
Fachhochschule Aachen

Standort Jülich

Fachbereich: Angewandte Naturwissenschaft und Technik

Studiengang: Technomathematik

**Entwicklung schneller Löser für
Randwertprobleme in der geoelektrischen
Widerstandstomographie**

Diplomarbeit von Stefan Jongen

Jülich, Februar 2006

Die vorliegende Diplomarbeit wurde in Zusammenarbeit mit der Forschungszentrum Jülich GmbH, Institut für Festkörperforschung (IFF), angefertigt.

Diese Diplomarbeit wurde betreut von:

Referent: Prof. Dr. M. Reißel
Koreferent: Privatdozent Dr. H. Lustfeld

Diese Arbeit wurde von mir selbständig angefertigt und verfasst. Es sind keine anderen als die angegebenen Quellen und Hilfsmittel benutzt worden.

Jülich, den 6. Februar 2006

Inhaltsverzeichnis

1	Einleitung	1
1.1	Einführung	1
1.2	Problembeschreibung	2
1.3	Ziel der Arbeit	5
2	Grundlagen	7
2.1	Das Randwertproblem	7
2.2	Das Cole-Cole Modell	8
2.2.1	Ableitungen der Cole-Cole Funktion	8
2.3	Fréchet-Ableitungen	9
2.4	Implizite Funktionen	10
2.4.1	Ableitung impliziter Funktionen	10
2.5	Der adjungierte Operator	11
2.6	Der Satz von Gauß	11
2.7	Das ADI-Verfahren	13
3	Diskretisierung des Randwertproblems	15
3.1	Einleitung	15
3.2	Einteilung des Gebietes in ein äquidistantes Gitter	17
3.3	Diskretisierung des Gradienten	18
3.4	Diskretisierung der Leitfähigkeit	20
3.5	Diskretisierung der Divergenz	20
3.6	Diskretisierung der rechten Seite der Differentialgleichung	22
3.7	Die diskreten Randbedingungen	24
3.7.1	Neumann Bedingung	24
3.7.2	Dirichlet-Bedingung	27
3.8	Zusammenfassung aller Diskretisierungen	28
3.8.1	Zusammenhang Beispiel-Matrix und Randwertproblem	29
3.9	Approximierte Lösungen einiger Vorwärtsprobleme	29
3.9.1	Lösung eines Vorwärtsproblems bei grober Auflösung	30

3.9.2	Lösung eines Vorwärtsproblems bei feiner Auflösung	31
3.9.3	Positionierung einer δ -Quelle am Neumann-Rand	31
3.9.4	Mehrere δ -Quellen in Ω	32
3.10	Verfeinerung	33
3.10.1	Beliebige Positionierung der δ -Quellen	34
3.10.2	Position der Messpunkte	38
3.10.3	Verschiedene Frequenzen	39
3.10.4	Messprogramme	40
3.10.5	Grobes c-Gitter	42
4	Ableitung und adjungierte Ableitung des Randwertproblems	45
4.1	Bestimmung der gesuchten Ableitung $\partial_c u(c)(c_0)(c_r)$	45
4.1.1	Bestimmung der Ableitung im kontinuierlichen Fall	46
4.1.2	Bestimmung der Ableitung im diskreten Fall	48
4.2	Bestimmung des adjungierten Operators	50
4.2.1	Bedeutung des adjungierten Operators in der geoelektrischen Tomographie	50
4.2.2	Bestimmung des adjungierten Operators im kontinuierlichen Fall	51
4.2.3	Bestimmung des adjungierten Operators im diskreten Fall	53
4.3	Kostenvergleich: Ableitungsmatrix gegen Matrix-Vektor Produkt	54
5	Spezielle Löser der linearen Gleichungssysteme	57
5.1	Einführung	57
5.2	Optimierung	58
5.2.1	Verkleinerung des Gleichungssystems	59
5.2.2	Erzeugung einer symmetrischen Matrix	62
5.3	Die klassische LU-Zerlegung	64
5.4	Die Cholesky-Zerlegung	65
5.5	Lösen mit Hilfe des ADI-Verfahrens	66
5.5.1	Erzeugung der Richtungs-Matrizen	67
5.5.2	Steuerung der Schrittweite	68
5.6	Die UMFPACK Bibliothek	70
6	Numerische Ergebnisse	71
6.1	Steuerung des ADI-Verfahrens	71
6.1.1	Konvergenzgeschwindigkeiten in Abhängigkeit von t	71
6.1.2	Variation der Leitfähigkeit	75
6.1.3	Optimale Steuerung	76
6.1.4	Unterschiedliche Steuerung bei gleicher Dimension	78
6.1.5	Aussagen über die Abschätzung des Relaxationsparameters	79
6.2	Vergleich der Löser	79
6.2.1	Vergleich des ursprünglichen Systems mit dem optimierten System	79

6.2.2	Vergleich aller optimierten Löser	82
7	Zusammenfassung und Ausblick	85

Kapitel 1

Einleitung

1.1 Einführung

In einem Projekt des Forschungszentrums Jülich versucht man, die Zusammensetzung eines Bodens mit möglichst geringem Aufwand zu ermitteln. Um dies zu erreichen, wird das Verfahren der Electric Resistance Tomography (ERT) genutzt. Dieses Verfahren gehört zu der Klasse der tomographischen Verfahren, bei denen die Verteilung eines Parameters in einem Objekt ermittelt werden soll. Das wohl bekannteste Verfahren ist die Computer-Tomographie, die oft in der Medizin eingesetzt wird.

Bei der ERT werden Elektroden in Bohrlöchern und auf der Oberfläche eines Bodens verteilt, die sowohl die Möglichkeit besitzen, Ströme abzugeben / aufzunehmen, als auch die Eigenschaft, Potentiale zu messen. Nun kann man durch eine der Elektroden Gleich- oder Wechselstrom injizieren und diesen an einer anderen wieder abfließen lassen. Es stellt sich ein Potential im Boden ein, das an den übrigen Elektroden gemessen werden kann. Diese Messungen können anhand verschiedener Ein- und Auslassstellen des Stroms und bei verschiedenen Frequenzen durchgeführt werden. Mit Hilfe dieser gewonnenen Daten versucht man, die Bodenstruktur möglichst exakt zu identifizieren.

Diese Aufgabe gehört zu der Klasse der inversen Probleme. Während man bei den direkten Problemen (Vorwärtsproblemen) aus bekannten Ursachen die unbekannten Wirkungen ermittelt, versucht man bei den inversen Problemen (Rückwärtsproblemen) anhand der beobachteten Wirkungen auf die Ursachen zu schließen. Auf das hier geschilderte Problem der Bodenidentifizierung angewandt bedeutet dies, dass versucht wird, aus einer gemessenen Potentialverteilung und einer bekannten Stromeinspeisung eine (eindeutige) Bodenstruktur zu ermitteln.

In der Regel lassen sich mehrere Zusammensetzungen bestimmen, die der Lösung des gemessenen Spannungsbildes genügen. Mit Hilfe vieler gemessenen Daten und sogenannten Regularisierungsverfahren (siehe [5]) wird versucht, die Menge der möglichen Lösungen einzugrenzen.

Zur Identifizierung des Bodens müssen daher sehr viele große Gleichungssysteme gelöst werden, die häufig mehr als tausend Unbekannte beinhalten. Da zu deren Bestimmung viel Arbeitsspeicher und Zeit benötigt werden, versucht man, diese Gleichungssysteme möglichst effizient zu lösen.

1.2 Problembeschreibung

Jeder Bodentyp besitzt eine frequenzabhängige Leitfähigkeit, die durch die vier sogenannten Cole-Cole Parameter (siehe Kapitel 2.2) dargestellt werden kann. Somit sieht die Potentialverteilung je nach Bodenzusammensetzung und Frequenz des injizierten Stroms unterschiedlich aus.

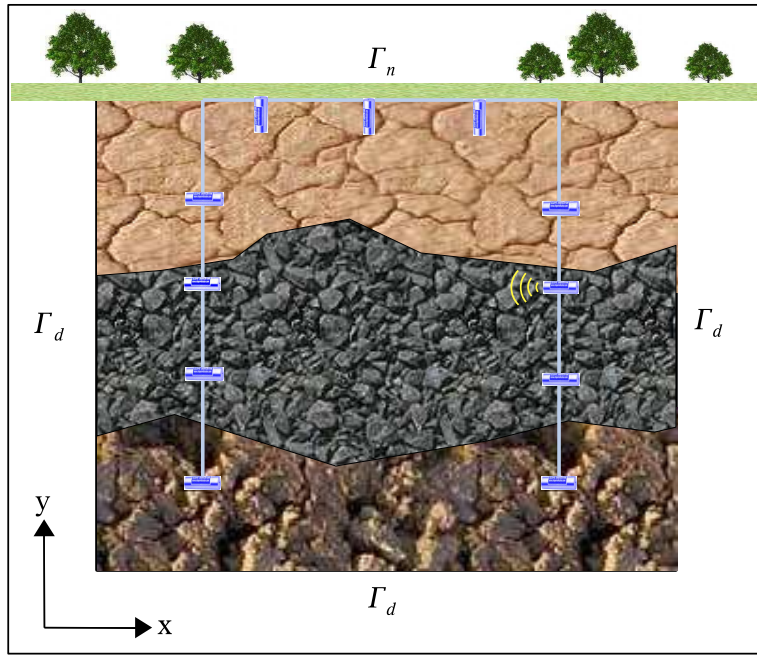


Abbildung 1.1: Gebiet Ω , Ränder Γ_d und Γ_n , Elektroden

In einem Versuchsaufbau wird bei einer gegebenen Frequenz ω und einer bekannten Strom-einspeisung I das sich einstellende Potential

$$\hat{u}(I, \omega) \quad (1.1)$$

an p Sonden gemessen

$$\hat{u}_k(I, \omega) \quad k = 1, \dots, p. \quad (1.2)$$

Nun kann man bei dem gleichen Versuchsaufbau (I, ω) das Vorwärtsproblem bei einer beliebigen Verteilung der Cole-Cole Parameter c numerisch berechnen. Hierbei ist zu beachten, dass die Cole-Cole Parameter c auf dem Gebiet nicht konstant sein müssen. Je nach Bodenschicht können sich die vier Parameter $(\rho_0, m, \tau, \gamma)$ unterscheiden. Es sollen daher nicht nur vier Parameter identifiziert

werden, sondern eine Verteilung der Parameter über Ω . Auch bei der numerischen Näherung

$$u(I, \omega, c) \quad \text{mit } c = (\rho_0, m, \tau, \gamma)^T. \quad (1.3)$$

werden nur die Potentiale $u_k(I, \omega, c)$ betrachtet, die sich an den Messsonden befinden. Für eine möglichst genaue Approximation der Bodenstruktur versucht man, die Differenzen

$$m_k(c) = \hat{u}_k - u_k(c) \quad k = 1, \dots, p \quad (1.4)$$

möglichst klein (betragsmäßig) zu halten. Je kleiner die Beträge der Differenzen sind, desto näher liegt die numerische Lösung $u_k(c)$ an der gemessenen Lösung $\hat{u}_k$ (bei festem I und ω).

Um mehr Informationen zu gewinnen, werden verschiedene Stromeinspeisungen und Frequenzen verwendet. Seien nun insgesamt n verschiedene Messungen mit den Eingangsdaten

$$(I_j, \omega_j) \quad j = 1, \dots, n \quad (1.5)$$

durchgeführt worden. Hierbei ist es durchaus möglich, dass bei einer Stromeinspeisung I mehrere Frequenzen ω benutzt werden. Somit bezeichnet

$$m_{k,j}(c) = \hat{u}_k(I_j, \omega_j) - u_k(I_j, \omega_j, c), \quad k = 1, \dots, p, \quad j = 1, \dots, n \quad (1.6)$$

einen Wert, der sich aus der Differenz der Messung im j -ten Versuch an dem k -ten Sensor mit dem entsprechenden numerisch angenäherten Wert ergibt.

Ziel ist es nun, eine Verteilung der Cole-Cole Parameter in Ω zu bestimmen, so dass alle $m_{k,j}(c)$ möglichst klein sind. Diese Verteilung wird von nun an c_m genannt. Um diese zu ermitteln, wird ein weiterer Index i definiert, der die Indizes k und j ersetzen soll. Das bedeutet, dass jedes i eindeutig für einen Messpunkt k in einem Versuch j steht, wie dies zum Beispiel der Ausdruck

$$i = (j - 1) \cdot p + k$$

verdeutlicht. Somit ergibt sich für eine Differenz in einem Messpunkt bei einem beliebigen Versuchsaufbau

$$m_i(c) = \hat{u}_i - u_i(c) \quad i = 1, \dots, s \quad \text{mit } s = n \cdot p. \quad (1.7)$$

Alle Differenzen werden im Vektor

$$m(c) := \begin{pmatrix} m_1(c) \\ \vdots \\ m_s(c) \end{pmatrix} \quad (1.8)$$

zusammengefasst.

Bemerkung:

Bei diesem Problem ist die Funktion $m_i(c)$ komplex. Lediglich bei Gleichstrom ($\omega = 0$) sind die Ergebnisse reell.

Mit Hilfe der Formulierung (1.7) kann man das inverse Problem in ein Least Squares Problem überführen.

Definition 1 Least Squares Problem

Finde c_m , einen lokalen Minimierer für

$$M(c) = \frac{1}{s} \cdot \frac{1}{2} \sum_{i=1}^s |m_i(c)|^2,$$

wobei die Funktionen $m_i : \mathbb{R}^z \rightarrow \mathbb{C}, i = 1, \dots, s$ gegeben sind und es gilt: $s \geq z$.

Um das Minimum der Funktion $M(c)$ zu finden, wird diese mit Hilfe der Taylorentwicklung angenähert. Es gilt

$$M(c+h) = M(c) + h^T \cdot g(c) + O(\|h\|^2), \quad (1.9)$$

wobei g der Gradient ist,

$$g(c) := M'(c) = \begin{pmatrix} \frac{\partial M}{\partial c_1}(c) \\ \vdots \\ \frac{\partial M}{\partial c_s}(c) \end{pmatrix}. \quad (1.10)$$

Ist nun c_m ein lokales Minimum und $\|h\|$ klein genug, dann existiert kein Punkt $c_m + h$ mit einem kleineren Funktionswert, d.h. $M(c_m + h) > M(c_m)$.

Theorem 1 Notwendige Bedingung für ein lokales Minimum

Wenn c_m ein lokaler Minimierer von $M(c)$ ist, dann gilt

$$g_m := M'(c_m) = 0.$$

Dieses notwendige Kriterium kann ebenso über die Funktion $m(c)$ dargestellt werden:

$$M(c) = \frac{1}{2} \sum_{i=1}^s |m_i(c)|^2 = \frac{1}{2} \|m(c)\|^2 = \frac{1}{2} m(c)^H m(c), \quad (1.11)$$

wobei $m(c)^H$ der hermitesche Vektor zu $m(c)$ ist.

Auch die Funktion $m(c)$ kann über die Taylor-Entwicklung angenähert werden:

$$m(c+h) = m(c) + J(c) \cdot h + O(\|h\|^2), \quad (1.12)$$

wobei J die Jacobi-Matrix repräsentiert. Diese Matrix beinhaltet die ersten partiellen Ableitungen der Funktionsparameter

$$(J(c))_{ij} = \frac{\partial m_i}{\partial c_j}(c). \quad (1.13)$$

Aus Gleichung (1.11) folgt:

$$\frac{\partial M}{\partial c_j}(c) = \sum_{i=1}^s m_i(c) \frac{\partial m_i}{\partial c_j}(c). \quad (1.14)$$

Daraus ergibt sich für den Gradienten der Funktion M :

$$g(c) = M'(c) = J(c)^H m(c) \quad (1.15)$$

Die in [5] beschriebenen Regularisierungsverfahren bauen auf die Linearisierung der Funktion $m(c)$ (Gleichung (1.12)) auf. Da zur Berechnung von $m(c)$, von $J(c)$ und somit von $g(c)$ mehrere Randwertprobleme gelöst werden müssen, wird das kontinuierliche Problem zunächst in eine diskrete Form überführt. Erst dann kann das Vorwärtsproblem numerisch gelöst werden. Mit diesen gewonnenen Daten und des *Satzes über implizite Funktionen* (vgl. Kapitel 2.4) kann man die Jacobi-Matrix $J(c_0)$ (vgl. 1.13) aufstellen.

Da einige Regularisierungsverfahren auch die Adjungierte der Jacobi-Matrix $J(c_0)^H$ benötigen, muss auch diese numerisch approximiert werden.

Alle diese numerischen Berechnungen führen zu der Aufgabe, große lineare Gleichungssysteme aufzustellen und zu lösen. Da die Matrizen neben ihrer dünnen Besetzung weitere spezielle Eigenschaften besitzen, sollen Lösungsverfahren angewandt werden, die diese Charakteristika ausnutzen und somit zu einer schnelleren Lösung führen.

1.3 Ziel der Arbeit

In dieser Arbeit wird nicht das gesamte inverse Problem gelöst; es werden die Schnittstellen zur Verfügung gestellt, die die in [5] beschriebenen Regularisierungsverfahren zur Lösung des Problems benötigen. Diese sind:

- Lösung des Vorwärtsproblems
- Berechnung des Produktes $J(c_0) \cdot c_r$
- Berechnung des Produktes $J^H(c_0) \cdot x$

Mit Hilfe der zentralen finiten Differenzen wird das elliptische Randwertproblem (siehe Kapitel 2.1) numerisch approximiert. Somit kann das Vorwärtsproblem mittels Lösen eines linearen Gleichungssystems berechnet werden.

Zum Aufstellen einer Spalte der Ableitungsmatrix $J(c_0)$ muss jeweils ein Randwertproblem gelöst werden. Um nun die gesamte Matrix aufzustellen, ist sehr viel Rechenzeit und ein hoher Aufwand (siehe Kapitel 4.3) nötig. In allen in [5] vorkommenden Regularisierungsverfahren wird jedoch die Ableitung in einem Punkt c_0 in Richtung c_r gesucht. Somit wäre es sinnvoll, nicht die gesamte Matrix $J(c_0)$ zu ermitteln, sondern direkt das Matrix-Vektor Produkt $J(c_0) \cdot c_r$. Dies erspart viele Berechnungen und beschleunigt die Verfahren. Da nun $J(c_0)$ nicht vollständig ermittelt werden muss, sollten auch die Matrix-Vektor Produkte $J(c_0)^H \cdot x$ direkt ermittelt werden.

Zur Bestimmung eines Produktes müssen sehr große Gleichungssysteme gelöst werden, die bestimmte Charakteristika aufweisen. Nutzt man das Wissen über diese Eigenschaften aus, so kann man die Gleichungssystem-Löser optimieren und eine Beschleunigung erzielen.

In dieser Arbeit wird neben den klassischen Lösern (LU- und Cholesky-Zerlegung) auch das iterative ADI-Verfahren (siehe Kapitel 2.7) untersucht, welches nicht ein Gleichungssystem über zwei Richtungen löst, sondern dieses Große auf zwei Kleinere überführt, die nur in eine Richtung gelöst werden müssen.

Diese Berechnungen sollen mit Hilfe eines C++ Programms realisiert werden, das speziell an dieses elliptische Randwertproblem angepasst ist. Es soll das Vorwärtsproblem sowie die beiden Matrix-Vektor Produkte $J(c_0) \cdot c$ und $J^H(c_0) \cdot x$ ermitteln können. Hierbei wird ein besonderes Augenmerk auf ein effizientes und schnelles Lösen der linearen Gleichungssysteme gelegt. Zum Vergleich wird eine C-Bibliothek herangezogen, die speziell für das optimierte Lösen großer Gleichungssysteme mit dünnbesetzten Matrizen entwickelt wurde.

Alle folgenden Herleitungen beschränken sich auf ein zwei-dimensionales Modell, das eine vorgegebene Breite (x-Richtung) und eine Tiefe (y-Richtung) (vgl. Abbildung 1.1) besitzt. Ziel ist es nicht, das Problem unter realen Bedingungen zu analysieren, sondern eine Optimierung der Rechenoperationen zu erreichen. Zur Bestimmung des inversen Problems müssen mehrere tausend Differentialgleichungen und somit viele Gleichungssysteme gelöst werden. Da diese Gleichungssysteme in diesem verkleinerten Modell die gleichen Charakteristika wie im drei-dimensionalen Fall aufweisen, wird die Optimierung der Löser auf das zwei-dimensionale Modell reduziert.

Kapitel 2

Grundlagen

2.1 Das Randwertproblem

Aufgrund der physikalischen Eigenschaften eines Bodens kann man die Potentialbildung u durch eine Differentialgleichung formulieren, deren Herleitung in [6] ausführlich beschrieben wird.

$$-\operatorname{div}(\sigma \cdot \operatorname{grad}(u)) = I \quad \text{in } \Omega \quad (2.1)$$

$$u = 0 \quad \text{auf } \Gamma_d \quad (2.2)$$

$$\partial_\nu u = 0 \quad \text{auf } \Gamma_n \quad (2.3)$$

Der Strom I wird eingeleitet und führt je nach Leitfähigkeit σ des Bodens zu dem bestimmten Potential u .

Die Neumann-Bedingung an der Oberfläche Γ_n (vgl. Abbildung 1.1) wurde aus physikalischen Gründen auf 0 gesetzt, da „an der Oberfläche bei ausreichender Luftfeuchtigkeit keine elektrischen Ladungen entstehen“¹. Weil Luft ein sehr guter Isolator ist, endet das elektrische Feld E an der Oberfläche ($\nu \cdot E = 0$).

Auf den Rändern Γ_d soll das Potential 0 sein. Da sich die Messsonden in der Regel nicht direkt an den Dirichlet-Rändern befinden, sind die numerischen "Messungen" unabhängig von dieser Randbedingung. Mit zunehmender Entfernung zur Stromquelle nimmt das Potential immer weiter ab, so dass es irgendwann (am Dirichlet-Rand) den Wert 0 besitzt.

Bemerkung:

Hierbei ist zu beachten, dass die Leitfähigkeit σ eines Bodens von den *Cole-Cole* Parametern c (siehe Kapitel 2.2) des Bodens und der Frequenz ω abhängt.

$$\sigma : (c, \omega) \rightarrow \sigma(c, \omega) \quad (2.4)$$

¹siehe [6]

Je nach Frequenz und Beschaffenheit des Bodens (Cole-Cole Parameter) ergibt sich eine andere Leitfähigkeit und somit ein anderes Potential.

2.2 Das Cole-Cole Modell

Zur Bestimmung der Zusammensetzung eines Bodens wird dessen Frequenzgang der Leitfähigkeit genauer betrachtet. *Cole and Cole* [7] entwickelten ein Modell der Gesteinsleitfähigkeit, das neben der Frequenz von vier Cole-Cole Parametern abhängig ist. Mit Hilfe dieser Parameter kann man einen Boden eindeutig identifizieren. Eine Herleitung dieser Formel wird in [4] genauer formuliert.

$$\sigma(\omega) = \frac{1}{\rho(c, \omega)} \quad \text{mit } c = (\rho_0, m, \tau, \gamma)^T \quad (2.5)$$

$$\rho(c, \omega) = \rho_0 \cdot \left[1 - m \cdot \left(1 - \frac{1}{1 + (i\omega\tau)^\gamma} \right) \right] \quad (2.6)$$

Erläuterung:

Die Leitfähigkeit $\sigma(c, \omega)$ ist der Kehrwert des Widerstandes $\rho(c, \omega)$. Dieser ist frequenzabhängig und lässt sich durch die vier Cole-Cole Parameter charakterisieren:

- ρ_0 : Widerstand im Gleichstromfall ($\omega = 0$)
- m : Aufladefähigkeit des Bodens
- γ : Frequenzabhängigkeit
- τ : Zeitkonstante

In der Regel ist die Leitfähigkeit komplex, da eine Frequenz $\omega \neq 0$ immer einen imaginären Anteil in den Widerstand bringt ($\tau, \gamma \neq 0$).

In Tabelle 2.1 sind typische Parameter für Böden aufgelistet.

	ρ_0	m	τ	γ
polarisierbarer Boden	100	0.25	10^{-2}	0.1
steiniger Boden	1000	0.1	$10^{-1.5}$	0.4
Bruchzone	500	0.4	10	0.3

Tabelle 2.1: Cole-Cole Parameter verschiedener Böden

2.2.1 Ableitungen der Cole-Cole Funktion

Um das Rückwärtsproblem lösen zu können, werden die Fréchet-Ableitungen (siehe Kapitel 2.3) des Potentials nach den Cole-Cole Parametern benötigt. Um diese berechnen zu können, müssen

unter anderem die partiellen Ableitungen der Cole-Cole Formel ermittelt werden.

$$\frac{\partial \sigma}{\partial \rho_0}(\omega) = \frac{-1}{\rho_0^2 \cdot \left[1 - m \cdot \left(1 - \frac{1}{1+(i\omega\tau)^\gamma}\right)\right]} \quad (2.7)$$

$$\frac{\partial \sigma}{\partial m}(\omega) = \frac{1 - \frac{1}{1+(i\omega\tau)^\gamma}}{\rho_0 \cdot \left[1 - m \cdot \left(1 - \frac{1}{1+(i\omega\tau)^\gamma}\right)\right]^2} \quad (2.8)$$

$$\frac{\partial \sigma}{\partial \tau}(\omega) = \frac{1}{\rho_0 \cdot \left[1 - m \cdot \left(1 - \frac{1}{1+(i\omega\tau)^\gamma}\right)\right]^2} \cdot \frac{m\gamma \cdot (i\omega\tau)^\gamma}{\tau(1 + (i\omega\tau)^\gamma)^2} \quad (2.9)$$

$$\frac{\partial \sigma}{\partial \gamma}(\omega) = \frac{1}{\rho_0 \cdot \left[1 - m \cdot \left(1 - \frac{1}{1+(i\omega\tau)^\gamma}\right)\right]^2} \cdot \frac{m \cdot (i\omega\tau)^\gamma \cdot \ln(i\omega\tau)}{(1 + (i\omega\tau)^\gamma)^2} \quad (2.10)$$

2.3 Fréchet-Ableitungen

Im Kapitel 1.2 wird die Linearisierung einer Funktion $M(c)$ (vgl. Ausdruck (1.9)) bzw. $m(c)$ (vgl. Ausdruck (1.12)) beschrieben, um deren Minima zu ermitteln. Als wichtiges Hilfsmittel zur Berechnung einer Ableitung ist die Definition der Fréchet-Differenzierbarkeit, die einen Differenzierbarkeitsbegriff für Operatoren auf Banachräumen beschreibt.

Definition 2 Die Fréchet-Differenzierbarkeit

Seien X und Y Banachräume. Eine Abbildung $f : X \rightarrow Y$ heißt Fréchet-differenzierbar an der Stelle $x_0 \in X$, falls es eine beschränkte, lineare Abbildung $L_{x_0} : X \rightarrow Y$ gibt, so dass

$$f(x_0 + h) = f(x_0) + L_{x_0}h + O(\|h\|_X) \quad \text{für } h \rightarrow 0.$$

Die Abbildung L_{x_0} heißt Fréchet-Ableitung oder Linearisierung von f in x_0 . Man schreibt

$$L_{x_0} := f'(x_0).$$

Beispiel:

Es sei A eine lineare, stetige Abbildung zwischen den normierten Räumen X und Y . Dann gilt für ein beliebiges $x_0 \in X$:

$$A(x_0 + h) = Ax_0 + Ah$$

und damit $A'(x_0) = A$. Die Fréchet-Ableitung eines linearen stetigen Operators ist also dieser Operator selbst.

Ist insbesondere $X = \mathbb{R}^n$ und $Y = \mathbb{R}^m$, so lässt sich die Abbildung $A \in \mathbb{L}(\mathbb{R}^n, \mathbb{R}^m)$ durch eine

Matrix

$$a = \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \cdots & a_{mn} \end{pmatrix}$$

vom Typ (m, n) darstellen. Analog gilt dies für die Ableitung: $A'(x_0) = A = a$.

2.4 Implizite Funktionen

Eine implizite Funktion ist eine Funktion, die nicht in der einfachen Zuordnungsvorschrift $y = f(x)$ gegeben ist, sondern implizit definiert ist: $F(x, y) = 0$. Diese Gleichung definiert eine Funktion implizit dadurch, dass jedem Wert x derjenige Wert $y(x)$ zugeordnet wird, der die Gleichung $F(x, y(x)) = 0$ erfüllt.

Auch die in Kapitel 2.1 beschriebene Differentialgleichung kann durch eine implizite Funktion dargestellt werden.

$$g(\sigma, u(\sigma)) := \begin{pmatrix} -\operatorname{div}(\sigma \cdot \operatorname{grad}(u)) - I \\ u|_{\Gamma_d} \\ \partial_\nu u|_{\Gamma_n} \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix} \quad (2.11)$$

2.4.1 Ableitung impliziter Funktionen

Bisher ist es lediglich möglich, aus einer vorgegebenen Cole-Cole Parameter Verteilung c und einer Frequenz ω die Leitfähigkeit und somit das sich einstellende Potential zu berechnen. Für das inverse Problem wird jedoch die Ableitung des Potentials nach den Cole-Cole Parametern benötigt. Da die Funktion $u(\sigma(c, \omega))$ nur implizit gegeben ist, wird der *Satz über implizite Funktionen* zu Hilfe genommen.

Satz 1 Implizite Funktionen

Seien X, Y und Z Banachräume über $\mathbb{C}$, $(x_0, y_0) \in X \times Y$ und $V(x_0, y_0)$ eine offene Umgebung mit

$$g : V(x_0, y_0) \subset X \times Y \rightarrow Z \quad \text{Fréchet-differenzierbar.}$$

Ist $\partial_y g(x_0, y_0) \in L(X \times Y, Z)$ bijektiv, dann gibt es eine Umgebung V , für die gilt:

- 1) $\forall x \in V$ hat $g(x, y) = 0$ genau eine Lösung $y(x)$
- 2) $y(x)$ ist Fréchet-differenzierbar auf V
- 3) $\partial_x y(x_0) \in L(X, Y)$ ist gegeben durch

$$\partial_x y(x_0) = -\partial_y g(x_0, y(x_0))^{-1} \circ \partial_x g(x_0, y(x_0))$$

2.5 Der adjungierte Operator

Sei A derjenige Operator, der zu dem gegebenen Randwertproblem bei einer festen Frequenz ω die Cole-Cole Parameter c_0 auf die Ableitung des Potentials u nach den Cole-Cole Parametern in c_0 abbildet.

$$A : c_0 \rightarrow \frac{\partial u}{\partial c}(c_0) \quad (2.12)$$

Für einige in [5] beschriebenen Regularisierungsverfahren wird jedoch auch noch der adjungierte Operator A^* zu A benötigt, der sich wie folgt definieren lässt.

Definition 3 Der adjungierte Operator

Für einen beschränkten linearen Operator L auf einem separablen Hilbert-Raum H wird durch

$$\langle Lv, w \rangle = \langle v, L^*w \rangle, \quad v, w \in H,$$

der sogenannte adjungierte Operator L^* definiert.

Gilt $L^* = L$, so bezeichnet man L als selbstadjungiert.

Im diskreten Fall des hier untersuchten Problems ist der Operator A eine (komplexe) Matrix. Der adjungierte Operator A^* ist die hermitesche Matrix zu A :

$$A^* = A^H \quad (2.13)$$

2.6 Der Satz von Gauß

Mit Hilfe des *Satzes von Gauß* lässt sich das Integral über die Divergenz eines Vektorfeldes leichter berechnen. Er stellt eine Beziehung zwischen einem n -dimensionalen und einem $n-1$ -dimensionalen Integral her. In dieser Arbeit wird ein zwei-dimensionales Integral auf ein ein-dimensionales zurückgeführt.

Satz 2 Integralsatz von Gauß

Gegeben seien eine vektorielle Ortsfunktion $\vec{p}(\vec{r})$ auf einem geschlossenen zwei-dimensionalen Gebiet Ω , das durch dessen stückweise stetig differenzierbaren Rand $\partial\Omega$ begrenzt ist, sowie der äußere Normalenvektor ν auf diesem Rand. Dann gilt:

$$\int_{\Omega} \operatorname{div}(\vec{p}) \, dA = \int_{\partial\Omega} \langle \vec{p}, \vec{\nu} \rangle \, ds$$

Die in einem Volumen aus den Quellen des Vektorfeldes erzeugten Vektoren müssen durch dessen

Oberfläche hindurchtreten. Integriert man die aus dem Volumen austretenden Vektoren über die gesamte Oberfläche, dann muss man gerade den im Innern des Volumens erzeugten Strom erhalten.

Beispiel:

Gegeben sei das Gebiet Ω und das Vektorfeld p mit

$$\Omega = (0, 2) \times (0, 1) \quad , \quad p(x, y) = \begin{pmatrix} x^2 \cdot y \\ x \cdot y^2 \end{pmatrix}$$

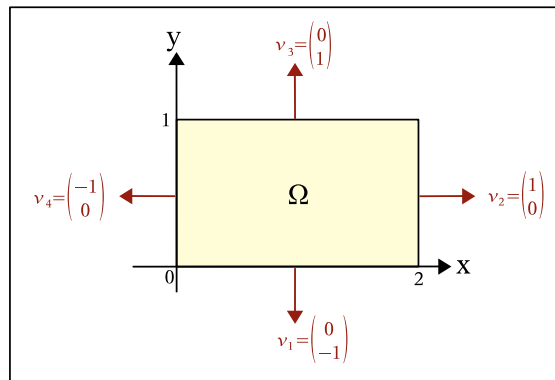


Abbildung 2.1: Gebiet Ω mit Normalenvektoren ν_i auf den Rändern

Gesucht ist das Volumen des Vektorfeldes über dem Gebiet.

• **direkte Berechnung**

Mit

$$\operatorname{div}(p(x, y)) = \partial_x p_x(x, y) + \partial_y p_y(x, y) = 2 \cdot xy + 2 \cdot xy = 4 \cdot xy$$

folgt für das Integral:

$$\begin{aligned} \int_{\Omega} \operatorname{div}(p) \, dA &= \int_{\Omega} \operatorname{div}(p) \, dx \, dy = \int_0^2 \int_0^1 \operatorname{div} \begin{pmatrix} x^2 \cdot y \\ x \cdot y^2 \end{pmatrix} \, dx \, dy \\ &= \int_0^2 \int_0^1 4 \cdot x \cdot y \, dx \, dy = 4 \cdot \int_0^2 x \, dx \cdot \int_0^1 y \, dy \\ &= 4 \cdot \left[\frac{x^2}{2} \right]_0^2 \cdot \left[\frac{y^2}{2} \right]_0^1 = 4 \end{aligned}$$

• **Berechnung mit Hilfe des Integralsatzes von Gauß**

$$\begin{aligned}
\int_{\Omega} \operatorname{div}(p) \, dA &= \int_{\partial\Omega} \langle \vec{p}, \vec{\nu} \rangle \, ds = \sum_{i=1}^4 \int_{\partial\Omega_i} \langle \vec{p}|_{\partial\Omega_i}, \vec{\nu}_i \rangle \, ds \\
&= \int_0^2 \langle p(x, 0), \nu_1 \rangle \, dx + \int_0^1 \langle p(2, y), \nu_2 \rangle \, dy \\
&\quad + \int_0^2 \langle p(x, 1), \nu_3 \rangle \, dx + \int_0^1 \langle p(0, y), \nu_4 \rangle \, dy \\
&= \int_0^2 \begin{pmatrix} 0 \\ 0 \end{pmatrix}^T \cdot \begin{pmatrix} 0 \\ -1 \end{pmatrix} \, dx + \int_0^1 \begin{pmatrix} 4y \\ 2y^2 \end{pmatrix}^T \cdot \begin{pmatrix} 1 \\ 0 \end{pmatrix} \, dy \\
&\quad + \int_0^2 \begin{pmatrix} x^2 \\ x \end{pmatrix}^T \cdot \begin{pmatrix} 0 \\ 1 \end{pmatrix} \, dx + \int_0^1 \begin{pmatrix} 0 \\ 0 \end{pmatrix}^T \cdot \begin{pmatrix} -1 \\ 0 \end{pmatrix} \, dy \\
&= \int_0^1 4 \cdot y \, dy + \int_0^2 x \, dx \\
&= 2 + 2 = 4
\end{aligned}$$

2.7 Das ADI-Verfahren

Das ADI-Verfahren (alternating direction implicit) wurde speziell für die Lösung von Differentialgleichungen im Mehrdimensionalen entwickelt. Liegt ein großes zwei-dimensionales Problem vor, zerlegt das ADI-Verfahren dieses in zwei kleinere ein-dimensionale Probleme. Dabei wird ein Zeitschritt des Problems in zwei Halbschritte aufgeteilt. Im ersten Halbschritt wird in x-Richtung explizit und in y-Richtung implizit gerechnet und im zweiten Halbschritt umgekehrt.

Beispiel:

Gegeben ist die 2D-Diffusionsgleichung

$$u_t = u_{xx} + u_{yy} \quad (2.14)$$

und deren diskrete Form

$$u_t = A_1 u + A_2 u. \quad (2.15)$$

Zerlege das zwei-dimensionale Problem in zwei (einfachere) ein-dimensionale Probleme. Die Näherungen u^n sind zeitabhängig ($t^n = t^0 + n \cdot \Delta t$) und werden iterativ bestimmt:

$$\frac{u^{n+1} - u^n}{\Delta t} = \frac{1}{2}(A_1 u^{n+1} + A_2 u^{n+1}) + \frac{1}{2}(A_1 u^n + A_2 u^n) + O(\Delta t^2) \quad (2.16)$$

bzw.

$$\left(I - \frac{\Delta t}{2}A_1 - \frac{\Delta t}{2}A_2\right)u^{n+1} = \left(I + \frac{\Delta t}{2}A_1 + \frac{\Delta t}{2}A_2\right)u^n + O(\Delta t^3) \quad (2.17)$$

Durch Umformen ergibt sich:

$$\left(I - \frac{\Delta t}{2}A_1\right) \cdot \left(I - \frac{\Delta t}{2}A_2\right)u^{n+1} = \left(I + \frac{\Delta t}{2}A_1\right) \cdot \left(I + \frac{\Delta t}{2}A_2\right)u^n + O(\Delta t^2) \quad (2.18)$$

Zur Lösung dieses Systems wird die Näherung v in zwei Schritten ermittelt:

$$\left(I - \frac{\Delta t}{2}A_1\right)v^{n+\frac{1}{2}} = \left(I + \frac{\Delta t}{2}A_2\right)v^n \quad (2.19)$$

und

$$\left(I - \frac{\Delta t}{2}A_2\right)v^{n+1} = \left(I + \frac{\Delta t}{2}A_1\right)v^{n+\frac{1}{2}} \quad (2.20)$$

Kapitel 3

Diskretisierung des Randwertproblems

3.1 Einleitung

Aus Kapitel 2.1 ist bekannt, dass das Randwertproblem im 2D-Fall folgende Form besitzt:

$$-\operatorname{div}(\sigma \cdot \operatorname{grad}(u)) = I \quad \text{in } \Omega \quad (3.1)$$

$$u = 0 \quad \text{auf } \Gamma_d \quad (3.2)$$

$$\partial_\nu u = 0 \quad \text{auf } \Gamma_n \quad (3.3)$$

Die Operatoren *Divergenz* und *Gradient* werden mit Hilfe der zentralen finiten Differenzen approximiert. Die einzelnen Berechnungen werden auf verschiedenen äquidistanten Gittern durchgeführt. Hierbei wird zwischen drei verschiedenen Gittern unterschieden:

- u-Gitter: diskrete Form des Potentials (vgl. Abbildung 3.1)
- p-Gitter: Produkt aus diskretem Gradienten und der diskreten Leitfähigkeit (vgl. Abbildung 3.3)
- f-Gitter: diskrete Form des injizierten Stroms (vgl. Abbildung 3.4)

Die Auflösung aller Gitter wird durch den Approximationsgrad des u-Gitters bestimmt, welches in $n_x \times n_y$ Gitterpunkte unterteilt werden soll. Die Abstände der einzelnen Punkte sind bei allen drei Gittern gleich,

$$h_x := \frac{x_2 - x_1}{n_x} \quad \text{und} \quad h_y := \frac{y_2 - y_1}{n_y} \quad (3.4)$$

wobei x_1 und x_2 die Begrenzungen des Gebiets in x- bzw. y_1 und y_2 die Begrenzungen in y-Richtung sind.

Da jeder Operator diskretisiert werden muss, werden sowohl der Gradient als auch die Divergenz durch ein Matrix-Vektor Produkt dargestellt. Auch die Multiplikation mit der Leitfähigkeit und die rechte Seite des Randwertproblems werden durch Matrizen ausgedrückt. Somit wird die kontinuierliche Gleichung (3.1) in die diskrete Gleichung

$$-D \cdot A \cdot G \cdot u = F \quad (3.5)$$

überführt.

Erläuterung:

- D : Matrix für die diskrete Divergenz
- A : Matrix, in der die Leitfähigkeiten eingetragen sind
- G : Matrix für den diskreten Gradienten
- F : Matrix der diskreten rechten Seite

Um die Randbedingungen zu erfüllen, wird eine weitere Matrix B benötigt, die die Dirichlet-Ränder erzwingen soll. Es ergibt sich das vollständige Gleichungssystem

$$-D \cdot A \cdot G \cdot u + B \cdot u = F. \quad (3.6)$$

Die Matrizen D , G und B sind bei einer festen Gittereinteilung (siehe Kapitel 3.2) konstant. Die Matrizen A und F ändern sich je nach Leitfähigkeit σ bzw. rechter Seite I . Somit lässt sich das diskrete Problem

$$\underbrace{(-D \cdot A(\sigma) \cdot G + B)}_{:=R(\sigma)} \cdot u = F(I) \quad (3.7)$$

in die vereinfachte Gleichung

$$R(\sigma) \cdot u(\sigma) = F(I) \quad (3.8)$$

umformen.

Im Folgenden wird nun erläutert, wie die einzelnen Matrizen aufgestellt werden. Um ein besseres Verständnis zu erlangen, werden alle Matrizen anhand eines kleinen Beispiels berechnet. Hierzu wird das Gebiet in ein äquidistantes Gitter der Größe $n_x = 4 \times 3 = n_y$ eingeteilt (vgl. Kapitel 3.2).

Bemerkung:

Zur Vereinfachung der Diskretisierung werden zunächst alle Stromquellen genau auf den Gitterpunkten positioniert. Dies hat den Vorteil, dass sich die diskrete rechte Seite $F(I)$ sehr leicht ermitteln lässt. In Kapitel 3.10.1 wird die beliebige Positionierung der Stromquellen in Ω erläutert.

3.2 Einteilung des Gebietes in ein äquidistantes Gitter

Um eine diskrete Variante der in Kapitel 2.1 beschriebenen Differentialgleichung zu erhalten, wird das Gebiet Ω zunächst durch ein äquidistantes Gitter beschrieben. Dieses Gitter wird durch die Anzahl der Punkte n_x in x-Richtung und n_y in y-Richtung definiert.

Somit wird das kontinuierliche Potential u durch $n_x \cdot n_y$ diskrete u-Gitterpunkte (vgl. Abbildung 3.1) approximiert. Je größer die Anzahl der Punkte ist, desto feiner wird das Gitter und desto besser wird das Potential angenähert.

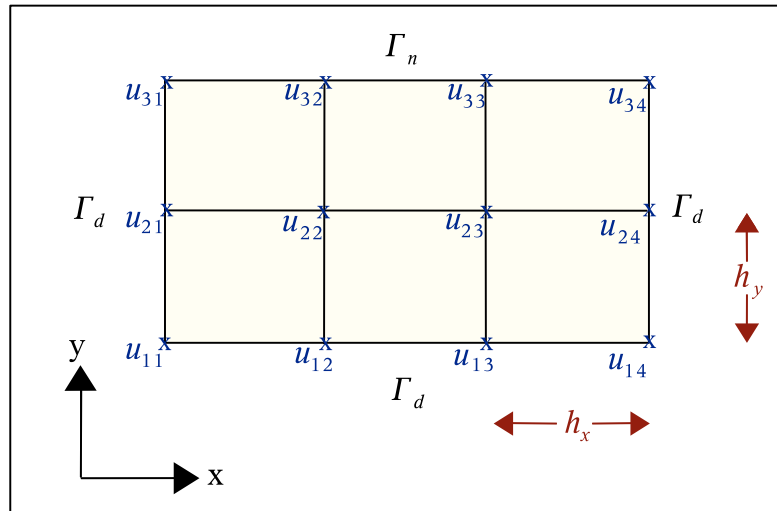


Abbildung 3.1: äquidistantes u-Gitter auf Ω , $n_x = 4, n_y = 3$

Bemerkung:

Sowohl bei den Zeichnungen als auch bei den Herleitungen der Matrizen werden die Gitterpunkte ähnlich einer Matrix durchnummeriert (siehe Abbildung 3.1). Dies erleichtert das Verständnis beim Aufstellen der diskreten Gleichungen.

Intern wird das u-Gitter jedoch auf einem Vektor abgespeichert:

$$\begin{pmatrix} u_{31} & u_{32} & u_{33} & u_{34} \\ u_{21} & u_{22} & u_{23} & u_{24} \\ u_{11} & u_{12} & u_{13} & u_{14} \end{pmatrix} \Rightarrow \begin{pmatrix} u_9 & u_{10} & u_{11} & u_{12} \\ u_5 & u_6 & u_7 & u_8 \\ u_1 & u_2 & u_3 & u_4 \end{pmatrix} \Rightarrow \begin{pmatrix} u_1 \\ u_2 \\ \dots \\ u_{11} \\ u_{12} \end{pmatrix} \quad (3.9)$$

Somit werden auch alle weiteren Gittergrößen, wie z.B. der Gradient des Potentials, durch einen Vektor dargestellt. Im Folgenden werden daher zunächst alle Herleitungen und Bilder zum besseren Verständnis an einer Gitterstruktur verdeutlicht, die Matrizen-Aufstellung jedoch anhand des durchnummerierten Vektors.

3.3 Diskretisierung des Gradienten

Der Gradient repräsentiert alle partiellen Ableitungen einer Funktion. Da das Gebiet zwei-dimensional ist, muss eine partielle Ableitung sowohl in x- als auch in y-Richtung bestimmt werden. Auch die diskreten Gradienten werden bezüglich ihrer Richtung unterschieden. Da der Gradient aus den zwei benachbarten u-Gitterpunkten ermittelt wird, wird dieser genau in der Mitte der beiden Punkte positioniert (vgl. Abbildung 3.2).

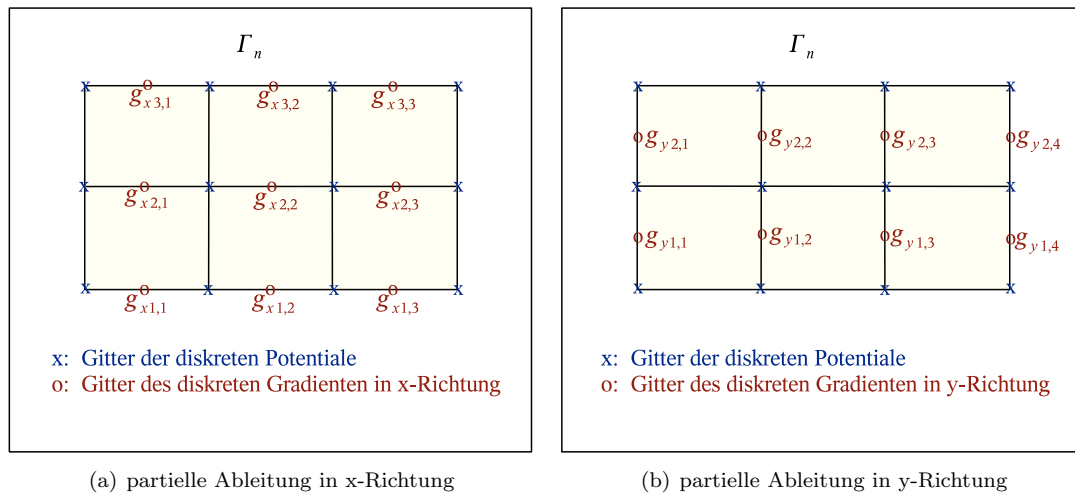


Abbildung 3.2: Aufteilung des diskreten Gradienten

Die partielle Ableitung in x- bzw. in y-Richtung wird mit Hilfe der zentralen finiten Differenzen über u bestimmt und ist wie folgt definiert:

$$\partial_x u_{ij} \approx g_{x\,ij} = \frac{1}{h_x} \cdot (u_{i+1\,j} - u_{i\,j}) \quad (3.10)$$

$$\partial_y u_{ij} \approx g_{y\,ij} = \frac{1}{h_y} \cdot (u_{i\,j+1} - u_{i\,j}) \quad (3.11)$$

Somit lässt sich der Gradient durch eine einfache Matrix-Vektor Multiplikation approximieren:

$$\text{grad}(u) = \begin{pmatrix} \partial_x u \\ \partial_y u \end{pmatrix} \approx \begin{pmatrix} G_x \\ G_y \end{pmatrix} \cdot u = G \cdot u \quad (3.12)$$

Die beiden Matrizen G_x und G_y sind dünn besetzt. Es sind 2-Band Matrizen, da der Gradient nur von dem rechten und linken bzw. oberen und unteren Wert des u-Gitters (vgl. Abbildung 3.1) abhängt. Somit ist die zusammengefasste Matrix G eine 4-Band Matrix.

In der internen Schreibweise (Potential u als Vektor) sind die einzelnen Gitterpunkte des Gradi-

enten nun wie folgt angeordnet:

$$\begin{pmatrix} g_{x\ 31} & g_{x\ 32} & g_{x\ 33} \\ g_{x\ 21} & g_{x\ 22} & g_{x\ 23} \\ g_{x\ 11} & g_{x\ 12} & g_{x\ 13} \end{pmatrix} \Rightarrow \begin{pmatrix} g_7 & g_8 & g_9 \\ g_4 & g_5 & g_6 \\ g_1 & g_2 & g_3 \end{pmatrix} \Rightarrow \begin{pmatrix} g_1 \\ g_2 \\ \dots \\ g_8 \\ g_9 \end{pmatrix}$$

$$\begin{pmatrix} g_{y\ 21} & g_{y\ 22} & g_{y\ 23} & g_{y\ 24} \\ g_{y\ 11} & g_{y\ 12} & g_{y\ 13} & g_{y\ 14} \end{pmatrix} \Rightarrow \begin{pmatrix} g_{14} & g_{15} & g_{16} & g_{17} \\ g_{10} & g_{11} & g_{12} & g_{13} \end{pmatrix} \Rightarrow \begin{pmatrix} g_{10} \\ g_{11} \\ \dots \\ g_{16} \\ g_{17} \end{pmatrix}$$

Somit lässt sich der Vektor g in dem 4×3 Beispiel folgendermaßen berechnen:

$$\begin{pmatrix} g_1 \\ g_2 \\ \dots \\ g_{16} \\ g_{17} \end{pmatrix} = \begin{pmatrix} G_x \\ G_y \end{pmatrix} \cdot u = G \cdot \begin{pmatrix} u_1 \\ u_2 \\ \dots \\ u_{11} \\ u_{12} \end{pmatrix}$$

mit

$$G = \begin{pmatrix} -\frac{1}{hx} & \frac{1}{hx} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & -\frac{1}{hx} & \frac{1}{hx} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & -\frac{1}{hx} & \frac{1}{hx} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & -\frac{1}{hx} & \frac{1}{hx} & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & -\frac{1}{hx} & \frac{1}{hx} & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & -\frac{1}{hx} & \frac{1}{hx} & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & -\frac{1}{hx} & \frac{1}{hx} & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & -\frac{1}{hx} & \frac{1}{hx} & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & -\frac{1}{hx} & \frac{1}{hx} & 0 \\ \hline -\frac{1}{hy} & 0 & 0 & 0 & \frac{1}{hy} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & -\frac{1}{hy} & 0 & 0 & 0 & \frac{1}{hy} & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & -\frac{1}{hy} & 0 & 0 & 0 & \frac{1}{hy} & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & -\frac{1}{hy} & 0 & 0 & 0 & \frac{1}{hy} & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & -\frac{1}{hy} & 0 & 0 & 0 & \frac{1}{hy} & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & -\frac{1}{hy} & 0 & 0 & 0 & \frac{1}{hy} & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & -\frac{1}{hy} & 0 & 0 & 0 & \frac{1}{hy} & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & -\frac{1}{hy} & 0 & 0 & 0 & \frac{1}{hy} & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & -\frac{1}{hy} & 0 & 0 & 0 & \frac{1}{hy} \end{pmatrix}$$

Die Dimension dieser Matrix ist 17×12 . Es gibt neun partielle Ableitungen in x- und acht in y-Richtung, die mit Hilfe der zwölf u-Gitterpunkte berechnet werden.

3.4 Diskretisierung der Leitfähigkeit

Wie in der Abbildung 3.2 zu sehen ist, ist der Gradient in der Mitte der Gitterkanten zwischen zwei u-Gitterpunkten definiert. Da im Randwertproblem das Produkt aus der Leitfähigkeit σ und dem Gradienten gebildet werden muss, ist es demnach sinnvoll, die diskreten Leitfähigkeiten an exakt den gleichen Koordinaten zu positionieren. Nun kann punktweise das Produkt berechnet werden, so dass sich ein p-Gitter ergibt, das genau wie das Gradienten-Gitter angeordnet ist und wie folgt definiert wird:

$$p_{x\ i,j} := \sigma_{x\ i,j} \cdot g_{x\ i,j} \quad \begin{array}{l} i = 1, \dots, n_x - 1 \\ j = 1, \dots, n_y \end{array} \quad (3.13)$$

$$p_{y\ k,l} := \sigma_{y\ k,l} \cdot g_{y\ k,l} \quad \begin{array}{l} k = 1, \dots, n_x \\ l = 1, \dots, n_y - 1 \end{array} \quad (3.14)$$

Um diese Produkte zu ermitteln, werden zwei Hilfsmatrizen A_x und A_y definiert, auf deren Diagonalen die Leitfähigkeiten der x- bzw. y-Kanten eingetragen sind:

$$\sigma \cdot \text{grad}(u) \approx \begin{pmatrix} A_x & 0 \\ 0 & A_y \end{pmatrix} \cdot \begin{pmatrix} G_x \\ G_y \end{pmatrix} \cdot u = A \cdot G \cdot u \quad (3.15)$$

Für das mitgeführte 4×3 Beispiel wird der Vektor p durch

$$\begin{pmatrix} p_1 \\ p_2 \\ \dots \\ p_{16} \\ p_{17} \end{pmatrix} = \begin{pmatrix} \text{diag}(\sigma_1, \dots, \sigma_9) & 0 \\ 0 & \text{diag}(\sigma_{10}, \dots, \sigma_{17}) \end{pmatrix} \cdot G \cdot u$$

berechnet.

3.5 Diskretisierung der Divergenz

Die Divergenz ist ein Differential-Operator eines Vektorfeldes. Im zwei-dimensionalen Fall lässt sie sich durch die vier benachbarten Punkte berechnen. Da ein Divergenz-Punkt d_{ij} genau in der Mitte der vier umliegenden p-Gitterpunkte liegt (vgl. Abbildung 3.3), wird das Divergenz-Gitter wieder auf das ursprüngliche u-Gitter abgebildet.

Hierbei ist jedoch zu beachten, dass die Divergenz (zunächst)¹ nur in den inneren Gitterpunkten von Ω berechnet werden kann, da die Randpunkte des Gitters keine vier Nachbarn besitzen. Daher wird in diesen Punkten die diskrete Divergenz auf den Wert 0 gesetzt.

¹In diesem Kapitel wird lediglich die Divergenz im Inneren des Gebietes Ω berechnet. Um die Neumann-Bedingung zu erfüllen, muss auch an Γ_n die Divergenz approximiert werden. Genauer wird in Kapitel 3.7.1 erläutert.

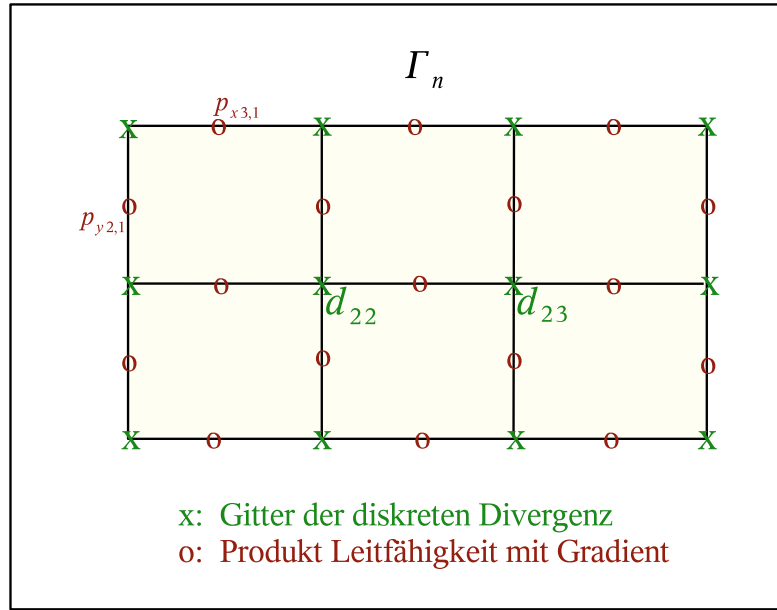


Abbildung 3.3: diskretes Divergenz-Gitter

Die Divergenz lässt sich wie folgt approximieren:

$$d_{ij} = \begin{cases} \frac{1}{h_x}(p_{x\ i,j} - p_{x\ i,j-1}) + \frac{1}{h_y}(p_{y\ i,j} - p_{y\ i-1,j}) & \begin{matrix} i = 2 \dots n_y - 1 \\ j = 2 \dots n_x - 1 \end{matrix} \\ 0 & \text{sonst} \end{cases} \quad (3.16)$$

Da alle Gitterpunkte äquidistant angeordnet sind, werden sowohl zur Berechnung des Gradienten als auch zur Bestimmung der Divergenz die zentralen Differenzenquotienten ermittelt. Dies führt dazu, dass der Ausdruck $-\text{div}(\sigma \cdot \text{grad}(u))$ durch eine Approximation 2. Ordnung diskretisiert wird.

$$D \cdot p = \begin{pmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \frac{-1}{h_x} & \frac{1}{h_x} & 0 & 0 & 0 & 0 & 0 & \frac{-1}{h_y} & 0 & 0 & 0 & \frac{1}{h_y} & 0 & 0 \\ 0 & 0 & 0 & 0 & \frac{-1}{h_x} & \frac{1}{h_x} & 0 & 0 & 0 & 0 & 0 & \frac{-1}{h_y} & 0 & 0 & 0 & \frac{1}{h_y} & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix} \cdot \begin{pmatrix} p_1 \\ p_2 \\ \dots \\ p_{16} \\ p_{17} \end{pmatrix}$$

Jeder innere Divergenz-Gitterpunkt d_{ij} lässt sich aus vier Gradientenkomponenten darstellen. Da diese vier Gradienten aus nur fünf u-Gitterpunkten ermittelt werden, ist

$$L(\sigma) := -D \cdot A(\sigma) \cdot G$$

eine 5 Band Matrix.

Bei einer Leitfähigkeit $\hat{\sigma}$ konstant $1 \Omega^{-1} \left[\frac{1}{\text{Ohm}} \right]$ und den Abmessungen $x = 1m$, $y = 2m$ besitzt $L(\hat{\sigma})$ im mitgeführten Beispiel folgende Form:

$$L(\hat{\sigma}) = \begin{pmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 & -9 & 20 & -9 & 0 & 0 & -1 & 0 & 0 \\ 0 & 0 & -1 & 0 & 0 & -9 & 20 & -9 & 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix}$$

Bemerkung:

Um die Neumann-Bedingungen in die diskrete Variante einzubringen, wird die Matrix D und somit auch $L(\sigma)$ um einige Einträge erweitert. Eine detaillierte Beschreibung der Änderungen wird in Kapitel 3.7.1 angegeben.

3.6 Diskretisierung der rechten Seite der Differentialgleichung

Die rechte Seite I der Differentialgleichung (siehe Kapitel 2.1) repräsentiert den eingeleiteten Strom. Die einzelnen Stromquellen sind δ -Quellen und sind (zunächst)¹ genau auf einem Gitterpunkt positioniert. Hierbei ist zu beachten, dass das f-Gitter der diskreten rechten Seite (vgl. Abbildung 3.4) identisch ist mit dem u-Gitter.

¹Eine allgemeine Positionierung ist in Kapitel 3.10.1 beschrieben.

Eine δ -Quelle ist dadurch charakterisiert, dass das Integral über diese Quelle identisch 1 ist.

$$\int_{\Omega} \delta = 1 \quad (3.17)$$

Diese Eigenschaft soll auch die diskrete rechte Seite F besitzen. Nur an den Gitterpunkten, an denen sich eine δ -Quelle befindet, soll ein Wert $\neq 0$ stehen. Um die Bedingung (3.17) zu erfüllen, muss dieser Wert die Größe $\frac{1}{h_x \cdot h_y}$ besitzen (siehe Kapitel 3.10.1).

Somit kann man bei einer beliebigen Anzahl an δ -Quellen auf dem Gitter eine diskrete rechte Seite F ermitteln. Dies wird in dem mitgeführten Beispiel anhand zwei frei angeordneter δ -Quellen verdeutlicht.

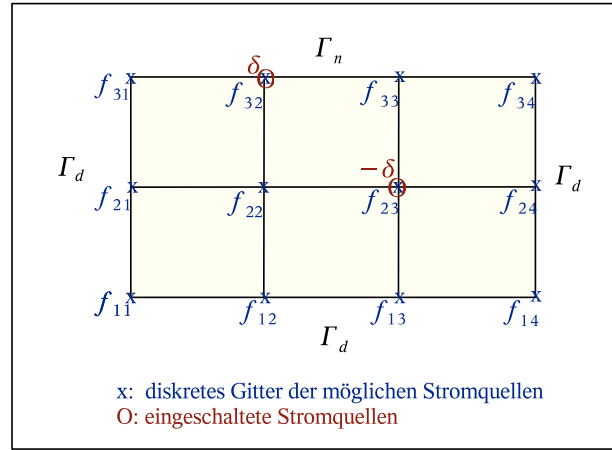


Abbildung 3.4: Diskretisierung der rechten Seite des RWP bei zwei δ -Quellen

Bei dieser Konstellation der zwei δ -Quellen und einer Durchnummerierung des Vektors entsprechend Ausdruck (3.9) besitzt die diskrete rechte Seite F bei dem 4×3 Beispiel folgende Form:

$$F = \begin{pmatrix} f_{11} \\ f_{12} \\ f_{13} \\ f_{14} \\ f_{21} \\ f_{22} \\ f_{23} \\ f_{24} \\ f_{31} \\ f_{32} \\ f_{33} \\ f_{34} \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ -\frac{1}{h_x \cdot h_y} \\ 0 \\ 0 \\ \frac{1}{h_x \cdot h_y} \\ 0 \\ 0 \end{pmatrix}$$

3.7 Die diskreten Randbedingungen

Im letzten Schritt der Diskretisierung des elliptischen Randwertproblems müssen noch die Randbedingungen eingebracht werden. Diese lauten wie folgt:

$$u = 0 \quad \text{auf } \Gamma_d \quad (3.18)$$

$$\partial_\nu u = 0 \quad \text{auf } \Gamma_n \quad (3.19)$$

Hierbei ist zu beachten, dass die Dirichlet-Bedingung (3.18) numerisch anders behandelt wird als die Neumann-Bedingung (3.19).

An der Oberfläche von Ω kommt die Frage auf, ob die Eckpunkte der Neumann- oder der Dirichlet-Bedingung unterliegen. In dieser Arbeit werden diese Eckpunkte als Dirichlet-Bedingung behandelt, da diese das Lösen der Gleichungssysteme vereinfachen (vgl. Kapitel 5.2).

3.7.1 Neumann Bedingung

Die Neumann-Bedingung besagt, dass die Normalen-Ableitung des Potentials am oberen Rand des Gebietes Ω gleich 0 ist. Um diese Bedingung diskret erfüllen zu können, wird das Randwertproblem gedanklich gespiegelt und auf ein doppelt so großes Randwertproblem überführt:

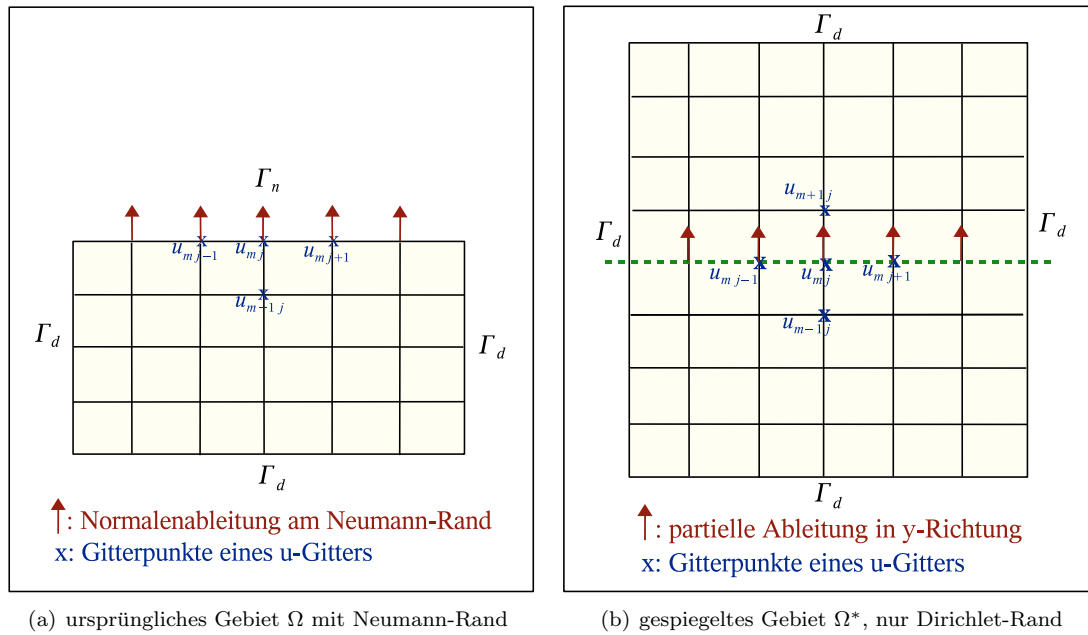


Abbildung 3.5: Umgestaltung des Randwertproblems

Das Gebiet Ω und die rechte Seite I des Randwertproblems werden am oberen Neumann-Rand gespiegelt. Somit existiert an allen Rändern des doppelt so großen Gebietes Ω^* nur noch die Dirichlet-Bedingung. Die Normalen-Ableitung $\partial_\nu u = 0$ ist nunmehr eine partielle Ableitung in-

3.7.2 Dirichlet-Bedingung

Die Dirichlet-Bedingung besagt, dass das Potential am Rand den Wert 0 besitzt. Somit muss auch das lineare Gleichungssystem der diskreten Lösung so aufgebaut sein, dass auch bei der Lösung

$$u = R(\sigma)^{-1} \cdot F \quad (3.27)$$

alle Werte auf dem Dirichlet-Rand diese Bedingung erfüllen. Die in dem Gleichungssystem

$$R(\sigma) \cdot u = (L(\sigma) + B) \cdot u = F \quad (3.28)$$

befindliche Matrix B ist für die Einhaltung der Dirichlet-Bedingung zuständig. Da die Divergenzen nur innerhalb des Gebietes Ω und auf dem Neumann-Rand definiert sind, sind deren diskrete Werte auf dem Dirichlet-Rand 0. Dies führt (vereinfacht) im diskreten Fall zu einer Gleichung:

$$L(\sigma) \cdot u_k = 0 \cdot u_k = f_k \quad , \quad u_k \in \Gamma_d \quad (3.29)$$

Da es numerisch nicht sinnvoll ist, eine δ -Quelle in die Nähe des Dirichlet-Randes zu setzen (vgl. 2.1), kann man den Strom auf Γ_d gleich 0 setzen.

$$f_k = 0 \quad , \quad f_k \in \Gamma_d \quad (3.30)$$

Somit ergibt sich ein Problem, da für jeden Wert von u_k die Gleichung

$$L(\sigma) \cdot u_k = 0 \cdot u_k = 0 \quad , \quad u_k \in \Gamma_d \quad (3.31)$$

erfüllt ist. Um nun den Wert u_k gleich 0 zu erzwingen, muss der Eintrag aus der Matrix B ins Spiel kommen. Vereinfacht sieht das folgendermaßen aus:

$$R(\sigma) \cdot u_k = L(\sigma) \cdot u_k + b \cdot u_k = (0 + b) \cdot u_k = 0 \quad , \quad u_k \in \Gamma_d \quad (3.32)$$

Ist $b \neq 0$, muss u_k den Wert 0 annehmen, damit die Gleichung (3.32) erfüllt ist.

Somit kann man die Matrix B des diskreten Randwertproblems als Diagonalmatrix aufstellen, auf deren Diagonalen b_k der Wert 1 steht, wo sich das Potential u_k auf dem Dirichlet-Rand befindet.

$$B = \text{diag}(b_k) \quad \text{mit} \quad b_k = \begin{cases} 1, & u_k \in \Gamma_d \\ 0, & \text{sonst} \end{cases} \quad (3.33)$$

Anders ausgedrückt bedeutet dies, dass jede Nullzeile z_k von $L(\sigma)$ eine 1 im Diagonalelement L_{kk} eingetragen bekommt. Somit ist die Matrix $R(\sigma)$ regulär.

Im 4×3 Beispiel sieht die Matrix B wie folgt aus:

$$B = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{pmatrix}$$

3.8 Zusammenfassung aller Diskretisierungen

Nachdem jeder Operator als Matrix aufgebaut wurde, kann nun bei gegebener Leitfähigkeit σ die gesamte Matrix $R(\sigma)$ berechnet werden. Diese Matrix ist immer regulär und ist (im zweidimensionalen Modell) durch fünf Bänder charakterisiert. Durch die gewählte Nummerierung des u-Gitters ist eine Matrix entstanden, die dünn besetzt ist und deren Nebendiagonalen jeweils paarweise gleich sind. Lediglich die Neumann-Bedingung verfälscht diese Aussagen ein wenig. Somit besitzt die linke Seite des Gleichungssystems bei einer konstanten Leitfähigkeit $\hat{\sigma}$ von 1 Ω^{-1} , den Abmessungen $x = 1m, y = 2m$ und einer Einteilung in $n_x = 4 \times 3 = n_y$ folgende Form:

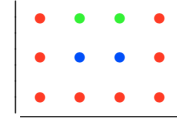
$$R(\hat{\sigma}) \cdot u = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 & -9 & 20 & -9 & 0 & 0 & -1 & 0 & 0 & 0 \\ 0 & 0 & -1 & 0 & 0 & -9 & 20 & -9 & 0 & 0 & -1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & -2 & 0 & 0 & -9 & 20 & -9 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & -2 & 0 & 0 & -9 & 20 & -9 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 \end{pmatrix} \cdot \begin{pmatrix} u_1 \\ u_2 \\ u_3 \\ u_4 \\ u_5 \\ u_6 \\ u_7 \\ u_8 \\ u_9 \\ u_{10} \\ u_{11} \\ u_{12} \end{pmatrix}.$$

3.8.1 Zusammenhang Beispiel-Matrix und Randwertproblem

Für das mitgeführte Beispiel existieren nun folgende Zusammenhänge zwischen der Matrix $R(\hat{\sigma})$ und den Gitterpunkten des Randwertproblems:

$$R(\hat{\sigma}) = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 & -9 & 20 & -9 & 0 & 0 & -1 & 0 & 0 & 0 \\ 0 & 0 & -1 & 0 & 0 & -9 & 20 & -9 & 0 & 0 & -1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & -2 & 0 & 0 & -9 & 20 & -9 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & -2 & 0 & 0 & -9 & 20 & -9 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \end{pmatrix}$$

- Dirichlet-Randbedingung
- Berechnung $-\operatorname{div}(\sigma \cdot \operatorname{grad}(u))$ innerhalb von Ω
- Berechnung $-\operatorname{div}(\sigma \cdot \operatorname{grad}(u))$ am Neumann-Rand



3.9 Approximierte Lösungen einiger Vorwärtsprobleme

Die in den vorherigen Kapiteln hergeleitete diskrete Lösung

$$u = R(\hat{\sigma})^{-1} \cdot F$$

des elliptischen Randwertproblems wird anhand einiger Beispiele vorgeführt. Das zwei-dimensionale Gebiet Ω besitzt in jedem Problem die Abmessungen $(x = 50m) \times (30m = y)$ und eine konstante Leitfähigkeit von $1 \Omega^{-1}$. Am oberen Rand gilt die Neumann-Bedingung, ansonsten Dirichlet (vgl. Abbildung 3.1).

Bemerkung:

Die Positionen der δ -Quellen sind in jedem Beispiel so gewählt, dass diese genau auf einem Gitterpunkt liegen. Beliebige Positionen werden im Kapitel 3.10.1 behandelt.

3.9.1 Lösung eines Vorwärtsproblems bei grober Auflösung

Im ersten Beispiel werden zwei δ -Quellen mitten in Ω positioniert.

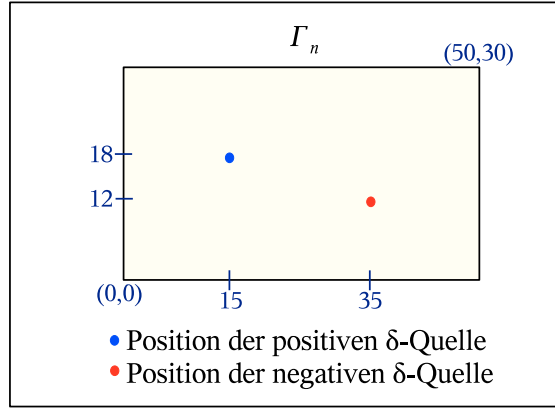


Abbildung 3.7: Positionierung der δ -Quellen in Ω

Bei einer groben Gitterauflösung von $n_x = 11 \times 6 = n_y$ kann das kontinuierliche Potential u nur sehr grob angenähert werden.

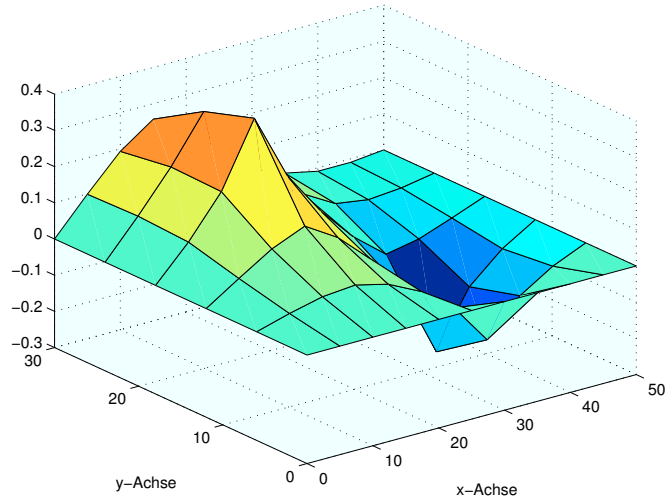


Abbildung 3.8: diskretes Potential bei grober Auflösung

Trotz der geringen Auflösung kann man erkennen, dass die Potentiale am Dirichlet-Rand den Wert 0 besitzen, während am Neumann-Rand Werte $\neq 0$ ermittelt wurden. Da sich die beiden δ -Quellen lediglich in ihrem Vorzeichen unterscheiden, erkennt man in den approximierten Lösungen zwei etwa gleich hohe Erhebungen.

3.9.2 Lösung eines Vorwärtsproblems bei feiner Auflösung

In diesem Beispiel werden die gleichen Positionen der δ -Quellen wie im Beispiel 3.9.1 verwendet. Eine feinere Auflösung des Gitters ($n_x = 51 \times 31 = n_y$) liefert eine deutlich exaktere Approximation der Lösung des kontinuierlichen Vorwärtsproblems.

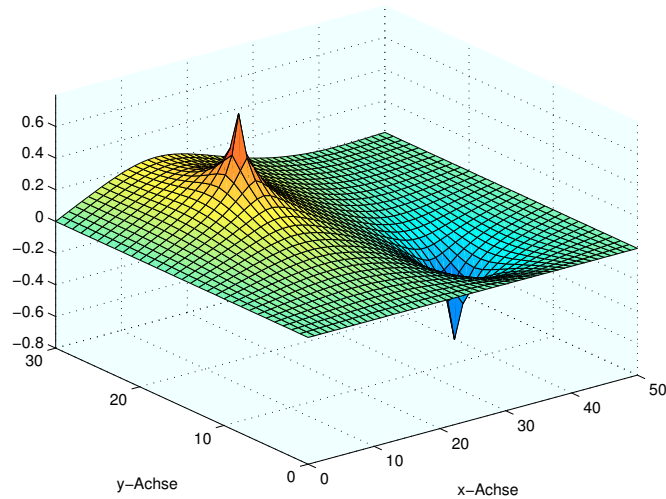


Abbildung 3.9: diskretes Potential bei feiner Auflösung

3.9.3 Positionierung einer δ -Quelle am Neumann-Rand

In diesem Beispiel wird bei einer Auflösung von $n_x = 51 \times 31 = n_y$ eine δ -Quelle direkt auf dem Neumann-Rand positioniert.

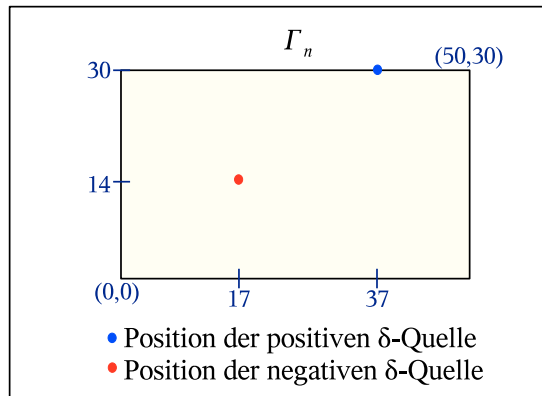


Abbildung 3.10: Positionierung einer δ -Quelle am Neumann-Rand

Hier wird gezeigt, dass auch am Neumann-Rand die diskrete Lösung die Randbedingung erfüllt.

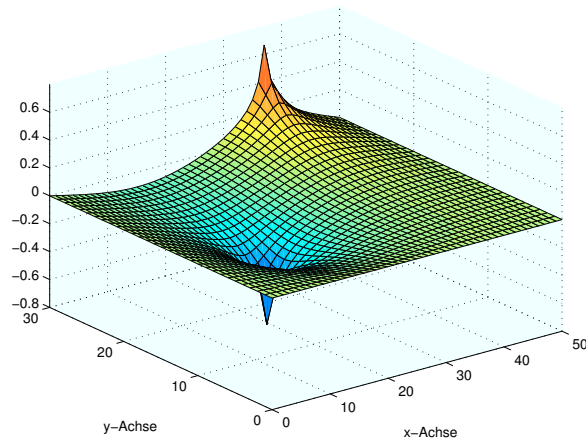


Abbildung 3.11: diskretes Potential mit δ -Quelle am Neumann-Rand

3.9.4 Mehrere δ -Quellen in Ω

Auch wenn mehrere (vier) δ -Quellen eingeschaltet sind, ermittelt das diskrete Randwertproblem eine gute Approximation.

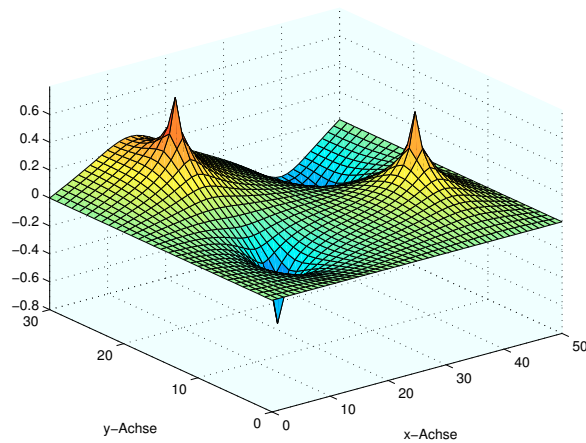


Abbildung 3.12: diskretes Potential bei vier δ -Quellen

3.10 Verfeinerung

Bei dem Feldversuch des Forschungszentrums Jülich werden die Elektroden/Messsonden an festen Stellen angebracht und zunächst nicht mehr bewegt. Damit stellt sich das Problem, dass je nach Auflösung des äquidistanten Gitters sowohl die Elektroden als auch die Messsonden nicht zwingend auf den Gitterpunkten liegen.

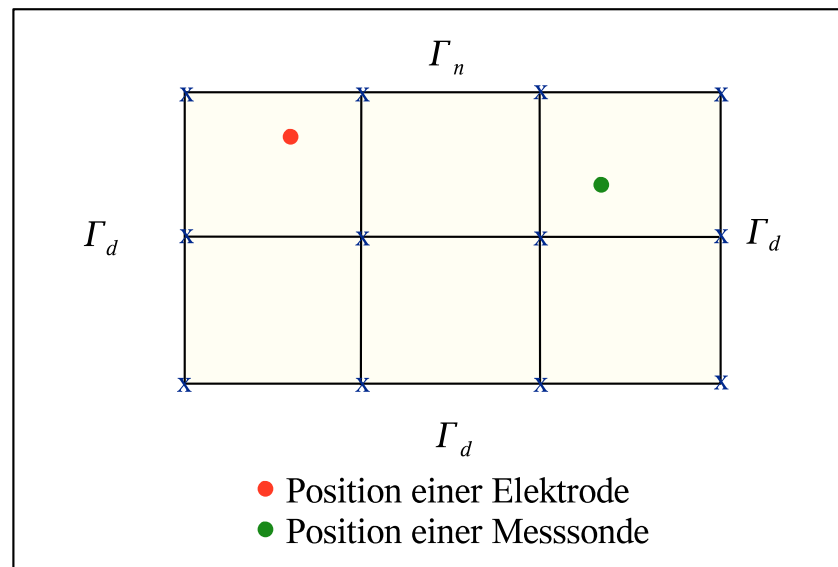


Abbildung 3.13: diskretes Gitter mit Elektrode (δ -Quelle) und Messsonde

Somit kann zum Beispiel ein numerisch berechnetes Potential nicht ein Potential an einer Messsonde repräsentieren. Das Potential an dieser Messsonde muss aus den vier umliegenden u-Gitterpunkten (vgl. Abbildung 3.13) gemittelt werden. Die Berechnung der diskret gemessenen Spannungen wird im Kapitel 3.10.2 erläutert.

Ebenso muss der abgegebene Strom auf das diskrete f-Gitter (vgl. 3.6) aufgeteilt werden. Liegt die Elektrode nicht genau auf einem Gitterpunkt (siehe Abbildung 3.13), wird die δ -Quelle entsprechend auf die vier umliegenden Gitterpunkte verteilt. Genaueres ist in Kapitel 3.10.1 zu lesen. Des Weiteren kann man mit Hilfe mehrerer verschiedener Versuchsdurchführungen (Versuchsaufbauten) mehr Informationen erhalten. Durch die Wahl verschiedener Frequenzen (Kapitel 3.10.3) und einer Reihe verschiedener Messungen (Kapitel 3.10.4) kann man Daten generieren, die später zu einer genaueren Lösung des inversen Problems führen sollen.

Zuletzt wird versucht, die Kosten der Berechnungen des inversen Problems zu senken. Während das Vorwärtsproblem auf dem (in der Regel) sehr feinen u-Gitter berechnet werden soll, kann das inverse Problem (unter Umständen) auf einem gröberen c-Gitter (vgl. Abbildung 3.24) ausgerechnet werden. Diese Überlegungen werden im Kapitel 3.10.5 fortgeführt.

3.10.1 Beliebige Positionierung der δ -Quellen

In der geoelektrischen Widerstandstomographie ist die Gewinnung von relevanten Daten zunächst sehr gering. Demnach ist es sehr wichtig, die Eingangsdaten möglichst genau zu behandeln. So werden die Elektroden im realen Versuch an festen Stellen positioniert und dann nicht mehr verändert.

Da die Elektrode in der Regel nicht genau auf einem f-Gitterpunkt (vgl. Kapitel 3.6) liegt, muss der abgegebene / aufgenommene Strom auf die vier umliegenden Punkte des diskreten Gitters aufgeteilt werden.

Sei nun eine δ -Quelle nicht auf einem Gitterpunkt gelegen (vgl. Abbildung 3.13). Bei der Verteilung des Stromes auf die vier entsprechenden Gitterpunkte muss beachtet werden, dass das Integral über diese verteilte Quelle identisch 1 ist.

$$\int_{\Omega} \delta_{diskret} = 1 \quad (3.34)$$

Um dies auch in der Approximation zu erreichen, werden die diskreten Werte mit Hilfe einer Hütchen-Funktion (siehe Abbildung 3.16) ausgewertet.

Zum besseren Verständnis wird diese Verteilung zunächst im ein-dimensionalen Fall dargestellt und später auf das zwei-dimensionale Problem erweitert.

- 1 dimensional

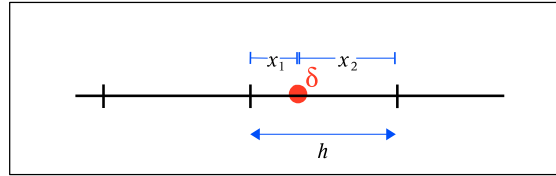


Abbildung 3.14: δ -Quelle beliebig in Ω (1 D)

Um die Verteilung der δ -Quelle auf die beiden f-Gitterpunkte zu ermitteln, wird über jedem dieser Punkte ein Rechteck der Breite h aufgebaut.

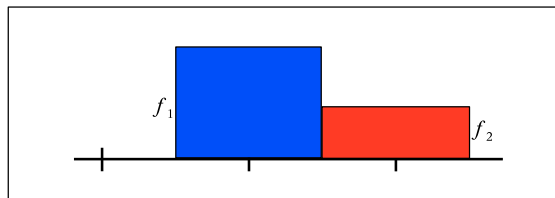


Abbildung 3.15: Rechtecke über diskrete Gitterpunkte definieren

Die Höhen f_1 und f_2 sind abhängig von dem Abstand der δ -Quelle zu den Gitterpunkten.

Demnach wird als erstes eine Gewichtung von

$$\hat{f}_1 = \frac{h - x_1}{h} \quad , \quad \hat{f}_2 = \frac{h - x_2}{h}$$

vermutet. Berechnet man mit diesen Werten nun die Fläche

$$A := \hat{f}_1 \cdot h + \hat{f}_2 \cdot h = \left(\frac{x_2}{h} + \frac{x_1}{h} \right) \cdot h = h \quad (3.35)$$

ist die Bedingung (3.34) nicht erfüllt. Setzt man nun für die Höhen

$$f_1 = \frac{1}{h} \cdot \frac{h - x_1}{h} = \frac{1}{h} \cdot \frac{x_2}{h} \quad (3.36)$$

$$f_2 = \frac{1}{h} \cdot \frac{h - x_2}{h} = \frac{1}{h} \cdot \frac{x_1}{h} \quad (3.37)$$

so ist die Fläche der Rechtecke identisch 1.

Diese Höhen können mit Hilfe einer Hütchen-Funktion berechnet werden.

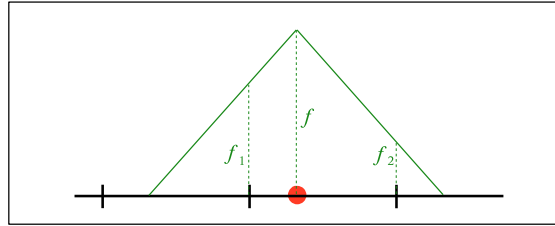


Abbildung 3.16: Hütchen-Funktion über δ -Quelle (1 D)

Die Breite der Hütchen-Funktion ist $2 \cdot h$, da der Wert der δ -Quelle auf zwei diskrete Gitterpunkte aufgeteilt werden soll. Lediglich die Höhe f und somit die gesuchten Werte f_1 und f_2 des Gitters sind unbekannt. Doch diese Höhe f kann mit Hilfe des Strahlensatzes ermittelt werden:

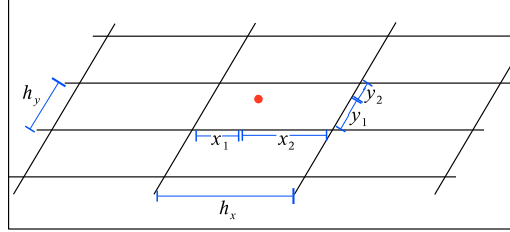
$$\frac{f_1}{x_2} = \frac{f}{h} \quad \text{bzw.} \quad \frac{f_2}{x_1} = \frac{f}{h}. \quad (3.38)$$

Es ergibt sich

$$f = \frac{f_1}{x_2} \cdot h = \frac{\frac{1}{h} \cdot \frac{x_2}{h}}{x_2} \cdot h = \frac{1}{h}. \quad (3.39)$$

Die gewichteten Werte f_1 und f_2 lassen sich somit mit Hilfe der Hütchen-Funktion der Breite $2 \cdot h$ und der Höhe $\frac{1}{h}$ über der δ -Quelle berechnen. Der Wert eines diskreten Gitterpunktes ist daher umgekehrt proportional zu dem Abstand zu der δ -Quelle. Je kleiner der Abstand, desto größer ist die Gewichtung des Gitterpunktes.

- 2 dimensional

Abbildung 3.17: δ -Quelle beliebig in Ω (2 D)

Auch im zwei-dimensionalen Fall wird zunächst eine Hilfe benötigt. Über jedem Gitterpunkt wird ein Quader mit der Grundfläche $h_x \cdot h_y$ aufgebaut.

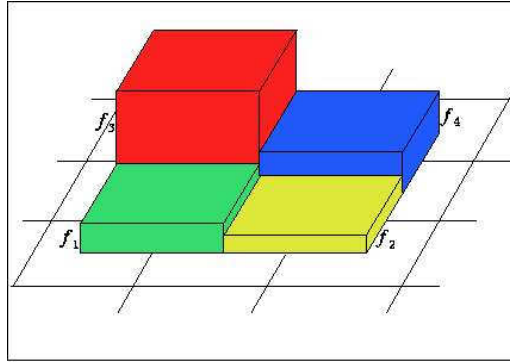


Abbildung 3.18: Quadrate über diskrete Gitterpunkte definieren

Hierbei sind die Höhen jedoch nicht mehr nur von **einem** Abstand der δ -Quelle zum entsprechenden Gitterpunkt, sondern von den beiden Abständen, in x- bzw. in y-Richtung, abhängig. Daraus ergibt sich analog zu (3.36) und (3.37) für die einzelnen Höhen:

$$f_1 = \left(\frac{1}{h_x} \cdot \frac{h_x - x_1}{h_x} \right) \cdot \left(\frac{1}{h_y} \cdot \frac{h_y - y_1}{h_y} \right) \quad (3.40)$$

$$f_2 = \left(\frac{1}{h_x} \cdot \frac{h_x - x_2}{h_x} \right) \cdot \left(\frac{1}{h_y} \cdot \frac{h_y - y_1}{h_y} \right) \quad (3.41)$$

$$f_3 = \left(\frac{1}{h_x} \cdot \frac{h_x - x_1}{h_x} \right) \cdot \left(\frac{1}{h_y} \cdot \frac{h_y - y_2}{h_y} \right) \quad (3.42)$$

$$f_4 = \left(\frac{1}{h_x} \cdot \frac{h_x - x_2}{h_x} \right) \cdot \left(\frac{1}{h_y} \cdot \frac{h_y - y_2}{h_y} \right) \quad (3.43)$$

Berechnet man das Volumen V der vier Quader, so ist mit

$$\begin{aligned} V &= (h_x \cdot h_y) \cdot (f_1 + f_2 + f_3 + f_4) \\ &= \frac{1}{h_x \cdot h_y} \cdot (x_2 \cdot y_2 + x_1 \cdot y_2 + x_2 \cdot y_1 + x_1 \cdot y_1) = 1 \end{aligned}$$

die Bedingung (3.34) erfüllt.
Auch hier können die Höhen

$$f_i \quad , \quad i = 1, \dots, 4$$

über zwei Hütchen-Funktionen berechnet werden. Während eine Funktion entlang der x-Kante (siehe Abbildung 3.19) aufgestellt wird, wird die andere entlang der y-Kante (siehe Abbildung 3.20) definiert. Die Breiten sind (analog zur ein-dimensionalen Herleitung) $2 \cdot h_x$ bzw. $2 \cdot h_y$, während mit $f_x = \frac{1}{h_x}$ bzw. $f_y = \frac{1}{h_y}$ die Höhen angesetzt werden.

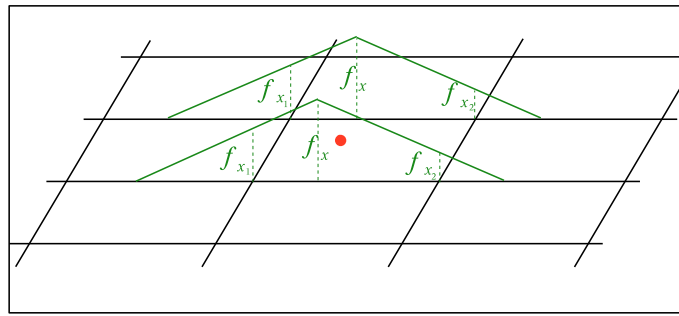


Abbildung 3.19: Hütchen-Funktion entlang der x-Kante

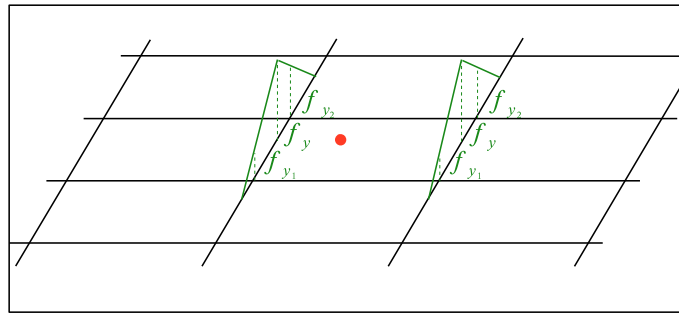


Abbildung 3.20: Hütchen-Funktion entlang der y-Kante

Somit lassen sich die einzelnen Höhen durch die Produkte der beiden Hütchen-Funktionen (ausgewertet an den entsprechenden Kanten) berechnen.

$$f_1 = f_{x_1} \cdot f_{y_1} = \left(\frac{1}{h_x} \cdot \frac{h_x - x_1}{h_x} \right) \cdot \left(\frac{1}{h_y} \cdot \frac{h_y - y_1}{h_y} \right) \quad (3.44)$$

$$f_2 = f_{x_2} \cdot f_{y_1} = \left(\frac{1}{h_x} \cdot \frac{h_x - x_2}{h_x} \right) \cdot \left(\frac{1}{h_y} \cdot \frac{h_y - y_1}{h_y} \right) \quad (3.45)$$

$$f_3 = f_{x_1} \cdot f_{y_2} = \left(\frac{1}{h_x} \cdot \frac{h_x - x_1}{h_x} \right) \cdot \left(\frac{1}{h_y} \cdot \frac{h_y - y_2}{h_y} \right) \quad (3.46)$$

$$f_4 = f_{x_2} \cdot f_{y_2} = \left(\frac{1}{h_x} \cdot \frac{h_x - x_2}{h_x} \right) \cdot \left(\frac{1}{h_y} \cdot \frac{h_y - y_2}{h_y} \right) \quad (3.47)$$

Bemerkung:

Im ein-dimensionalen Fall ist die Höhe eines Rechtecks umgekehrt proportional zum Verhältnis vom Abstand der δ -Quelle zum Gitterpunkt und der Gitterbreite:

$$f_i = \frac{1}{h} \cdot \frac{h - x_i}{h}. \quad (3.48)$$

Bei der Berechnung der Höhe im zwei-Dimensionalen sind es nun zwei Abstände, über die die Höhen der Quader definiert werden:

$$f_i = \frac{1}{h_x \cdot h_y} \cdot \frac{(h_x - x_i) \cdot (h_y - y_i)}{h_x \cdot h_y}. \quad (3.49)$$

Diese Verteilung der δ -Quelle(n) auf das diskrete Gitter kann durch eine Matrix dargestellt werden, in der die einzelnen Gewichte eingetragen sind.

3.10.2 Position der Messpunkte

Auch die Sensoren müssen nicht zwingend auf einem Gitterpunkt liegen. Im Gegensatz zu den Elektroden, bei denen eine δ -Quelle **auf** die vier benachbarten f-Gitterpunkte verteilt wird, muss bei den Messungen ein Potential **aus** den vier umliegenden Potentiale des u-Gitters berechnet werden.

Befindet sich die Sonde exakt in der Mitte eines Gitter-Rechtecks, so ist der numerisch gemessene Wert genau $\frac{1}{4}$ der Summe der vier angrenzenden Gitterpunkte. Bei einem Messpunkt genau auf einem Gitterpunkt wird nur dieser Wert gewichtet. Somit ist auch der Anteil eines Gitterpunktes umgekehrt proportional zur Entfernung zur Messsonde.

Auch diese Gewichtung der Potentialpunkte zu den Messstellen kann man über eine Matrix definieren. Dabei repräsentiert jede Zeile der Matrix die Gewichtung von maximal vier u-Gitterpunkten zu einer Messstelle. Die Summe aller Zeilenelemente muss 1 sein, damit die Gewichtung richtig berechnet wurde. Durch die wenigen Messstellen bei einem Versuch ist auch diese Matrix sehr dünn besetzt.

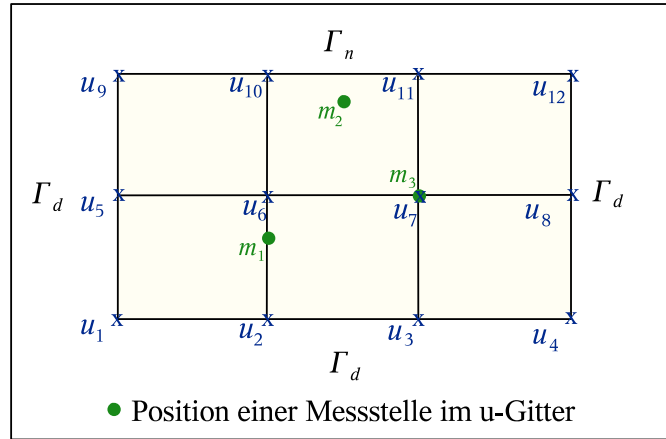


Abbildung 3.21: Beispiel-Anordnung der Messsonden

Für diese Anordnung der Messstellen im 4×3 Gitter gilt:

$$\begin{pmatrix} m_1 \\ m_2 \\ m_3 \end{pmatrix} = \begin{pmatrix} 0 & 0.4 & 0 & 0 & 0 & 0.6 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0.1 & 0.1 & 0 & 0 & 0.4 & 0.4 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \end{pmatrix} \cdot u \quad (3.50)$$

3.10.3 Verschiedene Frequenzen

Eine Möglichkeit, noch mehr Informationen zu gewinnen, ohne die Anordnung der Elektroden und Messsensoren neu zu gestalten, ist die Messung bei verschiedenen Frequenzen. Da die Leitfähigkeit frequenzabhängig ist, kann man bei verschiedenen Frequenzen, aber gleichen Versuchsaufbauten, einige neue Daten erfassen.

Am bedeutsamsten ist der Übergang vom Gleichstrom ($\omega = 0$) zum Wechselstrom ($\omega > 0$). Während die Leitfähigkeit beim Gleichstrom reell ist, erhält man bei den übrigen Frequenzen komplexe Leitfähigkeiten (siehe Kapitel 2.2).

Hierbei sollte man jedoch beachten, dass man neue Frequenzen geschickt wählt. Zwei zu nah aneinander liegende Frequenzen liefern in etwa das gleiche Potentialbild und somit wenig neue Informationen (siehe [6]).

In der numerischen Lösung muss für jede Frequenz das Vorwärtsproblem gelöst und an den entsprechenden Messstellen ausgewertet werden. Diese Daten werden für das inverse Problem mit ihren im Versuch zu der Frequenz passenden Potentialen verglichen (siehe Ausdruck (1.6)) und führen somit zu einer genaueren Lösung.

3.10.4 Messprogramme

Eine weitere Möglichkeit zur Informationserweiterung ist die Wahl verschiedener Messprogramme. Es ist nicht zwingend notwendig, immer alle Elektroden und Messsonden einzuschalten. Das Ein- und Ausschalten der jeweiligen Geräte kann nach Belieben eingestellt werden.

- **Schaltung der Elektroden**

Bei mehreren Elektroden gibt es sehr viele Möglichkeiten, diese zu schalten. Einige könnten den Strom injizieren, einige diesen absorbieren und einige sogar ausgeschaltet bleiben. Man sollte jedoch darauf achten, dass der abgegebene Strom auch wieder wegfließen kann.

Somit erhält man je nach Messprogramm weitere nützliche Daten, da jede Schaltung der Elektroden ein anderes Potentialbild erzeugt.

Diese Schaltung der Elektroden kann auch mit Hilfe einer Matrix definiert werden. Jede Spalte der Matrix repräsentiert eine Messung. Eine 1 bzw. -1 steht für Stromabgabe bzw. Stromaufnahme. Eine ausgeschaltete Elektrode wird durch eine 0 dargestellt. Angenommen es gäbe vier Elektroden und es sollen sechs verschiedene Schaltungen durchgeführt werden, dann kann die Schalt-Matrix J der Elektroden folgende Form besitzen:

$$J = \begin{pmatrix} 1 & -1 & 0 & -1 & 1 & 0 \\ 1 & 0 & 1 & 1 & 0 & 0 \\ -1 & 1 & -1 & 0 & 0 & -1 \\ -1 & 0 & 0 & 0 & -1 & 1 \end{pmatrix} \quad (3.51)$$

- **Messprogramm der Sonden**

In der Regel sollte man so viele Daten wie möglich sammeln. Daher ist es eigentlich sinnvoll, alle Sonden eingeschaltet zu haben.

In dem realen Versuchsaufbau besitzen die im Boden befindlichen Geräte sowohl die Fähigkeit, den Stromfluss zu regulieren, als auch die Möglichkeit, diesen zu messen. Im Versuch ist es daher nicht möglich, eine Elektrode auch gleichzeitig als Messsonde zu nutzen. Diese muss ausgeschaltet bleiben.

Somit gibt es bei der Mess-Matrix P exakt gleich viele Spalten wie bei der Schalt-Matrix J . Die Spaltengröße wird durch die Anzahl der Messsonden festgelegt. Hierbei steht die 1 in der Spalte für eine eingeschaltete Sonde und die 0 für einen Messpunkt, der ausgeschaltet bleibt.

Bei vier Messpunkten und obiger Schalt-Matrix J kann die Mess-Matrix folgende Form haben:

$$P = \begin{pmatrix} 0 & 0 & 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 & 1 & 0 \\ 0 & 1 & 1 & 1 & 0 & 0 \end{pmatrix} \quad (3.52)$$

Bemerkung:

Sei eine feste Frequenz und **eine** Schaltung der Messsonden bei einem Versuch vorgegeben. Nun gibt es n Elektroden im Gebiet, die nach Belieben Strom abgeben, aufnehmen oder ausgeschaltet sein können. Somit sind 3^n verschiedene Schaltungen der Elektroden möglich. Diese liefern aber nicht immer neue Informationen. Viele gemessene Potentiale sind aus Linearkombinationen anderer darstellbar. Dies liegt am linearen Superpositionsprinzip der Potentialverteilung:

Seien im Gebiet zwei δ -Quellen an beliebigen Stellen r_i, r_j $i \neq j$ positioniert. Die Elektrode an der Position r_i gebe den Strom ab, und die andere nehme ihn auf. Dann ist $\phi_{i,j}$ die Potentialverteilung, die das Randwertproblem

$$-\operatorname{div}(\sigma \cdot \operatorname{grad}(\phi_{i,j})) = \delta(r - r_i) - \delta(r - r_j) + \text{Randbedingungen} \quad (3.53)$$

bei einer gegebenen Leitfähigkeit $\hat{\sigma}$ erfüllt.

Seien nun $\phi_{0,1}$ und $\phi_{2,3}$ zwei Lösungen der Randwertprobleme

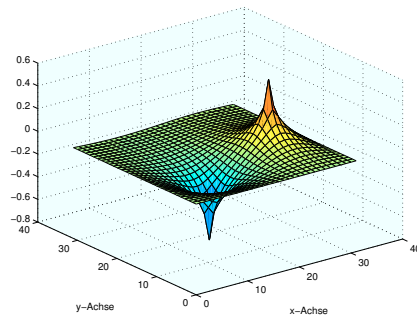
$$-\operatorname{div}(\hat{\sigma} \cdot \operatorname{grad}(\phi_{0,1})) = \delta(r - r_0) - \delta(r - r_1) + \text{RB} \quad (3.54)$$

$$-\operatorname{div}(\hat{\sigma} \cdot \operatorname{grad}(\phi_{2,3})) = \delta(r - r_2) - \delta(r - r_3) + \text{RB}. \quad (3.55)$$

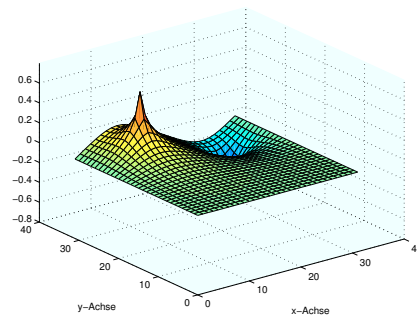
Schaltet man nun diese vier δ -Quellen gleichzeitig und ermittelt dann das Potential $\phi_{02,13}$, so ist die neue Potentialverteilung die zusammengefasste Variante der beiden ursprünglichen.

$$\begin{array}{rcl} \left(\begin{array}{l} -\operatorname{div}(\hat{\sigma} \cdot \operatorname{grad}(\phi_{0,1})) \\ + \left(-\operatorname{div}(\hat{\sigma} \cdot \operatorname{grad}(\phi_{2,3})) \right) \end{array} \right) & = & \begin{array}{l} \delta(r - r_0) - \delta(r - r_1) \\ \delta(r - r_2) - \delta(r - r_3) \end{array} + \text{RB} \\ \hline -\operatorname{div}(\hat{\sigma} \cdot \operatorname{grad}(\underbrace{\phi_{0,1} + \phi_{2,3}}_{\phi_{02,13}})) & = & \delta(r - r_0) + \delta(r - r_2) - \delta(r - r_1) - \delta(r - r_3) + \text{RB} \end{array}$$

Auch in der diskreten Lösung kann man die Linearität erkennen. Es werden zwei Potentialverteilungen auf einem Gebiet mit zwei δ -Quellen berechnet.



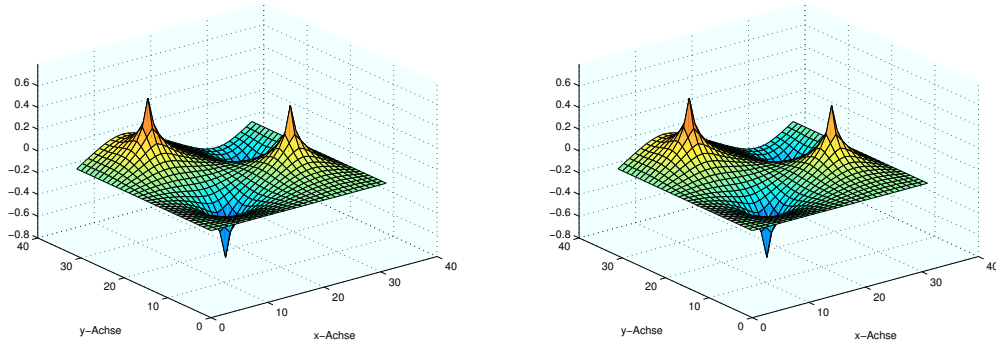
(a) Potentialverteilung $\phi_{0,1}$



(b) Potentialverteilung $\phi_{2,3}$

Abbildung 3.22: Potentialverteilungen bei zwei verschiedenen Stromeinspeisungen

Schaltet man alle δ -Quellen gleichzeitig und berechnet die daraus resultierende Potentialverteilung $\phi_{02,13}$, so ist diese gleich der Summe der beiden einzelnen Verteilungen $\phi_{0,1} + \phi_{2,3}$:



(a) Lösung des zusammengefassten Problems

(b) Summe der einzelnen Potentialverteilungen

Abbildung 3.23: Gegenüberstellung beider Potentialverteilungen

Nun werde eine Elektrode (e_1) von insgesamt n Elektroden immer als Stromausgang genutzt. Somit gibt es $n - 1$ Tupel, so dass nur zwei Stromquellen geschaltet sind:

$$(e_1, e_i) \quad i = 2, \dots, n.$$

Da man alle anderen Schaltungen durch eine Summe dieser $n - 1$ Tupel realisieren kann, gibt es insgesamt $n - 1$ linear unabhängige Potentialverteilungen.

Diese Bemerkung ist in der Numerik ein wenig kritisch zu betrachten. Sind die Eingangsdaten gestört, kann es dennoch sinnvoll sein, mehr als $n - 1$ Schaltungen durchzuführen. Der Grund liegt in der Nicht-Linearität der Störungen, so dass das lineare Superpositionsprinzip nicht angewandt werden kann.

3.10.5 Grobes c-Gitter

Um den Aufwand zur Berechnung der Ableitungsmatrix (siehe Kapitel 4) möglichst gering zu halten, kann man das in der Regel sehr feine u-Gitter der Größe $n_x \times n_y$ durch ein gröberes c-Gitter (vgl. Abbildung 3.24) ersetzen. Da für jedes Rechteck des c-Gitters die Ableitung ausgerechnet werden muss, erspart man sich durch eine gröbere Einteilung eine Vielzahl an Rechnungen. Je größer jedoch das c-Gitter ist, desto größer ist auch die Auflösung des inversen Problems.

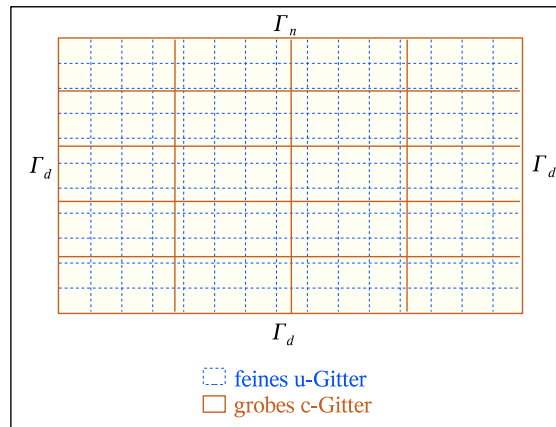


Abbildung 3.24: grobes c-Gitter auf feinem u-Gitter gelegt

Jedes Rechteck des c-Gitters soll die vier Cole-Cole Parameter dieser Region enthalten und somit den Boden in diesem Teilgebiet charakterisieren. Mit einer vorgegebenen Frequenz kann man nun die Leitfähigkeit für dieses Gebiet berechnen und diese auf das feinere u-Gitter übertragen. Dieser Übergang vom groben auf das feinere Gitter kann wieder durch eine Matrix dargestellt werden. Somit wird für die Vorwärtsberechnung das feinere u-Gitter und für die Rückwärtsberechnung das gröbere c-Gitter verwendet.

Kapitel 4

Ableitung und adjungierte Ableitung des Randwertproblems

4.1 Bestimmung der gesuchten Ableitung $\partial_c u(c)(c_0)(c_r)$

Bei dem gegebenen elliptischen Randwertproblem ist die Leitfähigkeit σ eine Funktion $\sigma(c, \omega)$, wobei c die Verteilung der vier Cole-Cole Parameter (siehe Kapitel 2.2) einer Bodenschicht in Ω repräsentiert. Zur Lösung des Vorwärtsproblems muss lediglich die Leitfähigkeit σ betrachtet werden, die sich aus den Cole-Cole Parametern und der Frequenz berechnen lässt.

Die Ableitung hingegen soll bzgl. der Cole-Cole Parameter berechnet werden, weil diese im inversen Problem identifiziert werden sollen. Da bei der Funktion $m(c)$ (siehe Kapitel 1.2) lediglich die Funktion u von den Cole-Cole Parametern abhängt (vgl. Ausdruck (1.7)), werden die partiellen Ableitungen des Potentials u nach den Cole-Cole Parametern in einer aktuellen Verteilung c_0 in Richtung c_r gesucht, also

$$\partial_c u(c)(c_0)(c_r). \tag{4.1}$$

Zur Bestimmung der gesuchten Ableitung (4.1) kann man sowohl bei der kontinuierlichen Gleichung, als auch bei der diskreten Lösung, den *Satz der impliziten Funktionen* (vgl. Satz 1) anwenden.

4.1.1 Bestimmung der Ableitung im kontinuierlichen Fall

Wie in Kapitel 2.4.1 beschrieben, lässt sich das kontinuierliche Randwertproblem als eine *implizite Funktion* darstellen.

$$g(\sigma(c, \omega), u(c, \omega)) := \begin{pmatrix} -\operatorname{div}(\sigma \cdot \operatorname{grad}(u)) - I \\ u|_{\Gamma_d} \\ \partial_\nu u|_{\Gamma_n} \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix} \quad (4.2)$$

Ohne Einschränkung der Allgemeinheit wird nun eine feste Frequenz ω_0 vorgegeben. Somit kann die komplexe Gleichung (4.2) in eine vereinfachte Gleichung überführt werden, da die Leitfähigkeit σ und somit auch das Potential u nur noch von den Cole-Cole Parametern c abhängt.

$$g(c, u(c)) := \begin{pmatrix} -\operatorname{div}(\sigma(c) \cdot \operatorname{grad}(u)) - I \\ u|_{\Gamma_d} \\ \partial_\nu u|_{\Gamma_n} \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix} \quad (4.3)$$

Wendet man nun den *Satz über implizite Funktionen* (Kapitel 2.4.1) auf Gleichung (4.3) an, lässt sich die gesuchte Ableitung $\partial_c u(c)(c_0)(c_r)$ folgendermaßen berechnen:

$$\underbrace{\partial_c u(c)(c_0)(c_r)}_{:=s} = -\partial_u g(c, u(c))(c_0)^{-1} \circ \partial_c g(c, u(c))(c_0)(c_r) \quad (4.4)$$

$$\Leftrightarrow \partial_u g(c, u(c))(c_0) \circ s = -\partial_c g(c, u(c))(c_0)(c_r) \quad (4.5)$$

Mit

$$L(\sigma(c), u) := -\operatorname{div}(\sigma(c) \cdot \operatorname{grad}(u)) \quad (4.6)$$

ist ein Operator definiert, der sowohl in u als auch in σ linear ist ¹. Über diesen Operator lassen sich die partiellen Ableitungen berechnen:

- **Berechnung von $\partial_u g(c, u(c))(c_0)$**

$$\partial_u g(c, u(c))(c_0) = \partial_u \begin{pmatrix} L(\sigma(c), u) - I \\ u|_{\Gamma_d} \\ \partial_\nu u|_{\Gamma_n} \end{pmatrix} (c_0) \quad (4.7)$$

Da der Operator $L(\sigma(c), u)$ linear in u ist, ist die Ableitung $\partial_u L(\sigma(c), u)$ wieder der Operator $L(\sigma(c), u)$ selbst. Auch die Abbildung von u bzw. $\partial_\nu u$ auf die Ränder Γ_d bzw. Γ_n sind durch

¹siehe [6]

lineare Operatoren ausdrückbar. Es folgt

$$\partial_u \begin{pmatrix} L(\sigma(c), u) - I \\ u | \Gamma_d \\ \partial_\nu u | \Gamma_n \end{pmatrix} (c_0) = \begin{pmatrix} L(\sigma(c), u) \\ u | \Gamma_d \\ \partial_\nu u | \Gamma_n \end{pmatrix} (c_0) = \begin{pmatrix} L(\sigma(c_0), u) \\ u | \Gamma_d \\ \partial_\nu u | \Gamma_n \end{pmatrix} \quad (4.8)$$

Es gilt

$$\partial_u g(c, u(c))(c_0) = \begin{pmatrix} L(\sigma(c_0), u) \\ u | \Gamma_d \\ \partial_\nu u | \Gamma_n \end{pmatrix} \quad (4.9)$$

• **Berechnung von $\partial_c g(c, u(c))(c_0)(c_r)$**

$$\partial_c g(c, u(c))(c_0)(c_r) = \partial_c \begin{pmatrix} L(\sigma(c), u(c)) - I \\ u | \Gamma_d \\ \partial_\nu u | \Gamma_n \end{pmatrix} (c_0)(c_r) \quad (4.10)$$

Die Abbildung von u auf $u | \Gamma_d$ bzw. $\partial_\nu u | \Gamma_n$ ist unabhängig von den Cole-Cole Parametern c und somit als konstant zu betrachten. Es folgt:

$$\partial_c \begin{pmatrix} L(\sigma(c), u(c)) - I \\ u | \Gamma_d \\ \partial_\nu u | \Gamma_n \end{pmatrix} (c_0)(c_r) = \begin{pmatrix} \partial_c (L(\sigma(c), u(c)) - I) \\ 0 \\ 0 \end{pmatrix} (c_0)(c_r) \quad (4.11)$$

Nun muss für den Ausdruck $\partial_c (L(\sigma(c), u(c)) - I) (c_0)(c_r)$ die Kettenregel angewandt werden, da die Leitfähigkeit von den Cole-Cole Parametern abhängt.

$$\partial_c (L(\sigma(c), u(c)) - I) (c_0)(c_r) = \partial_c L(\sigma(c), u(c_0)) (c_0)(c_r) \quad (4.12)$$

$$= (\partial_\sigma L(\sigma(c), u(c_0)) \circ \partial_c \sigma(c)) (c_0)(c_r) \quad (4.13)$$

Nun wird wiederum die Linearität des Operators (in σ) ausgenutzt und

$$\partial_\sigma L(\sigma(c), u(c)) = L(\sigma(c), u(c)) \quad (4.14)$$

in Gleichung (4.13) eingesetzt.

$$(L(\sigma(c), u(c)) \circ \partial_c \sigma(c)) (c_0)(c_r) = L(\partial_c \sigma(c), u(c)) (c_0)(c_r) \quad (4.15)$$

$$= L(\partial_c \sigma(c) (c_0)(c_r), u(c_0)) \quad (4.16)$$

Somit gilt für die zweite partielle Ableitung des Ausdrucks (4.5):

$$\partial_c g(c, u(c))(c_0)(c_r) = \begin{pmatrix} L(\partial_c \sigma(c_0)(c_r), u(c_0)) \\ 0 \\ 0 \end{pmatrix} \quad (4.17)$$

Verwendet man nun die Ergebnisse aus den Gleichungen (4.9) und (4.17), lässt sich die Gleichung (4.5) in

$$\begin{pmatrix} L(\sigma(c_0), s) \\ s|_{\Gamma_d} \\ \partial_\nu s|_{\Gamma_n} \end{pmatrix} = - \begin{pmatrix} L(\partial_c \sigma(c_0)(c_r), u(c_0)) \\ 0 \\ 0 \end{pmatrix} \quad (4.18)$$

überführen.

Somit kann die gesuchte Ableitung s über zwei Randwertprobleme bestimmt werden:

1. Randwertproblem:

$$-\operatorname{div}(\sigma(c_0) \cdot \operatorname{grad}(s)) = \operatorname{div}(\partial_c \sigma(c) \cdot \operatorname{grad}(u_0)) \quad \text{in } \Omega \quad (4.19)$$

$$s = 0 \quad \text{auf } \Gamma_d \quad (4.20)$$

$$\partial_\nu s = 0 \quad \text{auf } \Gamma_n \quad (4.21)$$

wobei u_0 Lösung ist von

2. Randwertproblem:

$$-\operatorname{div}(\sigma(c_0) \cdot \operatorname{grad}(u_0)) = f \quad \text{in } \Omega \quad (4.22)$$

$$u_0 = 0 \quad \text{auf } \Gamma_d \quad (4.23)$$

$$\partial_\nu u_0 = 0 \quad \text{auf } \Gamma_n \quad (4.24)$$

4.1.2 Bestimmung der Ableitung im diskreten Fall

Die im Kapitel 4.1.1 ermittelte Gleichung

$$\begin{pmatrix} L(c_0, s) \\ s|_{\Gamma_d} \\ \partial_\nu s|_{\Gamma_n} \end{pmatrix} = - \begin{pmatrix} L(\partial_c \sigma(c)(c_0)(c_r), u_0) \\ 0 \\ 0 \end{pmatrix} \quad (4.25)$$

kann auch für die diskrete Berechnung der Ableitung s genutzt werden. Die linke Seite der Gleichung (4.25) lässt sich, wie in Kapitel 3 beschrieben, durch

$$(-D \cdot A(\sigma(c_0)) \cdot G + B) \cdot s = R(\sigma(c_0)) \cdot s \quad (4.26)$$

approximieren, während

$$-(-D \cdot A(\partial_c \sigma(c_0)(c_r)) \cdot G) \cdot u_0 \quad (4.27)$$

eine Näherung der rechten Seite darstellt. Hierbei wird der Ausdruck $\partial_c \sigma(c)(c_0)$ mit den diskreten c_0 Parametern berechnet und kann durch eine Matrix dargestellt werden, die fortan S_{c_0} genannt wird.

$$\partial_c \sigma(c)(c_0) \xrightarrow{\text{diskret}} S_{c_0} \quad (4.28)$$

Somit ist der Ausdruck $\partial_c \sigma(c)(c_0)(c_r)$ lediglich ein Produkt aus der gerade definierten Matrix S_{c_0} und dem Richtungsvektor c_r .

Gleichung (4.27) wird in die Form

$$(D \cdot A(S_{c_0} \cdot c_r) \cdot G) \cdot u_0 \quad (4.29)$$

überführt, wobei u_0 die Lösung des diskreten Vorwärtsproblems

$$R(\sigma(c_0)) \cdot u = F \quad (4.30)$$

ist. Zusammenfassend lässt sich die diskrete Ableitung s durch

$$\begin{aligned} R(\sigma(c_0)) \cdot s &= (D \cdot A(S_{c_0} \cdot c_r) \cdot G) \cdot u_0 \\ &= (D \cdot A(S_{c_0} \cdot c_r) \cdot G) \cdot (R^{-1}(\sigma(c_0)) \cdot F) \\ \Leftrightarrow s &= R^{-1}(\sigma(c_0)) \cdot (D \cdot A(S_{c_0} \cdot c_r) \cdot G) \cdot (R^{-1}(\sigma(c_0)) \cdot F) \\ R_0 := R(\sigma(c_0)) \Leftrightarrow s &= R_0^{-1} \cdot (D \cdot A(S_{c_0} \cdot c_r) \cdot G) \cdot (R_0^{-1} \cdot F) \end{aligned} \quad (4.31)$$

berechnen.

• Bestimmung der Ableitungsmatrix

Die komplette Ableitungsmatrix J repräsentiert die Ableitung in dem Punkt c_0 in jede beliebige Richtung c_r . Da sich alle Richtungen durch die Basisvektoren $c_{e_1}, \dots, c_{e_n}$ darstellen lassen, muss die Ableitung in Richtung jedes Basisvektors berechnet werden. Diese berechneten Ableitungsvektoren werden hintereinander als Spalte der Matrix abgespeichert, so dass man nach n Schritten die gesamte Ableitungsmatrix J aufgestellt hat. Nun kann man die Ableitung in eine beliebige Richtung c_r durch das Matrix-Vektor Produkt $J \cdot c_r$ bestimmen.

Bemerkung:

Natürlich kann man die Richtung direkt in Gleichung (4.31) einsetzen. Muss man jedoch sehr

viele Richtungsableitungen (mehr als n) im Punkt c_0 ermitteln, ist es sinnvoller, erst die gesamte Matrix aufzustellen.

4.2 Bestimmung des adjungierten Operators

Im Kapitel 4.1 wurde der Operator ermittelt, der die Ableitung eines Potentials in einer Cole-Cole Verteilung c_0 in Richtung einer anderen Verteilung c_r berechnet. Gesucht ist nun der adjungierte Operator zu dem gerade beschriebenen, sowohl kontinuierlich als auch diskret.

Es ist jedoch nicht zwingend, den adjungierten Operator komplett aufzustellen. Lediglich das Ergebnis von dem adjungierten Operator angewandt auf einen Vektor w muss berechnet werden. Im diskreten Fall wird der Operator J durch eine Matrix angenähert. Nicht die ganze Matrix J^* muss aufgestellt werden, das Matrix-Vektor Produkt $J^* \cdot w$ reicht aus.

4.2.1 Bedeutung des adjungierten Operators in der geoelektrischen Tomographie

In der geoelektrischen Tomographie beschreibt der adjungierte Operator den Einzugsbereich eines Sensors in Ω . Angenommen im Punkt x_0 des Gebietes befindet sich ein Sensor. Dort wird bei einer Cole-Cole Verteilung c_k ein Potential $u(c_k)$ ermittelt.

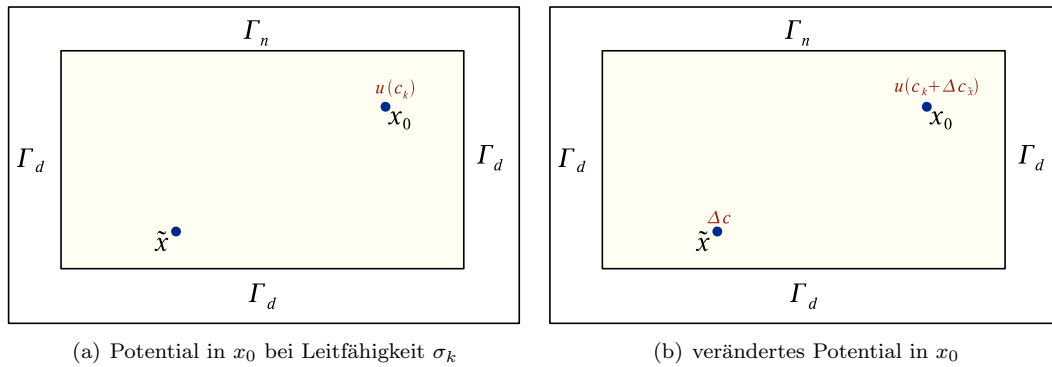


Abbildung 4.1: Änderung der Leitfähigkeit lokal in $\tilde{x}$

Nun werde in einem in Ω befindlichen Punkt $\tilde{x}$ die Cole-Cole Parameter lokal ein wenig verändert. Dies führt zu einem geänderten Potentialwert an der Stelle x_0 . Je weiter der Punkt $\tilde{x}$ von dem Messpunkt x_0 entfernt ist, desto weniger macht sich die Veränderung der Leitfähigkeit in $\tilde{x}$ bemerkbar.

Daher wird in einigen Regularisierungsverfahren (siehe [5]) der adjungierte Operator verwendet, um die ursprünglichen Gleichungen auf einen kleineren Unterraum zu projizieren, bei dem nur die relevanten Werte in Betracht gezogen werden. Das heißt, dass diejenigen Punkte $\tilde{x}_i$ bei einem Sensor x_0 mehr Gewichtung erhalten, die im Einzugsgebiet dieses Messpunktes liegen.

4.2.2 Bestimmung des adjungierten Operators im kontinuierlichen Fall

Sei A_{c_0} der Operator, der die Cole-Cole Parameter c_r auf die Ableitung $\partial_c u(c)(c_0)(c_r)$ abbildet, d.h. in Richtung c_r bei einer Potentialableitung in einem festen Punkt c_0 . Gesucht ist nun der adjungierte Operator $A_{c_0}^*$, so dass gilt:

$$\langle A_{c_0} c_r, w \rangle_{L^2} = \langle c_r, A_{c_0}^* w \rangle_{L^2} \quad (4.32)$$

Mit $s = \partial_c u(c)(c_0)(c_r)$ lässt sich die Ableitung berechnen durch (siehe (4.19)):

$$-\underbrace{\operatorname{div}(\sigma(c_0) \cdot \operatorname{grad}(s))}_{:=R(s)} = \operatorname{div}(\partial_c \sigma(c)(c_0)(c_r) \cdot \operatorname{grad}(u_0)) \quad (4.33)$$

$$\Leftrightarrow s = R^{-1}(-\operatorname{div}(\partial_c \sigma(c)(c_0)(c_r) \cdot \operatorname{grad}(u_0))) \quad (4.34)$$

Für die linke Seite der Gleichung (4.32) lässt sich

$$\langle A_{c_0} c_r, w \rangle_{L^2} = \langle s, w \rangle_{L^2} \quad (4.35)$$

$$= \langle R^{-1}(-\operatorname{div}(\partial_c \sigma(c)(c_0)(c_r) \cdot \operatorname{grad}(u_0))), w \rangle_{L^2} \quad (4.36)$$

bestimmen. Da der Operator R eine Abbildung zwischen zwei (in der Regel komplexen) Räumen darstellt, gilt nach [9]:

$$\langle R^{-1}a, b \rangle_{L^2} = \langle a, R^{-1}\bar{b} \rangle_{L^2} \quad (4.37)$$

Setzt man die Gleichung (4.37) in die Gleichung (4.36) ein, erhält man

$$\left\langle -\operatorname{div}(\underbrace{\partial_c \sigma(c)(c_0)(c_r) \cdot \operatorname{grad}(u_0)}_{:=r}), \underbrace{R^{-1}(\bar{w})}_{:=t} \right\rangle_{L^2} = \langle \operatorname{div}(r), t \rangle_{L^2} \quad (4.38)$$

Nun kann man die L^2 -Norm berechnen:

$$\langle -\operatorname{div}(r), t \rangle_{L^2} = \int_{\Omega} -\operatorname{div}(r) \cdot t \, dx \quad (4.39)$$

Im nächsten Schritt wird der Ausdruck $-\operatorname{div}(r) \cdot t$ umgeformt und anschließend der *Satz von Gauß* (Kapitel 2.6) angewandt:

$$\int_{\Omega} \operatorname{div}(r) \cdot t \, dx = -\int_{\Omega} \operatorname{div}(r \cdot t) - r \cdot \operatorname{grad}(t) \, dx \quad (4.40)$$

$$= -\int_{\Gamma} \nu \cdot r \cdot t \, dx + \int_{\Omega} r \cdot \operatorname{grad}(t) \, dx \quad (4.41)$$

Da $t = R^{-1}\bar{w}$ eine Lösung des ursprünglichen Randwertproblems (siehe Kapitel 2.1) ist, gilt mit

$$t|_{\Gamma_d} = 0 \quad (4.42)$$

$$\partial_\nu t|_{\Gamma_n} = 0 \quad (4.43)$$

für das Integral über den Rändern:

$$\int_{\Gamma} \nu \cdot r \cdot t \, dx = 0 \quad (4.44)$$

Somit bleibt nur noch das zweite Integral übrig, das wieder als ein Skalarprodukt dargestellt werden kann.

$$\langle A_{c_0} c_r, w \rangle_{L^2} = \int_{\Omega} r \cdot \text{grad}(t) \, dx \quad (4.45)$$

$$= \int_{\Omega} (\partial_c \sigma(c)(c_0)(c_r) \cdot \text{grad}(u_0)) \cdot (\text{grad}(R^{-1}(\bar{w}))) \, dx \quad (4.46)$$

$$= \langle \partial_c \sigma(c)(c_0)(c_r), \text{grad}(u_0) \cdot \text{grad}(R^{-1}(\bar{w})) \rangle_{L^2} \quad (4.47)$$

Zuletzt muss die Richtung c_r in der linken Seite des Skalarproduktes isoliert werden. Da die Leitfähigkeit σ von vier Cole-Cole Parametern abhängt (vgl. Kapitel 2.2), gilt:

$$\partial_c \sigma(c)(c_0)(c_r) = \sum_{i=1}^4 \partial_{c_i} \sigma(c)(c_0)(c_{r_i}) \quad (4.48)$$

$$= \sum_{i=1}^4 \partial_i \sigma(c)(c_0) \cdot c_{r_i} \quad (4.49)$$

$$= \begin{pmatrix} \partial_1 \sigma(c)(c_0) \\ \vdots \\ \partial_4 \sigma(c)(c_0) \end{pmatrix}^T \cdot c_r \quad (4.50)$$

Es ergibt sich

$$\langle A_{c_0} c_r, w \rangle_{L^2} = \left\langle c_r, \begin{pmatrix} \overline{\partial_1 \sigma(c)(c_0)} \\ \vdots \\ \overline{\partial_4 \sigma(c)(c_0)} \end{pmatrix} \text{grad}(u_0) \cdot \text{grad}(R^{-1}(\bar{w})) \right\rangle_{L^2} \quad (4.51)$$

$$= \langle c_r, A_{c_0}^* w \rangle_{L^2} \quad (4.52)$$

Somit ist

$$A_{c_0}^* w = \begin{pmatrix} \overline{\partial_1 \sigma(c)(c_0)} \\ \vdots \\ \overline{\partial_4 \sigma(c)(c_0)} \end{pmatrix} \text{grad}(u_0) \cdot \text{grad}(R^{-1}(\bar{w})). \quad (4.53)$$

4.2.3 Bestimmung des adjungierten Operators im diskreten Fall

Im diskreten Fall ist es zunächst einfach, den adjungierten Operator zu der Matrix J zu bestimmen. Man muss lediglich die gesamte Ableitungsmatrix aufstellen und anschließend die hermitesche Matrix J^H berechnen.

Da dies jedoch sehr kostenaufwendig (siehe Kapitel 4.3) und demnach nicht empfehlenswert ist, soll auch hier das Produkt $J^H \cdot w$ direkt berechnet werden. Die Herleitung verläuft ähnlich der kontinuierlichen Herleitung.

$$- \begin{pmatrix} \text{div}(\sigma(c_0) \cdot \text{grad}(s)) \\ s \mid \Gamma_d \\ \partial_\nu s \mid \Gamma_n \end{pmatrix} = \begin{pmatrix} \text{div}(\partial_c \sigma(c)(c_0)(c_r) \cdot \text{grad}(u_0)) \\ 0 \\ 0 \end{pmatrix} \quad (4.54)$$

$$\Rightarrow R_0 \cdot s = D \cdot A(S_{c_0} \cdot c_r) \cdot G \cdot u_0 \quad (4.55)$$

Demnach ergibt sich:

$$\langle J_{c_0} c_r, w \rangle = \langle s, w \rangle \quad (4.56)$$

$$= \langle R_0^{-1} \cdot D \cdot A(S_{c_0} \cdot c_r) \cdot G \cdot u_0, w \rangle \quad (4.57)$$

$$= \left\langle A(S_{c_0} \cdot c_r) \cdot \underbrace{G \cdot u_0}_{:=v}, \underbrace{D^H \cdot R_0^{-H} \cdot w}_{:=b} \right\rangle \quad (4.58)$$

Nun ist es problematisch, im Ausdruck

$$\langle A(S_{c_0} \cdot c_r) \cdot v, b \rangle \quad (4.59)$$

die Richtung c_r zu isolieren, da die Matrix A eine Diagonalmatrix ist, auf deren Einträge die Ableitungen des Potentials in c_0 in Richtung c_r stehen.

$$\langle A(S_{c_0} \cdot c_r) \cdot v, b \rangle = \left\langle \begin{pmatrix} (S_{c_0} \cdot c_r)_1 & & \\ & \ddots & \\ & & (S_{c_0} \cdot c_r)_n \end{pmatrix} \cdot \begin{pmatrix} v_1 \\ | \\ v_n \end{pmatrix}, \begin{pmatrix} b_1 \\ | \\ b_n \end{pmatrix} \right\rangle \quad (4.60)$$

$$= \left\langle \begin{pmatrix} (S_{c_0} \cdot c_r)_1 \cdot v_1 \\ | \\ (S_{c_0} \cdot c_r)_n \cdot v_n \end{pmatrix}, \begin{pmatrix} b_1 \\ | \\ b_n \end{pmatrix} \right\rangle \quad (4.61)$$

Im nächsten Schritt kann man das Skalarprodukt ausschreiben. Hierbei ist wieder zu beachten, dass die einzelnen Komponenten komplex sind.

$$\left\langle \begin{pmatrix} (S_{c_0} \cdot c_r)_1 \cdot v_1 \\ | \\ (S_{c_0} \cdot c_r)_n \cdot v_n \end{pmatrix}, \begin{pmatrix} b_1 \\ | \\ b_n \end{pmatrix} \right\rangle = \sum_{i=1}^n \overline{(S_{c_0} \cdot c_r)_i} \cdot v_i \cdot b_i \quad (4.62)$$

$$= \sum_{i=1}^n \overline{(S_{c_0} \cdot c_r)_i} \cdot (\overline{v_i} \cdot b_i) \quad (4.63)$$

$$= \left\langle S_{c_0} \cdot c_r, \begin{pmatrix} \overline{v_1} b_1 \\ | \\ \overline{v_n} b_n \end{pmatrix} \right\rangle \quad (4.64)$$

$$= (S_{c_0} \cdot c_r)^H \cdot \begin{pmatrix} \overline{v_1} b_1 \\ | \\ \overline{v_n} b_n \end{pmatrix} \quad (4.65)$$

$$= c_r^H \cdot \left(S_{c_0}^H \cdot \begin{pmatrix} \overline{v_1} b_1 \\ | \\ \overline{v_n} b_n \end{pmatrix} \right) \quad (4.66)$$

$$= \left\langle c_r, S_{c_0}^H \cdot \begin{pmatrix} \overline{v_1} b_1 \\ | \\ \overline{v_n} b_n \end{pmatrix} \right\rangle \quad (4.67)$$

Ersetzt man nun die Substitutionen wieder durch ihre ursprüngliche Form, erhält man zu der Matrix J_{c_0} den adjungierten Operator $J_{c_0}^*$.

$$\langle J_{c_0} c_r, w \rangle = \langle c_r, J_{c_0}^* w \rangle \quad (4.68)$$

$$= \left\langle c_r, S_{c_0}^H \cdot \begin{pmatrix} \overline{v_1} \cdot b_1 \\ | \\ \overline{v_n} \cdot b_n \end{pmatrix} \right\rangle \quad (4.69)$$

$$= \left\langle c_r, S_{c_0}^H \cdot \begin{pmatrix} (\overline{G \cdot u_0})_1 \cdot (D^H \cdot R_0^{-H} \cdot w)_1 \\ | \\ (\overline{G \cdot u_0})_n \cdot (D^H \cdot R_0^{-H} \cdot w)_n \end{pmatrix} \right\rangle \quad (4.70)$$

4.3 Kostenvergleich: Ableitungsmatrix gegen Matrix-Vektor Produkt

In diesem Kapitel werden die Kosten verglichen, die bei der Aufstellung der kompletten Ableitungsmatrix J bzw. dem Matrix-Vektor Produkt $J \cdot c$ entstehen. Unter Kosten wird hier die Anzahl der zu lösenden Gleichungssysteme verstanden.

Bei beiden Varianten muss zunächst das Vorwärtsproblem, also ein Gleichungssystem gelöst wer-

den.

Mit der ersten Variante berechnet man die gesamte Ableitungsmatrix J und kann somit problemlos die hermitsche J^H ermitteln. Dies ist jedoch sehr aufwendig, da für jeden Einheitsvektor des Gebietes die Ableitung ermittelt werden muss. Also hat man bei n Einheitsvektoren auch n Gleichungssysteme zu lösen.

In der zweiten Variante werden diese Matrix-Vektor Produkte auf direktem Weg berechnet. Somit müssen auch nur zwei Gleichungssysteme gelöst werden. Hierbei ist jedoch zu beachten, dass diese direkten Produkte bei vielen iterativen Regularisierungsverfahren (siehe [5]) angewandt werden. Demnach muss pro Iterationsschritt ein Gleichungssystem gelöst werden.

• **direktes Aufstellen der Matrix:**

- 1 Randwertproblem für das Lösen des Vorwärtsproblems
- für **jeden** Einheitsvektor das Gleichungssystem lösen
- Berechnung der hermitsche Matrix ohne Kosten

• **Matrix-Vektor Produkt:**

- 1 Randwertproblem für das Lösen des Vorwärtsproblems
- für jeden Iterationsschritt:
 - 1 Randwertproblem für $J \cdot v$
 - 1 Randwertproblem für $J^H \cdot v$

Da bei diesen Versuchen auch noch weitere Messprogramme und verschiedene Frequenzen in die numerischen Berechnungen eingehen, wird somit das Aufstellen der gesamten Ableitungsmatrix sehr teuer.

n_x	n_y	Anzahl Frequenzen	Anzahl Messungen pro Frequenz	Kosten	
				J berechnen	$J \cdot c$ direkt ¹
2	2	1	1	5	3
8	8	1	1	65	3
32	32	1	1	1065	3
32	32	5	25	133125	375

Tabelle 4.1: Kosten der verschiedenen Varianten bei gegebenen Beispielen

¹Kosten bei nur einem Iterationsschritt

Kapitel 5

Spezielle Löser der linearen Gleichungssysteme

5.1 Einführung

Wird das Gebiet in ein u-Gitter der Größe $n_x \times n_y$ eingeteilt (vgl. Kapitel 3.2), so ergibt sich ein Gleichungssystem mit $n := n_x \cdot n_y$ Unbekannten, die sich durch

$$R \cdot u = f, \quad R \in \mathbb{R}^{n \times n} \quad (5.1)$$

berechnen lassen. Neben (5.1) muss (bei gleicher Matrix R) für die Berechnung des adjungierten Operators das System

$$R^H \cdot v = s, \quad R \in \mathbb{R}^{n \times n} \quad (5.2)$$

gelöst werden.

Ohne Einschränkung der Allgemeinheit wird eine feste Frequenz ω_0 und eine Cole-Cole Verteilung c_0 vorgegeben, so dass sich eine Matrix R_0 ergibt, die in sehr vielen Gleichungssystemen verwendet wird. Es ändern sich lediglich die rechten Seiten.

Ziel ist es nun, einen/mehrere Löser zu entwickeln, der einerseits die Eigenschaften von R_0 berücksichtigt, jedoch auch die Lösung der hermiteschen Gleichung (5.2) mit sich bringt.

Bei diesem Randwertproblem ist es möglich, das große Gleichungssystem auf ein kleineres zu reduzieren, ohne dabei das Ergebnis des inversen Problems zu beeinflussen. Die Matrix des verkleinerten Systems wird mit R_{opt} bezeichnet. Es ist sogar möglich, die Matrix R_0 auf eine symmetrische Form zu bringen, die fortan R_{sym} genannt wird. Diese Thematik wird im Kapitel Optimierung (5.2) genauer behandelt.

Da die Matrix R_0 (und deren optimierte Matrix) eine reguläre 5-Band Matrix ist, sollen zunächst die *klassischen* Verfahren das Gleichungssystem lösen. Hierbei wird die *LU-Zerlegung* (Kapitel

5.3) bei beliebigen und die *Cholesky-Zerlegung* (Kapitel 5.4) auf symmetrischen Matrizen durchgeführt. Desweiteren werden speziell angepasste Löser, wie das *ADI-Verfahren* (Kapitel 5.5) und *UMFPACK* (Kapitel 5.6), genauer betrachtet.

Bemerkung:

Auch hier werden alle Matrizen anhand eines kleinen Beispiels veranschaulicht. Es wird eine Gittereinteilung von $n_x = 4 \times n_y = 3$, sowie eine Leitfähigkeit konstant $\frac{1}{100} \Omega^{-1}$ vorgegeben. Es ergibt sich für die Matrix R_0 :

$$R_0 = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & -0.04 & 0 & 0 & -0.09 & 0.26 & -0.09 & 0 & 0 & -0.04 & 0 & 0 \\ 0 & 0 & -0.04 & 0 & 0 & -0.09 & 0.26 & -0.09 & 0 & 0 & -0.04 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & -0.08 & 0 & 0 & -0.09 & 0.26 & -0.09 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & -0.08 & 0 & 0 & -0.09 & 0.26 & -0.09 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{pmatrix}$$

5.2 Optimierung

Um eine Optimierung der Gleichungssysteme durchführen zu können, müssen vorab zwei Bedingungen an das diskrete Problem gestellt werden:

1. Es darf kein Strom auf den Dirichlet-Rand abgebildet werden.

$$f_i = 0, \forall f_i \in \Gamma_d$$

2. Die Messsonden dürfen keine Messungen am Dirichlet-Rand durchführen, d.h. alle Werte werden aus diskreten Punkten innerhalb des Gebietes oder am Neumann-Rand ermittelt. Somit sind die numerischen Ergebnisse am Dirichlet-Rand für die Lösung uninteressant.

Treffen diese Bedingungen nicht ein, müssen entweder die δ -Quellen bzw. Messsonden neu positioniert oder die Auflösung des Vorwärtsproblems verfeinert werden.

5.2.1 Verkleinerung des Gleichungssystems

Will man die Gleichungen (5.1) und (5.2) mit Hilfe eines kleineren Gleichungssystems lösen, muss jede Gleichung einzeln analysiert werden.

- **Berechnung von $R_0 \cdot u = f$**

Wird kein Strom auf dem Rand abgebildet, werden auch alle Potentiale u_i , die auf dem Dirichlet-Rand liegen, den Wert 0 annehmen (vgl. Kapitel 3.7.2). Dies hingegen bedeutet, dass alle Spalten der Matrix R_0 gestrichen werden können, die mit einem $u_i \in \Gamma_d$ multipliziert werden, da das Produkt ohnehin den Wert 0 annehmen wird.

$$\begin{pmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & -0.04 & 0 & 0 & -0.09 & 0.26 & -0.09 & 0 & 0 & -0.04 & 0 & 0 \\ 0 & 0 & -0.04 & 0 & 0 & -0.09 & 0.26 & -0.09 & 0 & 0 & -0.04 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & -0.08 & 0 & 0 & -0.09 & 0.26 & -0.09 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & -0.08 & 0 & 0 & -0.09 & 0.26 & -0.09 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{pmatrix}$$

Es bleiben *NULL-Zeilen* stehen, die keine weitere Verwendung für das Gleichungssystem besitzen. Werden auch diese Zeilen (und die $f_i \in \Gamma_d$) gestrichen, bleibt ein kleineres Gleichungssystem mit der Matrix R_{opt} übrig, welches nur noch die Potentiale u_{in} innerhalb von Ω und am Neumann-Rand identifiziert.

$$R_{opt} \cdot u_{in} = R_{opt} \cdot \begin{pmatrix} u_6 \\ u_7 \\ u_{10} \\ u_{11} \end{pmatrix} = \begin{pmatrix} f_6 \\ f_7 \\ f_{10} \\ f_{11} \end{pmatrix} := f_{in}$$

mit

$$R_{opt} = \begin{pmatrix} 0.26 & -0.09 & -0.04 & 0 \\ -0.09 & 0.26 & 0 & -0.04 \\ -0.08 & 0 & 0.26 & -0.09 \\ 0 & -0.08 & -0.09 & 0.26 \end{pmatrix} \quad (5.3)$$

Die vollständige Lösung u kann nun wieder aus der Lösung von u_{in} bestimmt werden.

$$u = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ u_6 \\ u_7 \\ 0 \\ 0 \\ u_{10} \\ u_{11} \\ 0 \end{pmatrix}$$

• **Berechnung von $R_0^H \cdot v = s$**

Auch das Gleichungssystem mit der hermiteschen Matrix R_0^H kann auf ein kleineres System reduziert werden. Bei der hermiteschen Matrix werden die inneren Punkte des Gebietes und die Werte am Neumann-Rand unabhängig von den Werten des Dirichlet-Randes identifiziert.

$$R_0^H = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 & -0.04 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & -0.04 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & -0.09 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0.26 & -0.09 & 0 & 0 & -0.08 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & -0.09 & 0.26 & 0 & 0 & 0 & -0.08 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & -0.09 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & -0.09 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & -0.04 & 0 & 0 & 0 & 0.26 & -0.09 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & -0.04 & 0 & 0 & -0.09 & 0.26 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & -0.09 & 1 \end{pmatrix}$$

Somit kann zur Bestimmung der diskreten Werte v_i in Ω und am Neumann-Rand die verkleinerte Matrix

$$(R_0^H)_{opt} = \begin{pmatrix} 0.26 & -0.09 & -0.08 & 0 \\ -0.09 & 0.26 & 0 & -0.08 \\ -0.04 & 0 & 0.26 & -0.09 \\ 0 & -0.04 & -0.09 & 0.26 \end{pmatrix} \quad (5.4)$$

verwendet werden. Die Werte am Dirichlet-Rand werden wiederum auf 0 gesetzt. Dieses Ergebnis der optimierten Lösung stimmt jedoch nicht exakt mit der Lösung des großen Problems $R_0^H \cdot v = s$ überein, da diese Werte ungleich 0 am Dirichlet-Rand ermittelt. Bei einer rechten Seite

$$s = (0, 0, 0, 0, 0, 6, 0, 0, 0, 0, -6, 0)^T$$

werden die Lösungen

$$v = \begin{pmatrix} 0 \\ 0.8528 \\ -0.0123 \\ 0 \\ 1.9187 \\ 21.3189 \\ -0.3072 \\ -0.0277 \\ -0.4831 \\ -5.3679 \\ -24.9823 \\ -2.2484 \end{pmatrix} \quad \text{bzw.} \quad v_{in} = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 21.3189 \\ -0.3072 \\ 0 \\ 0 \\ -5.3679 \\ -24.9823 \\ 0 \end{pmatrix}$$

berechnet. Hierbei kann man erkennen, dass in den inneren Punkten von Ω und am Neumann-Rand die gleichen Werte identifiziert werden. Es unterscheiden sich lediglich die Ergebnisse am Dirichlet-Rand. Hier greift jedoch die zweite Bedingung ein, die es verbietet, Werte des Dirichlet-Randes zu 'messen'. Da nur die Ergebnisse innerhalb Ω und am Neumann-Rand von Bedeutung sind, spielt es keine Rolle, ob nun das große System

$$R_0^H \cdot v = s$$

oder das reduzierte

$$(R_0^H)_{opt} \cdot v_{in} = s_{in}$$

gelöst wird.

Bemerkung:

Für die Berechnung der verkleinerten Matrizen gilt:

$$(R_{opt})^H = (R^H)_{opt} \tag{5.5}$$

5.2.2 Erzeugung einer symmetrischen Matrix

- Berechnung von $R_0 \cdot u = f$

Die Umwandlung der großen Matrix R_0 in die symmetrische Matrix R_{sym} baut auf die zuvor erklärte Verkleinerung des Systems auf. Die ursprüngliche Matrix R_0 ist eine $n \times n$ -Matrix ($n = n_x \cdot n_y$), die auf eine kleinere Matrix R_{opt} reduziert wird. Bei dieser Matrix ist die Symmetrie fast gegeben. Lediglich die Diskretisierung der Differentialgleichung an den $n_x - 2$ Punkten des Neumann-Randes verfälschen (zunächst) diese Bedingung (vgl. (5.3)). Werden nun die letzten $n_x - 2$ Zeilen von R_{opt} mit dem Faktor $\frac{1}{2}$ multipliziert, ist die Symmetrie gegeben.

$$\begin{aligned} R_{sym} &= \frac{1}{2} \cdot \begin{pmatrix} 0.26 & -0.09 & -0.04 & 0 \\ -0.09 & 0.26 & 0 & -0.04 \\ -0.08 & 0 & 0.26 & -0.09 \\ 0 & -0.08 & -0.09 & 0.26 \end{pmatrix} \\ &= \begin{pmatrix} 0.26 & -0.09 & -0.04 & 0 \\ -0.09 & 0.26 & 0 & -0.04 \\ -0.04 & 0 & 0.13 & -0.045 \\ 0 & -0.04 & -0.045 & 0.13 \end{pmatrix} \end{aligned}$$

Um nun wieder das gleiche Ergebnis u_{in} zu erhalten, muss auch die rechte Seite f_{in} entsprechend multipliziert werden.

$$f_{sym} := \begin{pmatrix} f_6 \\ f_7 \\ \frac{1}{2} \cdot f_{10} \\ \frac{1}{2} \cdot f_{11} \end{pmatrix}$$

Nun ist die Lösung der Gleichung

$$R_{sym} \cdot u_{in} = f_{sym}$$

gleich der Lösung von

$$R_{opt} \cdot u_{in} = f_{in}.$$

• **Berechnung von $R_0^H \cdot v = s$**

Multipliziert man die letzten $n_x - 2$ Spalten der Matrix $(R_0^H)_{opt}$ (vgl. Gleichung (5.4)) mit dem Faktor $\frac{1}{2}$, so erhält man die symmetrische Matrix

$$(R^H)_{sym} = \begin{pmatrix} 0.26 & -0.09 & -0.04 & 0 \\ -0.09 & 0.26 & 0 & -0.04 \\ -0.04 & 0 & 0.13 & -0.045 \\ 0 & -0.04 & -0.045 & 0.13 \end{pmatrix}$$

Bei diesem Gleichungssystem wird die rechte Seite s_{in} nicht verändert. Nach der Ermittlung der Lösung von

$$(R^H)_{sym} \cdot \hat{v}_{in} = s_{in}$$

müssen die letzten $n_x - 2$ Elemente von $\hat{v}_{in}$ mit dem Faktor $\frac{1}{2}$ multipliziert werden, um wieder die ursprüngliche Lösung von

$$(R^H)_{opt} \cdot v_{in} = s_{in}$$

zu erhalten.

Bemerkung:

Die Berechnung des hermiteschen Gleichungssystems kann bei der symmetrischen Matrix auch mit Hilfe der ursprünglichen Matrix durchgeführt werden. Da

$$\overline{R_{sym}} = R_{sym}^H \quad (5.6)$$

gilt, lässt sich für das hermitesche Problem folgendes schreiben:

$$\begin{aligned} R_{sym}^H \cdot \hat{v} &= \overline{R_{sym}} \cdot \hat{v} = s \\ \Leftrightarrow \overline{R_{sym}} \cdot \hat{v} &= \bar{s} \\ \Leftrightarrow R_{sym} \cdot \bar{\hat{v}} &= \bar{s} \end{aligned}$$

Somit muss nur einmal die Matrix R_{sym} z.B. mit Hilfe der Cholesky-Zerlegung gesplittet werden und kann anschließend ebenfalls das hermitesche Problem lösen.

Nutzt man zur Lösung des linearen Gleichungssystems die klassische LU-Zerlegung, wird ein Aufwand der Ordnung $\frac{2}{3}n^3$ benötigt. Zwar sind bei einer dünn besetzten Matrix R_0 die Matrizen L und U auch dünn besetzt, sie behalten jedoch nicht die Eigenschaft als Band-Matrizen. Man kann im Vorhinein nicht erkennen, an welcher Position die *fill-in* erzeugt werden. Das Gleichungssystem (5.1) kann nun über die Matrizen L und U berechnet werden:

Durch die Eigenschaft, dass L eine untere und U eine obere Dreiecksmatrix ist, lässt sich die Lösung sehr schnell in zwei Schritten berechnen

Wurde die LU-Zerlegung für die Matrix R_0 durchgeführt, kann mit diesen beiden Matrizen die Lösung der Gleichung (5.2) ermittelt werden.

Nun ist U^H eine untere und L^H eine obere Dreiecksmatrix, so dass das Problem durch

gelöst werden kann.

Für das mitgeführte Beispiel lässt sich R_0 durch das Produkt aus

[illegible]

und

$$U = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0.26 & -0.09 & 0 & 0 & -0.04 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0.23 & -0.09 & 0 & -0.01 & -0.04 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0.25 & -0.09 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0.21 & -0.09 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{pmatrix}$$

berechnen.

5.4 Die Cholesky-Zerlegung

Um eine Cholesky-Zerlegung durchführen zu können, muss die Matrix des Gleichungssystems symmetrisch positiv definit sein. Demnach können die im Kapitel 5.2.2 beschriebenen Matrizen R_{sym} mit Hilfe einer unteren Dreiecksmatrix dargestellt werden

$$R_{sym} = L \cdot L^T$$

Zur Berechnung dieser Dreiecksmatrix wird ein Aufwand von $\frac{1}{3}n^3$ benötigt. Das Gleichungssystem

$$R_{sym} \cdot u = L \cdot L^T \cdot u = f$$

kann wiederum in zwei Schritten berechnet werden

$$\begin{aligned} L \cdot x &= f \\ L^T \cdot u &= x \end{aligned}$$

Da die Lösung des hermiteschen Problems auf die Lösung der Gleichung mit der ursprünglichen symmetrischen Matrix R_{sym} zurückgeführt werden kann (vgl. Kapitel 5.2.2), ist kein weiterer Aufwand für

$$(R_0^H)_{sym} \cdot \hat{v}_{in} = f_{in}$$

nötig.

Für das Beispiel gilt:

$$R_{sym} = L \cdot L^T$$

mit

$$L = \begin{pmatrix} 0.5099 & 0 & 0 & 0 \\ -0.1765 & 0.4784 & 0 & 0 \\ -0.0784 & -0.0289 & 0.3507 & 0 \\ 0 & -0.0836 & -0.1352 & 0.3236 \end{pmatrix}$$

5.5 Lösen mit Hilfe des ADI-Verfahrens

Das ADI-Verfahren wurde ursprünglich zum Lösen instationärer Differentialgleichungen entwickelt. Die Lösung u der Gleichung (5.1) ist jedoch eine stationäre Lösung eines elliptischen Randwertproblems. Aus diesem Grund wird diese Gleichung in eine *zeitabhängige* Variante umgeschrieben, die von einem Startpunkt u_{t_0} iterativ exaktere Lösungen u_{t_i} ermittelt. Für $t \rightarrow \infty$ soll das Ergebnis die stationäre Lösung u sein.

Um aus dem Gleichungssystem ein iteratives Verfahren zu erhalten, wird die Matrix R_{sym} durch die Summe zweier Matrizen R_x und R_y dargestellt.

$$R_{sym} = R_x + R_y \quad (5.7)$$

Hierbei repräsentieren die beiden Matrizen die Approximation des Randwertproblems in x- bzw. in y-Richtung, so dass diese (durch die Wahl der zentralen finiten Differenzen) jeweils nur drei Bänder beinhalten. Die Bestimmung der Matrizen R_x und R_y wird in dem Kapitel 5.5.1 erläutert. Mit (5.7) kann man Gleichung (5.1) in eine der beiden (äquivalenten) Formen

$$(R_x + t \cdot I) \cdot u + (R_y - t \cdot I) \cdot u = f \quad (5.8)$$

$$(R_x - t \cdot I) \cdot u + (R_y + t \cdot I) \cdot u = f \quad (5.9)$$

umwandeln. In der *Peaceman-Rachford Methode* [2] lässt sich nun der Vektor $u^{(n+1)}$ in zwei Schritten aus dem Vektor $u^{(n)}$ berechnen.

$$(R_x + t \cdot I) \cdot u^{(n+\frac{1}{2})} + (R_y - t \cdot I) \cdot u^{(n)} = f \quad (5.10)$$

$$(R_x - t \cdot I) \cdot u^{(n+\frac{1}{2})} + (R_y + t \cdot I) \cdot u^{(n+1)} = f \quad (5.11)$$

Da nun anstatt einem Gleichungssystem mit einer 5-Band Matrix mehrere Gleichungssysteme mit 3-Band Matrizen gelöst werden müssen, wird eine schnellere Ermittlung der Lösung u erhofft. Durch die Wahl des Relaxationsparameters t wird die Schrittweite des instationären Problems

gesteuert. Eine optimale Wahl von t wird in Kapitel 5.5.2 beschrieben.

Als Abbruchkriterium der Iteration wird die neue Approximation $u^{(n+1)}$ mit dem vorigen Wert $u^{(n)}$ verglichen. Unterschreitet der Quotient aus dem Maximum von $|u^{(n+1)} - u^{(n)}|$ und dem Maximum von $|u^{(n+1)}|$ eine Schranke ϵ , wird die Iteration beendet. Da bei den anderen Verfahren Genauigkeiten von $\epsilon = 10^{-15}$ erreicht werden, wird diese Schranke auch für das ADI-Verfahren übernommen.

wähle:

$$R_{sym} = R_x + R_y$$

$$u^{(n+1)} = 0$$

do:

$$u^{(n)} = u^{(n+1)}$$

$$(R_x + t \cdot I) \cdot u^{(n+\frac{1}{2})} = f - (R_y - t \cdot I) \cdot u^{(n)}$$

$$(R_y + t \cdot I) \cdot u^{(n+1)} = f - (R_x - t \cdot I) \cdot u^{(n+\frac{1}{2})}$$

while $\left(\frac{\max |u^{(n+1)} - u^{(n)}|}{\max |u^{(n+1)}|} > 10^{-15} \right)$

Abbildung 5.1: ADI-Verfahren

Bemerkung:

Für die Anwendung des ADI-Verfahrens müssen die Matrizen R_x bzw. R_y eine Bedingung erfüllen. Sie müssen die Eigenschaften einer *Stieltjes-Matrix* besitzen, d.h. dass sie symmetrisch sind und auf den Diagonalen (und nur dort) Werte größer 0 stehen. Aus diesem Grund wurde auch die Matrix R_{sym} als Ausgangspunkt des ADI-Verfahrens genommen, da deren Aufspaltung in x- und y-Richtung diese Bedingung erfüllt.

5.5.1 Erzeugung der Richtungs-Matrizen

Die Matrix R_{sym} ist eine 5-Band Matrix, da diese mit Hilfe der zentralen finiten Differenzen in x- bzw. in y-Richtung ermittelt wurde. Diese soll durch die Summe zweier Matrizen dargestellt werden.

$$R_{sym} = R_x + R_y$$

Da die Nummerierung des Gitters entlang der x-Achse verläuft (vgl. Kapitel 3.2), sind die Bänder links und rechts neben der Diagonalen nur von der Approximation des Randwertproblems in x-Richtung abhängig. Die zwei Bänder, die sich weiter außen befinden, machen die Approximation in y-Richtung aus. Lediglich auf der Diagonalen von R_{sym} befinden sich gemischte Terme. Diese werden (in dieser Arbeit) je nach Auflösung ($n_x \times n_y$) gewichtet. Ist die Anzahl der Einteilungen in

x-Richtung größer als die in y-Richtung, so erhält die Diagonale von R_x eine stärkere Gewichtung als die von R_y . Die Gewichte lassen sich berechnen durch:

$$g_x = \frac{(n_x - 1)^2}{(n_x - 1)^2 + (n_y - 1)^2}$$

$$g_y = \frac{(n_y - 1)^2}{(n_x - 1)^2 + (n_y - 1)^2}$$

Somit gilt für die Aufteilung der Beispiel-Matrix:

$$R_{sym} = \begin{pmatrix} 0.26 & -0.09 & -0.04 & 0 \\ -0.09 & 0.26 & 0 & -0.04 \\ -0.04 & 0 & 0.13 & -0.045 \\ 0 & -0.04 & -0.045 & 0.13 \end{pmatrix}$$

$$= \begin{pmatrix} 0.18 & -0.09 & 0 & 0 \\ -0.09 & 0.18 & 0 & 0 \\ 0 & 0 & 0.09 & -0.045 \\ 0 & 0 & -0.045 & 0.09 \end{pmatrix} + \begin{pmatrix} 0.08 & 0 & -0.04 & 0 \\ 0 & 0.08 & 0 & -0.04 \\ -0.04 & 0 & 0.04 & 0 \\ 0 & -0.04 & 0 & 0.04 \end{pmatrix}$$

5.5.2 Steuerung der Schrittweite

Das System

$$R \cdot u = f$$

wird in zwei Systeme der Form

$$\begin{aligned} (R_x + t \cdot I) \cdot u^{(n+\frac{1}{2})} &= f - (R_y - t \cdot I) \cdot u^{(n)} \\ (R_y + t \cdot I) \cdot u^{(n+1)} &= f - (R_x - t \cdot I) \cdot u^{(n+\frac{1}{2})} \end{aligned}$$

überführt, wobei R_x und R_y die *Stieltjes-Bedingung* erfüllen.

Zur Analyse der Schrittweitensteuerung durch t wird das System zunächst stark vereinfacht. Ist R eine skalare Größe, kann man die Lösung zwar direkt bestimmen, sie lässt sich jedoch auch durch das ADI-Verfahren ermitteln. Es gilt:

$$u^{(n+\frac{1}{2})} = \frac{f - (R_y - t) \cdot u^{(n)}}{R_x + t}$$

$$u^{(n+1)} = \frac{f - (R_x - t) \cdot u^{(n+\frac{1}{2})}}{R_y + t}$$

Somit kann man $u^{(n+1)}$ direkt aus $u^{(n)}$ ermitteln mit

$$u^{(n+1)} = \underbrace{\frac{(t - R_x) \cdot (t - R_y)}{(t + R_x) \cdot (t + R_y)}}_{:=V} \cdot u^{(n)} + \underbrace{\frac{2 \cdot t}{(t + R_y) \cdot (t + R_x)}}_{:=\tilde{f}} \cdot f$$

Das Verfahren konvergiert, falls der skalare Wert V zwischen -1 und 1 liegt. Dies ist aber für jedes $t > 0$ erfüllt, da

$$\begin{aligned} \left| \frac{t - R_x}{t + R_x} \right| &< 1 \quad \forall t > 0 \\ \text{und} \quad \left| \frac{t - R_y}{t + R_y} \right| &< 1 \quad \forall t > 0. \end{aligned}$$

Ziel ist es nun, den Relaxationsparameter t so zu wählen, dass $|V|$ möglichst klein wird. Im skalaren Beispiel nimmt V mit $t = R_x$ oder $t = R_y$ den Wert 0 an. Somit wird die iterative Lösung nach nur einem Schritt ermittelt.

$$u^{(n+1)} = 0 \cdot u^{(n)} + \hat{f} = \frac{1}{R_x + R_y} \cdot f = \frac{f}{R} \quad (5.12)$$

Für mehrdimensionale Matrizen R ist die Argumentation ähnlich. Die Konvergenz hängt nun nicht mehr von zwei skalaren Größen R_x und R_y ab, sondern von den Eigenwerten der Matrizen R_x und R_y . Kennt man diese, so kann man t in jeder Iteration so wählen, dass ein Eigenwert und somit auch der zugehörige Eigenvektor des Gleichungssystems wegfällt. Demnach ist es beim ADI-Verfahren sinnvoll, in jedem Iterationsschritt einen anderen Wert für t zu nehmen, so dass jeder Eigenwert berücksichtigt werden kann.

Da es jedoch mit hohen Kosten verbunden ist, Eigenwerte einer Matrix zu ermitteln, versucht man, je nach Randwertproblem, die Eigenwerte zu schätzen. In der hier vorliegenden Problemstellung ist es jedoch nicht ohne hohen Aufwand möglich, gute Abschätzungen für die Eigenwerte zu bestimmen.

Das ADI-Verfahren wird oft anhand einfacher Randwertprobleme erläutert, wie z.B.

$$\begin{aligned} -u'' &= f && \text{in } \Omega \\ u &= 0 && \text{auf } \Gamma \end{aligned}$$

Für dieses Randwertproblem lassen sich explizit die Eigenfunktionen und somit auch die Eigenwerte für die Operatoren in x- bzw. in y-Richtung bestimmen. Somit hat man auch die Eigenwerte der Matrizen berechnet und kann ein optimales t bestimmen.

Das hier vorliegende Randwertproblem hat (vereinfacht) die Form

$$\begin{aligned} -(\sigma \cdot u')' &= f && \text{in } \Omega \\ u &= 0 && \text{auf } \Gamma_d \\ \partial_\nu u &= 0 && \text{auf } \Gamma_n \end{aligned}$$

Eine Bestimmung der Eigenfunktion ist nun nicht so einfach, da die Funktion $\sigma(x, y)$ in Ω verschiedene Werte annehmen und daher nicht abgeschätzt werden kann.

Somit besteht nur noch die Möglichkeit, ein optimales t mittels experimentellen Analysen zu ermitteln. Die Ergebnisse dieser Untersuchungen werden im Kapitel 6.1 vorgestellt.

5.6 Die UMFPACK Bibliothek

UMFPACK ist eine Bibliothek (geschrieben in der Programmiersprache C), die für das Lösen großer Gleichungssysteme mit dünnbesetzten Matrizen entwickelt wurde. Intern bearbeitet ein *Multifrontal-Solver* das Problem. Dies bedeutet, dass die Matrix zunächst einer *symbolischen Analyse* unterzogen wird, so dass eine Permutationsmatrix erzeugt werden kann, mit deren Hilfe möglichst wenig *fill-in* in den Zerlegungsmatrizen L und U erzeugt werden. Anschließend wird die permutierte Matrix in Teilmatrizen aufgeteilt, für die dann jeweils die LU-Zerlegung durchgeführt wird. Diese LU-Zerlegung der Untermatrizen kann parallel laufen, da sich die einzelnen Eliminationsschritte nicht beeinflussen. Die Teilergebnisse werden darauf hin wieder zusammengefasst und das Gleichungssystem mittels Vorwärts- und Rückwärts-Substitution gelöst.

Bei den in dieser Arbeit vorkommenden Gleichungssystemen treten sehr häufig die gleichen Matrizen auf. Diese müssen nur einmal *symbolisch analysiert* werden und können dann beliebig oft die Gleichungssysteme mit Hilfe der *numerischen Analyse* lösen. Da man dem Löser sogar mitteilen kann, ob das ursprüngliche oder das hermitsche Problem gelöst werden soll, muss für beide Gleichungssysteme nur eine *symbolische Analyse* durchgeführt werden.

Kapitel 6

Numerische Ergebnisse

6.1 Steuerung des ADI-Verfahrens

Dieses Unterkapitel befasst sich mit der Bestimmung des Relaxationsparameters t des ADI-Verfahrens. Es wird anhand von Experimenten untersucht, wie stark dieser Parameter die Konvergenz beeinflussen kann. Des Weiteren wird analysiert, inwiefern externe Faktoren (Frequenz, Cole-Cole Parameter) das Verfahren und somit die Bestimmung von t beeinflussen.

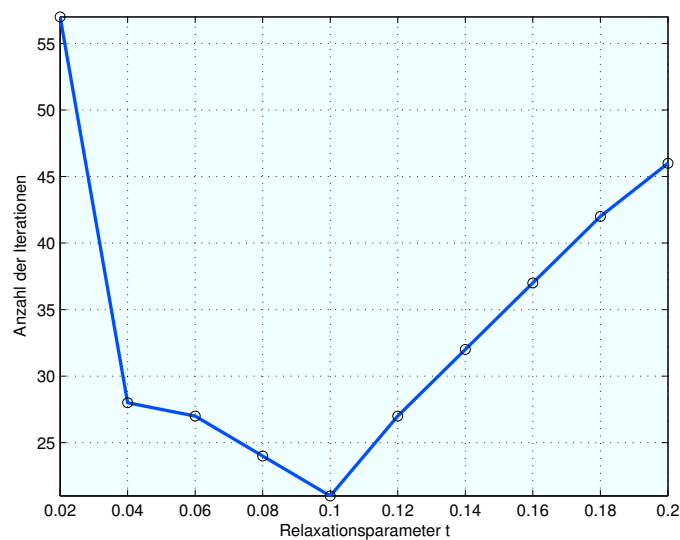
6.1.1 Konvergenzgeschwindigkeiten in Abhängigkeit von t

In diesen Testbeispielen wird gezeigt, dass die Steuerung des Relaxationsparameters t einen starken Einfluss auf die Konvergenzgeschwindigkeit des ADI-Verfahrens hat. Hierbei wird in allen Durchführungen eine Frequenz von 0 Hz und ein Gleichstromwiderstand von 1 Ω gesetzt. Durch die Veränderung der Auflösung des Problems werden verschieden große Gleichungssysteme aufgestellt, zu deren Lösung jeweils ein anderer Relaxationsparameter t verwendet wird. Das optimale t wird an der Anzahl der Iteration gemessen, die das ADI-Verfahren durchlaufen muss.

In den Versuchsdurchführungen wird daher bei verschiedenen Relaxationsparametern die Anzahl der Iterationen dokumentiert und mit Hilfe einer Graphik veranschaulicht.

- 4×3 Auflösung

t	# Iterationen
0.02	57
0.04	28
0.06	27
0.08	24
0.10	21
0.12	27
0.14	32
0.16	37
0.18	42
0.20	46

Tabelle 6.1: Konvergenzgeschwindigkeit bei 4×3 AuflösungAbbildung 6.1: Konvergenzgeschwindigkeit bei 4×3 Auflösung

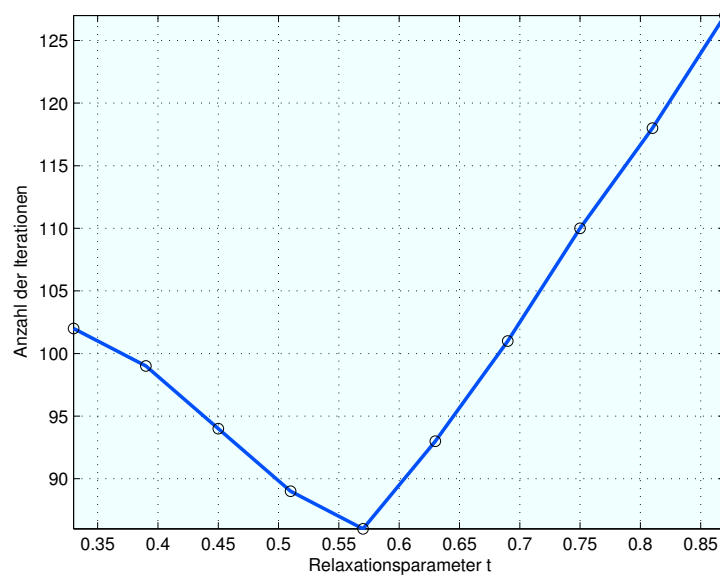
Ergebnis:

Bei einer geringen Auflösung hat der Relaxationsparameter t einen großen Einfluss auf die Konvergenzgeschwindigkeit. Wählt man den Wert zu klein, werden viele unnötige Iterationen durchlaufen, um ein akzeptables Ergebnis zu erhalten. Die Kurve fällt bis zur Schrittweite $t = 0.1$ ab, bei der bei nur noch 21 Iterationen durchgeführt werden müssen. Entfernt man sich wieder von diesem optimalen t , steigt die Anzahl der Iterationen.

Will man nur einen groben Wert von t_{opt} bestimmen, bewirken Abweichungen von 20% zum optimalen Parameter zwar mehr Iterationsschritte, diese Anzahl liegt dennoch in einer Toleranzgrenze von 20%.

- 12×8 Auflösung

t	# Iterationen
0.33	102
0.39	99
0.45	94
0.51	89
0.57	86
0.63	93
0.69	101
0.75	110
0.81	118
0.87	127

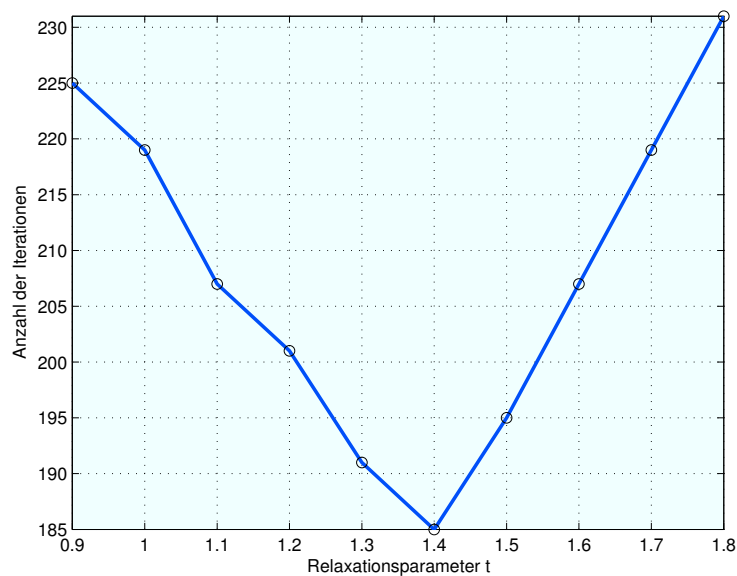
Tabelle 6.2: Konvergenzgeschwindigkeit bei 12×8 AuflösungAbbildung 6.2: Konvergenzgeschwindigkeit bei 12×8 Auflösung**Ergebnis:**

Auch bei einer feineren Auflösung gibt es einen Relaxationsparameter, der eine minimale Anzahl an Iterationen erzielt. Je weiter man sich von diesem Wert entfernt, desto rechenintensiver ist die Bestimmung der Lösung des ADI-Verfahrens.

Des Weiteren kann man erkennen, dass das optimale t bei dieser Auflösung größer ist als im Beispiel der gröberen Auflösung.

- 25×25 Auflösung

t	# Iterationen
0.90	225
1.00	219
1.10	207
1.20	201
1.30	191
1.40	185
1.50	195
1.60	207
1.70	219
1.80	231

Tabelle 6.3: Konvergenzgeschwindigkeit bei 25×25 AuflösungAbbildung 6.3: Konvergenzgeschwindigkeit bei 25×25 Auflösung

Ergebnis:

Bei einer sehr feinen Auflösung steigt die Anzahl der Unbekannten und somit auch die Größe der Matrix des Gleichungssystems. Die Kurve besitzt gleiche Charakteristika wie bei den zuvorigen Beispielen. Mit einem Relaxationsparameter von 1.40 wird das Ergebnis nach 185 Iterationen berechnet.

6.1.2 Variation der Leitfähigkeit

Bei dem inversen Problem werden verschiedene Cole-Cole Verteilungen miteinander verglichen. Diese Parameter und die Frequenz bestimmen die Leitfähigkeit. Demnach ist es von Bedeutung, ob die Variation der Leitfähigkeit die Steuerung des ADI-Verfahrens beeinflusst. Diese Änderungen werden bei einer 4×3 Auflösung durch eine konstante Cole-Cole Verteilung bei verschiedenen Frequenzen erreicht.

Frequenz	t_{opt}	# Iterationen
10^{-4}	0.3728	55
10^{-3}	0.3788	55
10^{-2}	0.3797	55
10^{-1}	0.3875	55
10^0	0.3916	55
10^1	0.3916	56
10^2	0.4070	55
10^3	0.4107	56
10^4	0.4129	56

Tabelle 6.4: Variation der Leitfähigkeit

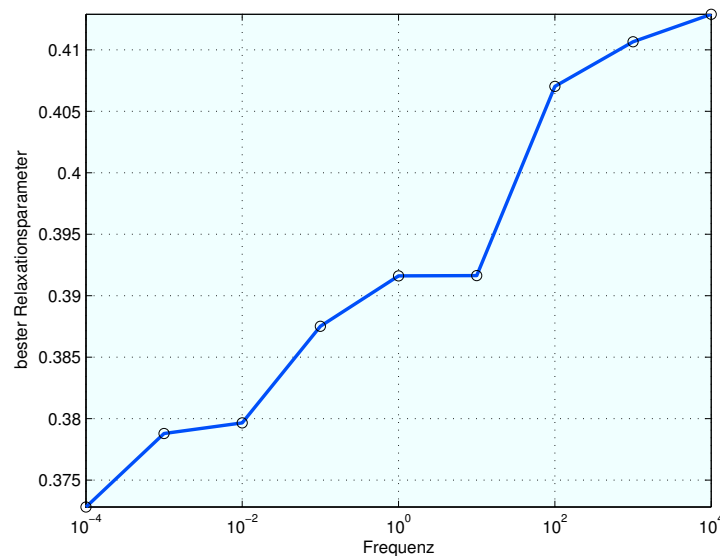


Abbildung 6.4: Variation der Leitfähigkeit

Ergebnis:

Bei der Wahl des optimalen Relaxationswertes darf die Leitfähigkeit nicht vernachlässigt werden, da bei verschiedenen Leitfähigkeiten verschiedene beste Schrittweiten t ermittelt werden. Dies zeigt, dass die Leitfähigkeit die Eigenwerte der Matrix beeinflusst und deren Approximation somit komplizierter wird.

6.1.3 Optimale Steuerung

Bei den folgenden Untersuchungen wird analysiert, wie das optimale t und die Anzahl der Iterationen von der Auflösung und somit von der Anzahl der Unbekannten des Gleichungssystems abhängen.

Um diese Werte ermitteln zu können, wurde das Optimum mit Hilfe einer Intervall-Schachtelung bestimmt. Es wurde ein Intervall vorgegeben, in dem sich der beste Relaxationsparameter befindet. Dieses wurde anschließend gedrittelt und das ADI-Verfahren mit den beiden inneren Relaxationsparametern gestartet.

Der Wert mit der kleineren Anzahl an Iterationen wird als Mittelpunkt des neuen Intervalls definiert, das durch den rechten und linken Nachbarn des vorigen Intervalls begrenzt wird. Wird keine Verbesserung erreicht, ist der beste Wert ermittelt.

Auflösung		bester Relaxationspar.	Anzahl Iterationen
n_x	n_y		
4	3	0.0953	21
4	4	0.1037	21
5	4	0.1502	28
5	5	0.1579	28
6	6	0.2226	40
10	7	0.4469	71
12	8	0.5417	86
15	10	0.7051	104
16	16	0.8318	118
18	17	0.9491	130
20	18	1.0427	143
20	20	1.0541	145
22	22	1.1933	165
25	25	1.3439	185
30	25	1.5360	208
30	30	1.6668	228
35	40	1.9959	258

Tabelle 6.5: optimaler Relaxationsparameter in Abhängigkeit der Auflösung

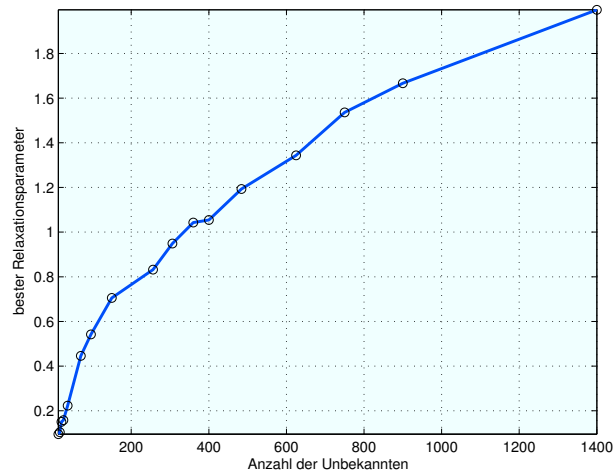


Abbildung 6.5: optimaler Relaxationsparameter in Abhängigkeit der Auflösung

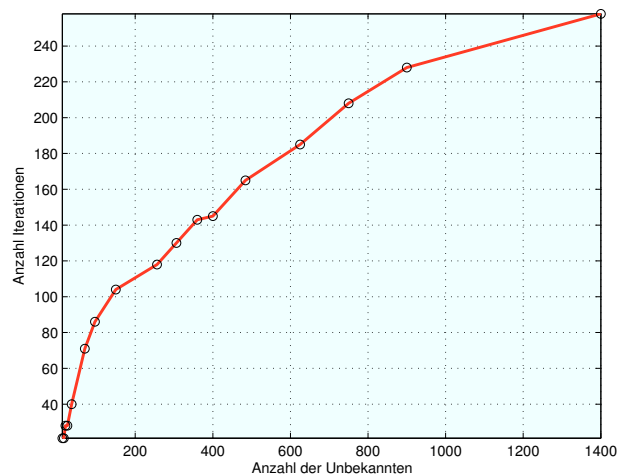


Abbildung 6.6: Anzahl der Iterationen bei optimaler Steuerung

Ergebnis:

Je größer die Anzahl der Unbekannten wird, desto größer werden auch die optimale Schrittweite und die Anzahl der Iterationen. Die ermittelten Ergebnisse hängen auch von der Vorgabe des Intervalls ab. Wird ein gröberes Intervall angegeben, werden zwar die gleiche Anzahl an minimalen Iterationen ausgegeben, der Relaxationsparameter kann sich jedoch ein wenig unterscheiden. So gibt es zum Beispiel bei einer Auflösung von 4×3 ein ganzes Intervall für t , in dem 21 Iterationen erreicht werden.

6.1.4 Unterschiedliche Steuerung bei gleicher Dimension

Graphik 6.5 lässt die Vermutung aufkommen, dass die Steuerung des Relaxationswertes t mit Hilfe der Wurzelfunktion abgeschätzt werden kann.

$$t_{opt} \approx a \cdot \sqrt{\text{Anzahl der Unbekannten}}$$

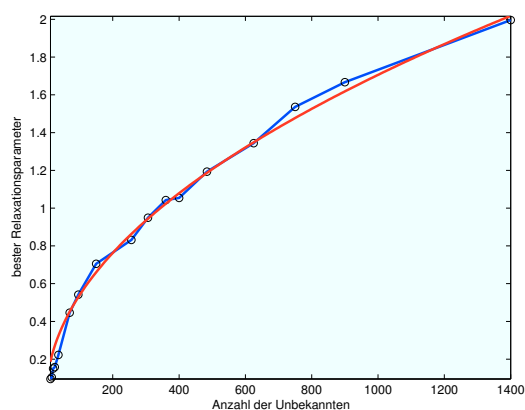


Abbildung 6.7: Approximation der Kurve durch Wurzelfunktion mit $a = 0.051$

Hierbei ist diese Aussage jedoch sehr kritisch zu betrachten, da die Auflösungen n_x und n_y immer im gleichen Verhältnis vergrößert wurden (vgl. Tabelle 6.5).

Bei dem folgenden Versuch ist daher die Anzahl der Unbekannten konstant 240. Es werden diejenigen Auflösungen untersucht, für die gilt: $n_x \cdot n_y = 240$.

Auflösung		bester Relaxationspar.	Anzahl Iterationen
n_x	n_y		
4	60	0.1140	38
5	48	0.1762	54
6	40	0.2149	63
8	30	0.3270	75
10	24	0.4685	81
12	20	0.6183	87
15	16	0.7447	106
16	15	0.8403	117
20	12	1.0256	144
24	10	1.1943	173

Tabelle 6.6: optimaler Relaxationsparameter bei gleicher Dimension

Ergebnis:

Der optimale Relaxationsparameter hängt nicht ausschließlich von der Anzahl der Unbekannten ab. Die Auflösung spielt dabei eine sehr große Rolle.

6.1.5 Aussagen über die Abschätzung des Relaxationsparameters

Es ist nicht möglich, aus der Auflösung und der Leitfähigkeit ein optimales t abzuschätzen. Aus der Versuchsreihe des Kapitels 6.1.1 erkennt man, dass kleine Abweichungen (20%) zu dem optimalen Relaxationsparameter immer noch akzeptable Ergebnisse liefern, da die Anzahl der Iterationen auch um ca. 20% steigen.

Das Problem liegt jedoch darin, dass man dieses Intervall nicht abschätzen kann. Verschiedene Leitfähigkeiten und im Besonderen die Auflösung des Gitters beeinflussen das t_{opt} sehr stark.

6.2 Vergleich der Löser

In diesem Kapitel werden die einzelnen Löser miteinander verglichen. Es soll die gesamte Ableitungsmatrix und deren hermitesche Matrix berechnet werden. Das Testbeispiel beinhaltet drei verschiedene Frequenzen bei einer konstanten Cole-Cole Verteilung und zwei unterschiedliche Stromeinspeisungen.

6.2.1 Vergleich des ursprünglichen Systems mit dem optimierten System

Im Folgenden wird getestet, ob die Optimierung der Matrix einen Vorteil erbracht hat. Nur die LU-Zerlegung und der Black Box Löser können sowohl das große als auch das verkleinerte Gleichungssystem lösen.

Um die anderen Verfahren (Cholesky und ADI) anwenden zu können, sind symmetrische Matrizen und somit die Optimierung vorausgesetzt.

- Optimierung bei der LU-Zerlegung

Auflösung		Dauer (sek)	
n_x	n_y	großes GLS	optimiertes GLS
5	5	0.01	0.01
6	6	0.01	0.01
10	7	0.07	0.06
12	8	0.14	0.12
15	10	0.42	0.35
16	16	1.33	1.16
18	17	2.19	1.85
20	18	3.36	2.84
20	20	4.03	3.56
22	22	6.65	5.85
25	25	13.39	11.55
25	30	23.02	20.91
30	30	33.60	30.84
35	40	106.17	94.22

Tabelle 6.7: LU-Zerlegung auf dem optimierten System

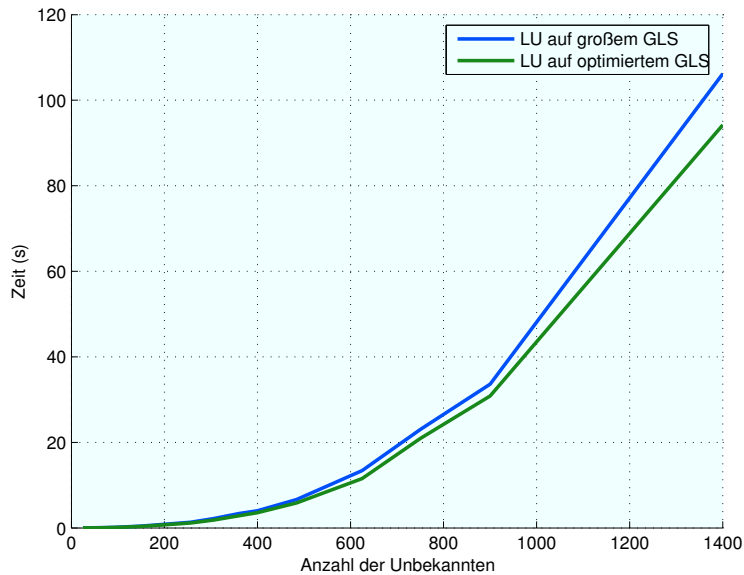


Abbildung 6.8: LU-Zerlegung auf dem optimierten System

- Optimierung bei UMFPACK

Auflösung		Dauer (sek)	
n_x	n_y	großes GLS	optimiertes GLS
5	5	0.01	0.01
6	6	0.01	0.01
10	7	0.04	0.04
12	8	0.08	0.08
15	10	0.20	0.20
16	16	0.55	0.55
18	17	0.79	0.79
20	18	1.09	1.10
20	20	1.35	1.36
22	22	1.98	1.99
25	25	3.35	3.36
25	30	4.79	4.82
30	30	6.95	6.95
35	40	16.68	16.95

Tabelle 6.8: Optimierung mit UMFPACK

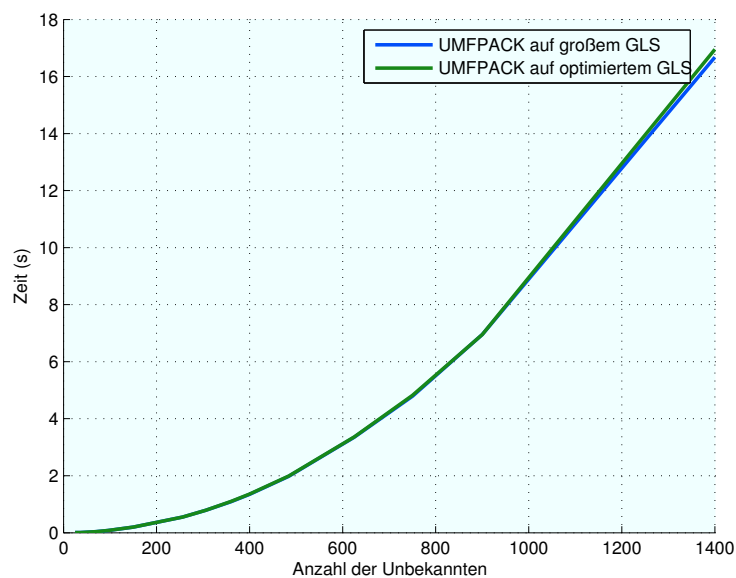


Abbildung 6.9: Optimierung mit UMFPACK

Ergebnis:

Eine Optimierung liefert für die LU-Zerlegung eine schnellere Lösung. Durch die Verkleinerung der Matrix müssen auch weniger Schritte der LU-Zerlegung durchgeführt werden.

Bei dem *UMFPACK* Löser ist keine Verbesserung der Geschwindigkeit zu erkennen. Dies liegt

u.a. daran, dass auch UMFPACK die ursprüngliche Matrix optimiert und dass das Umwandeln des eigenen Matrix-Formates in das UMFPACK-Format viel Zeit in Anspruch nimmt. Das C++ Programm speichert die Matrizen in einer anderen Form als UMFPACK. Somit müssen die Daten für jede Berechnung so umformatiert werden, dass UMFPACK diese Gleichungssysteme lösen kann.

6.2.2 Vergleich aller optimierten Löser

In diesem Test sollen alle Löser das gleiche Problem bewältigen. Hierbei werden ausschließlich die optimierten Varianten angewandt.

Auflösung		Dauer (sek)			
n_x	n_y	LU-Zerlegung	ADI-Verfahren	Cholesky-Zerlegung	UMFPACK
5	5	0.01	0.02	0.01	0.01
6	6	0.01	0.03	0.01	0.01
10	7	0.06	0.14	0.04	0.04
12	8	0.12	0.24	0.08	0.08
15	10	0.35	0.51	0.20	0.20
16	16	1.16	1.22	0.56	0.55
18	17	1.85	1.67	0.83	0.79
20	18	2.84	2.26	1.17	1.10
20	20	3.56	2.63	1.45	1.36
22	22	5.85	3.83	2.19	1.99
25	25	11.55	6.12	3.75	3.36
25	30	20.91	8.77	5.62	4.82
30	30	30.84	12.26	8.14	6.95
35	40	94.22	26.44	20.67	16.95

Tabelle 6.9: Vergleich aller optimierten Löser

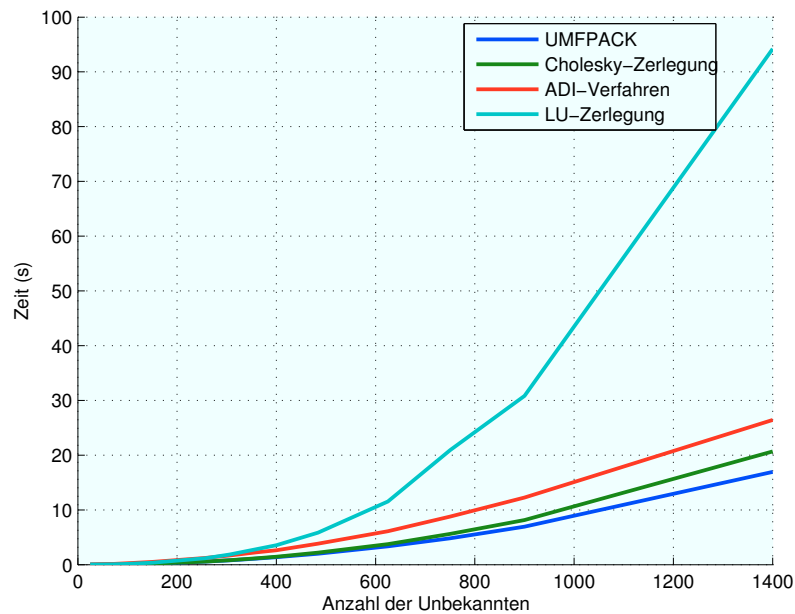


Abbildung 6.10: Vergleich aller optimierten Löser

Ergebnis:

Abbildung 6.10 zeigt, dass die LU-Zerlegung für das gegebene Randwertproblem unbrauchbar ist. Dies liegt an dem Grund, dass zur Berechnung der hermiteschen Ableitungsmatrix auch die hermiteschen L und U benötigt werden. Bei der Cholesky-Zerlegung und dem ADI-Verfahren wird sowohl für das ursprüngliche als auch für das hermitesche Problem nur eine Matrix betrachtet. Demnach ist das Ergebnis wesentlich schneller berechnet.

Die schnellste Art, diese spezielle Gleichungssysteme zu lösen, wird mit Hilfe der UMFPACK-Bibliothek erreicht.

Diese Ergebnisse sind anhand eines zwei-dimensionalen Modells berechnet worden. Wird das Problem auf die dritte Dimension erweitert, kann das ADI-Verfahren von Vorteil sein. Die Matrix des Gleichungssystems ist dann eine 7-Band Matrix, die in drei 3-Band Matrizen aufgeteilt werden kann.

Kapitel 7

Zusammenfassung und Ausblick

Die Aufgabe der geoelektrischen Widerstandstomographie ist in zwei Schritte geteilt, das Lösen des Vorwärts- und des Rückwärtsproblems. In dieser Arbeit wird das direkte Problem (Vorwärtsproblem) zunächst diskretisiert und anschließend gelöst. Die Potentialbildung lässt sich durch eine Differentialgleichung beschreiben, die mit Hilfe der zentralen finiten Differenzen approximiert werden kann. Hat man auch die Randbedingungen in eine diskrete Form gebracht, lässt sich das Vorwärtsproblem mit Hilfe eines Gleichungssystems lösen.

Für das inverse Problem muss sehr oft das Produkt aus der Jacobi-Matrix mit einem Vektor bzw. der Adjungierten der Jacobi-Matrix mit einem Vektor berechnet werden. Auch diese Berechnungen kann man diskret formulieren, so dass die Bestimmung der Produkte wiederum Lösungen von Gleichungssystemen sind.

Diese Gleichungssysteme sind in der Regel sehr groß und besitzen spezielle Charakteristika, die man zur Beschleunigung der Löser ausnutzen kann. Geht man davon aus, dass kein Strom auf dem Rand abgebildet wird und keine Potentiale am Rand gemessen werden, können alle Gleichungssysteme verkleinert und die Matrizen sogar auf eine symmetrische Form gebracht werden. Zur Lösung dieser Gleichungssysteme ist es nicht ratsam, die klassische LU-Zerlegung zu verwenden. Am schnellsten berechnet die UMFPACK-Bibliothek das Problem, gefolgt von der Cholesky-Zerlegung und dem ADI-Verfahren. Hierbei ist es nicht möglich, eine optimale Steuerung des Relaxationsparameters des ADI-Verfahrens ohne großen Aufwand zu ermitteln. Abschätzungen über die Eigenwerte der Matrizen können auf Grund der Komplexität der Differentialgleichungen ohne genauerer Betrachtung nicht durchgeführt werden.

Alle Berechnungen wurden an einem zwei-dimensionalen Modell durchgeführt. Bei dem realen drei-dimensionalen Problem kann das ADI-Verfahren Vorteile haben, da hier noch der Operator in Richtung der dritten Dimension herausgefiltert werden kann, so dass anstatt einer großen 7-Band Matrix nur noch drei 3-Band Matrizen die Lösung des Gleichungssystems bestimmen.

Des Weiteren sollten Überlegungen zur Bestimmung des optimalen Relaxationsparameters des ADI-Verfahrens gemacht werden. Die Leitfähigkeit beeinflusst stark die Eigenwerte der Matrizen.

Abschätzungen oder Vereinfachungen der Differentialgleichung könnten zu einer besseren Bestimmung des Relaxationswertes führen.

Bei sehr großen Gleichungssystemen gewinnen die *Multigrid-Löser* immer mehr an Bedeutung, da deren Konvergenzgeschwindigkeit proportional zu $c \cdot n$ ist. Hierbei ist c jedoch eine sehr große Konstante, so dass diese Verfahren erst bei der Erweiterung auf das drei-dimensionale Problem vorteilhaft sind.

Ein anderer Ansatz wäre es, andere Diskretisierungs-Methoden zu verwenden. Anstelle eines äquidistanten Gitters kann das Gebiet sowohl an den Messsonden als auch an den Elektroden sehr fein und auf dem restlichen Gebiet gröber approximiert werden. Die diskreten Werte können dann neben den *finiten Differenzen* über die *finiten Elemente* oder *finiten Volumen* berechnet werden.

Abbildungsverzeichnis

1.1	Gebiet Ω , Ränder Γ_d und Γ_n , Elektroden	2
2.1	Gebiet Ω mit Normalenvektoren ν_i auf den Rändern	12
3.1	äquidistantes u-Gitter auf Ω , $n_x = 4, n_y = 3$	17
3.2	Aufteilung des diskreten Gradienten	18
3.3	diskretes Divergenz-Gitter	21
3.4	Diskretisierung der rechten Seite des RWP bei zwei δ -Quellen	23
3.5	Umgestaltung des Randwertproblems	24
3.6	diskretes Gitter am Neumann-Rand (vergrößert)	25
3.7	Positionierung der δ -Quellen in Ω	30
3.8	diskretes Potential bei grober Auflösung	30
3.9	diskretes Potential bei feiner Auflösung	31
3.10	Positionierung einer δ -Quelle am Neumann-Rand	31
3.11	diskretes Potential mit δ -Quelle am Neumann-Rand	32
3.12	diskretes Potential bei vier δ -Quellen	32
3.13	diskretes Gitter mit Elektrode (δ -Quelle) und Messsonde	33
3.14	δ -Quelle beliebig in Ω (1 D)	34
3.15	Rechtecke über diskrete Gitterpunkte definieren	34
3.16	Hütchen-Funktion über δ -Quelle (1 D)	35
3.17	δ -Quelle beliebig in Ω (2 D)	36
3.18	Quadrate über diskrete Gitterpunkte definieren	36
3.19	Hütchen-Funktion entlang der x-Kante	37
3.20	Hütchen-Funktion entlang der y-Kante	37
3.21	Beispiel-Anordnung der Messsonden	39
3.22	Potentialverteilungen bei zwei verschiedenen Stromeinspeisungen	41
3.23	Gegenüberstellung beider Potentialverteilungen	42
3.24	grobes c-Gitter auf feinem u-Gitter gelegt	43
4.1	Änderung der Leitfähigkeit lokal in $\tilde{x}$	50

5.1	ADI-Verfahren	67
6.1	Konvergenzgeschwindigkeit bei 4×3 Auflösung	72
6.2	Konvergenzgeschwindigkeit bei 12×8 Auflösung	73
6.3	Konvergenzgeschwindigkeit bei 25×25 Auflösung	74
6.4	Variation der Leitfähigkeit	75
6.5	optimaler Relaxationsparameter in Abhängigkeit der Auflösung	77
6.6	Anzahl der Iterationen bei optimaler Steuerung	77
6.7	Approximation der Kurve durch Wurzelfunktion mit $a = 0.051$	78
6.8	LU-Zerlegung auf dem optimierten System	80
6.9	Optimierung mit UMFPACK	81
6.10	Vergleich aller optimierten Löser	83

Tabellenverzeichnis

2.1	Cole-Cole Parameter verschiedener Böden	8
4.1	Kosten der verschiedenen Varianten bei gegebenen Beispielen	55
6.1	Konvergenzgeschwindigkeit bei 4×3 Auflösung	72
6.2	Konvergenzgeschwindigkeit bei 12×8 Auflösung	73
6.3	Konvergenzgeschwindigkeit bei 25×25 Auflösung	74
6.4	Variation der Leitfähigkeit	75
6.5	optimaler Relaxationsparameter in Abhängigkeit der Auflösung	76
6.6	optimaler Relaxationsparameter bei gleicher Dimension	78
6.7	LU-Zerlegung auf dem optimierten System	80
6.8	Optimierung mit UMFPACK	81
6.9	Vergleich aller optimierten Löser	82

Literaturverzeichnis

- [1] Doss, S.: Dynamic ADI methods for elliptic equations with gradient dependent coefficients, 1978
- [2] Young, D. M.: Iterative Solution Of Large Linear Systems, Orlando, 1971
- [3] Reißel, M.: Vorlesung Spezielle Methoden der angewandten Mathematik, FH Aachen, 2005
- [4] Lustfeld, H.: Vorlesung Moderne Probleme der Geophysik mit Methoden der theoretischen Physik II, Forschungszentrum Jülich, 2004
- [5] Kleefeld, A.: Numerische Untersuchungen zur Bestimmung von Bodenstrukturen durch elektrische Widerstandstomographie (ERT), Diplomarbeit, FH Aachen - Abteilung Jülich - in Zusammenarbeit mit Forschungszentrum Jülich GmbH - Institut für Festkörperforschung, 2005
- [6] Erven, A.: Untersuchungen zur Informationsoptimierung in der geoelektrischen Widerstandstomografie, Diplomarbeit, FH Aachen - Abteilung Jülich - in Zusammenarbeit mit Forschungszentrum Jülich GmbH - Institut für Festkörperforschung, 2005
- [7] Cole, K.S. and Cole, R. H.: Dispersion and Absorption In Dielectrics, 1941
- [8] Reißel, M.: Vorlesungsskript Numerische Mathematik III, FH Aachen, 2004
- [9] Zeidler, E.: Applied functional analysis, Vol. 108 Applications to Mathematical Physics, Vol. 109 Main principles and their applications, Springer-Verlag 1995
- [10] Stoer, J. / Bulirsch, R.: Numerische Mathematik 2, 4. Auflage, Springer-Verlag 2000
- [11] Hockney, R. W. / Eastwood R. W.: Computer Simulation Using Particles, 1999
- [12] <http://de.wikipedia.org>

Danksagungen

Mein besonderer Dank gilt

Herrn Prof. Dr. Reißel für die Übernahme des Hauptreferats und

Herrn Dr. Lustfeld als Koreferent sowie

beiden vorher genannten für ihre fachliche Unterstützung.