_m \end{pmatrix}}_{=: \mathbf{U}_m} \quad (4.3.40)$$

zerlegen, dann folgt

$$\begin{aligned}\delta_j &= c_j \text{ für } j = 2, \dots, m \\ \beta_1 &= a_1,\end{aligned}\tag{4.3.41}$$

$$\gamma_j = \frac{c_j}{\beta_{j-1}} \text{ für } j = 2, \dots, m,\tag{4.3.42}$$

$$\beta_j = a_j - \gamma_j c_j \text{ für } j = 2, \dots, m.\tag{4.3.43}$$

Analog zur FOM schreiben wir die Näherungslösung unter Verwendung des ersten Einheitsvektors $\mathbf{e}_1 \in \mathbb{R}^m$ in der Form

$$\begin{aligned}\mathbf{x}_m &= \mathbf{x}_0 + \mathbf{V}_m \mathbf{H}_m^{-1} \underbrace{\|\mathbf{r}_0\|_2 \mathbf{e}_1}_{=: \tilde{\mathbf{e}}_1} \stackrel{(4.3.40)}{=} \mathbf{x}_0 + \mathbf{V}_m \mathbf{U}_m^{-1} \mathbf{L}_m^{-1} \tilde{\mathbf{e}}_1 \\ &= \mathbf{x}_0 + \mathbf{P}_m \mathbf{z}_m\end{aligned}$$

mit $\mathbf{P}_m := \mathbf{V}_m \mathbf{U}_m^{-1} \in \mathbb{R}^{n \times m}$ und $\mathbf{z}_m := \mathbf{L}_m^{-1} \tilde{\mathbf{e}}_1$. Seien $\mathbf{p}_1, \dots, \mathbf{p}_m \in \mathbb{R}^n$ die Spaltenvektoren von $\mathbf{P}_m$, dann folgt mit

$$\mathbf{P}_m \begin{pmatrix} \beta_1 & \delta_2 & & \\ & \ddots & \ddots & \\ & & \ddots & \delta_m \\ & & & \beta_m \end{pmatrix} = (\mathbf{v}_1 \dots \mathbf{v}_m)$$

die Gleichung $\delta_m \mathbf{p}_{m-1} + \beta_m \mathbf{p}_m = \mathbf{v}_m$, wodurch sich die Vorschrift zur Berechnung der m -ten Spalte mittels

$$\mathbf{p}_m = \frac{1}{\beta_m} (\mathbf{v}_m - \delta_m \mathbf{p}_{m-1})\tag{4.3.44}$$

unter Berücksichtigung der Festlegung $\delta_0 := 0$ respektive $\mathbf{p}_0 := \mathbf{0}$ ergibt. Gilt die Beziehung $\mathbf{P}_{m-1} \mathbf{U}_{m-1} = \mathbf{V}_{m-1}$, dann folgt mit (4.3.44) die Gleichung $\mathbf{P}_m \mathbf{U}_m = \mathbf{V}_m$, wodurch sich $\mathbf{P}_m = \mathbf{V}_m \mathbf{U}_m^{-1}$ ergibt. Sei desweiteren $\mathbf{z}_{m-1} = (\xi_1, \dots, \xi_{m-1})^T \in \mathbb{R}^{m-1}$ die Lösung von $\mathbf{L}_{m-1} \mathbf{z}_{m-1} = \tilde{\mathbf{e}}_1 \in \mathbb{R}^{m-1}$, so erhalten wir durch Inkrementierung der Raumdimension

$$\mathbb{R}^m \ni \tilde{\mathbf{e}}_1 = \mathbf{L}_m \mathbf{z}_m = \begin{pmatrix} & & 0 \\ & \mathbf{L}_{m-1} & \vdots \\ & & 0 \\ 0 & \dots & 0 & \gamma_m & 1 \end{pmatrix} \begin{pmatrix} \mathbf{z}_{m-1} \\ \xi_m \end{pmatrix} = \begin{pmatrix} \tilde{\mathbf{e}}_1 \\ \gamma_m \xi_{m-1} + \xi_m \end{pmatrix}$$

und somit

$$\xi_m = -\gamma_m \xi_{m-1}.\tag{4.3.45}$$

Hierbei liefert die erste Gleichung direkt $\xi_1 = 1$. Für die Näherungslösung folgt daher

$$\begin{aligned}\mathbf{x}_m &= \mathbf{x}_0 + (\mathbf{P}_{m-1}, \mathbf{p}_m) \begin{pmatrix} \mathbf{z}_{m-1} \\ \xi_m \end{pmatrix} = \mathbf{x}_0 + \mathbf{P}_{m-1} \mathbf{z}_{m-1} + \xi_m \mathbf{p}_m \\ &= \mathbf{x}_{m-1} + \xi_m \mathbf{p}_m\end{aligned}\tag{4.3.46}$$

und wir erhalten zusammenfassend das D(Direct)-Lanczos-Verfahren:

D-Lanczos-Algorithmus —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$ und $\varepsilon > 0$		
$\mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0$		
Y	$\ \mathbf{r}_0\ _2 > \varepsilon$	N
$\mathbf{v}_1 := \frac{\mathbf{r}_0}{\ \mathbf{r}_0\ _2}$		
$c_1 := \gamma_1 := 0, \ \xi_1 := 1, \ \mathbf{p}_0 := \mathbf{v}_0 := \mathbf{0}$		
Für $j = 1, \dots$		
$\mathbf{w}_j := \mathbf{A}\mathbf{v}_j - c_j\mathbf{v}_{j-1}$		
$a_j := (\mathbf{v}_j, \mathbf{w}_j)_2$		
Y	$j > 1$	N
$\gamma_j \stackrel{(4.3.42)}{:=} \frac{c_j}{\beta_{j-1}}, \ \xi_j \stackrel{(4.3.45)}{:=} -\gamma_j \xi_{j-1}$		
$\beta_j \stackrel{(4.3.43)}{:=} a_j - \gamma_j c_j$		
$\mathbf{p}_j \stackrel{(4.3.44)}{:=} \frac{1}{\beta_j}(\mathbf{v}_j - c_j \mathbf{p}_{j-1}), \ \mathbf{x}_j \stackrel{(4.3.46)}{:=} \mathbf{x}_{j-1} + \xi_j \mathbf{p}_j$		
Y	$\ \mathbf{b} - \mathbf{A}\mathbf{x}_j\ _2 > \varepsilon$	N
$\mathbf{w}_j := \mathbf{w}_j - a_j \mathbf{v}_j$		$\mathbf{v}_{j+1} := \mathbf{0}$
$c_{j+1} := \ \mathbf{w}_j\ _2$		STOP
Y	$c_{j+1} \neq 0$	N
$\mathbf{v}_{j+1} := \frac{\mathbf{w}_j}{c_{j+1}}$		$\mathbf{v}_{j+1} := \mathbf{0}$
		STOP

Die Anwendung einer LQ-Zerlegung anstelle einer LU-Zerlegung führt auf das Verfahren SYMMLQ, welches 1975 von Paige und Saunders [54] vorgestellt wurde. Die naheliegende Idee der Nutzung einer QR-Zerlegung wird zudem in [24] vorgestellt.

Der D-Lanczos-Algorithmus wurde auf der Basis der Umformung (4.3.38) hergeleitet und stellt somit analog zur FOM eine orthogonale Krylov-Unterraum-Methode dar. Betrachten wir nun den Spezialfall einer positiv definiten und symmetrischen Matrix $\mathbf{A}$. Da $\mathbf{V}_m$

orthonormale Spaltenvektoren besitzt, erhalten wir $\mathbf{V}_m \mathbf{x} \neq \mathbf{0}$ für alle $\mathbf{x} \in \mathbb{R}^n \setminus \{\mathbf{0}\}$, wodurch sich mit

$$(\mathbf{x}, \mathbf{H}_m \mathbf{x})_2 = (\mathbf{V}_m \mathbf{x}, \mathbf{A} \mathbf{V}_m \mathbf{x})_2$$

neben der Symmetrie auch die positive Definitheit auf die Matrix $\mathbf{H}_m \in \mathbb{R}^{m \times m}$ überträgt. Mit $\mathbf{H}_m$ stellen auch alle Hauptabschnittsmatrizen $\mathbf{H}_m[k]$ für $k = 1, \dots, m$ positiv definite Matrizen und folglich auch reguläre Matrizen dar. Hierdurch folgt mit Satz 3.8 die Existenz und Eindeutigkeit der Zerlegung (4.3.40). In diesem Fall ist der Algorithmus wohldefiniert und als orthogonale Krylov-Unterraum-Methode zudem äquivalent zum CG-Verfahren. Somit können wir den D-Lanczos-Algorithmus als Verallgemeinerung des CG-Verfahrens auf indefinite symmetrische Matrizen ansehen.

4.3.2.3 Der Bi-Lanczos-Algorithmus

Die aus dem Arnoldi- und Lanczos-Algorithmus hervorgegangenen iterativen Verfahren FOM und D-Lanczos weisen unterschiedliche Vorteile und Nachteile auf.

Der Vorteil der Full Orthogonalization Method liegt in der Anwendbarkeit für beliebige reguläre Matrizen $\mathbf{A} \in \mathbb{R}^{n \times n}$, während sich der D-Lanczos-Algorithmus durch einen geringen Speicherplatzbedarf für die Orthonormalbasis $\{\mathbf{v}_1, \dots, \mathbf{v}_m\}$ des Krylov-Unterraums K_m auszeichnet. Betrachten wir die Nachteile der beiden Algorithmen, so weist die FOM einen hohen Speicherplatzbedarf für die Orthonormalbasis auf, der bei schwachbesetzten Matrizen $\mathbf{A} \in \mathbb{R}^{n \times n}$ oftmals keine n Iterationen zuläßt. Abhilfe liefert die „Restarted-Version“, die jedoch nach n Schritten auch theoretisch keine Konvergenz gewährleistet. Das D-Lanczos-Verfahren ist dagegen auf die Lösung von Gleichungssystemen mit einer symmetrischen Matrix $\mathbf{A}$ beschränkt und kann als Spezialform der FOM für symmetrische Matrizen interpretiert werden. Beide Verfahren sind für positiv definite Matrizen formal äquivalent zum Verfahren der konjugierten Gradienten.

Das Ziel des Bi-Lanczos-Algorithmus liegt in der Ermittlung einer Basis des Krylov-Unterraums K_m derart, daß hierauf beruhende iterative Verfahren den Vorteil eines geringen Speicherplatzes mit der Anwendbarkeit auf unsymmetrische Matrizen verknüpfen.

Sei $\mathbf{A} \in \mathbb{R}^{n \times n}$ eine beliebige reguläre Matrix, dann definieren wir neben

$$K_m = K_m(\mathbf{A}, \mathbf{r}_0) = \text{span} \{ \mathbf{r}_0, \mathbf{A} \mathbf{r}_0, \dots, \mathbf{A}^{m-1} \mathbf{r}_0 \}$$

den zu K_m transponierten Krylov-Unterraum

$$K_m^T = K_m(\mathbf{A}^T, \mathbf{r}_0) = \text{span} \left\{ \mathbf{r}_0, \mathbf{A}^T \mathbf{r}_0, \dots, \left(\mathbf{A}^T \right)^{m-1} \mathbf{r}_0 \right\}.$$

Definition 4.69 Seien $\mathbf{v}_1, \dots, \mathbf{v}_m, \mathbf{w}_1, \dots, \mathbf{w}_m \in \mathbb{R}^n$, dann heißen die beiden Mengen $\{\mathbf{v}_1, \dots, \mathbf{v}_m\}$ und $\{\mathbf{w}_1, \dots, \mathbf{w}_m\}$ bi-orthogonal, wenn

$$(\mathbf{v}_i, \mathbf{w}_j)_2 = \delta_{ij} c_{ij} \text{ mit } c_{ij} \in \mathbb{R} \setminus \{0\} \text{ für } i, j = 1, \dots, m \quad (4.3.47)$$

gilt. Im Fall $c_{ii} = 1$ für $i = 1, \dots, m$ heißen die Mengen bi-orthonormal.

Eine wesentliche Eigenschaft bi-orthogonaler Mengen wird durch das folgende Lemma beschrieben.

Lemma 4.70 Seien die Mengen $\{\mathbf{v}_1, \dots, \mathbf{v}_m\}$ und $\{\mathbf{w}_1, \dots, \mathbf{w}_m\}$ bi-orthogonal, dann gilt

$$\dim \text{span} \{\mathbf{v}_1, \dots, \mathbf{v}_m\} = m = \dim \text{span} \{\mathbf{w}_1, \dots, \mathbf{w}_m\}$$

Beweis:

Aus Symmetriegründen ist es hierbei ausreichend, die Behauptung für die Vektoren $\mathbf{v}_1, \dots, \mathbf{v}_m$ nachzuweisen. Gelte $\mathbf{0} = \sum_{i=1}^m \lambda_i \mathbf{v}_i$, so erhalten wir für $j = 1, \dots, m$ die Darstellung

$$\lambda_j = \frac{1}{c_{jj}} \sum_{i=1}^m \lambda_i c_{ij} \delta_{ij} = \frac{1}{c_{jj}} \sum_{i=1}^m \lambda_i (\mathbf{v}_i, \mathbf{w}_j)_2 = \frac{1}{c_{jj}} \left(\sum_{i=1}^m \lambda_i \mathbf{v}_i, \mathbf{w}_j \right)_2 = \frac{1}{c_{jj}} (\mathbf{0}, \mathbf{w}_j)_2 = 0.$$

□

Der Bi-Lanczos-Algorithmus basiert auf der folgenden Idee: Anstelle eine Orthonormalbasis des K_m zu ermitteln, wird simultan eine Basis $\{\mathbf{v}_1, \dots, \mathbf{v}_m\}$ des K_m und eine Basis $\{\mathbf{w}_1, \dots, \mathbf{w}_m\}$ des K_m^T berechnet, die der Bi-Orthogonalitätsbedingung (4.3.47) genügen.

Bi-Lanczos-Algorithmus —

$h_{1,0} = h_{0,1} := 0$	
$\mathbf{v}_0 = \mathbf{w}_0 := \mathbf{0}$	
$\mathbf{v}_1 = \mathbf{w}_1 := \frac{\mathbf{r}_0}{\ \mathbf{r}_0\ _2} \quad (4.3.48)$	
für $j = 1, \dots, m$	
$h_{jj} := (\mathbf{w}_j, \mathbf{A} \mathbf{v}_j)_2 \quad (4.3.49)$	
$\mathbf{v}_{j+1}^* := \mathbf{A} \mathbf{v}_j - h_{jj} \mathbf{v}_j - h_{j-1,j} \mathbf{v}_{j-1} \quad (4.3.50)$	
$\mathbf{w}_{j+1}^* := \mathbf{A}^T \mathbf{w}_j - h_{jj} \mathbf{w}_j - h_{j,j-1} \mathbf{w}_{j-1} \quad (4.3.51)$	
$h_{j+1,j} := (\mathbf{v}_{j+1}^*, \mathbf{w}_{j+1}^*)_2 ^{1/2}$	
Y $\backslash$ $h_{j+1,j} \neq 0$ $\backslash$ N	
$h_{j,j+1} := \frac{(\mathbf{v}_{j+1}^*, \mathbf{w}_{j+1}^*)_2}{h_{j+1,j}} \quad (4.3.52)$	$h_{j,j+1} := 0$
$\mathbf{v}_{j+1} := \frac{\mathbf{v}_{j+1}^*}{h_{j+1,j}} \quad (4.3.53)$	$\mathbf{v}_{j+1} := \mathbf{0}$
$\mathbf{w}_{j+1} := \frac{\mathbf{w}_{j+1}^*}{h_{j,j+1}} \quad (4.3.54)$	$\mathbf{w}_{j+1} := \mathbf{0}$
	STOP

Mit dem nächsten Satz werden wir nachweisen, daß der vorliegende Bi-Lanczos-Algorithmus tatsächlich bi-orthonormale Basen der Krylov-Unterräume K_m und K_m^T generiert.

Satz 4.71 *Vorausgesetzt, der Bi-Lanczos-Algorithmus bricht nicht vor der Berechnung von $\mathbf{w}_m \neq \mathbf{0}$ ab, dann stellen $\mathcal{V}_j = \{\mathbf{v}_1, \dots, \mathbf{v}_j\}$ und $\mathcal{W}_j = \{\mathbf{w}_1, \dots, \mathbf{w}_j\}$ eine Basis des Krylov-Unterraums K_j beziehungsweise K_j^T für $j = 1, \dots, m$ dar. Zudem erfüllen die Basisvektoren die Bi-Orthogonalitätsbedingung (4.3.47) mit $c_{ii} = 1$ für $i = 1, \dots, m$.*

Beweis:

Den Nachweis der Bi-Orthonormalität erhalten wir durch eine vollständige Induktion über j .

Für $j = 1$ folgt

$$(\mathbf{w}_1, \mathbf{v}_1)_2 = \frac{(\mathbf{r}_0, \mathbf{r}_0)_2}{\|\mathbf{r}_0\|_2^2} = 1.$$

Sei die Bedingung für $\ell = 1, \dots, j < m$ erfüllt, dann erhalten wir

$$\begin{aligned} (\mathbf{v}_{j+1}, \mathbf{w}_j)_2 &\stackrel{(4.3.50)}{=} \frac{1}{h_{j+1,j}} \left(\underbrace{(\mathbf{A}\mathbf{v}_j, \mathbf{w}_j)_2}_{=h_{jj}} - h_{jj} \underbrace{(\mathbf{v}_j, \mathbf{w}_j)_2}_{=1} - h_{j-1,j} \underbrace{(\mathbf{v}_{j-1}, \mathbf{w}_j)_2}_{=0} \right) \\ &= 0, \end{aligned}$$

wobei sich im Fall $j = 1$ die Eigenschaft $(\mathbf{v}_{j-1}, \mathbf{w}_j)_2 = (\mathbf{v}_0, \mathbf{w}_1)_2 = 0$ durch die Definition von $\mathbf{v}_0 := 0$ ergibt. Zudem gilt für $0 < i < j$

$$\begin{aligned} (\mathbf{A}\mathbf{v}_j, \mathbf{w}_i)_2 &= (\mathbf{v}_j, \mathbf{A}^T \mathbf{w}_i)_2 \\ &\stackrel{(4.3.51)}{=} \stackrel{(4.3.54)}{=} h_{i,i+1} (\mathbf{v}_j, \mathbf{w}_{i+1})_2 + \underbrace{(\mathbf{v}_j, h_{ii} \mathbf{w}_i + h_{i,i-1} \mathbf{w}_{i-1})_2}_{=0} \\ &= \begin{cases} h_{j-1,j} & \text{für } i = j-1, \\ 0 & \text{für } i < j-1. \end{cases} \end{aligned}$$

Die Verwendung der beiden obigen Aussagen liefert für $0 < i < j$ die Gleichung

$$\begin{aligned} (\mathbf{v}_{j+1}, \mathbf{w}_i)_2 &\stackrel{(4.3.50)}{=} \stackrel{(4.3.53)}{=} \frac{1}{h_{j+1,j}} \left(\underbrace{(\mathbf{A}\mathbf{v}_j, \mathbf{w}_i)_2}_{= \begin{cases} h_{j-1,j}, & i=j-1 \\ 0, & j < j-1 \end{cases}} - h_{jj} \underbrace{(\mathbf{v}_j, \mathbf{w}_i)_2}_{=0} - h_{j-1,j} \underbrace{(\mathbf{v}_{j-1}, \mathbf{w}_i)_2}_{= \begin{cases} 1, & i=j-1 \\ 0, & i < j-1 \end{cases}} \right) \\ &= 0. \end{aligned}$$

Analog ergibt sich $(\mathbf{w}_{j+1}, \mathbf{v}_i)_2 = 0$ für $i \leq j$. Mit

$$\begin{aligned}
(\mathbf{v}_{j+1}, \mathbf{w}_{j+1})_2 &= \left(\frac{\mathbf{v}_{j+1}^*}{h_{j+1,j}}, \frac{\mathbf{w}_{j+1}^*}{h_{j,j+1}} \right)_2 \\
&\stackrel{(4.3.52)}{=} \frac{1}{h_{j+1,j}} \cdot \frac{h_{j+1,j}}{(\mathbf{v}_{j+1}^*, \mathbf{w}_{j+1}^*)_2} \cdot (\mathbf{v}_{j+1}^*, \mathbf{w}_{j+1}^*)_2 \\
&= 1
\end{aligned}$$

erhalten wir abschließend die behauptete Bi-Orthonormalität der Vektoren.

Seien $q_0, \dots, q_{m-1}$ beliebige Polynome mit $\text{grad}(q_i) = i$ für $i = 0, \dots, m-1$ und $q_0 \not\equiv 0$, dann läßt sich jedes $\mathbf{x} \in K_m$ in der Form

$$\mathbf{x} = \sum_{i=0}^{m-1} \alpha_i \mathbf{A}^i \mathbf{r}_0 = \sum_{i=0}^{m-1} \beta_i q_i(\mathbf{A}) \mathbf{r}_0$$

darstellen, wodurch

$$\text{span} \{q_0(\mathbf{A}) \mathbf{r}_0, \dots, q_{m-1}(\mathbf{A}) \mathbf{r}_0\} = K_m \quad (4.3.55)$$

folgt.

Basierend auf der Gleichung (4.3.55) erhalten wir den Nachweis der Erzeugendeneigenschaften durch eine vollständige Induktion über j . Für $j = 1$ folgt

$$\mathbf{v}_1 = \frac{\mathbf{r}_0}{\|\mathbf{r}_0\|_2} = q_0(\mathbf{A}) \mathbf{r}_0 \quad \text{mit} \quad q_0(\mathbf{A}) = \frac{1}{\|\mathbf{r}_0\|_2} \mathbf{A}^0.$$

Mit $\text{grad}(q_0) = 0$ und $q_0(\mathbf{A}) \neq \mathbf{0}$ erhalten wir folglich $\text{span}\{\mathbf{v}_1\} = K_1$.

Für $j = 2$ ergibt sich aus (4.3.49) und (4.3.53) unter Berücksichtigung von $\mathbf{v}_0 = \mathbf{0}$ und $h_{j+1,j} \neq 0$ die Darstellung

$$\mathbf{v}_2 = \frac{1}{h_{2,1}} (\mathbf{A} \mathbf{v}_1 - h_{2,2} \mathbf{v}_1 - h_{1,2} \mathbf{v}_0) = \frac{1}{h_{2,1}} \underbrace{((\mathbf{A} - h_{2,2} \mathbf{I}) q_0(\mathbf{A}))}_{q_1(\mathbf{A}) :=} \mathbf{r}_0$$

mit $\text{grad}(q_1) = 1$. Sei die Erzeugendeneigenschaft für $\ell = 1, \dots, j < m$ mit $j \geq 2$ erfüllt, dann folgt für die im Bi-Lanczos-Algorithmus berechneten skalaren Größen $h_{j+1,j}$ aufgrund des vorausgesetzten Nichtabbruchs des Verfahrens unter Berücksichtigung von $j < m$ die Eigenschaft $h_{j+1,j} \neq 0$. Hierdurch ergibt sich

$$\mathbf{v}_{j+1} \stackrel{(4.3.50)}{=} \frac{1}{h_{j+1,j}} (\mathbf{A} \mathbf{v}_j - h_{j,j} \mathbf{v}_j - h_{j-1,j} \mathbf{v}_{j-1}) = q_j(\mathbf{A}) \mathbf{r}_0$$

mit

$$q_j(\mathbf{A}) = \frac{1}{h_{j+1,j}} ((\mathbf{A} - h_{j,j} \mathbf{I}) q_{j-1}(\mathbf{A}) - h_{j-1,j} q_{j-2}(\mathbf{A})).$$

Somit erhalten wir $\text{grad}(q_j) = j$, wodurch

$$\text{span}\{\mathbf{v}_1, \dots, \mathbf{v}_{j+1}\} = \text{span}\{q_0(\mathbf{A}) \mathbf{r}_0, \dots, q_j(\mathbf{A}) \mathbf{r}_0\} = K_{j+1}$$

folgt. Der Beweis für K_m^T ergibt sich analog. Die lineare Unabhängigkeit folgt durch Lemma 4.70 aus der Bi-Orthonormalität der Vektoren. $\square$

Die in den Sätzen 4.67 und 4.68 beschriebenen Eigenschaften der Basisvektoren haben sich bereits bei der Beschreibung der FOM als hilfreich erwiesen und werden auch bei der Herleitung des GMRES-Verfahrens von entscheidender Bedeutung sein. Analoge und für die weiteren Verfahren grundlegende Aussagen hinsichtlich der vorliegenden bi-orthonormalen Basisvektoren liefert der folgende Satz.

Satz 4.72 *Vorausgesetzt, der Bi-Lanczos-Algorithmus bricht nicht vor der Berechnung von $\mathbf{w}_m \neq \mathbf{0}$ ab, dann seien $\mathbf{V}_m = (\mathbf{v}_1 \dots \mathbf{v}_m) \in \mathbb{R}^{n \times m}$ und $\mathbf{W}_m = (\mathbf{w}_1 \dots \mathbf{w}_m) \in \mathbb{R}^{n \times m}$ die aus den ermittelten Basisvektoren des K_m und K_m^T gebildeten Matrizen. Desweiteren seien h_{ij} die im Algorithmus auftretenden skalaren Größen, dann gelten die Gleichungen*

$$\mathbf{T}_m = \mathbf{W}_m^T \mathbf{A} \mathbf{V}_m \quad (4.3.56)$$

$$\mathbf{A} \mathbf{V}_m = \mathbf{V}_m \mathbf{T}_m + h_{m+1,m} \mathbf{v}_{m+1} \mathbf{e}_m^T \quad (4.3.57)$$

$$\mathbf{A}^T \mathbf{W}_m = \mathbf{W}_m \mathbf{T}_m^T + h_{m,m+1} \mathbf{w}_{m+1} \mathbf{e}_m^T \quad (4.3.58)$$

mit $\mathbf{e}_m = (0, \dots, 0, 1)^T \in \mathbb{R}^m$ sowie der durch

$$(\mathbf{T}_m)_{ij} := \begin{cases} h_{ij} & \text{für } j-1 \leq i \leq j+1, \\ 0 & \text{sonst} \end{cases}$$

gegebenen Tridiagonalmatrix $\mathbf{T}_m \in \mathbb{R}^{m \times m}$ und die Matrizen $\mathbf{V}_m$ und $\mathbf{W}_m$ besitzen maximalen Rang.

Beweis:

Aus (4.3.49) folgt direkt

$$(\mathbf{T}_m)_{jj} = h_{jj} = (\mathbf{w}_j, \mathbf{A} \mathbf{v}_j)_2.$$

Die Definition des Vektors $\mathbf{v}_{j+1}^*$ innerhalb des Bi-Lanczos-Verfahrens liefert mit (4.3.53) für $i > j$ die Gleichung

$$\begin{aligned} (\mathbf{w}_i, \mathbf{A} \mathbf{v}_j)_2 &= \underbrace{(\mathbf{w}_i, \mathbf{v}_{j+1}^*)_2}_{= h_{j+1,j} (\mathbf{w}_i, \mathbf{v}_{j+1})_2} + h_{jj} \underbrace{(\mathbf{w}_i, \mathbf{v}_j)_2}_{= 0} + h_{j,j-1} \underbrace{(\mathbf{w}_i, \mathbf{v}_{j-1})_2}_{= 0} \\ &= \begin{cases} h_{j+1,j} & \text{für } i = j+1, \\ 0 & \text{für } i > j+1. \end{cases} \end{aligned}$$

Analog erhalten wir mit der Definition von $\mathbf{w}_{j+1}^*$ und (4.3.51) für $i < j$

$$(\mathbf{w}_i, \mathbf{A} \mathbf{v}_j) = \left(\mathbf{A}^T \mathbf{w}_i, \mathbf{v}_j \right) = \begin{cases} h_{j-1,j} & \text{für } i = j-1, \\ 0 & \text{für } i < j-1, \end{cases}$$

womit die behauptete Darstellung der Matrix $\mathbf{T}_m$ gemäß (4.3.56) gezeigt ist.

Berücksichtigen wir die Gestalt der Matrix $\mathbf{T}_m$, dann stellen die Gleichungen (4.3.57) und (4.3.58) die Matrixdarstellungen der Gleichungen (4.3.50) und (4.3.51) unter Verwendung von (4.3.53) respektive (4.3.54) dar.

Der Nachweis $\text{rang}(\mathbf{V}_m) = \text{rang}(\mathbf{W}_m) = m$ folgt unmittelbar aus der Basiseigenschaft der entsprechenden Spaltenvektoren. $\square$

Der Bi-Lanczos-Algorithmus terminiert in der präsentierten Form im Fall $h_{j+1,j} = 0$. Eine derartige Situation tritt auf, wenn einer der beiden Basisvektoren identisch verschwindet. Gilt $\mathbf{v}_{j+1}^* = \mathbf{0}$, so liegt mit K_j ein $\mathbf{A}$ -invarianter Krylov-Unterraum vor, d. h. es gilt $K_j = K_{j+1}$, und hierauf basierende orthogonale als auch schiefe Krylov-Unterraum-Methoden liefern die exakte Lösung des Gleichungssystems, siehe [61]. Eine analoge Aussage ergibt sich für $\mathbf{w}_{j+1}^* = \mathbf{0}$. Diese beiden Fälle werden als reguläre Verfahrensabbrüche beschrieben. Bei orthogonalen Vektoren $\mathbf{v}_{j+1}^* \neq \mathbf{0} \neq \mathbf{w}_{j+1}^*$ liegt ein sogenannter ernsthafter Abbruch (serious breakdown) vor, da hierauf beruhende Projektionsmethoden in der Regel ohne Berechnung einer akzeptablen Näherungslösung terminieren. Diese Problematik kann oftmals durch eine vorausschauende Variante des Verfahrens (Look-ahead Lanczos algorithm) behoben werden. Hierbei wird ausgenutzt, daß häufig die Vektoren $\mathbf{v}_{j+2}^*$ und $\mathbf{w}_{j+2}^*$ auch ohne explizite Verwendung von $\mathbf{v}_{j+1}^*$ und $\mathbf{w}_{j+1}^*$ definiert werden können. Sollten auch $\mathbf{v}_{j+2}^*$ und $\mathbf{w}_{j+2}^*$ nicht ermittelt werden können, dann versucht der Algorithmus die Vektoren $\mathbf{v}_{j+3}^*$ und $\mathbf{w}_{j+3}^*$ zu berechnen, usw.. Gelingt die Berechnung eines solchen Vektorpaares, so kann der Algorithmus in seiner ursprünglichen Form weitergeführt werden. Die auf dem Bi-Lanczos-Algorithmus beruhenden Iterationsverfahren haben sich innerhalb praktischer Anwendungen auch bei Nutzung der beschriebenen Lanczos-Version als stabil erwiesen [48, 50]. Wir verzichten daher auf eine detaillierte Beschreibung der vorausschauenden Variante und verweisen den interessierten Leser auf die Arbeiten von Brezinski, Zaglia, Sadok [13] und Parlett, Taylor, Liu [55], wobei in letzterem auch sogenannte unheilbare Verfahrensabbrüche diskutiert werden.

4.3.2.4 Das GMRES-Verfahren

Das von Saad und Schultz [62] 1986 vorgestellte GMRES-Verfahren (Generalized Minimal Residual) ist bei beliebigen regulären Matrizen anwendbar. Analog zum CG-Algorithmus kann das Verfahren formal als direktes und iteratives Verfahren aufgefaßt werden. Die Verwendung des GMRES-Verfahrens in einer direkten Form ist jedoch in der Regel aufgrund des benötigten Speicherplatzes nicht praktikabel. Der Algorithmus wird daher bei praxisrelevanten Problemstellungen zumeist in einer Restarted-Version genutzt, die bereits bei der FOM vorgestellt wurde. Das Verfahren kann auf zwei unterschiedliche Arten betrachtet werden.

Einen Zugang erhalten wir, indem wir GMRES als Krylov-Unterraum-Methode betrachten, bei der die Projektion durch eine Petrov-Galerkin-Bedingung mit $L_m = \mathbf{A}K_m$ gegeben ist und das Verfahren somit eine schiefe Projektionsmethode darstellt.

Die zweite Möglichkeit zur Herleitung des Verfahrens besteht in der Umformulierung des linearen Gleichungssystems in ein Minimierungsproblem. Wir definieren hierzu im Gegensatz zum Verfahren der konjugierten Gradienten die Funktion F durch

$$\begin{aligned} F : \mathbb{R}^n &\rightarrow \mathbb{R} \\ \mathbf{x} &\mapsto \|\mathbf{b} - \mathbf{A}\mathbf{x}\|_2^2. \end{aligned} \quad (4.3.59)$$

Lemma 4.73 Seien $\mathbf{A} \in \mathbb{R}^{n \times n}$ regulär und $\mathbf{b} \in \mathbb{R}^n$, dann gilt mit der durch (4.3.59) gegebenen Funktion F

$$\hat{\mathbf{x}} = \arg \min_{\mathbf{x} \in \mathbb{R}^n} F(\mathbf{x})$$

genau dann, wenn $\hat{\mathbf{x}} = \mathbf{A}^{-1}\mathbf{b}$ gilt.

Beweis:

Es gilt $F(\mathbf{x}) = \|\mathbf{b} - \mathbf{A}\mathbf{x}\|_2^2 \geq 0$ für alle $\mathbf{x} \in \mathbb{R}^n$. Somit erhalten wir durch

$$F(\hat{\mathbf{x}}) = 0 \quad \Leftrightarrow \quad \mathbf{A}\hat{\mathbf{x}} = \mathbf{b}$$

die Aussage

$$\hat{\mathbf{x}} = \arg \min_{\mathbf{x} \in \mathbb{R}^n} F(\mathbf{x}) = \mathbf{A}^{-1}\mathbf{b}.$$

□

Zum Abschluß dieser kurzen Einordnung der GMRES-Methode müssen wir noch nachweisen, daß beide Herleitungen zum gleichen Verfahren führen, das heißt, daß die bei den unterschiedlichen Bedingungen erzielten Näherungslösungen übereinstimmen. Diese Aussage liefert uns das folgende Lemma.

Lemma 4.74 Seien $F : \mathbb{R}^n \rightarrow \mathbb{R}$ durch (4.3.59) gegeben und $\mathbf{x}_0 \in \mathbb{R}^n$ beliebig. Dann folgt

$$\tilde{\mathbf{x}} = \arg \min_{\mathbf{x} \in \mathbf{x}_0 + K_m} F(\mathbf{x}) \quad (4.3.60)$$

genau dann, wenn

$$\mathbf{b} - \mathbf{A}\tilde{\mathbf{x}} \perp L_m = \mathbf{A}K_m \quad (4.3.61)$$

gilt.

Beweis:

Seien $\mathbf{x}_m, \tilde{\mathbf{x}} \in \mathbf{x}_0 + K_m$ mit $\tilde{\mathbf{x}} = \mathbf{x}_0 + \mathbf{z}$ und $\mathbf{x}_m = \mathbf{x}_0 + \mathbf{z}_m$, dann erhalten wir

$$\begin{aligned} & F(\mathbf{x}_m) - F(\tilde{\mathbf{x}}) \\ &= (\mathbf{A}\mathbf{x}_m - \mathbf{b}, \mathbf{A}\mathbf{x}_m - \mathbf{b})_2 - (\mathbf{A}\tilde{\mathbf{x}} - \mathbf{b}, \mathbf{A}\tilde{\mathbf{x}} - \mathbf{b})_2 \\ &= (\mathbf{A}(\mathbf{x}_0 + \mathbf{z}_m) - \mathbf{b}, \mathbf{A}(\mathbf{x}_0 + \mathbf{z}_m) - \mathbf{b})_2 - (\mathbf{A}(\mathbf{x}_0 + \mathbf{z}) - \mathbf{b}, \mathbf{A}(\mathbf{x}_0 + \mathbf{z}) - \mathbf{b})_2 \\ &= (\mathbf{A}\mathbf{z}_m, \mathbf{A}\mathbf{z}_m)_2 + 2(\mathbf{A}\mathbf{x}_0 - \mathbf{b}, \mathbf{A}\mathbf{z}_m)_2 + (\mathbf{A}\mathbf{x}_0 - \mathbf{b}, \mathbf{A}\mathbf{x}_0 - \mathbf{b})_2 \\ &\quad - (\mathbf{A}\mathbf{z}, \mathbf{A}\mathbf{z})_2 - 2(\mathbf{A}\mathbf{x}_0 - \mathbf{b}, \mathbf{A}\mathbf{z})_2 - (\mathbf{A}\mathbf{x}_0 - \mathbf{b}, \mathbf{A}\mathbf{x}_0 - \mathbf{b})_2 \\ &= \underbrace{(\mathbf{A}\mathbf{z}_m, \mathbf{A}\mathbf{z}_m)_2 - 2(\mathbf{A}\mathbf{z}, \mathbf{A}\mathbf{z}_m)_2 + (\mathbf{A}\mathbf{z}, \mathbf{A}\mathbf{z})_2}_{= (\mathbf{A}(\mathbf{z}_m - \mathbf{z}), \mathbf{A}(\mathbf{z}_m - \mathbf{z}))_2} \\ &\quad + 2 \underbrace{\left\{ (\mathbf{A}\mathbf{x}_0 - \mathbf{b}, \mathbf{A}\mathbf{z}_m)_2 + (\mathbf{A}\mathbf{z}, \mathbf{A}\mathbf{z}_m)_2 - (\mathbf{A}\mathbf{x}_0 - \mathbf{b}, \mathbf{A}\mathbf{z})_2 - (\mathbf{A}\mathbf{z}, \mathbf{A}\mathbf{z})_2 \right\}}_{= (\mathbf{A}\tilde{\mathbf{x}} - \mathbf{b}, \mathbf{A}(\mathbf{z}_m - \mathbf{z}))_2} \\ &= (\mathbf{A}(\mathbf{z}_m - \mathbf{z}), \mathbf{A}(\mathbf{z}_m - \mathbf{z}))_2 + 2(\mathbf{A}\tilde{\mathbf{x}} - \mathbf{b}, \mathbf{A}(\mathbf{z}_m - \mathbf{z}))_2. \end{aligned} \quad (4.3.62)$$

„(4.3.61) $\Rightarrow$ (4.3.60)“

Gelte $\mathbf{b} - \mathbf{A}\tilde{\mathbf{x}} \perp \mathbf{A}K_m$.

Dann ergibt sich $(\mathbf{b} - \mathbf{A}\tilde{\mathbf{x}}, \mathbf{A}\mathbf{z})_2 = 0$ für alle $\mathbf{z} \in K_m$ und wir erhalten

$$F(\mathbf{x}_m) - F(\tilde{\mathbf{x}}) \stackrel{(4.3.62)}{=} \|\mathbf{A}(\mathbf{z}_m - \mathbf{z})\|_2^2.$$

Unter Berücksichtigung der Regularität der Matrix $\mathbf{A}$ folgt somit

$$F(\mathbf{x}_m) > F(\tilde{\mathbf{x}}) \text{ für alle } \mathbf{x}_m \in \{\mathbf{x}_0 + K_m\} \setminus \{\tilde{\mathbf{x}}\}.$$

„(4.3.60) $\Rightarrow$ (4.3.61)“

Gelte $\tilde{\mathbf{x}} = \arg \min_{\mathbf{x} \in \mathbf{x}_0 + K_m} F(\mathbf{x})$.

Wir führen einen Widerspruchsbeweis durch und nehmen hierzu an, daß ein $\mathbf{z}_m \in K_m$ derart existiert, daß

$$(\mathbf{b} - \mathbf{A}\tilde{\mathbf{x}}, \mathbf{A}\mathbf{z}_m)_2 = \varepsilon \neq 0$$

gilt. O.B.d.A. können wir $\varepsilon > 0$ voraussetzen. Mit der Regularität der Matrix $\mathbf{A}$ gilt $\mathbf{z}_m \neq \mathbf{0}$ und wir definieren

$$\eta := (\mathbf{A}\mathbf{z}_m, \mathbf{A}\mathbf{z}_m)_2 > 0.$$

Desweiteren betrachten wir für gegebenes $\xi \in \mathbb{R}$ mit $0 < \xi < \frac{2\varepsilon}{\eta}$ die Vektoren

$$\mathbf{z}_m^\xi := \xi \mathbf{z}_m + \mathbf{z} \in K_m \tag{4.3.63}$$

und

$$\mathbf{x}_m^\xi := \mathbf{x}_0 + \mathbf{z}_m^\xi.$$

Somit folgt

$$\begin{aligned} F(\mathbf{x}_m^\xi) - F(\tilde{\mathbf{x}}) &\stackrel{(4.3.62)}{=} (\mathbf{A}(\mathbf{z}_m^\xi - \mathbf{z}), \mathbf{A}(\mathbf{z}_m^\xi - \mathbf{z}))_2 + 2(\mathbf{A}\tilde{\mathbf{x}} - \mathbf{b}, \mathbf{A}(\mathbf{z}_m^\xi - \mathbf{z}))_2 \\ &\stackrel{(4.3.63)}{=} (\mathbf{A}(\xi \mathbf{z}_m), \mathbf{A}(\xi \mathbf{z}_m))_2 + 2(\mathbf{A}\tilde{\mathbf{x}} - \mathbf{b}, \mathbf{A}(\xi \mathbf{z}_m))_2 \\ &= \xi^2 \eta - 2\xi \varepsilon \\ &= \xi(\xi \eta - 2\varepsilon) \\ &< 0. \end{aligned}$$

Hiermit liegt ein Widerspruch zur Minimalitätseigenschaft von $\tilde{\mathbf{x}}$ vor. $\square$

Die GMRES-Methode basiert auf der vom Arnoldi-Algorithmus berechneten Orthonormalbasis $\{\mathbf{v}_1, \dots, \mathbf{v}_m\}$ des K_m . Prinzipiell kann die Ermittlung der Orthonormalbasis auch mittels hiervon abweichender Verfahren wie zum Beispiel einer Householder Transformation durchgeführt werden. Eine derartige Implementierung wird in [73] beschrieben. Bei der Herleitung gehen wir von der Darstellung des Verfahrens als Minimierungsproblem aus. Sei $\mathbf{V}_m = (\mathbf{v}_1 \dots \mathbf{v}_m) \in \mathbb{R}^{n \times m}$, dann läßt sich jedes $\mathbf{x}_m \in \mathbf{x}_0 + K_m$ in der Form

$$\mathbf{x}_m = \mathbf{x}_0 + \mathbf{V}_m \boldsymbol{\alpha}_m \text{ mit } \boldsymbol{\alpha}_m \in \mathbb{R}^m$$

darstellen. Mit

$$\begin{aligned} J_m : \mathbb{R}^m &\rightarrow \mathbb{R} \\ \boldsymbol{\alpha} &\mapsto \|\mathbf{b} - \mathbf{A}(\mathbf{x}_0 + \mathbf{V}_m \boldsymbol{\alpha})\|_2 \end{aligned} \quad (4.3.64)$$

ist die Minimierung gemäß (4.3.60) äquivalent zu

$$\boldsymbol{\alpha}_m := \arg \min_{\boldsymbol{\alpha} \in \mathbb{R}^m} J_m(\boldsymbol{\alpha}), \quad (4.3.65)$$

$$\mathbf{x}_m := \mathbf{x}_0 + \mathbf{V}_m \boldsymbol{\alpha}_m. \quad (4.3.66)$$

Zwei zentrale Ziele der weiteren Untersuchung sind nun eine möglichst einfache Berechnung von $\boldsymbol{\alpha}_m$ zu finden und $\boldsymbol{\alpha}_m$ erst dann berechnen zu müssen, wenn

$$\|\mathbf{b} - \mathbf{A}\mathbf{x}_m\|_2 \leq \varepsilon$$

mit einer vorgegebenen Genauigkeitsschranke $\varepsilon > 0$ gilt.

Mit dem Residuenvektor $\mathbf{r}_0 = \mathbf{b} - \mathbf{A}\mathbf{x}_0$ schreiben wir unter Verwendung des ersten Einheitsvektors $\mathbf{e}_1 = (1, 0, \dots, 0)^T \in \mathbb{R}^{m+1}$

$$\begin{aligned} J_m(\boldsymbol{\alpha}) &= \|\mathbf{b} - \mathbf{A}(\mathbf{x}_0 + \mathbf{V}_m \boldsymbol{\alpha})\|_2 \\ &= \|\mathbf{r}_0 - \mathbf{A}\mathbf{V}_m \boldsymbol{\alpha}\|_2 \\ &= \|\|\mathbf{r}_0\|_2 \mathbf{v}_1 - \mathbf{A}\mathbf{V}_m \boldsymbol{\alpha}\|_2 \\ &\stackrel{\text{Satz 4.68}}{=} \|\|\mathbf{r}_0\|_2 \mathbf{v}_1 - \mathbf{V}_{m+1} \overline{\mathbf{H}}_m \boldsymbol{\alpha}\|_2 \\ &= \|\mathbf{V}_{m+1} (\|\mathbf{r}_0\|_2 \mathbf{e}_1 - \overline{\mathbf{H}}_m \boldsymbol{\alpha})\|_2. \end{aligned} \quad (4.3.67)$$

Hierbei stellt $\overline{\mathbf{H}}_m$ die Matrix

$$\overline{\mathbf{H}}_m = \begin{pmatrix} \mathbf{H}_m & \\ 0 \dots 0 & h_{m+1,m} \end{pmatrix} \in \mathbb{R}^{(m+1) \times m}$$

mit einer rechten oberen Hessenbergmatrix $\mathbf{H}_m$ dar.

Der Sinn der vorgenommenen äquivalenten Umformung der Minimierungsaufgabe liegt in der speziellen Gestalt der $(m+1) \times m$ Matrix $\overline{\mathbf{H}}_m$. Die erste Formulierung des Minimierungsproblems beinhaltet den Nachteil, daß wir bei einer vorgegebenen Genauigkeit eine sukzessive Berechnung der Folge

$$\mathbf{x}_m = \arg \min_{\mathbf{x} \in \mathbf{x}_0 + K_m} F(\mathbf{x}), \quad m = 1, \dots$$

explizit bis zum Erreichen der Genauigkeitsschranke durchführen müssen. Wie wir im folgenden nachweisen werden, ermöglicht uns die durch (4.3.65) und (4.3.66) gegebene Formulierung des Problems, die Berechnung des minimalen Fehlers

$$\min_{\mathbf{x} \in \mathbf{x}_0 + K_m} F(\mathbf{x})$$

vorzunehmen, ohne $\mathbf{x}_m$ explizit ermitteln zu müssen. Damit haben wir die Möglichkeit, erst dann $\mathbf{x}_m$ zu bestimmen, wenn der zu erwartende Fehler die geforderte Genauigkeit erfüllt.

Lemma 4.75 *Es sei vorausgesetzt, daß der Arnoldi-Algorithmus nicht vor der Berechnung von $\mathbf{v}_{m+1}$ abbricht und die Matrizen $\mathbf{G}_{i+1,i} \in \mathbb{R}^{(m+1) \times (m+1)}$ für $i = 1, \dots, m$ durch*

$$\mathbf{G}_{i+1,i} := \begin{pmatrix} 1 & & & & & \\ & \ddots & & & & \\ & & 1 & & & \\ & & & c_i & s_i & \\ & & & -s_i & c_i & \\ & & & & & 1 & \\ & & & & & & \ddots & \\ & & & & & & & 1 \end{pmatrix}$$

gegeben sind, wobei c_i und s_i gemäß

$$c_i := \frac{a}{\sqrt{a^2 + b^2}} \quad \text{und} \quad s_i := \frac{b}{\sqrt{a^2 + b^2}} \quad (4.3.68)$$

mit

$$a := (\mathbf{G}_{i,i-1} \cdots \mathbf{G}_{3,2} \cdot \mathbf{G}_{2,1} \overline{\mathbf{H}}_m)_{i,i}$$

und

$$b := (\mathbf{G}_{i,i-1} \cdots \mathbf{G}_{3,2} \cdot \mathbf{G}_{2,1} \overline{\mathbf{H}}_m)_{i+1,i}$$

definiert sind. Dann stellt

$$\mathbf{Q}_m := \mathbf{G}_{m+1,m} \cdots \mathbf{G}_{2,1}$$

eine orthogonale Matrix dar, für die

$$\mathbf{Q}_m \overline{\mathbf{H}}_m = \overline{\mathbf{R}}_m$$

mit

$$\overline{\mathbf{R}}_m = \begin{pmatrix} \bar{r}_{11} & \cdots & \cdots & \bar{r}_{1m} \\ 0 & \ddots & & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ \vdots & & \ddots & \bar{r}_{mm} \\ 0 & \cdots & \cdots & 0 \end{pmatrix} =: \begin{pmatrix} \mathbf{R}_m \\ 0 \cdots 0 \end{pmatrix} \in \mathbb{R}^{(m+1) \times m} \quad (4.3.69)$$

gilt und $\mathbf{R}_m$ regulär ist.

Beweis:

Gilt $\mathbf{v}_{m+1} \neq \mathbf{0}$, dann folgt $h_{j+1,j} \neq 0$ für $j = 1, \dots, m$, wodurch alle Spaltenvektoren der Matrix $\overline{\mathbf{H}}_m$ linear unabhängig sind und somit $\text{rang} \overline{\mathbf{H}}_m = m$ gilt. Im Fall $\mathbf{v}_{m+1} = \mathbf{0}$ gilt $h_{m+1,m} = 0$, so daß mit (4.3.35) $\mathbf{A}\mathbf{V}_m = \mathbf{V}_m \mathbf{H}_m$ und wegen $\text{rang}(\mathbf{A}\mathbf{V}_m) = \text{rang}(\mathbf{V}_m \mathbf{H}_m) = m$ die Eigenschaft $\text{rang} \overline{\mathbf{H}}_m = \text{rang} \mathbf{H}_m = m$ folgt.

Wir führen den Beweis mittels einer vollständigen Induktion über i durch.

Für $i = 1$ erhalten wir mit $\text{rang} \overline{\mathbf{H}}_m = m$ direkt

$$h_{11}^2 + h_{21}^2 \neq 0$$

und somit die Wohldefiniertheit der Rotationsmatrix $\mathbf{G}_{2,1}$. Durch elementares Nachrechnen wird leicht ersichtlich, daß eine orthogonale Transformation der Matrix $\overline{\mathbf{H}}_m$ durch $\mathbf{G}_{2,1}$

$$\mathbf{G}_{2,1}\overline{\mathbf{H}}_m = \begin{pmatrix} h_{11}^{(1)} & h_{12}^{(1)} & \dots & \dots & h_{1m}^{(1)} \\ 0 & h_{22}^{(1)} & \dots & \dots & h_{2m}^{(1)} \\ 0 & h_{32} & & & \vdots \\ 0 & 0 & \ddots & & \vdots \\ \vdots & \vdots & & \ddots & \vdots \\ 0 & 0 & \dots & 0 & h_{m+1,m} \end{pmatrix}$$

liefert. Gelte nun für $i < m$

$$\mathbf{G}_{i+1,i} \cdot \dots \cdot \mathbf{G}_{2,1}\overline{\mathbf{H}}_m = \begin{pmatrix} h_{11}^{(i)} & \dots & \dots & \dots & \dots & \dots & h_{1m}^{(i)} \\ 0 & \ddots & & & & & \vdots \\ \vdots & \ddots & \ddots & & & & \vdots \\ \vdots & & 0 & h_{i+1,i+1}^{(i)} & \dots & \dots & h_{i+1,m}^{(i)} \\ \vdots & & \vdots & h_{i+2,i+1} & \dots & \dots & h_{i+2,m} \\ \vdots & & \vdots & 0 & \ddots & & \vdots \\ \vdots & & \vdots & \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & 0 & \dots & 0 & h_{m+1,m} \end{pmatrix}.$$

Da alle Givensrotationen $\mathbf{G}_{j+1,j}$, $j = 1, \dots, i$ orthogonale Drehmatrizen darstellen, folgt

$$\text{rang}(\mathbf{G}_{i+1,i} \cdot \dots \cdot \mathbf{G}_{2,1}\overline{\mathbf{H}}_m) = m,$$

wodurch sich

$$\left(h_{i+1,i+1}^{(i)}\right)^2 + (h_{i+2,i+1})^2 \neq 0$$

ergibt. Somit ist auch die Matrix $\mathbf{G}_{i+2,i+1}$ wohldefiniert und es folgt

$$\mathbf{G}_{i+2,i+1} \cdot \dots \cdot \mathbf{G}_{2,1}\overline{\mathbf{H}}_m = \begin{pmatrix} h_{11}^{(i+1)} & \dots & \dots & \dots & \dots & \dots & h_{1m}^{(i+1)} \\ 0 & \ddots & & & & & \vdots \\ \vdots & \ddots & \ddots & & & & \vdots \\ \vdots & & 0 & h_{i+2,i+2}^{(i+1)} & \dots & \dots & h_{i+2,m}^{(i+1)} \\ \vdots & & \vdots & h_{i+3,i+2} & \dots & \dots & h_{i+3,m} \\ \vdots & & \vdots & 0 & \ddots & & \vdots \\ \vdots & & \vdots & \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & 0 & \dots & 0 & h_{m+1,m} \end{pmatrix}.$$

Die Matrix $\mathbf{Q}_m = \mathbf{G}_{m+1,m} \cdot \dots \cdot \mathbf{G}_{2,1} \in \mathbb{R}^{(m+1) \times (m+1)}$ ist mit Lemma 2.27 orthogonal und es gilt

$$\mathbf{Q}_m \overline{\mathbf{H}}_m = \overline{\mathbf{R}}_m \quad \text{mit} \quad \bar{r}_{ij} = h_{ij}^{(m)}$$

für $i = 1, \dots, m, j = 1, \dots, m$. Aus $\text{rang} \overline{\mathbf{R}}_m = \text{rang}(\mathbf{Q}_m \overline{\mathbf{H}}_m) = \text{rang} \overline{\mathbf{H}}_m = m$ folgt abschließend die Regularität der Matrix $\mathbf{R}_m$. $\square$

Definieren wir mit der durch Lemma 4.75 gegebenen Matrix $\mathbf{Q}_m \in \mathbb{R}^{(m+1) \times (m+1)}$ unter Verwendung von $\mathbf{e}_1 = (1, 0, \dots, 0)^T \in \mathbb{R}^{m+1}$ den Vektor

$$\bar{\mathbf{g}}_m := \|\mathbf{r}_0\|_2 \mathbf{Q}_m \mathbf{e}_1 = \left(\gamma_1^{(m)}, \dots, \gamma_m^{(m)}, \gamma_{m+1} \right)^T = (\mathbf{g}_m^T, \gamma_{m+1})^T \in \mathbb{R}^{m+1}, \quad (4.3.70)$$

dann folgt mit (4.3.67) im Fall $\mathbf{v}_{m+1} \neq \mathbf{0}$ die Darstellung

$$\begin{aligned} \min_{\alpha \in \mathbb{R}^m} J_m(\alpha) &= \min_{\alpha \in \mathbb{R}^m} \|\mathbf{V}_{m+1} (\|\mathbf{r}_0\|_2 \mathbf{e}_1 - \overline{\mathbf{H}}_m \alpha)\|_2 \\ &= \min_{\alpha \in \mathbb{R}^m} \|\|\mathbf{r}_0\|_2 \mathbf{e}_1 - \overline{\mathbf{H}}_m \alpha\|_2 \\ &= \min_{\alpha \in \mathbb{R}^m} \|\mathbf{Q}_m (\|\mathbf{r}_0\|_2 \mathbf{e}_1 - \overline{\mathbf{H}}_m \alpha)\|_2 \\ &\stackrel{\text{Lemma 4.75}}{=} \min_{\alpha \in \mathbb{R}^m} \|\bar{\mathbf{g}}_m - \overline{\mathbf{R}}_m \alpha\|_2 \\ &= \min_{\alpha \in \mathbb{R}^m} \sqrt{|\gamma_{m+1}|^2 + \|\mathbf{g}_m - \mathbf{R}_m \alpha\|_2^2}. \end{aligned} \quad (4.3.71)$$

Durch die Regularität der Matrix $\mathbf{R}_m$ ergibt sich somit

$$\min_{\alpha \in \mathbb{R}^m} J_m(\alpha) = |\gamma_{m+1}|. \quad (4.3.72)$$

Liegt $\mathbf{v}_{m+1} = \mathbf{0}$ vor, so erhalten wir

$$\min_{\alpha \in \mathbb{R}^m} J_m(\alpha) = \min_{\alpha \in \mathbb{R}^m} \|\mathbf{V}_m (\|\mathbf{r}_0\|_2 \mathbf{e}_1 - \mathbf{H}_m \alpha)\|_2 = \min_{\alpha \in \mathbb{R}^m} \|\mathbf{g}_m - \mathbf{R}_m \alpha\|_2 = 0.$$

Bemerkung:

Für die durch (4.3.70) gegebenen Fehlergrößen $\gamma_1, \dots, \gamma_{m+1}$ gilt

$$\|\mathbf{r}_j\|_2 = |\gamma_{j+1}| \leq |\gamma_j| = \|\mathbf{r}_{j-1}\|_2$$

für $j = 1, \dots, m$.

Zusammenfassend erhalten wir den GMRES-Algorithmus in der folgenden Form. Auch zu diesem Verfahren kann dem Anhang A eine mögliche Umsetzung in MATLAB entnommen werden.

GMRES-Algorithmus —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$	
$\mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0$	
Y $\mathbf{r}_0 = \mathbf{0}$ N	
$\mathbf{x} := \mathbf{x}_0$	$\mathbf{v}_1 := \frac{\mathbf{r}_0}{\ \mathbf{r}_0\ _2}, \gamma_1 := \ \mathbf{r}_0\ _2$
	Für $j = 1, \dots, n$
	Für $i = 1, \dots, j$ setze $h_{ij} := (\mathbf{v}_i, \mathbf{A}\mathbf{v}_j)_2$
	$\mathbf{w}_j := \mathbf{A}\mathbf{v}_j - \sum_{i=1}^j h_{ij}\mathbf{v}_i, h_{j+1,j} := \ \mathbf{w}_j\ _2$
	Für $i = 1, \dots, j-1$
	$\begin{pmatrix} h_{ij} \\ h_{i+1,j} \end{pmatrix} := \begin{pmatrix} c_i & s_i \\ -s_i & c_i \end{pmatrix} \begin{pmatrix} h_{ij} \\ h_{i+1,j} \end{pmatrix}$
	$\beta := \sqrt{h_{jj}^2 + h_{j+1,j}^2}; s_j := \frac{h_{j+1,j}}{\beta}$
	$c_j := \frac{h_{jj}}{\beta}; h_{jj} := \beta$
	$\gamma_{j+1} := -s_j\gamma_j; \gamma_j := c_j\gamma_j$
	Y $\gamma_{j+1} = 0$ N
	Für $i = j, \dots, 1$
	$\alpha_i := \frac{1}{h_{ii}} \left(\gamma_i - \sum_{k=i+1}^j h_{ik}\alpha_k \right)$
	$\mathbf{x} := \mathbf{x}_0 + \sum_{i=1}^j \alpha_i \mathbf{v}_i$
	STOP
	$\mathbf{v}_{j+1} := \frac{\mathbf{w}_j}{h_{j+1,j}}$

Untersucht man das Verfahren, so erkennt man, daß die einzige Abbruchmöglichkeit im Verschwinden von $h_{j+1,j}$ liegt. Ein wesentliches Merkmal des Verfahrens besteht darin, daß GMRES an dieser Stelle nicht zusammenbricht, sondern die exakte Lösung liefert. Eine mathematische Beschreibung dieser Eigenschaft liefert uns der folgende Satz.

Satz 4.76 Seien $\mathbf{A} \in \mathbb{R}^{n \times n}$ eine reguläre Matrix sowie $h_{j+1,j}$ und $\mathbf{w}_j$ durch den Arnoldi-Algorithmus gegeben und gelte $j < n$. Dann sind die folgenden Aussagen äquivalent:

(1) Für die Folge der Krylov-Unterräume gilt

$$K_1 \subset K_2 \subset \dots \subset K_j = K_{j+1} = \dots$$

(2) Das GMRES-Verfahren liefert im j -ten Schritt die exakte Lösung.

(3) $\mathbf{w}_j = \mathbf{0} \in \mathbb{R}^n$.

(4) $h_{j+1,j} = 0$.

Beweis:

„(1) $\Leftrightarrow$ (3)“

Wir erhalten die Aussage direkt durch die folgenden äquivalenten Umformulierungen:

$$\begin{aligned} K_{j+1} &= K_j \\ \Leftrightarrow \text{span} \{ \mathbf{r}_0, \dots, \mathbf{A}^j \mathbf{r}_0 \} &= \text{span} \{ \mathbf{r}_0, \dots, \mathbf{A}^{j-1} \mathbf{r}_0 \} \\ \Leftrightarrow \text{span} \{ \mathbf{v}_1, \dots, \mathbf{v}_j, \mathbf{w}_j \} &= \text{span} \{ \mathbf{v}_1, \dots, \mathbf{v}_j \} \\ \Leftrightarrow \mathbf{w}_j &\in \text{span} \{ \mathbf{v}_1, \dots, \mathbf{v}_j \} \\ \Leftrightarrow \mathbf{w}_j = \mathbf{0} \in \mathbb{R}^n, &\text{ da } (\mathbf{v}_i, \mathbf{w}_j)_2 = 0 \text{ für } i = 1, \dots, j \text{ gilt.} \end{aligned}$$

„(3) $\Leftrightarrow$ (4)“

Der Nachweis folgt unmittelbar aus $h_{j+1,j} = \|\mathbf{w}_j\|_2$.

„(2) $\Rightarrow$ (4)“

Sei $\mathbf{x}_j$ die exakte Lösung von $\mathbf{Ax} = \mathbf{b}$, dann folgt mit (4.3.72) $\gamma_{j+1} = 0$. O.B.d.A. setzen wir voraus, daß $\mathbf{x}_{j-1} \neq \mathbf{x}_j$ und somit $\gamma_j \neq 0$ gilt. Wir nutzen die Gleichung

$$\begin{aligned} 0 = \gamma_{j+1} &= \mathbf{e}_{j+1}^T \mathbf{G}_{j+1,j} \begin{pmatrix} \mathbf{Q}_{j-1} & \mathbf{0} \\ \mathbf{0}^T & 1 \end{pmatrix} \|\mathbf{r}_0\|_2 \mathbf{e}_1 \\ &= \mathbf{e}_{j+1}^T \begin{pmatrix} 1 & & & & \\ & \ddots & & & \\ & & 1 & & \\ & & & c_j & s_j \\ & & & -s_j & c_j \end{pmatrix} \begin{pmatrix} \gamma_1^{(j-1)} \\ \vdots \\ \gamma_{j-1}^{(j-1)} \\ \gamma_j \\ 0 \end{pmatrix} \\ &= -s_j \gamma_j \end{aligned}$$

und erhalten mit $\gamma_j \neq 0$ die Aussage $s_j = 0$. Unter Verwendung der Gleichung (4.3.68) ergibt sich hiermit $h_{j+1,j} = 0$.

„(4) $\Rightarrow$ (2)“

Da $h_{j+1,j} = 0$ berechnet wurde, hat kein Abbruch des Arnoldi-Algorithmus vor der Berechnung von $\mathbf{v}_j \neq \mathbf{0}$ stattgefunden. Mit Lemma 4.75 erhalten wir folglich die Existenz einer orthogonalen Matrix

$$\mathbf{Q}_j := \mathbf{G}_{j+1,j} \cdot \dots \cdot \mathbf{G}_{2,1} \in \mathbb{R}^{(j+1) \times (j+1)}$$

mit

$$\overline{\mathbf{R}}_j = \begin{pmatrix} \mathbf{R}_j & \\ 0 & \dots & 0 \end{pmatrix} = \mathbf{Q}_j \overline{\mathbf{H}}_j, \quad (4.3.73)$$

wobei $\mathbf{R}_j$ eine reguläre obere Dreiecksmatrix darstellt. Die Voraussetzung $h_{j+1,j} = 0$ liefert mit (4.3.68) $s_j = 0$, so daß

$$\mathbf{G}_{j+1,j} = \begin{pmatrix} 1 & & & & \\ & \ddots & & & \\ & & 1 & & \\ & & & c_j & \\ & & & & c_j \end{pmatrix} \quad (4.3.74)$$

folgt. Desweiteren stellen die im Arnoldi-Algorithmus ermittelten Vektoren $\mathbf{v}_1, \dots, \mathbf{v}_j$ eine Orthonormalbasis des K_j dar.

Sei $\mathbf{V}_{j+1} = (\mathbf{V}_j, \tilde{\mathbf{v}}_{j+1}) \in \mathbb{R}^{n \times (j+1)}$ mit einem normierten und zu $\mathbf{v}_1, \dots, \mathbf{v}_j$ orthogonalen Vektor $\tilde{\mathbf{v}}_{j+1} \in \mathbb{R}^n$, dann folgt mit $h_{j+1,j} = 0$ durch Satz 4.67 die Gleichung

$$\mathbf{A}\mathbf{V}_j = \mathbf{V}_j \mathbf{H}_j = \mathbf{V}_{j+1} \overline{\mathbf{H}}_j. \quad (4.3.75)$$

Sei $\overline{\mathbf{g}}_j$ gemäß (4.3.70) in der Form

$$\overline{\mathbf{g}}_j = \|\mathbf{r}_0\|_2 \mathbf{Q}_j \mathbf{e}_1 = (\mathbf{g}_j^T, \gamma_{j+1})^T \in \mathbb{R}^{j+1}$$

gegeben, dann definieren wir

$$\boldsymbol{\alpha}_j := \mathbf{R}_j^{-1} \mathbf{g}_j \in \mathbb{R}^j$$

und

$$\mathbf{x}_j := \mathbf{x}_0 + \mathbf{V}_j \boldsymbol{\alpha}_j \in \mathbb{R}^n. \quad (4.3.76)$$

Hiermit folgt unter Verwendung von

$$\begin{aligned} \overline{\mathbf{g}}_j &= \|\mathbf{r}_0\|_2 \mathbf{G}_{j+1,j} \begin{pmatrix} \mathbf{Q}_{j-1} & \mathbf{0} \\ \mathbf{0}^T & 1 \end{pmatrix} \mathbf{e}_1 \\ &= \mathbf{G}_{j+1,j} \begin{pmatrix} \mathbf{g}_{j-1} \\ \gamma_j \\ 0 \end{pmatrix} \\ &\stackrel{(4.3.74)}{=} \begin{pmatrix} \mathbf{g}_j \\ 0 \end{pmatrix} \end{aligned}$$

die Gleichung

$$\begin{aligned}
 \|\mathbf{b} - \mathbf{A}\mathbf{x}_j\|_2 &= \|\mathbf{r}_0 - \mathbf{A}\mathbf{V}_j\boldsymbol{\alpha}_j\|_2 \\
 &= \|\mathbf{V}_{j+1} (\|\mathbf{r}_0\|_2 \mathbf{e}_1 - \overline{\mathbf{H}}_j \boldsymbol{\alpha}_j)\|_2 \\
 &= \|\mathbf{Q}_j (\|\mathbf{r}_0\|_2 \mathbf{e}_1 - \overline{\mathbf{H}}_j \boldsymbol{\alpha}_j)\|_2 \\
 &= \left\| \begin{pmatrix} \mathbf{g}_j \\ 0 \end{pmatrix} - \begin{pmatrix} \mathbf{R}_j \\ 0 \dots 0 \end{pmatrix} \boldsymbol{\alpha}_j \right\|_2 \\
 &= 0.
 \end{aligned}$$

Somit stellt $\mathbf{x}_j$ die exakte Lösung der Gleichung $\mathbf{A}\mathbf{x} = \mathbf{b}$ dar. □

Das GMRES-Verfahren weist zwei grundlegende Nachteile auf. Die erste Problematik liegt im Rechenaufwand zur Bestimmung der Orthonormalbasis, der mit der Dimension des Krylov-Unterraums anwächst. Desweiteren ergibt sich ein hoher Speicherplatzbedarf für die Basisvektoren. Im Extremfall muß auch bei einer schwachbesetzten Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ eine vollbesetzte Matrix $\mathbf{V}_n \in \mathbb{R}^{n \times n}$ abgespeichert werden. Bei praxisrelevanten Problemen übersteigt der Speicherplatzbedarf daher oftmals die vorhandenen Ressourcen. Aus diesem Grund wird oft ein GMRES-Verfahren mit Restart verwendet, bei dem die maximale Krylov-Unterraumdimension beschränkt wird. Weist das Residuum $\|\mathbf{r}_m\|_2$ bei Erreichen dieser Obergrenze nicht eine vorgegebene Genauigkeit $\|\mathbf{r}_m\|_2 \leq \varepsilon > 0$ auf, so wird dennoch die zur Zeit optimale Näherungslösung $\mathbf{x}_m$ bestimmt und als Startvektor innerhalb eines Restarts verwendet. Das folgende Diagramm beschreibt eine GMRES-Version mit Restart und einer maximalen Krylov-Unterraumdimension von m . Das Verfahren wird oftmals als Restarted GMRES(m) bezeichnet. Die theoretisch mögliche Interpretation des GMRES-Verfahrens als direkte Methode geht in dieser Formulierung zwar verloren, aber das Residuum ist dennoch monoton fallend.

Im Rahmen einer impliziten Finite-Volumen-Methode zur Simulation reibungsfreier und reibungsbehafteter Luftströmungen erwies sich bei schwachbesetzten Gleichungssystemen mit etwa 10.000 Unbekannten eine maximale Krylov-Unterraumdimension im Bereich 10 bis 15 als ausreichend, um eine Genauigkeit von $\varepsilon = 10^{-6}$ ohne Restart zu erreichen, falls eine geeignete Präkonditionierung implementiert wurde. Hierbei hat sich gezeigt, daß das Konvergenzverhalten des GMRES-Verfahrens sehr stark vom gewählten Präkonditionierer abhängt. Mit einer unvollständigen LU-Zerlegung (siehe Abschnitt 5.4) erwies sich die Methode als robust, während eine einfache Skalierung (siehe Abschnitt 5.1) analog zum Verfahren ohne Vorkonditionierung häufig auf nicht akzeptable Rechenzeiten führte. Eine ausführliche Diskussion der angesprochenen Ergebnisse kann der Arbeit [48] entnommen werden.

Das Ersetzen der Arnoldi-Methode durch eine unvollständige Orthogonalisierungsvorschrift innerhalb des GMRES-Verfahrens führt auf den sogenannten Quasi-GMRES-Algorithmus (QGMRES). Dieser Ansatz kann auch in einer direkten Version der DQGMRES verwendet werden. Beide Verfahren sind allerdings nicht so stark verbreitet und werden detailliert in [61] beschrieben.

GMRES(m)-Algorithmus mit begrenzter Anzahl von Restarts —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$ und $\varepsilon > 0$	
$restart := 0$; $\mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0$	
Y $\mathbf{r}_0 = \mathbf{0}$ N	
$\mathbf{x} := \mathbf{x}_0$	$\mathbf{v}_1 := \frac{\mathbf{r}_0}{\ \mathbf{r}_0\ _2}, \gamma_1 := \ \mathbf{r}_0\ _2$
	Für $j = 1, \dots, m$
	Für $i = 1, \dots, j$ setze $h_{ij} := (\mathbf{v}_i, \mathbf{A}\mathbf{v}_j)_2$
	$\mathbf{w}_j := \mathbf{A}\mathbf{v}_j - \sum_{i=1}^j h_{ij}\mathbf{v}_i, \quad h_{j+1,j} := \ \mathbf{w}_j\ _2$
	Für $i = 1, \dots, j-1$
	$\begin{pmatrix} h_{ij} \\ h_{i+1,j} \end{pmatrix} := \begin{pmatrix} c_i & s_i \\ -s_i & c_i \end{pmatrix} \begin{pmatrix} h_{ij} \\ h_{i+1,j} \end{pmatrix}$
	$\beta := \sqrt{h_{jj}^2 + h_{j+1,j}^2}; \quad s_j := \frac{h_{j+1,j}}{\beta}; \quad c_j := \frac{h_{jj}}{\beta}$
	$h_{jj} := \beta; \quad \gamma_{j+1} := -s_j\gamma_j; \quad \gamma_j := c_j\gamma_j$
	Y $\gamma_{j+1} \leq \varepsilon$ N
	Für $i = j, \dots, 1$
	$\alpha_i := \frac{1}{h_{ii}} \left(\gamma_i - \sum_{k=i+1}^j h_{ik}\alpha_k \right)$
	$\mathbf{x} := \mathbf{x}_0 + \sum_{i=1}^j \alpha_i \mathbf{v}_i$
	STOP
	$\mathbf{v}_{j+1} := \frac{\mathbf{w}_j}{h_{j+1,j}}$
	Für $i = m, \dots, 1$
	$\alpha_i := \frac{1}{h_{ii}} \left(\gamma_i - \sum_{k=i+1}^m h_{ik}\alpha_k \right)$
	$\mathbf{x} := \mathbf{x}_0 + \sum_{i=1}^m \alpha_i \mathbf{v}_i$
	$\mathbf{x}_0 := \mathbf{x}; \quad \mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0; \quad restarts := restarts + 1$
solange $restart \leq$ Maximale Anzahl von Restarts	

Unter Verwendung von $\mathbf{W}_m = (\mathbf{A}\mathbf{v}_1 \dots \mathbf{A}\mathbf{v}_m)$ erfüllt das GMRES-Verfahren die Voraussetzung des Satzes 4.49. Folglich gilt

$$\|\mathbf{r}_m\| \leq \|\mathbf{P}_m\| \min_{p \in \mathcal{P}_m^1} \|p(\mathbf{A})\mathbf{r}_0\|$$

bezüglich jeder Norm, wobei $\mathbf{P}_m$ die in (4.3.4) aufgeführte Projektion darstellt. Für Projektionen gilt mit $\mathbf{P}_m^2 = \mathbf{P}_m$ stets $\|\mathbf{P}_m\| \geq 1$, wodurch sich für den Spezialfall der euklidischen Norm schon aus der Definition der Funktion F gemäß (4.3.59) für die durch das GMRES-Verfahren ermittelte Näherungslösung $\mathbf{x}_m$ mit

$$\|\mathbf{r}_m\|_2 = \min_{p \in \mathcal{P}_m^1} \|p(\mathbf{A})\mathbf{r}_0\|_2 \quad (4.3.77)$$

eine Verbesserung der Aussage ergibt. Unter zusätzlichen Voraussetzungen an die Matrix ergeben sich mit den anschließenden Aussagen weitere Abschätzungen des Residuums.

Zu einer gegebenen symmetrischen Matrix $\mathbf{B} \in \mathbb{R}^{n \times n}$ bezeichne $\lambda_{\min}(\mathbf{B})$ und $\lambda_{\max}(\mathbf{B})$ im folgenden stets den betragskleinsten beziehungsweise betragsgrößten Eigenwert von $\mathbf{B}$.

Satz 4.77 *Sei $\mathbf{A} \in \mathbb{R}^{n \times n}$ positiv definit und $\mathbf{r}_m$ der im GMRES-Verfahren ermittelte m -te Residuenvektor, dann konvergiert das GMRES-Verfahren, und es gilt*

$$\|\mathbf{r}_m\|_2 \leq \left(1 - \frac{\lambda_{\min}^2\left(\frac{\mathbf{A}^T + \mathbf{A}}{2}\right)}{\lambda_{\max}\left(\mathbf{A}^T \mathbf{A}\right)} \right)^{\frac{m}{2}} \|\mathbf{r}_0\|_2. \quad (4.3.78)$$

Beweis:

Aus (4.3.77) erhalten wir

$$\|\mathbf{r}_m\|_2 \leq \min_{p \in \mathcal{P}_m^1} \|p(\mathbf{A})\|_2 \|\mathbf{r}_0\|_2. \quad (4.3.79)$$

Wir werden nun den Nachweis der Abschätzung (4.3.78) durch eine explizite Angabe eines Polynoms $p \in \mathcal{P}_m^1$ erbringen.

Für $\alpha > 0$ definieren wir $p_1(\mathbf{A}) = \mathbf{I} - \alpha\mathbf{A} \in \mathcal{P}_1^1$, wodurch $(p_1(\mathbf{A}))^m \in \mathcal{P}_m^1$ und

$$\min_{p \in \mathcal{P}_m^1} \|p(\mathbf{A})\|_2 \leq \|(p_1(\mathbf{A}))^m\|_2 \leq \|p_1(\mathbf{A})\|_2^m \quad (4.3.80)$$

gilt. Aus

$$\begin{aligned} \|p_1(\mathbf{A})\|_2^2 &\stackrel{(2.2.2)}{=} \sup_{\|\mathbf{x}\|_2=1} \|(\mathbf{I} - \alpha\mathbf{A})\mathbf{x}\|_2^2 = \sup_{\mathbf{x} \neq \mathbf{0}} \frac{\|(\mathbf{I} - \alpha\mathbf{A})\mathbf{x}\|_2^2}{\|\mathbf{x}\|_2^2} \\ &= \sup_{\mathbf{x} \neq \mathbf{0}} \left\{ 1 - 2\alpha \frac{(\mathbf{x}, \mathbf{A}\mathbf{x})_2}{(\mathbf{x}, \mathbf{x})_2} + \alpha^2 \frac{(\mathbf{A}\mathbf{x}, \mathbf{A}\mathbf{x})_2}{(\mathbf{x}, \mathbf{x})_2} \right\} \end{aligned}$$

ergibt sich für $\mathbf{x} \neq \mathbf{0}$ unter Verwendung von

$$0 < \frac{(\mathbf{A}\mathbf{x}, \mathbf{A}\mathbf{x})_2}{(\mathbf{x}, \mathbf{x})_2} = \frac{(\mathbf{x}, \mathbf{A}^T \mathbf{A}\mathbf{x})_2}{(\mathbf{x}, \mathbf{x})_2} \leq \lambda_{\max}(\mathbf{A}^T \mathbf{A})$$

und

$$\frac{(\mathbf{x}, \mathbf{A}\mathbf{x})_2}{(\mathbf{x}, \mathbf{x})_2} = \frac{\left(\mathbf{x}, \frac{\mathbf{A}^T + \mathbf{A}}{2}\mathbf{x}\right)_2}{(\mathbf{x}, \mathbf{x})_2} \geq \lambda_{\min}\left(\frac{\mathbf{A}^T + \mathbf{A}}{2}\right) \stackrel{\mathbf{A} \text{ pos. def.}}{>} 0$$

die Ungleichung

$$\|p_1(\mathbf{A})\|_2^2 \leq 1 - 2\alpha\lambda_{\min}\left(\frac{\mathbf{A}^T + \mathbf{A}}{2}\right) + \alpha^2\lambda_{\max}(\mathbf{A}^T\mathbf{A}). \quad (4.3.81)$$

Das Minimum der rechten Seite bezüglich α erhalten wir im Punkt

$$\alpha_{\min} = \frac{\lambda_{\min}\left(\frac{\mathbf{A}^T + \mathbf{A}}{2}\right)}{\lambda_{\max}(\mathbf{A}^T\mathbf{A})} > 0.$$

Durch Einsetzen der Größe $\alpha_{\min}$ in die Abschätzung (4.3.81) folgt für $p_1(\mathbf{A}) = \mathbf{I} - \alpha_{\min}\mathbf{A}$ die Ungleichung

$$0 \leq \|p_1(\mathbf{A})\|_2^2 \leq 1 - \frac{\lambda_{\min}^2\left(\frac{\mathbf{A}^T + \mathbf{A}}{2}\right)}{\lambda_{\max}(\mathbf{A}^T\mathbf{A})} < 1. \quad (4.3.82)$$

Somit konvergiert das GMRES-Verfahren, und es gilt

$$\begin{aligned} \|\mathbf{r}_m\|_2 &\stackrel{(4.3.79)}{\leq} \min_{p \in \mathcal{P}_m^1} \|p(\mathbf{A})\|_2 \|\mathbf{r}_0\|_2 \\ &\stackrel{(4.3.80)}{\leq} \|p_1(\mathbf{A})\|_2^m \|\mathbf{r}_0\|_2 \\ &\stackrel{(4.3.82)}{\leq} \left(1 - \frac{\lambda_{\min}^2\left(\frac{\mathbf{A}^T + \mathbf{A}}{2}\right)}{\lambda_{\max}(\mathbf{A}^T\mathbf{A})}\right)^{\frac{m}{2}} \|\mathbf{r}_0\|_2. \end{aligned}$$

□

Für eine symmetrische Matrix ergibt sich hieraus zudem eine Abschätzung unter Verwendung der Konditionszahl der Matrix $\mathbf{A}$, die den Vorteil einer kleineren Konditionszahl verdeutlicht.

Korollar 4.78 Sei $\mathbf{A} \in \mathbb{R}^{n \times n}$ positiv definit und symmetrisch. Zudem sei $\mathbf{r}_m$ der im GMRES-Verfahren ermittelte m -te Residuenvektor, dann konvergiert das GMRES-Verfahren und es gilt

$$\|\mathbf{r}_m\|_2 \leq \left(\frac{\text{cond}_2^2(\mathbf{A}) - 1}{\text{cond}_2^2(\mathbf{A})}\right)^{\frac{m}{2}} \|\mathbf{r}_0\|_2.$$

Zum Abschluß dieses Abschnitts wenden wir uns nochmals dem GMRES(m)-Verfahren zu. Während mit Satz 4.76 das GMRES-Verfahren spätestens nach n Iterationen die exakte Lösung liefert, kann aufgrund des Restarts eine solche Aussage für die GMRES(m)-Methode ($m < n$) nicht nachgewiesen werden. Zwar liegt auch bei der Restarted-Version ein monoton Abfallen des Residuums vor, es bleibt jedoch zu befürchten, daß sich ein konstanter Residuenverlauf einstellt. Wir sind daher an Konvergenzaussagen für das GMRES(m)-Verfahren interessiert, die zudem aus praktischen Gesichtspunkten eine möglichst kleine untere Schranke für die maximale Krylov-Unterraumdimension m fordern sollten. Eine derartige Aufgabenstellung kann in Spezialfällen durch die anschließenden Sätze positiv beantwortet werden.

Satz 4.79 Sei $\mathbf{A} \in \mathbb{R}^{n \times n}$ positiv definit, dann konvergiert das GMRES(m)-Verfahren für $m \geq 1$.

Beweis:

Durch die Ungleichung (4.3.82) ergibt sich mit Satz 4.77 stets

$$\|\mathbf{r}_1\|_2 \leq \gamma \|\mathbf{r}_0\|_2 \quad (4.3.83)$$

mit $\gamma < 1$. Da γ unabhängig vom Startvektor $\mathbf{x}_0$ ist, folgt aus (4.3.83) die Behauptung. $\square$

Satz 4.80 Sei $\mathbf{A} \in \mathbb{R}^{n \times n}$ regulär und symmetrisch, dann konvergiert das GMRES(m)-Verfahren für $m \geq 2$.

Beweis:

Ausgehend von der Abschätzung (4.3.79) wählen wir $p(\mathbf{A}) = \mathbf{I} - \alpha \mathbf{A}^2$ mit $\alpha > 0$. Analog zum Beweis des Satzes 4.77 folgt unter Verwendung der Symmetrie der Matrix $\mathbf{A}$ die Ungleichung

$$\|p(\mathbf{A})\|_2^2 \leq 1 - 2\alpha \lambda_{\min}(\mathbf{A}^2) + \alpha^2 \lambda_{\max}(\mathbf{A}^4).$$

Für den optimalen Parameter ergibt sich

$$\alpha_{\min} = \frac{\lambda_{\min}(\mathbf{A}^2)}{\lambda_{\max}(\mathbf{A}^4)}$$

wodurch mit $p_1(\mathbf{A}) = \mathbf{I} - \alpha_{\min} \mathbf{A}^2$ die Abschätzung

$$\|\mathbf{r}_2\|_2 \leq \underbrace{\left(1 - \frac{\lambda_{\min}^2(\mathbf{A}^2)}{\lambda_{\max}(\mathbf{A}^4)}\right)^{\frac{1}{2}}}_{< 1} \|\mathbf{r}_0\|_2$$

folgt. $\square$

Weitere Konvergenzaussagen für das GMRES-Verfahren wie auch für die Methoden CGN und CGS werden in [52] hergeleitet.

Liegt mit $\mathbf{A}$ eine symmetrische Matrix vor, dann degeneriert der Arnoldi-Algorithmus zum Lanczos-Verfahren (siehe Abschnitt 4.3.2.2) und wir erhalten mit $\mathbf{H}_m$ eine Tridiagonalmatrix. Aufgrund der speziellen Gestalt der Matrix läßt sich das GMRES-Verfahren in der Form einer Dreitermrekursion schreiben, wodurch das angesprochene Speicherplatzproblem entfällt und keine Notwendigkeit zur Betrachtung einer Version mit Restart vorliegt. Diese spezielle Form des GMRES-Algorithmus wurde bereits 1975 von Paige und Saunders [54] vorgestellt und trägt den Namen MINRES (Minimal Residual). Eine Herleitung dieses Verfahrens wird in [32] präsentiert.

4.3.2.5 Das BiCG-Verfahren

Das BiCG-Verfahren (Bi-Conjugate Gradient) wurde ursprünglich bereits 1952 von Lanczos [45] vorgeschlagen. Seine große Verbreitung erzielte er jedoch erst in der von Fletcher [25] präsentierten Formulierung. Das Verfahren stellt eine auf dem Bi-Lanczos-Algorithmus basierende Krylov-Unterraum-Methode dar, wobei die Orthogonalität durch eine Petrov-Galerkin-Bedingung $L_m = K_m^T = \text{span}\{\mathbf{r}_0, \mathbf{A}^T \mathbf{r}_0, \dots, (\mathbf{A}^T)^{m-1} \mathbf{r}_0\}$ gegeben ist. Im Vergleich zur GMRES-Methode zeichnet sich der Algorithmus durch einen deutlich geringeren Speicherplatzbedarf aus. Jedoch weist das Verfahren zwei entscheidende Nachteile auf. Vererbt durch den zugrundeliegenden Bi-Lanczos-Algorithmus werden Multiplikationen mit der zu $\mathbf{A}$ transponierten Matrix benötigt, und das Verfahren kann, wie in Abschnitt 4.3.2.3 beschrieben, bei einer regulären Matrix abbrechen, ohne die exakte Lösung berechnet zu haben.

Satz 4.81 *Vorausgesetzt, der Bi-Lanczos-Algorithmus bricht nicht vor der Berechnung von $\mathbf{v}_m \neq \mathbf{0}$ ab, und die in Satz 4.72 definierte Tridiagonalmatrix $\mathbf{T}_m$ ist regulär, dann stellt*

$$\mathbf{x}_m = \mathbf{x}_0 + \mathbf{V}_m \mathbf{T}_m^{-1} (\|\mathbf{r}_0\|_2 \mathbf{e}_1) \in \mathbf{x}_0 + K_m \quad (4.3.84)$$

mit $\mathbf{e}_1 = (1, 0, \dots, 0)^T \in \mathbb{R}^m$ die eindeutig bestimmte Lösung der auf der Petrov-Galerkin-Bedingung

$$\mathbf{b} - \mathbf{A}\mathbf{x}_m \perp K_m^T \quad (4.3.85)$$

basierenden Krylov-Unterraum-Methode dar.

Beweis:

Da die Vektoren der Matrix $\mathbf{V}_m$ laut Satz 4.71 eine Basis des K_m darstellen, folgt direkt

$$\mathbf{x}_m \in \mathbf{x}_0 + K_m.$$

Laut Satz 4.71 und Satz 4.72 gilt zudem

$$\mathbf{T}_m = \mathbf{W}_m^T \mathbf{A} \mathbf{V}_m \quad (4.3.86)$$

sowie

$$\mathbf{W}_m^T \mathbf{V}_m = \mathbf{V}_m^T \mathbf{W}_m = \mathbf{I}.$$

Somit erhalten wir für $i = 1, \dots, m$

$$\begin{aligned} (\mathbf{b} - \mathbf{A}\mathbf{x}_m, \mathbf{w}_i)_2 &= (\mathbf{r}_0 - \mathbf{A}\mathbf{V}_m \mathbf{T}_m^{-1} (\|\mathbf{r}_0\|_2 \mathbf{e}_1), \mathbf{w}_i)_2 \\ &= \|\mathbf{r}_0\|_2 \left[\underbrace{\left(\frac{\mathbf{r}_0}{\|\mathbf{r}_0\|_2}, \mathbf{w}_i \right)_2}_{=\delta_{i1}} - (\mathbf{A}\mathbf{V}_m \mathbf{T}_m^{-1} \mathbf{e}_1, \mathbf{w}_i)_2 \right] \\ &= \|\mathbf{r}_0\|_2 \left[\delta_{i1} - \underbrace{\mathbf{w}_i^T \mathbf{A} \mathbf{V}_m \mathbf{T}_m^{-1} \mathbf{e}_1}_{\substack{(4.3.86) \\ = \\ \mathbf{e}_i^T}} \right] \\ &= 0, \end{aligned}$$

wodurch $\mathbf{b} - \mathbf{A}\mathbf{x}_m \perp \mathbf{w}_i$ für $i = 1, \dots, m$ folgt, und daher

$$\mathbf{b} - \mathbf{A}\mathbf{x}_m \perp K_m^T = \text{span}\{\mathbf{w}_1, \dots, \mathbf{w}_m\}$$

gilt.

Wir kommen nun zum Nachweis der Eindeutigkeit der Lösung. Seien $\mathbf{x}_m, \tilde{\mathbf{x}}_m \in \mathbf{x}_0 + K_m$ zwei Vektoren, die der Bedingung (4.3.85) genügen, dann folgt

$$\mathbf{x}_m = \mathbf{x}_0 + \mathbf{V}_m \boldsymbol{\alpha}_m \quad \text{und} \quad \tilde{\mathbf{x}}_m = \mathbf{x}_0 + \mathbf{V}_m \tilde{\boldsymbol{\alpha}}_m$$

mit $\boldsymbol{\alpha}_m, \tilde{\boldsymbol{\alpha}}_m \in \mathbb{R}^m$. Aus der Orthogonalitätsbedingung ergibt sich

$$\begin{aligned} \mathbf{0} &= \mathbf{W}_m^T (\mathbf{b} - \mathbf{A}\tilde{\mathbf{x}}_m - \mathbf{b} + \mathbf{A}\mathbf{x}_m) \\ &= \mathbf{W}_m^T \mathbf{A} (\mathbf{x}_m - \tilde{\mathbf{x}}_m) \\ &= \mathbf{W}_m^T \mathbf{A} \mathbf{V}_m (\boldsymbol{\alpha}_m - \tilde{\boldsymbol{\alpha}}_m) \\ &= \mathbf{T}_m (\boldsymbol{\alpha}_m - \tilde{\boldsymbol{\alpha}}_m). \end{aligned}$$

Laut Voraussetzung ist $\mathbf{T}_m \in \mathbb{R}^{m \times m}$ regulär, und wir erhalten somit $\boldsymbol{\alpha}_m = \tilde{\boldsymbol{\alpha}}_m$ und folglich $\mathbf{x}_m = \tilde{\mathbf{x}}_m$. $\square$

Zur Herleitung des Verfahrens setzen wir voraus, daß sich die Matrix $\mathbf{T}_m$ in der Form

$$\mathbf{T}_m = \underbrace{\begin{pmatrix} 1 & & & & \\ \ell_{21} & \ddots & & & \\ & \ddots & \ddots & & \\ & & \ddots & \ddots & \\ & & & \ell_{m,m-1} & 1 \end{pmatrix}}_{=: \mathbf{L}_m} \underbrace{\begin{pmatrix} u_{11} & h_{12} & & & \\ & \ddots & \ddots & & \\ & & \ddots & h_{m-1,m} & \\ & & & u_{mm} & \end{pmatrix}}_{=: \mathbf{U}_m}$$

zerlegen läßt. Mit

$$\tilde{\mathbf{P}}_m = (\tilde{\mathbf{p}}_0 \dots \tilde{\mathbf{p}}_{m-1}) := \mathbf{V}_m \mathbf{U}_m^{-1} \in \mathbb{R}^{n \times m}$$

und

$$\mathbf{z}_m = (\xi_1, \dots, \xi_m)^T := \mathbf{L}_m^{-1} (\|\mathbf{r}_0\|_2 \mathbf{e}_1) \in \mathbb{R}^m$$

folgt

$$\begin{aligned} \mathbf{x}_m &= \mathbf{x}_0 + \mathbf{V}_m \mathbf{T}_m^{-1} (\|\mathbf{r}_0\|_2 \mathbf{e}_1) = \mathbf{x}_0 + \mathbf{V}_m \mathbf{U}_m^{-1} \mathbf{L}_m^{-1} (\|\mathbf{r}_0\|_2 \mathbf{e}_1) \\ &= \mathbf{x}_0 + \tilde{\mathbf{P}}_m \mathbf{z}_m \end{aligned} \tag{4.3.87}$$

mit dem ersten Einheitsvektor $\mathbf{e}_1 \in \mathbb{R}^m$. Wegen $\mathbf{V}_m = \tilde{\mathbf{P}}_m \mathbf{U}_m$ gilt unter Berücksichtigung der Festlegung $\tilde{\mathbf{p}}_{-1} := \mathbf{0}$ respektive $h_{0,1} := 0$ für den m -ten Spaltenvektor der Matrix $\mathbf{V}_m$ die Darstellung $\mathbf{v}_m = h_{m-1,m} \tilde{\mathbf{p}}_{m-2} + u_{mm} \tilde{\mathbf{p}}_{m-1}$ und somit

$$\tilde{\mathbf{p}}_{m-1} = \frac{1}{u_{mm}} (\mathbf{v}_m - h_{m-1,m} \tilde{\mathbf{p}}_{m-2}). \tag{4.3.88}$$

Zudem erhalten wir aus $\mathbf{L}_m \mathbf{z}_m = \|\mathbf{r}_0\|_2 \mathbf{e}_1$ für $m > 1$ die Gleichung

$$\xi_m + \ell_{m,m-1} \xi_{m-1} = 0. \tag{4.3.89}$$

Zudem gilt $\xi_1 = \|\mathbf{r}_0\|_2$, und aus (4.3.87) ergibt sich daher

$$\begin{aligned}\mathbf{x}_m &= \mathbf{x}_0 + \tilde{\mathbf{P}}_m \mathbf{z}_m = \mathbf{x}_0 + \tilde{\mathbf{P}}_{m-1} \mathbf{z}_{m-1} + \tilde{\mathbf{p}}_{m-1} \xi_m \\ &= \mathbf{x}_{m-1} + \xi_m \tilde{\mathbf{p}}_{m-1}\end{aligned}\tag{4.3.90}$$

und

$$\mathbf{r}_m = \mathbf{b} - \mathbf{A}\mathbf{x}_m = \mathbf{r}_{m-1} - \xi_m \mathbf{A}\tilde{\mathbf{p}}_{m-1}.\tag{4.3.91}$$

Für den Residuenvektor erhalten wir somit die Darstellung

$$\begin{aligned}\mathbf{r}_m &= \mathbf{b} - \mathbf{A}\mathbf{x}_m \\ &\stackrel{(4.3.84)}{=} \mathbf{r}_0 - \mathbf{A}\mathbf{V}_m \mathbf{T}_m^{-1} (\|\mathbf{r}_0\|_2 \mathbf{e}_1) \\ &\stackrel{\text{Satz 4.72}}{=} \mathbf{r}_0 - (\mathbf{V}_m \mathbf{T}_m + h_{m+1,m} \mathbf{v}_{m+1} \mathbf{e}_m^T) \mathbf{T}_m^{-1} (\|\mathbf{r}_0\|_2 \mathbf{e}_1) \\ &= \underbrace{\mathbf{r}_0 - \mathbf{V}_m \|\mathbf{r}_0\|_2 \mathbf{e}_1}_{= 0} - h_{m+1,m} \mathbf{v}_{m+1} \underbrace{\mathbf{e}_m^T \mathbf{T}_m^{-1} (\|\mathbf{r}_0\|_2 \mathbf{e}_1)}_{\in \mathbb{R}} \\ &= \sigma_m \mathbf{v}_{m+1}\end{aligned}\tag{4.3.92}$$

mit $\sigma_m \in \mathbb{R}$. Der Vektor $\tilde{\mathbf{p}}_{m-1}$ läßt sich folglich unter Verwendung der Gleichung (4.3.88) in der Form

$$\tilde{\mathbf{p}}_{m-1} = \frac{1}{u_{mm}} \left(\frac{1}{\sigma_{m-1}} \mathbf{r}_{m-1} - h_{m-1,m} \tilde{\mathbf{p}}_{m-2} \right).$$

schreiben. Analog erhalten wir für das transponierte Problem

$$\mathbf{b} - \mathbf{A}^T \mathbf{x}_m^* \perp K_m$$

mit $\mathbf{x}_m^* = \mathbf{x}_0^* + K_m^T$ und $\mathbf{x}_0^* = \mathbf{A}^{-T} \mathbf{A} \mathbf{x}_0$ bei vorausgesetzter Regularität der Matrix $\mathbf{T}_m$ die eindeutig bestimmte Lösung in der Form

$$\mathbf{x}_m^* = \mathbf{x}_0^* + \mathbf{W}_m \mathbf{T}_m^{-T} (\|\mathbf{r}_0\|_2 \mathbf{e}_1)$$

mit $\mathbf{W}_m$ gemäß Satz 4.72. Für den Residuenvektor $\mathbf{r}_m^* = \mathbf{b} - \mathbf{A}^T \mathbf{x}_m^*$ ergibt sich entsprechend der Gleichung (4.3.92) die Darstellung

$$\mathbf{r}_m^* = \sigma_m^* \mathbf{w}_{m+1}$$

mit $\sigma_m^* \in \mathbb{R}$. Definieren wir

$$\tilde{\mathbf{P}}_m^* = (\tilde{\mathbf{p}}_0^* \cdots \tilde{\mathbf{p}}_{m-1}^*) := \mathbf{W}_m \mathbf{L}_m^{-T},$$

so gilt

$$\tilde{\mathbf{p}}_{m-1}^* = \mathbf{w}_m - \ell_{m,m-1} \tilde{\mathbf{p}}_{m-2}^* = \sigma_{m-1}^* \mathbf{r}_{m-1}^* - \ell_{m,m-1} \tilde{\mathbf{p}}_{m-2}^*,$$

wobei analog zur Bestimmung der Vektoren $\tilde{\mathbf{p}}_{m-1}$ laut (4.3.88) die Festlegung $\tilde{\mathbf{p}}_{-1}^* := \mathbf{0}$ respektive $\ell_{1,0} := 0$ berücksichtigt werden muß. Die Näherungslösung läßt sich in der Form

$$\mathbf{x}_m^* = \mathbf{x}_0^* + \tilde{\mathbf{P}}_m^* \mathbf{z}_m^*$$

mit

$$\mathbf{z}_m^* = (\xi_1^*, \dots, \xi_m^*)^T := \mathbf{U}_m^{-T} (\|\mathbf{r}_0\|_2 \mathbf{e}_1)$$

schreiben. Analog zu (4.3.89) gilt für $m > 1$

$$\xi_m^* = -\frac{h_{m-1,m}}{u_{mm}} \xi_{m-1}^*$$

und

$$\xi_1 = -\frac{\|\mathbf{r}_0\|_2}{u_{11}}.$$

Daher ergibt sich

$$\mathbf{x}_m^* = \mathbf{x}_{m-1}^* + \xi_m^* \tilde{\mathbf{p}}_{m-1}^*$$

sowie

$$\mathbf{r}_m^* = \mathbf{r}_{m-1}^* - \xi_m^* \mathbf{A}^T \tilde{\mathbf{p}}_{m-1}^*.$$

Skalieren wir die Suchvektoren gemäß

$$\mathbf{p}_{m-1} = \frac{\xi_m}{\alpha_{m-1}} \tilde{\mathbf{p}}_{m-1} \quad \text{und} \quad \mathbf{p}_{m-1}^* = \frac{\xi_m^*}{\alpha_{m-1}} \tilde{\mathbf{p}}_{m-1}^*,$$

dann folgt

$$\mathbf{x}_m = \mathbf{x}_{m-1} + \alpha_{m-1} \mathbf{p}_{m-1},$$

$$\mathbf{x}_m^* = \mathbf{x}_{m-1}^* + \alpha_{m-1} \mathbf{p}_{m-1}^*$$

und

$$\mathbf{p}_m = \tau_m \mathbf{r}_m + \beta_{m-1} \mathbf{p}_{m-1},$$

$$\mathbf{p}_m^* = \tau_m^* \mathbf{r}_m^* + \beta_{m-1}^* \mathbf{p}_{m-1}^*$$

mit $\tau_m, \tau_m^*, \beta_{m-1}, \beta_{m-1}^* \in \mathbb{R}$. Hierbei bleibt die Berechnung der skalaren Größen anzugeben. Es erweist sich als übersichtlich, zunächst den BiCG-Algorithmus in der folgenden Form aufzuführen und anschließend nachzuweisen, daß in der gewählten Formulierung mit $\mathbf{x}_{j+1}$ für $j \in \mathbb{N}_0$ stets die Lösung der auf der Petrov-Galerkin-Bedingung (4.3.85) basierenden Krylov-Unterraum-Methode vorliegt. Ein direkt ausführbares Programm wird im Anhang A dargestellt.

BiCG-Algorithmus —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$ und $\varepsilon > 0$
$\mathbf{r}_0 = \mathbf{r}_0^* = \mathbf{p}_0 = \mathbf{p}_0^* := \mathbf{b} - \mathbf{A} \mathbf{x}_0$
$j := 0$
Solange $\ \mathbf{r}_j\ _2 > \varepsilon$
$\alpha_j := \frac{(\mathbf{r}_j, \mathbf{r}_j^*)_2}{(\mathbf{A} \mathbf{p}_j, \mathbf{p}_j^*)_2}$
$\mathbf{x}_{j+1} := \mathbf{x}_j + \alpha_j \mathbf{p}_j$
$\mathbf{r}_{j+1} := \mathbf{r}_j - \alpha_j \mathbf{A} \mathbf{p}_j$
$\mathbf{r}_{j+1}^* := \mathbf{r}_j^* - \alpha_j \mathbf{A}^T \mathbf{p}_j^*$
$\beta_j := \frac{(\mathbf{r}_{j+1}, \mathbf{r}_{j+1}^*)_2}{(\mathbf{r}_j, \mathbf{r}_j^*)_2}$
$\mathbf{p}_{j+1} := \mathbf{r}_{j+1} + \beta_j \mathbf{p}_j$
$\mathbf{p}_{j+1}^* := \mathbf{r}_{j+1}^* + \beta_j \mathbf{p}_j^*$
$j := j + 1$

Lemma 4.82 *Vorausgesetzt, der BiCG-Algorithmus bricht nicht vor der Berechnung von $\mathbf{p}_m^*$ ab, dann gilt für $j = 0, \dots, m$*

$$\mathbf{r}_j = \sum_{i=1}^{j+1} \gamma_i \mathbf{v}_i, \quad \mathbf{r}_j^* = \sum_{i=1}^{j+1} \gamma_i^* \mathbf{w}_i, \quad (4.3.93)$$

$$\mathbf{p}_j = \sum_{i=1}^{j+1} \lambda_i \mathbf{v}_i, \quad \mathbf{p}_j^* = \sum_{i=1}^{j+1} \lambda_i^* \mathbf{w}_i \quad (4.3.94)$$

mit $\gamma_{j+1}, \gamma_{j+1}^*, \lambda_{j+1}, \lambda_{j+1}^* \neq 0$ und den aus dem Bi-Lanczos-Algorithmus resultierenden bi-orthonormalen Vektoren $\mathbf{v}_1, \dots, \mathbf{v}_{j+1}$ und $\mathbf{w}_1, \dots, \mathbf{w}_{j+1}$.

Beweis:

Der Nachweis ergibt sich durch eine Induktion über j . Für $j = 0$ gilt $\mathbf{v}_1 = \mathbf{w}_1 = \mathbf{r}_0 = \mathbf{p}_0 = \mathbf{r}_0^* = \mathbf{p}_0^*$, so daß die Darstellung mit $\gamma_1 = \gamma_1^* = \lambda_1 = \lambda_1^* = 1$ erfüllt ist.

Vorausgesetzt, die Behauptung sei für $\ell = 0, \dots, j < m$ erfüllt, dann folgt aus dem BiCG-Verfahren

$$\begin{aligned} \mathbf{r}_{j+1} &= \mathbf{r}_j - \alpha_j \mathbf{A} \mathbf{p}_j \\ &\stackrel{(4.3.93)}{=} -\alpha_j \lambda_{j+1} \underbrace{\mathbf{A} \mathbf{v}_{j+1}}_{\in K_{j+2}} + \underbrace{\sum_{i=1}^{j+1} \gamma_i \mathbf{v}_i - \alpha_j \sum_{i=1}^j \lambda_i \mathbf{A} \mathbf{v}_i}_{\in K_{j+1}}. \end{aligned}$$

Laut Voraussetzung konnte $\mathbf{p}_{j+1}$ berechnet werden, wodurch $\alpha_j \neq 0$ folgt und mit $\lambda_{j+1} \neq 0$ sich der Residuenvektor in der Form

$$\mathbf{r}_{j+1} = \gamma_{j+2} \mathbf{v}_{j+2} + \underbrace{\psi_{j+1}}_{\in K_{j+1}} \in K_{j+2}$$

mit $\gamma_{j+2} \neq 0$ schreiben läßt. Für den Suchvektor ergibt sich entsprechend

$$\mathbf{p}_{j+1} = \mathbf{r}_{j+1} + \underbrace{\beta_j \mathbf{p}_j}_{\in K_{j+1}} \stackrel{(4.3.93)}{=} \lambda_{j+2} \mathbf{v}_{j+2} + \underbrace{\phi_{j+1}}_{\in K_{j+1}} \in K_{j+2}.$$

mit $\lambda_{j+2} \neq 0$. Die Aussage für die verbleibenden zwei Vektoren folgt analog. $\square$

Lemma 4.83 *Vorausgesetzt, der BiCG-Algorithmus bricht nicht vor der Berechnung von $\mathbf{p}_{m+1}^*$ ab, dann gilt*

$$(\mathbf{r}_j, \mathbf{r}_i^*)_2 = 0 \quad \text{für } i \neq j \leq m+1, \quad (4.3.95)$$

$$(\mathbf{A} \mathbf{p}_j, \mathbf{p}_i^*)_2 = 0 \quad \text{für } i \neq j \leq m+1. \quad (4.3.96)$$

Beweis:

Wir führen den Beweis mittels einer vollständigen Induktion über m durch. Für $m = 0$ ergibt sich die Behauptung gemäß

$$(\mathbf{r}_1, \mathbf{r}_0^*)_2 = (\mathbf{r}_0, \mathbf{r}_0^*)_2 - \frac{(\mathbf{r}_0, \mathbf{r}_0^*)_2}{(\mathbf{A} \mathbf{r}_0, \mathbf{r}_0^*)_2} (\mathbf{A} \mathbf{r}_0, \mathbf{r}_0^*)_2 = 0 \quad (4.3.97)$$

und

$$\begin{aligned} (\mathbf{A} \mathbf{p}_1, \mathbf{p}_0^*)_2 &= (\mathbf{A} \mathbf{r}_1, \mathbf{p}_0^*)_2 + \beta_0 (\mathbf{A} \mathbf{p}_0, \mathbf{p}_0^*)_2 \\ &= \left(\mathbf{r}_1, \mathbf{A}^T \mathbf{p}_0^* \right)_2 + \frac{\beta_0}{\alpha_0} (\mathbf{r}_0, \mathbf{r}_0^*)_2 \\ &= \left(\mathbf{r}_1, \frac{\mathbf{r}_0^* - \mathbf{r}_1^*}{\alpha_0} \right)_2 + \frac{(\mathbf{r}_1, \mathbf{r}_1^*)_2}{\alpha_0} \\ &= 0. \end{aligned} \quad (4.3.98)$$

Analog ergibt sich die behauptete Aussage für $(\mathbf{r}_0, \mathbf{r}_1^*)_2$ sowie $(\mathbf{A}\mathbf{p}_0, \mathbf{p}_1^*)_2$.

Sei die Behauptung für $\ell = 0, \dots, m$ erfüllt, dann folgt der Nachweis in den folgenden Schritten. Für $j < m$ erhalten wir unter Berücksichtigung von $\beta_{-1} := 0$ respektive $\mathbf{p}_{-1}^* := \mathbf{0}$ die Gleichungen

$$\begin{aligned}
 (\mathbf{r}_{m+1}, \mathbf{r}_j^*)_2 &= \underbrace{(\mathbf{r}_m, \mathbf{r}_j^*)_2}_{=0} - \alpha_m (\mathbf{A}\mathbf{p}_m, \mathbf{r}_j^*)_2 \\
 &= -\alpha_m (\mathbf{A}\mathbf{p}_m, \mathbf{p}_j^* - \beta_{j-1}\mathbf{p}_{j-1}^*)_2 \\
 &\stackrel{(4.3.96)}{=} 0
 \end{aligned} \tag{4.3.99}$$

und

$$\begin{aligned}
 (\mathbf{r}_{m+1}, \mathbf{r}_m^*)_2 &= (\mathbf{r}_m, \mathbf{r}_m^*)_2 - \alpha_m (\mathbf{A}\mathbf{p}_m, \mathbf{r}_m^*)_2 \\
 &= (\mathbf{r}_m, \mathbf{r}_m^*)_2 - \alpha_m (\mathbf{A}\mathbf{p}_m, \mathbf{p}_m^*)_2 + \alpha_m \underbrace{(\mathbf{A}\mathbf{p}_m, \beta_{m-1}\mathbf{p}_{m-1}^*)_2}_{=0} \\
 &= (\mathbf{r}_m, \mathbf{r}_m^*)_2 - \frac{(\mathbf{r}_m, \mathbf{r}_m^*)_2}{(\mathbf{A}\mathbf{p}_m, \mathbf{p}_m^*)_2} (\mathbf{A}\mathbf{p}_m, \mathbf{p}_m^*)_2 \\
 &= 0.
 \end{aligned} \tag{4.3.100}$$

Desweiteren ergibt sich für $j < m$

$$\begin{aligned}
 (\mathbf{A}\mathbf{p}_{m+1}, \mathbf{p}_j^*)_2 &= (\mathbf{A}\mathbf{r}_{m+1}, \mathbf{p}_j^*)_2 + \beta_m \underbrace{(\mathbf{A}\mathbf{p}_m, \mathbf{p}_j^*)_2}_{=0} = (\mathbf{r}_{m+1}, \mathbf{A}^T \mathbf{p}_j^*)_2 \\
 &= \left(\mathbf{r}_{m+1}, \frac{\mathbf{r}_j^* - \mathbf{r}_{j+1}^*}{\alpha_j} \right)_2 \\
 &\stackrel{(4.3.99)}{=} \stackrel{(4.3.100)}{=} 0
 \end{aligned}$$

und

$$\begin{aligned}
 (\mathbf{A}\mathbf{p}_{m+1}, \mathbf{p}_m^*)_2 &= (\mathbf{A}\mathbf{r}_{m+1}, \mathbf{p}_m^*)_2 + \beta_m (\mathbf{A}\mathbf{p}_m, \mathbf{p}_m^*)_2 \\
 &= \left(\mathbf{r}_{m+1}, \mathbf{A}^T \mathbf{p}_m^* \right)_2 + \frac{\beta_m}{\alpha_m} (\mathbf{r}_m, \mathbf{r}_m^*)_2 \\
 &= \left(\mathbf{r}_{m+1}, \frac{\mathbf{r}_m^* - \mathbf{r}_{m+1}^*}{\alpha_m} \right)_2 + \frac{(\mathbf{r}_{m+1}, \mathbf{r}_{m+1}^*)_2}{\alpha_m} \\
 &= 0.
 \end{aligned}$$

Eine analoge Vorgehensweise liefert die verbleibenden Gleichungen $(\mathbf{r}_j, \mathbf{r}_{m+1}^*)_2 = 0$ und $(\mathbf{A}\mathbf{p}_j, \mathbf{p}_{m+1}^*)_2 = 0$ für $j < m+1$. $\square$

Satz 4.84 *Unter der Voraussetzung, daß der BiCG-Algorithmus nicht vor der Berechnung von $\mathbf{r}_m^*$ abbricht, stellt $\mathbf{x}_m$ die Lösung des Problems*

$$\mathbf{b} - \mathbf{A}\mathbf{x} \perp K_m^T$$

mit $\mathbf{x} \in \mathbf{x}_0 + K_m$ dar.

Beweis:

Das Einsetzen von (4.3.93) in (4.3.95) liefert

$$0 \stackrel{(4.3.95)}{=} \mathbf{r}_m^T \mathbf{r}_i^* \stackrel{(4.3.93)}{=} \sum_{j=1}^{m+1} \gamma_j \mathbf{v}_j^T \mathbf{r}_i^* \quad \text{für } i = 0, \dots, m-1.$$

Startend von $i = 0$ erhalten wir mittels der Bi-Orthogonalitätsbedingung (4.3.47) und der Darstellung (4.3.93) für $\mathbf{r}_i^*$ sukzessive die Gleichung

$$\gamma_1 = \dots = \gamma_m = 0,$$

wodurch

$$\mathbf{b} - \mathbf{A}\mathbf{x}_m = \mathbf{r}_m = \gamma_{m+1} \mathbf{v}_{m+1} \perp K_m^T$$

gilt und folglich mit $\mathbf{x}_m$ die Lösung der Krylov-Unterraum-Methode vorliegt. $\square$

Bemerkung:

Der im BiCG-Verfahren implizit vorliegende Vektor $\mathbf{x}_m^* = \mathbf{A}^{-T}(\mathbf{b} - \mathbf{r}_m^*)$ stellt analog zur Näherungslösung $\mathbf{x}_m$ die Lösung des transponierten Problems

$$\mathbf{b} - \mathbf{A}^T \mathbf{x} \perp K_m$$

mit $\mathbf{x} \in \mathbf{x}_0^* + K_m^T$, $\mathbf{x}_0^* = \mathbf{A}^{-T} \mathbf{A} \mathbf{x}_0$ dar.

Der BiCG-Algorithmus weist im wesentlichen drei Nachteile auf. Zunächst können Residuenvektoren $\mathbf{r}_j \neq 0$ und $\mathbf{r}_j^* \neq 0$ mit $(\mathbf{r}_j, \mathbf{r}_j^*) = 0$ auftreten, die zu einem Verfahrensabbruch führen, obwohl die Lösung der betrachteten Gleichung nicht berechnet wurde. Desweiteren werden bei jeder Iteration Matrix-Vektor-Multiplikationen mit $\mathbf{A}$ und $\mathbf{A}^T$ benötigt, und die BiCG-Iterierten sind im Gegensatz zu den GMRES-Iterierten in der Regel durch keine Minimierungseigenschaft charakterisiert, wodurch sich Oszillationen im Konvergenzverhalten zeigen können [26, 28, 32, 41, 74]. Im Fall einer symmetrischen, positiv definiten Matrix $\mathbf{A}$ stimmt der Bi-Lanczos-Algorithmus mit der Lanczos-Methode und entsprechend das BiCG-Verfahren mit dem CG-Verfahren überein.

4.3.2.6 Das CGS-Verfahren

Das CGS-Verfahren (Conjugate Gradient Squared) wurde 1989 von Sonneveld [67] präsentiert und stellt eine Weiterentwicklung des BiCG-Verfahrens dar, bei dem die Berechnung der skalaren Größen α_j und β_j ohne Verwendung des Residuenvektors $\mathbf{r}_m^*$ und des Suchvektors $\mathbf{p}_m^*$ möglich ist, wodurch der Zugriff auf $\mathbf{A}^T$ entfällt.

Wir betrachten $\mathcal{P}_j$ als die Menge der Polynome vom Höchstgrad j und schreiben

$$\begin{aligned} \mathbf{r}_j &= \varphi_j(\mathbf{A})\mathbf{r}_0, & \mathbf{r}_j^* &= \varphi_j(\mathbf{A}^T)\mathbf{r}_0 \\ \mathbf{p}_j &= \psi_j(\mathbf{A})\mathbf{r}_0, & \mathbf{p}_j^* &= \psi_j(\mathbf{A}^T)\mathbf{r}_0 \end{aligned}$$

mit $\varphi_j, \psi_j \in \mathcal{P}_j$. Auf der Basis des BiCG-Algorithmus sind in unserem Fall die Polynome φ_j und ψ_j rekursiv durch

$$\psi_j(\lambda) = \varphi_j(\lambda) + \beta_{j-1}\psi_{j-1}(\lambda) \quad (4.3.101)$$

und

$$\varphi_{j+1}(\lambda) = \varphi_j(\lambda) - \alpha_j \lambda \psi_j(\lambda) \quad (4.3.102)$$

mit $\varphi_0(\lambda) = \psi_0(\lambda) = 1$ gegeben. Die grundlegende Idee des CGS-Verfahrens liegt in der Nutzung der Eigenschaft

$$\left(\phi_j(\mathbf{A})\mathbf{r}_0, \phi_j(\mathbf{A}^T)\mathbf{r}_0 \right)_2 = \left(\phi_j^2(\mathbf{A})\mathbf{r}_0, \mathbf{r}_0 \right)_2,$$

die für alle Polynome $\phi_j \in \mathcal{P}_j$ gültig ist. Definieren wir

$$\hat{\mathbf{r}}_j := \varphi_j^2(\mathbf{A})\mathbf{r}_0 \quad \text{und} \quad \hat{\mathbf{p}}_j := \psi_j^2(\mathbf{A})\mathbf{r}_0$$

dann folgt

$$\begin{aligned} \alpha_j &= \frac{\left(\varphi_j(\mathbf{A})\mathbf{r}_0, \varphi_j(\mathbf{A}^T)\mathbf{r}_0 \right)_2}{\left(\mathbf{A}\psi_j(\mathbf{A})\mathbf{r}_0, \psi_j(\mathbf{A}^T)\mathbf{r}_0 \right)_2} = \frac{\left(\varphi_j^2(\mathbf{A})\mathbf{r}_0, \mathbf{r}_0 \right)_2}{\left(\mathbf{A}\psi_j^2(\mathbf{A})\mathbf{r}_0, \mathbf{r}_0 \right)_2} \\ &= \frac{(\hat{\mathbf{r}}_j, \mathbf{r}_0)_2}{(\mathbf{A}\hat{\mathbf{p}}_j, \mathbf{r}_0)_2} \end{aligned} \quad (4.3.103)$$

und

$$\beta_j = \frac{\left(\varphi_{j+1}(\mathbf{A})\mathbf{r}_0, \varphi_{j+1}(\mathbf{A}^T)\mathbf{r}_0 \right)_2}{\left(\varphi_j(\mathbf{A})\mathbf{r}_0, \varphi_j(\mathbf{A}^T)\mathbf{r}_0 \right)_2} = \frac{(\hat{\mathbf{r}}_{j+1}, \mathbf{r}_0)_2}{(\hat{\mathbf{r}}_j, \mathbf{r}_0)_2}. \quad (4.3.104)$$

Mit (4.3.101) und (4.3.102) erhalten wir

$$\psi_j^2(\lambda) = \varphi_j^2(\lambda) + 2\beta_{j-1}\varphi_j(\lambda)\psi_{j-1}(\lambda) + \beta_{j-1}^2\psi_{j-1}^2(\lambda), \quad (4.3.105)$$

$$\varphi_{j+1}^2(\lambda) = \varphi_j^2(\lambda) - \alpha_j \lambda \left(2\varphi_j^2(\lambda) + 2\beta_{j-1}\varphi_j(\lambda)\psi_{j-1}(\lambda) - \alpha_j \lambda \psi_j^2(\lambda) \right) \quad (4.3.106)$$

und für die gemischten Terme

$$\varphi_{j+1}(\lambda)\psi_j(\lambda) = \varphi_j^2(\lambda) + \beta_{j-1}\varphi_j(\lambda)\psi_{j-1}(\lambda) - \alpha_j \lambda \psi_j^2(\lambda). \quad (4.3.107)$$

Wir definieren

$$\hat{\mathbf{q}}_j := \varphi_{j+1}(\mathbf{A})\psi_j(\mathbf{A})\mathbf{r}_0 \quad (4.3.108)$$

und

$$\begin{aligned} \hat{\mathbf{u}}_j &:= \varphi_j(\mathbf{A})\psi_j(\mathbf{A})\mathbf{r}_0 \\ &\stackrel{(4.3.101)}{=} \varphi_j^2(\mathbf{A})\mathbf{r}_0 + \beta_{j-1}\varphi_j(\mathbf{A})\psi_{j-1}(\mathbf{A})\mathbf{r}_0 \\ &= \hat{\mathbf{r}}_j + \beta_{j-1}\hat{\mathbf{q}}_{j-1}. \end{aligned} \quad (4.3.109)$$

Damit ergibt sich

$$\begin{aligned} \hat{\mathbf{q}}_j &\stackrel{(4.3.107)}{=} \varphi_j^2(\mathbf{A})\mathbf{r}_0 + \beta_{j-1}\varphi_j(\mathbf{A})\psi_{j-1}(\mathbf{A})\mathbf{r}_0 - \alpha_j\mathbf{A}\psi_j^2(\mathbf{A})\mathbf{r}_0 \\ &= \hat{\mathbf{r}}_j + \beta_{j-1}\hat{\mathbf{q}}_{j-1} - \alpha_j\mathbf{A}\hat{\mathbf{p}}_j \\ &\stackrel{(4.3.109)}{=} \hat{\mathbf{u}}_j - \alpha_j\mathbf{A}\hat{\mathbf{p}}_j, \end{aligned} \quad (4.3.110)$$

$$\begin{aligned} \hat{\mathbf{r}}_{j+1} &= \varphi_{j+1}^2(\mathbf{A})\mathbf{r}_0 \\ &\stackrel{(4.3.106)}{=} \varphi_j^2(\mathbf{A})\mathbf{r}_0 - \alpha_j\mathbf{A}(2\varphi_j^2(\mathbf{A})\mathbf{r}_0 + 2\beta_{j-1}\varphi_j(\mathbf{A})\psi_{j-1}(\mathbf{A})\mathbf{r}_0 - \alpha_j\mathbf{A}\psi_j^2(\mathbf{A})\mathbf{r}_0) \\ &= \hat{\mathbf{r}}_j - \alpha_j\mathbf{A}(2\hat{\mathbf{r}}_j + 2\beta_{j-1}\hat{\mathbf{q}}_{j-1} - \alpha_j\mathbf{A}\hat{\mathbf{p}}_j) \\ &= \hat{\mathbf{r}}_j - \alpha_j\mathbf{A}\left(\underbrace{\hat{\mathbf{r}}_j + \beta_{j-1}\hat{\mathbf{q}}_{j-1}}_{=\hat{\mathbf{u}}_j} + \underbrace{\hat{\mathbf{r}}_j + \beta_{j-1}\hat{\mathbf{q}}_{j-1} - \alpha_j\mathbf{A}\hat{\mathbf{p}}_j}_{=\hat{\mathbf{q}}_j}\right) \\ &= \hat{\mathbf{r}}_j - \alpha_j\mathbf{A}(\hat{\mathbf{u}}_j + \hat{\mathbf{q}}_j) \end{aligned} \quad (4.3.111)$$

und

$$\begin{aligned} \hat{\mathbf{p}}_{j+1} &= \psi_{j+1}^2(\mathbf{A})\mathbf{r}_0 \\ &\stackrel{(4.3.105)}{=} \varphi_{j+1}^2(\mathbf{A})\mathbf{r}_0 + 2\beta_j\varphi_{j+1}(\mathbf{A})\psi_j(\mathbf{A})\mathbf{r}_0 + \beta_j^2\psi_j^2(\mathbf{A})\mathbf{r}_0 \\ &= \hat{\mathbf{r}}_{j+1} + 2\beta_j\hat{\mathbf{q}}_j + \beta_j^2\hat{\mathbf{p}}_j \\ &= \hat{\mathbf{u}}_{j+1} + \beta_j(\hat{\mathbf{q}}_j + \beta_j\hat{\mathbf{p}}_j). \end{aligned} \quad (4.3.112)$$

Mit dem Residuenvektor $\hat{\mathbf{r}}_{j+1}$ erhalten wir

$$\hat{\mathbf{x}}_{j+1} = \hat{\mathbf{x}}_j + \alpha_j(\hat{\mathbf{u}}_j + \hat{\mathbf{q}}_j). \quad (4.3.113)$$

Vernachlässigen wir das Superskript $\hat{\cdot}$, so folgt der CGS-Algorithmus durch (4.3.103), (4.3.104) sowie (4.3.109) bis (4.3.113) in der folgenden Form. Die Umsetzung eines derartigen Pseudocodes in ein ausführbares MATLAB-Programm liefert Anhang A.

CGS-Algorithmus —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$ und $\varepsilon > 0$
$\mathbf{u}_0 = \mathbf{r}_0 = \mathbf{p}_0 := \mathbf{b} - \mathbf{A} \mathbf{x}_0$, $j := 0$
Solange $\ \mathbf{r}_j\ _2 > \varepsilon$
$\mathbf{v}_j := \mathbf{A} \mathbf{p}_j, \quad \alpha_j := \frac{(\mathbf{r}_j, \mathbf{r}_0)_2}{(\mathbf{v}_j, \mathbf{r}_0)_2}$
$\mathbf{q}_j := \mathbf{u}_j - \alpha_j \mathbf{v}_j$
$\mathbf{x}_{j+1} := \mathbf{x}_j + \alpha_j (\mathbf{u}_j + \mathbf{q}_j)$
$\mathbf{r}_{j+1} := \mathbf{r}_j - \alpha_j \mathbf{A} (\mathbf{u}_j + \mathbf{q}_j)$
$\beta_j := \frac{(\mathbf{r}_{j+1}, \mathbf{r}_0)_2}{(\mathbf{r}_j, \mathbf{r}_0)_2}$
$\mathbf{u}_{j+1} := \mathbf{r}_{j+1} + \beta_j \mathbf{q}_j$
$\mathbf{p}_{j+1} := \mathbf{u}_{j+1} + \beta_j (\mathbf{q}_j + \beta_j \mathbf{p}_j)$, $j := j + 1$

4.3.2.7 Das BiCGSTAB-Verfahren

Der vorgestellte CGS-Algorithmus zeichnet sich im Vergleich zum BiCG-Verfahren durch zwei Vorteile aus. Zum einen benötigt die Methode keinerlei Operationen mit $\mathbf{A}^T$ und zum anderen weist sie aufgrund der Quadrierung des Polynoms zur Definition des Residuenvektors ein schnelleres Konvergenzverhalten bei gleichbleibendem Rechenaufwand auf. Jedoch beinhaltet auch das CGS-Verfahren die bereits im Bi-Lanczos-Algorithmus vorliegende Problematik eines vorzeitigen Abbruchs, und es zeigen sich teilweise Oszillationen im Residuenverlauf [27, 30, 50, 52, 70]. Die auftretenden vorzeitigen Verfahrensabbrüche können analog zum Look-ahead Ansatz teilweise behoben werden [14].

Mit der BiCGSTAB-Methode (BiCG Stabilized) hat van der Vorst [71] eine Variante des CGS-Verfahrens vorgestellt, die ein wesentlich glatteres Konvergenzverhalten aufweist. Hierzu werden im Gegensatz zum CGS-Algorithmus gezielt unterschiedliche Polynome bei der Definition der Suchrichtungen und Residuenvektoren betrachtet, womit bei jeder

Iteration ein zusätzlicher Freiheitsgrad zur Verfügung steht, der zur Minimierung des Residuums genutzt wird.

Zunächst werden die Residuenvektoren $\mathbf{r}_j$, $\mathbf{r}_j^*$ und die Suchvektoren $\mathbf{p}_j$, $\mathbf{p}_j^*$ entsprechend dem CGS-Verfahren definiert, wobei die hierbei benötigten Polynome φ_j und ψ_j gemäß (4.3.101) und (4.3.102) gegeben sind. Wir setzen weiterhin

$$\tilde{\mathbf{p}}_j^* := \tilde{\mathbf{r}}_j^* := \phi_j \left(\mathbf{A}^T \right) \mathbf{r}_0,$$

wobei das Polynom durch die einfache Rekursionsvorschrift

$$\phi_{j+1}(\lambda) := (1 - \omega_j \lambda) \phi_j(\lambda) \quad (4.3.114)$$

mit $\phi_0(\lambda) := 1$ festgelegt ist. Als Folgerung des Satzes 4.84 erhalten wir

$$\varphi_j(\mathbf{A}) \mathbf{r}_0 = \mathbf{r}_j \perp K_j^T$$

und damit

$$\left(\varphi_j(\mathbf{A}) \mathbf{r}_0, \pi_{j-1} \left(\mathbf{A}^T \right) \mathbf{r}_0 \right)_2 = 0, \quad (4.3.115)$$

für alle $\pi_{j-1} \in \mathcal{P}_{j-1}$. Hiermit folgt die Darstellung

$$\begin{aligned} \alpha_j &= \alpha_j \frac{\left(\varphi_j(\mathbf{A}) \mathbf{r}_0, \phi_j \left(\mathbf{A}^T \right) \mathbf{r}_0 \right)_2}{\left(\varphi_j(\mathbf{A}) \mathbf{r}_0, \phi_j \left(\mathbf{A}^T \right) \mathbf{r}_0 \right)_2} \\ &\stackrel{(4.3.115)}{=} \frac{(\phi_j(\mathbf{A}) \varphi_j(\mathbf{A}) \mathbf{r}_0, \mathbf{r}_0)_2}{\left(\frac{\varphi_j(\mathbf{A}) - \varphi_{j+1}(\mathbf{A})}{\alpha_j} \mathbf{r}_0, \phi_j \left(\mathbf{A}^T \right) \mathbf{r}_0 \right)_2} \\ &\stackrel{(4.3.102)}{=} \frac{(\phi_j(\mathbf{A}) \varphi_j(\mathbf{A}) \mathbf{r}_0, \mathbf{r}_0)_2}{(\phi_j(\mathbf{A}) \mathbf{A} \psi_j(\mathbf{A}) \mathbf{r}_0, \mathbf{r}_0)_2}, \end{aligned}$$

und wir definieren daher

$$\hat{\mathbf{r}}_j := \phi_j(\mathbf{A}) \varphi_j(\mathbf{A}) \mathbf{r}_0 \quad (4.3.116)$$

sowie

$$\hat{\mathbf{p}}_j := \phi_j(\mathbf{A}) \psi_j(\mathbf{A}) \mathbf{r}_0. \quad (4.3.117)$$

Damit folgt

$$\begin{aligned} \hat{\mathbf{s}}_j &:= \phi_j(\mathbf{A}) \varphi_{j+1}(\mathbf{A}) \mathbf{r}_0 \\ &\stackrel{(4.3.102)}{=} \phi_j(\mathbf{A}) \varphi_j(\mathbf{A}) \mathbf{r}_0 - \alpha_j \mathbf{A} \phi_j(\mathbf{A}) \psi_j(\mathbf{A}) \mathbf{r}_0 \\ &= \hat{\mathbf{r}}_j - \alpha_j \mathbf{A} \hat{\mathbf{p}}_j \end{aligned} \quad (4.3.118)$$

und wir erhalten die Rekursionsformeln

$$\begin{aligned} \hat{\mathbf{r}}_{j+1} &= \phi_{j+1}(\mathbf{A}) \varphi_{j+1}(\mathbf{A}) \mathbf{r}_0 \\ &= (\mathbf{I} - \omega_j \mathbf{A}) \phi_j(\mathbf{A}) \varphi_{j+1}(\mathbf{A}) \mathbf{r}_0 \\ &= (\mathbf{I} - \omega_j \mathbf{A}) \hat{\mathbf{s}}_j \end{aligned} \quad (4.3.119)$$

und

$$\begin{aligned}
 \hat{\mathbf{p}}_{j+1} &= \phi_{j+1}(\mathbf{A}) \psi_{j+1}(\mathbf{A}) \mathbf{r}_0 \\
 &\stackrel{(4.3.101)}{=} \phi_{j+1}(\mathbf{A}) \varphi_{j+1}(\mathbf{A}) \mathbf{r}_0 + \beta_j \phi_{j+1}(\mathbf{A}) \psi_j(\mathbf{A}) \mathbf{r}_0 \\
 &= \hat{\mathbf{r}}_{j+1} + \beta_j (\mathbf{I} - \omega_j \mathbf{A}) \phi_j(\mathbf{A}) \psi_j(\mathbf{A}) \mathbf{r}_0 \\
 &= \hat{\mathbf{r}}_{j+1} + \beta_j (\mathbf{I} - \omega_j \mathbf{A}) \hat{\mathbf{p}}_j.
 \end{aligned} \tag{4.3.120}$$

Um eine vollständige Darstellung des Verfahrens zu erhalten, müssen wir abschließend die Berechnung der skalaren Größen β_j und ω_j erläutern. Wir schreiben

$$\begin{aligned}
 \varphi_j(\lambda) &\stackrel{(4.3.101)}{=} -\alpha_{j-1} \lambda \varphi_{j-1}(\lambda) + \pi_{j-1}(\lambda) \\
 &\stackrel{(4.3.102)}{=} \alpha_{j-1} \alpha_{j-2} \lambda^2 \varphi_{j-2}(\lambda) + \tilde{\pi}_{j-1}(\lambda) \\
 &= (-1)^j \alpha_{j-1} \dots \alpha_0 \lambda^j + \bar{\pi}_{j-1}(\lambda)
 \end{aligned} \tag{4.3.121}$$

mit $\pi_{j-1}, \tilde{\pi}_{j-1}, \bar{\pi}_{j-1} \in \mathcal{P}_{j-1}$. Analog ergibt sich die Darstellung

$$\phi_j(\lambda) = (-1)^j \omega_{j-1} \dots \omega_0 \lambda^j + \hat{\pi}_{j-1}(\lambda) \tag{4.3.122}$$

mit $\hat{\pi}_{j-1} \in \mathcal{P}_{j-1}$, wodurch mit $\mathbf{r}_j^* = \varphi_j(\mathbf{A}^T) \mathbf{r}_0$ die Gleichung

$$\begin{aligned}
 \frac{(\mathbf{r}_j, \mathbf{r}_j^*)_2}{(\mathbf{r}_j, \hat{\mathbf{r}}_j^*)_2} &= \frac{(\varphi_j(\mathbf{A}) \mathbf{r}_0, \varphi_j(\mathbf{A}^T) \mathbf{r}_0)_2}{(\varphi_j(\mathbf{A}) \mathbf{r}_0, \phi_j(\mathbf{A}^T) \mathbf{r}_0)_2} \\
 &\stackrel{(4.3.121)}{=} \frac{(\varphi_j(\mathbf{A}) \mathbf{r}_0, (-1)^j \alpha_{j-1} \dots \alpha_0 (\mathbf{A}^T)^j \mathbf{r}_0 + \bar{\pi}_{j-1}(\mathbf{A}^T) \mathbf{r}_0)_2}{(\varphi_j(\mathbf{A}) \mathbf{r}_0, (-1)^j \omega_{j-1} \dots \omega_0 (\mathbf{A}^T)^j \mathbf{r}_0 + \hat{\pi}_{j-1}(\mathbf{A}^T) \mathbf{r}_0)_2} \\
 &\stackrel{(4.3.115)}{=} \frac{\alpha_{j-1} \dots \alpha_0}{\omega_{j-1} \dots \omega_0} \frac{(\varphi_j(\mathbf{A}) \mathbf{r}_0, (\mathbf{A}^T)^j \mathbf{r}_0)_2}{(\varphi_j(\mathbf{A}) \mathbf{r}_0, (\mathbf{A}^T)^j \mathbf{r}_0)_2} \\
 &= \frac{\alpha_{j-1} \dots \alpha_0}{\omega_{j-1} \dots \omega_0}
 \end{aligned}$$

folgt. Somit gilt

$$\begin{aligned}
\beta_j &= \frac{(\mathbf{r}_{j+1}, \mathbf{r}_{j+1}^*)_2}{(\mathbf{r}_j, \mathbf{r}_j^*)_2} \\
&= \frac{(\mathbf{r}_{j+1}, \mathbf{r}_{j+1}^*)_2}{(\mathbf{r}_{j+1}, \tilde{\mathbf{r}}_{j+1}^*)_2} \frac{(\mathbf{r}_{j+1}, \tilde{\mathbf{r}}_{j+1}^*)_2}{(\mathbf{r}_j, \tilde{\mathbf{r}}_j^*)_2} \frac{(\mathbf{r}_j, \tilde{\mathbf{r}}_j^*)_2}{(\mathbf{r}_j, \mathbf{r}_j^*)_2} \\
&= \frac{\alpha_j}{\omega_j} \frac{(\mathbf{r}_{j+1}, \tilde{\mathbf{r}}_{j+1}^*)_2}{(\mathbf{r}_j, \tilde{\mathbf{r}}_j^*)_2} \\
&= \frac{\alpha_j}{\omega_j} \frac{(\varphi_{j+1}(\mathbf{A}) \mathbf{r}_0, \phi_{j+1}(\mathbf{A}^T) \mathbf{r}_0)_2}{(\varphi_j(\mathbf{A}) \mathbf{r}_0, \phi_j(\mathbf{A}^T) \mathbf{r}_0)_2} \\
&= \frac{\alpha_j}{\omega_j} \frac{(\hat{\mathbf{r}}_{j+1}, \mathbf{r}_0)_2}{(\hat{\mathbf{r}}_j, \mathbf{r}_0)_2}. \tag{4.3.123}
\end{aligned}$$

Mit ω_j steht uns ein Parameter zur gezielten Minimierung des Residuums zur Verfügung. Wir betrachten die Funktion $f_j : \mathbb{R} \rightarrow \mathbb{R}$

$$f_j(\omega) := \|(\mathbf{I} - \omega \mathbf{A}) \hat{\mathbf{s}}_j\|_2^2.$$

Für $\hat{\mathbf{s}}_j \neq \mathbf{0}$ ergibt sich $f_j''(\omega) = 2\|\mathbf{A}\hat{\mathbf{s}}_j\|_2^2 > 0$, wodurch die Funktion in diesem Fall strikt konvex ist. Wie wir im folgenden sehen werden, kann im Fall $\hat{\mathbf{s}}_j = \mathbf{0}$ der Parameter ω_j beliebig gewählt werden und das Verfahren liefert die exakte Lösung. Wir beschränken uns daher auf den Fall $\hat{\mathbf{s}}_j \neq \mathbf{0}$, wodurch wir

$$\omega_j := \arg \min_{\omega \in \mathbb{R}} f_j(\omega)$$

mittels

$$0 = f_j'(\omega_j) = -2(\mathbf{A}\hat{\mathbf{s}}_j, \hat{\mathbf{s}}_j)_2 + 2\omega_j(\mathbf{A}\hat{\mathbf{s}}_j, \mathbf{A}\hat{\mathbf{s}}_j)_2$$

in der Form

$$\omega_j = \frac{(\mathbf{A}\hat{\mathbf{s}}_j, \hat{\mathbf{s}}_j)_2}{(\mathbf{A}\hat{\mathbf{s}}_j, \mathbf{A}\hat{\mathbf{s}}_j)_2} \tag{4.3.124}$$

erhalten. Die BiCGSTAB-Iterierte läßt sich folglich gemäß

$$\begin{aligned}
\hat{\mathbf{x}}_{j+1} &= \mathbf{A}^{-1}(\mathbf{b} - \hat{\mathbf{r}}_{j+1}) \\
&\stackrel{(4.3.119)}{=} \mathbf{A}^{-1}(\mathbf{b} - \hat{\mathbf{s}}_j + \omega_j \mathbf{A}\hat{\mathbf{s}}_j) \\
&\stackrel{(4.3.118)}{=} \mathbf{A}^{-1}(\mathbf{b} - \hat{\mathbf{r}}_j + \alpha_j \mathbf{A}\hat{\mathbf{p}}_j + \omega_j \mathbf{A}\hat{\mathbf{s}}_j) \\
&= \hat{\mathbf{x}}_j + \alpha_j \hat{\mathbf{p}}_j + \omega_j \hat{\mathbf{s}}_j \tag{4.3.125}
\end{aligned}$$

schreiben. Vernachlässigen wir wiederum das Superskript $\wedge$, so ergibt sich mit (4.3.118) bis (4.3.120), (4.3.123) bis (4.3.125) und der Definition der skalaren Größe α_j zusammenfassend der BiCGSTAB-Algorithmus in der folgenden Form. Zur Implementierung siehe Anhang A.

BiCGSTAB-Algorithmus —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$ und $\varepsilon > 0$
$\mathbf{r}_0 := \mathbf{p}_0 := \mathbf{b} - \mathbf{A} \mathbf{x}_0$, $\rho_0 := (\mathbf{r}_0, \mathbf{r}_0)_2$, $j := 0$
Solange $\ \mathbf{r}_j\ _2 > \varepsilon$
$\mathbf{v}_j := \mathbf{A} \mathbf{p}_j$, $\alpha_j := \frac{\rho_j}{(\mathbf{v}_j, \mathbf{r}_0)_2}$
$\mathbf{s}_j := \mathbf{r}_j - \alpha_j \mathbf{v}_j$, $\mathbf{t}_j := \mathbf{A} \mathbf{s}_j$
$\omega_j := \frac{(\mathbf{t}_j, \mathbf{s}_j)_2}{(\mathbf{t}_j, \mathbf{t}_j)_2}$
$\mathbf{x}_{j+1} := \mathbf{x}_j + \alpha_j \mathbf{p}_j + \omega_j \mathbf{s}_j$
$\mathbf{r}_{j+1} := \mathbf{s}_j - \omega_j \mathbf{t}_j$
$\rho_{j+1} := (\mathbf{r}_{j+1}, \mathbf{r}_0)_2$, $\beta_j := \frac{\alpha_j}{\omega_j} \frac{\rho_{j+1}}{\rho_j}$
$\mathbf{p}_{j+1} := \mathbf{r}_{j+1} + \beta_j (\mathbf{p}_j - \omega_j \mathbf{v}_j)$, $j := j + 1$

Eine Verallgemeinerung des vorgestellten Verfahrens im Sinne einer ℓ -dimensionalen anstelle einer eindimensionalen Minimierung stellt die BiCGSTAB(ℓ)-Methode dar, die 1993 von Sleijpen und Fokkema [64] hergeleitet wurde. Der Algorithmus kann als Kombination der GMRES(ℓ)-Methode und des BiCG-Verfahrens interpretiert werden und stimmt im Fall $\ell = 1$ mit der BiCGSTAB-Methode überein. Dabei werden zunächst implizit ℓ BiCG-Residuenvektoren berechnet und anschließend eine Minimierung über den hierdurch festgelegten ℓ -dimensionalen Unterraum durchgeführt.

Wir haben bereits im Abschnitt 4.3.2.3 die möglichen Verfahrensabbrüche des Bi-Lanczos-Algorithmus diskutiert. Mit der folgenden Variante des BiCGSTAB-Verfahrens werden wir eine Stabilisierung der Methode hinsichtlich dieser Problematik präsentieren. Zunächst werden wir vor der Berechnung des Skalars α_j die Größe des auftretenden Nenners überprüfen, um bei großen, aber annähernd orthogonalen Vektoren $\mathbf{v}_j$ und $\mathbf{r}_0$ eine Divisionen durch sehr kleine Werte zu vermeiden. Liegt der Betrag des Nenners relativ zur Länge der Vektoren $\mathbf{v}_j$ und $\mathbf{r}_0$ unterhalb einer vorgegebenen Schranke, so wird ein Restart des Verfahrens durchgeführt. Eine weitere mögliche Division durch Null tritt bei der Berechnung von ω_j auf. Zur Behebung dieser Problematik wird in [8] eine Kontrolle des Vektors $\mathbf{s}_j$ empfohlen, der dem Residuenvektor im BiCG-Verfahren entspricht. Gilt $\|\mathbf{s}_j\|_2 \leq \varepsilon$, so erhalten wir mit

$$\mathbf{x}_{j+1} = \mathbf{x}_j + \alpha_j \mathbf{p}_j$$

bereits

$$\|\mathbf{r}_{j+1}\|_2 = \|\mathbf{r}_j - \alpha_j \mathbf{v}_j\|_2 = \|\mathbf{s}_j\|_2 \leq \varepsilon,$$

wodurch keine Notwendigkeit zur Ermittlung des Vektors $\mathbf{t}_j$ und damit der Berechnung der skalaren Größe ω_j vorliegt. Durch die Berücksichtigung der obigen Verbesserungen ergibt sich die folgende Stabilisierung des BiCGSTAB-Verfahrens.

Stabilisierter BiCGSTAB-Algorithmus —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$ und $\varepsilon, \tilde{\varepsilon} > 0$		
$\mathbf{r}_0 := \mathbf{p}_0 := \mathbf{b} - \mathbf{A} \mathbf{x}_0$, $\rho_0 := (\mathbf{r}_0, \mathbf{r}_0)_2$, $j := 0$		
Solange $\ \mathbf{r}_j\ _2 > \varepsilon$		
$\mathbf{v}_j := \mathbf{A} \mathbf{p}_j$		
$\sigma_j := (\mathbf{v}_j, \mathbf{r}_0)_2$		
Y $\sigma_j > \tilde{\varepsilon} \ \mathbf{v}_j\ _2 \ \mathbf{r}_0\ _2$ N		
$\alpha_j := \frac{\rho_j}{\sigma_j}$		Restart mit $\mathbf{x}_0 = \mathbf{x}_j$
$\mathbf{s}_j := \mathbf{r}_j - \alpha_j \mathbf{v}_j$		
Y $\ \mathbf{s}_j\ _2 > \varepsilon$ N		
$\mathbf{t}_j := \mathbf{A} \mathbf{s}_j$	$\mathbf{x}_{j+1} := \mathbf{x}_j + \alpha_j \mathbf{p}_j$	
$\omega_j := \frac{(\mathbf{t}_j, \mathbf{s}_j)_2}{(\mathbf{t}_j, \mathbf{t}_j)_2}$	$\mathbf{r}_{j+1} := \mathbf{s}_j$	
	$j := j + 1$	
$\mathbf{x}_{j+1} := \mathbf{x}_j + \alpha_j \mathbf{p}_j + \omega_j \mathbf{s}_j$		
$\mathbf{r}_{j+1} := \mathbf{s}_j - \omega_j \mathbf{t}_j$		
$\rho_{j+1} := (\mathbf{r}_{j+1}, \mathbf{r}_0)_2$		
$\beta_j := \frac{\alpha_j}{\omega_j} \frac{\rho_{j+1}}{\rho_j}$		
$\mathbf{p}_{j+1} := \mathbf{r}_{j+1} + \beta_j (\mathbf{p}_j - \omega_j \mathbf{v}_j)$		
$j := j + 1$		

Natürlich können zudem weitere Modifikationen zu einem besseren Verhalten der Methode beitragen. Eine Möglichkeit besteht in der Kontrolle der Variation der relevanten Größen. Wegen

$$\frac{\|\mathbf{r}_{j+1} - \mathbf{r}_j\|_2}{\|\mathbf{r}_j\|_2} \leq \frac{\|\mathbf{s}_j - \mathbf{r}_j\|_2}{\|\mathbf{r}_j\|_2} = \frac{|\alpha_j| \|\mathbf{v}_j\|_2}{\|\mathbf{r}_j\|_2}$$

liegt im Fall

$$\frac{|\alpha_j| \|\mathbf{v}_j\|_2}{\|\mathbf{r}_j\|_2} \leq \gamma$$

die relative Änderung des Residuums unterhalb einer zu wählenden Grenze γ , so daß ein Restart sinnvoll erscheint. Zur Kontrolle von Rundungsfehlern kann in regelmäßigen Abständen der im Verfahren ermittelte Residuenvektor $\mathbf{r}_{j+1}$ mit dem tatsächlichen Residuenvektor $\mathbf{b} - \mathbf{A}\mathbf{x}_{j+1}$ verglichen werden. Größere Abweichungen weisen auf Rundungsfehler hin, wodurch ebenfalls ein Restart angebracht erscheint. Eine Übertragung der Look-ahead-Strategie auf das BiCGSTAB-Verfahren wird in [12] vorgestellt.

4.3.2.8 Das TFQMR-Verfahren

Mit der QMR-Methode (Quasi Minimal Residual) präsentierten Freund und Nachtigal [28] ein neuartiges BiCG-ähnliches Verfahren, das als Kombination zwischen der GMRES-Methode und dem BiCG-Verfahren angesehen werden kann. Das Verfahren verknüpft die Vorteile des geringen Speicher- und Rechenaufwandes der BiCG-Methode, die auf die Verwendung des Bi-Lanczos-Algorithmus zurückgeht, mit der Eigenschaft des glatten Residuenverlaufes des GMRES-Verfahrens, das auf der Minimierung des Residuums beruht. Hierzu wurde die vorausschauende Variante des Bi-Lanczos-Verfahrens, der Look-ahead Lanczos-Algorithmus, zur Berechnung einer Basis $\{\mathbf{v}_1, \dots, \mathbf{v}_m\}$ des K_m genutzt. Die aus den Basisvektoren bestehende Matrix $\mathbf{V}_m = (\mathbf{v}_1 \dots \mathbf{v}_m)$ ist aufgrund des genutzten Bi-Lanczos-Algorithmus nicht notwendigerweise orthogonal, weshalb beim QMR-Verfahren im Gegensatz zur GMRES-Methode lediglich eine Quasiminimierung der durch (4.3.59) gegebenen Funktion vorgenommen wird. In natürlicher Weise impliziert der BiCG-Algorithmus einen notwendigen Zugriff auf die Matrix $\mathbf{A}^T$ im QMR-Verfahren. Die im folgenden beschriebene TFQMR-Methode (Transpose-Free QMR) wurde von Freund [27] auf der Basis des CGS-Verfahrens anstelle der BiCG-Methode entwickelt, wodurch Matrix-Vektor-Multiplikationen mit $\mathbf{A}^T$ entfallen. Bevor wir zur direkten Herleitung des Verfahrens kommen, werden wir mit den folgenden beiden Lemmata eine Basis des Krylov-Unterraums K_m definieren und eine Eigenschaft der Basisvektoren beweisen. Hierzu sei die Abbildung $\lfloor \cdot \rfloor : \mathbb{R} \rightarrow \mathbb{Z}$ durch

$$x \xrightarrow{\lfloor \cdot \rfloor} \lfloor x \rfloor \in]x - 1, x] \quad (4.3.126)$$

definiert.

Lemma 4.85 *Vorausgesetzt, das CGS-Verfahren bricht nicht vor der Berechnung der Vektoren $\mathbf{q}_{\lfloor \frac{m-1}{2} \rfloor}$ und $\mathbf{u}_{\lfloor \frac{m-1}{2} \rfloor}$ ab, dann stellen die durch*

$$\mathbf{y}_k := \begin{cases} \mathbf{u}_{\lfloor \frac{k-1}{2} \rfloor} & , \text{ falls } k \text{ ungerade,} \\ \mathbf{q}_{\lfloor \frac{k-1}{2} \rfloor} & , \text{ falls } k \text{ gerade} \end{cases} \quad (4.3.127)$$

definierten Vektoren $\mathbf{y}_1, \dots, \mathbf{y}_m$ eine Basis des K_m dar.

Beweis:

Sei $n \in \mathbb{N}_0$, dann folgt für $k = 2n + 1$

$$\begin{aligned} \mathbf{y}_k &= \mathbf{u}_n = \varphi_n(\mathbf{A}) \psi_n(\mathbf{A}) \mathbf{r}_0 \\ &= \alpha_0^2 \cdot \dots \cdot \alpha_{n-1}^2 \mathbf{A}^{k-1} \mathbf{r}_0 + \pi(\mathbf{A}) \mathbf{r}_0 \quad \text{mit } \pi \in \mathcal{P}_{k-2} \end{aligned}$$

und für $k = 2n > 0$

$$\begin{aligned} \mathbf{y}_k &= \mathbf{q}_{n-1} = \varphi_n(\mathbf{A}) \psi_{n-1}(\mathbf{A}) \mathbf{r}_0 \\ &= -\alpha_0^2 \cdot \dots \cdot \alpha_{n-2}^2 \alpha_{n-1} \mathbf{A}^{k-1} \mathbf{r}_0 + \pi(\mathbf{A}) \mathbf{r}_0 \quad \text{mit } \pi \in \mathcal{P}_{k-2}. \end{aligned}$$

Laut Voraussetzung gilt $\alpha_i \neq 0$ für $i = 0, \dots, n-1$ und damit

$$\text{span}\{\mathbf{y}_1, \dots, \mathbf{y}_m\} = K_m.$$

□

Lemma 4.86 *Sei*

$$\mathbf{w}_k := \begin{cases} \varphi_{\lfloor \frac{k-1}{2} \rfloor}^2(\mathbf{A}) \mathbf{r}_0 & , \quad \text{falls } k \text{ ungerade,} \\ \varphi_{\lfloor \frac{k-1}{2} \rfloor}(\mathbf{A}) \varphi_{\lfloor \frac{k}{2} \rfloor}(\mathbf{A}) \mathbf{r}_0 & , \quad \text{falls } k \text{ gerade,} \end{cases} \quad (4.3.128)$$

dann gilt

$$\mathbf{w}_{k+1} = \mathbf{w}_k - \alpha_{\lfloor \frac{k-1}{2} \rfloor} \mathbf{A} \mathbf{y}_k \quad \text{für } k = 1, \dots, m$$

mit den durch (4.3.127) gegebenen Vektoren $\mathbf{y}_1, \dots, \mathbf{y}_m$.

Beweis:

Aus (4.3.102) erhalten wir

$$\lambda \psi_j(\lambda) = \frac{1}{\alpha_j} (\varphi_j(\lambda) - \varphi_{j+1}(\lambda)). \quad (4.3.129)$$

Sei $n \in \mathbb{N}_0$, dann folgt für $k = 2n + 1$

$$\begin{aligned} \mathbf{A} \mathbf{y}_k &= \mathbf{A} \mathbf{u}_n \\ &\stackrel{(4.3.109)}{=} \frac{1}{\alpha_n} (\varphi_n^2(\mathbf{A}) \mathbf{r}_0 - \varphi_n(\mathbf{A}) \varphi_{n+1}(\mathbf{A}) \mathbf{r}_0) \\ &= \frac{1}{\alpha_{\lfloor \frac{k-1}{2} \rfloor}} \left(\varphi_{\lfloor \frac{k-1}{2} \rfloor}^2(\mathbf{A}) \mathbf{r}_0 - \varphi_{\lfloor \frac{(k+1)-1}{2} \rfloor}(\mathbf{A}) \varphi_{\lfloor \frac{k+1}{2} \rfloor}(\mathbf{A}) \mathbf{r}_0 \right) \\ &= \frac{1}{\alpha_{\lfloor \frac{k-1}{2} \rfloor}} (\mathbf{w}_k - \mathbf{w}_{k+1}) \end{aligned}$$

und für $k = 2n > 0$

$$\begin{aligned}
\mathbf{A}\mathbf{y}_k &= \mathbf{A}\mathbf{q}_{n-1} \\
&\stackrel{(4.3.108)}{=} \stackrel{(4.3.129)}{=} \frac{1}{\alpha_{n-1}} (\varphi_{n-1}(\mathbf{A}) \varphi_n(\mathbf{A}) \mathbf{r}_0 - \varphi_n^2(\mathbf{A}) \mathbf{r}_0) \\
&= \frac{1}{\alpha_{\lfloor \frac{k-1}{2} \rfloor}} \left(\varphi_{\lfloor \frac{k-1}{2} \rfloor}(\mathbf{A}) \varphi_{\lfloor \frac{k}{2} \rfloor}(\mathbf{A}) \mathbf{r}_0 - \varphi_{\lfloor \frac{(k+1)-1}{2} \rfloor}^2(\mathbf{A}) \mathbf{r}_0 \right) \\
&= \frac{1}{\alpha_{\lfloor \frac{k-1}{2} \rfloor}} (\mathbf{w}_k - \mathbf{w}_{k+1}).
\end{aligned}$$

□

Aufgrund der vorgenommenen Definition des Vektors $\mathbf{w}_k$ gilt stets $\mathbf{w}_{2k+1} = \mathbf{r}_k$. Mit

$$\mathbf{Y}_m = (\mathbf{y}_1 \dots \mathbf{y}_m), \quad (4.3.130)$$

$$\mathbf{W}_{m+1} = (\mathbf{w}_1 \dots \mathbf{w}_{m+1}) \quad (4.3.131)$$

und

$$\overline{\mathbf{B}}_m = \begin{pmatrix} \alpha_{\lfloor \frac{1-1}{2} \rfloor}^{-1} & & & & \\ -\alpha_{\lfloor \frac{1-1}{2} \rfloor}^{-1} & \alpha_{\lfloor \frac{2-1}{2} \rfloor}^{-1} & & & \\ & \ddots & \ddots & & \\ & & -\alpha_{\lfloor \frac{m-2}{2} \rfloor}^{-1} & \alpha_{\lfloor \frac{m-1}{2} \rfloor}^{-1} & \\ & & & -\alpha_{\lfloor \frac{m-1}{2} \rfloor}^{-1} & \end{pmatrix} \in \mathbb{R}^{(m+1) \times m}$$

läßt sich Lemma 4.86 in der Form

$$\mathbf{A}\mathbf{Y}_m = \mathbf{W}_{m+1} \overline{\mathbf{B}}_m \quad (4.3.132)$$

schreiben. Zudem existiert aufgrund des Lemmas 4.85 für jeden Vektor $\mathbf{x}_m \in \mathbf{x}_0 + K_m$ die Darstellung

$$\mathbf{x}_m = \mathbf{x}_0 + \mathbf{Y}_m \boldsymbol{\alpha}_m$$

mit $\boldsymbol{\alpha}_m \in \mathbb{R}^m$. Durch (4.3.128) folgt $\mathbf{w}_1 = \mathbf{r}_0$ und wir erhalten daher für den zu $\mathbf{x}_m$ gehörenden Residuenvektor unter Verwendung der Gleichung (4.3.132)

$$\begin{aligned}
\mathbf{r}_m &= \mathbf{b} - \mathbf{A}\mathbf{x}_m \\
&= \mathbf{r}_0 - \mathbf{A}\mathbf{Y}_m \boldsymbol{\alpha}_m \\
&= \mathbf{W}_{m+1} (\mathbf{e}_1 - \overline{\mathbf{B}}_m \boldsymbol{\alpha}_m)
\end{aligned} \quad (4.3.133)$$

mit $\mathbf{e}_1 = (1, 0, \dots, 0)^T \in \mathbb{R}^{m+1}$. Die Matrix $\mathbf{W}_{m+1}$ ist in der Regel nicht normerhaltend, weshalb bei der Quasiminimierung keine vollständige Minimierung des Residuums betrachtet wird. Unter zusätzlicher Verwendung der Skalierungsmatrix

$$\mathbf{S}_{m+1} := \begin{pmatrix} \|\mathbf{w}_1\|_2 & & & \\ & \ddots & & \\ & & & \|\mathbf{w}_{m+1}\|_2 \end{pmatrix} \in \mathbb{R}^{(m+1) \times (m+1)}$$

folgt

$$\|\mathbf{r}_m\|_2 \leq \left\| \mathbf{W}_{m+1} \mathbf{S}_{m+1}^{-1} \right\|_2 \underbrace{\left\| \|\mathbf{w}_1\|_2 \mathbf{e}_1 - \mathbf{S}_{m+1} \overline{\mathbf{B}}_m \boldsymbol{\alpha}_m \right\|_2}_{=:\tau_m}. \quad (4.3.134)$$

Die Skalierung dient in der obigen Formulierung als Vorkonditionierer, da die Konditionszahl der Matrix $\mathbf{W}_{m+1} \mathbf{S}_{m+1}^{-1}$ ein Kriterium für die Güte der Abschätzung des Residuums durch die Größe τ_m darstellt. Wir erhalten die im folgenden Lemma bewiesene Abschätzung:

Lemma 4.87 Sei $\tau_m := \left\| \|\mathbf{w}_1\|_2 \mathbf{e}_1 - \mathbf{S}_{m+1} \overline{\mathbf{B}}_m \boldsymbol{\alpha}_m \right\|_2$, dann folgt

$$\|\mathbf{r}_m\|_2 \leq \sqrt{m+1} \tau_m$$

für den durch (4.3.133) gegebenen Residuenvektor $\mathbf{r}_m$.

Beweis:

Da die euklidische Norm einer beliebigen Matrix $\mathbf{A} = (a_{ij}) \in \mathbb{R}^{m \times n}$ der Abschätzung

$$\|\mathbf{A}\|_2 \leq \left(\sum_{i=1}^m \sum_{j=1}^n |a_{ij}|^2 \right)^{1/2} = \|\mathbf{A}\|_F$$

genügt, folgt die Behauptung durch Einsetzen der Ungleichung

$$\left\| \mathbf{W}_{m+1} \mathbf{S}_{m+1}^{-1} \right\|_2 \leq \left(\sum_{i=1}^{m+1} \left\| \frac{\mathbf{w}_i}{\|\mathbf{w}_i\|_2} \right\|_2^2 \right)^{1/2} = \sqrt{m+1}$$

in die Gleichung (4.3.134). □

In Anlehnung an die im GMRES-Verfahren durch (4.3.64) eingeführte Abbildung nutzen wir aufgrund der Abschätzung (4.3.134) die Funktion $J_m : \mathbb{R}^m \rightarrow \mathbb{R}$ mit

$$J_m(\boldsymbol{\alpha}) := \left\| \|\mathbf{w}_1\|_2 \mathbf{e}_1 - \mathbf{S}_{m+1} \overline{\mathbf{B}}_m \boldsymbol{\alpha} \right\|_2 \quad (4.3.135)$$

und beschreiben mit den folgenden zwei Sätzen ein Iterationsverfahren zur Lösung der Minimierungsaufgabe

$$\boldsymbol{\alpha}_m = \arg \min_{\boldsymbol{\alpha} \in \mathbb{R}^m} J_m(\boldsymbol{\alpha}),$$

die die TFQMR-Iterierte mittels

$$\mathbf{x}_m = \mathbf{x}_0 + \mathbf{Y}_m \boldsymbol{\alpha}_m \quad (4.3.136)$$

liefert.

Satz 4.88 Sei $m \geq 1$ und gelte

$$\overline{\mathbf{T}}_m = \begin{pmatrix} & \mathbf{T}_m \\ h_{m+1,1} \dots h_{m+1,m} \end{pmatrix} := \mathbf{S}_{m+1} \overline{\mathbf{B}}_m \in \mathbb{R}^{(m+1) \times m}$$

mit einer regulären Matrix $\mathbf{T}_m$. Für $k = m-1, m$ seien desweiteren

$$\boldsymbol{\alpha}_k = \arg \min_{\boldsymbol{\alpha} \in \mathbb{R}^k} J_k(\boldsymbol{\alpha}),$$

und

$$\tau_k := J_k(\boldsymbol{\alpha}_k)$$

mit der durch (4.3.135) gegebenen Funktion J_k . Dann gilt

$$\boldsymbol{\alpha}_m = (1 - c_m^2) \begin{pmatrix} \boldsymbol{\alpha}^{m-1} \\ 0 \end{pmatrix} + c_m^2 \tilde{\boldsymbol{\alpha}}_m \quad (4.3.137)$$

mit

$$\tilde{\boldsymbol{\alpha}}_m := \left(\alpha_{\lfloor \frac{j-1}{2} \rfloor} \right)_{j=1, \dots, m}^T = \left(\alpha_0, \alpha_0, \alpha_1, \dots, \alpha_{\lfloor \frac{m-1}{2} \rfloor} \right)^T \in \mathbb{R}^m,$$

sowie

$$\tau_m = \tau_{m-1} \vartheta_m c_m$$

mit

$$\vartheta_m = \frac{\|\mathbf{w}_{m+1}\|_2}{\tau_{m-1}}$$

und

$$c_m = (1 + \vartheta_m^2)^{-1/2}. \quad (4.3.138)$$

Beweis:

Einfaches Nachrechnen liefert

$$\overline{\mathbf{T}}_m = \begin{pmatrix} \frac{\|\mathbf{w}_1\|_2}{\alpha_0} & & & & \\ -\frac{\|\mathbf{w}_2\|_2}{\alpha_0} & \frac{\|\mathbf{w}_2\|_2}{\alpha_0} & & & \\ & \ddots & \ddots & & \\ & & -\frac{\|\mathbf{w}_m\|_2}{\alpha_{\lfloor \frac{m-2}{2} \rfloor}} & \frac{\|\mathbf{w}_m\|_2}{\alpha_{\lfloor \frac{m-1}{2} \rfloor}} & \\ & & & -\frac{\|\mathbf{w}_{m+1}\|_2}{\alpha_{\lfloor \frac{m-1}{2} \rfloor}} & \end{pmatrix},$$

wodurch

$$\tilde{\boldsymbol{\alpha}}_m = \|\mathbf{w}_1\|_2 \mathbf{T}_m^{-1} \mathbf{e}_1 \quad (4.3.139)$$

und somit

$$\|\mathbf{w}_{m+1}\|_2 = \left\| \|\mathbf{w}_1\|_2 \mathbf{e}_1 - \overline{\mathbf{T}}_m \tilde{\boldsymbol{\alpha}}_m \right\|_2 \quad (4.3.140)$$

gilt. Wir führen analog zum Lemma 4.75 eine orthogonale Transformation der Matrix $\overline{\mathbf{T}}_m$ mittels Givens-Rotationen durch und erhalten die Darstellung

$$\mathbf{Q}_m \overline{\mathbf{T}}_m = \overline{\mathbf{R}}_m = \begin{pmatrix} \mathbf{R}_m \\ \mathbf{0}^T \end{pmatrix}. \quad (4.3.141)$$

Zudem definieren wir

$$\overline{\mathbf{g}}_m := \|\mathbf{w}_1\|_2 \mathbf{Q}_m \mathbf{e}_1.$$

Aufgrund der besonderen Gestalt der Matrix $\bar{\mathbf{T}}_m$ stellt die Abbildung $\mathbf{R}_m$ eine rechte obere Dreiecksmatrix mit Bandbreite zwei dar. Für den Vektor $\bar{\mathbf{g}}_m$ ergibt sich dabei mit

$$\bar{\mathbf{g}}_m = \begin{pmatrix} \mathbf{g}_m \\ \tilde{\gamma}_{m+1} \end{pmatrix} = \begin{pmatrix} \gamma_1 \\ \vdots \\ \gamma_m \\ \tilde{\gamma}_{m+1} \end{pmatrix} \quad \text{und} \quad \bar{\mathbf{g}}_{m-1} = \begin{pmatrix} \mathbf{g}_{m-1} \\ \tilde{\gamma}_m \end{pmatrix}$$

unter Verwendung der durch (4.3.141) gegebenen Transformation die Form

$$\begin{aligned} \bar{\mathbf{g}}_m &= \begin{pmatrix} \mathbf{I} & & \\ & c_m & s_m \\ & -s_m & c_m \end{pmatrix} \begin{pmatrix} \mathbf{Q}_{m-1} & \mathbf{0} \\ \mathbf{0}^T & 1 \end{pmatrix} \begin{pmatrix} \|\mathbf{w}_1\|_2 \\ \mathbf{0} \end{pmatrix} \\ &= \begin{pmatrix} \mathbf{I} & & \\ & c_m & s_m \\ & -s_m & c_m \end{pmatrix} \begin{pmatrix} \mathbf{g}_{m-1} \\ \tilde{\gamma}_m \\ 0 \end{pmatrix} = \begin{pmatrix} \mathbf{g}_{m-1} \\ c_m \tilde{\gamma}_m \\ -s_m \tilde{\gamma}_m \end{pmatrix}, \end{aligned} \quad (4.3.142)$$

wobei o.B.d.A. $c_m \geq 0$ gelten soll. Analog zum GMRES-Verfahren liegt mit $\mathbf{R}_m$ eine reguläre Matrix vor, und wir erhalten

$$\boldsymbol{\alpha}_m = \arg \min_{\boldsymbol{\alpha} \in \mathbb{R}^m} J_m(\boldsymbol{\alpha}) = \mathbf{R}_m^{-1} \mathbf{g}_m.$$

Desweiteren gilt

$$\begin{pmatrix} \mathbf{R}_m \\ \mathbf{0}^T \end{pmatrix} = \bar{\mathbf{R}}_m = \begin{pmatrix} \mathbf{I} & & \\ & c_m & s_m \\ & -s_m & c_m \end{pmatrix} \begin{pmatrix} \mathbf{Q}_{m-1} \mathbf{T}_m \\ 0 \dots 0 h_{m+1,m} \end{pmatrix},$$

wodurch mit $c_m^2 + s_m^2 = 1$ die Gleichung

$$\mathbf{Q}_{m-1} \mathbf{T}_m = \begin{pmatrix} \mathbf{I} & \mathbf{0} \\ \mathbf{0}^T & c_m \end{pmatrix} \mathbf{R}_m$$

folgt. Mit (4.3.139) erhalten wir daher unter Berücksichtigung der vorausgesetzten Regularität der Matrix $\mathbf{T}_m$ die Eigenschaft $c_m \neq 0$ und folglich die Darstellung

$$\begin{aligned} \tilde{\boldsymbol{\alpha}}_m &= \|\mathbf{w}_1\|_2 \mathbf{T}_m^{-1} \mathbf{e}_1 \\ &= (\mathbf{Q}_{m-1} \mathbf{T}_m)^{-1} \bar{\mathbf{g}}_{m-1} \\ &= \mathbf{R}_m^{-1} \begin{pmatrix} \mathbf{I} & \mathbf{0} \\ \mathbf{0}^T & c_m^{-1} \end{pmatrix} \begin{pmatrix} \mathbf{g}_{m-1} \\ \tilde{\gamma}_m \end{pmatrix} \\ &= \mathbf{R}_m^{-1} \begin{pmatrix} \mathbf{R}_{m-1} \boldsymbol{\alpha}_{m-1} \\ \tilde{\gamma}_m c_m^{-1} \end{pmatrix}. \end{aligned} \quad (4.3.143)$$

Elementares Nachrechnen liefert die Form

$$\mathbf{R}_m = \begin{pmatrix} \mathbf{R}_{m-1} & \mathbf{0} \\ \mathbf{0}^T & r_{m,m} \end{pmatrix}.$$

Aus (4.3.142) erhalten wir

$$\mathbf{g}_m = \begin{pmatrix} \mathbf{g}_{m-1} \\ c_m \tilde{\gamma}_m \end{pmatrix},$$

wodurch mit $\mathbf{g}_{m-1} = \mathbf{R}_{m-1} \boldsymbol{\alpha}_{m-1}$ unter Berücksichtigung der obigen Gestalt der oberen Dreiecksmatrix $\mathbf{R}_m$ die behauptete Gestalt des Koeffizientenvektors durch

$$\begin{aligned} \boldsymbol{\alpha}_m &= \mathbf{R}_m^{-1} \mathbf{g}_m = \mathbf{R}_m^{-1} \begin{pmatrix} \mathbf{R}_{m-1} \boldsymbol{\alpha}_{m-1} \\ \tilde{\gamma}_m c_m \end{pmatrix} \\ &= \mathbf{R}_m^{-1} \left((1 - c_m^2) \begin{pmatrix} \mathbf{R}_{m-1} \boldsymbol{\alpha}_{m-1} \\ 0 \end{pmatrix} + c_m^2 \begin{pmatrix} \mathbf{R}_{m-1} \boldsymbol{\alpha}_{m-1} \\ \tilde{\gamma}_m c_m^{-1} \end{pmatrix} \right) \\ &= (1 - c_m^2) \begin{pmatrix} \boldsymbol{\alpha}_{m-1} \\ 0 \end{pmatrix} + c_m^2 \tilde{\boldsymbol{\alpha}}_m \end{aligned} \quad (4.3.144)$$

folgt. Wenden wir die Funktion J_m auf $\tilde{\boldsymbol{\alpha}}_m$ an, so erhalten wir unter Verwendung der Identität $c_m^{-2} - 1 = c_m^{-2} s_m^2$ die Gleichung

$$\begin{aligned} \|\mathbf{w}_{m+1}\|_2 &\stackrel{(4.3.140)}{=} J_m(\tilde{\boldsymbol{\alpha}}_m) = \left\| \mathbf{Q}_m \left[\begin{pmatrix} \|\mathbf{w}_1\|_2 \\ \mathbf{0} \end{pmatrix} - \overline{\mathbf{T}}_m \tilde{\boldsymbol{\alpha}}_m \right] \right\|_2 \\ &\stackrel{(4.3.141)}{=} \left\| \begin{pmatrix} \mathbf{g}_{m-1} \\ c_m \tilde{\gamma}_m \\ -s_m \tilde{\gamma}_m \end{pmatrix} - \begin{pmatrix} \mathbf{R}_m \tilde{\boldsymbol{\alpha}}_m \\ 0 \end{pmatrix} \right\|_2 \\ &\stackrel{(4.3.143)}{=} \left\| \begin{pmatrix} \mathbf{g}_{m-1} \\ c_m \tilde{\gamma}_m \\ -s_m \tilde{\gamma}_m \end{pmatrix} - \begin{pmatrix} \mathbf{g}_{m-1} \\ \tilde{\gamma}_m c_m^{-1} \\ 0 \end{pmatrix} \right\|_2 \\ &= \sqrt{\tilde{\gamma}_m^2 (c_m^2 - 2 + c_m^{-2} + s_m^2)} = |\tilde{\gamma}_m| |s_m| c_m^{-1}, \end{aligned}$$

wodurch mit (4.3.142) und (4.3.144)

$$\begin{aligned} \tau_m &= \left\| \begin{pmatrix} \mathbf{g}_{m-1} \\ c_m \tilde{\gamma}_m \\ -s_m \tilde{\gamma}_m \end{pmatrix} - \begin{pmatrix} \mathbf{R}_m \boldsymbol{\alpha}_m \\ 0 \end{pmatrix} \right\|_2 = |\tilde{\gamma}_m s_m| \\ &= \tau_{m-1} \frac{\|\mathbf{w}_{m+1}\|_2}{\tau_{m-1}} c_m = \tau_{m-1} \vartheta_m c_m \end{aligned}$$

mit $\vartheta_m = \|\mathbf{w}_{m+1}\|_2 \tau_{m-1}^{-1}$ folgt, und zudem wegen $s_m^2 + c_m^2 = 1$ unter Einbeziehung von $\tau_{m-1} = |\tilde{\gamma}_m|$

$$c_m = \sqrt{\frac{1}{1 + \frac{s_m^2}{c_m^2}}} = \sqrt{\frac{1}{1 + \vartheta_m^2}}$$

gilt. □

Satz 4.89 Seien $c_m, \vartheta_m, \mathbf{y}_m$ durch Lemma 4.85 und Satz 4.88 gegeben und $m \geq 1$. Dann läßt sich die TFQMR-Iterierte mittels

$$\eta_m := \alpha_{\lfloor \frac{m-1}{2} \rfloor} c_m^2 \quad (4.3.145)$$

und

$$\mathbf{d}_m := \mathbf{y}_m + \frac{\vartheta_{m-1}^2 \eta_{m-1}}{\alpha_{\lfloor \frac{m-1}{2} \rfloor}} \mathbf{d}_{m-1}$$

in der rekursiven Form

$$\mathbf{x}_m = \mathbf{x}_{m-1} + \eta_m \mathbf{d}_m$$

schreiben.

Beweis:

Definieren wir den Vektor

$$\tilde{\mathbf{x}}_m = \mathbf{x}_0 + \mathbf{Y}_m \tilde{\boldsymbol{\alpha}}_m$$

mit dem durch Satz 4.88 gegebenen Vektor $\tilde{\boldsymbol{\alpha}}_m$ und der Matrix $\mathbf{Y}_m$ gemäß (4.3.130), so folgt

$$\tilde{\mathbf{x}}_m = \mathbf{x}_0 + \mathbf{Y}_{m-1} \tilde{\boldsymbol{\alpha}}_{m-1} + \alpha_{\lfloor \frac{m-1}{2} \rfloor} \mathbf{y}_m = \tilde{\mathbf{x}}_{m-1} + \alpha_{\lfloor \frac{m-1}{2} \rfloor} \mathbf{y}_m. \quad (4.3.146)$$

Für die TFQMR-Iterierte $\mathbf{x}_m$ ergibt sich mit (4.3.136) und (4.3.137) die Rekursionsvorschrift

$$\begin{aligned} \mathbf{x}_m &= (1 - c_m^2) \left(\mathbf{x}_0 + \mathbf{Y}_m \begin{pmatrix} \boldsymbol{\alpha}_{m-1} \\ 0 \end{pmatrix} \right) + c_m^2 (\mathbf{x}_0 + \mathbf{Y}_m \tilde{\boldsymbol{\alpha}}_m) \\ &= (1 - c_m^2) \mathbf{x}_{m-1} + c_m^2 \tilde{\mathbf{x}}_m \\ &\stackrel{(4.3.145)}{=} \mathbf{x}_{m-1} + \frac{\eta_m}{\alpha_{\lfloor \frac{m-1}{2} \rfloor}} (\tilde{\mathbf{x}}_m - \mathbf{x}_{m-1}). \end{aligned} \quad (4.3.147)$$

Legen wir den Vektor $\tilde{\mathbf{d}}_m$ gemäß

$$\tilde{\mathbf{d}}_m := \frac{\tilde{\mathbf{x}}_m - \mathbf{x}_{m-1}}{\alpha_{\lfloor \frac{m-1}{2} \rfloor}} \quad (4.3.148)$$

fest, dann folgt die Gleichung

$$\begin{aligned}
\tilde{\mathbf{d}}_m &\stackrel{(4.3.146)}{=} \mathbf{y}_m + \frac{1}{\alpha_{\lfloor \frac{m-1}{2} \rfloor}} \left(\tilde{\mathbf{x}}_{m-1} - \mathbf{x}_{m-2} - \eta_{m-1} \tilde{\mathbf{d}}_{m-1} \right) \\
&\stackrel{(4.3.145)}{=} \stackrel{(4.3.148)}{=} \mathbf{y}_m + \frac{1}{\alpha_{\lfloor \frac{m-1}{2} \rfloor}} \left(\frac{1}{c_{m-1}^2} \eta_{m-1} \tilde{\mathbf{d}}_{m-1} - \eta_{m-1} \tilde{\mathbf{d}}_{m-1} \right) \\
&\stackrel{(4.3.138)}{=} \mathbf{y}_m + \frac{\vartheta_{m-1}^2 \eta_{m-1}}{\alpha_{\lfloor \frac{m-1}{2} \rfloor}} \tilde{\mathbf{d}}_{m-1},
\end{aligned}$$

wodurch wir $\mathbf{d}_m$ mit $\tilde{\mathbf{d}}_m$ gleichsetzen können. Das Einsetzen von (4.3.148) in (4.3.147) liefert somit die behauptete Darstellung der TFQMR-Iterierten. $\square$

Betrachten wir die im CGS-Algorithmus auftretenden Vektoren $\mathbf{q}_j, \mathbf{u}_{j+1}$ und $\mathbf{p}_{j+1}$, so erhalten wir unter Verwendung der Lemmata 4.85 und 4.86 sowie der Hilfsvariablen

$$\mathbf{v}_j := \mathbf{A} \mathbf{p}_j$$

die Rekursionsvorschriften

$$\begin{aligned}
\mathbf{y}_{2j} &= \mathbf{q}_{j-1} = \mathbf{u}_{j-1} - \alpha_{j-1} \mathbf{v}_{j-1} = \mathbf{y}_{2j-1} - \alpha_{j-1} \mathbf{v}_{j-1}, \\
\mathbf{y}_{2j+1} &= \mathbf{u}_j = \mathbf{r}_j + \beta_{j-1} \mathbf{q}_{j-1} = \mathbf{w}_{2j+1} + \beta_{j-1} \mathbf{y}_{2j}.
\end{aligned}$$

Zudem gilt

$$\mathbf{p}_j = \mathbf{u}_j + \beta_{j-1} (\mathbf{q}_{j-1} + \beta_{j-1} \mathbf{p}_{j-1}) = \mathbf{y}_{2j+1} + \beta_{j-1} (\mathbf{y}_{2j} + \beta_{j-1} \mathbf{p}_{j-1}),$$

so daß für die Hilfsvariable

$$\mathbf{v}_j = \mathbf{A} \mathbf{y}_{2j+1} + \beta_{j-1} (\mathbf{A} \mathbf{y}_{2j} + \beta_{j-1} \mathbf{v}_{j-1})$$

folgt. Mit den Lemmata 4.86 und 4.87 sowie den Sätzen 4.88 und 4.89 erhalten wir zusammenfassend den TFQMR-Algorithmus in der folgenden Formulierung, dessen Umsetzung in MATLAB im Anhang A aufgeführt ist.

TFQMR-Algorithmus —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$ und $\varepsilon > 0$	
$\mathbf{w}_1 = \mathbf{y}_1 = \mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0, \tau_0 := \ \mathbf{r}_0\ _2$	
Y	$\tau_0 > \varepsilon$ N
$\mathbf{v}_0 := \mathbf{A}\mathbf{y}_1, \mathbf{d}_0 := \mathbf{0}, \eta_0 = \vartheta_0 := 0, j := 1$	
$\alpha_{j-1} := \frac{(\mathbf{w}_{2j-1}, \mathbf{r}_0)_2}{(\mathbf{v}_{j-1}, \mathbf{r}_0)_2}, \mathbf{y}_{2j} := \mathbf{y}_{2j-1} - \alpha_{j-1} \mathbf{v}_{j-1}$	
$m := 2j - 1$	
$\mathbf{w}_{m+1} := \mathbf{w}_m - \alpha_{j-1} \mathbf{A}\mathbf{y}_m$	
$\vartheta_m := \frac{\ \mathbf{w}_{m+1}\ _2}{\tau_{m-1}}, c_m := (1 + \vartheta_m^2)^{-1/2}$	
$\mathbf{d}_m := \mathbf{y}_m + \frac{\vartheta_{m-1}^2 \eta_{m-1}}{\alpha_{j-1}} \mathbf{d}_{m-1}, \eta_m := c_m^2 \alpha_{j-1}$	
$\mathbf{x}_m := \mathbf{x}_{m-1} + \eta_m \mathbf{d}_m, \tau_m := \tau_{m-1} \vartheta_m c_m$	
Y	$\sqrt{m+1} \tau_m > \varepsilon$ N
$m := m + 1$	STOP
solange $m < 2j + 1$	
$\beta_{j-1} := \frac{(\mathbf{w}_{2j+1}, \mathbf{r}_0)_2}{(\mathbf{w}_{2j-1}, \mathbf{r}_0)_2}$	
$\mathbf{y}_{2j+1} := \mathbf{w}_{2j+1} + \beta_{j-1} \mathbf{y}_{2j}$	
$\mathbf{v}_j := \mathbf{A}\mathbf{y}_{2j+1} + \beta_{j-1} (\mathbf{A}\mathbf{y}_{2j} + \beta_{j-1} \mathbf{v}_{j-1})$	
$j := j + 1$	
solange $1 = 1$	

4.3.2.9 Das QMRCGSTAB-Verfahren

Motiviert durch die Entwicklung der TFQMR-Methode aus dem CGS-Algorithmus veröffentlichten Chan et al. das QMRCGSTAB-Verfahren [16], das sich unter Verwendung des Prinzips der Quasi-Minimierung aus dem BiCGSTAB-Verfahren herleiten läßt. Hierzu

definieren wir mit dem folgenden Lemma eine Basis des Krylov-Unterraums K_m unter Nutzung der durch (4.3.126) gegebenen Abbildung.

Lemma 4.90 *Vorausgesetzt, das BiCGSTAB-Verfahren bricht nicht vor der Berechnung der Vektoren $\mathbf{p}_{\lfloor \frac{m-1}{2} \rfloor}$ und $\mathbf{s}_{\lfloor \frac{m-1}{2} \rfloor}$ ab, dann stellen die durch*

$$\mathbf{y}_k := \begin{cases} \mathbf{p}_{\lfloor \frac{k-1}{2} \rfloor} & , \text{ falls } k \text{ ungerade,} \\ \mathbf{s}_{\lfloor \frac{k-1}{2} \rfloor} & , \text{ falls } k \text{ gerade} \end{cases} \quad (4.3.149)$$

definierten Vektoren $\mathbf{y}_1, \dots, \mathbf{y}_m$ eine Basis des K_m dar.

Beweis:

Der Beweis ergibt sich analog zum Nachweis des Lemmas 4.85 unter Verwendung der Gleichungen (4.3.117) und (4.3.118). $\square$

Dem Lemma 4.86 entsprechend erhalten wir das folgende Lemma.

Lemma 4.91 *Seien*

$$\mathbf{w}_k := \begin{cases} \mathbf{r}_{\lfloor \frac{k-1}{2} \rfloor} & , \text{ falls } k \text{ ungerade,} \\ \mathbf{s}_{\lfloor \frac{k-1}{2} \rfloor} & , \text{ falls } k \text{ gerade} \end{cases} \quad (4.3.150)$$

und

$$\delta_k := \begin{cases} \alpha_{\lfloor \frac{k-1}{2} \rfloor} & , \text{ falls } k \text{ ungerade,} \\ \omega_{\lfloor \frac{k-1}{2} \rfloor} & , \text{ falls } k \text{ gerade,} \end{cases}$$

dann gilt

$$\mathbf{w}_{k+1} = \mathbf{w}_k - \delta_k \mathbf{A} \mathbf{y}_k \text{ f\"ur } k = 1, \dots, m$$

mit den durch (4.3.149) gegebenen Vektoren $\mathbf{y}_1, \dots, \mathbf{y}_m$.

Beweis:

Mit den Gleichungen (4.3.114), (4.3.116) und (4.3.118) läßt sich die Behauptung durch einfaches Nachrechnen beweisen. $\square$

Definieren wir die Matrizen $\mathbf{Y}_m$ und $\mathbf{W}_{m+1}$ gemäß (4.3.130) und (4.3.131) mit den durch (4.3.149) und (4.3.150) vorliegenden Vektoren, so läßt sich der Residuenvektor in der Form

$$\mathbf{r}_m = \mathbf{W}_{m+1} \mathbf{S}_{m+1}^{-1} \mathbf{S}_{m+1} (\mathbf{e}_1 - \overline{\mathbf{B}}_m \boldsymbol{\alpha}_m)$$

mit $\mathbf{S}_{m+1} = \text{diag}\{\|\mathbf{w}_1\|_2, \dots, \|\mathbf{w}_{m+1}\|_2\}$,

$$\overline{\mathbf{B}}_m = \begin{pmatrix} \delta_1^{-1} & & & & \\ -\delta_1^{-1} & \delta_2^{-1} & & & \\ & \ddots & \ddots & & \\ & & & -\delta_{m-1}^{-1} & \delta_m^{-1} \\ & & & & -\delta_m^{-1} \end{pmatrix} \in \mathbb{R}^{(m+1) \times m}$$

und $\mathbf{e}_1 = (1, 0, \dots, 0)^T \in \mathbb{R}^{m+1}$ darstellen. Der QMRCGSTAB-Algorithmus unterscheidet sich vom TFQMR-Verfahren einzig in der vorliegenden Matrix $\overline{\mathbf{B}}_m$, die bei beiden Verfahren zwar die gleiche Struktur, jedoch unterschiedliche Einträge aufweist. Beim QMRCGSTAB-Verfahren wird daher ebenfalls die Abbildung J_m gemäß (4.3.135) eingeführt und die QMRCGSTAB-Iterierte mittels (4.3.136) definiert, wodurch für die Norm des Residuenvektors ebenfalls die durch Lemma 4.87 gegebene Abschätzung gilt. Zudem können die Sätze 4.88 und 4.89 direkt für das vorliegende QMRCGSTAB-Verfahren umgeschrieben werden, und wir erhalten mit den beiden folgenden Sätzen die mathematische Grundlage der Methode.

Satz 4.92 *Sei $m \geq 1$ und desweiteren für $k = m - 1, m$*

$$\boldsymbol{\alpha}_k = \arg \min_{\boldsymbol{\alpha} \in \mathbb{R}^k} J_k(\boldsymbol{\alpha}) \text{ und } \tau_k := J_k(\boldsymbol{\alpha}_k)$$

mit der durch (4.3.135) gegebenen Funktion J_k . Dann gilt

$$\boldsymbol{\alpha}_m = (1 - c_m^2) \begin{pmatrix} \boldsymbol{\alpha}_{m-1} \\ 0 \end{pmatrix} + c_m^2 \tilde{\boldsymbol{\alpha}}_m$$

mit $\tilde{\boldsymbol{\alpha}}_m := (\delta_1, \delta_2, \dots, \delta_m)^T \in \mathbb{R}^m$ sowie $\tau_m = \tau_{m-1} \vartheta_m c_m$ mit $\vartheta_m = \frac{\|\mathbf{w}_{m+1}\|_2}{\tau_{m-1}}$

und $c_m = (1 + \vartheta_m^2)^{-1/2}$.

Beweis:

Der Beweis ergibt sich analog zum Nachweis des Satzes 4.88. □

Satz 4.93 *Seien $c_m, \vartheta_m, \mathbf{y}_m$ durch Lemma 4.90 und Satz 4.92 gegeben und $m \geq 1$. Dann läßt sich die QMRCGSTAB-Iterierte mittels*

$$\eta_m := \delta_m c_m^2$$

und

$$\mathbf{d}_m := \mathbf{y}_m + \frac{\vartheta_{m-1}^2 \eta_{m-1}}{\delta_m} \mathbf{d}_{m-1}$$

in der rekursiven Form

$$\mathbf{x}_m = \mathbf{x}_{m-1} + \eta_m \mathbf{d}_m$$

schreiben.

Beweis:

Die Behauptung läßt sich analog zum Beweis des Satzes 4.89 nachweisen. □

Berücksichtigen wir die spezielle Darstellung der Vektoren $\mathbf{y}_1, \dots, \mathbf{y}_m$ und $\mathbf{w}_1, \dots, \mathbf{w}_m$, so läßt sich der QMRCGSTAB-Algorithmus in der folgenden Form schreiben. Dem Anhang A kann eine Implementierung in MATLAB entnommen werden.

QMRCGSTAB-Algorithmus —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$ und $\varepsilon > 0$		
$\mathbf{p}_0 = \mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0, \tau_0 := \ \mathbf{r}_0\ _2, \mathbf{v}_0 := \mathbf{A}\mathbf{p}_0$		
Y	$\tau_0 > \varepsilon$	N
$\mathbf{d}_0 := \mathbf{0}, \eta_0 = \vartheta_0 := 0, j := 1$		
$\alpha_{j-1} := \frac{(\mathbf{r}_{j-1}, \mathbf{r}_0)_2}{(\mathbf{v}_{j-1}, \mathbf{r}_0)_2}, \mathbf{s}_{j-1} := \mathbf{r}_{j-1} - \alpha_{j-1} \mathbf{v}_{j-1}$		
$\tilde{\vartheta}_j := \frac{\ \mathbf{s}_{j-1}\ _2}{\tau_{j-1}}, \tilde{c}_j := \left(1 + \tilde{\vartheta}_j^2\right)^{-1/2}$		
$\tilde{\mathbf{d}}_j := \mathbf{p}_{j-1} + \frac{\vartheta_{j-1}^2 \eta_{j-1}}{\alpha_{j-1}} \mathbf{d}_{j-1}, \tilde{\eta}_j := \tilde{c}_j^2 \alpha_{j-1}$		
$\tilde{\mathbf{x}}_j := \mathbf{x}_{j-1} + \tilde{\eta}_j \tilde{\mathbf{d}}_j, \tilde{\tau}_j := \tau_{j-1} \tilde{\vartheta}_j \tilde{c}_j$		
$\omega_{j-1} := \frac{(\mathbf{A}\mathbf{s}_{j-1}, \mathbf{s}_{j-1})_2}{(\mathbf{A}\mathbf{s}_{j-1}, \mathbf{A}\mathbf{s}_{j-1})_2}, \mathbf{r}_j := \mathbf{s}_{j-1} - \omega_{j-1} \mathbf{A}\mathbf{s}_{j-1}$		
$\vartheta_j := \frac{\ \mathbf{r}_j\ _2}{\tilde{\tau}_j}, c_j := \left(1 + \vartheta_j^2\right)^{-1/2}$		
$\mathbf{d}_j := \mathbf{s}_{j-1} + \frac{\tilde{\vartheta}_j^2 \tilde{\eta}_j}{\omega_{j-1}} \tilde{\mathbf{d}}_j, \eta_j := c_j^2 \omega_{j-1}$		
$\mathbf{x}_j := \tilde{\mathbf{x}}_j + \eta_j \mathbf{d}_j, \tau_j := \tilde{\tau}_j \vartheta_j c_j$		
Y	$\sqrt{2j+1} \tau_j > \varepsilon$	N
$\beta_{j-1} := \frac{\alpha_{j-1}}{\omega_{j-1}} \frac{(\mathbf{r}_j, \mathbf{r}_0)_2}{(\mathbf{r}_{j-1}, \mathbf{r}_0)_2}$		STOP
$\mathbf{p}_j := \mathbf{r}_j + \beta_{j-1} (\mathbf{p}_{j-1} - \omega_{j-1} \mathbf{A}\mathbf{p}_{j-1})$		
$\mathbf{v}_j := \mathbf{A}\mathbf{p}_j, j = j + 1$		
solange 1 = 1		

4.3.2.10 Konvergenzanalysen

In diesem Abschnitt untersuchen wir das Konvergenzverhalten der für reguläre Matrizen anwendbaren Verfahren GMRES, CGS, BiCGSTAB, TFQMR und QMRCGSTAB. Zur

Analyse nutzen wir die im Beispiel 1.4 hergeleitete unsymmetrische Matrix (1.0.9), wobei für den Diffusionsparameter der Wert $\varepsilon = 0.1$ gewählt wurde und mit $N = 100$ stets 10.000 Unbekannte vorliegen.

Alle präsentierten Abbildungen zeigen den Logarithmus des Residuums zur Basis 10 aufgetragen über der Iterationzahl. Aufgrund der vorliegenden Gestalt des Gleichungssystems benötigt eine Matrix-Vektor-Multiplikation stets $5N^2 - 4N$ Operationen, die folglich stets in der Größenordnung einer Skalarmultiplikation mit N^2 Operationen liegt. Die Verfahren CGS, BiCGSTAB, TFQMR und QMRGSTAB benötigen hierbei pro Iteration jeweils zwei Matrix-Vektor-Multiplikationen. Dagegen weist das GMRES-Verfahren lediglich eine derartige Operationen pro Iteration auf. Jedoch zeigt sich durch den im GMRES-Verfahren enthaltenen Arnoldi-Algorithmus zur Berechnung der Orthonormalbasis eine von der Dimension des Krylov-Unterraums abhängige Anzahl an Skalarprodukten. Das GMRES-Verfahren wurde in unserem Fall in einer Restarted-Version mit einer maximalen Krylov-Unterraumdimension von $m = 30$ verwendet. Eine kleinere obere Grenze führte zu einer Erhöhung der Iterationszahl und wurde daher vermieden. Innerhalb eines Restarts müssen demzufolge 465 Skalarprodukte im Rahmen des Arnoldi-Algorithmus berechnet werden, die etwa 3 Matrix-Vektor-Multiplikationen pro Iteration entsprechen. Alle weiteren Verfahren beinhalten keinerlei iterationsabhängige Anzahl von Skalarprodukten. Je Iteration werden 3 (CGS), 4 (BiCGSTAB, TFQMR) beziehungsweise 6 (QMRGSTAB) Skalarprodukte berechnet, so daß diese Verfahren einen vergleichbaren Aufwand innerhalb eines Schleifendurchlaufs aufweisen. Dagegen ergibt sich bei GMRES-Verfahren durchschnittlich ein deutlicher höherer Rechenbedarf.

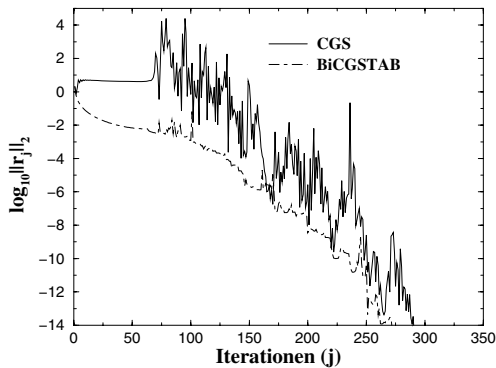


Bild 4.21 Konvergenzverläufe des CGS- und BiCGSTAB-Verfahrens

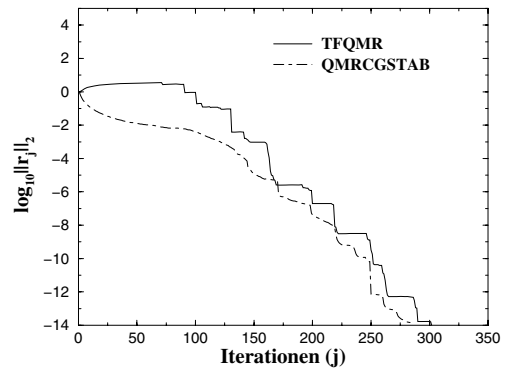


Bild 4.22 Konvergenzverläufe des TFQMR- und QMRGSTAB-Verfahrens

Die dargestellten Konvergenzverläufe zeigen das zu erwartende typische Verhalten der Gleichungssystemlöser. Während das CGS-Verfahren sehr starke Oszillationen aufweist, ergibt sich bei der BiCGSTAB-Methode durch die eingeführte eindimensionale Residuenminimierung ein deutlich glatteres Abfallverhalten (siehe Bild 4.21). Die Nutzung der Quasiminimierung innerhalb des TFQMR- und QMRGSTAB-Verfahrens führt, wie in Bild 4.22 zu sehen, auf ein annähernd monotonen und oszillationsfreies Konvergenzverhalten. Ein durchgehend monotonen Abklingverhalten des Residuums ergibt sich aus theoretischer Sicht für das GMRES-Verfahren und auch dessen GMRES(m)-Variante. Diese Eigenschaft wird auch durch die Abbildung 4.23 wiedergegeben.

Aus den ermittelten Konvergenzverläufen wird ersichtlich, daß bei allen betrachteten Gleichungssystemlösern eine Reduktion des Residuums um 14 Größenordnungen in weniger als 1000 Schleifendurchläufen erzielt wurde. Hierbei ergab sich für die Verfahren BiCGSTAB (272 Iterationen), CGS (291 Iterationen), TFQMR (302 Iterationen), QMRGCGSTAB (286 Iterationen) ein vergleichbarer Aufwand. Einzig das GMRES-Verfahren benötigte mit 838 Schleifendurchläufen eine deutlich höhere Iterationszahl.

Im Kontext der Euler- und Navier-Stokes-Gleichungen wird in [48] von einem ähnlichen Verhalten der Gleichungssystemlöser berichtet, wobei sich das GMRES-Verfahren ohne Nutzung einer geeigneten Präkonditionierung sogar teilweise als unpraktikabel erwies.

Die auf dem Bi-Lanczos-Algorithmus beruhenden Verfahren CGS, BiCGSTAB, TFQMR und QMRGCGSTAB haben sich in vielen praxisrelevanten Problemstellungen als robust und stabil herausgestellt. Es bleibt hierbei allerdings zu erwähnen, daß diese vier Methoden die möglichen vorzeitigen Verfahrensabbrüche des Bi-Lanczos-Algorithmus erben, und daher die Konvergenz auch bei Vernachlässigung möglicher Rundungsfehler nicht garantiert werden kann.

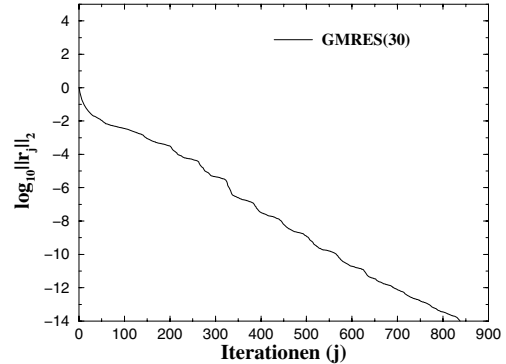


Bild 4.23 Konvergenzverlauf des GMRES(30)-Verfahrens

4.4 Übungsaufgaben

Aufgabe 1:

Zeigen Sie, dass das Jacobi-Verfahren und das Gauss-Seidel-Verfahren bei einer strikt diagonal dominanten Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ konvergieren.

Hinweis: Eine Matrix $\mathbf{A} = (a_{ij})_{i,j=1,\dots,n} \in \mathbb{R}^{n \times n}$ heißt streng diagonal dominant, wenn

$$|a_{kk}| > \sum_{\substack{i=1 \\ i \neq k}}^n |a_{ki}|, \quad k = 1, \dots, n$$

gilt.

Aufgabe 2:

Seien $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$ mit $\mathbf{y} \neq \mathbf{0}$. Sei $\psi: \mathbb{R} \rightarrow \mathbb{R}$ definiert durch

$$\psi(a) := \|\mathbf{x} - a\mathbf{y}\|_2.$$

Zeigen Sie:

$$\frac{\mathbf{x}^T \mathbf{y}}{\mathbf{y}^T \mathbf{y}} = \arg \min_{a \in \mathbb{R}} \psi(a).$$

Aufgabe 3:

Sei $\mathbf{A} \in \mathbb{C}^{n \times n}$ hermitesch und positiv definit. Zeigen Sie, dass

$$\begin{aligned} \|\cdot\|_{\mathbf{A}} : \mathbb{C}^n &\rightarrow \mathbb{R} \\ \mathbf{x} &\mapsto \|\mathbf{x}\|_{\mathbf{A}} := \|\mathbf{A}^{1/2}\mathbf{x}\|_2 \end{aligned}$$

eine Norm auf $\mathbb{C}^n$ darstellt.

Aufgabe 4:

Zeigen Sie, dass die Matrix

$$\mathbf{A} = \begin{pmatrix} 2 & -1 & & & \\ -1 & \ddots & \ddots & & \\ & \ddots & \ddots & \ddots & \\ & & \ddots & \ddots & -1 \\ & & & -1 & 2 \end{pmatrix} \in \mathbb{R}^{n \times n}$$

irreduzibel ist, ohne die vorgestellte graphentheoretische Vorgehensweise auszunutzen.

Aufgabe 5:

Gegeben sei die Matrix

$$\mathbf{A} = \begin{pmatrix} 1 - \frac{7}{8} \cos \frac{\pi}{4} & \frac{7}{8} \sin \frac{\pi}{4} \\ -\frac{7}{8} \sin \frac{\pi}{4} & 1 - \frac{7}{8} \cos \frac{\pi}{4} \end{pmatrix}$$

und der Vektor

$$\mathbf{b} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}.$$

Zeigen Sie, dass das lineare Iterationsverfahren

$$\mathbf{x}_{n+1} = (\mathbf{I} - \mathbf{A})\mathbf{x}_n + \mathbf{b}, \quad n = 0, 1, 2, \dots$$

für beliebigen Startvektor $\mathbf{x}_0 \in \mathbb{R}^2$ gegen die eindeutig bestimmte Lösung $\mathbf{x} = (0, 0)^T$ konvergiert, obwohl eine induzierte Matrixnorm mit

$$\|\mathbf{I} - \mathbf{A}\| > 1$$

existiert. Veranschaulichen Sie beide Sachverhalte zudem graphisch.

Aufgabe 6:

Geben Sie ein Beispiel an, bei dem die linearen Iterationsverfahren ψ und ϕ nicht konvergieren und die Produktiteration $\psi \circ \phi$ konvergiert.

Aufgabe 7:

Das Einzelschrittverfahren zur Lösung des Gleichungssystems $\mathbf{A}\mathbf{x} = \mathbf{y}$ ist für symmetrische, reelle, positiv definite $n \times n$ Matrizen $\mathbf{A}$ konvergent. Beweisen Sie die Aussage für $n = 2$ ohne auf den allgemeinen Beweis (Satz 4.16) zurückzugreifen.

Aufgabe 8:

Eine Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ heißt M -Matrix, falls

- $a_{ii} > 0, i = 1, \dots, n,$
- $a_{ij} \leq 0, i, j = 1, \dots, n, i \neq j$ und
- $\mathbf{A}$ ist regulär mit $\mathbf{A}^{-1} \geq \mathbf{0}$

gilt.

- (a) Geben Sie ein Beispiel für eine M -Matrix an, die keine Diagonalmatrix darstellt.
- (b) Zeigen Sie für eine M -Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$:
 - (1) Sei $\mathbf{A}\mathbf{x} = \mathbf{b}$ und $\mathbf{A}\mathbf{x}' = \mathbf{b}'$, dann folgt aus $\mathbf{b} \leq \mathbf{b}'$ auch $\mathbf{x} \leq \mathbf{x}'$.
 - (2) Für die zu $\mathbf{A}$ gehörige Iterationsmatrix des Jacobi-Verfahrens gilt

$$\mathbf{M}_J \geq \mathbf{0}$$

und

$$\rho(\mathbf{M}_J) < 1.$$

Hinweis: Die Relationen $\mathbf{A} \geq \mathbf{0}$ respektive $\mathbf{x} \geq \mathbf{0}$ sind stets komponentenweise zu verstehen. Nutzen Sie die Aussage, dass für $\mathbf{A} \geq \mathbf{0}$ stets zu $\lambda = \rho(\mathbf{A})$ ein Eigenvektor $\mathbf{x} \geq \mathbf{0}$ gehört.

Aufgabe 9:

Wir betrachten Matrizen $\mathbf{A} = (a_{ij})_{i,j=1,\dots,n} \in \mathbb{R}^{n \times n}$ ($n \geq 2$) mit $a_{ii} = 1, i = 1, \dots, n$ und $a_{ij} = a$ für $i \neq j$.

- (a) Wie sieht die Iterationsmatrix des Jacobi-Verfahrens zu $\mathbf{A}$ aus? Berechnen Sie ihre Eigenwerte und die Eigenwerte von $\mathbf{A}$.
- (b) Für welche a konvergiert das Jacobi-Verfahren? Für welche a ist $\mathbf{A}$ positiv definit?
- (c) Gibt es positiv definite Matrizen $\mathbf{A} \in \mathbb{R}^{n \times n}$ für die das Jacobi-Verfahren nicht konvergiert?

Aufgabe 10:

Sei $U := \text{span}\{\mathbf{u}_1, \dots, \mathbf{u}_m\}$ ein Unterraum des $\mathbb{R}^n$

- (a) Zeigen Sie: Gilt für $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n, \mathbf{x} - \mathbf{y} \in U$, dann folgt

$$\text{span}\{U, \mathbf{x}\} = \text{span}\{U, \mathbf{y}\}.$$

- (b) Sei $\mathbf{A} \in \mathbb{R}^{n \times n}$, dann gilt

$$\mathbf{A}U := \{\mathbf{A}\mathbf{x} \mid \mathbf{x} \in U\}.$$

Zeigen Sie:

$$\mathbf{A}U = \text{span}\{\mathbf{A}\mathbf{u}_1, \dots, \mathbf{A}\mathbf{u}_m\}.$$

Aufgabe 11:

Beweisen Sie:

- (a) Wenn $\mathbf{A}$ mindestens einen positiven und einen negativen reellen Eigenwert besitzt, dann divergiert das Richardson-Verfahren für jede Wahl von $\Theta \in \mathbb{C}$.
- (b) Wenn $\mathbf{A}$ unter anderem zwei komplexe Eigenwerte λ_1, λ_2 mit entgegengesetztem Vorzeichen besitzt ($\lambda_1/|\lambda_1| = -\lambda_2/|\lambda_2|$), so divergiert das Richardson-Verfahren für jede Wahl von $\Theta \in \mathbb{C}$.
- (c) Das Spektrum von $\mathbf{A}$ liege in einem abgeschlossenen Kreis um $\mu \in \mathbb{C} \setminus \{0\}$ mit dem Radius $r < |\mu|$, dann führt die Wahl $\Theta = 1/\mu$ zur Konvergenz des Richardson-Verfahrens mit

$$\rho(\mathbf{I} - \Theta \mathbf{A}) \leq r/|\mu| < 1.$$

Aufgabe 12:

Bestimmen Sie für das System

$$\begin{pmatrix} 4 & 0 & 2 \\ 0 & 5 & 2 \\ 5 & 4 & 10 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} 4 \\ -3 \\ 2 \end{pmatrix}$$

die Spektralradien der Iterationsmatrizen für das Gesamt- und Einzelschrittverfahren und schreiben Sie beide Verfahren in Komponenten. Zeigen Sie, dass die Matrix konsistent geordnet ist und bestimmen Sie den optimalen Relaxationsparameter für das Einzelschrittverfahren.

Aufgabe 13:

Gegeben sei die Iterationsvorschrift

$$\mathbf{x}_{k+1} = \mathbf{x}_k + a_k(\mathbf{b} - \mathbf{A}\mathbf{x}_k), \quad a_k \in \mathbb{R} \quad (4.4.1)$$

zur Lösung der Gleichung $\mathbf{A}\mathbf{x} = \mathbf{b}$ mit $\mathbf{A} \in \mathbb{R}^{n \times n}$ und $\mathbf{b} \in \mathbb{R}^n$.

- (a) Schreiben Sie den Residuenvektor $\mathbf{r}_{k+1} = \mathbf{b} - \mathbf{A}\mathbf{x}_{k+1}$ als Funktion von $\mathbf{A}, \mathbf{r}_k$ und a_k , das heißt

$$\mathbf{r}_{k+1} = \mathbf{r}_{k+1}(\mathbf{A}, \mathbf{r}_k, a_k).$$

- (b) Berechnen Sie $a_k \in \mathbb{R}$ derart, dass $\mathbf{r}_{k+1}$ minimal bezüglich der euklidischen Norm ist.

- (c) Sei $\mathcal{F}(\mathbf{A}) := \left\{ \frac{\mathbf{y}^T \mathbf{A} \mathbf{y}}{\mathbf{y}^T \mathbf{y}} : \mathbf{y} \in \mathbb{R}^n \setminus \{0\} \right\}$.

Zeigen Sie: Für die Iterationsvorschrift (4.4.1) mit $a_k = \frac{(\mathbf{r}_k, \mathbf{A}\mathbf{r}_k)}{(\mathbf{A}\mathbf{r}_k, \mathbf{A}\mathbf{r}_k)}$ gilt für beliebiges $\mathbf{r}_k \neq \mathbf{0}$

$$\|\mathbf{r}_{k+1}\|_2 \leq \|\mathbf{r}_k\|_2$$

genau dann, wenn $\mathbf{0} \notin \mathcal{F}(\mathbf{A}^T)$ gilt.

Aufgabe 14:

Gegeben sei das Gleichungssystem $\mathbf{Ax} = \mathbf{b}$ mit

$$\mathbf{A} = \begin{pmatrix} 0.78 & 0.563 \\ 0.913 & 0.659 \end{pmatrix}, \quad \mathbf{b} = \begin{pmatrix} 0.217 \\ 0.254 \end{pmatrix}.$$

- (a) Wie lautet die exakte Lösung?
- (b) Berechnen Sie für die beiden Näherungslösungen $\mathbf{u} = (0.999, -1.001)^T$ und $\mathbf{v} = (0.341, -0.087)^T$ die Residuen $r(\mathbf{u})$ und $r(\mathbf{v})$. Hat die „genauere“ Lösung $\mathbf{u}$ das kleinere Residuum? Erklären Sie die Diskrepanz durch Betrachtung der Residuenfunktion r .

Hinweis: Das Residuum $r(\mathbf{z})$ einer Näherungslösung $\mathbf{z}$ der Gleichung $\mathbf{Ax} = \mathbf{b}$ ist durch $r(\mathbf{z}) = \|\mathbf{b} - \mathbf{Az}\|_2$ definiert.

Aufgabe 15:

Das Gleichungssystem

$$\begin{pmatrix} 3 & -1 \\ -1 & 3 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} 1 \\ -1 \end{pmatrix}$$

soll mit dem Jacobi- und Gauß-Seidel-Verfahren gelöst werden. Wieviele Iterationen sind jeweils ungefähr erforderlich, um den Fehler $\|\mathbf{x}_n - \mathbf{x}\|_2$ um den Faktor 10^{-6} zu reduzieren?

Aufgabe 16:

Zeigen Sie, dass das Verfahren des steilsten Abstiegs konsistent jedoch nicht linear ist.

Aufgabe 17:

Betrachten Sie das Verfahren der konjugierten Richtungen und konstruieren Sie eine Matrix $\mathbf{A} \in \mathbb{R}^{4 \times 4}$ sowie vier $\mathbf{A}$ -orthogonale Suchrichtungen $\mathbf{p}_0, \dots, \mathbf{p}_3$ derart, dass für gegebene rechte Seite $\mathbf{b} = (10, 9, 8, 7)^T$ und gegebenen Startvektor $\mathbf{x}_0 = (1, 2, 3, 4)^T$ der durch das Verfahren ermittelte Fehlervektor

$$\mathbf{e}_m = \mathbf{x}_m - \mathbf{A}^{-1}\mathbf{b}$$

die Bedingung

$$\|\mathbf{e}_m\|_2 \geq 2,749 \quad \text{für } m = 0, 1, 2, 3$$

erfüllt.

Aufgabe 18:

Zeigen Sie: Sei $\mathbf{A} \in \mathbb{R}^{n \times n}$ symmetrisch und positiv definit und gelte $\mathbf{A}^j = \mathbf{A}$ für ein $j \in \mathbb{N}$ mit $n > j > 1$, dann liefert das CG-Verfahren spätestens mit $\mathbf{x}_j$ die exakte Lösung der Gleichung $\mathbf{Ax} = \mathbf{b}$ für beliebiges $\mathbf{b} \in \mathbb{R}^n$.

Aufgabe 19:

Zeigen Sie: Seien $\mathbf{A} \in \mathbb{R}^{n \times n}$ und $\boldsymbol{\xi} \in \mathbb{R}^n$ gegeben. Desweiteren seien $\phi_j, \psi_j \in \mathcal{P}_j$, dann gilt

$$\left(\phi_j(\mathbf{A}^T)\boldsymbol{\eta}, \psi_j(\mathbf{A})\boldsymbol{\xi} \right)_2 = 0 \quad \forall \boldsymbol{\eta} \in \mathbb{R}^n,$$

falls $\boldsymbol{\xi} \in \text{Kern}(\phi_j(\mathbf{A})\psi_j(\mathbf{A}))$ ist.

Aufgabe 20:

Gegeben sei das lineare Gleichungssystem

$$\begin{pmatrix} a & 14 \\ 7 & 50 \end{pmatrix} \mathbf{x} = \begin{pmatrix} 1 \\ 1 \end{pmatrix} \quad \text{mit } a \in \mathbb{R}^+ \setminus \left\{ \frac{49}{25} \right\}.$$

- (a) Für welche Werte von a konvergiert das Jacobi-Verfahren respektive das Gauss-Seidel-Verfahren?
- (b) Lässt sich die Konvergenzgeschwindigkeit durch Relaxation des Gauß-Seidel-Verfahrens erhöhen?

5 Präkonditionierer

Den Einfluß der Konditionszahl auf die Eigenschaften des Gleichungssystems und den Zusammenhang zwischen dem Fehler- und Residuenvektor haben wir bereits ausführlich im Abschnitt 2.3 studiert. Die erzielten Resultate (Sätze 2.39 und 2.40), wie auch die für das Verfahren des steilsten Abstiegs (Satz 4.53), das CG-Verfahren (Satz 4.65) und die GMRES-Methode (Korollar 4.78) vorliegenden Konvergenzaussagen, verdeutlichen nachdrücklich den Vorteil einer kleinen Konditionszahl der Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ des linearen Gleichungssystems $\mathbf{Ax} = \mathbf{b}$.

Motiviert durch diese Sachverhalte liegt die Nutzung einer äquivalenten Umformulierung des Gleichungssystems mit dem Ziel der Verringerung der Konditionszahl der Matrix nahe. Solche Techniken werden als Präkonditionierung des Systems bezeichnet und haben sich im Kontext von Modellproblemen [18, 32, 69, 74] wie auch bei praktischen Anwendungsfällen [2, 21, 39, 48, 49, 50, 56] als effizientes Mittel zur Beschleunigung und Stabilisierung von Krylov-Unterraum-Verfahren erwiesen. Neben der Verwendung eines geeigneten Gleichungssystemlösers ist die Wahl des Präkonditionierers von ausschlaggebender Bedeutung für das resultierende Gesamtverfahren. Aus diesem Grund werden wir uns in diesem Abschnitt ausführlich mit der Beschreibung und Untersuchung möglicher Präkonditionierungstechniken befassen. Einen aktuellen Überblick findet man zudem in der Arbeit von Benzi [9]. Nach einer einführenden Definition werden wir unterschiedliche Varianten zur Präkonditionierung vorstellen. Anschließend präsentieren wir neben dem präkonditionierten CG-Verfahren auch eine spezielle Formulierung einer präkonditionierten BiCGSTAB-Methode, die einen Vergleich von Links- und Rechtspräkonditionierungen ermöglicht. Zur Konvergenzanalyse nutzen wir die Poisson-Gleichung und die Konvektions-Diffusions-Gleichung.

Definition 5.1 Seien $\mathbf{P}_L$ und $\mathbf{P}_R \in \mathbb{R}^{n \times n}$ regulär. Dann heißt

$$\mathbf{P}_L \mathbf{A} \mathbf{P}_R \mathbf{x}^P = \mathbf{P}_L \mathbf{b} \quad (5.0.1)$$

$$\mathbf{x} = \mathbf{P}_R \mathbf{x}^P \quad (5.0.2)$$

das zum Gleichungssystem $\mathbf{Ax} = \mathbf{b}$ gehörige präkonditionierte System. Gilt $\mathbf{P}_L \neq \mathbf{I}$, so heißt $\mathbf{P}_L$ Linkspräkonditionierer und das System linkspräkonditioniert. Ist $\mathbf{P}_R \neq \mathbf{I}$, so heißt $\mathbf{P}_R$ Rechtsprekonditionierer und das System rechtspräkonditioniert. Ein links- und rechtspräkonditioniertes System heißt beidseitig präkonditioniert.

Aufgrund der Invertierbarkeit der Matrix $\mathbf{A}$ stellt zum Beispiel $\mathbf{A}^T$ eine mögliche Präkonditionierungsmatrix dar. Mit $\mathbf{P}_L = \mathbf{A}^T$ und $\mathbf{P}_R = \mathbf{I}$ erhalten wir die Normalengleichungen

$$\mathbf{A}^T \mathbf{Ax} = \mathbf{A}^T \mathbf{b}. \quad (5.0.3)$$

Die Matrix $\mathbf{B} = \mathbf{A}^T \mathbf{A}$ ist symmetrisch und positiv definit, wodurch die Verwendung des CG-Verfahrens möglich ist. Diese Vorgehensweise wird als CGNR (CG Normal equations Residual minimizing) bezeichnet. Analog führt $\mathbf{P}_L = \mathbf{I}$ und $\mathbf{P}_R = \mathbf{A}^T$ auf das

CGNE-Verfahren (CG Normal equations Error minimizing). Beide Verfahren können also als präkonditionierte CG-Verfahren interpretiert werden. Oftmals erweisen sich diese Ansätze jedoch im Fall schlecht konditionierter Matrizen aufgrund von $\text{cond}_2(\mathbf{A}\mathbf{A}^T) \gg \text{cond}_2(\mathbf{A})$ als ineffizient. Besitzt $\mathbf{A}$ dagegen eine kleine Konditionszahl, so stellen beide Verfahren gute Alternativen zu den im Abschnitt 4.3.2 diskutierten Methoden dar, da sich die positiven Eigenschaften des CG-Algorithmus auf diese Verfahren übertragen. Besonders im Spezialfall orthogonaler Matrizen liefern CGNR und CGNE wegen $\mathbf{A}^T\mathbf{A} = \mathbf{I}$ die exakte Lösung des Gleichungssystems im ersten Iterationsschritt.

Mit $\mathbf{P}_L\mathbf{A}\mathbf{P}_R = \mathbf{I}$ liegt eine im Sinne der Konvergenz optimale Wahl der Matrizen $\mathbf{P}_L$ und $\mathbf{P}_R$ vor. Eine derartige Forderung kommt jedoch einer Invertierung der Matrix $\mathbf{A}$ gleich und ist daher aus Speicherplatz- und Rechenzeitgründen üblicherweise nicht realisierbar. Das Ziel der Präkonditionierung liegt in der Definition einfach berechenbarer Matrizen $\mathbf{P}_L$ und $\mathbf{P}_R$, die einen geringen Speicherplatzbedarf aufweisen und mit $\mathbf{P}_L\mathbf{A}\mathbf{P}_R$ eine gute Approximation der Einheitsmatrix liefern, so daß $\text{cond}(\mathbf{P}_L\mathbf{A}\mathbf{P}_R) \ll \text{cond}(\mathbf{A})$ gilt. Die beidseitige Variante wird, wie im Abschnitt 4.1.5 erwähnt, bei Verfahren für positiv definite und symmetrische Matrizen und der Wahl $\mathbf{P}_L^T = \mathbf{P}_R$ verwendet, damit mit $\mathbf{P}_L\mathbf{A}\mathbf{P}_R$ wiederum eine positiv definite und symmetrische Matrix vorliegt.

In der Praxis unterscheidet man oft zwischen expliziten und impliziten Präkonditionieren. Mit dieser Sprechweise will man betonen, daß man entweder die Präkonditionierungsmatrix explizit zur Verfügung hat oder aber den Präkonditionierer nur implizit in seiner Wirkungsweise bei Matrix-Multiplikationen kennt. Ein Beispiel für die implizite Variante ist die unvollständige LU-Zerlegung, bei der die Inverse des Präkonditionierers in faktorisierter Form als Produkt von zwei Dreiecksmatrizen vorliegt. Das Produkt der Matrizen $\mathbf{LU}$ kann auch bei schwachbesetzten Matrizen $\mathbf{L}$ und $\mathbf{U}$ eine vollbesetzte Matrix darstellen. Die explizite Berechnung der Produktmatrix würde daher eventuell die zur Verfügung stehenden Speicherressourcen übersteigen. Desweiteren müßte zur Bestimmung der Präkonditionierungsmatrix die hieraus resultierende Matrix aufwendig invertiert werden, während die Invertierung des Produktes $\mathbf{LU}$ durch einfache Rückwärts- und Vorwärtselimination implizit gegeben ist.

5.1 Skalierungen

Skalierungen stellen die einfachste Form der Präkonditionierung dar. Den Vorteilen des geringen Speicherplatzbedarfs und der leichten Berechnung steht der Nachteil entgegen, daß mit der Skalierung in der Regel nur eine sehr grobe Approximation der Inversen der Matrix $\mathbf{A}$ vorliegt, und folglich auch zumeist nur eine geringfügige Beschleunigung erzielt wird.

Definition 5.2 Eine reguläre Diagonalmatrix $\mathbf{D} = \text{diag}\{d_{11}, \dots, d_{nn}\} \in \mathbb{R}^{n \times n}$ heißt Skalierung.

Mit der folgenden Auflistung werden wir einige gebräuchliche Formen der Skalierung auf der Basis einer regulären Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ präsentieren:

- (a) Skalierung mit dem Diagonalelement:

Unter der Voraussetzung $a_{ii} \neq 0$ für $i = 1, \dots, n$ wählt man

$$d_{ii} := \frac{1}{a_{ii}} \text{ fur } i = 1, \dots, n. \quad (5.1.1)$$

(b) Zeilen-/Spaltenskalierung bezuglich der Betragssummennorm:

$$d_{ii} := \frac{1}{\sum_{j=1}^n |a_{ij}|} \text{ beziehungsweise } d_{jj} := \frac{1}{\sum_{i=1}^n |a_{ij}|} \quad (5.1.2)$$

fur $i = 1, \dots, n$ respektive $j = 1, \dots, n$.

(c) Zeilen-/Spaltenskalierung bezuglich der euklidischen Norm:

$$d_{ii} := \frac{1}{\left(\sum_{j=1}^n |a_{ij}|^2\right)^{\frac{1}{2}}} \text{ beziehungsweise } d_{jj} := \frac{1}{\left(\sum_{i=1}^n |a_{ij}|^2\right)^{\frac{1}{2}}} \quad (5.1.3)$$

fur $i = 1, \dots, n$ respektive $j = 1, \dots, n$.

(d) Zeilen-/Spaltenskalierung bezuglich der Maximumsnorm:

$$d_{ii} := \frac{1}{\max_{j=1, \dots, n} |a_{ij}|} \text{ beziehungsweise } d_{jj} := \frac{1}{\max_{i=1, \dots, n} |a_{ij}|} \quad (5.1.4)$$

fur $i = 1, \dots, n$ respektive $j = 1, \dots, n$.

Wie zu erwarten, hangt die Gute einer Skalierung auch entscheidend von der gewahlten Norm ab. In der Betragssummennorm und der Maximumsnorm sind sogar optimale Skalierungen moglich.

Satz 5.3 Seien $A \in \mathbb{R}^{n \times n}$ regular und $\tilde{D} \in \mathbb{R}^{n \times n}$ als Zeilenskalierung bezuglich der Betragssummennorm gema (5.1.2) gegeben. Dann gilt

$$\text{cond}_{\infty}(\tilde{D}A) \leq \text{cond}_{\infty}(DA) \quad (5.1.5)$$

fur jede Skalierung $D \in \mathbb{R}^{n \times n}$.

Beweis:

Aufgrund der vorausgesetzten Regularitat der Matrix A ist $\tilde{D}$ wohldefiniert, und es liegt mit $C = \tilde{D}A$ eine Matrix vor, deren Zeilen die Lange 1 bezuglich der Betragssummennorm aufweisen. Sei eine beliebige Skalierung $D \in \mathbb{R}^{n \times n}$ gegeben, dann definieren wir $\bar{D} = \text{diag}\{\bar{d}_{11}, \dots, \bar{d}_{nn}\} = D\tilde{D}^{-1}$. Aus dem Zeilengleichgewicht der Matrix C folgt

$$\|DA\|_{\infty} = \|\bar{D}C\|_{\infty} = \max_{i=1, \dots, n} |\bar{d}_{ii}| \|C\|_{\infty} = \max_{i=1, \dots, n} |\bar{d}_{ii}| \|\tilde{D}A\|_{\infty},$$

wohingegen stets

$$\|(DA)^{-1}\|_{\infty} = \|(\bar{D}C)^{-1}\|_{\infty} = \|C^{-1}\bar{D}^{-1}\|_{\infty} \geq \frac{\|C^{-1}\|_{\infty}}{\|\bar{D}\|_{\infty}} = \frac{1}{\max_{i=1, \dots, n} |\bar{d}_{ii}|} \|(\tilde{D}A)^{-1}\|_{\infty}$$

gilt. Zusammenfassend ergibt sich somit die Behauptung. $\square$

Berucksichtigt man $\|A\|_1 = \|A^T\|_{\infty}$, so erhalt man das folgende Korollar.

Korollar 5.4 Seien $A \in \mathbb{R}^{n \times n}$ regulär und $\tilde{D} \in \mathbb{R}^{n \times n}$ als Spaltenskalierung bezüglich der Betragssummennorm gemäß (5.1.2) gegeben. Dann gilt

$$\text{cond}_1(A\tilde{D}) \leq \text{cond}_1(AD) \quad (5.1.6)$$

für jede Skalierung $D \in \mathbb{R}^{n \times n}$.

Bemerkung:

Die Konditionszahl einer gegebenen Matrix ist invariant gegenüber Multiplikationen mit $\lambda \in \mathbb{R} \setminus \{0\}$. Somit sind alle optimalen Skalierungen stets nur bis auf einen multiplikativen Faktor eindeutig bestimmt.

Für die Konditionszahl in der euklidischen Norm gilt folgende Abschätzung:

Satz 5.5 Seien A regulär und die Skalierung $\tilde{D}$ derart gewählt, daß alle Spalten von $A\tilde{D}$ die gleiche Länge bezüglich der euklidischen Norm besitzen. Dann gilt

$$\text{cond}_2(A\tilde{D}) \leq \sqrt{n} \inf_D \text{cond}_2(AD), \quad (5.1.7)$$

wobei das Infimum über alle regulären Diagonalmatrizen $D \in \mathbb{R}^{n \times n}$ gebildet wird.

Beweis:

Es sei B eine beliebige Matrix und b_j bezeichne die j -te Spalte von B , so gilt die folgende Kette von Ungleichungen

$$\max_{j=1, \dots, n} \|b_j\|_2 \leq \|B\|_2 \leq \left(\sum_{i,j=1}^n |b_{ij}|^2 \right)^{\frac{1}{2}} \leq \sqrt{n} \max_{j=1, \dots, n} \|b_j\|_2. \quad (5.1.8)$$

Wir definieren $C = A\tilde{D}$. Da für jede Skalierung D eine Skalierung $\overline{D}$ mit $D = \tilde{D}^{-1}\overline{D}$ existiert, kann die Aussage in der äquivalenten Form

$$\text{cond}_2(C) \leq \sqrt{n} \inf_{\overline{D}} \text{cond}_2(C\overline{D})$$

formuliert werden, wobei das Infimum wiederum über alle regulären Diagonalmatrizen $\overline{D} \in \mathbb{R}^{n \times n}$ gebildet wird. Aufgrund der obigen Bemerkung können wir zudem für die Spalten c_j der Matrix C die Eigenschaft

$$\|c_j\|_2 = 1 \quad (5.1.9)$$

o.B.d.A. annehmen. Sei nun eine beliebige Skalierung $\overline{D}$ gegeben, so definieren wir $\hat{D} = \text{diag}\{\hat{d}, \dots, \hat{d}\}$ mit $\hat{d} = \max_{j=1, \dots, n} |\overline{d}_{jj}|$. Diese Festlegung impliziert

$$\|\hat{D}^{-1}\|_2 = \frac{1}{\|\overline{D}\|_2},$$

woraus unter Verwendung der Submultiplikativität laut Satz 2.17 die Ungleichung

$$\|C^{-1}\|_2 = \hat{d} \|(C\hat{D})^{-1}\|_2 = \hat{d} \|(\hat{D}^{-1}\overline{D}\overline{D}^{-1}C^{-1})\|_2 \leq \hat{d} \|(C\overline{D})^{-1}\|_2 \quad (5.1.10)$$

folgt. Aus (5.1.8) erhalten wir

$$\frac{\hat{d}}{\sqrt{n}} \|C\|_2 = \frac{1}{\sqrt{n}} \|C\hat{D}\|_2 \leq \max_{j=1,\dots,n} \|(C\hat{D})_j\|_2 \stackrel{(5.1.9)}{=} \max_{j=1,\dots,n} \|(C\overline{D})_j\|_2 \leq \|C\overline{D}\|_2. \quad (5.1.11)$$

Eine Kombination der Ungleichungen (5.1.10) und (5.1.11) liefert die Behauptung. $\square$

Natürlich ist diese Aussage für die Praxis unbefriedigend, da dort die Dimension des Gleichungssystems im allgemeinen so groß ist, daß auch der Faktor $\sqrt{n}$ untragbar hoch ist. Es ist aber möglich, Verschärfungen dieses Satzes nachzuweisen, wie sie den Arbeiten [65] und [66] entnommen werden können.

Es bleibt zu bemerken, daß eine Skalierung nicht unbedingt zu einer Verbesserung der Konditionszahl der Matrix führen muß. So besagt der Satz 5.3, daß im Fall einer regulären Matrix A deren Zeilen bereits die gleiche Länge bezüglich der Betragssummennorm aufweisen, mittels einer Linkspräkonditionierung in der Form einer Skalierung keine Verbesserung der Konditionszahl hinsichtlich der Zeilensummennorm erzielt werden kann. Eine analoge Aussage ergibt sich für die Rechtspräkonditionierung mit Korollar 5.4 für reguläre Matrizen deren Spalten die gleiche Länge bezüglich der Betragssummennorm besitzen. Betrachten wir die Matrix

$$A = \begin{pmatrix} 2 & 100 \\ 100 & 100 \end{pmatrix},$$

so liefert eine Linkspräkonditionierung des Gleichungssystems mittels einer Skalierung gemäß (5.1.1) die Matrix

$$B = \begin{pmatrix} 1 & 50 \\ 1 & 1 \end{pmatrix}$$

mit $\text{cond}_2(B) = 51,062$, so daß die Skalierung wegen $\text{cond}_2(A) = 2,6899$ zu einer deutlichen Erhöhung der Konditionszahl geführt hat.

5.2 Polynomiale Präkonditioner

Die grundlegende Idee polynomialer Präkonditioner beruht auf der Darstellung der Inversen einer Matrix in der Form einer Neumannschen Reihe.

Satz 5.6 *Sei $\rho(I - A) < 1$, dann ist die Matrix $A \in \mathbb{R}^{n \times n}$ regulär, und die Inverse A^{-1} besitzt die Darstellung in der Form einer Neumannschen Reihe*

$$A^{-1} = \sum_{k=0}^{\infty} (I - A)^k. \quad (5.2.1)$$

Beweis:

Mit $\rho(I - A) < 1$ existiert laut Satz 2.35 eine induzierte Matrixnorm $\|\cdot\|$ mit $\|I - A\| < 1$. Hiermit folgt aufgrund der geometrischen Reihe

$$\left\| \sum_{k=0}^{\infty} (I - A)^k \right\| \leq \sum_{k=0}^{\infty} \|(I - A)^k\| \leq \sum_{k=0}^{\infty} \|I - A\|^k = \frac{1}{1 - \|I - A\|},$$

so daß mit $B = \sum_{k=0}^{\infty} (I - A)^k$ ein beschränkter linearer Operator vorliegt, der wegen $\|(I - A)^{m+1}\| \leq \|I - A\|^{m+1} \rightarrow 0$ für $m \rightarrow \infty$ der Eigenschaft

$$BA = \lim_{m \rightarrow \infty} \left[\sum_{k=0}^m (I - A)^k \right] A = \lim_{m \rightarrow \infty} (I - (I - A)^{m+1}) = I$$

genügt und folglich die Inverse der Matrix A darstellt. $\square$

Die Voraussetzung über den Spektralradius ist natürlich einschneidend. Wie der folgende Satz zeigt, ist diese Bedingung für gewisse Klassen von Matrizen jedoch stets durch eine geeignete Skalierung erfüllbar.

Satz 5.7 *Sei $A = (a_{ij})_{i,j=1,\dots,n} \in \mathbb{R}^{n \times n}$ regulär und strikt diagonaldominant, die Skalierung $D = \text{diag}\{d_{11}, \dots, d_{nn}\}$ sei gegeben durch*

$$d_{ii} := \frac{1}{a_{ii}} \text{ für } i = 1, \dots, n. \quad (5.2.2)$$

Dann ist

$$\rho(I - DA) < 1, \quad (5.2.3)$$

und es gilt

$$A^{-1} = \left[\sum_{k=0}^{\infty} (I - DA)^k \right] D. \quad (5.2.4)$$

Beweis:

Die Existenz und Regularität der Matrix D folgt aus der strikten Diagonaldominanz von A . Weiter gilt $\|I - DA\|_{\infty} < 1$, wodurch (5.2.3) erfüllt ist, und die Darstellung der Inversen A^{-1} aus dem Satz 5.6 folgt. $\square$

Natürlich kann bei Verwendung der Matrix D in Form einer geeigneten Approximation der Inversen A^{-1} die Bedingung (5.2.3) prinzipiell gewährleistet werden. Diese Matrix D stellt dann aber nicht notwendigerweise eine Diagonalmatrix dar, und wir sind wieder beim Ausgangsproblem der Präkonditionierung angelangt.

Die Verwendung einer abgeschnittenen Neumannreihe als Präkonditioner ergibt sich auf ganz natürliche Weise aus dem Satz 5.6 und wurde erstmalig in [20] vorgeschlagen.

Definition 5.8 Sei $\rho(I - A) < 1$ und $m \in \mathbb{N}$ beliebig, dann heißt

$$P^{(m)} := \sum_{k=0}^m (I - A)^k \quad (5.2.5)$$

die abgeschnittene Neumannsche Reihe m -ter Stufe zur Matrix A .

Bemerkung:

Die so gebildeten Matrizen $P^{(m)}$ können sowohl zur Links- als auch zur Rechtspräkonditionierung verwendet werden. Bei der Implementierung ist zu beachten, daß $P^{(m)}$ nur bei Matrix-Vektor-Multiplikationen benötigt wird, weshalb die Verwendung des Horner-schemas naheliegt.

Realisierung einer Matrix-Vektor-Multiplikation $P^{(m)}z$ —

$b_0 := z$
für $i = 1, \dots, m$
$b_i := z + (I - A) b_{i-1}$
$P^{(m)}z = b_m$

Der Polynomansatz (5.2.5) legt es nahe, auch andere Polynome als die abgeschnittene Neumannsche Reihe zu betrachten. So erhält man zum Beispiel durch den Ansatz

$$P := \sum_{k=0}^m \theta_k (I - A)^k \quad (5.2.6)$$

ein parameterabhängiges Polynom, wobei die Gewichte θ_k dazu benutzt werden können, die Konvergenzgeschwindigkeit des Iterationsverfahrens gezielt zu steuern.

Ist die Matrix $A \in \mathbb{C}^{n \times n}$ hermitesch, so kann man aus dem Raum $\mathcal{P}_m$ aller Polynome vom Höchstgrad m die Bestapproximation p^* bestimmen, welche in einer vorgegebenen Norm $\|\cdot\|$ die Bedingung

$$\|I - A p^*(A)\| \leq \|I - A p(A)\| \quad \forall p \in \mathcal{P}_m \quad (5.2.7)$$

erfüllt. Diese Vorgehensweise werden wir anhand der euklidischen Norm demonstrieren. Laut Satz 2.36 ist A unitär diagonalisierbar, und wir erhalten die zu (5.2.7) äquivalente Forderung

$$\|I - D p^*(D)\|_2 \leq \|I - D p(D)\|_2 \quad \forall p \in \mathcal{P}_m,$$

wobei $D = \text{diag}\{\lambda_1, \dots, \lambda_n\} \in \mathbb{R}^{n \times n}$ die Diagonalmatrix mit den Eigenwerten von A repräsentiert. Folglich minimiert p^* unter allen Polynomen $p \in \mathcal{P}_m$ den Ausdruck

$$\max_{i=1, \dots, n} |1 - \lambda_i p(\lambda_i)|.$$

Für die unbekannten Eigenwerte kann mit dem Satz von Gerschgorin [68] ein Intervall $I = [a, b]$ mit $a < 0 < b$ festgelegt werden, welches alle Eigenwerte enthält. Anschließend wird p^* derart bestimmt, daß der Ausdruck

$$\max_{x \in I} |1 - x p(x)|$$

über alle $p \in \mathcal{P}_m$ minimiert wird. An dieser Stelle stellt sich natürlich die Frage nach der Existenz eines derartigen Polynoms. Zunächst liegt mit $\|q\|_\infty := \max_{x \in I} |q(x)|$ eine Norm auf dem linearen Raum $\mathcal{P}_{m+1}(I) := \{q : I \mapsto \mathbb{R} \mid q \in \mathcal{P}_{m+1}\}$ vor. Desweiteren stellt $\mathcal{P}_{m+1}^0(I) := \{q \in \mathcal{P}_{m+1}(I) \mid q(0) = 0\}$ einen endlichdimensionalen Unterraum von $\mathcal{P}_{m+1}(I)$ dar. Wählen wir ein festes aber beliebiges Polynom $\tilde{q} \in \mathcal{P}_{m+1}^0(I)$, so erhalten wir mit

$$\mathcal{M} := \{q \in \mathcal{P}_{m+1}^0(I) \mid \|1 - q\|_\infty \leq \|1 - \tilde{q}\|_\infty\}$$

eine abgeschlossene und beschränkte Teilmenge von $\mathcal{P}_{m+1}^0(I)$, die folglich kompakt ist. Somit stellt $\mathcal{M}$ zudem eine kompakte Teilmenge von $\mathcal{P}_{m+1}(I)$ dar, wodurch zu jedem $\hat{q} \in \mathcal{P}_{m+1}(I)$ mindestens ein $q \in \mathcal{M}$ mit

$$\|\hat{q} - q\|_\infty \leq \|\hat{q} - \bar{q}\|_\infty \quad \forall \bar{q} \in \mathcal{M}$$

existiert. Die Wahl $\hat{q}(x) = 1$ liefert

$$\|1 - q\|_\infty \leq \|1 - \bar{q}\|_\infty \leq \|1 - \check{q}\|_\infty$$

für alle $\check{q} \in \mathcal{P}_{m+1}^0(I)$. Mit

$$p(x) = q(x)/x \text{ für } x \neq 0$$

und $p(0) = \lim_{x \rightarrow 0} q(x)/x$ erhalten wir das gesuchte Polynom. Für Details zur praktischen Berechnung sei auf [5] verwiesen.

Die Bestimmung polynomialer Präkonditionierer als Bestapproximation in einer gewichteten L_2 -Norm findet man zum Beispiel in [40] und [59].

5.3 Splitting-assozierte Präkonditionierer

Im Abschnitt 4.1 haben wir uns ausführlich mit der Herleitung unterschiedlicher Splitting-Methoden und der Diskussion ihrer Eigenschaften beschäftigt. Diese Verfahren basieren auf einer Aufteilung der Matrix $\mathbf{A}$ in der Form $\mathbf{A} = \mathbf{B} + (\mathbf{A} - \mathbf{B})$. Hierbei soll $\mathbf{B}$ eine leicht invertierbare Approximation der Matrix $\mathbf{A}$ darstellen, so daß die resultierende Iterationsmatrix $\mathbf{M} = \mathbf{B}^{-1}(\mathbf{B} - \mathbf{A})$ einen möglichst kleinen Spektralradius besitzt. Diese Forderung weist eine sehr große Ähnlichkeit zu den gewünschten Eigenschaften von Präkonditionierern auf, wodurch sich eine Verwendung der Matrix $\mathbf{N} = \mathbf{B}^{-1}$ als Präkonditionierer anbietet.

Definition 5.9 Sei durch $\mathbf{x}_{j+1} = \mathbf{M}\mathbf{x}_j + \mathbf{N}\mathbf{b}$ eine Splitting-Methode zur Lösung von $\mathbf{A}\mathbf{x} = \mathbf{b}$ mit einer regulären Matrix $\mathbf{N}$ gegeben, dann heißt $\mathbf{P} = \mathbf{N}$ der zur Splitting-Methode assoziierte Präkonditionierer.

Die durch die obige Definition eingeführten Präkonditionierungsmatrizen lassen sich zur Links- und Rechtspräkonditionierung verwenden. Die im folgenden aufgeführten Möglichkeiten zur Wahl der Matrix $\mathbf{N}$ basieren ausschließlich auf den im Abschnitt 4.1 vorgestellten iterativen Verfahren, wodurch sich auch die zur Invertierung der Matrix $\mathbf{B}$ notwendige Eigenschaft der Regularität der Diagonalmatrix $\mathbf{D} = \text{diag}\{a_{11}, \dots, a_{nn}\}$ überträgt. Desweiteren stellen die Matrizen $\mathbf{L}$ und $\mathbf{R}$ gemäß (4.1.11) und (4.1.12) die strikte linke untere beziehungsweise strikte rechte obere Dreiecksmatrix bezüglich $\mathbf{A}$ dar. Die Tabelle 5.1 liefert eine Auswahl möglicher Splitting-assoziierter Präkonditionierer.

Bei der Implementierung nutzt man aus, daß bei den vorgestellten Präkonditionierern die zu invertierenden Matrizen stets eine Diagonal- oder Dreiecksgestalt aufweisen und folglich bei der Matrix-Vektor-Multiplikation die entsprechende Eliminationstechnik genutzt werden kann. Der dargestellte Jacobi-assozierte Präkonditionierer ist äquivalent

Splitting-Methode	Assoziierter Präkonditionierer
Jacobi-Verfahren	$\mathbf{P}_{Jac} = \mathbf{D}^{-1}$
Gauß-Seidel-Verfahren	$\mathbf{P}_{GS} = (\mathbf{D} + \mathbf{L})^{-1}$
SOR-Verfahren	$\mathbf{P}_{GS}(\omega) = \omega (\mathbf{D} + \omega \mathbf{L})^{-1}$
Symm. Gauß-Seidel-Verfahren	$\mathbf{P}_{SGS} = (\mathbf{D} + \mathbf{R})^{-1} \mathbf{D} (\mathbf{D} + \mathbf{L})^{-1}$
SSOR-Verfahren	$\mathbf{P}_{SGS}(\omega) = \omega(2 - \omega) (\mathbf{D} + \omega \mathbf{R})^{-1} \mathbf{D} (\mathbf{D} + \omega \mathbf{L})^{-1}$

Tabelle 5.1 Splitting-assoziierte Präkonditionierer

zur Skalierung mit dem Diagonalelement (5.1.1). Die Eigenschaften der im Bezug zu einem Relaxationsverfahren stehenden Präkonditionierer $\mathbf{P}_{GS}(\omega)$ und $\mathbf{P}_{SGS}(\omega)$ hängen entscheidend von der Wahl des Relaxationsparameters ω ab. Im Fall einer symmetrischen, positiv definiten Matrix $\mathbf{A}$ werden in [6] Abschätzungen für das Spektrum des präkonditionierten Systems und ein im Sinne von cond_2 optimales ω angegeben.

5.4 Die unvollständige LU-Zerlegung

Liegt eine große schwachbesetzte Matrix vor, so erweist sich die Berechnung einer LU-Zerlegung unter den Gesichtspunkten der Rechenzeit und des Speicherplatzbedarfs auch bei Vernachlässigung der Rundungsfehler als ineffizient und häufig sogar unpraktikabel.

Die Idee der unvollständigen LU-Zerlegung ist es, eine Zerlegung der Form $\mathbf{A} = \mathbf{LU} + \mathbf{F}$ mit einer linken unteren Dreiecksmatrix $\mathbf{L}$ und einer rechten oberen Dreiecksmatrix $\mathbf{U}$ durchzuführen, wobei der gesamte Speicherplatzbedarf der Matrizen $\mathbf{L}$ und $\mathbf{U}$ identisch zu dem der Matrix $\mathbf{A}$ sein soll. Das Vernachlässigen der Restmatrix $\mathbf{F}$ liefert eine leicht invertierbare Matrix $\tilde{\mathbf{A}} = \mathbf{LU}$, deren Inverse als Approximation von $\mathbf{A}^{-1}$ zur Präkonditionierung des Gleichungssystems $\mathbf{Ax} = \mathbf{b}$ verwendet werden kann. Um die oben genannten Eigenschaften zu gewährleisten, müssen $n^2 + n$ Bedingungen an die unvollständige LU-Zerlegung gestellt werden.

Bevor wir einen leicht programmierbaren Algorithmus herleiten werden, führen wir die hilfreichen Begriffe Matrix-, Spalten- sowie Zeilenmuster und Besetzungsstruktur ein.

Definition 5.10 Jede Menge

$$\mathcal{M} \subset \{(i, j) \mid i, j \in \{1, \dots, n\}\}$$

heißt Matrixmuster im Raum $\mathbb{R}^{n \times n}$. Zu gegebenem Matrixmuster $\mathcal{M}$ heißt

$$\mathcal{M}_S(j) := \{i \mid (i, j) \in \mathcal{M}\}$$

das zu $\mathcal{M}$ gehörige j -te Spaltenmuster und

$$\mathcal{M}_Z(j) := \{i \mid (j, i) \in \mathcal{M}\}$$

das zu $\mathcal{M}$ gehörige j -te Zeilenmuster. Zu gegebener Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ bezeichnet

$$\mathcal{M}^{\mathbf{A}} := \{(i, j) \mid a_{ij} \neq 0\}$$

die Besetzungsstruktur von $\mathbf{A}$.

Unter Zuhilfenahme der Besetzungsstruktur legen wir nun die Bedingungen zur Berechnung der unvollständigen LU-Zerlegung mittels der folgenden Definition fest.

Definition 5.11 Sei $\mathbf{A} \in \mathbb{R}^{n \times n}$. Die Zerlegung

$$\mathbf{A} = \mathbf{L}\mathbf{U} + \mathbf{F} \quad (5.4.1)$$

existiere unter den Bedingungen

- 1) $u_{ii} = 1$ für $i = 1, \dots, n$,
- 2) $\ell_{ij} = u_{ij} = 0$, falls $(i, j) \notin \mathcal{M}^{\mathbf{A}}$,
- 3) $(\mathbf{L}\mathbf{U})_{ij} = a_{ij}$, falls $(i, j) \in \mathcal{M}^{\mathbf{A}}$,

und es seien $\mathbf{L} = (\ell_{ij})_{i,j=1,\dots,n}$ und $\mathbf{U} = (u_{ij})_{i,j=1,\dots,n}$ eine reguläre linke untere beziehungsweise rechte obere Dreiecksmatrix, dann heißt (5.4.1) unvollständige LU-Zerlegung (incomplete LU, ILU) der Matrix $\mathbf{A}$ zum Muster $\mathcal{M}^{\mathbf{A}}$.

Die im Abschnitt 3.1 gewonnenen Existenzaussagen für die LR-Zerlegung lassen sich nicht auf die unvollständige Formulierung übertragen. So existiert aufgrund der Bedingung 2 für

$$\mathbf{A} = \begin{pmatrix} 1 & -1 \\ 1 & 0 \end{pmatrix}$$

keine unvollständige LU-Zerlegung, obwohl mit $\det \mathbf{A}[k] \neq 0$ ($k = 1, 2$) durch die Sätze 3.7 und 3.8 die Existenz und eine Form der Eindeutigkeit der LR-Zerlegung gewährleistet ist. Dennoch stellt die unvollständige LU-Zerlegung eine sehr effiziente Form der Präkonditionierung dar, die bei Simulationen technischer Problemstellungen häufig erfolgreich angewendet wird, da die hierbei auftretenden Gleichungssysteme in der Regel keinen Widerspruch zu den Bedingungen 1 und 2 liefern (siehe die Beispiele 1.1 und 1.4). Existenzaussagen können für den Fall einer M-Matrix der Arbeit von Meijerink und van der Vorst [47] und für H-Matrizen der Veröffentlichung von Manteuffel [46] entnommen werden.

Zur Herleitung der Zerlegung entwickeln wir sukzessive für $i = 1, \dots, n$ zunächst die i -te Spalte von $\mathbf{L}$ und anschließend die i -te Zeile von $\mathbf{U}$. Aus den Bedingungen 1 und 3 erhalten wir wegen $u_{mi} = 0$ ($m > i$) die Gleichung

$$a_{ki} = \sum_{m=1}^n \ell_{km} u_{mi} = \sum_{m=1}^i \ell_{km} u_{mi} = \sum_{m=1}^{i-1} \ell_{km} u_{mi} + \ell_{ki},$$

woraus sich sofort die i -te Spalte von $\mathbf{L}$ in der Form

$$\ell_{ki} = a_{ki} - \sum_{m=1}^{i-1} \ell_{km} u_{mi} \quad \text{für } k = i, \dots, n \quad (5.4.2)$$

ergibt. Mit $\ell_{im} = 0$ ($m > i$) folgt für $k = i + 1, \dots, n$ die Gleichung

$$a_{ik} = \sum_{m=1}^n \ell_{im} u_{mk} = \sum_{m=1}^i \ell_{im} u_{mk},$$

und wir erhalten die Darstellung der i -ten Zeile von $\mathbf{U}$ gemäß

$$u_{ik} = \frac{1}{\ell_{ii}} \left(a_{ik} - \sum_{m=1}^{i-1} \ell_{im} u_{mk} \right) \quad \text{für } k = i + 1, \dots, n. \quad (5.4.3)$$

Die Berücksichtigung der Bedingung 2 liefert abschließend die unvollständige LU-Zerlegung in der folgenden Schreibweise:

Die unvollständige LU-Zerlegung —

Für $i = 1, \dots, n$	
Für $k = i, \dots, n, k \in \mathcal{M}_S^A(i)$	
$\ell_{ki} = a_{ki} - \sum_{m=1}^{i-1} \ell_{km} u_{mi}$ $m \in \mathcal{M}_Z^A(k) \cap \mathcal{M}_S^A(i)$	
Für $k = i + 1, \dots, n, k \in \mathcal{M}_Z^A(i)$	
$u_{ik} = \frac{1}{\ell_{ii}} \left(a_{ik} - \sum_{m=1}^{i-1} \ell_{im} u_{mk} \right)$ $m \in \mathcal{M}_Z^A(i) \cap \mathcal{M}_S^A(k)$	

Definition 5.12 Es sei

$$\mathbf{A} = \mathbf{L}\mathbf{U} + \mathbf{F}$$

eine unvollständige LU-Zerlegung der Matrix $\mathbf{A}$, dann ist

$$\mathbf{P} := \mathbf{U}^{-1} \mathbf{L}^{-1} \quad (5.4.4)$$

der zugehörige ILU-Präkonditionierer.

Wir haben bereits angemerkt, daß die Matrix $\mathbf{P}$ nicht explizit berechnet wird, sondern nur implizit in ihrer Wirkung auf der Basis einer Vorwärtselimination und einer anschließenden Rückwärtselimination verwendet wird. Zudem gibt es keine Notwendigkeit zur Speicherung der Diagonalelemente der Matrix $\mathbf{U}$, so daß die beiden Matrizen $\mathbf{L}$ und $\mathbf{U}$ in der Summe einen zu $\mathbf{A}$ identischen Speicherplatzbedarf aufweisen.

Bei Matrizen mit einer relativ kleinen Bandbreite erweist sich häufig die unvollständige LU-Zerlegung mit *fill-in* als effektiv. Hierbei wird ausgenutzt, daß die Wahl der Besetzungsstruktur $\mathcal{M}^A$ innerhalb der Bedingungen 2 und 3 nicht zwingend ist. Ersetzt man bei diesen Forderungen die Besetzungsstruktur $\mathcal{M}^A$ durch ein Muster $\mathcal{M} \supset \mathcal{M}^A$, so erhält die Fehlermatrix F weniger von Null verschiedene Einträge, wodurch mit LU zumeist eine bessere Approximation der Matrix A vorliegt. Diese Erweiterung wird mit $ILU(p)$ beschrieben, wobei die natürliche Zahl p einen Grad für die Anzahl der *fill-in*-Elemente darstellt. Für eine ausführliche Beschreibung der angesprochenen Verallgemeinerung sei auf Saad [60] verwiesen.

5.5 Die unvollständige Cholesky-Zerlegung

Bereits im Abschnitt 3.2 hatten wir erkannt, daß im Fall einer symmetrischen und positiv definiten Matrix A der Aufwand der LR-Zerlegung reduziert werden kann, indem die Eigenschaften der Matrix gezielt ausgenutzt werden. Die daraufhin hergeleitete Cholesky-Zerlegung werden wir nun auch zur Präkonditionierung in einer unvollständigen Formulierung betrachten.

Definition 5.13 Sei $A \in \mathbb{R}^{n \times n}$ symmetrisch und positiv definit. Die Zerlegung

$$A = LL^T + F \quad (5.5.1)$$

existiere unter den Bedingungen

- 1) $\ell_{ij} = 0$, falls $(i, j) \notin \mathcal{M}^A$,
- 2) $(LL^T)_{ij} = a_{ij}$, falls $(i, j) \in \mathcal{M}^A$,

und es sei $L = (\ell_{ij})_{i,j=1,\dots,n}$ eine reguläre linke untere Dreiecksmatrix, dann heißt (5.5.1) unvollständige Cholesky-Zerlegung (incomplete Cholesky, IC) der Matrix A zum Muster $\mathcal{M}^A$.

Die wesentlichen Vorteile einer solchen Zerlegung liegen im Rechen- und Speicheraufwand und in der Eigenschaft, daß sich im Fall einer positiv definiten, symmetrischen Matrix A wegen

$$(LAL^T)^T = LAL^T$$

und

$$(LAL^T x, x)_2 = (AL^T x, L^T x)_2 > 0, \quad \forall x \in \mathbb{R}^n \setminus \{0\}$$

die Symmetrie und positive Definitheit auch auf das Produkt LAL^T übertragen, wodurch die Verfahren für symmetrische, positiv definite Matrizen auch auf das präkonditionierte System angewendet werden können.

Aus den Gleichungen (3.2.2) und (3.2.3) ergibt sich unter Berücksichtigung der an die unvollständige Zerlegung gestellten Bedingungen die Berechnung der unvollständigen Cholesky-Zerlegung in der folgenden Form:

Die unvollständige Cholesky-Zerlegung —

Für $k = 1, \dots, n$	
$\ell_{kk} = \sqrt{a_{kk} - \sum_{\substack{j=1 \\ j \in \mathcal{M}_{\mathcal{Z}}^{\mathbf{A}}(k)}}^{k-1} \ell_{kj}^2}$	
Für $i = k + 1, \dots, n, i \in \mathcal{M}_{\mathcal{Z}}^{\mathbf{A}}(k)$	
$\ell_{ik} = \frac{1}{\ell_{kk}} \left(a_{ik} - \sum_{\substack{j=1 \\ j \in \mathcal{M}_{\mathcal{Z}}^{\mathbf{A}}(i) \cap \mathcal{M}_{\mathcal{Z}}^{\mathbf{A}}(k)}}^{k-1} \ell_{ij} \ell_{kj} \right)$	

Im Gegensatz zur ILU benötigt die IC keine zusätzliche Bedingung an die Diagonalelemente (siehe Bedingung 1 in der Definition 5.11), da aufgrund der positiven Definitheit stets $(k, k) \in \mathcal{M}^{\mathbf{A}}$ gilt.

5.6 Die unvollständige QR-Zerlegung

Bei der Konstruktion der unvollständigen LU-Zerlegung war der Hauptgedanke, den Speicherplatz für den Präkonditionierer a priori einzuschränken, indem eine zu $\mathbf{A}$ äquivalente Besetzungsstruktur gefordert wurde. Dieser Gedanke läßt sich auch auf andere Präkonditionierer übertragen.

Definition 5.14 Wird mittels einer modifizierten Variante eines beliebigen Algorithmus zur Bestimmung einer QR-Zerlegung eine Darstellung der Matrix $\mathbf{A}$ in der Form

$$\mathbf{A} = \mathbf{Q}\mathbf{R} + \mathbf{F} \quad (5.6.1)$$

berechnet, wobei die obere Dreiecksmatrix $\mathbf{R}$ und die Matrix $\mathbf{Q}$ den Bedingungen

- 1) $r_{ij} = 0$, falls $(i, j) \notin \mathcal{M}^{\mathbf{A}}$
- 2) $q_{ij} = 0$, falls $(i, j) \notin \mathcal{M}^{\mathbf{A}}$

genügen, und liegt bei Ersetzen der Besetzungsstruktur $\mathcal{M}^{\mathbf{A}}$ durch $\{(i, j) | i, j = 1, \dots, n\}$ mit $\mathbf{Q}$ eine orthogonale Matrix und mit $\mathbf{F}$ eine Nullmatrix vor, dann heißt (5.6.1) unvollständige QR-Zerlegung (incomplete QR, IQR) der Matrix $\mathbf{A}$ zum Muster $\mathcal{M}^{\mathbf{A}}$.

Mit dem Gram-Schmidt-Verfahren haben wir bereits in Abschnitt 3.3.1 eine Methode zur Ermittlung einer QR-Zerlegung kennengelernt. Unter Berücksichtigung der aufgestellten Bedingungen 1 und 2 läßt sich das Gram-Schmidt-Verfahren in der folgenden Form zur Berechnung einer unvollständigen QR-Zerlegung nutzen.

Die unvollständige QR-Zerlegung —

Für $j = 1, \dots, n$	
$\mathbf{q}_j := \mathbf{a}_j$	
Für $i = 1, \dots, j - 1$	
Y	$(i, j) \in \mathcal{M}^A$ N
$r_{ij} := (\mathbf{q}_i, \mathbf{a}_j)_2$	$r_{ij} := 0$
Für $k = 1, \dots, n$	
Y	$(j, k) \in \mathcal{M}^A$ N
$q_{kj} := q_{kj} - r_{ij}q_{ki}$	$q_{kj} := 0$
$r_{jj} := \ \mathbf{q}_j\ _2$	
$\mathbf{q}_j := \mathbf{q}_j / r_{jj}$	

Definition 5.15 Es sei

$$\mathbf{A} = \mathbf{Q}\mathbf{R} + \mathbf{F}$$

eine unvollständige QR-Zerlegung der Matrix $\mathbf{A}$. Sind $\mathbf{Q}$ und $\mathbf{R}$ regulär, dann ist

$$\mathbf{P} := \mathbf{R}^{-1}\mathbf{Q}^T \quad (5.6.2)$$

der zugehörige IQR-Präkonditionierer.

Der Speicheraufwand der Matrizen $\mathbf{Q}$ und $\mathbf{R}$ liegt zusammen geringfügig über dem 1,5-fachen der Matrix $\mathbf{A}$. Dabei stellt $\mathbf{Q}$ nicht notwendigerweise eine orthogonale Matrix dar, so daß sich mit $\mathbf{Q}^T$ auch nur eine Approximation an $\mathbf{Q}^{-1}$ ergibt.

5.7 Die unvollstandige Frobenius-Inverse

Mit der unvollstandigen Frobenius-Inversen werden wir innerhalb dieses Abschnitts eine Prakonditionierungsmatrix vorstellen, die bei vorgegebenem Muster eine Optimalitat hinsichtlich der Frobenius-Norm aufweist. Wir werden hierbei stets als Muster die Besetzungsstruktur der Matrix $\mathbf{A}$ betrachten. Es existieren jedoch auch Strategien zur adaptiven Wahl des Musters [10].

Definition 5.16 Ein Muster $\mathcal{M}$ heit regular im $\mathbb{R}^{n \times n}$, wenn

$$\mathcal{P}_{\mathcal{M}} := \{\mathbf{P} \in \mathbb{R}^{n \times n} \mid \mathbf{P} \text{ ist regular und } \mathcal{M}^{\mathbf{P}} = \mathcal{M}\} \neq \emptyset$$

gilt. Vorausgesetzt, da zu gegebenem regularen Muster $\mathcal{M}$ das folgende Minimum existiert, dann heit

$$\mathbf{P}_{\mathcal{M}}^R \in \arg \min_{\mathbf{P} \in \mathcal{P}_{\mathcal{M}}} \|\mathbf{A}\mathbf{P} - \mathbf{I}\|_F^2$$

unvollstandige Frobenius-Inverse (Frobenius-Rechtsinverse) der Matrix $\mathbf{A}$ zum Muster $\mathcal{M}$.

Wie das folgende Lemma belegt, lat sich die unvollstandige Frobenius-Inverse spaltenweise bestimmen. Da jede Spalte von $\mathbf{P}_{\mathcal{M}}^R$ unabhangig von den anderen berechnet werden kann, eignet sich die Konstruktion sehr gut zur parallelen Implementierung, siehe zum Beispiel [33].

Lemma 5.17 Sei $\mathcal{M}$ ein regulares Muster im $\mathbb{R}^{n \times n}$ und

$$\mathcal{V}_{\mathcal{M}_S(j)} := \{\mathbf{v} \in \mathbb{R}^n \mid v_i = 0 \text{ fur } i \notin \mathcal{M}_S(j)\},$$

dann ist $\mathbf{P}_{\mathcal{M}}^R = (\mathbf{p}_1, \dots, \mathbf{p}_n)$ genau dann eine unvollstandige Frobenius-Inverse zum Muster $\mathcal{M}$, wenn jeder Spaltenvektor $\mathbf{p}_j$ der Bedingung

$$\mathbf{p}_j \in \arg \min_{\mathbf{p} \in \mathcal{V}_{\mathcal{M}_S(j)} \setminus \{\mathbf{0}\}} \|\mathbf{A}\mathbf{p} - \mathbf{e}_j\|_2^2 \quad (5.7.1)$$

genugt, wobei $\mathbf{e}_j$ den j -ten Einheitsvektor im $\mathbb{R}^n$ reprasentiert.

Beweis:

Die Behauptung folgt unmittelbar aus $\|\mathbf{A}\mathbf{P}_{\mathcal{M}}^R - \mathbf{I}\|_F^2 = \sum_{j=1}^n \|\mathbf{A}\mathbf{p}_j - \mathbf{e}_j\|_2^2$. $\square$

Zur Berechnung der unvollstandigen Frobenius-Inversen mu mit Lemma 5.17 fur jeden Spaltenvektor $\mathbf{p}_j$ ein restringiertes Minimierungsproblem (5.7.1) gelost werden. Mit dem anschlieenden Satz werden wir hierzu ein aquivalentes lineares Ausgleichsproblem herleiten.

Satz 5.18 Seien $\mathbf{A} \in \mathbb{R}^{n \times n}$ regular und $m_j = |\mathcal{M}_S^{\mathbf{A}}(j)|$ fur $j = 1, \dots, n$. Desweiteren gelte $\mathcal{M}_S^{\mathbf{A}}(j) = \{i_1^j, \dots, i_{m_j}^j\}$ und fur $j = 1, \dots, n$ sei die Matrix $\mathbf{Z}_j \in \mathbb{R}^{m_j \times n}$ durch

$$(z_j)_{m\ell} := \delta_{\ell, i_m^j} \text{ fur } m = 1, \dots, m_j \text{ und } \ell = 1, \dots, n$$

unter Verwendung des Kronecker-Symbols

$$\delta_{i,j} = \begin{cases} 0, & i \neq j \\ 1, & i = j \end{cases}$$

gegeben. Dann ist $\mathbf{P}_{\mathcal{M}^A}^R = (\mathbf{p}_1, \dots, \mathbf{p}_n)$ genau dann eine Frobenius-Inverse (Frobenius-Rechtsinverse) der Matrix $\mathbf{A}$ zum Muster $\mathcal{M}^A$, wenn für jeden Spaltenvektor $\mathbf{p}_j$ der zugehörige projizierte Vektor $\tilde{\mathbf{p}}_j = \mathbf{Z}_j \mathbf{p}_j$ die Bedingung

$$\tilde{\mathbf{p}}_j \in \arg \min_{\mathbf{p} \in \mathbb{R}^{m_j} \setminus \{\mathbf{0}\}} \|\tilde{\mathbf{A}} \mathbf{Z}_j \mathbf{p} - \mathbf{e}_j\|_2^2 \quad (5.7.2)$$

erfüllt. Außerdem existiert höchstens eine Frobenius-Inverse der Matrix $\mathbf{A}$ zum Muster $\mathcal{M}^A$.

Beweis:

Die Matrix $\mathbf{Z}_j \in \mathbb{R}^{m_j \times n}$ stellt einen Isomorphismus zwischen $\mathbb{R}^{m_j} \setminus \{\mathbf{0}\}$ und $\mathcal{V}_{\mathcal{M}_S^A(j)} \setminus \{\mathbf{0}\}$ dar. Zudem gilt $\mathbf{Z}_j^T \mathbf{Z}_j \mathbf{p} = \mathbf{p}$ für alle $\mathbf{p} \in \mathcal{V}_{\mathcal{M}_S^A(j)} \setminus \{\mathbf{0}\}$ und $\mathbf{Z}_j \mathbf{Z}_j^T = \mathbf{I} \in \mathbb{R}^{m_j \times m_j}$. Somit folgt

$$\begin{aligned} \mathbf{Z}_j \left(\arg \min_{\mathbf{p} \in \mathcal{V}_{\mathcal{M}_S^A(j)} \setminus \{\mathbf{0}\}} \|\mathbf{A} \mathbf{p} - \mathbf{e}_j\|_2^2 \right) &= \mathbf{Z}_j \mathbf{Z}_j^T \left(\arg \min_{\mathbf{p} \in \mathbb{R}^{m_j} \setminus \{\mathbf{0}\}} \|\mathbf{A} \mathbf{Z}_j^T \mathbf{p} - \mathbf{e}_j\|_2^2 \right) \\ &= \arg \min_{\mathbf{p} \in \mathbb{R}^{m_j} \setminus \{\mathbf{0}\}} \|\mathbf{A} \mathbf{Z}_j^T \mathbf{p} - \mathbf{e}_j\|_2^2 \end{aligned}$$

und

$$\begin{aligned} \mathbf{Z}_j^T \left(\arg \min_{\mathbf{p} \in \mathbb{R}^{m_j} \setminus \{\mathbf{0}\}} \|\mathbf{A} \mathbf{Z}_j^T \mathbf{p} - \mathbf{e}_j\|_2^2 \right) &= \mathbf{Z}_j^T \mathbf{Z}_j \left(\arg \min_{\mathbf{p} \in \mathcal{V}_{\mathcal{M}_S^A(j)} \setminus \{\mathbf{0}\}} \|\mathbf{A} \mathbf{Z}_j^T \mathbf{Z}_j \mathbf{p} - \mathbf{e}_j\|_2^2 \right) \\ &= \arg \min_{\mathbf{p} \in \mathcal{V}_{\mathcal{M}_S^A(j)} \setminus \{\mathbf{0}\}} \|\mathbf{A} \mathbf{p} - \mathbf{e}_j\|_2^2, \end{aligned}$$

wodurch $\mathbf{Z}_j$ und $\mathbf{Z}_j^T$ Isomorphismen auf den in (5.7.1) und (5.7.2) aufgeführten Mengen darstellen und hiermit der erste Teil der Behauptung folgt.

Wir kommen nun zur Eindeutigkeit der Frobenius-Inversen. Die Matrix $\mathbf{A}_j := \mathbf{A} \mathbf{Z}_j^T \in \mathbb{R}^{n \times m_j}$ besteht aus den m_j linear unabhängigen Spalten $\mathbf{a}_i$, $i \in \mathcal{M}_S^A(j)$ der Matrix $\mathbf{A}$. Somit existiert mit Satz 3.13 eine Zerlegung der Matrix $\mathbf{A}_j$ in der Form

$$\mathbf{A}_j = \mathbf{Q} \begin{pmatrix} \mathbf{R} \\ \mathbf{0} \end{pmatrix}$$

mit einer orthogonalen Matrix $\mathbf{Q} \in \mathbb{R}^{n \times n}$ und einer regulären rechten oberen Dreiecksmatrix $\mathbf{R} \in \mathbb{R}^{m_j \times m_j}$. Mit

$$\mathbf{f}_j = \begin{pmatrix} \mathbf{f}_{j1} \\ \mathbf{f}_{j2} \end{pmatrix} = \mathbf{Q}^T \mathbf{e}_j$$

erhalten wir

$$\|\mathbf{A} \mathbf{Z}_j^T \mathbf{p} - \mathbf{e}_j\|_2^2 = \|\mathbf{Q}^T (\mathbf{A}_j \mathbf{p} - \mathbf{e}_j)\|_2^2 = \left\| \begin{pmatrix} \mathbf{R} \\ \mathbf{0} \end{pmatrix} \mathbf{p} - \begin{pmatrix} \mathbf{f}_{j1} \\ \mathbf{f}_{j2} \end{pmatrix} \right\|_2^2.$$

Aufgrund der geforderten Existenz des Minimums $\min_{\mathbf{P} \in \mathcal{P}_{\mathcal{M}}} \|\mathbf{A}\mathbf{P} - \mathbf{I}\|_F^2$ gilt $\mathbf{f}_{j1} \neq \mathbf{0}$. Unter Berucksichtigung von

$$\|\mathbf{A}\mathbf{Z}_j^T \mathbf{p} - \mathbf{e}_j\|_2^2 = \|\mathbf{R}\mathbf{p} - \mathbf{f}_{j1}\|_2^2 + \|\mathbf{f}_{j2}\|_2^2$$

erhalten wir die eindeutig bestimmte j -te Spalte der Frobenius-Inversen durch

$$\mathbf{p}_j = \mathbf{Z}_j^T \left(\arg \min_{\mathbf{p} \in \mathbb{R}^{m_j} \setminus \{\mathbf{0}\}} \|\mathbf{R}\mathbf{p} - \mathbf{f}_{j1}\|_2^2 \right) = \mathbf{Z}_j^T \mathbf{R}^{-1} \mathbf{f}_{j1} \neq \mathbf{0}.$$

□

Somit laßt sich die Frobenius-Inverse auf der Basis der im Abschnitt 3.3 beschriebenen Verfahren zur QR-Zerlegung bestimmen, wobei naturlich auch weitere Algorithmen (zum Beispiel die Householder-Transformation) genutzt werden konnen.

Die Forderung nach der Existenz des Minimums innerhalb der Definition 5.16 ist wesentlich, da die Menge $\mathcal{P}_{\mathcal{M}}$ hinsichtlich der Norm nicht abgeschlossen ist. Betrachten wir das Beispiel $\mathbf{A} = (\mathbf{e}_3, \mathbf{e}_1, \mathbf{e}_2) \in \mathbb{R}^{3 \times 3}$, so folgt $\mathcal{M}_{\mathcal{S}}^{\mathbf{A}}(1) = 3$ und wir erhalten fur $\mathbf{p} \in \mathcal{V}_{\mathcal{M}_{\mathcal{S}}^{\mathbf{A}}(1)} \setminus \{\mathbf{0}\} = \{(0, 0, \lambda)^T \mid \lambda \in \mathbb{R} \setminus \{0\}\}$ die Gleichung

$$\|\mathbf{A}\mathbf{p} - \mathbf{e}_1\|_2^2 = \left\| \begin{pmatrix} 0 \\ \lambda \\ 0 \end{pmatrix} - \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} \right\|_2^2 = 1 + \lambda^2.$$

Somit existiert $\min_{\mathbf{p} \in \mathcal{V}_{\mathcal{M}_{\mathcal{S}}^{\mathbf{A}}(1)} \setminus \{\mathbf{0}\}} \|\mathbf{A}\mathbf{p} - \mathbf{e}_1\|_2^2$ nicht, und wir erhalten lediglich

$$\arg \inf_{\mathbf{p} \in \mathcal{V}_{\mathcal{M}_{\mathcal{S}}^{\mathbf{A}}(1)} \setminus \{\mathbf{0}\}} \|\mathbf{A}\mathbf{p} - \mathbf{e}_1\|_2^2 = \mathbf{0} \notin \mathcal{V}_{\mathcal{M}_{\mathcal{S}}^{\mathbf{A}}(1)} \setminus \{\mathbf{0}\}.$$

5.8 Das prakonditionierte CG-Verfahren

Die Nutzung der vorgestellten Prakonditionierer werden wir in diesem Abschnitt am Beispiel des CG-Verfahrens studieren. Sei mit $\mathbf{A} \in \mathbb{R}^{n \times n}$ eine symmetrische und positiv definite Matrix des Gleichungssystems

$$\mathbf{A}\mathbf{x} = \mathbf{b} \tag{5.8.1}$$

gegeben, so transformieren wir das System (5.8.1) unter Verwendung einer regularen Matrix $\mathbf{P}_L$ in die aquivalente Form

$$\mathbf{A}^P \mathbf{x}^P = \mathbf{b}^P \tag{5.8.2}$$

mit $\mathbf{A}^P = \mathbf{P}_L \mathbf{A} \mathbf{P}_L^T$, $\mathbf{x}^P = \mathbf{P}_L^{-T} \mathbf{x}$ und $\mathbf{b}^P = \mathbf{P}_L \mathbf{b}$. Die Matrix $\mathbf{P}_L$ kann hierbei durch die unvollstandige Cholesky-Zerlegung oder eine symmetrische Splitting-Methode (zum Beispiel das Jacobi-Verfahren oder das symmetrische Gauß-Seidel-Verfahren) festgelegt werden, wobei das Jacobi-Verfahren ausschlielich bei Matrizen mit unterschiedlichen Diagonalelementen zu einer Veranderung der Konditionszahl fuhrt und somit im Fall der zweidimensionalen Poisson-Gleichung keinen Einflu auf die Konditionszahl besitzt.

Kennzeichnen wir, wie bereits oben angedeutet, alle präkonditionierten Größen mit dem Superskript P , so läßt sich das präkonditionierte CG-Verfahren (PCG) in der Form des ursprünglichen Verfahrens der konjugierten Gradienten durch Einfügen des Superskriptes schreiben, wobei eine abschließende Ermittlung der Näherungslösung $\mathbf{x}_{m+1} = \mathbf{P}_L^T \mathbf{x}_{m+1}^P$ ergänzt werden muß.

Wir werden nun eine spezielle Darstellung des PCG-Verfahrens herleiten, die in den ursprünglichen Variablen formuliert ist. Eine formale Nutzung der ursprünglichen Veränderlichen liefert unter Verwendung der Definitionen

$$\hat{\mathbf{p}}_m := \mathbf{P}_L^T \mathbf{p}_m^P \quad \text{und} \quad \hat{\mathbf{v}}_m := \mathbf{P}_L^{-1} \mathbf{v}_m^P$$

das präkonditionierte CG-Verfahren in der vorläufigen Form:

Vorläufiges PCG-Verfahren —

Wähle $\mathbf{P}_L^{-T} \mathbf{x}_0 \in \mathbb{R}^n$		
$\mathbf{P}_L \mathbf{r}_0 := \mathbf{P}_L \mathbf{b} - \mathbf{P}_L \mathbf{A} \mathbf{P}_L^T \mathbf{P}_L^{-T} \mathbf{x}_0$		(5.8.3)
$\mathbf{P}_L^{-T} \hat{\mathbf{p}}_0 := \mathbf{P}_L \mathbf{r}_0, \alpha_0^P := \mathbf{r}_0^T \mathbf{P}_L^T \mathbf{P}_L \mathbf{r}_0$		
Für $m = 0, \dots, n-1$		
Y	$\alpha_m^P \neq 0$	N
$\mathbf{P}_L \hat{\mathbf{v}}_m := \mathbf{P}_L \mathbf{A} \mathbf{P}_L^T \mathbf{P}_L \mathbf{p}_m = \mathbf{P}_L \mathbf{A} \hat{\mathbf{p}}_m$		STOP
$\lambda_m^P := \frac{\alpha_m^P}{(\mathbf{P}_L \hat{\mathbf{v}}_m, \mathbf{P}_L \mathbf{p}_m)_2} = \frac{\alpha_m^P}{(\hat{\mathbf{v}}_m, \hat{\mathbf{p}}_m)_2}$		
$\mathbf{P}_L^{-T} \mathbf{x}_{m+1} := \mathbf{P}_L^{-T} \mathbf{x}_m + \lambda_m^P \mathbf{P}_L \mathbf{p}_m = \mathbf{P}_L^{-T} \mathbf{x}_m + \lambda_m^P \mathbf{P}_L^{-T} \hat{\mathbf{p}}_m$		
$\mathbf{P}_L \mathbf{r}_{m+1} := \mathbf{P}_L \mathbf{r}_m - \lambda_m^P \mathbf{P}_L \hat{\mathbf{v}}_m$		
$\alpha_{m+1}^P := \mathbf{r}_{m+1}^T \mathbf{P}_L^T \mathbf{P}_L \mathbf{r}_{m+1}$		
$\underbrace{\mathbf{P}_L \mathbf{p}_{m+1}}_{=\mathbf{P}_L^{-T} \hat{\mathbf{p}}_{m+1}} := \mathbf{P}_L \mathbf{r}_{m+1} + \frac{\alpha_{m+1}^P}{\alpha_m^P} \underbrace{\mathbf{P}_L \mathbf{p}_m}_{=\mathbf{P}_L^{-T} \hat{\mathbf{p}}_m}$		
		(5.8.7)

Da $\mathbf{P}_L^T \in \mathbb{R}^{n \times n}$ regulär und $\mathbf{x}_0 \in \mathbb{R}^n$ beliebig wählbar ist, kann die im Algorithmus auftretende Festlegung $\mathbf{P}_L^{-T} \mathbf{x}_0 \in \mathbb{R}^n$ durch die Wahl von $\mathbf{x}_0 \in \mathbb{R}^n$ ersetzt werden.

Linksseitige Multiplikation der Gleichungen (5.8.3), (5.8.4) und (5.8.6) mit $\mathbf{P}_L^{-1}$ und der Gleichungen (5.8.5) und (5.8.7) mit $\mathbf{P}_L^T$ liefert bei Verwendung von $\mathbf{P} = \mathbf{P}_L^T \mathbf{P}_L$ das präkonditionierte CG-Verfahren:

Algorithmus PCG-Verfahren —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$.		
$\mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0$		
$\hat{\mathbf{p}}_0 := \mathbf{P}\mathbf{r}_0$, $\alpha_0^P := (\mathbf{r}_0, \hat{\mathbf{p}}_0)_2$		
Für $m = 0, \dots, n-1$		
Y	$\alpha_m^P \neq 0$	
	N	
	$\hat{\mathbf{v}}_m := \mathbf{A}\hat{\mathbf{p}}_m$, $\lambda_m := \frac{\alpha_m^P}{(\hat{\mathbf{v}}_m, \hat{\mathbf{p}}_m)_2}$	STOP
	$\mathbf{x}_{m+1} := \mathbf{x}_m + \lambda_m^P \hat{\mathbf{p}}_m$	
	$\mathbf{r}_{m+1} := \mathbf{r}_m - \lambda_m^P \hat{\mathbf{v}}_m$	
	$\mathbf{z}_{m+1} := \mathbf{P}\mathbf{r}_{m+1}$, $\alpha_{m+1}^P := (\mathbf{r}_{m+1}, \mathbf{z}_{m+1})_2$	
	$\hat{\mathbf{p}}_{m+1} := \mathbf{z}_{m+1} + \frac{\alpha_{m+1}^P}{\alpha_m^P} \hat{\mathbf{p}}_m$	

Die Kennzeichnungen P und $\hat{}$ verdeutlichen den verbleibenden Unterschied zum ursprünglichen CG-Verfahren. Auch das PCG-Verfahren wird in der Regel als iterative Methode genutzt. Im CG-Algorithmus ermöglicht die skalare Größe α_m eine direkte Kontrolle des Residuums, wodurch ein Abbruchkriterium in der Form $\sqrt{\alpha_m} \leq \varepsilon$ anstelle $\alpha_m = 0$ verwendet werden kann. Eine analoge Vorgehensweise zur Kontrolle des Residuums ist im PCG-Verfahren nicht möglich, da $\alpha_m^P = \|\mathbf{P}_L \mathbf{r}_m\|_2^2$ keine direkte Evaluation des Residuums liefert. Daher sollte in diesem Fall stets $\|\mathbf{r}_m\|_2 \leq \varepsilon$ als Abbruchbedingung implementiert werden.

Beispiel 5.19 Wir nutzen das bereits im Beispiel 4.64 verwendete und aus der Diskretisierung der zweidimensionalen Poisson-Gleichung resultierende Gleichungssystem zum Vergleich des CG-Verfahrens mit dem PCG-Algorithmus. Es handelt sich wiederum um ein Gleichungssystem mit 40.000 Unbekannten. Zur Präkonditionierung wird hierbei das symmetrische Gauß-Seidel-Verfahren mit

$$\mathbf{P} = \mathbf{P}_{SGS} = (\mathbf{D} + \mathbf{R})^{-1} \mathbf{D} (\mathbf{D} + \mathbf{L})^{-1}$$

verwendet. Die in der folgenden Tabelle präsentierten Ergebnisse zeigen die zu erwartende Beschleunigung des CG-Verfahrens.

	CG-Verfahren	PCG-Verfahren
Iterationen	Residuenverlauf	
0	140.348	140.348
50	491.151	8.58174
100	150.025	0.0105147
150	1.83245	4.23371e-05
200	0.148948	5.42568e-08
250	0.00307128	1.69676e-11
300	2.40822e-05	6.69697e-15
336	5.07545e-07	9.04322e-17

5.9 Das präkonditionierte BiCGSTAB-Verfahren

In diesem Abschnitt werden wir einige Aspekte der Präkonditionierung am Beispiel des BiCGSTAB-Verfahrens studieren. Im Prinzip kann man aus einem beliebigen Iterationsverfahren zur Lösung von $\mathbf{Ax} = \mathbf{b}$ eine präkonditionierte Version erhalten, indem die Methode auf die Gleichung (5.0.1) angewendet wird. Im Algorithmus wird dann die Matrix $\mathbf{A}$ durch das Produkt $\mathbf{P}_L \mathbf{A} \mathbf{P}_R$ und der Vektor $\mathbf{b}$ durch $\mathbf{P}_L \mathbf{b}$ ersetzt und nach Abbruch der Iteration die Näherungslösung $\mathbf{x}_m$ zum Ausgangsproblem gemäß (5.0.2) durch $\mathbf{x}_m = \mathbf{P}_R \mathbf{x}_m^P$ aus der Approximation $\mathbf{x}_m^P$ gewonnen. Innerhalb des Iterationsverfahrens ergeben sich daher, verglichen zum Basisalgorithmus, unterschiedliche Werte, die wir mit dem Superskript P kennzeichnen. Es ist natürlich wünschenswert, daß bei vorgegebenem Startvektor $\mathbf{x}_0$ und festgelegter Genauigkeit $\varepsilon > 0$ nach Abbruch des Basisverfahrens und der präkonditionierten Variante stets Näherungslösungen $\mathbf{x}_m$ vorliegen, die der Bedingung $\|\mathbf{b} - \mathbf{Ax}_m\| \leq \varepsilon$ genügen. Betrachtet man das im obigen Sinne durch Verwendung von $\mathbf{P}_L \mathbf{A} \mathbf{P}_R$ und $\mathbf{P}_L \mathbf{b}$ präkonditionierte Verfahren, so wird der Algorithmus stets auf der Grundlage der Norm des Residuenvektors $\mathbf{r}_m^P = \mathbf{P}_L \mathbf{b} - \mathbf{P}_L \mathbf{A} \mathbf{P}_R \mathbf{x}_m^P$ terminiert, wodurch sich für $\mathbf{x}_m = \mathbf{P}_R \mathbf{x}_m^P$ das Residuum $\|\mathbf{P}_L (\mathbf{b} - \mathbf{Ax}_m)\|$ ergibt. Folglich liegt bei Nutzung einer Linkspräkonditionierung innerhalb des Verfahrens keine Kontrolle des relevanten Residuums vor. Deshalb kann auch bei Terminierung des Verfahrens keine direkte Aussage über die Güte der Approximation getroffen werden. Eine Rechtspräkonditionierung ergibt hingegen $\mathbf{r}_m^P = \mathbf{r}_m$, so daß zudem ein zuverlässiger Vergleich von Rechts- und Linkspräkonditionierern in dieser Form nicht möglich ist.

In [49] wird eine spezielle Formulierung eines präkonditionierten BiCGSTAB-Verfahrens vorgestellt, die eine Kontrolle des relevanten Residuums auch im Fall einer linksseitigen Präkonditionierung ermöglicht. Diese Variante werden wir im folgenden beschreiben. Die Vorgehensweise ist hierbei allgemeingültig und kann daher analog auf weitere Verfahren angewendet werden. Grundlegend für die Herleitung ist die Tatsache, daß die explizite Berechnung der Matrix $\mathbf{P}_L \mathbf{A} \mathbf{P}_R$ aus zwei Gründen nicht empfehlenswert ist. Zum einen sind die Präkonditionierer oftmals nur implizit gegeben (zum Beispiel die ILU und IQR), und zum anderen stellt $\mathbf{P}_L \mathbf{A} \mathbf{P}_R$ eventuell eine vollbesetzte Matrix dar, obwohl $\mathbf{A}$ schwachbesetzt ist. Die Matrix-Vektor-Multiplikationen werden daher sukzessive durchgeführt, und die Berechnung

$$\mathbf{v}^P = \mathbf{P}_L \mathbf{A} \mathbf{P}_R \mathbf{p}^P$$

kann ohne zusätzlichen Rechenaufwand in zwei Schritten gemäß

$$\begin{aligned} \mathbf{v} &= \mathbf{A}\mathbf{P}_R\mathbf{p}^P \\ \mathbf{v}^P &= \mathbf{P}_L\mathbf{v} \end{aligned} \quad (5.9.1)$$

ausgeführt werden.

Eine rechtsseitige Präkonditionierung weist keinen Einfluß auf die Kontrolle des Residuums auf, sodaß wir uns bei den weiteren Betrachtungen im Sinne der Übersichtlichkeit auf den interessanten Fall einer linksseitigen Präkonditionierung beschränken und anschließend die rechtsseitige Präkonditionierung einbeziehen werden.

Mit dem Startvektor $\mathbf{x}_0$ ergeben sich die Residuenvektoren $\mathbf{r}_0 = \mathbf{b} - \mathbf{A}\mathbf{x}_0$ und $\mathbf{r}_0^P = \mathbf{P}_L(\mathbf{b} - \mathbf{A}\mathbf{x}_0) = \mathbf{P}_L\mathbf{r}_0$. Wegen $\mathbf{p}_0^P = \mathbf{r}_0^P$ können die Vektoren $\mathbf{r}_j$, $\mathbf{r}_j^P$ und $\mathbf{p}_j^P$ im $j + 1$ -ten Iterationsschritt als bekannt vorausgesetzt werden. Das Einfügen der Größe

$$\mathbf{v}_j = \mathbf{A}\mathbf{p}_j^P \quad (5.9.2)$$

liefert

$$\mathbf{v}_j^P = \mathbf{P}_L\mathbf{v}_j.$$

Somit ergibt sich mit

$$\mathbf{s}_j = \mathbf{r}_j - \alpha_j^P \mathbf{v}_j \quad (5.9.3)$$

die Gleichung

$$\mathbf{s}_j^P = \mathbf{r}_j^P - \alpha_j^P \mathbf{v}_j^P = \mathbf{P}_L\mathbf{s}_j,$$

und es folgen durch

$$\mathbf{t}_j = \mathbf{A}\mathbf{s}_j^P \quad (5.9.4)$$

die Zusammenhänge

$$\mathbf{t}_j^P = \mathbf{P}_L\mathbf{t}_j \text{ und } \mathbf{r}_{j+1}^P = \mathbf{s}_j^P - \omega_j^P \mathbf{t}_j^P = \mathbf{P}_L(\mathbf{s}_j - \omega_j^P \mathbf{t}_j).$$

Wegen der Regularität der Matrix $\mathbf{P}_L$ ergibt sich aus $\mathbf{r}_{j+1}^P = \mathbf{P}_L\mathbf{r}_{j+1}$ direkt

$$\mathbf{r}_{j+1} = \mathbf{s}_j - \omega_j^P \mathbf{t}_j.$$

Durch die Einführung der Hilfsvektoren $\mathbf{v}_j$, $\mathbf{s}_j$ und $\mathbf{t}_j$ gemäß (5.9.2), (5.9.3) und (5.9.4) sind wir in der Lage das Residuum $\|\mathbf{r}_j\|$ anstelle $\|\mathbf{r}_j^P\|$ ohne zusätzlichen Rechenaufwand zu kontrollieren. Somit ergibt sich der präkonditionierte BiCGSTAB-Algorithmus in der im folgenden Diagramm dargestellten Form.

Aufgrund der Darstellung

$$\begin{aligned} \mathbf{x}_{j+1} &= \mathbf{P}_R\mathbf{x}_{j+1}^P = \mathbf{P}_R\mathbf{x}_j^P + \alpha_j^P \mathbf{P}_R\mathbf{p}_j^P + \omega_j^P \mathbf{P}_R\mathbf{s}_j^P \\ &= \mathbf{x}_j + \mathbf{P}_R(\alpha_j^P \mathbf{p}_j^P + \omega_j^P \mathbf{s}_j^P) \end{aligned} \quad (5.9.5)$$

kann die Gleichung (5.9.6) durch (5.9.5) ersetzt werden, wodurch die Gleichung (5.9.7) entfällt, und der Algorithmus direkt in der gesuchten Näherungslösung $\mathbf{x}_j$ formuliert wird. Während die Kontrolle des Residuums $\|\mathbf{r}_j\|$ keinen zusätzlichen Rechenaufwand verursacht, benötigt die Nutzung der Gleichung (5.9.5) pro Iteration eine Matrix-Vektor-Multiplikation mit $\mathbf{P}_R$, die in der im Diagramm dargestellten Form nur einmal außerhalb der Iteration durchgeführt werden muß.

Das präkonditionierte BiCGSTAB-Verfahren —

Wähle $\mathbf{x}_0 \in \mathbb{R}^n$ und $\varepsilon > 0$
$\mathbf{r}_0 = \mathbf{p}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0, \mathbf{r}_0^P = \mathbf{p}_0^P := \mathbf{P}_L \mathbf{r}_0$
$\rho_0^P := (\mathbf{r}_0^P, \mathbf{r}_0^P)_2, j := 0$
Solange $\ \mathbf{r}_j\ _2 > \varepsilon$
$\mathbf{v}_j := \mathbf{A} \mathbf{P}_R \mathbf{p}_j^P, \mathbf{v}_j^P := \mathbf{P}_L \mathbf{v}_j$
$\alpha_j^P := \frac{\rho_j^P}{(\mathbf{v}_j^P, \mathbf{r}_0^P)_2}, \mathbf{s}_j := \mathbf{r}_j - \alpha_j^P \mathbf{v}_j, \mathbf{s}_j^P := \mathbf{P}_L \mathbf{s}_j$
$\mathbf{t}_j := \mathbf{A} \mathbf{P}_R \mathbf{s}_j^P, \mathbf{t}_j^P := \mathbf{P}_L \mathbf{t}_j$
$\omega_j^P := \frac{(\mathbf{t}_j^P, \mathbf{s}_j^P)_2}{(\mathbf{t}_j^P, \mathbf{t}_j^P)_2}$
$\mathbf{x}_{j+1}^P := \mathbf{x}_j^P + \alpha_j^P \mathbf{p}_j^P + \omega_j^P \mathbf{s}_j^P$ (5.9.6)
$\mathbf{r}_{j+1} := \mathbf{s}_j - \omega_j^P \mathbf{t}_j, \mathbf{r}_{j+1}^P := \mathbf{s}_j^P - \omega_j^P \mathbf{t}_j^P$
$\rho_{j+1}^P := (\mathbf{r}_{j+1}^P, \mathbf{r}_0^P)_2, \beta_j^P := \frac{\alpha_j^P}{\omega_j^P} \frac{\rho_{j+1}^P}{\rho_j^P}$
$\mathbf{p}_{j+1}^P := \mathbf{r}_{j+1}^P + \beta_j^P (\mathbf{p}_j^P - \omega_j^P \mathbf{v}_j^P), j := j + 1$
$\mathbf{x}_j := \mathbf{P}_R \mathbf{x}_j^P$ (5.9.7)

5.10 Vergleich der Präkonditionierer

Zur Analyse der Präkonditionierungstechniken betrachten wir wiederum das Modellproblem der Konvektions-Diffusions-Gleichung mit $N^2 = 10.000$ Punkten. Die Bilder 5.1 bis 5.5 zeigen die bereits im Abschnitt 4.3.2.10 diskutierten Konvergenzverläufe der einzelnen Projektionsverfahren für einen Diffusionsparameter $\varepsilon = 0.1$. Die zweite Kurve zeigt jeweils die durch eine Rechtspräkonditionierung mittels einer unvollständigen LU-Zerlegung erzielte Beschleunigung des Grundalgorithmus. Hierbei wurde die benötigte

Iterationszahl auf ca. 30% des jeweiligen Basisverfahrens reduziert. Obwohl die Verwendung einer unvollständigen LU-Zerlegung neben ihrer Berechnung stets zu einer Verdopplung der Matrix-Vektor-Multiplikationen führt, ergibt sich insgesamt bei allen Gleichungssystemlöstern eine Rechenzeitersparnis verglichen mit dem nicht präkonditionierten Verfahren.

In der Arbeit [48] wird ein Rechenzeitvergleich unterschiedlich präkonditionierter Krylov-Unterraum-Verfahren bei der Simulation reibungsfreier und reibungsbehafteter Strömungsfelder präsentiert. Bei der betrachteten großen Bandbreite verschiedenster Problemstellungen zeigte sich dabei das BiCGSTAB-Verfahren in Kombination mit einer unvollständigen LU-Zerlegung als sehr effizient, wobei die Nutzung der Präkonditionierung zu einer immensen Beschleunigung des Gesamtverfahrens führte.

Die Abbildungen 5.6 bis 5.10 korrespondieren mit der im Abschnitt 5.9 vorgestellten präkonditionierten Variante des BiCGSTAB-Verfahrens, wodurch ein Vergleich von links- und rechtsseitiger Präkonditionierung ermöglicht wird. Im Unterschied zu den vorangegangenen Abbildungen wurde die Konvektions-Diffusions-Gleichung mit einem Diffusionsparameter $\varepsilon = 0.01$ genutzt, wodurch der unsymmetrische Anteil der Matrix (1.0.9) einen größeren Einfluß verglichen zu $\varepsilon = 0.1$ besitzt. Zur Unterscheidung der Präkonditionierer führen wir die in der Tabelle 5.2 aufgelisteten Abkürzungen ein.

Abkürzung	Präkonditionierer
ILU	Unvollständige LU-Zerlegung
Jacobi	Jacobi-Präkonditionierer
GS	Gauß-Seidel-Präkonditionierer
SGS	Symm. Gauß-Seidel-Präkonditionierer
POLY(3)	Polynomialer Präkonditionierer gemäß (5.2.5) mit $m = 3$
Z_1	Zeilenskalierung gemäß Betragssummennorm
Z_2	Zeilenskalierung gemäß euklidischer Norm
S_1	Spaltenskalierung gemäß Betragssummennorm
S_2	Spaltenskalierung gemäß euklidischer Norm

Tabelle 5.2 Präkonditionierer und deren Abkürzung

Zudem wird durch den Zusatz L beziehungsweise R die Verwendung des jeweiligen Präkonditionierers in einer links- respektive rechtsseitigen Formulierung gekennzeichnet.

Die Präkonditionierer lassen sich aus dem Gesichtspunkt der Effizienz in drei Gruppen unterteilen. Zunächst fassen wir hierzu die in der Abbildung 5.6 aufgeführte unvollständige LU-Zerlegung und den ebenfalls im Bild 5.6 enthaltenen symmetrischen Gauß-Seidel-Präkonditionierer in der ersten Gruppe zusammen. Die zweite Gruppe beinhaltet alle in den Abbildungen 5.7, 5.9 und 5.10 dargestellten Präkonditionierer. Bezogen auf die Algorithmen der ersten Gruppe führen diese Methoden auf etwa die dreifache Anzahl an Iterationen. Die letzte Gruppe der für diesen Testfall ineffizienten Verfahren beinhaltet bei den untersuchten Techniken den in Abbildung 5.8 aufgeführten polynomialen Präkonditionierer. Diese Methode ergab eine deutlich höhere Iterationszahl und benötigt in der

genutzten Form ($m = 3$) bezogen auf den Basisalgorithmus viermal so viele Matrix-Vektor-Multiplikationen. Desweiteren lassen sich die unvollständige QR-Zerlegung und die unvollständige Frobenius-Inverse in diese Gruppe einordnen. Beide Prädiktionierer erwiesen sich bereits bei ihrer Berechnung als unakzeptabel aufwendig und wurden daher in den Abbildungen nicht berücksichtigt. Eine Rechenzeiterparnis bei der Berechnung der Frobenius-Inversen kann allerdings durch die gute Parallelisierbarkeit dieser Methode erzielt werden.

Beim Vergleich der links- und rechtsseitigen Prädiktionierung ergaben sich keine gravierenden Unterschiede, wobei zumeist bei der rechtsseitigen Prädiktionierung eine geringfügig kleinere Iterationszahl vorlag.

Abschließend können wir feststellen, daß sich eine rechtsseitige Prädiktionierung auf der Basis einer unvollständigen LU-Zerlegung beziehungsweise des symmetrischen Gauß-Seidel-Verfahrens als eine sehr effiziente Vorgehensweise herausgestellt hat. Diese Ergebnisse decken sich hervorragend mit der in [49] präsentierten Studie unterschiedlicher Prädiktionierer im Kontext der Euler-Gleichungen. Der zum symmetrischen Gauß-Seidel-Verfahren assoziierte Prädiktionierer benötigt im Gegensatz zur unvollständigen LU-Zerlegung keinen Speicherplatz und ist zudem direkt durch die Matrix des Gleichungssystems gegeben, wodurch die Berechnung des Prädiktionierers entfällt. Im Sinne eines möglichst geringen Speicherplatzbedarfs stellt folglich das symmetrische Gauß-Seidel-Verfahren eine effiziente Alternative zur bewährten unvollständigen LU-Zerlegung dar.

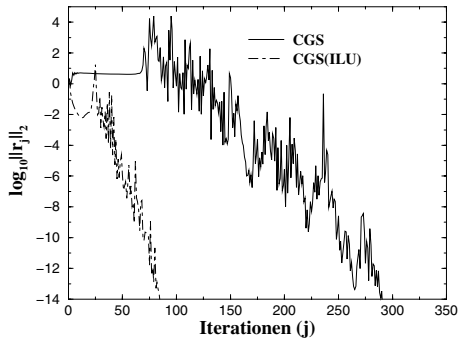


Bild 5.1 Konvergenzverlauf des CGS-Verfahrens

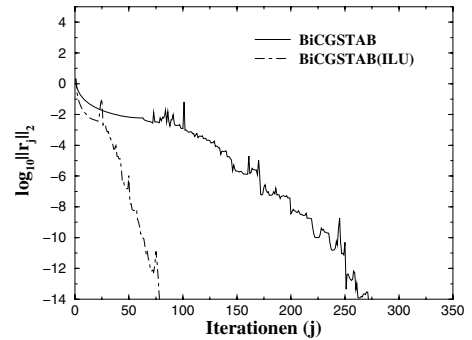


Bild 5.2 Konvergenzverlauf des BiCGSTAB-Verfahrens

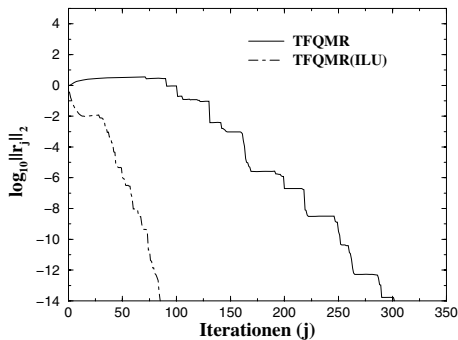


Bild 5.3 Konvergenzverlauf des TFQMR-Verfahrens

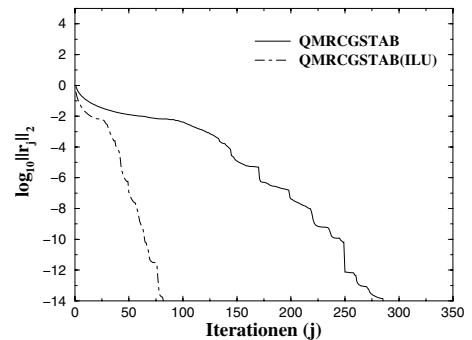


Bild 5.4 Konvergenzverlauf des QMRGSTAB-Verfahrens

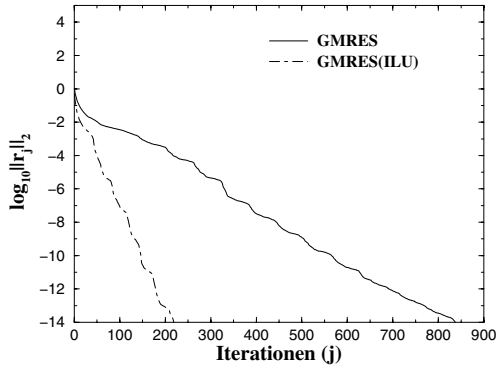


Bild 5.5 Konvergenzverlauf des GMRES-Verfahrens

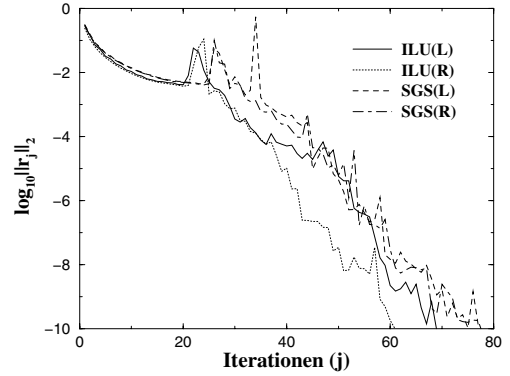


Bild 5.6 Konvergenzverläufe des präkonditionierten BiCGSTAB-Verfahrens

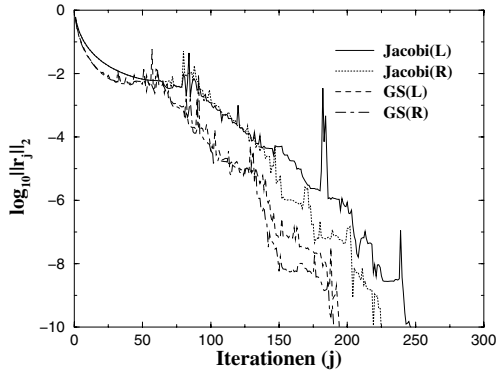


Bild 5.7 Konvergenzverläufe des präkonditionierten BiCGSTAB-Verfahrens

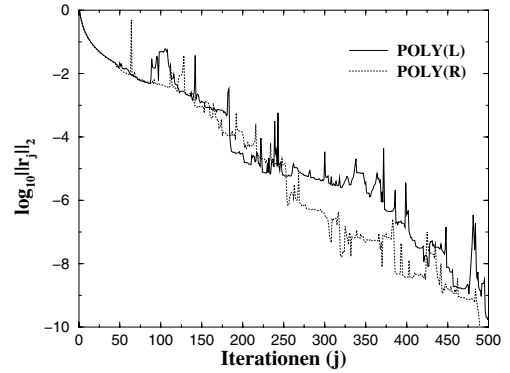


Bild 5.8 Konvergenzverläufe des präkonditionierten BiCGSTAB-Verfahrens

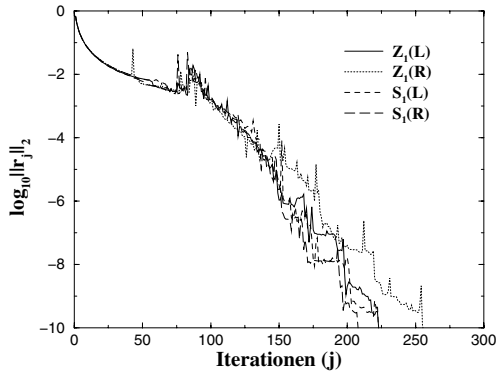


Bild 5.9 Konvergenzverläufe des präkonditionierten BiCGSTAB-Verfahrens

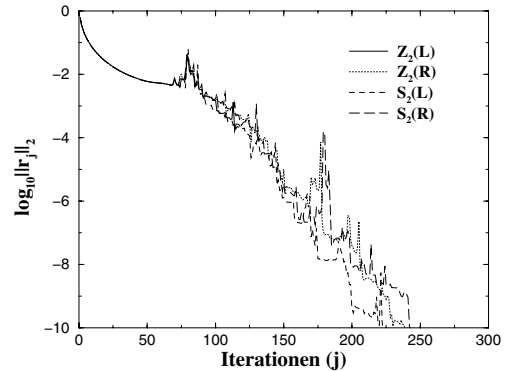


Bild 5.10 Konvergenzverläufe des präkonditionierten BiCGSTAB-Verfahrens

5.11 Übungsaufgaben

Aufgabe 1:

Zeigen Sie: Sei $\mathbf{P}_L = \mathbf{P}_R^T \in \mathbb{R}^{n \times n}$ regulär und $\mathbf{A} \in \mathbb{R}^{n \times n}$ symmetrisch und positiv definit, dann ist auch die Matrix

$$\mathbf{A}^P := \mathbf{P}_L \mathbf{A} \mathbf{P}_R$$

symmetrisch und positiv definit.

Aufgabe 2:

Berechnen Sie eine unvollständige LU-Zerlegung der Matrix

$$\mathbf{A} = \begin{pmatrix} 3 & 13 & 1 & 0 & 0 \\ 0 & 1 & 0 & 0 & 2 \\ 0 & 0 & 1 & 1 & 0 \\ 1 & 0 & 1 & 3 & 1 \\ 1 & 4 & 0 & 1 & 4 \end{pmatrix}$$

und vergleichen Sie diese mit einer vollständigen LU-Zerlegung, das heisst der LR-Zerlegung.

Aufgabe 3:

Bestimmen Sie die Präkonditionierer $\mathbf{P}_{Jac}$, $\mathbf{P}_{GS}$ und $\mathbf{P}_{SGS}$ zur Matrix

$$\mathbf{A} = \begin{pmatrix} 10 & 7 & 1 \\ 4 & 6 & 2 \\ 8 & 5 & 12 \end{pmatrix}$$

und ermitteln Sie die Konditionszahlen

$$\text{cond}_\infty(\mathbf{A}), \text{cond}_\infty(\mathbf{P}_{Jac}\mathbf{A}), \text{cond}_\infty(\mathbf{P}_{GS}\mathbf{A}) \text{ sowie } \text{cond}_\infty(\mathbf{P}_{SGS}\mathbf{A}).$$

Aufgabe 4:

Zeigen Sie, dass für eine symmetrische, positiv definite Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ die beidseitige Präkonditionierung

$$\mathbf{A}^P = \mathbf{P}_L \mathbf{A} \mathbf{P}_R$$

mit

$$\mathbf{P}_L = \mathbf{P}_R^T = (\mathbf{D} + \mathbf{L})^{-1} \mathbf{D}^{1/2}$$

exisitiert. Berechnen Sie $\mathbf{A}^P$ für

$$\mathbf{A} = \begin{pmatrix} 10 & 4 & 1 \\ 4 & 12 & 2 \\ 1 & 2 & 6 \end{pmatrix}$$

und vergleichen Sie die Konditionszahlen $\text{cond}_\infty(\mathbf{A})$ und $\text{cond}_\infty(\mathbf{A}^P)$.

A Implementierungen in MATLAB

Innerhalb dieses Anhangs werden mögliche Implementierungen der häufig verwendeten Verfahren in MATLAB [1] vorgestellt. Die Algorithmen wurden von C. Vömel entwickelt und sind über

<http://www.vieweg.de/tu/8y>

öffentlich zugänglich.

Mit Ausnahme des Arnoldi-Algorithmus können alle aufgeführten Verfahren zur Lösung des Problems $\mathbf{Ax} = \mathbf{b}$ mit $\mathbf{A} \in \mathbb{R}^{n \times n}$, $\mathbf{b} \in \mathbb{R}^n$ unter Verwendung der Notationen

```
x = Name(A,b);
x = Name(A,b,tol);
x = Name(A,b,tol,maxit);
x = Name(A,b,tol,maxit,x0);
```

aufgerufen werden, wobei **Name** jeweils durch die gewünschte Verfahrensbezeichnung ersetzt werden muß. Es gelten hierbei für die optionalen Angaben die in der folgenden Tabelle angegebenen Einstellungen.

Eingangsvariable	Beschreibung	Vorgabe
tol	Genauigkeitsvorgabe (Toleranz)	10^{-6}
maxint	Maximale Anzahl an Iterationen	$\min\{n, 30\}$
x0	Startvektor	Nullvektor

Bei jeder dieser Methoden wird eine Überprüfung der Eingangsparameter und der Vorlage einer trivialen Lösung mittels der Programme Argumente.m respektive Trivial.m durchgeführt.

Programm SteilsterAbstieg.m: Verfahren des steilsten Abstiegs

```
function x = SteilsterAbstieg(A,b,tol,maxit,x0)
%
% CHECK THE INPUT ARGUMENTS
%   Argumente
%
% CHECK FOR TRIVIAL SOLUTION
%   Trivial
%
% MAIN ALGORITHM
```

```

%
    tolb = tol * normb;
    r = b - A * x;
% iterate
    for i = 1:maxit
        normr = norm(r);
        if (normr <= tolb)
%            konvergiert
            break
        end
        q = A*r;
        if (q == 0)
            error('Matrix ist singulaer.');
```

Programm KonjugGrad.m: CG-Verfahren

```

function x = KonjugGrad(A,b,tol,maxit,x0)
%
% CHECK THE INPUT ARGUMENTS
    Argumente

% CHECK FOR TRIVIAL SOLUTION
    Trivial

% MAIN ALGORITHM
%
    tolb = tol * normb;
    r = b - A * x;
    p = r;
    normr = norm(r);
    alpha = normr^2;
% iterate
    for i = 1:maxit
        if (normr <= tolb)
```

```

%    converged
    break
end
v = A*p;
if (v == 0)
    error('Matrix ist singulaer.');
```

end

```

vr = (v'*r);
if(vr <=0)
    error('Matrix ist nicht positiv definit.');
```

end

```

lambda = alpha/vr;
x = x + lambda * p;
r = r - lambda * v;
alpha2 = alpha;
normr = norm(r);
alpha = normr^2;
p = r + alpha/alpha2 * p;
end
normr = norm(b-A*x);
if (normr > tol)
    warning('Verfahren hat nicht konvergiert')
```

end

```

return
```

Programm Arnoldi.m: Arnoldi-Algorithmus

```

function [V,H] = Arnoldi(A,r0,m);
%
% INPUT VARIABLEN:
% -----
% V -- n*m orthogonale Matrix
% H -- m*m obere Hessenberg Matrix
%
% OUTPUT VARIABLEN:
% -----
% A: quadratische n*n Matrix.
% r0: Spaltenvektor
% m: Anzahl der Spalten von Q und H

% CHECK THE INPUT ARGUMENTS
if (nargin < 3)
    error('Funktion braucht mehr Parameter.');
```

else

```

    [mm,n] = size(A);
    if (mm ~= n)
        error('Matrix ist nicht quadratisch.');
```

end

A Implementierungen in MATLAB

Innerhalb dieses Anhangs werden mögliche Implementierungen der häufig verwendeten Verfahren in MATLAB [1] vorgestellt. Die Algorithmen wurden von C. Vömel entwickelt und sind über

<http://www.vieweg.de/tu/8y>

öffentlich zugänglich.

Mit Ausnahme des Arnoldi-Algorithmus können alle aufgeführten Verfahren zur Lösung des Problems $\mathbf{Ax} = \mathbf{b}$ mit $\mathbf{A} \in \mathbb{R}^{n \times n}$, $\mathbf{b} \in \mathbb{R}^n$ unter Verwendung der Notationen

```
x = Name(A,b);
x = Name(A,b,tol);
x = Name(A,b,tol,maxit);
x = Name(A,b,tol,maxit,x0);
```

aufgerufen werden, wobei **Name** jeweils durch die gewünschte Verfahrensbezeichnung ersetzt werden muß. Es gelten hierbei für die optionalen Angaben die in der folgenden Tabelle angegebenen Einstellungen.

Eingangsvariable	Beschreibung	Vorgabe
tol	Genauigkeitsvorgabe (Toleranz)	10^{-6}
maxint	Maximale Anzahl an Iterationen	$\min\{n, 30\}$
x0	Startvektor	Nullvektor

Bei jeder dieser Methoden wird eine Überprüfung der Eingangsparameter und der Vorlage einer trivialen Lösung mittels der Programme Argumente.m respektive Trivial.m durchgeführt.

Programm SteilsterAbstieg.m: Verfahren des steilsten Abstiegs

```
function x = SteilsterAbstieg(A,b,tol,maxit,x0)
%
% CHECK THE INPUT ARGUMENTS
%   Argumente
%
% CHECK FOR TRIVIAL SOLUTION
%   Trivial
%
% MAIN ALGORITHM
```

```

%
    tolb = tol * normb;
    r = b - A * x;
% iterate
    for i = 1:maxit
        normr = norm(r);
        if (normr <= tolb)
%            konvergiert
            break
        end
        q = A*r;
        if (q == 0)
            error('Matrix ist singulaer.');
```

Programm KonjugGrad.m: CG-Verfahren

```

function x = KonjugGrad(A,b,tol,maxit,x0)
%
% CHECK THE INPUT ARGUMENTS
    Argumente

% CHECK FOR TRIVIAL SOLUTION
    Trivial

% MAIN ALGORITHM
%
    tolb = tol * normb;
    r = b - A * x;
    p = r;
    normr = norm(r);
    alpha = normr^2;
% iterate
    for i = 1:maxit
        if (normr <= tolb)
```

```

%    converged
    break
end
v = A*p;
if (v == 0)
    error('Matrix ist singulaer.');
```

end

```

vr = (v'*r);
if(vr <=0)
    error('Matrix ist nicht positiv definit.');
```

end

```

lambda = alpha/vr;
x = x + lambda * p;
r = r - lambda * v;
alpha2 = alpha;
normr = norm(r);
alpha = normr^2;
p = r + alpha/alpha2 * p;
end
normr = norm(b-A*x);
if (normr > tol)
    warning('Verfahren hat nicht konvergiert')
```

end

```

return
```

Programm Arnoldi.m: Arnoldi-Algorithmus

```

function [V,H] = Arnoldi(A,r0,m);
%
% INPUT VARIABLEN:
% -----
% V -- n*m orthogonale Matrix
% H -- m*m obere Hessenberg Matrix
%
% OUTPUT VARIABLEN:
% -----
% A: quadratische n*n Matrix.
% r0: Spaltenvektor
% m: Anzahl der Spalten von Q und H

% CHECK THE INPUT ARGUMENTS
if (nargin < 3)
    error('Funktion braucht mehr Parameter.');
```

else

```

    [mm,n] = size(A);
    if (mm ~= n)
        error('Matrix ist nicht quadratisch.');
```

end

```

    if ~isequal(size(r0),[mm,1])
        error('r0 hat nicht die richtige Dimension.');
```

end

```

    if ((m > n) || (m <= 0))
        error('Falsche Spaltenanzahl.');
```

end

end

```

%
% MAIN ALGORITHM
%
```

```

V = zeros(n,m);
H = zeros(m,m);
w = r0;
normw = norm(w);
if (normw == 0)
    error('Arnoldi break-down.');
```

end

```

V(:,1) = w/normw;
for j = 1:m
    w = A*V(:,j);
    H(1:j,j) = V(:,1:j)'*w;
    w = w - V(:,1:j)*H(1:j,j);
    if(j == m), break; end
    normw = norm(w);
    if (normw == 0)
        error('Arnoldi break-down.');
```

end

```

    V(:,j+1) = w/normw;
    H(j+1,j) = normw;
end
return
```

Programm RestartedFOM.m: Restarted FOM
--

```

function x = RestartedFOM(A,b,tol,maxit,x0)
%
% CHECK THE INPUT ARGUMENTS
    Argumente

% CHECK FOR TRIVIAL SOLUTION
    Trivial

% MAIN ALGORITHM
%
    tolb = tol * normb;
    r = b - A * x;
    normr = norm(r);
% iterate
```

```

for i = 1:maxit
    if (normr <= tolb)
%       converged
        break
    end
    [V,H] = Arnoldi(A,r,i);
    e = normr*eye(i,1);
    alpha = H\e;
    x = x + V * alpha;
    r = b - A * x;
    normr = norm(r);
end
if (normr > tolb)
    warning('Verfahren hat nicht konvergiert')
end
return

```

Programm Gmres.m: GMRES-Algorithmus

```

function x = Gmres(A,b,tol,maxit,x0)
%
% CHECK THE INPUT ARGUMENTS
    Argumente

% CHECK FOR TRIVIAL SOLUTION
    Trivial

% MAIN ALGORITHM
%
    tolb = tol * normb;
    r = b - A * x;
    normr = norm(r);
    Gamma=normr*eye(maxit+1,1);
% Hessenberg Matrix
    H = zeros(maxit+1);
% orthogonale Basis
    V = zeros(n,maxit+1);
% Givensparameter
    c=zeros(maxit+1,1);
    s=zeros(maxit+1,1);
% Dimension des Krylovraums
    kdim = 0;

    V(:,1) = r/normr;
% iterate
    for j = 1:maxit
        if (normr <= tolb)
%           converged

```

```

%    berechne
%    break
else
    kdim = kdim+1;
end
w = A*V(:,j);
% Modified Gram-Schmidt
for i =1:j
    H(i,j) = w'*V(:,i);
    w = w - H(i,j)*V(:,i);
end
H(j+1,j) = norm(w);
% transformiere j-te Spalte der Hessenbergmatrix mit akkumulierten
% Givensrotationen (H wird mit R ueberschrieben)
for i = 1:j-1
    H(i:i+1,j) = [c(i),s(i);-s(i),c(i)] * H(i:i+1,j);
end
beta=norm(H(j:j+1,j));
if(beta ~= 0)
    V(:,j+1)=w/H(j+1,j);
else
%    'happy breakdown'
end
if beta~=0
%    Berechne neue Givensrotation
    s(j)=H(j+1,j)/beta;
    c(j)=H(j,j)/beta;
%    Anwendung auf H(j:j+1,j)
    H(j,j) = beta;
    H(j+1,j)=0;
%    Anwendung auf b(j:j+1,j)
    Gamma(j+1) = -s(j)*Gamma(j);
    Gamma(j) = c(j)*Gamma(j);
end
normr = abs(Gamma(j+1));
end
% Berechne die Basiskoeffizienten, loese H*Alpha = Gamma
Alpha = H(1:kdim,1:kdim)\Gamma(1:kdim);
x = x + V(:,1:kdim) * Alpha;
r = b - A*x;
normr = norm(r);
if (normr > tol)
    warning('Verfahren hat nicht konvergiert')
end
return

```

Programm Bicg.m: BiCG-Algorithmus

```

function x = Bicg(A,b,tol,maxit,x0)
%
% CHECK THE INPUT ARGUMENTS
  Argumente

% CHECK FOR TRIVIAL SOLUTION
  Trivial

% MAIN ALGORITHM
%
  tolb = tol * normb;
  r = b - A * x;
  p = r;
  rstar = r;
  pstar = r;
  w = r'*rstar;
% iterate
  for i = 1:maxit
    normr = norm(r);
    if (normr <= tolb)
%      converged
      break
    end
    v = A*p;
    vp = (v'*pstar);
    if(vp ==0)
      error('Bicg break-down. ');
    end
    if(w ==0)
      error('Bicg Loesung stagniert. ');
    end
    alpha = w/vp;
    x = x + alpha * p;
    r = r - alpha * v;
    rstar = rstar - alpha * A' * pstar;
    w1 = r'*rstar;
    beta = w1/w;
    p = r + beta * p;
    pstar = rstar + beta * pstar;
    w = w1;
  end
  normr = norm(b-A*x);
  if (normr > tolb)
    warning('Verfahren hat nicht konvergiert')
  end
  return

```

Programm Cgs.m: CGS-Algorithmus

```

function x = Cgs(A,b,tol,maxit,x0)
%
% CHECK THE INPUT ARGUMENTS
    Argumente
% CHECK FOR TRIVIAL SOLUTION
    Trivial

% MAIN ALGORITHM
%
    tolb = tol * normb;
    r0 = b - A * x;
    r = r0;
    rr0 = r'*r0;
    p = r;
    u = r;
% iterate
    for i = 1:maxit
        normr = norm(r);
        if (normr <= tolb)
%            konvergiert
            break
        end
        v = A*p;
        vr0 = v'*r0;
        if(vr0 ==0)
            error('Cgs break-down.');


```

 end
 if(rr0 ==0)
 error('Cgs Loesung stagniert.');
```



```

 end
 alpha = rr0/vr0;
 q = u - alpha * v;
 t = u + q;
 x = x + alpha * t;
 r = r - alpha * A * t;
 r1r0 = r'*r0;
 beta = r1r0/rr0;
 u = r + beta * q;
 p = u + beta * (q + beta * p);
 rr0 = r1r0;
 end
 normr = norm(b-A*x);
 if (normr > tolb)
 warning('Verfahren hat nicht konvergiert')
 end
 return

```


```

Programm Bicgstab.m: BiCGSTAB-Algorithmus

```

function x = Bicgstab(A,b,tol,maxit,x0)
%
% CHECK THE INPUT ARGUMENTS
  Argumente

% CHECK FOR TRIVIAL SOLUTION
  Trivial

% MAIN ALGORITHM
%
  tolb = tol * normb;
  r0 = b - A * x;
  r = r0;
  rr0 = r'*r0;
  p = r;
% iterate
  for i = 1:maxit
    normr = norm(r);
    if (normr <= tolb)
%      converged
      break
    end
    v = A*p;
    vr0 = v'*r0;
    if(vr0 ==0)
      error('Bicgstab break-down. ');
    end
    if(rr0 ==0)
      error('Bicgstab Loesung stagniert. ');
    end
    alpha = rr0/vr0;
    s = r - alpha * v;
    t = A * s;
    ts = s'*t;
    tt = t'*t;
    if((tt ==0) || (ts==0))
      error('Bicgstab break-down. ');
    end
    omega = ts/tt;
    x = x + alpha * p + omega * s;
    r = s - omega * t;
    r1r0 = r'*r0;
    beta = (alpha*r1r0)/(omega*rr0);
    p = r + beta * (p - omega * v);
    rr0 = r1r0;
  end
end

```

```

normr = norm(b-A*x);
if (normr > tolb)
    warning('Verfahren hat nicht konvergiert')
end
return

```

Programm Tfqmr.m: TFQMR-Algorithmus

```

function x = Tfqmr(A,b,tol,maxit,x0)
%
% CHECK THE INPUT ARGUMENTS
  Argumente

% CHECK FOR TRIVIAL SOLUTION
  Trivial

% MAIN ALGORITHM
%
  tolb = tol * normb;
  r0 = b - A * x;
  y1 = r0;
  w = r0;
  tau = norm(r0);

  if(tau > tolb)
    v = A * y1;
    d = zeros(n,1);
    eta = 0;
    theta = 0;
    rho1 = tau^2;

    for j = 1:maxit
      rho2 = v'*r0;
      if(rho2 ==0)
        error('Tfqmr break-down.');
```

```

        x = x + eta * d;
        tau = tau * theta * c;
        if (sqrt(m+1)*tau <= tolb)
%           konvergiert
            break
        end
        y1 = y2;
    end
    rho2 = w'*r0;
    beta = rho2/rho1;
    y1 = w + beta * y2;
    v = A*y1 + beta *(A*y2+beta*v);
    rho1 = rho2;
end
end
normr = norm(b-A*x);
if (normr > tolb)
    warning('Verfahren hat nicht konvergiert')
end
return

```

Programm Qmrcgstab.m: QMRCGSTAB-Algorithmus

```

function x = Qmrcgstab(A,b,tol,maxit,x0)
%
% CHECK THE INPUT ARGUMENTS
    Argumente

% CHECK FOR TRIVIAL SOLUTION
    Trivial

% MAIN ALGORITHM
%
    tolb = tol * normb;
    r0 = b - A * x;
    p = r0;
    tau = norm(r0);
    v = A * p;
    r = r0;
    if(tau > tolb)
        d = zeros(n,1);
        eta = 0;
        theta = 0;
        rho1 = tau^2;
        for j = 1:maxit
            rho2 = v'*r0;
            if(rho2 ==0)
                error('Qmrcgstab break-down.');
```

```

end
if(rho1 ==0)
    error('Qmrcgstab Loesung stagniert.');
```

end

```

alpha = rho1/rho2;
s = r - alpha * v ;
% Erster Quasi-Minimierungsschritt
theta2 = norm(s)/tau;
c = 1/sqrt(1+(theta2)^2);
tau2 = tau * theta2 * c;
eta2 = c^2 * alpha;
d2 = p + ((theta2)^2*eta/alpha) * d;
x2 = x + eta2 * d2;
% Berechne t, omega, und das update des Residuums
t = A * s;
uu = s'*t;
vv = t'*t;
if( vv==0)
    error('Matrix ist singulaer.');
```

end

```

omega = uu/vv;
if( omega==0 )
    error('Qmrcgstab Loesung stagniert.');
```

end

```

r = s - omega * t;
% Zweiter Quasi-Minimierungsschritt
theta = norm(r)/tau2;
c = 1/sqrt(1+theta^2);
tau = tau2 * theta * c;
eta = c^2 * omega;
d = s + (((theta2)^2)*eta2/omega) * d2;
x = x2 + eta * d;
if (sqrt(2*j+1)*abs(tau) <= tolB)
%     konvergiert
    break
end
rho2 = r'*r0;
if(rho2 ==0)
    error('Qmrcgstab Loesung stagniert.');
```

end

```

beta = (alpha*rho2)/(omega*rho1);
p = r + beta * (p - omega * v);
v = A * p;
rho1 = rho2;
end
end
normr = norm(b-A*x);
if (normr > tolB)
```

```
    warning('Verfahren hat nicht konvergiert')
end
return
```

Programm Argumente.m: Überprüfung der Eingangsparameter

```
if (nargin < 2)
    error('Funktion braucht mehr Parameter.');
```

else

```
    [m,n] = size(A);
    if (m ~= n)
        error('Matrix ist nicht quadratisch.');
```

end

```
    if ~isequal(size(b),[m,1])
        error('Rechte Seite hat nicht die richtige Dimension.');
```

end

```
end
if (nargin < 3) | isempty(tol)
    tol = 1.0e-6;
end
if (nargin < 4) | isempty(maxit)
    maxit = min(n,30);
end
if (nargin < 5) | isempty(x0)
    x = zeros(n,1);
else
    if ~isequal(size(x0),[n,1])
        error('Startvector x0 hat nicht die richtige Dimension.');
```

end

```
    x = x0;
end
```

Programm Trivial.m: Überprüfung auf triviale Lösung

```
normb = norm(b);
if (normb == 0)
    x = zeros(n,1);
    return
end
```

Literaturverzeichnis

- [1] *MATLAB*. <http://www.mathworks.com>.
- [2] K. AJMANI, W.-F. NG, M. LIOU. Preconditioned Conjugate Gradient Methods for the Navier-Stokes Equations. *J. Comput. Phys.*, 110: 68–81, 1994.
- [3] E. ANDERSON, Z. BAI, C. BISCHOF, J. DEMMEL, J. DONGARRA, J. DU CROZ, A. GREENBAUM, S. HAMMARLING, A. MCKENNEY, S. OSTROUCHOV, D. SORENSEN. *LAPACK User's Guide*. SIAM, Philadelphia, 1995.
- [4] A. ARNONE, M.-S. LIOU, L. A. POVINELLI. Integration of Navier-Stokes Equations Using Dual Time Stepping and a Multigrid. *AIAA Journal*, 33 (6): 985–990, 1995.
- [5] S. F. ASHBY, T. A. MANTEUFFEL, P. E. SAYLOR. Adaptive Polynomial Preconditioning for Hermitian Indefinite Linear Systems. *BIT*, 29: 583–609, 1989.
- [6] O. AXELSSON. A survey of preconditioned iterative methods for linear systems of algebraic equations. *BIT*, 25: 166–187, 1985.
- [7] O. AXELSSON. *Iterative Solution Methods*. Cambridge University Press, Cambridge, 1996.
- [8] R. BARRETT, M. BERRY, T. CHAN, J. DEMMEL, J. DONATO, J. DONGARRA, V. EIJKHOUT, R. POZO, C. ROMINE, H. VAN DER VORST. *Templates for the Solution of Linear Systems*. SIAM, Philadelphia, 1994.
- [9] M. BENZI. Preconditioning Techniques for Large Linear Systems: A Survey. *J. Comput. Phys.*, 182: 418–477, 2002.
- [10] M. BENZI, M. TUMA. A comparative study of Sparse Approximate Inverse Preconditioners. *Los Alamos National Laboratory Technical Report*, LA-UR-98-24, 1998.
- [11] A. BRANDT. Guide to Multigrid Development. In *Multigrid Methods*, (Eds.) W. HACKBUSCH, U. TROTTEBERG, number 960 in Lecture Notes in Mathematics, pp. 220–312, Berlin, Heidelberg, New York, 1981. Springer.
- [12] C. BREZINSKI, M. REDIVO-ZAGLIA. Look-ahead in BiCGSTAB and other product methods for linear systems. *BIT*, 35: 169–201, 1995.
- [13] C. BREZINSKI, M. REDIVO-ZAGLIA, H. SADOK. A breakdown-free Lanczos type algorithm for solving linear systems. *Numer. Math.*, 63: 29–38, 1992.
- [14] C. BREZINSKI, H. SADOK. Avoiding breakdown in the cgs algorithm. *Numerical Algorithms*, 1: 199–206, 1991.
- [15] W. BUNSE, A. BUNSE-GERSTNER. *Numerische lineare Algebra*. Teubner, Stuttgart, 1985.

- [16] T. F. CHAN, E. GALLOPOULOS, V. SIMONCINI, T. SZETO, C. H. TONG. A quasi-minimal residual variant of the Bi-CGSTAB algorithm for nonsymmetric systems. *SIAM J. Sci. Comput.*, 15 (2): 338–347, 1994.
- [17] A. J. CHORIN. A Numerical Method for Solving Incompressible Viscous Flow Problems. *J. Comput. Phys.*, 2: 12–26, 1967.
- [18] J. DEMMEL. *Applied Numerical Linear Algebra*. SIAM, Philadelphia, 1997.
- [19] P. DEUFLHARD, A. HOHMANN. *Numerische Mathematik I*. de Gruyter, Berlin, New York, 2. überarb. edition , 1993.
- [20] P. F. DUBOIS, A. GREENBAUM, G. H. RODRIGUE. Approximating the inverse of a matrix for use in iterative algorithms on vector processors. *Computing*, 22: 257–268, 1979.
- [21] W. S. EDWARDS, L. S. TUCKERMAN, R. A. FRIESNER, D. C. SORENSEN. Krylov Methods for the Incompressible Navier-Stokes Equations. *J. Comput. Phys.*, 110: 82–102, 1994.
- [22] J. H. FERZIGER, M. PERIĆ. *Computational Methods for Fluid Dynamics*. Springer, Berlin, Heidelberg, 1996.
- [23] K. GRAF F. VON FINCKENSTEIN. *Einführung in die Numerische Mathematik*. Hanser, München, 1. edition , 1977.
- [24] B. FISCHER. *Polynomial Based Iteration Methods for Symmetric Linear Systems*. Advances in Numerical Mathematics. Wiley-Teubner, Stuttgart, 1996.
- [25] R. W. FLETCHER. Conjugate Gradients Methods for Indefinite Systems. In *Dundee Biennial Conference on Numerical Analysis*, (Ed.) G. A. WATSON, pp. 73–89, New York, 1975. Springer.
- [26] R. W. FREUND. Conjugate Gradient-Type Methods for Linear Systems with Complex Symmetric Coefficient Matrices. *SIAM J. Sci. Stat. Comput.*, 13 (1): 425–448, 1992.
- [27] R. W. FREUND. A Transpose-Free Quasi-Minimal Residual Algorithm for Non-Hermitian Linear Systems. *SIAM J. Sci. Comput.*, 14 (2): 470–482, 1993.
- [28] R. W. FREUND, N. M. NACHTIGAL. QMR: A quasi-minimal residual method for non-Hermitian linear systems. *Numer. Math.*, 60: 315–339, 1991.
- [29] A. L. GAITONDE. A dual-time method for the solution of the unsteady Euler equations. *The Aeronautical Journal of the Royal Aeronautical Society*, 98: 283–291, 1994.
- [30] (Eds.) G. GOLUB, A. GREENBAUM, M. LUSKIN. Transpose-Free Quasi-Minimal Residual Methods for Non-Hermitian Linear Systems, volume 60 of *The IMA Volumes in Mathematics and its Applications, Recent Advances in Iterative Methods*, New York, 1994. Springer.
- [31] G. H. GOLUB, C. VAN LOAN. *Matrix Computations*. The John Hopkins University Press, Baltimore, Maryland, 2. edition , 1989.

-
- [32] A. GREENBAUM. *Iterative Methods for Solving Linear Systems*. SIAM, Philadelphia, 1997.
 - [33] M. J. GROTE, T. HUCKLE. Parallel Preconditioning with Sparse Approximate Inverses. *SIAM J. Sci. Comput.*, 18(3): 838–853, 1997.
 - [34] W. HACKBUSCH. *Multi-Grid Methods and Applications*, volume 4 of *Springer Series in Computational Mathematics*. Springer, Berlin, Heidelberg, New York, Tokio, 1985.
 - [35] W. HACKBUSCH. *Integralgleichungen: Theorie und Numerik*. Teubner, Stuttgart, 1989.
 - [36] W. HACKBUSCH. *Iterative Lösung großer schwachbesetzter Gleichungssysteme*. Teubner, Stuttgart, 2., überarb. und erw. edition , 1993.
 - [37] M. R. HESTENES, E. STIEFEL. Methods of Conjugate Gradients for Solving Linear Systems. *NBS J. Res.*, 49: 409–436, 1952.
 - [38] M. HOCHBRUCK, C. LUBICH. Error analysis of Krylov methods in a nutshell. *SIAM J. Sci. Comp.*, 19: 695–701, 1998.
 - [39] E. ISSMAN, G. DEGREZ. Acceleration of compressible flow solvers by Krylov subspace methods. Von Karman Institute for Fluid Dynamics, Lecture Series 1995-14, Rhode-Saint-Genève, March 1995.
 - [40] O. G. JOHNSON, C. A. MICCHELLI, G. PAUL. Polynomial preconditioners for conjugate gradient calculations. *SIAM J. Numer. Anal.*, 20: 362–376, 1983.
 - [41] W. JOUBERT. Lanczos Methods for the Solution of Nonsymmetric Systems of Linear Equations. *SIAM J. Matrix Anal. Appl.*, 13 (3): 926–943, 1992.
 - [42] C.T. KELLEY. *Iterative Methods for Linear and Nonlinear Equations*. SIAM, Philadelphia, PA, 1995.
 - [43] R. KRESS. *Linear Integral Equations*, volume 82 of *Applied Mathematical Sciences*. Springer, New York, Berlin, Heidelberg, 1989.
 - [44] C. LANCZOS. An iterative method for the solution of the eigenvalue problem of linear differential and integral operators. *J. Res. Nat. Bur. Standards*, 45: 255–282, 1950.
 - [45] C. LANCZOS. Solution of systems of linear equations by minimized iterations. *J. Res. Nat. Bur. Standards*, 49: 33–53, 1952.
 - [46] T. A. MANTEUFFEL. An incomplete factorization technique for positive definite linear systems. *Math. Comp.*, 34: 473–497, 1980.
 - [47] J. A. MEIJERNINK, H. A. VAN DER VORST. An Iterative Solution Method for Linear Systems of which the Coefficient Matrix is a Symmetric M-Matrix. *Mathematics of Computation*, 31 (137): 148–162, 1977.
 - [48] A. MEISTER. Comparison of Different Krylov Subspace Methods Embedded in an Implicit Finite Volume Scheme for the Computation of Viscous and Inviscid Flow Fields on Unstructured Grids. *J. Comput. Phys.*, 140: 311–345, 1998.

- [49] A. MEISTER, C. VÖMEL. Efficient Preconditioning of Linear Systems arising from the Discretization of Hyperbolic Conservation Laws. *Advances in Computational Mathematics*, Vol.14 (1): 49–73, 2001.
- [50] A. MEISTER, J. WITZEL. Krylov Subspace Methods in Computational Fluid Dynamics. *Surv. Math. Ind.*, 10 (3): 231–267, 2002.
- [51] G. MEURANT. *Computer Solution of Large Linear Systems*. North-Holland, Amsterdam, 1999.
- [52] N. M. NACHTIGAL, S. C. REDDY, L. N. TREFETHEN. How fast are nonsymmetric matrix iterations? *SIAM J. Matrix Anal. Appl.*, 13 (3): 778–795, 1992.
- [53] G. OPFER. *Numerische Mathematik für Anfänger*. Vieweg, Wiesbaden, 4. edition , 2002.
- [54] C. C. PAIGE, M. A. SAUNDERS. Solution of Sparse Indefinite Systems of Linear Equations. *SIAM J. Num. Anal.*, 12 (4): 617–629, 1975.
- [55] B. N. PARLETT, D. R. TAYLOR, Z. A. LIOU. A Look-Ahead Lanczos Algorithm for Unsymmetric Matrices. *Math Comp.*, 44: 105–124, 1985.
- [56] G. RADICATI, Y. ROBERT, S. SUCCI. Iterative Algorithms for the Solution of Nonsymmetric Systems in the Modelling of Weak Plasma Turbulence. *J. Comput. Phys.*, 80: 489–497, 1989.
- [57] J. K. REID. On the method of conjugate gradients for the solution of large sparse systems of linear equations. In *Large sparse sets of linear equations*, (Ed.) J. K. REID, London, New York, 1971. Academic Press.
- [58] Y. SAAD. Practical use of some Krylov Subspace Methods for Solving Indefinite and Nonsymmetric Linear Systems. *SIAM J. Sci. Stat. Comput.*, 5 (1): 203–228, 1984.
- [59] Y. SAAD. Practical use of polynomial preconditionings for the conjugate gradient method. *SIAM J. Sci. Stat. Comput.*, 6(4): 865–881, 1985.
- [60] Y. SAAD. Krylov subspace techniques, conjugate gradients, preconditioning and sparse matrix solvers. *von Karman Institute of Fluid Dynamics, Lecture Series*, 1994–05, 1994.
- [61] Y. SAAD. *Iterative Methods for Sparse Linear Systems*. PWS Publishing Company, Boston, 1996.
- [62] Y. SAAD, M. H. SCHULTZ. GMRES: A generalized minimal residual algorithm for solving nonsymmetric linear systems. *SIAM J. Sci. Stat. Comput.*, 7: 856–869, 1986.
- [63] H. R. SCHWARZ, N. KÖCKLER. *Numerische Mathematik*. Teubner, Stuttgart, 6. edition , 2006.
- [64] G. L. G. SLEIJPEN, D. R. FOKKEMA. BiCGSTAB(ℓ) for Linear Systems involving Unsymmetric Matrices with Complex Spectrum. *Electronic Transactions on Numerical Analysis*, 1: 11–32, 1993.

-
- [65] A. VAN DER SLUIS. Condition Numbers and Equilibration of Matrices. *Numerische Mathematik*, 14: 14–23, 1969.
 - [66] A. VAN DER SLUIS. Condition, Equilibration and Pivoting in Linear Algebraic Systems. *Numerische Mathematik*, 15: 74–86, 1970.
 - [67] P. SONNEVELD. CGS: A fast Lanczos-Type Solver for Nonsymmetric Linear Systems. *SIAM J. Sci. Stat. Comput.*, 10 (1): 36–52, 1989.
 - [68] J. STOER, R. BULIRSCH. *Numerische Mathematik II*. Springer, Berlin, Heidelberg, New York, 3. edition , 1990.
 - [69] L. N. TREFETHEN, D. BAU. *Numerical Linear Algebra*. SIAM, Philadelphia, 1997.
 - [70] H. A. VAN DER VORST. The convergence behaviour of preconditioned CG and CG-S in the presence of rounding errors. *Lecture Notes in Mathematics*, 1457: 126–136, 1990.
 - [71] H. A. VAN DER VORST. BI-CGSTAB: A fast and smoothly converging variant of BI-CG for the solution of nonsymmetric linear systems. *SIAM J. Sci. Stat. Comput.*, 13: 631–644, 1992.
 - [72] H. A. VAN DER VORST. *Iterative Krylov Methods for Large Linear Systems*, volume 13 of *Cambridge Monographs on Applied and Computational Mathematics*. Cambridge University Press, Cambridge, 2003.
 - [73] H. F. WALKER. Implementation of the GMRES Method using Householder Transformations. *SIAM J. Sci. Comput.*, 9 (1): 152–163, 1988.
 - [74] R. WEISS. *Parameter-Free Iterative Linear Solvers*. Akademie Verlag, Berlin, 1. edition , 1996.
 - [75] R. WEISS, H. HÄFNER, W. SCHÖNAUER. LINSOL (LINear SOLver) Description and User’s Guide for parallelized version. IB Nr. 61/95, Rechenzentrum der Universität Karlsruhe, 1995.
 - [76] J.H. WILKINSON, C. REINSCH. *Linear Algebra. Handbook for automatic computation, Vol 2. Grundlehren der mathematischen Wissenschaften in Einzeldarstellungen*, volume 186. Springer, Berlin, Heidelberg, New York, 1971.

Index

Äquivalenzkonstanten, 10, 12, 29
 Überrelaxationsverfahren, 78
 äquivalente Gleichungssysteme, 17

Arnoldi-Algorithmus, 135, 140
 — modifizierter, 140

Banachscher Fixpunktsatz, 30
 beidseitig präkonditioniert, 209
 Besetzungsstruktur, 209
 Bi-Lanczos-Algorithmus, 144
 bi-orthogonal, 144
 bi-orthonormal, 144
 BiCG-Verfahren, 116, 164, 180
 — Algorithmus, 168
 BiCGSTAB-Verfahren, 116, 174, 189
 — Algorithmus, 178
 — Algorithmus des präkonditionierten, 221
 — präkonditioniertes, 219, 222
 — Stabilisierter, 179

CG-Verfahren, 116, 127, 144, 200, 218
 — Algorithmus, 130
 — Algorithmus des präkonditionierten, 218
 — präkonditioniertes, 216
 CGN, 163
 CGNE, 201
 CGNR, 200
 CGS-Verfahren, 116, 163, 171, 174, 180
 — Algorithmus, 174
 Cholesky-Zerlegung, 46
 — Algorithmus, 48
 — Algorithmus der unvollständigen, 212
 — unvollständige, 211, 216

D-Lanczos-Methode, 140, 143
 Defekt, 105
 Diagonalmatrix, 17
 DIOM, 140
 DQGMRES, 159
 Dreiecksmatrix, 17

Eigenvektor, 20
 Eigenwert, 20
 Einzelschrittverfahren, 75
 Energienorm, 94

Fehlerabschätzung
 — a posteriori, 30
 — a priori, 30, 67
 Fixpunkt, 30, 63
 Folge

— Cauchy-, 8
 — divergente, 8
 — konvergente, 8
 FOM, 138, 143
 — Restarted, 139
 Frobenius-Inverse
 — unvollständige, 223
 Frobeniusnorm, 19

Galerkin-Bedingung, 112
 Gaußsches Eliminationsverfahren, 36
 Gaußsches Eliminationsverfahren
 — Algorithmus, 37, 42
 Gauß-Seidel-Verfahren, 74
 — symmetrisches, 93, 216
 gerichteter Graph, 71
 gerichteter Weg, 71
 Gesamtschrittverfahren, 69, 80
 Givens-Verfahren, 55
 — Algorithmus, 58
 Glätter, 108
 GMRES-Verfahren, 116, 149, 180, 200
 — Algorithmus, 156
 — restarted Algorithmus, 159, 160
 Gradientenverfahren, 119
 Gram-Schmidt-Verfahren, 49
 — Algorithmus, 52
 — mit Nachorthogonalisierung, 54
 — modifiziertes, 53
 Grobitterkorrekturverfahren, 106

Hauptabschnittsdeterminante, 38
 Hauptabschnittsmatrix, 38
 Hessenbergmatrix, 137

Integralgleichung, 3
 IOM, 140
 Iterationsmatrix, 63
 Iterationsverfahren, 62
 — Fixpunkt eines, 63
 — konsistentes, 63
 — konvergentes, 63
 — lineares, 62

Jacobi-Verfahren, 69, 103, 132

Konditionszahl, 25, 28
 — äquivalente, 29
 Kontraktionszahl, 30
 Konvektions-Diffusions-Gleichung, 5, 73, 221
 Krylov-Unterraum, 113
 — transponierter, 144

- Krylov-Unterraum-Methode, 113
- Lanczos-Algorithmus, 140
 - Look-ahead-, 149, 174, 180
- LAPACK, viii
- Laplace-Operator, 1
- LINSOL, viii
- LU-Zerlegung, 140
 - Algorithmus der unvollständigen, 210
 - mit fill-in, 211
 - unvollständige, 159, 209, 221
- MATLAB, viii
- Matrix
 - ähnliche, 17
 - adjungierte, 17
 - diagonaldominante, 70
 - Frobenius-, 37
 - Hauptabschnitts-, 38
 - hermitesche, 17
 - Hessenberg-, 137
 - irreduzible, 70, 73
 - Konditionszahl einer, 25
 - konsistent geordnete, 82
 - muster, 208
 - negativ definite, 17
 - negativ semidefinite, 17
 - normale, 17
 - orthogonale, 17
 - Permutations-, 36
 - positiv definite, 17
 - positiv semidefinite, 17
 - reduzible, 70
 - strikt diagonaldominante, 70, 76
 - symmetrische, 17
 - transponierte, 16
 - Tridiagonal-, 83
 - unitäre, 17
 - unzerlegbare, 70
 - zerlegbare, 70
- Matrixnorm
 - induzierte, 14
- Mehrgitterverfahren, 97
 - Algorithmus, 110
 - vollständiges, 111
- Methode der kleinsten Quadrate, 4
- Methode der konjugierten Richtungen, 124, 127
 - Algorithmus, 126
- Methode des steilsten Abstiegs, 118, 127, 200
 - Algorithmus, 119
- MINRES, 163
- Muster
 - Matrix-, 208
 - reguläres, 214
 - Spalten-, 208
 - Zeilen-, 209
- Neumannsche Reihe, 204
 - m-ter Stufe, 205
- Norm, 7
 - äquivalente, 9
 - Betragssummen-, 7
 - Energie-, 94
 - euklidische, 8
 - Frobenius-, 19
 - Maximum-, 8
 - Operator-, 14
 - Spaltensummen-, 19
 - Spektral-, 22
 - Vektor-, 7
 - Zeilensummen-, 19
- Operator
 - beschränkter, 14
 - Fixpunkt eines, 30
 - kontrahierender, 30
 - linearer, 14
 - stetiger, 14
- Operatornorm, 14
- PCG-Verfahren
 - Algorithmus, 218
- Petrov-Galerkin-Bedingung, 112
- Pivotisierung, 45
 - Spalten-, 45
 - vollständige, 45
 - Zeilen-, 45
- Poisson-Gleichung, 1, 73, 97, 132, 218
- Präkonditionierer, 200
 - Gauß-Seidel-, 208, 222
 - ILU-, 210, 221
 - IQR-, 213, 223
 - Jacobi-, 208, 222
 - Links-, 200, 223
 - Polynomiale, 204, 222
 - Rechts-, 200, 223
 - SOR-, 208
 - Splitting-assozierte, 207
 - SSOR-, 208
 - symmetrischer Gauß-Seidel-, 208, 222
- Präkonditionierungen, 27
- Produktiteration, 107
- Projektionsmethode, 112
 - orthogonale, 112
 - schiefe, 112
- Prolongation, 102
- QGMRES, 159
- QMR-Verfahren, 116
- QMRCGSTAB-Verfahren, 116, 189

- Algorithmus, 192
- QR-Zerlegung, 49
- Algorithmus der unvollständigen, 213
- unvollständige, 212, 223
- Raum
 - Banach-, 12
 - normierter, 7
 - Prä-Hilbert-, 8
- Relaxationsparameter, 78, 79
- Relaxationsverfahren, 77
 - symmetrisches Gauß-Seidel-, 93
 - Gauß-Seidel-, 80, 84
 - Jacobi-, 78
- Residuum, 118
- Restriktion, 102
- Richardson-Verfahren, 89
- singuläre Werte, 27, 28
- Skalarprodukt, 8
- Skalierung, 201
 - Spalten-, 202
 - Zeilen-, 202
- SOR-Verfahren, 87
- Spaltensummennorm, 19
- Spektralnorm, 22
- Spektralradius, 20, 67, 68
- Spektrum, 20
- Splitting-Methode, 65, 113, 132
 - symmetrische, 66, 92, 216
- SSOR-Verfahren, 93
- Submultiplikativität, 16
- Templates, viii
- TFQMR-Verfahren, 116, 180
 - Algorithmus, 189
- Unterrelaxationsverfahren, 78
- unvollständige Frobenius-Inverse, 214
- V-Zyklus, 111
- Vektor
 - Fehler-, 26
 - norm, 7
 - Residuen-, 26, 118
 - Singulär-, 27
- Verfahren der konjugierten Gradienten, 116, 127, 144, 200, 218
 - Algorithmus, 130
 - Algorithmus des präkonditionierten, 218
 - präkonditioniertes, 216
- Verfahren der konjugierten Richtungen, 124, 127
 - Algorithmus, 126
- Verfahren des steilsten Abstiegs, 118, 127, 200
 - Algorithmus, 119
- W-Zyklus, 111
- Zeilensummennorm, 19
- Zweigitterverfahren, 105
 - Algorithmus, 109