

Mathematische Methoden für Computer Vision, Robotik und Grafik

Kursnotizen für CS 205A, Herbst 2013

Justin Solomon
Abteilung für Computerwissenschaften
Universität in Stanford

Inhalt

I Vorbereitungen	9
0 Mathematik-Rezension	11
0.1 Vorbereitungen: Zahlen und Mengen	11
0.2 Vektorräume	12
0.2.1 Vektorräume definieren	12
0.2.2 Spanne, lineare Unabhängigkeit und Basen	13
0.2.3 Unser Fokus: $\mathbb{R}^n$	14
0.3 Linearität.	16
0.3.1 Matrizen	17
0.3.2 Skalare, Vektoren und Matrizen	19
0.3.3 Modellproblem: $Ax = b$	21
0.4 Nichtlinearität: Differentialrechnung	22
0.4.1 Differenzierung	22
0.4.2 Optimierung	25
0.5 Probleme.	27
1 Numerik und Fehleranalyse 1.1	29
Speichern von Zahlen mit Bruchteilen	29
1.1.1 Fixpunktdarstellungen	30
1.1.2 Gleitkommadarstellungen	31
1.1.3 Weitere exotische Optionen	32
1.2 Fehler verstehen	33
1.2.1 Klassifizierungsfehler	34
1.2.2 Konditionierung, Stabilität und Genauigkeit.	35
1.3 Praktische Aspekte	36
1.3.1 Beispiel im größeren Maßstab: Summation	37
1.4 Probleme	39
II Lineare Algebra	41
2 Lineare Systeme und die LU-Zerlegung 2.1	43
Lösbarkeit linearer Systeme	43
2.2 Ad-hoc-Lösungsstrategien	45
2.3 Zeilenoperationen kodieren	46

2.3.1 Permutation	46
2.3.2 Zeilenskalierung	47
2.3.3 Eliminierung	47
2.4 Gaußsche Eliminierung	49
2.4.1 Stürmerwechsel	49
2.4.2 Rückenwechsel	50
2.4.3 Analyse der Gaußschen Eliminierung	50
2.5 LU-Faktorisierung	52
2.5.1 Konstruktion der Faktorisierung	53
2.5.2 LU implementieren	54
2.6 Probleme	55
3 Entwurf und Analyse linearer Systeme 3.1	57
Lösung quadratischer Systeme	57
3.1.1 Regression	57
3.1.2 Kleinste Quadrate	59
3.1.3 Weitere Beispiele	60
3.2 Besondere Eigenschaften linearer Systeme	61
3.2.1 Positiv definite Matrizen und die Cholesky-Faktorisierung	61
3.2.2 Sparsity	64
3.3 Sensitivitätsanalyse	66
3.3.1 Matrix- und Vektornormen	66
3.3.2 Konditionsnummern	68
3.4 Probleme	70
4 Spaltenräume und QR 4.1	73
Die Struktur der Normalgleichungen	73
4.2 Orthogonalität	74
4.2.1 Strategie für nicht-orthogonale Matrizen	75
4.3 Gram-Schmidt-Orthogonalisierung	76
4.3.1 Projektionen	76
4.3.2 Gram-Schmidt-Orthogonalisierung	77
4.4 Haushaltstransformationen	78
4.5 Reduzierte QR-Faktorisierung	80
4.6 Probleme	81
5 Eigenvektoren	83
5.1 Motivation	83
5.1.1 Statistik	83
5.1.2 Differentialgleichungen	84
5.2 Spektrale Einbettung	85
5.3 Eigenschaften von Eigenvektoren	86
5.3.1 Symmetrische und positiv definite Matrizen	87
5.3.2 Spezialisierte Eigenschaften	89
5.4 Eigenwerte berechnen	90
5.4.1 Power-Iteration	90

5.4.2 Inverse Iteration	. 92
5.4.3 Schalten	. 92
5.4.4 Finden mehrerer Eigenwerte	. 93
5.5 Sensibilität und Konditionierung	. 97
5.6 Probleme	. 98
6 Singulärwertzerlegung 6.1	99
Ableitung der SVD	. 99
6.1.1 Berechnung der SVD	. 101
6.2 Anwendungen des SVD	. 101
6.2.1 Lösung linearer Systeme und der Pseudoinversen	. 101
6.2.2 Zerlegung in äußere Produkte und Approximationen mit niedrigem Rang.	. 102
6.2.3 Matrixnormen	. 104
6.2.4 Das Procrustes-Problem und die Ausrichtung	. 105
6.2.5 Hauptkomponentenanalyse (PCA)	. 106
6.3 Probleme	. 107
III Nichtlineare Techniken	109
7 Nichtlineare Systeme	111
7.1 Einvariablenprobleme	. 111
7.1.1 Charakterisierung von Problemen	. 111
7.1.2 Kontinuität und Halbierung	. 112
7.1.3 Analyse der Wurzelfindung	. 112
7.1.4 Fixpunktiteration	. 113
7.1.5 Newtons Methode	. 114
7.1.6 Sekantenmethode	. 115
7.1.7 Hybridtechniken	. 116
7.1.8 Einzelvariablenfall: Zusammenfassung	. 116
7.2 Multivariable Probleme	. 117
7.2.1 Newtons Methode	. 117
7.2.2 Newton schneller machen: Quasi-Newton und Broyen	. 117
7.3 Konditionierung	. 119
7.4 Probleme	. 119
8 Uneingeschränkte Optimierung	121
8.1 Uneingeschränkte Optimierung: Motivation	. 121
8.2 Optimalität	. 122
8.2.1 Differenzielle Optimalität	. 123
8.2.2 Optimalität über Funktionseigenschaften	. 125
8.3 Eindimensionale Strategien	. 126
8.3.1 Newtons Methode	. 126
8.3.2 Suche im Goldenen Schnitt	. 127
8.4 Multivariable Strategien	. 129
8.4.1 Gradientenabstieg	. 129

8.4.2 Newtons Methode	129
8.4.3 Optimierung ohne Derivate: BFGS	130
8.5 Probleme	132
9 Eingeschränkte Optimierung	135
9.1 Motivation	135
9.2 Theorie der eingeschränkten Optimierung	137
9.3 Optimierungsalgorithmen	140
9.3.1 Sequentielle quadratische Programmierung (SQP)	140
9.3.2 Barrieremethoden	141
9.4 Konvexe Programmierung	141
9.5 Probleme	143
10 iterative lineare Löser	145
10.1 Gradientenabstieg	146
10.1.1 Ableitung des iterativen Schemas	146
10.1.2 Konvergenz	147
10.2 Konjugierte Farbverläufe	149
10.2.1 Motivation	149
10.2.2 Suboptimalität des Gradientenabstiegs	151
10.2.3 A-konjugierte Richtungen erzeugen	152
10.2.4 Formulieren des Algorithmus für konjugierte Gradienten	154
10.2.5 Konvergenz- und Stoppbedingungen	155
10.3 Vorkonditionierung	156
10.3.1 CG mit Vorkonditionierung	157
10.3.2 Gemeinsame Vorkonditionierer	158
10.4 Andere iterative Schemata	159
10.5 Probleme	160
IV Funktionen, Ableitungen und Integrale	161
11 Interpolation	163
11.1 Interpolation in einer einzelnen Variablen	163
11.1.1 Polynominterpolation	164
11.1.2 Alternative Basen	165
11.1.3 Stückweise Interpolation	167
11.1.4 Gaußsche Prozesse und Kriging	168
11.2 Multivariable Interpolation	168
11.3 Theorie der Interpolation	171
11.3.1 Lineare Algebra von Funktionen	171
11.3.2 Approximation über stückweise Polynome	173
11.4 Probleme	174

12 Numerische Integration und Differenzierung	175
12.1 Motivation	176
12.2 Quadratur	177
12.2.1 Interpolatorische Quadratur	177
12.2.2 Quadraturregeln	178
12.2.3 Newton-Cotes-Quadratur	179
12.2.4 Gaußsche Quadratur	182
12.2.5 Adaptive Quadratur	183
12.2.6 Mehrere Variablen	183
12.2.7 Konditionierung	184
12.3 Differenzierung	185
12.3.1 Basisfunktionen differenzieren	185
12.3.2 Endliche Differenzen	185
12.3.3 Auswahl der Schrittgröße	187
12.3.4 Integrierte Größen	187
12.4 Probleme	188
13 Gewöhnliche Differentialgleichungen	189
13.1 Motivation	190
13.2 Theorie der ODEs	190
13.2.1 Grundbegriffe	191
13.2.2 Existenz und Einzigartigkeit	192
13.2.3 Modellgleichungen	193
13.3 Zeitschrittverfahren	194
13.3.1 Vorwärts-Euler	194
13.3.2 Rückwärts-Euler	195
13.3.3 Trapezmethode	196
13.3.4 Runge-Kutta-Methoden	197
13.3.5 Exponentielle Integratoren	198
13.4 Mehrwertige Methoden	199
13.4.1 Newmark-Systeme	199
13.4.2 Versetztes Raster	202
13.5 Zu erledigen	203
13.6 Probleme	203
14 Partielle Differentialgleichungen	205
14.1 Motivation	205
14.2 Grundlegende Definitionen	209
14.3 Modellgleichungen	209
14.3.1 Elliptische PDEs	210
14.3.2 Parabolische PDEs	211
14.3.3 Hyperbolische PDEs	211
14.4 Derivate als Operatoren	212
14.5 PDEs numerisch lösen	214
14.5.1 Elliptische Gleichungen lösen	214
14.5.2 Parabolische und hyperbolische Gleichungen lösen	215

14.6 Methode der Finiten Elemente	. 217
14.7 Beispiele aus der Praxis	. 217
14.7.1 Bildverarbeitung im Gradientenbereich	. 217
14.7.2 Kantenerhaltende Filterung	. 217
14.7.3 Gitterbasierte Flüssigkeiten	. 217
14.8 Zu erledigen	. 217
14.9 Probleme	. 218

Teil I

Vorrunden

Kapitel 0

Mathematik-Rezension

In diesem Kapitel werden wir relevante Begriffe aus der linearen Algebra und der Multivariablenrechnung besprechen, die in unsere Diskussion über Rechentechniken einfließen. Es ist als Überblick über Hintergrundmaterial gedacht, wobei der Schwerpunkt auf Ideen und Interpretationen liegt, die in der Praxis häufig anzutreffen sind. Das Kapitel kann bedenkenlos übersprungen oder als Nachschlagewerk für Schüler mit ausgeprägteren Mathematikkenntnissen verwendet werden.

0.1 Vorbereitungen: Zahlen und Mengen

Anstatt algebraische (und manchmal philosophische) Diskussionen wie „Was ist eine Zahl?“ zu berücksichtigen, werden wir uns auf Intuition und mathematischen gesunden Menschenverstand verlassen, um einige Mengen zu

definieren: • Die natürlichen Zahlen $\mathbf{N} = \{1, 2, 3, \dots\}$

• Die ganzen Zahlen $\mathbf{Z} = \{\dots, -2, -1, 0, 1, 2, \dots\}$

• Die rationalen Zahlen $\mathbf{Q} = \{a/b : a, b \in \mathbf{Z}, b \neq 0\}$ • Die

reellen Zahlen $\mathbf{R}$, die $\mathbf{Q}$ sowie irrationale Zahlen wie $\sqrt{2}$ und $\sqrt{5}$ umfassen • Die komplexen Zahlen $\mathbf{C} =$

$\{a + bi : a, b \in \mathbf{R}\}$, wobei wir davon ausgehen, dass $i^2 = -1$ erfüllt.

Es muss anerkannt werden, dass unsere Definition von $\mathbf{R}$ alles andere als streng ist. Die Konstruktion der reellen Zahlen kann ein wichtiges Thema für Praktiker von Kryptographietechniken sein, die alternative Zahlensysteme verwenden, aber diese Feinheiten sind für die vorliegende Diskussion irrelevant.

Wie alle anderen Mengen können $\mathbf{N}$, $\mathbf{Z}$, $\mathbf{Q}$, $\mathbf{R}$ und $\mathbf{C}$ mithilfe generischer Operationen manipuliert werden, um neue Zahlenmengen zu erzeugen. Denken Sie insbesondere daran, dass wir das „euklidische Produkt“ zweier Mengen A und B definieren können als

$$A \times B = \{(a, b) : a \in A \text{ und } b \in B\}.$$

Wir können Potenzen von Mengen durch Schreiben ermitteln

$$A^n = \underbrace{A \times A \times \dots \times A}_{n \text{ mal}}.$$

¹Dies ist das erste von vielen Malen, dass wir die Notation $\{A : B\}$ verwenden; Die geschweiften Klammern sollten eine Menge bezeichnen und der Doppelpunkt kann als „so dass“ gelesen werden. Beispielsweise kann die Definition von $\mathbf{Q}$ als „die Menge der Brüche a/b , sodass a und b ganze Zahlen sind“ gelesen werden. Als zweites Beispiel könnten wir $\mathbf{N} = \{n \in \mathbf{Z} : n > 0\}$ schreiben.

Diese Konstruktion ergibt das, was in den kommenden Kapiteln zu unserem Lieblingszahlsatz werden wird:

$$\mathbf{R}^N = \{(a_1, a_2, \dots, a_n) : a_i \in \mathbf{R} \text{ für alle } i\}$$

0,2 Vektorräume

Einführungskurse in die lineare Algebra könnten problemlos den Titel „Einführung in endlichdimensionale Vektorräume“ tragen. Obwohl die Definition eines Vektorraums abstrakt erscheinen mag, werden wir viele konkrete Anwendungen finden, die alle die formalen Aspekte erfüllen und daher von der Maschinerie, die wir entwickeln werden, profitieren können.

0.2.1 Vektorräume definieren

Wir beginnen mit der Definition eines Vektorraums und geben eine Reihe von Beispielen:

Definition 0.1 (Vektorraum). Ein Vektorraum ist eine Menge V , die durch Skalarmultiplikation und -addition abgeschlossen ist.

Für unsere Zwecke ist ein Skalar eine Zahl in $\mathbf{R}$, und die Additions- und Multiplikationsoperationen erfüllen die üblichen Axiome (Kommutativität, Assoziativität usw.). Normalerweise ist es einfach, Vektorräume in freier Wildbahn zu erkennen, einschließlich der folgenden Beispiele:

Beispiel 0.1 ($\mathbf{R}^n$ als Vektorraum). Das häufigste Beispiel für einen Vektorraum ist $\mathbf{R}^n$. Hier erfolgen Addition und Skalarmultiplikation Komponente für Komponente:

$$(1, 2) + (3, 4) = (1+3, 2+4) = (4, 6)$$

$$10 \cdot (1, 1) = (10 \cdot 1, 10 \cdot 1) = (10, 10)$$

Beispiel 0.2 (Polynome). Ein zweites wichtiges Beispiel eines Vektorraums ist der „Ring“ von Polynomen mit reellen Zahleneingaben, bezeichnet mit $\mathbf{R}[x]$. Ein Polynom $p \in \mathbf{R}[x]$ ist eine Funktion $p : \mathbf{R} \rightarrow \mathbf{R}$ der Form²

$$p(x) = \sum_k a_k x^k$$

Addition und Skalarmultiplikation werden auf die übliche Weise durchgeführt, z. B. wenn $p(x) = x^2 + 2x + 1$ und $q(x) = x^3 + 3x^2 + 5x + 1$, dann ist $(p+q)(x) = x^3 + 4x^2 + 7x + 2$. Beachten Sie, dass $x^2 + 6x + 3$, was ein weiteres Polynom ist. Abgesehen davon sind Funktionen wie $p(x) = (x+1)(x+1) + 1$ gelten, obwohl sie nicht explizit in der Form $\sum_k a_k x^k$ (für $k \leq 5$) immer noch gerade Polynome obigen Form geschrieben sind.

Elemente $v \in V$ eines Vektorraums V werden Vektoren genannt, und eine gewichtete Summe der Form $\sum_i a_i v_i$, wobei $a_i \in \mathbf{R}$ und $v_i \in V$ ist, wird als Linearkombination der v_i bezeichnet. In unserem zweiten Beispiel sind die „Vektoren“ Funktionen, obwohl wir diese Sprache normalerweise nicht verwenden, um $\mathbf{R}[x]$ zu diskutieren. Eine Möglichkeit, Standpunkte mit der Sequenz zu verknüpfen, wäre die Identifizierung des Polynoms $\sum_k a_k x^k$ mit dieser beiden $(a_0, a_1, a_2, \dots)$; Denken Sie daran, dass Polynome endlich viele Terme haben, sodass die Folge letztendlich in einer Folge von Nullen endet.

²Die Notation $f : A \rightarrow B$ bedeutet, dass f eine Funktion ist, die als Eingabe ein Element der Menge A nimmt und ein Element der Menge B ausgibt. Beispielsweise nimmt $f : \mathbf{R} \rightarrow \mathbf{Z}$ als Eingabe eine reelle Zahl in $\mathbf{R}$ und gibt eine ganze Zahl Z aus, wie es bei $f(x) = x$ der Fall sein könnte, der „Abrundungs“-Funktion.

0.2.2 Spanne, lineare Unabhängigkeit und Basen

Angenommen, wir beginnen mit den Vektoren $v_1, \dots, v_k \in V$ für den Vektorraum V . Nach Definition 0.1 haben wir zwei Möglichkeiten, mit diesen Vektoren zu beginnen und neue Elemente von V zu konstruieren: Addition und Skalarmultiplikation. Die Idee von span besteht darin, dass es alle Vektoren beschreibt, die Sie über diese beiden Operationen erreichen können:

Definition 0.2 (Span). Die Spanne einer Menge $S \subseteq V$ von Vektoren ist die Menge

$$\text{span } S = \{a_1 v_1 + \dots + a_k v_k : k \in \mathbb{N}, v_i \in S \text{ für alle } i \text{ und } a_i \in \mathbb{R} \text{ für alle } i\}.$$

Beachten Sie, dass die Spanne S ein Unterraum von V ist, also eine Teilmenge von V , die selbst ein Vektorraum ist. Wir können einige Beispiele nennen:

Beispiel 0.3 (Mixologie). Der typische „Brunnen“ einer Cocktailbar enthält mindestens vier Zutaten, die dem Barkeeper zur Verfügung stehen: Wodka, Tequila, Orangensaft und Grenadine. Vorausgesetzt, wir haben diesen einfachen Brunnen, können wir Getränke als Punkte in $\mathbb{R}^4$ mit einem Platz für jede Zutat darstellen. Beispielsweise kann ein typischer „Tequila-Sonnenaufgang“ mit dem Punkt $(0, 1, 5, 6, 0, 75)$ dargestellt werden, der die Mengen an Wodka, Tequila, Orangensaft und Grenadine (in Unzen) darstellt.

Das Getränkeset, das mit dem typischen Brunnen zubereitet werden kann, ist in enthalten

$$\text{span} \{(1, 0, 0, 0), (0, 1, 0, 0), (0, 0, 1, 0), (0, 0, 0, 1)\},$$

das heißt, alle Kombinationen der vier Grundzutaten. Ein Barkeeper, der Zeit sparen möchte, bemerkt jedoch möglicherweise, dass viele Getränke das gleiche Verhältnis von Orangensaft zu Grenadine haben, und mischt die Flaschen. Der neue, vereinfachte Brunnen lässt sich zwar leichter einschenken, kann aber wesentlich weniger Getränke zubereiten:

$$\text{Spanne} \{(1, 0, 0, 0), (0, 1, 0, 0), (0, 0, 6, 0, 75)\}$$

Beispiel 0.4 (Polynome). Definieren Sie $p_k(x) \in \mathbb{R}[x]$. Dann ist das leicht zu erkennen

$$\mathbb{R}[x] = \text{span} \{p_k : k \in \mathbb{N}\}.$$

Stellen Sie sicher, dass Sie die Notation gut genug verstehen, um zu verstehen, warum dies der Fall ist.

Das Hinzufügen eines weiteren Elements zu einer Menge von Vektoren führt nicht immer zu einer Vergrößerung der Spanne. Für $\mathbb{R}^2$ ist dies beispielsweise eindeutig der Fall

$$\text{Spanne} \{(1, 0), (0, 1)\} = \text{Spanne} \{(1, 0), (0, 1), (1, 1)\}.$$

In diesem Fall sagen wir, dass die Menge $\{(1, 0), (0, 1), (1, 1)\}$ linear abhängig ist:

Definition 0.3 (Lineare Abhängigkeit). Wir bieten drei äquivalente Definitionen. Eine Menge $S \subseteq V$ von Vektoren ist linear abhängig, wenn:

1. Eines der Elemente von S kann als lineare Kombination der anderen Elemente geschrieben werden, oder S enthält null.
2. Es existiert eine nichtleere Linearkombination von Elementen $v_k \in S$, die $\sum_{k=1}^M c_k v_k = 0$ für alle k ergibt, wobei $\sum_{k=1}^M c_k v_k = 0$ wobei
3. Es gibt $v \in S$, so dass $\text{Spanne } S = \text{Spanne } S \setminus \{v\}$. Das heißt, wir können einen Vektor ohne entfernen aus S

Auswirkungen auf seine Spannweite.

Wenn S nicht linear abhängig ist, dann sagen wir, dass es linear unabhängig ist.

Für Schüler, die mit Notation und abstrakter Mathematik weniger vertraut sind, ist es eine lohnenswerte Übung, Beweise oder informelle Nachweise dafür zu erbringen, dass jede Definition ihren Entsprechungen entspricht (auf die Art und Weise, „wenn und nur wenn“).

Das Konzept der linearen Abhängigkeit führt zu einer Vorstellung von „Redundanz“ in einer Menge von Vektoren. In diesem Sinne stellt sich natürlich die Frage, wie groß die Menge sein kann, die wir wählen können, bevor das Hinzufügen eines weiteren Vektors die Spanne unmöglich vergrößern kann. Nehmen wir insbesondere an, wir haben eine linear unabhängige Menge $S \subseteq V$ und wählen nun einen zusätzlichen Vektor $v \in V$. Das Hinzufügen von v zu S führt zu einem von zwei möglichen Ergebnissen:

1. Die Spannweite von $S \cup \{v\}$ ist größer als die Spannweite von S .

2. Das Hinzufügen von v zu S hat keinen Einfluss auf die Spanne.

Die Dimension von V ist nichts anderes als die maximale Häufigkeit, mit der wir Ergebnis 1 erhalten, v zu S addieren und wiederholen können.

Definition 0,4 (Dimension und Basis). Die Dimension von V ist die maximale Größe $|S|$ einer linear unabhängigen Menge $S \subseteq V$ mit der Spanne $S = V$. Jede Menge S , die diese Eigenschaft erfüllt, heißt Basis für V .

Beispiel 0,5 ($\mathbb{R}^n$). Die Standardbasis für $\mathbb{R}^n$ ist die Menge der Vektoren der Form

$$e_k \in (0, \dots, \underbrace{0}_{k-1 \text{ Slots}}, \underbrace{1}_{k \text{ Slot}}, \underbrace{0, \dots, 0}_{n-k \text{ Slots}}).$$

Das heißt, e_k hat alle Nullen bis auf eine einzelne Eins im k -ten Slot. Es ist klar, dass diese Vektoren linear unabhängig sind und eine Basis bilden; Beispielsweise kann in $\mathbb{R}^3$ jeder Vektor (a, b, c) als $ae_1 + be_2 + ce_3$ geschrieben werden.

Daher ist die Dimension von $\mathbb{R}^n$ erwartungsgemäß n .

Beispiel 0.6 (Polynome). Es ist klar, dass die Menge $\{1, x, x^2, \dots\}$ ist eine linear unabhängige Menge von Polynomen, die $\mathbb{R}[x]$ aufspannen. Beachten Sie, dass diese Menge unendlich groß ist und daher die Dimension von $\mathbb{R}[x]$ ∞ ist.

0.2.3 Unser Fokus: $\mathbb{R}^n$

Von besonderer Bedeutung für unsere Zwecke ist der Vektorraum $\mathbb{R}^n$, der klassische, das sogenannte n -dimensionale Euklidische Raum. Dies ist nichts anderes als der Satz von Koordinatenachsen, die man im Mathematikunterricht der Oberstufe antrifft:

- $\mathbb{R}^1 \subseteq \mathbb{R}$ ist der Zahlenstrahl
- $\mathbb{R}^2$ ist die zweidimensionale Ebene mit den Koordinaten (x, y)
- $\mathbb{R}^3$ stellt den dreidimensionalen Raum mit den Koordinaten (x, y, z) dar

Fast alle Methoden in diesem Kurs befassen sich mit Transformationen und Funktionen auf dem $\mathbb{R}^n$.

Der Einfachheit halber schreiben wir Vektoren im $\mathbb{R}^n$ normalerweise wie folgt in „Spaltenform“.

$$(a_1, \dots, a_n) \in \mathbb{R}^n \quad \begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_n \end{pmatrix}$$

Diese Notation umfasst Vektoren als Sonderfälle von Matrizen, die unten besprochen werden.

Im Gegensatz zu einigen Vektorräumen hat $\mathbb{R}^n$ nicht nur eine Vektorraumstruktur, sondern auch eine zusätzliche Konstruktion, die den entscheidenden Unterschied macht: das Skalarprodukt.

Definition 0,5 (Skalarprodukt). Das Skalarprodukt zweier Vektoren $a = (a_1, \dots, a_n)$ und $b = (b_1, \dots, b_n)$ im $\mathbb{R}^n$ ist gegeben durch

$$a \cdot b = \sum_{k=1}^n a_k b_k.$$

Beispiel 0,7 ($\mathbb{R}^2$). Das Skalarprodukt von $(1, 2)$ und $(\sqrt{2}, 6)$ ist $1 \cdot \sqrt{2} + 2 \cdot 6 = \sqrt{2} + 12 = 10$.

Das Skalarprodukt ist ein Beispiel für eine Metrik, und seine Existenz gibt $\mathbb{R}^n$ eine Vorstellung von der Geometrie. Beispielsweise können wir den Satz des Pythagoras verwenden, um die Norm oder Länge eines Vektors als Quadratwurzel zu definieren

$$\sqrt{a_1^2 + \dots + a_n^2} = \sqrt{a \cdot a}.$$

Dann ist der Abstand zwischen zwei Punkten $a, b \in \mathbb{R}^n$ einfach $\sqrt{a \cdot a - 2a \cdot b + b \cdot b}$.

Punktprodukte liefern nicht nur Vorstellungen von Längen und Abständen, sondern auch von Winkeln. Erinnern Sie sich an die folgende Identität aus der Trigonometrie für $\mathbb{R}^3$:

$$a \cdot b = ab \cos \gamma$$

wobei γ der Winkel zwischen a und b ist. Für $n \geq 4$ ist der Begriff „Winkel“ für $\mathbb{R}^n$ jedoch viel schwieriger zu visualisieren. Wir könnten den Winkel γ zwischen a und b als den Wert γ definieren, der durch gegeben ist

$$\gamma = \arccos \frac{a \cdot b}{ab}.$$

Wir müssen unsere Hausaufgaben machen, bevor wir eine solche Definition machen! Denken Sie insbesondere daran, dass die Kosinusfunktion Werte im Intervall $[-1, 1]$ ausgibt. Daher müssen wir überprüfen, ob die Eingabe für den Arcuskosinus (auch als $\cos^{-1}$ bezeichnet) in diesem Intervall liegt. Glücklicherweise garantiert die bekannte Cauchy-Schwarz-Ungleichung $a \cdot b \leq ab$ genau diese Eigenschaft.

Wenn $a = cb$ für ein $c \in \mathbb{R}$, gilt $\gamma = \arccos 1 = 0$, wie wir erwarten würden: Der Winkel zwischen parallelen Vektoren ist Null. Was bedeutet es, dass Vektoren senkrecht zueinander stehen? Ersetzen wir $\gamma = 90^\circ$. Dann haben wir

$$0 = \cos 90^\circ = \frac{a \cdot b}{ab}$$

Die Multiplikation beider Seiten mit ab motiviert die Definition:

Definition 0,6 (Orthogonalität). Zwei Vektoren sind senkrecht oder orthogonal, wenn $a \cdot b = 0$.

Diese Definition ist aus geometrischer Sicht etwas überraschend. Insbesondere ist es uns gelungen, zu definieren, was es bedeutet, senkrecht zu sein, ohne explizit Winkel zu verwenden. Diese Konstruktion wird es einfacher machen, bestimmte Probleme zu lösen, bei denen die Nichtlinearität von Sinus und Cosinus in einfacheren Situationen möglicherweise Kopfschmerzen bereitet hätte.

Neben 0,1. Hier gibt es viele theoretische Fragen, über die wir nachdenken müssen, einige davon werden wir in späteren Kapiteln behandeln, wenn sie motivierter sind:

- Lassen alle Vektorräume Skalarprodukte oder ähnliche Strukturen zu?
- Lassen alle endlichdimensionalen Vektorräume Skalarprodukte zu?
- Was könnte ein sinnvolles Skalarprodukt zwischen Elementen von $R[x]$ sein?

Interessierte Studierende können Texte zur Real- und Funktionsanalyse zu Rate ziehen.

0,3 Linearität

Eine Funktion zwischen Vektorräumen, die die Struktur bewahrt, wird als lineare Funktion bezeichnet:

Definition 0,7 (Linearität). Angenommen, V und V sind Vektorräume. Dann ist $L : V \rightarrow V$ linear, wenn es die folgenden zwei Kriterien für alle $v, v_1, v_2 \in V$ und $c \in R$ erfüllt:

- L bewahrt Summen: $L[v_1 + v_2] = L[v_1] + L[v_2]$
- L bewahrt Skalarprodukte: $L[cv] = cL[v]$

Es ist einfach, lineare Abbildungen zwischen Vektorräumen zu erstellen, wie wir in den folgenden Beispielen sehen können:

Beispiel 0.8 (Linearität in R^n). Die folgende Abbildung $f : R^2 \rightarrow R^3$ ist linear:

$$f(x, y) = (3x, 2x + y, y)$$

Wir können die Linearität wie folgt überprüfen:

- Summenerhaltung:

$$\begin{aligned} f(x_1 + x_2, y_1 + y_2) &= (3(x_1 + x_2), 2(x_1 + x_2) + (y_1 + y_2), y_1 + y_2) \\ &= (3x_1, 2x_1 + y_1, y_1) + (3x_2, 2x_2 + y_2, y_2) \\ &= f(x_1, y_1) + f(x_2, y_2) \end{aligned}$$

- Skalare Produktkonservierung:

$$\begin{aligned} f(cx, cy) &= (3cx, 2cx + cy, cy) = c(3x, \\ &2x + y, y) = c f(x, y) \end{aligned}$$

Im Gegensatz dazu ist $g(x, y) = xy^2$ nicht linear. Zum Beispiel ist $g(1, 1) = 1$, aber $g(2, 2) = 8 = 2 \cdot g(1, 1)$, daher behält diese Form keine Skalarprodukte bei.

Beispiel 0.9 (Integration). Das folgende „funktionale“ L von $R[x]$ nach R ist linear:

$$L[p(x)] = \int_0^1 p(x) dx.$$

Dieses etwas abstraktere Beispiel bildet Polynome $p(x)$ auf reelle Zahlen $L[p(x)]$ ab. Wir können zum Beispiel schreiben

$$L[3x^2 + x - 1] = \int_0^1 (3x^2 + x - 1) dx = 2.$$

Die Linearität ergibt sich aus den folgenden bekannten Fakten aus der Analysis:

$$\begin{aligned} \int_0^1 c \cdot f(x) dx &= c \int_0^1 f(x) dx \\ \int_0^1 [f(x) + g(x)] dx &= \int_0^1 f(x) dx + \int_0^1 g(x) dx \end{aligned}$$

Wir können eine besonders schöne Form für lineare Abbildungen auf $\mathbb{R}^n$ schreiben. Denken Sie daran, dass der Vektor $a = (a_1, \dots, a_n)$ gleich der Summe $\sum_k a_k e_k$ ist, wobei e_k der k -te Standardbasisvektor ist. Wenn L dann linear ist, wissen wir:

$$\begin{aligned} L[a] &= L \sum_k a_k e_k \text{ für die Standardbasis } e_k \\ &= \sum_k a_k L[e_k] \text{ durch Summenerhaltung} \\ &= \sum_k a_k L[e_k] \text{ durch Skalarprodukterhaltung} \end{aligned}$$

Diese Ableitung zeigt die folgende wichtige Tatsache:

L wird vollständig durch seine Wirkung auf die Standardbasis vectorse bestimmt.

Das heißt, für jeden Vektor a , Wir können die obige Summe verwenden, um $L[a]$ durch lineare Kombination zu bestimmen $\mathbb{R}^n$ $L[e_1], \dots, L[e_n]$.

Beispiel 0.10 (Erweitern einer linearen Karte). Erinnern Sie sich an die Karte in Beispiel 0.8, gegeben durch $f(x, y) = (3x, 2x + y, y)$. Es gilt $f(e_1) = f(1, 0) = (3, 2, 0)$ und $f(e_2) = f(0, 1) = (0, 1, 1)$. Somit zeigt die obige Formel:

$$f(x, y) = x f(e_1) + y f(e_2) = x \begin{pmatrix} 3 \\ 2 \\ 0 \end{pmatrix} + y \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}$$

0.3.1 Matrizen

Die obige Erweiterung linearer Karten deutet auf einen von vielen Kontexten hin, in denen es nützlich ist, mehrere Vektoren in derselben Struktur zu speichern. Allgemeiner gesagt, wir haben n Vektoren $v_1, \dots, v_n \in \mathbb{R}^m$. Wir können jeden als Spaltenvektor schreiben:

$$v_1 = \begin{pmatrix} v_{11} \\ v_{21} \\ \vdots \\ v_{m1} \end{pmatrix}, v_2 = \begin{pmatrix} v_{12} \\ v_{22} \\ \vdots \\ v_{m2} \end{pmatrix}, \dots, v_n = \begin{pmatrix} v_{1n} \\ v_{2n} \\ \vdots \\ v_{mn} \end{pmatrix}$$

Diese getrennt mit sich herumzutragen kann in der Notation umständlich sein. Um die Sache zu vereinfachen, kombinieren wir sie einfach in einer einzigen $m \times n$ -Matrix:

$$\begin{pmatrix} \vec{y} & | & \\ \hline v_1 & v_2 & \dots & v_n \\ \hline \vec{y} & | & \end{pmatrix} = \begin{pmatrix} \vec{y} & v_{11} & v_{12} & \dots & v_{1n} \\ v_{21} & v_{22} & \dots & v_{2n} \\ \vdots & \vdots & & \vdots \\ \hline v_{m1} & v_{m2} & \dots & v_{mn} \end{pmatrix}$$

Wir nennen den Raum solcher Matrizen $R_{m \times n}$.

Beispiel 0.11 (Identitätsmatrix). Wir können die Standardbasis für R_n in der $n \times n$ „Identitätsmatrix“ speichern.

$I_{n \times n}$ gegeben durch:

$$I_{n \times n} = \begin{pmatrix} \vec{y} & | & \\ \hline e_1 & e_2 & \dots & e_n \\ \hline \vec{y} & | & \end{pmatrix} = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & & \vdots \\ 0 & 0 & \dots & 1 \end{pmatrix}$$

Da wir Matrizen als bequeme Möglichkeit zum Speichern von Vektormengen konstruiert haben, können wir mithilfe der Multiplikation ausdrücken, wie sie linear kombiniert werden können. Insbesondere kann eine Matrix in $R_{m \times n}$ mit einem Spaltenvektor in R_n wie folgt multipliziert werden:

$$\begin{pmatrix} \vec{y} & | & \\ \hline v_1 & v_2 & \dots & v_n \\ \hline \vec{y} & | & \end{pmatrix} \begin{pmatrix} c_1 \\ c_2 \\ \vdots \\ c_n \end{pmatrix} = \begin{pmatrix} c_1 v_1 + c_2 v_2 + \dots + c_n v_n \end{pmatrix}$$

Die Erweiterung dieser Summe ergibt die folgende explizite Formel für Matrix-Vektor-Produkte:

$$\begin{pmatrix} \vec{y} & v_{11} & v_{12} & \dots & v_{1n} \\ v_{21} & v_{22} & \dots & v_{2n} \\ \vdots & \vdots & & \vdots \\ \hline v_{m1} & v_{m2} & \dots & v_{mn} \end{pmatrix} \begin{pmatrix} c_1 \\ c_2 \\ \vdots \\ c_n \end{pmatrix} = \begin{pmatrix} c_1 v_{11} + c_2 v_{12} + \dots + c_n v_{1n} \\ c_1 v_{21} + c_2 v_{22} + \dots + c_n v_{2n} \\ \vdots \\ c_1 v_{m1} + c_2 v_{m2} + \dots + c_n v_{mn} \end{pmatrix}$$

Beispiel 0.12 (Identitätsmatrixmultiplikation). Es ist eindeutig wahr, dass für jedes $x \in R_n$, wir können schreiben $x = I_{n \times n} x$ ist, wobei $I_{n \times n}$ die Identitätsmatrix aus Beispiel 0.11 ist.

Beispiel 0.13 (Lineare Karte). Wir kehren noch einmal zum Ausdruck aus Beispiel 0.8 zurück, um eine weitere alternative Form zu zeigen:

$$f(x, y) = \begin{pmatrix} 3 & 0 & 2 \\ 1 & & \\ 0 & y_1 & \end{pmatrix} \begin{pmatrix} x \\ y \\ j \end{pmatrix}$$

Auf ähnliche Weise definieren wir ein Produkt zwischen einer Matrix in $M \in R_{m \times n}$ und einer anderen Matrix in $R_{n \times p}$, indem wir einzelne Matrix-Vektor-Produkte verketteten:

$$M \begin{pmatrix} \vec{y} & | & \\ \hline c_1 & c_2 & \dots & c_n \\ \hline \vec{y} & | & \end{pmatrix} = \begin{pmatrix} M c_1 & M c_2 & \dots & M c_n \end{pmatrix}$$

Beispiel 0.14 (Mixologie). Nehmen wir in Fortsetzung von Beispiel 0.3 an, dass wir in unserem vereinfachten Brunnen einen Tequila Sunrise und eine zweite Mischung mit gleichen Teilen der beiden Liköre herstellen. Um herauszufinden, wie viele Grundzutaten in jeder Bestellung enthalten sind, könnten wir die Rezepte für jede Spalte spaltenweise kombinieren und eine Matrix verwenden Multiplikation:

$$\begin{array}{ccccccc}
 & \text{Gut 1} & \text{Gut 2} & \text{Gut 3} & & \text{Trinken Sie 1, trinken Sie 2} & \\
 \text{Wodka} & 1 & 0 & 0 & & 0 & 0,75 \\
 \text{Tequila} & 0 & 1 & 0 & \cdot & 1,5 & 0,75 \\
 \text{ABI} & 0 & 0 & 6 & & 1 & 2 \\
 \text{Grenadine} & 0 & 0 & 0,75 & & 0,75 & 1,5
 \end{array} = \begin{array}{ccc}
 \text{Trinken Sie 1, trinken Sie 2} & & \\
 0 & 0,75 & 1,5 \\
 6 & 12 & 1,5 \\
 0,75 & 1,5 & 1,5
 \end{array}$$

Im Allgemeinen verwenden wir Großbuchstaben zur Darstellung von Matrizen, wie z. B. $A \in \mathbb{R}^{m \times n}$. Wir werden das verwenden Notation $A_{ij} \in \mathbb{R}$, um das Element von A in Zeile i und Spalte j zu bezeichnen.

0.3.2 Skalare, Vektoren und Matrizen

Es ist keine Überraschung, dass wir einen Skalar als 1×1 -Vektor $c \in \mathbb{R}^{1 \times 1}$ schreiben können. Ähnlich, wie wir schon Wie in §0.2.3 vorgeschlagen, können Vektoren in $\mathbb{R}^n$, wenn wir sie in Spaltenform schreiben, als $n \times 1$ -Matrizen betrachtet werden $v \in \mathbb{R}^{n \times 1}$. Beachten Sie, dass Matrix-Vektor-Produkte in diesem Zusammenhang leicht interpretiert werden können; Zum Beispiel, wenn $A \in \mathbb{R}^{m \times n}$, $x \in \mathbb{R}^n$, und $b \in \mathbb{R}^m$, dann können wir Ausdrücke schreiben wie

$$\begin{array}{ccc}
 A & x & = b \\
 m \times n & n \times 1 & m \times 1
 \end{array}$$

Wir werden einen zusätzlichen Operator für Matrizen einführen, der in diesem Zusammenhang nützlich ist:

Definition 0.8 (Transponieren). Die Transponierte einer Matrix $A \in \mathbb{R}^{m \times n}$ ist eine Matrix $A^T \in \mathbb{R}^{n \times m}$ mit Elementen $(A^T)_{ij} = A_{ji}$.

Beispiel 0.15 (Transposition). Die Transponierte der Matrix

$$A = \begin{pmatrix} 1 & 2 \\ 3 & 4 \\ 5 & 6 \end{pmatrix}$$

ist gegeben durch

$$A^T = \begin{pmatrix} 1 & 3 & 5 \\ 2 & 4 & 6 \end{pmatrix}$$

Geometrisch gesehen können wir uns Transposition als Umkehrung einer Matrix auf ihrer Diagonale vorstellen.

Diese einheitliche Behandlung von Skalaren, Vektoren und Matrizen in Kombination mit Operationen wie Transposition und Multiplikation kann zu raffinierten Ableitungen bekannter Identitäten führen. Zum Beispiel,

Wir können die Skalarprodukte der Vektoren $s, a, b \in \mathbb{R}^n$ berechnen, indem wir die folgende Reihe von Schritten ausführen:

$$a \cdot b = \sum_{k=1}^n \tilde{y}_k a_k b_k$$

Viele wichtige Identitäten aus der linearen Algebra lassen sich ableiten, indem man diese Operationen mit ein paar Regeln verkettet:

$$(A) = A$$
$$(A + B) = A + B$$
$$(AB) = BA$$

Beispiel 0,16 (Restnorm). Angenommen, wir haben eine Matrix A und zwei Vektoren x und b . Wenn wir wissen möchten, wie gut Ax b annähert, könnten wir einen Residuum $r = b - Ax$ definieren; Dieses Residuum ist genau dann Null, wenn $Ax = b$. Andernfalls könnten wir die Norm r^2 als Proxy für die Beziehung zwischen Ax und b verwenden.

Wir können die obigen Identitäten zur Vereinfachung verwenden:

$$\begin{aligned} R^2 &= b^T A x \\ &= (b^T A x) \cdot (b^T A x) \text{ wie in §0.2.3 erklärt} = (b^T A x) \\ &= (b^T A x) \text{ durch unseren Ausdruck für das Skalarprodukt oben} = (b^T A x) \\ &= b^T A x \text{ durch Eigenschaften der Transposition} = b^T A x \\ &= b^T A x \text{ nach der Multiplikation} \end{aligned}$$

Alle vier Terme auf der rechten Seite sind Skalare oder entsprechend 1×1 -Matrizen. Als Matrizen gedachte Skalare genießen trivialerweise eine zusätzliche schöne Eigenschaft $c = c$, da es nichts zu transponieren gibt! So können wir schreiben

$$x \mathbin{\dot{A}} b = (x \mathbin{\dot{A}} b) = b \mathbin{\dot{A}} x$$

Dadurch können wir unseren Ausdruck noch weiter vereinfachen:

$$\begin{aligned} R_2^T &= b \tilde{y} - 2b A x + x^T A A x \\ &= A x^T \tilde{y} - 2b A x + b \end{aligned}$$

Wir hätten diesen Ausdruck mithilfe von Punktproduktidentitäten ableiten können, aber die obigen Zwischenschritte werden sich in unserer späteren Diskussion als nützlich erweisen.

0.3.3 Modellproblem: $Ax = b$

Im Einführungskurs in die Algebra verbringen die Schüler viel Zeit damit, lineare Systeme wie das folgende für Tripel (x, y, z) zu lösen:

$$\begin{aligned} 3x + 2y + 5z &= 0 \\ 4x + 9y - 3z &= 7 \\ 3y - 3z &= 1 \end{aligned}$$

Unsere Konstruktionen in §0.3.1 ermöglichen es uns, solche Systeme sauberer zu kodieren:

$$\begin{pmatrix} 3 & 2 & 5 \\ 4 & 9 & -3 \\ 0 & 3 & -3 \end{pmatrix} \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 0 \\ 7 \\ 1 \end{pmatrix}$$

Allgemeiner gesagt können wir lineare Gleichungssysteme in der Form $Ax = b$ schreiben, indem wir dem gleichen Muster wie oben folgen; hier ist der Vektor x unbekannt, während A und b bekannt sind. Es ist nicht immer garantiert, dass ein solches Gleichungssystem eine Lösung hat. Wenn A beispielsweise nur Nullen enthält, erfüllt eindeutig kein x $Ax = b$, wenn $b \neq 0$. Wir werden eine allgemeine Betrachtung darüber, wann eine Lösung existiert, auf unsere Diskussion linearer Löser in zukünftigen Kapiteln verschieben.

Eine Schlüsselinterpretation des Systems $Ax = b$ besteht darin, dass es folgende Aufgabe anspricht:

Schreiben Sie b als lineare Kombination der Spalten von A .

Warum? Erinnern Sie sich an §0.3.1, dass das Produkt Ax eine lineare Kombination der Spalten von A mit Gewichten kodiert, die in Elementen von x enthalten sind. Die Gleichung $Ax = b$ verlangt also, dass die Linearkombination Ax gleich dem gegebenen Vektor b ist. Angesichts dieser Interpretation definieren wir den Spaltenraum von A als den Raum der rechten Seiten b , für den das System eine Lösung hat:

Definition 0.9 (Spaltenraum). Der Spaltenraum einer Matrix $A \in \mathbb{R}^{m \times n}$ ist die Spanne der Spalten von A . Wir können schreiben als

$$\text{col } A = \{Ax : x \in \mathbb{R}^n\}.$$

Ein wichtiger Fall ist etwas einfacher zu betrachten. Angenommen, A ist quadratisch, also können wir $A \in \mathbb{R}^{n \times n}$ schreiben. Nehmen wir außerdem an, dass das System $Ax = b$ eine Lösung für alle Auswahlmöglichkeiten von b hat.

Die einzige Bedingung für b ist, dass es ein Mitglied von $\mathbb{R}^n$ ist. Daraus lässt sich schließen, dass die Spalten von $A \in \mathbb{R}^n$ überspannen.

Da in diesem Fall das lineare System immer lösbar ist, setzen wir die Standardbasis $e_1, \dots, e_n$ ein, um die Vektoren $x_1, \dots, x_n$, wobei $Ax_k = e_k$ für jedes k erfüllt ist. Dann können wir diese Ausdrücke „stapeln“, um Folgendes zu zeigen:

$$A \begin{pmatrix} | & | & \dots & | \\ x_1 & x_2 & \dots & x_n \\ | & | & \dots & | \end{pmatrix} = \begin{pmatrix} | & | & \dots & | \\ Ax_1 & Ax_2 & \dots & Ax_n \\ | & | & \dots & | \end{pmatrix} = \begin{pmatrix} | & | & \dots & | \\ e_1 & e_2 & \dots & e_n \\ | & | & \dots & | \end{pmatrix} = I_{n \times n},$$

wobei $I_{n \times n}$ die Identitätsmatrix aus Beispiel 0.11 ist. Wir nennen die Matrix mit den Spalten x_k die Umkehrung A^{-1} , was befriedigt

$$AA^{-1} = A^{-1}A = I_{n \times n}.$$

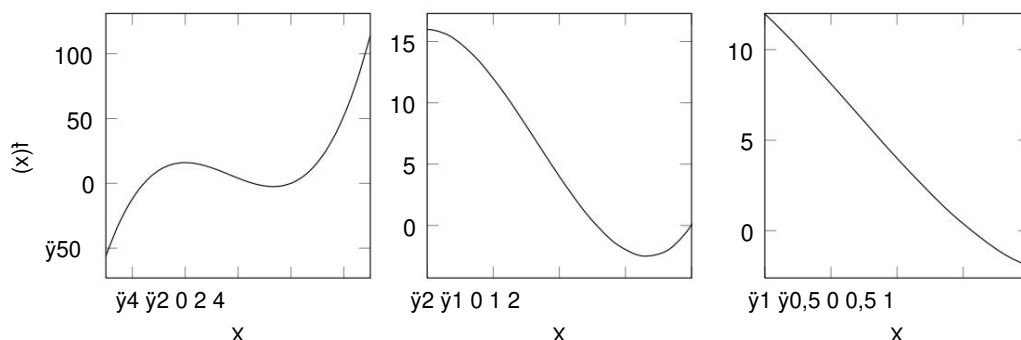


Abbildung 1: Je näher wir an $f(x) = x^3 + 2x + 4$ heranzoomen, desto mehr sieht es wie eine Linie aus.

Es ist auch leicht zu überprüfen, dass $(A^{-1})^{-1} = A$. Wenn eine solche Umkehrung existiert, ist es einfach, das System $Ax = b$ zu lösen. Insbesondere finden wir:

$$x = I_n x = (A^{-1}A)x = A^{-1}(Ax) = A^{-1}b$$

0,4 Nichtlinearität: Differentialrechnung

Während die Schönheit und Anwendbarkeit der linearen Algebra sie zu einem Hauptziel des Studiums macht, gibt es in der Natur viele Nichtlinearitäten und wir müssen oft Rechensysteme entwerfen, die mit dieser Tatsache des Lebens umgehen können. Denn auf der grundlegendsten Ebene macht das Quadrat in der berühmten Beziehung $E = mc^2$ sie einer linearen Analyse nicht zugänglich.

0.4.1 Differenzierung

Während viele Funktionen global nichtlinear sind, zeigen sie lokal ein lineares Verhalten. Diese Idee der „lokalen Linearität“ ist einer der Hauptmotive der Differentialrechnung. Abbildung 1 zeigt beispielsweise, dass eine glatte Funktion, wenn man nah genug heranzoomt, letztendlich wie eine Linie aussieht. Die Ableitung $f'(x)$ einer Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$ ist nichts anderes als die Steigung der Näherungsgeraden, berechnet durch Ermitteln der Steigung von Geraden durch immer näher an x liegende Punkte:

$$f'(x) = \lim_{y \rightarrow x} \frac{f(y) - f(x)}{y - x}$$

Wir können die lokale Linearität ausdrücken, indem wir $f(x + \Delta x) = f(x) + \Delta x \cdot f'(x) + O(\Delta x^2)$ schreiben.

Wenn die Funktion f mehrere Eingaben benötigt, kann sie wie folgt geschrieben werden: $f: \mathbb{R}^n \rightarrow \mathbb{R}$ für $x \in \mathbb{R}^n$; mit anderen Worten, jedem Punkt $x = (x_1, \dots, x_n)$ im n -dimensionalen Raum weist f eine einzelne Zahl $f(x_1, \dots, x_n)$ zu. Unsere Vorstellung von lokaler Linearität bricht hier etwas zusammen, da Linien eindimensionale Objekte sind. Die Festlegung aller Variablen bis auf eine reduziert sich jedoch auf den Fall der Einzelvariablenrechnung. Zum Beispiel könnten wir $g(t) = f(t, x_2, \dots, x_n)$ schreiben, wobei wir einfach die Konstanten $x_2, \dots, x_n$ festlegen. Dann ist $g(t)$ eine differenzierbare Funktion einer einzelnen Variablen. Natürlich hätten wir t in jeden der Eingabefelder für f einfügen können, daher machen wir im Allgemeinen die folgende Definition der partiellen Ableitung von f :

Definition 0,10 (partielle Ableitung). Die k -te partielle Ableitung von f , notiert als $\frac{\partial f}{\partial x_k}$, die f in $\frac{\partial f}{\partial x_k}$ wird durch different gegeben seiner k -ten Eingabevariablen bindet:

$$\frac{\partial f}{\partial x_k}(x_1, \dots, x_n) = \frac{D}{dt} f(x_1, \dots, x_k + t, x_{k+1}, \dots, x_n) \Big|_{t=0}$$

Die Notation „ $t=x_k$ “ sollte als „ausgewertet bei $t = x_k$ “ gelesen werden.

Beispiel 0,17 (Relativität). Die Beziehung $E = mc^2$ kann man sich als Funktion von m und c zu E vorstellen. Somit könnten wir $E(m, c) = mc^2$ schreiben und die Ableitungen ergeben

$$\begin{aligned} \frac{\partial E}{\partial m} &= c \\ \frac{\partial E}{\partial c} &= 2mc \end{aligned}$$

Unter Verwendung der Einvariablenrechnung können wir schreiben:

$$\begin{aligned} f(x + \vec{y}) &= f(x_1 + y_1, x_2 + y_2, \dots, x_n + y_n) \\ &= f(x_1, x_2 + y_2, \dots, x_n + y_n) + y_1 \frac{\partial f}{\partial x_1}(x_1, x_2 + y_2, \dots, x_n + y_n) + O(y_1^2) \text{ durch Einzelvariablenrechnung} \\ &= f(x_1, \dots, x_n) + \sum_{k=1}^n y_k \frac{\partial f}{\partial x_k}(x_1, \dots, x_n) + O(\|\vec{y}\|^2), \text{ indem Sie dies } n\text{-mal wiederholen} \\ &= f(x) + \vec{y} \cdot \nabla f(x) + O(\|\vec{y}\|^2) \end{aligned}$$

wobei wir den Gradienten von f als definieren

$$\nabla f = \left(\frac{\partial f}{\partial x_1}, \frac{\partial f}{\partial x_2}, \dots, \frac{\partial f}{\partial x_n} \right) \in \mathbb{R}^n$$

Aus dieser Beziehung ist leicht zu erkennen, dass f in jede Richtung v differenziert werden kann; wir können diese Ableitung $D_v f$ wie folgt auswerten:

$$D_v f(x) = \frac{D}{dt} f(x + tv) \Big|_{t=0} = \nabla f(x) \cdot v$$

Beispiel 0,18 (R2). Nehmen Sie $f(x, y) = x^2 y^3$. Dann,

$$\begin{aligned} \frac{\partial f}{\partial x} &= 2xy^3 \\ \frac{\partial f}{\partial y} &= 3x^2 y^2 \end{aligned}$$

Somit können wir schreiben: $\nabla f(x, y) = (2xy^3, 3x^2 y^2)$. Die Ableitung von f bei $(1, 2)$ in Richtung $(\vec{y}_1, 4)$ ist gegeben durch $(\vec{y}_1, 4) \cdot \nabla f(1, 2) = (\vec{y}_1, 4) \cdot (16, 12) = 32$.

Beispiel 0.19 (Lineare Funktionen). Es ist offensichtlich, aber erwähnenswert, dass der Gradient von $f(x) = a_1 x_1 + c_1 + \dots + a_n x_n + c_n$ ist.

Beispiel 0.20 (Quadratische Formen). Nehmen Sie eine beliebige Matrix $A \in \mathbb{R}^{n \times n}$, und definiere $f(x) = x^T A x$. Erweitern Sie f auf $\mathbb{R}^n$, die diese Funktion Element für Element zeigt

$$f(x) = \sum_{i,j} A_{ij} x_i x_j$$

Es lohnt sich, f zu erweitern und diesen Zusammenhang explizit zu überprüfen. Nehmen Sie ein $k \in \{1, \dots, n\}$. Dann können wir alle Terme heraustrennen, die x_k enthalten:

$$f(x) = A_{kk} x_k^2 + x_k \sum_{i=k}^n A_{ik} x_i + x_k \sum_{j=k}^n A_{kj} x_j + \sum_{i,j=k}^n A_{ij} x_i x_j$$

Mit dieser Faktorisierung ist es leicht zu erkennen

$$\begin{aligned} \frac{\partial f}{\partial x_k} &= 2A_{kk} x_k + \sum_{i=k}^n A_{ik} x_i + \sum_{j=k}^n A_{kj} x_j \\ &= \sum_{i=1}^n (A_{ik} + A_{ki}) x_i \end{aligned}$$

Diese Summe ist nichts anderes als die Definition der Matrix-Vektor-Multiplikation! So können wir schreiben

$$\frac{\partial f}{\partial x} = Ax + A^T x.$$

Wir haben von $f: \mathbb{R} \rightarrow \mathbb{R}$ auf $f: \mathbb{R}^n \rightarrow \mathbb{R}$ verallgemeinert. Um eine vollständige Allgemeingültigkeit zu erreichen, möchten wir $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$ betrachten. Mit anderen Worten: f nimmt n Zahlen auf und gibt m Zahlen aus. Glücklicherweise ist diese Erweiterung unkompliziert, da wir uns f als eine Sammlung einwertiger Funktionen $f_1, \dots, f_m$ vorstellen können. $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$ zu einem einzigen Vektor zusammengeschlagen. Das heißt, wir schreiben:

$$f(x) = \begin{pmatrix} f_1(x) \\ f_2(x) \\ \vdots \\ f_m(x) \end{pmatrix}$$

Jedes f_k kann wie zuvor differenziert werden, sodass wir am Ende eine Matrix partieller Ableitungen erhalten, die Jacobian von f genannt wird:

Definition 0.11 (Jakobianisch). Der Jacobi von $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$ ist die Matrix $Df \in \mathbb{R}^{m \times n}$ mit Einträgen

$$(Df)_{ij} = \frac{\partial f_i}{\partial x_j}.$$

Beispiel 0.21 (Einfache Funktion). Angenommen, $f(x, y) = (3x, x^2 y, x + y)$. Dann,

$$Df(x, y) = \begin{pmatrix} 3 & 0 & 0 \\ 2xy & x^2 & 0 \\ 1 & 1 & 1 \end{pmatrix}.$$

Stellen Sie sicher, dass Sie diese Berechnung manuell durchführen können.

Beispiel 0,22 (Matrixmultiplikation). Es überrascht nicht, dass der Jacobi-Wert von $f(x) = Ax$ für Matrix A durch $D f(x) = A$ gegeben ist.

Hier stoßen wir auf einen gemeinsamen Punkt der Verwirrung. Angenommen, eine Funktion verfügt über eine Vektoreingabe und eine Skalarausgabe, das heißt $f: \mathbb{R}^n \rightarrow \mathbb{R}$. Wir haben den Gradienten von f als Spaltenvektor definiert. Um diese Definition mit der des Jacobi-Werts in Einklang zu bringen, müssen wir schreiben

$$D f = \tilde{y} f.$$

0.4.2 Optimierung

Erinnern wir uns an Minima und Maxima von f aus der Einzelvariablenrechnung: $\mathbb{R} \rightarrow \mathbb{R}$ muss an Punkten x auftreten, die $f'(x) = 0$ erfüllen. Natürlich ist diese Bedingung eher notwendig als ausreichend: Es kann Punkte x mit $f'(x) = 0$ geben, die keine Maxima oder Minima sind. Allerdings kann das Finden solcher kritischer Punkte von f ein Schritt eines Funktionsminimierungsalgorithmus sein, solange der nächste Schritt sicherstellt, dass das resultierende x tatsächlich ein Minimum/Maximum ist.

Wenn $f: \mathbb{R}^n \rightarrow \mathbb{R}$ bei x minimiert oder maximiert ist, müssen wir sicherstellen, dass es keine einzige Richtung $\tilde{y}x$ von x gibt, in der f abnimmt bzw. zunimmt. Nach der Diskussion in §0.4.1 bedeutet dies, dass wir Punkte finden müssen, für die $\tilde{y} f = 0$ ist.

Beispiel 0.23 (Einfache Funktion). Angenommen, $f(x, y) = x^2 + 2xy + 4y^2$ und $\tilde{y}x \tilde{y}y = \tilde{y} f \tilde{y} f$. Dann ist $\tilde{y} f \tilde{y} f = 2x + 2y = 0$.
Somit erfüllen kritische Punkte von f :

$$\begin{aligned} 2x + 2y &= 0 \\ 2x + 8y &= 0 \end{aligned}$$

Offensichtlich wird dieses System bei $(x, y) = (0, 0)$ gelöst. Tatsächlich ist dies das Minimum von f , wie man deutlicher erkennen kann, wenn man $f(x, y) = (x + y)^2 + 3y^2$ schreibt.

Beispiel 0.24 (Quadratische Funktionen). Angenommen, $f(x) = x^T A x + b^T x + c$. Dann können wir aus den Beispielen im vorherigen Abschnitt schreiben: $\tilde{y} f(x) = (A + A)^T x + b$. Somit erfüllen kritische Punkte x von f $(A + A)^T x + b = 0$.

Im Gegensatz zur Einzelvariablenrechnung können wir bei der Berechnung von $\mathbb{R}^n$ Einschränkungen zu unserer Optimierung hinzufügen. Die allgemeinste Form eines solchen Problems sieht so aus:

Minimiere $f(x)$, so

dass $g(x) = 0$

Beispiel 0,25 (Rechteckflächen). Angenommen, ein Rechteck hat die Breite w und die Höhe h . Ein klassisches Geometrieproblem besteht darin, die Fläche mit einem festen Umfang zu maximieren 1:

$$\begin{aligned} &\text{maximieren } wh \\ &\text{so dass } 2w + 2h = 1 = 0 \end{aligned}$$

Wenn wir diese Einschränkung hinzufügen, können wir nicht mehr erwarten, dass kritische Punkte $\tilde{y} f(x) = 0$ erfüllen, da diese Punkte möglicherweise nicht $g(x) = 0$ erfüllen.

Angenommen, $g : \mathbb{R}^n \rightarrow \mathbb{R}$. Betrachten Sie die Punktmenge $S_0 = \{x : g(x) = 0\}$. Offensichtlich erfüllen zwei beliebige $x, y \in S_0$ die Beziehung $g(y) - g(x) = 0 - 0 = 0$. Angenommen, $y = x + \tilde{y}x$ für kleine $\tilde{y}x$. Dann ist $g(y) - g(x) = \tilde{y}g(x) \cdot \tilde{y}x + O(\tilde{y}x^2)$. Mit anderen Worten, wenn wir bei x beginnen und $g(x) = 0$ erfüllen, dann verschieben wir in der $\tilde{y}x$ -Richtung $\tilde{y}g(x) \cdot \tilde{y}x \approx 0$, um diese Beziehung weiterhin zu erfüllen.

Erinnern Sie sich nun daran, dass die Ableitung von f in der Richtung v bei x durch $\tilde{y}f \cdot v$ gegeben ist. Wenn x ein Minimum des obigen eingeschränkten Optimierungsproblems ist, sollte jede kleine Verschiebung von x zu $x + v$ einen Anstieg von $f(x)$ zu $f(x + v)$ bewirken. Da wir uns nur um Verschiebungen v kümmern, die die Einschränkung $g(x + v) = c$ wahren, wollen wir aus unserem obigen Argument $\tilde{y}f \cdot v = 0$ für alle v , die $\tilde{y}g(x) \cdot v = 0$ erfüllen. Mit anderen Worten, $\tilde{y}f$ und $\tilde{y}g$ müssen parallel sein, eine Bedingung, die wir als $\tilde{y}f = \tilde{y}g$ für ein $\tilde{y} \in \mathbb{R}$ schreiben können.

Definieren

$$\tilde{y}(x, \tilde{y}) = f(x) - \tilde{y}g(x).$$

Dann erfüllen kritische Punkte von $\tilde{y}$ ohne Einschränkungen:

$$\begin{aligned} 0 &= \frac{\tilde{y}\tilde{y}}{\tilde{y}\tilde{y}} = \tilde{y}g(x) \\ 0 &= \tilde{y}x\tilde{y} = \tilde{y}f(x) - \tilde{y}\tilde{y}g(x) \end{aligned}$$

Mit anderen Worten, kritische Punkte von $\tilde{y}$ erfüllen $g(x) = 0$ und $\tilde{y}f(x) = \tilde{y}\tilde{y}g(x)$, genau die Optimalitätsbedingungen, die wir abgeleitet haben!

Die Erweiterung auf multivariate Einschränkungen ergibt Folgendes:

Satz 0.1 (Methode der Lagrange-Multiplikatoren). Kritische Punkte des obigen eingeschränkten Optimierungsproblems sind uneingeschränkte kritische Punkte der Lagrange-Multiplikatorfunktion

$$\tilde{y}(x, \tilde{y}) = f(x) - \tilde{y}g(x),$$

sowohl bezüglich x als auch $\tilde{y}$.

Beispiel 0,26 (Bereich maximieren). In Fortsetzung von Beispiel 0.25 definieren wir die Lagrange-Multiplikatorfunktion $\tilde{y}(w, h, \tilde{y}) = wh - \tilde{y}(2w + 2h - 1)$. Differenzierend finden wir:

$$\begin{aligned} \tilde{y}w \frac{\tilde{y}\tilde{y}}{\tilde{y}\tilde{y}} &= h - 2\tilde{y} \quad 0 = \\ w - 2\tilde{y} \frac{\tilde{y}\tilde{y}}{\tilde{y}h} &= \\ \frac{\tilde{y}\tilde{y}}{\tilde{y}\tilde{y}} &= 1 - 2w - 2h \quad 0 = \end{aligned}$$

Kritische Punkte des Systems sind also erfüllt

$$\begin{array}{ccc} 0 & 1 & \tilde{y}^2 \\ \tilde{y} & 1 & 0 \tilde{y}^2 \\ \tilde{y} & 2 & 2 \quad 0 \end{array} \quad \begin{array}{ccc} w & & 0 \\ \tilde{y} & \tilde{y} & \tilde{y} \\ \tilde{y} & h & \tilde{y} \end{array} = \begin{array}{ccc} \tilde{y} & 0 & \tilde{y} \\ \tilde{y} & 1 & \tilde{y} \end{array}$$

Die Lösung des Systems ergibt $w = h = 1/4$ und $\tilde{y} = 1/8$. Mit anderen Worten: Bei einem festen Umfang ist das Rechteck mit der maximalen Fläche ein Quadrat.

Beispiel 0.27 (Eigenprobleme). Angenommen, A ist eine symmetrische positiv definite Matrix, was bedeutet, dass $A = A^T$ (Symmetrie) und $x^T A x > 0$ für alle $x \in \mathbb{R}^n \setminus \{0\}$ (positiv definit) ist. Oft möchten wir $x^T A x$ unter der Bedingung $x^T x = 1$ für eine gegebene Matrix $A \in \mathbb{R}^{n \times n}$ minimieren; Beachten Sie, dass ohne die Einschränkung das Minimum trivialerweise bei $x = 0$ auftritt. Wir definieren die Lagrange-Multiplikatorfunktion

$$\begin{aligned} \ddot{y}(x, \ddot{y}) &= x A x \ddot{y} \ddot{y}(x) \quad \frac{2}{2} \ddot{y} 1) \\ &= x A x \ddot{y} \ddot{y}(x \times \ddot{y} 1). \end{aligned}$$

Wenn wir nach x differenzieren, finden wir

$$0 = \ddot{y}x\ddot{y} = 2Ax \ddot{y} - 2\ddot{y}x$$

Mit anderen Worten, x ist ein Eigenvektor der Matrix A :

$$Ax = \ddot{y}x.$$

0,5 Probleme

Aufgabe 0.1. Sei $C^1(\mathbb{R})$ die Menge der Funktionen $f: \mathbb{R} \rightarrow \mathbb{R}$, die eine erste Ableitung $f'(x)$ zulassen. Warum ist $C^1(\mathbb{R})$ ein Vektorraum? Beweisen Sie, dass $C^1(\mathbb{R})$ die Dimension ∞ hat.

Aufgabe 0.2. Angenommen, die Zeilen von $A \in \mathbb{R}^{m \times n}$ seien durch die Transponierten von $r_1, \dots, r_m \in \mathbb{R}^n$ und die Spalten von $A \in \mathbb{R}^{m \times n}$ sind gegeben durch $c_1, \dots, c_n \in \mathbb{R}^m$. Das ist,

$$A = \begin{pmatrix} \ddot{y}_{r1} & & & \\ \ddot{y}_{r2} & & & \\ & \ddots & & \\ & & \ddot{y}_{rm} & \end{pmatrix} \quad \ddot{y} = \begin{pmatrix} \ddot{y}_1 & \ddot{y}_2 & \dots & \ddot{y}_n \\ \ddot{y}_1 & \ddot{y}_2 & \dots & \ddot{y}_n \\ \ddot{y}_1 & \ddot{y}_2 & \dots & \ddot{y}_n \end{pmatrix} =$$

Geben Sie Ausdrücke für die Elemente von AA und AA anhand dieser Vektoren an.

Aufgabe 0.3. Geben Sie ein lineares Gleichungssystem an, das durch Minima der Energie $f(x) = Ax \cdot y$ bezüglich $x \in \mathbb{R}^n$ und $y \in \mathbb{R}^m$ erfüllt wird, für $x \in \mathbb{R}^n$ und $b \in \mathbb{R}^m$. Dieses System wird „Normalgleichungen“ genannt und wird an anderer Stelle in diesen Anmerkungen erscheinen; Dennoch lohnt es sich, die Ableitung durchzuarbeiten und vollständig zu verstehen.

Aufgabe 0.4. Angenommen $A, B \in \mathbb{R}^{n \times n}$. Formulieren Sie eine Bedingung dafür, dass Vektoren $x, \tilde{y} \in \mathbb{R}^n$ kritische Punkte des $Ax^2 + Bx + \tilde{y}^T x$ Subjekts für $Bx = 1$ sind. Geben Sie außerdem eine alternative Form für die optimalen Werte von $Ax^2 + Bx + \tilde{y}^T x$ an.

Aufgabe 0.5. Legen Sie einen Vektor $\vec{y} \in \mathbb{R}^n$ fest und definieren Sie $f(x) = \vec{a} \cdot \vec{x}$. Geben Sie einen Ausdruck für das Maximum von $f(x)$ unter der Bedingung $\|\vec{x}\| = 1$ an.

Kapitel 1

Numerik und Fehleranalyse

Beim Studium der numerischen Analyse gehen wir vom Umgang mit Ints und Longs zu Floats und Doubles über. Dieser scheinbar harmlose Übergang bedeutet eine enorme Veränderung in der Art und Weise, wie wir über das Design und die Implementierung von Algorithmen nachdenken müssen. Im Gegensatz zu den Grundlagen diskreter Algorithmen können wir nicht mehr erwarten Unsere Algorithmen liefern in allen Fällen exakte Lösungen. „Big O“ und Vorgangszählung funktionieren nicht immer herrsche unumschränkt; Stattdessen sind wir gezwungen, selbst wenn wir die grundlegendsten Techniken verstehen, zu studieren der Kompromiss zwischen Timing, Näherungsfehler usw.

1.1 Speichern von Zahlen mit Bruchteilen

Denken Sie daran, dass Computer Daten im Allgemeinen im Binärformat speichern. Insbesondere ist jede Ziffer positiv Ganzzahl entspricht einer anderen Zweierpotenz. Beispielsweise könnten wir 463 in eine Binärzahl umwandeln anhand der folgenden Tabelle:

1	1	1 0	0 1			1	1	1
8 2	7 2	6 2	5 2	4 2	3 2	2 2	1 2	2 ⁰

Mit anderen Worten: Diese Notation kodiert die Tatsache, dass 463 in Zweierpotenzen zerlegt werden kann eindeutig als:

$$463 = 2^{8 \cdot 2} + 2^{7 \cdot 2} + 2^{6 \cdot 2} + 2^{5 \cdot 2} + 2^{4 \cdot 2} + 2^{3 \cdot 2} + 2^{2 \cdot 2} + 2^{1 \cdot 2} + 2^0$$

$$= 256 + 128 + 64 + 8 + 4 + 2 + 1$$

Abgesehen von Überlaufproblemen können alle positiven ganzen Zahlen in dieser Form mit einer endlichen Anzahl von geschrieben werden Ziffern. Auch negative Zahlen können auf diese Weise dargestellt werden, entweder durch Einführung eines Vorzeichens Bit oder mit dem „Zweierkomplement“-Trick.

Eine solche Zerlegung inspiriert zu einer einfachen Erweiterung auf Zahlen, die Brüche enthalten: einfach negative Zweierpotenzen einschließen. Beispielsweise ist die Zerlegung von 463,25 so einfach wie die Addition von zwei

Schlüssel:

1	1	1 0	0 1	3 2		1	1	1. 0	0 2	1
8 2	7 2	6 2	5 2	4 2		2 2	1 2		2 ⁻¹	2 ⁻²

Genau wie im Dezimalsystem ist jedoch auch die Darstellung gebrochener Teile von Zahlen auf diese Weise möglich nicht annähernd so brav wie die Darstellung von ganzen Zahlen. Zum Beispiel den Bruch 1/3 binär schreiben ergibt den Ausdruck:

$$\frac{1}{3} = 0,0101010101 \dots$$

Solche Beispiele zeigen, dass es auf allen Skalen Zahlen gibt, die nicht mit a dargestellt werden können endliche binäre Zeichenfolge. Tatsächlich sind Zahlen wie $\tilde{y} = 11,00100100001 \dots$ haben unendlich lange Entwicklungen unabhängig davon, welche (ganzzahlige) Basis Sie verwenden!

Aus diesem Grund müssen wir beim Entwerfen von Computersystemen, die mit **R** statt mit **Z** rechnen, Folgendes tun sind gezwungen, Näherungen für nahezu jede einigermaßen effiziente numerische Darstellung vorzunehmen.

Dies kann beim Codieren zu vielen Verwirrungen führen. Betrachten Sie beispielsweise das folgende C++

Ausschnitt:

```
doppeltes x = 1,0;
doppeltes y = x / 3,0;
if ( x == y * 3.0) cout << "Sie sind gleich! " " ";
sonst cout << "Sie sind NICHT gleich. " ;
```

Entgegen der Intuition gibt dieses Programm „Sie sind NICHT gleich“ aus. Warum? Die Definition von y macht eine Annäherung an $1/3$, da es nicht als abschließende Binärzeichenfolge geschrieben werden kann, Rundung zu einer nahegelegenen Zahl, die es darstellen kann. Somit multipliziert $y \cdot 3.0$ nicht mehr 3 mit $1/3$. Eine Möglichkeit, das Problem zu beheben Dieses Problem ist unten:

```
doppeltes x = 1,0;
doppeltes y = x / 3,0;
if ( fabs ( x - y * 3.0) < numeric_limits < double >:: epsilon ) cout << "Sie sind gleich! " " Sie sind NICHT gleich. " ;
```

Hier überprüfen wir, ob x und $y \cdot 3.0$ innerhalb einer gewissen Toleranz zueinander liegen, anstatt dies zu überprüfen exakte Gleichheit. Dies ist ein Beispiel für einen sehr wichtigen Punkt:

Der Operator `==` und seine Äquivalente **sollten selten, wenn überhaupt, für Bruchwerte verwendet werden.**

Stattdessen sollte eine gewisse Toleranz verwendet werden, um zu prüfen, ob die Zahlen gleich sind.

Natürlich gibt es hier einen Kompromiss: Die Größe der Toleranz definiert eine Grenze zwischen Gleichheit und „Nähe, aber nicht gleich“, die für eine bestimmte Anwendung sorgfältig ausgewählt werden müssen.

Im Folgenden betrachten wir einige Möglichkeiten zur Darstellung von Zahlen auf einem Computer.

1.1.1 Fixpunktdarstellungen

Die einfachste Möglichkeit zum Speichern von Brüchen ist das Hinzufügen eines festen Dezimalpunkts. Das heißt, wie in

Im obigen Beispiel stellen wir Werte dar, indem wir $0/1$ -Koeffizienten vor Zweierpotenzen speichern

Bereich von 2^k bis 2 für ein k , $\tilde{y} \in \mathbb{Z}$. Zum Beispiel alle nichtnegativen Werte dazwischen darstellen

0 und $127,75$ in Schritten von $1/4$ ist so einfach wie $k = 2$ und $\tilde{y} = 7$; In dieser Situation vertreten wir

Diese Werte werden mit 9 Binärziffern angegeben, von denen zwei nach dem Dezimalpunkt stehen.

Der Hauptvorteil dieser Darstellung besteht darin, dass nahezu alle arithmetischen Operationen möglich sind erfolgt mit den gleichen Algorithmen wie bei ganzen Zahlen. Das ist zum Beispiel leicht zu erkennen

$$a + b = (a \cdot 2^{-k} + b \cdot 2^{-k}) \cdot 2^{\tilde{y}k}.$$

Die Multiplikation unserer festen Darstellung mit 2^k garantiert, dass das Ergebnis ganzzahlig ist, also diese Beobachtung

zeigt im Wesentlichen, dass die Addition durch Ganzzahladdition im Wesentlichen durch „Ignorieren“ des Dezimalpunkts durchgeführt werden kann. Anstatt also spezielle Hardware zu verwenden, wird die bereits vorhandene Ganzzahl verwendet

Die Arithmetik-Logik-Einheit (ALU) führt Festkomma-Mathematik schnell aus.

Festkomma-Arithmetik ist zwar schnell, kann jedoch unter schwerwiegenden Präzisionsproblemen leiden. Insbesondere,

Es kommt häufig vor, dass die Ausgabe einer binären Operation wie Multiplikation oder Division erforderlich sein kann

mehr Bits als die Operanden. Nehmen wir zum Beispiel an, wir berücksichtigen eine Dezimalstelle für die Genauigkeit und

möchte das Produkt $1/2 \cdot 1/2 = 1/4$ ausführen. Wir schreiben $0,12 \times 0,12 = 0,012$, was auf 0 gekürzt wird. In diesem System ist es ziemlich einfach, Festkommazahlen auf sinnvolle Weise zu kombinieren und ein unvernünftiges Ergebnis zu erhalten.

Aufgrund dieser Nachteile enthalten die meisten großen Programmiersprachen standardmäßig keinen Festkomma-Dezimaldatentyp. Die Geschwindigkeit und Regelmäßigkeit der Festpunktarithmetik kann jedoch ein erheblicher Vorteil für Systeme sein, die Timing gegenüber Genauigkeit bevorzugen. Tatsächlich implementieren einige Grafikprozessoren (GPUs) der unteren Preisklasse nur diese Vorgänge, da für viele Grafikanwendungen eine Genauigkeit von wenigen Dezimalstellen ausreicht.

1.1.2 Gleitkommadarstellungen

Eine der vielen numerischen Herausforderungen beim Schreiben wissenschaftlicher Anwendungen ist die Vielfalt der möglichen Skalen. Nur Chemiker beschäftigen sich mit Werten irgendwo zwischen $9,11 \times 10^{31}$ und $6,022 \times 10^{23}$. Eine so harmlose Operation wie eine Änderung der Einheiten kann einen plötzlichen Übergang zwischen diesen Regimen verursachen: Dieselbe Beobachtung, geschrieben in Kilogramm pro Lichtjahr, sieht in Megatonnen erheblich anders aus pro Sekunde. Als numerische Analysten besteht unsere Aufgabe darin, Software zu schreiben, die einen reibungslosen Übergang zwischen diesen Skalen ermöglicht, ohne dem Kunden unnatürliche Einschränkungen seiner Techniken aufzuerlegen.

Hierzu sind einige Vorstellungen und Beobachtungen aus der Kunst der wissenschaftlichen Messung relevant Diskussion. Erstens ist offensichtlich eine der folgenden Darstellungen kompakter als die andere:

$$6,022 \times 10^{23} = 602, 200, 000, 000, 000, 000, 000$$

Darüber hinaus ist der Unterschied zwischen $6,022 \times 10^{23}$ und $6,022 \times 10^{23} + 9,11 \times 10^{31}$ mangels außergewöhnlicher wissenschaftlicher Ausrüstung vernachlässigbar. Eine Möglichkeit, zu dieser Schlussfolgerung zu gelangen, besteht darin, zu sagen, dass $6,022 \times 10^{23}$ nur eine dreistellige Genauigkeit hat und wahrscheinlich einen Bereich möglicher Messungen darstellt $[6,022 \times 10^{23} \pm \delta, 6,011 \times 10^{23} + \delta]$ für einige $\delta \ll 0,001 \times 10^{23}$.

Unsere erste Beobachtung war in der Lage, unsere Darstellung von $6,022 \times 10^{23}$ zu verdichten, indem wir sie in wissenschaftlicher Notation schrieben. Dieses Zahlensystem trennt die „interessanten“ Ziffern einer Zahl von ihrer Größenordnung, indem es sie in der Form $a \times 10^b$ für einige $a \approx 1$ und $b \approx Z$ schreibt. Wir nennen dieses Format die Gleitkommaform einer Zahl, weil im Gegensatz zur Festkomma-Einstellung in §1.1.1 „schwebt“ hier der Dezimalpunkt nach oben. Wir können Gleitkommasysteme mit wenigen Parametern beschreiben (CITE):

- Die Basis $\approx N$; Für die oben erläuterte wissenschaftliche Schreibweise ist die Basis 10
- Die Genauigkeit $p \approx N$, die die Anzahl der Stellen in der Dezimalentwicklung darstellt
- Der Bereich der Exponenten $[L, U]$, der die möglichen Werte von b darstellt

Eine solche Erweiterung sieht folgendermaßen aus:

$$\pm \underbrace{(d_0 + d_1 \cdot \tilde{r}^{-1} + d_2 \cdot \tilde{r}^{-2} + \dots + d_{p-1} \cdot \tilde{r}^{1-p})}_{\text{Mantisse}} \times \tilde{r}^{\text{Exponent}}$$

wobei jede Ziffer d_k im Bereich $[0, \tilde{r} - 1]$ liegt und $b \in [L, U]$.

Gleitkommadarstellungen haben eine merkwürdige Eigenschaft, die sich auf unerwartete Weise auf Software auswirken kann: Ihre Abstände sind ungleichmäßig. Zum Beispiel die Anzahl der darstellbaren Werte zwischen $\tilde{r}$ und $\tilde{r}^2$, obwohl das gleiche wie das zwischen $\tilde{r}$ und $\tilde{r}^2$ ist, ist mit dem gegebenen Zahlensystem möglich. Präzision zu verstehen, definieren wir die Maschinengenauigkeit $\tilde{m}$ als

kleinstes $\tilde{y}_m > 0$, so dass $1 + \tilde{y}_m$ darstellbar ist. Dann sind Zahlen wie $\tilde{y} + \tilde{y}_m$ im Zahlensystem nicht darstellbar, weil $\tilde{y}_m$ zu klein ist!

Der mit Abstand gebräuchlichste Standard zum Speichern von Gleitkommazahlen ist der IEEE 754-Standard. Dieser Standard spezifiziert mehrere Klassen von Gleitkommazahlen. Beispielsweise wird eine Gleitkommazahl mit doppelter Genauigkeit in der Basis $\tilde{y} = 2$ geschrieben (wie die meisten Zahlen auf dem Computer), mit einem einzelnen $\pm$ -Vorzeichenbit, 52 Ziffern für d und einem Exponentenbereich zwischen $\tilde{y}1022$ und 1023 . Der Standard legt auch fest, wie $\pm\tilde{y}$ und Werte wie NaN oder „keine Zahl“ gespeichert werden, die für die Ergebnisse von Berechnungen wie $10/0$ reserviert sind. Ein zusätzliches Maß an Präzision kann erreicht werden, indem man normalisierende Gleitkommawerte schreibt und davon ausgeht, dass die höchstwertige Ziffer $d_0 = 1$ ist, und sie nicht schreibt.

Der IEEE-Standard enthält auch vereinbarte Optionen für den Umgang mit der endlichen Anzahl von Werten, die bei einer endlichen Anzahl von Bits dargestellt werden können. Beispielsweise besteht eine gängige unverzerrte Strategie für Rundungsberechnungen darin, auf den nächsten Wert zu runden und auf gerade Werte zu binden. Dabei werden äquidistante Bindungen aufgehoben, indem auf den nächsten Gleitkommawert mit einem geraden niedrigstwertigen Bit (ganz rechts) gerundet wird. Beachten Sie, dass es viele gleichermaßen legitime Strategien zum Runden gibt; Die Auswahl einer einzigen Option garantiert, dass wissenschaftliche Software auf allen Client-Rechnern, die denselben Standard implementieren, identisch funktioniert.

1.1.3 Weitere exotische Optionen

In Zukunft gehen wir davon aus, dass Dezimalwerte im Gleitkommaformat gespeichert werden, sofern nicht anders angegeben. Dies bedeutet jedoch nicht, dass es keine anderen numerischen Systeme gibt, und für bestimmte Anwendungen könnte eine alternative Wahl erforderlich sein. Wir erkennen hier einige dieser Situationen an.

Der Aufwand, Toleranzen hinzuzufügen, um Rundungsfehlern Rechnung zu tragen, könnte für einige Anwendungen unakzeptabel sein. Diese Situation tritt bei Anwendungen der Computergeometrie auf, z. B. wenn der Unterschied zwischen nahezu und vollständig parallelen Linien schwierig zu unterscheiden sein kann. Eine Lösung könnte darin bestehen, Arithmetik mit beliebiger Genauigkeit zu verwenden, das heißt, Arithmetik ohne Rundung oder Fehler jeglicher Art zu implementieren.

Arithmetik mit beliebiger Genauigkeit erfordert eine spezielle Implementierung und eine sorgfältige Überlegung, welche Arten von Werten dargestellt werden müssen. Beispielsweise könnte es sein, dass für eine gegebene Anwendung rationale Zahlen $\mathbf{Q}$ ausreichend sind, die als Verhältnisse a/b für $a, b \in \mathbf{Z}$ geschrieben werden können.

Grundrechenarten können in $\mathbf{Q}$ ohne Präzisionsverlust ausgeführt werden. Es ist zum Beispiel leicht zu erkennen

$$\frac{A}{B} \times \frac{C}{D} = \frac{ac}{bd} \qquad \frac{A}{B} \div \frac{C}{D} = \frac{\text{Anzeige}}{\text{v. Chr.}}$$

Die Arithmetik in den rationalen Zahlen schließt die Existenz eines Quadratwurzeloperators aus, da Werte wie $\sqrt{2}$ irrational sind. Auch diese Darstellung ist nicht eindeutig, da $\sqrt{5a/b} = 5a/5b$.

In anderen Fällen kann es nützlich sein, Fehler einzuklammern, indem Werte neben Fehlerschätzungen als Paar $a, \tilde{y} \in \mathbf{R}$ dargestellt werden; Wir stellen uns das Paar $(a, \tilde{y})$ als den Bereich $a \pm \tilde{y}$ vor. Dann aktualisieren arithmetische Operationen nicht nur den Wert, sondern auch die Fehlerschätzung, wie in

$$(x \pm \tilde{y}_1) + (y \pm \tilde{y}_2) = (x + y) \pm (\tilde{y}_1 + \tilde{y}_2 + \text{Fehler}(x + y)),$$

wobei der letzte Term eine Schätzung des Fehlers darstellt, der durch die Addition von x und y verursacht wird.

1.2 Fehler verstehen

Mit Ausnahme der in §1.1.3 beschriebenen Systeme mit beliebiger Genauigkeit ist nahezu jede computergestützte Darstellung reeller Zahlen mit Bruchteilen gezwungen, Rundungs- und andere Näherungsschemata zu verwenden. Dieses Schema stellt eine von vielen Approximationsquellen dar, die typischerweise in numerischen Systemen anzutreffen sind:

- Der Kürzungsfehler ergibt sich aus der Tatsache, dass wir nur eine endliche Teilmenge aller möglichen Wertemengen in $\mathbb{R}$ darstellen können. Beispielsweise müssen wir lange oder unendliche Sequenzen über den Dezimalpunkt hinaus auf die Anzahl der Bits kürzen, die wir speichern möchten.
- Diskretisierungsfehler entstehen durch unsere computergestützten Anpassungen von Analysis, Physik und anderem Aspekten der kontinuierlichen Mathematik. Wir machen zum Beispiel eine Näherung

$$\frac{dy}{dx} \approx \frac{y(x + \bar{y}) - y(x)}{\bar{y}}.$$

Wir werden lernen, dass diese Näherung legitim und nützlich ist, aber abhängig von der Wahl von $\bar{y}$ möglicherweise nicht ganz korrekt ist.

- Modellierungsfehler entstehen durch unvollständige oder ungenaue Beschreibungen der Probleme, die wir lösen möchten. Beispielsweise könnte eine Simulation zur Vorhersage des Wetters in Deutschland den kollektiven Flügelschlag von Schmetterlingen in Malaysia vernachlässigen, obwohl die Luftverdrängung dieser Schmetterlinge ausreichen könnte, um die Wettermuster anderswo etwas zu stören.
- Empirische Konstantenfehler entstehen durch schlechte Darstellungen physikalischer oder mathematischer Konstanten. Zum Beispiel können wir $\bar{y}$ mithilfe einer Taylor-Folge berechnen, die wir vorzeitig beenden, und selbst Wissenschaftler kennen die Lichtgeschwindigkeit möglicherweise nicht einmal auf mehr als einige Stellen genau.
- Eingabefehler können von benutzergenerierten Näherungen von Parametern eines bestimmten Systems (und von Tippfehlern!) herrühren. Simulations- und numerische Techniken können zur Beantwortung von „Was-wäre-wenn“-Fragen eingesetzt werden, bei denen explorative Auswahlmöglichkeiten für Eingabeeinstellungen getroffen werden, um eine Vorstellung davon zu bekommen, wie sich ein System verhält.

Beispiel 1.1 (Computerphysik). Angenommen, wir entwerfen ein System zur Simulation von Planeten, die sich um die Erde drehen. Das System löst im Wesentlichen die Newtonsche Gleichung $F = ma$, indem es zeitlich vorwärts gerichtete Kräfte integriert. Beispiele für Fehlerquellen in diesem System könnten sein:

- Kürzungsfehler: Verwendung von IEEE-Gleitkomma zur Darstellung von Parametern und Ausgaben des Systems und Kürzung bei der Berechnung der Produktma
- Diskretisierungsfehler: Ersetzen der Beschleunigung a durch eine geteilte Differenz
- Modellierungsfehler: Es wurde versäumt, die Auswirkungen des Mondes auf die Bewegung der Erde innerhalb des Planeten zu simulieren System
- Empirischer Fehler: Die Masse des Jupiter wird nur vierstellig eingegeben
- Eingabefehler: Der Benutzer möchte möglicherweise die Kosten für die Verbringung von Müll in den Weltraum abschätzen, anstatt eine Wall-E-Ansammlung auf der Erde zu riskieren, kann aber nur die Menge an Müll schätzen, die die Regierung bereit ist, auf diese Weise abzuwerfen

1.2.1 Klassifizierungsfehler

Angeichts unserer vorherigen Diskussion könnte davon ausgegangen werden, dass die folgenden beiden Zahlen den gleichen potenziellen Fehler aufweisen:

$$1 \pm 0,01$$

$$105 \pm 0,01$$

Obwohl er die Größe des Bereichs $[1 - 0,01, 1 + 0,01]$ hat, scheint der Bereich $[105 - 0,01, 105 + 0,01]$ eine zuverlässigere Messung zu kodieren, da der Fehler 0,01 im Vergleich zu 105 viel kleiner ist als zu 1.

Die Unterscheidung zwischen diesen beiden Fehlerklassen wird durch die Unterscheidung zwischen absoluter Fehler und relativer Fehler beschrieben:

Definition 1.1 (Absoluter Fehler). Der absolute Fehler einer Messung ergibt sich aus der Differenz zwischen dem Näherungswert und dem zugrunde liegenden wahren Wert.

Definition 1.2 (Relativer Fehler). Der relative Fehler einer Messung ergibt sich aus dem absoluten Fehler dividiert durch den wahren Wert.

Eine Möglichkeit, zwischen diesen beiden Fehlerarten zu unterscheiden, ist die Verwendung von Einheiten gegenüber Prozentsätzen.

Beispiel 1.2 (Absoluter und relativer Fehler). Hier sind zwei äquivalente Aussagen in gegensätzlicher Form:

$$\text{Absolut: } 2 \text{ Zoll} \pm 0,02 \text{ Zoll}$$

$$\text{Relativ: } 2 \text{ Zoll} \pm 1 \%$$

In den meisten Anwendungen ist der wahre Wert unbekannt; Wäre dies nicht der Fall, wäre die Verwendung einer Näherung anstelle des wahren Werts möglicherweise ein zweifelhaftes Unterfangen. Es gibt zwei beliebte Möglichkeiten, dieses Problem zu lösen. Die erste besteht einfach darin, bei der Durchführung von Berechnungen konservativ zu sein: Nehmen Sie bei jedem Schritt die größtmögliche Fehlerschätzung vor und übertragen Sie diese Schätzungen nach Bedarf weiter. Solche konservativen Schätzungen sind insofern aussagekräftig, als wir, wenn sie klein sind, sehr sicher sein können, dass unsere Lösung nützlich ist.

Eine alternative Auflösung hängt davon ab, was Sie messen können. Nehmen wir zum Beispiel an, wir möchten die Gleichung $f(x) = 0$ für x bei gegebener Funktion $f: \mathbf{R} \rightarrow \mathbf{R}$ lösen. Wir wissen, dass es irgendwo eine Wurzel x_0 gibt, die genau $f(x_0) = 0$ erfüllt, aber wenn wir das wüssten, wären unsere Algorithmen gar nicht nötig. In der Praxis kann unser Rechensystem ein x_{est} ergeben, das $f(x_{\text{est}}) = \tilde{y}$ für ein $\tilde{y}$ mit $|\tilde{y}|$ erfüllt 1. Wir können die Differenz $x_0 - x_{\text{est}}$ möglicherweise nicht auswerten, da x_0 unbekannt ist. Andererseits können wir durch einfaches Auswerten von $f(x_{\text{est}}) - f(x_0) = f(x_{\text{est}})$ berechnen, da wir per Definition wissen, dass $f(x_0) = 0$ ist. Dieser Wert gibt einen Anhaltspunkt für den Fehler unserer Berechnung.

Dieses Beispiel veranschaulicht die Unterscheidung zwischen Vorwärts- und Rückwärtsfehler. Der durch eine Näherung verursachte Vorwärtsfehler definiert höchstwahrscheinlich unsere Intuition für die Fehleranalyse als die Differenz zwischen der angenäherten und der tatsächlichen Lösung, aber wie wir bereits besprochen haben, ist eine Berechnung nicht immer möglich. Der Rückwärtsfehler hat jedoch die Besonderheit, dass er berechenbar ist, aber nicht unser genaues Ziel bei der Lösung eines bestimmten Problems ist. Wir können unsere Definition und Interpretation des Rückwärtsfehlers anpassen, wenn wir uns verschiedenen Problemen nähern. Eine geeignete, wenn auch vage Definition lautet jedoch wie folgt:

Definition 1.3 (Rückwärtsfehler). Der Rückwärtsfehler wird durch den Betrag angegeben, um den sich eine Problemstellung ändern müsste, um eine gegebene Annäherung an ihre Lösung zu erreichen.

Diese Definition ist etwas unklar, daher veranschaulichen wir ihre Verwendung anhand einiger Beispiele.

Beispiel 1.3 (Lineare Systeme). Angenommen, wir möchten das $n \times n$ -lineare System $Ax = b$ lösen. Nennen Sie die wahre Lösung $x_0 \approx A^{-1}b$. In Wirklichkeit liefert unser System aufgrund von Kürzungsfehlern und anderen Problemen eine nahezu lösungsorientierte Lösung. Der Vorwärtsfehler dieser Näherung wird offensichtlich anhand der Differenz $x_{\text{est}} - x_0$ gemessen; In der Praxis ist es unmöglich, diesen Wert zu berechnen, da wir x_0 nicht kennen. In Wirklichkeit ist x_{est} die exakte Lösung eines modifizierten Systems $Ax = \text{best for best} \approx Ax_{\text{est}}$; Daher könnten wir den Rückwärtsfehler anhand der Differenz $b - Ax_{\text{est}}$ messen. Im Gegensatz zum Vorwärtsfehler lässt sich dieser Fehler leicht berechnen, ohne A umzukehren, und es ist leicht zu erkennen, dass x_{est} genau dann eine Lösung für das Problem ist, wenn der Rückwärtsfehler (oder Vorwärtsfehler) Null ist.

Beispiel 1.4 (Gleichungen lösen, CITE). Angenommen, wir schreiben eine Funktion zum Finden von Quadratwurzeln positiver Zahlen, die $\sqrt{2} \approx 1,4$ ausgibt. Der Vorwärtsfehler beträgt $|1,4 - 1,41421 \dots| \approx 0,0142$. Beachten Sie, dass $1,42 = 1,96$, der Rückwärtsfehler also $|1,96 - 2|$ beträgt $= 0,04$.

Die beiden obigen Beispiele zeigen ein größeres Muster, dass der Rückwärtsfehler viel einfacher zu berechnen ist als der Vorwärtsfehler. Beispielsweise erforderte die Auswertung des Vorwärtsfehlers in Beispiel 1.3 die Invertierung einer Matrix A , während die Auswertung des Rückwärtsfehlers nur eine Multiplikation mit A erforderte. In ähnlicher Weise ersetzte der Übergang vom Vorwärtsfehler zum Rückwärtsfehler in Beispiel 1.4 die Quadratwurzelberechnung durch Multiplikation.

1.2.2 Konditionierung, Stabilität und Genauigkeit

In fast jedem numerischen Problem bedeutet ein Null-Rückwärtsfehler einen Null-Vorwärtsfehler und umgekehrt.

Daher kann eine Software, die zur Lösung eines solchen Problems entwickelt wurde, sicherlich abbrechen, wenn sie feststellt, dass eine Kandidatenlösung keinen Rückwärtsfehler aufweist. Was aber, wenn der Rückwärtsfehler ungleich Null, aber klein ist?

Bedeutet dies unbedingt einen kleinen Vorwärtsfehler? Solche Fragen motivieren die Analyse der meisten numerischen Techniken, deren Ziel die Minimierung von Vorwärtsfehlern ist, in der Praxis jedoch nur Rückwärtsfehler messen können.

Wir möchten Änderungen im Rückwärtsfehler relativ zum Vorwärtsfehler analysieren, damit unsere Algorithmen mit Sicherheit sagen können, dass sie nur unter Verwendung des Rückwärtsfehlers akzeptable Lösungen hervorgebracht haben.

Dieser Zusammenhang kann für jedes Problem, das wir lösen möchten, unterschiedlich sein, daher nehmen wir am Ende folgende grobe Klassifizierung vor:

- Ein Problem ist unempfindlich oder gut konditioniert, wenn kleine Mengen an Rückwärtsfehlern kleine Mengen an Vorwärtsfehlern implizieren. Mit anderen Worten: Eine kleine Störung der Aussage eines gut konditionierten Problems führt nur zu einer kleinen Störung der wahren Lösung.
- Ein Problem ist sensibel oder schlecht konditioniert, wenn dies nicht der Fall ist.

Beispiel 1.5 ($ax = b$). Nehmen wir als Spielzeugbeispiel an, dass wir die Lösung $x_0 \approx b/a$ der linearen Gleichung $ax = b$ für $a, x, b \in \mathbb{R}$ finden wollen. Der Vorwärtsfehler einer möglichen Lösung x ist durch $x - x_0$ gegeben, der Rückwärtsfehler dagegen gegeben durch $b - ax = a(x - x_0)$. Wenn also $|a| \gg 1$ ist das Problem gut konditioniert, da kleine Werte des Rückwärtsfehlers $a(x - x_0)$ noch kleinere Werte von $x - x_0$ implizieren; im Gegensatz dazu, wenn $|a| \ll 1$ Das Problem ist schlecht konditioniert, denn selbst wenn $a(x - x_0)$ klein ist, kann der Vorwärtsfehler $x - x_0 \approx 1/a \cdot a(x - x_0)$ bei gegebenem $1/a$ -Faktor groß sein.

Wir definieren die Bedingungsahl als Maß für die Sensibilität eines Problems:

Definition 1.4 (Konditionsnummer). Die Bedingungsahl eines Problems ist das Verhältnis zwischen der Änderung seiner Lösung und der Änderung seiner Aussage bei kleinen Störungen. Alternativ ist es das Verhältnis von Vorwärts- zu Rückwärtsfehler bei kleinen Änderungen in der Problemstellung.

Beispiel 1.6 ($ax = b$, Teil zwei). Wenn wir Beispiel 1.5 fortsetzen, können wir die Bedingungsahl genau berechnen:

$$c = \frac{\text{Vorwärtsfehler}}{\text{Rückwärtsfehler}} = \frac{x - x_0}{a(x - x_0)} \approx \frac{1}{a}$$

Im Allgemeinen ist die Berechnung von Zustandszahlen fast so schwierig wie die Berechnung des Vorwärtsfehlers, und daher ist ihre genaue Berechnung wahrscheinlich unmöglich. Dennoch ist es oft möglich, Grenzen oder Näherungen für Bedingungsahlen zu finden, um zu beurteilen, wie vertrauenswürdig eine Lösung ist.

Beispiel 1.7 (Wurzelfindung). Angenommen, wir erhalten eine glatte Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ und möchten Werte x mit $f(x) = 0$ finden. Beachten Sie, dass $f(x + \delta) \approx f(x) + \delta f'(x)$. Somit könnte eine Annäherung an die Bedingungsahl zum Finden von x sein

$$\begin{aligned} \frac{\text{Änderung des Vorwärtsfehlers}}{\text{Änderung des Rückwärtsfehlers}} &= \frac{(x + \delta) - x f'(x)}{\delta f'(x)} \\ &= \frac{1}{f'(x)} \end{aligned}$$

Beachten Sie, dass diese Näherung mit der in Beispiel 1.6 übereinstimmt. Wenn wir x nicht kennen, können wir $f'(x)$ natürlich nicht auswerten, aber wenn wir die Form von f und die Grenze betrachten können, in der Nähe von x haben wir eine Vorstellung von der Worst-Case-Situation.

Vorwärts- und Rückwärtsfehler sind Maße für die Genauigkeit einer Lösung. Aus Gründen der wissenschaftlichen Wiederholbarkeit möchten wir außerdem stabile Algorithmen ableiten, die selbstkonsistente Lösungen für eine Klasse von Problemen liefern. Beispielsweise lohnt es sich möglicherweise nicht, einen Algorithmus zu implementieren, der nur in einem Fünftel der Fälle sehr genaue Lösungen generiert, selbst wenn wir die oben genannten Techniken anwenden können, um zu überprüfen, ob die Kandidatenlösung gut ist.

1.3 Praktische Aspekte

Die Unendlichkeit und Dichte der reellen Zahlen $\mathbb{R}$ kann bei der Implementierung numerischer Algorithmen schädliche Fehler verursachen. Während uns die in §1.2 vorgestellte Theorie der Fehleranalyse letztendlich dabei helfen wird, Garantien für die Qualität numerischer Techniken zu geben, die in zukünftigen Kapiteln vorgestellt werden, ist es erwähnenswert, bevor wir fortfahren, eine Reihe häufiger Fehler und „Fallstricke“ zu beachten, die bei Implementierungen numerischer Methoden vorkommen.

Wir haben den größten Übeltäter absichtlich zu Beginn in §1.1 eingeführt und wiederholen ihn zur wohlverdienten Hervorhebung

in größerer Schrift: Der Operator `==` und seine Äquivalente **sollten selten, wenn überhaupt, für Bruchwerte verwendet werden.**

Die Suche nach einem geeigneten Ersatz für $==$ und entsprechenden Bedingungen zum Beenden einer numerischen Methode hängt von der betrachteten Technik ab. Beispiel 1.3 zeigt, dass eine Methode zur Lösung von $Ax = b$ enden kann, wenn das Residuum $b - \tilde{y} Ax$ Null ist; Da wir nicht explizit prüfen wollen, ob $A\tilde{x} = b$ ist, prüfen Implementierungen in der Praxis $\text{norm}(A\tilde{x} - b) < \epsilon$. Beachten Sie, dass dieses Beispiel zwei Techniken demonstriert:

- Die Verwendung von Rückwärtsfehler $\tilde{y} Ax$ anstelle von Vorwärtsfehler, um zu bestimmen, wann beendet werden soll, Und
- Überprüfen, ob der Rückwärtsfehler kleiner als Epsilon ist, um das verbotene $==0$ -Prädikat zu vermeiden.

Der Parameter Epsilon hängt davon ab, wie genau die gewünschte Lösung sein muss, sowie von der Auflösung des verwendeten Zahlensystems.

Ein Programmierer, der diese Datentypen und Operationen verwendet, muss wachsam sein, wenn es darum geht, fehlerhafte numerische Operationen zu erkennen und zu verhindern. Betrachten Sie beispielsweise den folgenden Codeausschnitt zur Berechnung der Norm x_2 für einen Vektor $x \in \mathbb{R}^n$, der als 1D-Array $x[]$ dargestellt wird:

```
double normSquared = 0; for ( int i
= 0; i < n ; i ++ ) normSquared += x [ i ] * x
[ i ];
return sqrt ( normSquared );
```

Es ist leicht zu erkennen, dass in der Theorie $\min_i |x_i| \leq x_2 / \sqrt{n} \leq \max_i |x_i|$, das heißt, die Norm von x liegt in der Größenordnung der Werte der in x enthaltenen Elemente. In der Berechnung von x_2 verbirgt sich jedoch der Ausdruck $x[i] * x[i]$. Wenn es ein i gibt, so dass $x[i]$ in der Größenordnung von DOUBLE_MAX liegt, läuft das Produkt $x[i] * x[i]$ über, obwohl x_2 immer noch im Bereich der Doubles liegt. Ein solcher Überlauf lässt sich leicht verhindern, indem man x durch seinen Maximalwert dividiert, die Norm berechnet und zurückmultipliziert:

```
double maxElement = epsilon ; // Ich möchte nicht durch Null dividieren! for ( int i = 0; i < n ; i ++ )
    maxElement = max ( maxElement for ( int i , fabs ( x [ i ] ) );
i = 0; i < n ; i ++ ) {
    doppelt skaliert = x [ i ] / maxElement ; normSquared
    += skaliert * skaliert ;
}
return sqrt ( normSquared ) * maxElement;
```

Der Skalierungsfaktor beseitigt das Überlaufproblem, indem er sicherstellt, dass die summierten Elemente nicht größer als 1 sind.

Dieses kleine Beispiel zeigt einen von vielen Umständen, in denen ein einzelnes Codezeichen zu einem nicht offensichtlichen numerischen Problem führen kann. Während unsere Intuition aus der kontinuierlichen Mathematik ausreicht, um viele numerische Methoden zu generieren, müssen wir immer noch einmal überprüfen, ob die von uns verwendeten Operationen aus diskreter Sicht gültig sind.

1.3.1 Beispiel im größeren Maßstab: Summierung

Wir liefern nun ein Beispiel für ein numerisches Problem, das durch Arithmetik mit endlicher Genauigkeit verursacht wird und mit einem weniger offensichtlichen algorithmischen Trick gelöst werden kann.

Angenommen, wir möchten eine Liste von Gleitkommawerten zusammenfassen, eine Aufgabe, die von Systemen in der Buchhaltung, beim maschinellen Lernen, in der Grafik und in fast allen anderen Bereichen problemlos benötigt wird. Ein Codeausschnitt zur Lösung dieser Aufgabe, der zweifellos in unzähligen Anwendungen vorkommt, sieht wie folgt aus:

```
Doppelsumme = 0; for
( int i = 0; i < n ; i ++ ) sum += x [ i ];
```

Bevor wir fortfahren, ist es erwähnenswert, dass dies für die überwiegende Mehrheit der Anwendungen eine vollkommen stabile und sicherlich mathematisch gültige Technik ist.

Aber was kann schiefgehen? Betrachten Sie den Fall, in dem n groß ist und die meisten Werte $x[i]$ klein und positiv sind. Wenn i groß genug ist, ist in diesem Fall die Variablensumme im Verhältnis zu $x[i]$ groß. Letztendlich könnte sum so groß sein, dass $x[i]$ nur die niederwertigsten Bits von sum beeinflusst, und im Extremfall könnte sum so groß sein, dass das Hinzufügen von $x[i]$ überhaupt keine Auswirkung hat. Während ein einzelner solcher Fehler vielleicht keine große Sache ist, könnte die Gesamtwirkung, wenn dieser Fehler wiederholt gemacht wird, die Summe übersteigen, der wir überhaupt vertrauen können.

Um diesen Effekt mathematisch zu verstehen, nehmen wir an, dass die Berechnung einer Summe $a + b$ um bis zu $\tilde{y} > 0$ abweichen kann. Dann kann die obige Methode eindeutig einen Fehler in der Größenordnung von $n\tilde{y}$ induzieren, der linear mit n wächst. Wenn die meisten Elemente $x[i]$ tatsächlich in der Größenordnung von $\tilde{y}$ liegen, kann man der Summe überhaupt nicht trauen! Das ist ein enttäuschendes Ergebnis: Der Fehler kann so groß sein wie die Summe selbst.

Glücklicherweise gibt es viele Möglichkeiten, es besser zu machen. Wenn Sie beispielsweise die kleinsten Werte zuerst addieren, könnte dies dazu beitragen, deren kumulierten Effekt zu erklären. Paarweise Methoden, die Wertepaare von $x[]$ rekursiv addieren und eine Summe bilden, sind ebenfalls stabiler, es kann jedoch schwierig sein, sie so effizient zu implementieren wie die obige for-Schleife. Glücklicherweise bietet ein Algorithmus von Kahan (CITE) eine einfach zu implementierende Methode der „kompensierten Summierung“, die fast genauso schnell ist.

Die nützliche Beobachtung hier ist, dass wir tatsächlich eine Annäherung an den Fehler in der Summe während einer bestimmten Iteration verfolgen können. Betrachten Sie insbesondere den Ausdruck

$$((a + b) - a) - b.$$

Offensichtlich ist dieser Ausdruck algebraisch Null. Zahlenmäßig ist dies jedoch möglicherweise nicht der Fall.

Insbesondere kann die Summe $(a + b)$ das Ergebnis runden, um es im Bereich der Gleitkommawerte zu halten. Die Subtraktion von a und b nacheinander ergibt dann eine Näherung des durch diese Operation verursachten Fehlers; Beachten Sie, dass die Subtraktionsoperationen wahrscheinlich besser konditioniert sind, da der Übergang von großen zu kleinen Zahlen aufgrund der Aufhebung Ziffern an Genauigkeit hinzufügt.

Somit läuft die Kahan-Technik wie folgt ab:

```
Doppelsumme = 0;
Doppelkompensation = 0; // eine Annäherung an den Fehler

for ( int i = 0; i < n ; i ++ ) { // versuchen, sowohl
    zu x [ i ] als auch zum fehlenden Teil wieder hinzuzufügen double nextTerm = x [ i ] +
    Kompensation ;

    // Berechne das Summationsergebnis dieser Iteration double nextSum = sum +
    nextTerm ;

    // Berechnen Sie die Kompensation als Differenz zwischen dem Term, den Sie // hinzufügen wollten, und dem tatsächlichen
    Ergebnis. kompensation = nextTerm - (nextSum -
    sum);

    sum = nextSum ;
}
```

Anstatt einfach die Summe beizubehalten, verfolgen wir jetzt die Summe und führen eine Näherungskompensation der Differenz zwischen der Summe und dem gewünschten Wert durch. Bei jeder Iteration versuchen wir, etwas hinzuzufügen

Diese Kompensation zusätzlich zum aktuellen Element von $x[]$, und dann berechnen wir die Kompensation neu, um den letzten Fehler zu berücksichtigen.

Die Analyse des Kahan-Algorithmus erfordert eine sorgfältigere Buchführung als die Analyse der einfacheren inkrementellen Technik. Am Ende dieses Kapitels werden Sie eine Ableitung eines Fehlerausdrucks durchgehen. Das endgültige mathematische Ergebnis wird sein, dass sich der Fehler von $n\tilde{\gamma}$ auf $O(\tilde{\gamma} + n\tilde{\gamma})$ verbessert, erhebliche Verbesserung, wenn $0 < \tilde{\gamma}$.¹

Die Implementierung der Kahan-Summierung ist unkompliziert, verdoppelt aber die Operationsanzahl des resultierenden Programms mehr als. Auf diese Weise gibt es einen impliziten Kompromiss zwischen Geschwindigkeit und Genauigkeit, den Softwareentwickler eingehen müssen, wenn sie entscheiden, welche Technik am besten geeignet ist.

Im weiteren Sinne ist Kahans Algorithmus eine von mehreren Methoden, die die Anhäufung numerischer Fehler im Verlauf einer Berechnung, die aus mehr als einer Operation besteht, umgeht. Weitere Beispiele sind Bresenham's Algorithmus zur Rasterung von Linien (CITE), der nur ganzzahlige Arithmetik zum Zeichnen von Linien verwendet, selbst wenn sie Pixelzeilen und -spalten an nicht ganzzahligen Stellen schneiden, und die schnelle Fourier-Transformation (CITE), die effektiv die binäre Partition nutzt oben beschriebenen Summationstrick.

1.4 Probleme

Problem 1.1. Hier ist ein Problem.

Teil II

Lineare Algebra

Kapitel 2

Lineare Systeme und die LU Zersetzung

In Kapitel 0 haben wir verschiedene Situationen besprochen, in denen lineare Gleichungssysteme $Ax = b$ in der mathematischen Theorie und in der Praxis auftreten. In diesem Kapitel gehen wir das Grundproblem direkt an und untersuchen numerische Methoden zur Lösung solcher Systeme.

2.1 Lösbarkeit linearer Systeme

Wie in §0.3.3 eingeführt, sind lineare Gleichungssysteme wie

$$\begin{aligned} 3x + 2y &= 6 \\ 4x + y &= 7 \end{aligned}$$

kann in Matrixform geschrieben werden wie in

$$\begin{pmatrix} 3 & 2 \\ 4 & 1 \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 6 \\ 7 \end{pmatrix}.$$

Allgemeiner können wir Systeme der Form $Ax = b$ für $A \in \mathbb{R}^{m \times n}$ schreiben, $x \in \mathbb{R}^n$, und $b \in \mathbb{R}^m$.

Die Lösbarkeit des Systems muss in einen von drei Fällen fallen:

1. Das System lässt möglicherweise keine Lösungen zu, wie in:

$$\begin{pmatrix} 1 & 0 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 1 \\ 1 \end{pmatrix}.$$

Dieses System erfordert $x = 1$ und $x = 1$ gleichzeitig, offensichtlich zwei inkompatible Bedingungen.

2. Das System kann eine einzelne Lösung zulassen; Beispielsweise wird das System am Anfang dieses Abschnitts durch $(x, y) = (8/11, 45/11)$ gelöst.

3. Das System kann unendlich viele Lösungen zulassen, zB $0x = 0$. Beachten Sie, dass, wenn ein System $Ax = b$ zwei Lösungen x_0 und x_1 zulässt, es automatisch unendlich viele Lösungen der Form $cx_0 + (1 - c)x_1$ für $c \in \mathbb{R}$ hat, da

$$A(cx_0 + (1 - c)x_1) = cAx_0 + (1 - c)Ax_1 = cb + (1 - c)b = b.$$

Dieses lineare System würde als unterbestimmt bezeichnet.

Im Allgemeinen hängt die Lösbarkeit eines Systems sowohl von A als auch von b ab. Wenn wir zum Beispiel das obige unlösbare System so ändern, dass es ist

$$\begin{pmatrix} 1 & 0 \\ 1 & 0 \end{pmatrix} x = \begin{pmatrix} 1 \\ 1 \end{pmatrix},$$

dann bewegt sich das System von null zu unendlich vielen Lösungen der Form $(1, y)$. Tatsächlich lässt jede Matrix A eine rechte Seite b zu, so dass $Ax = b$ lösbar ist, da $Ax = 0$ immer durch $x = 0$ gelöst werden kann, unabhängig von A . Erinnern Sie sich an §0.3.1, dass die Matrix-Vektor-Multiplikation betrachtet werden kann als lineares Kombinieren der Spalten von A mit Gewichten von x . Daher können wir, wie in §0.3.3 erwähnt, erwarten, dass $Ax = b$ genau dann lösbar ist, wenn b im Spaltenraum von A liegt.

Im Großen und Ganzen hat die „Form“ der Matrix $A \in \mathbb{R}^{m \times n}$ erheblichen Einfluss auf die Lösbarkeit von $Ax = b$. Denken Sie daran, dass die Spalten von A m -dimensionale Vektoren sind. Betrachten Sie zunächst den Fall, dass A „breit“ ist, das heißt, wenn es mehr Spalten als Zeilen hat ($n > m$). Jede Spalte ist ein Vektor in $\mathbb{R}^m$, daher kann der Spaltenraum höchstens die Dimension m haben. Da $n > m$ ist, müssen die n Spalten von A dann linear abhängig sein; Dies impliziert, dass es ein $x_0 = 0$ gibt, so dass $Ax_0 = 0$ ist. Wenn wir dann $Ax = b$ nach x auflösen können, dann ist $A(x + yx_0) = Ax + yAx_0 = b + 0 = b$, was zeigt, dass es tatsächlich unendlich viele Lösungen x für $Ax = b$ gibt. Mit anderen Worten: Wir haben gezeigt, dass kein breites Matrixsystem eine eindeutige Lösung zulässt.

Wenn A „groß“ ist, das heißt, wenn es mehr Zeilen als Spalten hat ($m > n$), können die n Spalten $\mathbb{R}^m$ nicht überspannen. Somit existiert ein Vektor $b_0 \in \mathbb{R}^m \setminus \text{col } A$. Per Definition kann dieses $Ax = b_0$ für kein x erfüllen. Mit anderen Worten: Jede hohe Matrix A lässt Systeme $Ax = b_0$ zu, die nicht lösbar sind.

Beide oben genannten Situationen sind alles andere als günstig für den Entwurf numerischer Algorithmen. Wenn beispielsweise ein lineares System viele Lösungen zulässt, müssen wir zunächst definieren, welche Lösung vom Benutzer gewünscht wird: Schließlich ist die Lösung $x + 10^{31}x_0$ möglicherweise nicht so aussagekräftig wie $x \approx 0,1x_0$. Auf der anderen Seite ist im großen Fall jede kleine Störung $Ax = b + yb_0$ nicht mehr lösbar, selbst wenn $Ax = b$ für ein bestimmtes b lösbar ist; Diese Situation kann einfach deshalb auftreten, weil die im letzten Kapitel besprochenen Rundungsverfahren A und B überhaupt nur approximieren können.

Angesichts dieser Komplikationen werden wir in diesem Kapitel einige vereinfachende Annahmen treffen:

Wir werden nur das Quadrat $A \in \mathbb{R}^{n \times n}$ betrachten.

- Wir gehen davon aus, dass A nichtsingulär ist, das heißt, dass $Ax = b$ für jedes b lösbar ist.

Erinnern Sie sich an §0.3.3, dass die Nichtsingularitätsbedingung der Anforderung entspricht, dass die Spalten von A $\mathbb{R}^n$ überspannen und impliziert die Existenz einer Matrix A^{-1} erfüllt $A^{-1}A = AA^{-1} = I_{n \times n}$.

Eine irreführende Beobachtung ist die Annahme, dass das Lösen von $Ax = b$ gleichbedeutend mit der expliziten Berechnung der Matrix A^{-1} und der anschließenden Multiplikation ist, um $x = A^{-1}b$ zu finden. Obwohl diese Lösungsstrategie sicherlich gültig ist, kann sie einen erheblichen Overkill darstellen: Schließlich sind wir nur an den n Werten in x und nicht an n^2 Werten in A^{-1} interessiert. Selbst wenn A sich gut benimmt, kann es außerdem sein, dass A^{-1} geschrieben wird, was zu numerischen Schwierigkeiten, die umgangen werden können.

2.2 Ad-hoc-Lösungsstrategien

In der einführenden Algebra betrachten wir das Problem der Lösung eines linearen Gleichungssystems oft als eine Kunstform. Die Strategie besteht darin, Variablen zu „isolieren“ und iterativ alternative Formen des linearen Systems zu schreiben, bis jede Zeile die Form $x = \text{const}$ hat.

Bei der Formulierung systematischer Algorithmen zur Lösung linearer Systeme ist es aufschlussreich, ein Beispiel dieses Lösungsprozesses durchzuführen. Betrachten Sie das folgende System:

$$\begin{aligned}y - z &= 1 \\ 3x - y + z &= 4 \\ x + y - 2z &= 3\end{aligned}$$

Parallel dazu können wir eine Matrixversion dieses Systems pflegen. Anstatt $Ax = b$ explizit auszugeben, können wir etwas Platz sparen, indem wir die folgende „erweiterte“ Matrix schreiben:

$$\begin{array}{ccc|c}0 & 1 & -1 & 1 \\ -3 & 1 & 1 & 4 \\ 1 & 1 & -2 & 3\end{array}$$

Wir können lineare Systeme immer auf diese Weise schreiben, solange wir uns darauf einigen, dass die Variablen auf der linken Seite der Gleichungen und die Konstanten auf der rechten Seite bleiben.

Vielleicht möchten wir uns zuerst mit der Variablen x befassen. Der Einfachheit halber können wir die Zeilen des Systems so vertauschen, dass die dritte Gleichung zuerst erscheint:

$$\begin{aligned}x + y - 2z &= 3 \\ -z &= 1 - 3x \\ y + z &= 4\end{aligned} \quad \begin{array}{ccc|c}1 & 1 & -2 & 3 \\ -3 & 1 & 1 & 4 \\ 0 & 1 & -1 & 1\end{array}$$

Wir können dann die erste Gleichung in die dritte einsetzen, um den $3x$ -Term zu eliminieren. Dies ist dasselbe, als würde man die Beziehung $x + y - 2z = 3$ um -3 skalieren und das Ergebnis zur dritten Gleichung hinzufügen:

$$\begin{aligned}x + y - 2z &= 3 \\ -z &= 1 - 3x \\ + 7z &= 13\end{aligned} \quad \begin{array}{ccc|c}1 & 1 & -2 & 3 \\ -3 & 1 & 1 & 4 \\ 0 & 1 & -1 & 1\end{array}$$

Um y aus der dritten Gleichung zu eliminieren, können wir auf ähnliche Weise die zweite Gleichung mit 4 multiplizieren und das Ergebnis zur dritten addieren:

$$\begin{aligned}x + y - 2z &= 3 \\ -z &= 1 - 3x \\ &= 9\end{aligned} \quad \begin{array}{ccc|c}1 & 1 & -2 & 3 \\ -3 & 1 & 1 & 4 \\ 0 & 0 & 3 & 9\end{array}$$

Wir haben jetzt z isoliert! Daher können wir die dritte Zeile um $1/3$ skalieren, um einen Ausdruck für z zu erhalten:

$$\begin{aligned}x + y - 2z &= 3 \\ -z &= 1 - 3x \\ &= 3\end{aligned} \quad \begin{array}{ccc|c}1 & 1 & -2 & 3 \\ -3 & 1 & 1 & 4 \\ 0 & 0 & 1 & 3\end{array}$$

Jetzt können wir $z = 3$ in die anderen beiden Gleichungen einsetzen, um z aus allen bis auf die letzte Zeile zu entfernen:

$$\begin{array}{rcl} x + y = 3 & y = & \\ 2z = 3 & & \end{array} \quad \begin{array}{ccc|c} 1 & 1 & 0 & 3 \\ \check{y} & 0 & 1 & 0 & 2 \\ \check{y} & 0 & 0 & 1 & 3 \end{array} \quad \begin{array}{c} \check{y} \\ \check{y} \end{array}$$

Schließlich führen wir eine ähnliche Substitution für y durch, um die Lösung abzuschließen:

$$\begin{array}{rcl} x = 1 & y & \\ = 2 & z = 3 & \end{array} \quad \begin{array}{cccc|c} 1 & 0 & 0 & 1 & 0 & 1 & 0 \\ \check{y} & 2 & & & & & \check{y} \\ \check{y} & 0 & 0 & 1 & 3 & & \check{y} \end{array}$$

Dieses Beispiel mag etwas pedantisch sein, aber wenn wir auf unsere Strategie zurückblicken, ergeben sich einige wichtige Beobachtungen dazu, wie wir bei der Lösung linearer Systeme vorgehen können:

- Wir haben aufeinanderfolgende Systeme $A_i x = b_i$ geschrieben, die als Vereinfachungen des Originals angesehen werden können $Ax = b$.
- Wir haben das System gelöst, ohne jemals A aufzuschreiben $\check{y}^1$.
- Wir haben wiederholt einige einfache Operationen verwendet: Skalieren, Hinzufügen und Permutieren von Zeilen der System.
- Die gleichen Operationen wurden auf A und b angewendet. Wenn wir die k -te Zeile von A skaliert haben, haben wir auch skaliert die k -te Reihe von b . Wenn wir die Zeilen k und l von A hinzugefügt haben, haben wir die Zeilen k und l von b hinzugefügt.
- Weniger offensichtlich war, dass die Lösungsschritte nicht von b abhingen. Das heißt, dass alle unsere Lösungsentscheidungen dadurch motiviert waren, dass wir Werte ungleich Null in A eliminierten, anstatt die Werte in b ; b zu untersuchen, die einfach mit von der Partie waren.
- Wir schlossen ab, als wir das System auf $n \times n$ $Ax = b$ reduzierten.

Wir werden all diese allgemeinen Beobachtungen zur Lösung linearer Systeme zu unserem Vorteil nutzen.

2.3 Zeilenoperationen kodieren

Wenn wir uns das Beispiel in §2.2 noch einmal ansehen, sehen wir, dass die Lösung des linearen Systems eigentlich nur die Anwendung von drei Operationen erforderte: Permutation, Zeilenskalierung und Addition der Skalierung einer Zeile zur anderen. Tatsächlich können wir auf diese Weise jedes lineare System lösen, daher lohnt es sich, diese Operationen genauer zu untersuchen.

2.3.1 Permutation

Unser erster Schritt in §2.2 bestand darin, zwei Zeilen im Gleichungssystem zu vertauschen. Allgemeiner könnten wir die Zeilen einer Matrize mit den Zahlen $1, \dots, M$ indizieren. Dann kann eine Permutation dieser Zeilen als Funktion $\check{y}$ geschrieben werden, sodass die Liste $\check{y}(1), \dots, \check{y}(m)$ deckt den gleichen Satz von Indizes ab.

Wenn e_k die k -te Standardbasisfunktion ist, ist leicht zu erkennen, dass das Produkt $e_k A$ ergibt die k -te Zeile der Matrix A . Daher können wir diese Zeilenvektoren vertikal „stapeln“ oder verketteten, um eine Matrix zu erhalten, die die Zeilen gemäß $\tilde{y}$ permutiert :

$$P_{\tilde{y}} = \begin{pmatrix} \tilde{y}(1) \\ \tilde{y}(2) \\ \vdots \\ \tilde{y}(m) \end{pmatrix} A$$

Das heißt, das Produkt $P_{\tilde{y}} A$ ist genau die Matrix A mit gemäß $\tilde{y}$ permutierten Zeilen.

Beispiel 2.1 (Permutationsmatrizen). Angenommen, wir möchten Zeilen einer Matrix in 3×3 mit $\tilde{y}(1) = 2$, $\tilde{y}(2) = 3$ und $\tilde{y}(3) = 1$ permutieren. Nach unserer Formel hätten wir

$$P_{\tilde{y}} = \begin{pmatrix} 0 & 1 & 0 \\ \tilde{y} & 0 & 0 & 1 \\ \tilde{y} & 1 & 0 & 0 \end{pmatrix}$$

Aus Beispiel 2.1 können wir sehen, dass $P_{\tilde{y}}$ Einsen an Positionen der Form $(k, \tilde{y}(k))$ und an anderen Stellen Nullen hat. Das Paar $(k, \tilde{y}(k))$ stellt die Aussage dar: „Wir möchten, dass Zeile k der Ausgabematrix Zeile $\tilde{y}(k)$ aus der Eingabematrix ist.“ Basierend auf dieser Beschreibung einer Permutationsmatrix ist es leicht zu erkennen, dass die Umkehrung von $P_{\tilde{y}}$ die Transponierte P ist, da dadurch einfach die Rollen der Zeilen und Spalten vertauscht werden – jetzt nehmen wir Zeile $\tilde{y}(k)$ der Eingabe und fügen sie ein Zeile k der Ausgabe. Mit anderen Worten $P_{\tilde{y}}^{-1} = P_{\tilde{y}}^T$.

2.3.2 Zeilenskalierung

Angenommen, wir schreiben eine Liste von Konstanten $a_1, \dots, a_m$ und versuchen, die k -te Zeile einer Matrix A mit a_k zu skalieren. Dies wird offensichtlich durch die Anwendung der Skalierungsmatrix S_a erreicht, die gegeben ist durch:

$$S_a = \begin{pmatrix} a_1 & 0 & 0 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 & a_m \end{pmatrix}$$

Unter der Annahme, dass alle a_k die Bedingung $a_k \neq 0$ erfüllen, ist es einfach, S_a durch „Rückskalierung“ umzukehren :

$$S_a^{-1} = \begin{pmatrix} 1/a_1 & 0 & 0 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 & 1/a_m \end{pmatrix}$$

2.3.3 Eliminierung

Nehmen wir abschließend an, wir möchten Zeile k um eine Konstante c skalieren und das Ergebnis zu Zeile hinzufügen. Dieser Vorgang mag weniger natürlich erscheinen als die beiden vorherigen, ist aber tatsächlich recht praktisch: Es ist der einzige, den wir durchführen

müssen Gleichungen aus verschiedenen Zeilen des linearen Systems kombinieren! Wir werden diese Operation mithilfe einer „Eliminationsmatrix“ M realisieren, sodass das Produkt MA diese Operation auf Matrix A anwendet.

Denken Sie daran, dass das Produkt $e_k A$ wählt die k -te Zeile von A aus. Die Vormultiplikation mit e ergibt dann eine Matrix E , die Null ist, außer dass die k -te Zeile gleich der k -ten von A ist.

Beispiel 2.2 (Konstruktion der Eliminierungsmatrix). Nehmen

$$A = \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{pmatrix}$$

Angenommen, wir möchten die dritte Zeile von A (3×3) isolieren und in Zeile zwei verschieben. Wie oben erläutert, wird dieser Vorgang durch Schreiben erreicht:

$$\begin{aligned} e_2 e_3 A &= \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{pmatrix} \\ &= \begin{pmatrix} 1 & 2 & 3 \\ 0 & 0 & 0 \\ 7 & 8 & 9 \end{pmatrix} \\ &= \begin{pmatrix} 1 & 2 & 3 \\ 7 & 8 & 9 \\ 0 & 0 & 0 \end{pmatrix} \end{aligned}$$

Natürlich haben wir oben von rechts nach links multipliziert, hätten das Produkt aber genauso gut als $(e_2 e_3)$ gruppieren können. Der Aufbau dieses Produkts ist leicht zu erkennen:

$$e_2 e_3 = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 0 \end{pmatrix} = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 0 \end{pmatrix}$$

Es ist uns gelungen, Zeile k zu isolieren und in Zeile zu verschieben. Unsere ursprüngliche Eliminierungsoperation wollte c mal Zeile k zu Zeile $A = k$ ($1 \times n + c e_k$) addieren, was wir nun als Summe $A + c e_k$ erreichen können.

Beispiel 2.3 (Ein System lösen). Wir können nun jede unserer Operationen aus Abschnitt 2.2 mithilfe der Matrizen kodieren, die wir oben erstellt haben:

1. Permutieren Sie die Zeilen, um die dritte Gleichung in die erste Zeile zu verschieben:

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

2. Skalieren Sie Zeile eins um -3 und addieren Sie das Ergebnis zu Zeile drei: $E_1 = I_3 \times 3 - 3e_1$ 1

3. Skalieren Sie Zeile zwei mit 4 und addieren Sie das Ergebnis zu Zeile drei: $E_2 = I_3 \times 3 + 4e_2$ 2

4. Skalieren Sie Zeile drei um $1/3$: $S = \text{diag}(1, 1, 1/3)$

5. Skalieren Sie Zeile drei um 2 und fügen Sie sie zur ersten Zeile hinzu: $E_3 = I_3 \times 3 + 2e_1 e_3$

6. Fügen Sie Zeile drei zu Zeile zwei hinzu: $E_4 = I_3 \times 3 + e_2 e_3$

7. Skalieren Sie Zeile drei um -1 und addieren Sie das Ergebnis zu Zeile eins: $E_5 = I_3 \times 3 - e_1 e_3$

Somit erfüllt die Umkehrung von A in Abschnitt 2.2

$$A^{-1} = E_5 E_4 E_3 E_2 E_1 P.$$

Stellen Sie sicher, dass Sie verstehen, warum diese Matrizen in umgekehrter Reihenfolge erscheinen!

2.4 Gaußsche Eliminierung

Die in Abschnitt 2.2 gewählte Schrittfolge war keineswegs eindeutig: Es gibt viele verschiedene Wege, die zur Lösung von $Ax = b$ führen können. Unsere Schritte folgten jedoch der Strategie der Gaußschen Eliminierung, einem berühmten Algorithmus zur Lösung linearer Gleichungssysteme.

Nehmen wir allgemeiner an, dass unser System die folgende „Form“ hat:

$$A \vec{b} = \begin{array}{c|c} \vec{y} & \begin{array}{cccc} \times & \times & \times & \times \\ \times & \times & \times & \times \\ \times & \times & \times & \times \\ \times & \times & \times & \times \end{array} \\ \hline & \vec{y} \end{array}$$

Der Algorithmus läuft in den unten beschriebenen Phasen ab.

2.4.1 Stürmerwechsel

Betrachten Sie das obere linke Element unserer Matrix:

$$A \vec{b} = \begin{array}{c|c} \vec{y} & \begin{array}{cccc} \times & \times & \times & \times \\ \times & \times & \times & \times \\ \times & \times & \times & \times \\ \times & \times & \times & \times \end{array} \\ \hline & \vec{y} \end{array}$$

Wir nennen dieses Element unseren ersten Pivot und gehen davon aus, dass es ungleich Null ist; Wenn es Null ist, können wir Zeilen so permutieren, dass dies nicht der Fall ist. Wir wenden zunächst eine Skalierungsmatrix an, sodass der Pivot gleich eins ist:

$$\vec{y} \begin{array}{c|c} \begin{array}{c} 1 \\ \times \\ \times \\ \times \end{array} & \begin{array}{cccc} \times & \times & \times & \times \\ \times & \times & \times & \times \\ \times & \times & \times & \times \\ \times & \times & \times & \times \end{array} \\ \hline & \vec{y} \end{array}$$

Jetzt verwenden wir die Zeile, die den Pivot enthält, um alle anderen Werte darunter in derselben Spalte zu entfernen:

$$\vec{y} \begin{array}{c|c} \begin{array}{c} 1 \\ 0 \\ \times \end{array} & \begin{array}{cccc} \times & \times & \times & 0 \\ \times & \times & \times & 0 \\ \times & \times & \times & \times \end{array} \\ \hline & \vec{y} \end{array}$$

Wir verschieben nun unseren Pivot in die nächste Zeile und wiederholen eine ähnliche Reihe von Vorgängen:

$$\begin{array}{cccc|c} & 1 & x & x & x & \\ \ddot{y} & 0 & 1 & x & x & \ddot{y} \\ & 0 & 0 & x & x & \\ \text{???} & 0 & 0 & x & x & \text{???} \end{array}$$

Beachten Sie, dass hier etwas Schönes passiert. Nachdem der erste Pivot aus allen anderen Zeilen entfernt wurde, ist die erste Spalte unter Zeile 1 eine Null. Das bedeutet, dass wir problemlos Vielfache von Zeile zwei zu den darunter liegenden Zeilen addieren können, ohne dass sich dies auf die Nullen in Spalte eins auswirkt.

Wir wiederholen diesen Vorgang, bis die Matrix die obere Dreiecksform annimmt:

$$\begin{array}{cccc|c} & 1 & x & x & x & \\ \ddot{y} & 0 & 1 & x & x & 0 & 1 & x & \ddot{y} \\ & x & & & & & & & \\ \text{???} & 0 & 0 & 0 & 1 & x & & & \text{???} \end{array}$$

2.4.2 Rückenwechsel

Das Entfernen der verbleibenden x aus dem System ist nun ein unkomplizierter Prozess. Nun gehen wir in umgekehrter Reihenreihenfolge vor und eliminieren rückwärts. Nach der ersten Reihe von Rückersetzungsschritten erhalten wir beispielsweise die folgende Form:

$$\begin{array}{cccc|c} & 1 & x & x & 0 & x & 0 & 1 & x & 0 & \\ \ddot{y} & x & 0 & 0 & 1 & 0 & x & & & & \ddot{y} \\ & 0 & 0 & 0 & 1 & x & & & & & \text{???} \end{array}$$

In ähnlicher Weise ergibt die zweite Iteration:

$$\begin{array}{cccc|c} & 1 & x & 0 & 0 & x & 0 & 1 & 0 & 0 & \\ \ddot{y} & x & & & & & & & & & \ddot{y} \\ & 0 & 0 & 1 & 0 & x & 0 & 0 & 0 & 1 & \\ \text{???} & x & & & & & & & & & \text{???} \end{array}$$

Nach unserem letzten Eliminierungsschritt bleibt uns die gewünschte Form übrig:

$$\begin{array}{cccc|c} & 1 & 0 & 0 & 0 & x & 0 & 1 & 0 & 0 & \\ \ddot{y} & 0 & x & 0 & 0 & 1 & 0 & x & & & \ddot{y} \\ & 0 & 0 & 0 & 1 & x & & & & & \text{???} \end{array}$$

Die rechte Seite ist nun die Lösung des linearen Systems $Ax = b$.

2.4.3 Analyse der Gaußschen Eliminierung

Jede Zeilenoperation bei der Gaußschen Eliminierung – Skalierung, Eliminierung und Vertauschen zweier Zeilen – nimmt offensichtlich $O(n)$ Zeit in Anspruch, da Sie über alle n Elemente einer Zeile iterieren müssen (bzw

zwei) von A. Sobald wir einen Pivot ausgewählt haben, müssen wir n Vorwärts- oder Rückwärtsersetzungen in den Zeilen unter bzw. über diesem Drehpunkt; Dies bedeutet, dass die Arbeit für einen einzelnen Pivot insgesamt $O(n)^2$ beträgt vornehmen. Insgesamt wählen wir einen Pivot pro Zeile und fügen einen Endfaktor von n hinzu. Daher ist es leicht, diese Gaußsche Funktion zu erkennen. Eliminierungsläufe in $O(n^3)$ Zeit.

Eine Entscheidung, die während der Gaußschen Eliminierung stattfindet und die wir nicht besprochen haben, ist die Wahl der Drehpunkte. Denken Sie daran, dass wir Zeilen des linearen Systems nach Belieben permutieren können, bevor wir eine Rück- oder Vorwärtssubstitution durchführen. Diese Operation ist notwendig, um mit allen möglichen Matrizen A umgehen zu können. Überlegen Sie beispielsweise, was passieren würde, wenn wir im Folgenden keine Pivotierung verwenden würden Matrix:

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

Beachten Sie, dass das eingekreiste Element genau Null ist, daher können wir nicht damit rechnen, Zeile eins durch eine beliebige Zahl zu dividieren, um diese 0 durch eine 1 zu ersetzen. Das bedeutet nicht, dass das System nicht lösbar ist, sondern nur, dass wir eine Pivotierung durchführen müssen, die durch Vertauschen erreicht wird erste und zweite Reihe, um einen Wert ungleich Null in diesen Steckplatz einzufügen.

Nehmen wir allgemeiner an, dass unsere Matrix wie folgt aussieht:

$$A = \begin{pmatrix} \tilde{y} & 1 \\ 1 & 0 \end{pmatrix},$$

wobei $0 < \tilde{y}$ 1. Wenn wir nicht schwenken, ergibt die erste Iteration der Gaußschen Eliminierung:

$$A^- = \begin{pmatrix} 1 & 1/\tilde{y} & 0 \\ \tilde{y} & 1/\tilde{y} \end{pmatrix},$$

Wir haben eine Matrix A, die fast wie eine Permutationsmatrix aussieht (tatsächlich eine sehr einfache $\tilde{y}^{-1} \tilde{y} A$, A Möglichkeit, das System zu lösen!) in ein System mit potenziell **großen** Werten $1/\tilde{y}$ umgewandelt.

Dieses Beispiel zeigt, dass es Fälle gibt, in denen wir möglicherweise einen Wechsel wünschen, auch wenn dies streng genommen nicht notwendig ist. Da wir mit dem Kehrwert des Pivotwerts skalieren, besteht die numerisch stabilste Option eindeutig darin, einen großen Pivotwert zu haben: Kleine Pivotwerte haben große Kehrwerte, wodurch Zahlen in Regimen, die wahrscheinlich an Präzision verlieren, auf große Werte skaliert werden. Es gibt zwei bekannte Pivot-Strategien:

1. Bei der teilweisen Pivotierung wird die aktuelle Spalte durchsucht und die Zeilen der Matrix so permutiert der größte absolute Wert erscheint auf der Diagonale.
2. Beim vollständigen Pivotieren wird die **gesamte** Matrix durchlaufen und sowohl Zeilen als auch Spalten vertauscht, um den größtmöglichen Wert auf der Diagonale zu erhalten. Beachten Sie, dass das Permutieren von Spalten einer Matrix eine gültige Operation ist: Sie entspricht dem Ändern der Beschriftung der Variablen im System oder dem Nachmultiplizieren von A mit einer Permutation.

Beispiel 2.4 (Pivotieren). Angenommen, nach der ersten Iteration der Gaußschen Eliminierung erhalten wir die folgende Matrix:

$$\begin{pmatrix} 1 & 10 & \tilde{y} & 10 \\ \tilde{y} & 0 & 0,1 & 9 \\ \tilde{y}\tilde{y} & 0 & 4 & 6.2 \end{pmatrix} \tilde{y}$$

Wenn wir eine teilweise Pivotierung implementieren, schauen wir nur in der zweiten Spalte und vertauschen die zweite und dritte Zeile; Beachten Sie, dass wir die 10 in der ersten Zeile belassen, da diese bereits vom Algorithmus besucht wurde:

$$\begin{array}{cc|c} 1 & 10 & \ddot{y}_{10} \\ \ddot{y} & 0 & 4,6,2 \\ \ddot{y} & 0 & 0,1,9 \end{array}$$

Wenn wir das vollständige Pivotieren implementieren, verschieben wir die 9:

$$\begin{array}{cc|c} 1 & \ddot{y}_{10} & 10 \\ \ddot{y} & 0 & 9,0,1 \\ \ddot{y} & 0 & 6,2,4 \end{array}$$

Offensichtlich liefert die vollständige Pivotierung die bestmöglichen Zahlen, allerdings ist die Suche nach großen Elementen in der Matrix teurer.

2,5 LU-Faktorisierung

Es kommt oft vor, dass wir eine Folge von Problemen $Ax_1 = b_1$, $Ax_2 = b_2$, ... lösen möchten. . . . Wie wir bereits besprochen haben, hängen die Schritte der Gaußschen Eliminierung zur Lösung von $Ax = bk$ hauptsächlich von der Struktur von A und nicht von den Werten in einem bestimmten b_k ab. Da A hier konstant gehalten wird, möchten wir uns vielleicht an die Schritte „erinnern“, die wir zur Lösung des Systems unternommen haben, damit wir nicht jedes Mal, wenn uns ein neues Problem präsentiert wird, bei Null anfangen müssen.

Dies bestätigt den Verdacht, dass wir einige der $O(n)$ bewegen können³) für die Gaußsche Eliminierung in der Vorberechnungszeit erinnern wir uns an das obere Dreieckssystem, das sich nach der Vorwärtssubstitutionsstufe ergibt:

$$\begin{array}{cccc|c} 1 & x & x & x & \ddot{y} \\ \ddot{y} & 0 & 1 & x & x & 0 & 0 & 1 & x & \ddot{y} \\ & x & & & & & & & & \\ \text{???} & 0 & 0 & 0 & 1 & x & & & & \text{???} \end{array}$$

Tatsächlich benötigt die Lösung dieses Systems durch Rücksubstitution ²⁾ Zeit! Warum? Zurück Substitution nur $O(n)$, in diesem Fall ist dies dank der Struktur der Nullstellen im System viel einfacher. Beispielsweise erhalten wir in der ersten Reihe von Rücksubstitutionen die folgende Matrix:

$$\begin{array}{cccc|c} 1 & x & x & 0 & x & 0 & x & 0 & x & 0 & \ddot{y} \\ \ddot{y} & 1 & 0 & x & 1 & & & & & & \ddot{y} \\ & & 0 & & & & & & & & \\ \text{???} & 0 & 0 & 0 & 1 & x & & & & & \text{???} \end{array}$$

Da wir wissen, dass die (eingekreisten) Werte links vom Pivot konstruktionsbedingt Null sind, müssen wir sie nicht explizit kopieren. Somit dauerte dieser Schritt nur $O(n)$ Zeit und nicht $O(n^2)$ für die Vorwärtssubstitution.

Nun führt unser nächster Pivot eine ähnliche Ersetzung durch:

$$\begin{array}{cccc|c} 1 & x & 0 & 0 & x & \ddot{y} \\ \ddot{y} & 0 & 1 & 0 & 0 & x & & & & & \ddot{y} \\ & 0 & 0 & 1 & 0 & x & 0 & 0 & 1 & x & \\ \text{???} & 0 & 0 & 0 & 1 & x & & & & & \text{???} \end{array}$$

Auch hier müssen die Nullen auf beiden Seiten der 1 nicht explizit kopiert werden.

So haben wir herausgefunden:

Überwachung. Während die Gaußsche Eliminierung $O(n)$ benötigt³) Zeit, das Lösen von Dreieckssystemen dauert $O(n^2)$ Zeit.

2.5.1 Konstruktion der Faktorisierung

Erinnern Sie sich an §2.3, dass alle Operationen der Gaußschen Eliminierung als Vormultiplikation von $Ax = b$ mit verschiedenen Matrizen M betrachtet werden können, um ein einfacheres System $(MA)x = Mb$ zu erhalten. Wie wir in Beispiel 2.3 gezeigt haben, stellt von diesem Standpunkt aus jeder Schritt der Gaußschen Eliminierung ein System $(M_k \cdot \cdot \cdot M_2 M_1 A)x = M_k \cdot \cdot \cdot M_2 M_1 b$ dar. Natürlich ist die explizite Speicherung dieser Matrizen M_k als $n \times n$ Objekte übertrieben, aber wenn man diese Interpretation aus theoretischer Sicht im Hinterkopf behält, werden viele unserer Berechnungen vereinfacht.

Nach der Vorwärtssubstitutionsphase der Gaußschen Eliminierung bleibt uns eine obere Dreiecksmatrix übrig, die wir $U \in \mathbb{R}^{n \times n}$ nennen können. Aus der Perspektive der Matrixmultiplikation ist das möglich schreiben:

$$M_k \cdot \cdot \cdot M_1 A = U,$$

oder gleichwertig,

$$\begin{aligned} A &= (M_k \cdot \cdot \cdot M_1)^{-1} U \\ &= (M_1^{-1} \cdot \cdot \cdot M_k^{-1}) U \\ &\stackrel{\text{wenn wir die Definition } L \stackrel{\text{def}}{=} M_1^{-1} \cdot \cdot \cdot M_k^{-1} \text{ machen}}{=} LU, \end{aligned}$$

Wir wissen noch nichts über die Struktur von L , aber wir wissen, dass Systeme der Form $Uy = d$ einfacher zu lösen sind, da U eine obere Dreiecksform ist. Wenn L ebenso schön ist, könnten wir $Ax = b$ in zwei Schritten lösen, indem wir $(LU)x = b$ oder $x = U^{-1} L^{-1} b$ schreiben: 1. Lösen Sie $Ly = b$ nach y auf,

was $y = L^{-1} b$ ergibt.

2. Nachdem wir nun y haben, lösen Sie $Ux = y$, was $x = U^{-1} y = U^{-1} (L^{-1} b) = (LU)^{-1} b = A^{-1} b$ ergibt.

Wir wissen bereits, dass dieser Schritt nur $O(n)$ erfordert²) Zeit.

Unsere verbleibende Aufgabe besteht darin, sicherzustellen, dass L eine schöne Struktur hat, die das Lösen von $Ly = b$ einfacher macht als das Lösen von $Ax = b$. Glücklicherweise – und nicht überraschend – werden wir feststellen, dass L eine untere Dreiecksform ist und daher mit $O(n)$ gelöst werden kann²) Vorwärtssubstitution.

Um dies zu sehen, nehmen wir zunächst an, dass wir kein Pivotieren implementieren. Dann jede unserer Matrizen M_k ist entweder eine Skalierungsmatrix oder hat die Struktur

$$M_k = I_{n \times n} + c e e_k^T,$$

wobei $k > 1$, da wir nur eine Vorwärtssubstitution durchgeführt haben. Denken Sie daran, dass diese Matrix einem bestimmten Zweck dient: Zeile k um c skalieren und das Ergebnis zur Zeile $k-1$ hinzufügen. Diese Operation lässt sich offensichtlich leicht rückgängig machen: Skalieren Sie Zeile k um c und subtrahieren Sie das Ergebnis von Zeile $k-1$. Wir können dies formal überprüfen:

$$\begin{aligned} (I_{n \times n} + c e e_k^T)(I_{n \times n} - c e e_k^T) &= I_{n \times n} + (-c e e_k^T + c e e_k^T) - c^2 e e_k^T e e_k^T \\ &= I_{n \times n} - c^2 e (e_k^T e) e_k^T \\ &= I_{n \times n} \text{ seit } e_k^T e = 1, \text{ und } k > 1 \end{aligned}$$

Die L-Matrix ist also das Produkt von Skalierungsmatrizen und Matrizen der Form $M_k = I_{n \times n} + c_{kk} e_k e_k^T$. Sind unteres Dreieck, wenn $k > 1$. Skalierungsmatrizen sind diagonal und die Matrix M ist untere dreieckig. In Übung 2.1 werden Sie zeigen, dass das Produkt der unteren Dreiecksmatrizen die untere Dreiecksform hat, was wiederum zeigt, dass L bei Bedarf die untere Dreiecksform hat.

Wir haben gezeigt, dass Sie $A = LU$ in das Produkt der unteren und oberen Dreiecksmatrizen faktorisieren können, wenn es möglich ist, die Gaußsche Eliminierung von A ohne Pivotierung durchzuführen. Vorwärts- und Rückwärtssubstitution benötigen jeweils $O(n^2)$ Zeit. Wenn diese Faktorisierung also im Voraus berechnet werden kann, kann die lineare Lösung schneller als die vollständige $O(n^3)$ Gaußsche Eliminierung. Du wirst vorbeischaun durchgeführt werden. Übung 2.2: Was passiert, wenn wir LU mit Pivotieren durchführen? Keine größeren Änderungen sind notwendig.

2.5.2 LU implementieren

Eine einfache Implementierung der Gaußschen Eliminierung zur Lösung von $Ax = b$ ist recht einfach zu formulieren. Wie wir bereits besprochen haben, können wir insbesondere die erweiterte Matrix $(A | b)$ bilden und Zeilenoperationen nacheinander auf diesen $n \times (n + 1)$ -Block anwenden, bis er wie $(I_{n \times n} | A^{-1}b)$ aussieht. Dieser Prozess ist jedoch destruktiv, das heißt, am Ende kümmern wir uns nur um die letzte Spalte der erweiterten Matrix und haben keinen Beweis für unseren Lösungsweg aufbewahrt. Ein solches Verhalten ist für die LU-Faktorisierung eindeutig nicht akzeptabel.

Sehen wir uns an, was passiert, wenn wir zwei Eliminierungsmatrizen multiplizieren:

$$(I_{n \times n} + c_{kk} e_k e_k^T)(I_{n \times n} + c_{pe k} e_k e_k^T) = I_{n \times n} + c_{kk} e_k e_k^T + c_{pe k} e_k e_k^T$$

Wie bei unserer Konstruktion der Umkehrung einer Eliminierungsmatrix verschwindet das Produkt der beiden e_i -Terme, da die Standardbasis orthogonal ist. Diese Formel zeigt, dass nach der Skalierung des Pivots auf 1 das Produkt der Eliminierungsmatrizen, die zum Vorwärtssatz für diesen Pivot verwendet werden, die Form hat:

$$M = \begin{pmatrix} 1 & 0 & 0 & 0 \\ \tilde{y} & 0 & 1 & 0 \\ 0 & \times & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix},$$

wobei die Werte $\times$ die Werte sind, die zum Eliminieren des Rests der Spalte verwendet werden. Die Multiplikation von Matrizen dieser Form zeigt, dass die Elemente unterhalb der Diagonale von L lediglich von Koeffizienten stammen, die zur Durchführung der Substitution verwendet werden.

Wir können eine endgültige Entscheidung treffen, die Elemente entlang der Diagonale von L in der LU-Faktorisierung gleich 1 zu belassen. Diese Entscheidung ist legitim, da wir ein L immer mit einer Skalierungsmatrix S nachmultiplizieren können, um diese Elemente auf 1 zu bringen und schreiben Sie $LU = (LS)(S^{-1}U)$, ohne das Dreiecksmuster von L oder U zu beeinflussen. Mit dieser Entscheidung können wir unsere Speicherung von L und U in eine einzelne $n \times n$ -Matrix komprimieren, deren oberes Dreieck U ist und der gleich L unter der Diagonale ist; die fehlenden Diagonalelemente von L sind alle 1.

Wir sind nun bereit, Pseudocode für die einfachste LU-Faktorisierungsstrategie zu schreiben, bei der wir Zeilen oder Spalten nicht permutieren, um Pivots zu erhalten:

```
// Nimmt als Eingabe eine n-by-n-Matrix A [i, j]
// Bearbeitet A an Ort und Stelle, um die oben beschriebene kompakte LU-Faktorisierung zu erhalten

für Pivot von 1 nach n {
    PivotValue = A [Pivot, Pivot]; // Schlechte Annahme, dass dies ungleich Null ist!
```

```

for EliminationRow from ( Pivot +1) to n { // Werte unterhalb des Pivots eliminieren
    // Wie stark die Pivot-Zeile skaliert werden soll, um den Wert in der aktuellen Zeile zu eliminieren
    // Reihe ; Beachten Sie, dass wie die Skala der Pivot-Zeile nicht auf 1
    // Pivot-Wert
    Scale = A [EliminationRow, Pivot]/ PivotValue;

    // Da wir die Pivot-Zeile von der aktuellen Zeile multiplizieren / subtrahieren / skalieren
    // während der Gaußschen Eliminierung // , wir / addieren / es in der umgekehrten Operation
    // in L gespeichert
    A [EliminationRow, Pivot] = Scale;

    // Führen Sie nun wie bei der Gaußschen Eliminierung die Zeilenoperation für den Rest durch
    // der Zeile: daraus wird U
    für EliminationCol von ( Pivot +1) bis n
        Eine [ EliminationRow , EliminationCol] -= A [Pivot, EliminationCol]* Skala;
    }
}

```

2.6 Probleme

Problem 2.1. Das Produkt der unteren dreieckigen Dinge ist das untere Dreieck; Das Produkt der Pivotmatrizen sieht richtig aus

Problem 2.2. Implementieren Sie LU mit Pivotierung

Problem 2.3. Nicht quadratisches LU

Kapitel 3

Lineares Entwerfen und Analysieren Systeme

Da wir nun über einige Methoden zur Lösung linearer Gleichungssysteme verfügen, können wir diese zur Lösung einer Vielzahl von Problemen verwenden. In diesem Kapitel werden wir einige solcher Anwendungen und begleitende Analysetechniken untersuchen, um die Arten von Lösungen zu charakterisieren, die wir erwarten können.

3.1 Lösung quadratischer Systeme

Zu Beginn des vorherigen Kapitels haben wir mehrere Annahmen über die Arten linearer Systeme gemacht, die wir lösen wollten. Obwohl diese Einschränkung nicht trivial war, können tatsächlich viele, wenn nicht die meisten Anwendungen linearer Löser auf der Grundlage quadratischer, invertierbarer linearer Systeme dargestellt werden. Im Folgenden untersuchen wir einige gegensätzliche Anwendungen.

3.1.1 Regression

Wir beginnen mit einer einfachen Anwendung in der Datenanalyse, die als Regression bezeichnet wird. Angenommen, wir führen ein wissenschaftliches Experiment durch und möchten die Struktur unserer experimentellen Ergebnisse verstehen. Eine Möglichkeit, eine solche Beziehung zu modellieren, könnte darin bestehen, die unabhängigen Variablen des Experiments in einen Vektor $x \in \mathbb{R}^n$ zu schreiben und die abhängige Variable als Funktion $f(x) : \mathbb{R}^n \rightarrow \mathbb{R}$ zu betrachten. Unser Ziel ist es, die Ausgabe von f vorherzusagen ohne das komplette Experiment durchzuführen.

Beispiel 3.1 (Biologisches Experiment). Angenommen, wir möchten die Wirkung von Dünger, Sonnenlicht und Wasser auf das Pflanzenwachstum messen. Wir könnten eine Reihe von Experimenten durchführen, indem wir unterschiedliche Mengen an Dünger (in cm^3), Sonnenlicht (in Watt) und Wasser (in ml) verwenden und nach ein paar Tagen die Höhe der Pflanze messen. Wir könnten unsere Beobachtungen als Funktion $f : \mathbb{R}^3 \rightarrow \mathbb{R}$ modellieren, indem wir die drei Parameter, die wir testen möchten, berücksichtigen und die Höhe der Pflanze ausgeben.

Bei der parametrischen Regression machen wir eine vereinfachende Annahme über die Struktur von f . Zum Beispiel, Nehmen wir an, dass f linear ist:

$$f(x) = a_1x_1 + a_2x_2 + \dots + a_nx_n.$$

Dann wird unser Ziel konkreter: die Koeffizienten a_k zu schätzen.

Angenommen, wir führen eine Reihe von Experimenten durch, die die $x^{(k)}$ $y^{(k)}$ $\tilde{y} = f(x^{(k)})$. Durch den Anschluss an unsere Form für f zeigen, erhalten wir eine Reihe von Aussagen:

$$\begin{aligned} J^{(1)} = f(x^{(1)}) &= a_1 x_1^{(1)} + a_2 x_2^{(1)} + \dots + a_N x_N^{(1)} \\ J^{(2)} = f(x^{(2)}) &= a_1 x_1^{(2)} + a_2 x_2^{(2)} + \dots + a_N x_N^{(2)} \\ &\vdots \end{aligned}$$

Beachten Sie, dass im Gegensatz zu unserer Notation $Ax = b$ die Unbekannten hier die a_k -Variablen und nicht die x -Variablen sind. Wenn wir n Beobachtungen machen, können wir schreiben:

$$\begin{array}{ccc} \tilde{y} & x^{(1)} & \tilde{y} \\ \tilde{y} & x^{(2)} & \tilde{y} \\ \vdots & \vdots & \vdots \\ \tilde{y} & x^{(n)} & \tilde{y} \end{array} \quad \begin{array}{ccc} a_1 & & y^{(1)} \\ a_2 & & y^{(2)} \\ \vdots & & \vdots \\ \text{ein } \tilde{y} & = & y^{(n)} \end{array}$$

Mit anderen Worten: Wenn wir n Versuche unseres Experiments durchführen und diese in die Spalten einer Matrix $X \in \mathbb{R}^{n \times n}$ schreiben und die abhängigen Variablen in einen Vektor $y \in \mathbb{R}^n$ schreiben, können die Koeffizienten a durch Lösen von $Xa = y$ wiederhergestellt werden.

Tatsächlich können wir unseren Ansatz auf andere interessantere nichtlineare Formen für die Funktion f verallgemeinern. Wichtig ist hier, dass f eine lineare Kombination von Funktionen ist. Nehmen wir insbesondere an, dass $f(x)$ die folgende Form annimmt:

$$f(x) = a_1 f_1(x) + a_2 f_2(x) + \dots + a_m f_m(x),$$

wobei $f_k : \mathbb{R}^n \rightarrow \mathbb{R}$ und wir die Parameter a_k schätzen wollen. Dann können wir durch eine parallele Ableitung gegebenen m Beobachtungen der Form $x^{(k)}$ die Parameter durch Lösen von: $\tilde{y} = y$ finden

$$\begin{array}{ccc} f_1(x^{(1)}) & f_2(x^{(1)}) & \tilde{y} \\ f_1(x^{(2)}) & f_2(x^{(2)}) & \tilde{y} \\ \vdots & \vdots & \vdots \\ f_1(x^{(m)}) & f_2(x^{(m)}) & \tilde{y} \end{array} \quad \begin{array}{ccc} \dots & f_m(x^{(1)}) & \tilde{y} \\ \dots & f_m(x^{(2)}) & \tilde{y} \\ \dots & \vdots & \vdots \\ \dots & f_m(x^{(m)}) & \tilde{y} \end{array} \quad \begin{array}{ccc} a_1 & & y^{(1)} \\ a_2 & & y^{(2)} \\ \vdots & & \vdots \\ \text{Bin } \tilde{y} & = & y^{(M)} \end{array}$$

Das heißt, selbst wenn die f nichtlinear sind, können wir die Gewichte a_k mit rein linearen Techniken lernen.

Beispiel 3.2 (Lineare Regression). Das System $Xa = y$ kann aus der allgemeinen Formulierung wiederhergestellt werden, indem $f_k(x) = x_k$ angenommen wird.

Beispiel 3.3 (Polynomielle Regression). Angenommen, wir beobachten eine Funktion einer einzelnen Variablen $f(x)$ und möchten sie als Polynom n -ten Grades schreiben

$$f(x) = a_0 + a_1 x + a_2 x^2 + \dots + a_N x^N.$$

Gegeben sind n Paare $x^{(k)}$ $y^{(k)}$, Wir können die Parameter über das System ermitteln

$$\begin{array}{ccc} 1 & x^{(1)} & y^{(1)} \\ 1 & x^{(2)} & y^{(2)} \\ \vdots & \vdots & \vdots \\ 1 & x^{(n)} & y^{(n)} \end{array} \quad \begin{array}{ccc} x^{(1)^2} & \dots & x^{(1)^n} \\ x^{(2)^2} & \dots & x^{(2)^n} \\ \vdots & \vdots & \vdots \\ x^{(n)^2} & \dots & x^{(n)^n} \end{array} \quad \begin{array}{ccc} a_0 & & y^{(1)} \\ a_2 & & y^{(2)} \\ \vdots & & \vdots \\ \text{ein } \tilde{y} & = & y^{(n)} \end{array}$$

Mit anderen Worten, wir nehmen $f_k(x) = x^k$ in unserer obigen allgemeinen Form. Übrigens ist die Matrix auf der linken Seite dieser Beziehung als Vandermonde-Matrix bekannt, die viele besondere Eigenschaften aufweist, die für ihre Struktur spezifisch sind.

Beispiel 3.4 (Oszillation). Angenommen, wir möchten a und $\tilde{y}$ für eine Funktion $f(x) = a \cos(x + \tilde{y})$ finden. Erinnern Sie sich an die Trigonometrie, dass wir $\cos(x + \tilde{y}) = \cos x \cos \tilde{y} - \sin x \sin \tilde{y}$ schreiben können. Wenn wir also zwei Beispielpunkte gegeben haben, können wir die obige Technik verwenden, um $f(x) = a_1 \cos x + a_2 \sin x$ zu finden, und unter Anwendung dieser Identität können wir schreiben

$$2a = a_1^2 + a_2^2$$

$$a_2 \tilde{y} = \tilde{y} \arctan \frac{a_1}{a_2}$$

Diese Konstruktion kann erweitert werden, um $f(x) = \tilde{y}_k \cos(x + \tilde{y}_k)$ zu finden, was eine Möglichkeit bietet, die diskrete Fourier-Transformation von f zu motivieren.

3.1.2 Kleinste Quadrate

Die Techniken in §3.1.1 bieten wertvolle Methoden zum Finden eines stetigen f , das genau zu einer Menge von Datenpaaren $x_k \tilde{y}_k$ passt. Dieser Ansatz hat zwei damit verbundene Nachteile:

- Bei der Messung der Werte x_k und $\tilde{y}_k$ kann es zu Fehlern kommen. In diesem Fall kann eine ungefähre Beziehung $f(x_k) \approx \tilde{y}_k$ akzeptabel oder sogar einer exakten Beziehung $f(x_k) = \tilde{y}_k$ vorzuziehen sein.
- Beachten Sie, dass wir genau m Beobachtungen $x_k \tilde{y}_k$ benötigen, wenn es m Funktionen f_k total gäbe. Zusätzliche Beobachtungen müssten verworfen werden, oder wir müssten die Form von f ändern.

Beide Probleme hängen mit dem größeren Problem der Überanpassung zusammen: Die Anpassung einer Funktion mit n Freiheitsgraden an n Datenpunkte lässt keinen Raum für Messfehler.

Nehmen wir allgemeiner an, wir möchten das lineare System $Ax = b$ nach x lösen. Wenn wir Zeile k von A mit a_{rk} bezeichnen, dann sieht unser System aus

$$\begin{pmatrix} a_{r1} \\ a_{r2} \\ \vdots \end{pmatrix} x_1 + \begin{pmatrix} a_{r2} \\ a_{r3} \\ \vdots \end{pmatrix} x_2 + \dots + \begin{pmatrix} a_{rn} \\ a_{r,n+1} \\ \vdots \end{pmatrix} x_n = b_r$$

oder $\sum_{j=1}^n a_{rj} x_j = b_r$ nach Definition der Matrixmultiplikation.

Somit entspricht jede Zeile des Systems einer Beobachtung der Form $\sum_{j=1}^n a_{rj} x_j = b_r$. Das heißt, eine weitere Möglichkeit, das lineare System $Ax = b$ zu interpretieren, ist als n Aussagen der Form „Das Skalarprodukt von x mit a_{rk} ist b_k .“

Unter diesem Gesichtspunkt kodiert ein hohes System $Ax = b$ mit $A \in \mathbb{R}^{m \times n}$ und $m > n$ einfach mehr als n dieser Skalarproduktbeobachtungen. Wenn wir jedoch mehr als n Beobachtungen machen, können diese inkompatibel sein; Wie in Abschnitt 2.1 erläutert, werden große Systeme wahrscheinlich keine Lösung zulassen. In unserem oben erläuterten „experimentellen“ Aufbau könnte diese Situation auf Fehler bei der Messung der Paare $x_k \tilde{y}_k$ zurückzuführen sein.

Wenn wir $Ax = b$ nicht exakt lösen können, könnten wir das Problem etwas lockern, um $Ax \approx b$ anzunähern. Insbesondere können wir verlangen, dass das Residuum $b - Ax$ so klein wie möglich ist, indem wir es minimieren

Norm $\|b - Ax\|$. Beachten Sie, dass diese Norm bei Null minimiert wird, wenn es eine exakte Lösung für das lineare System gibt, da in diesem Fall $b - Ax = 0$ gilt. Die Minimierung von $\|b - Ax\|^2$ ist dasselbe wie die Minimierung von $\|b - Ax\|$, was wir in Beispiel 0.16 erweitert haben zu:

$$\|b - Ax\|^2 = (b - Ax)^T (b - Ax) = b^T b - 2b^T Ax + x^T A^T A x$$

Der Gradient dieses Ausdrucks in Bezug auf x muss im Minimum Null sein, was das folgende System ergibt:

$$0 = 2A^T Ax - 2A^T b$$

Oder äquivalent: $A^T Ax = A^T b$.

Diese berühmte Beziehung verdient einen Satz:

Satz 3.1 (Normalgleichungen). Minima des Residuums $\|b - Ax\|$ für $A \in \mathbb{R}^{m \times n}$ (ohne Einschränkung auf m oder n) erfüllen $A^T Ax = A^T b$.

Wenn mindestens n Zeilen von A linear unabhängig sind, dann ist die Matrix $A^T A \in \mathbb{R}^{n \times n}$ invertierbar. In diesem Fall tritt das minimale Residuum (eindeutig) bei $(A^T A)^{-1} A^T b$ auf, oder äquivalent dazu ist die Lösung des Problems der kleinsten Quadrate genauso einfach wie die Lösung des quadratischen linearen Systems $A^T Ax = A^T b$ aus Satz 3.1. Daher haben wir unseren Satz an Lösungsstrategien auf $A \in \mathbb{R}^{m \times n}$ mit $m \geq n$ erweitert, indem wir nur Techniken für quadratische Matrizen angewendet haben.

Der unterbestimmte Fall $m < n$ ist wesentlich schwieriger zu handhaben. Insbesondere verlieren wir die Möglichkeit einer eindeutigen Lösung für $Ax = b$. In diesem Fall müssen wir eine zusätzliche Annahme für x treffen, um eine eindeutige Lösung zu erhalten, z. B. dass es eine kleine Norm hat oder viele Nullstellen enthält. Jede dieser Regularisierungsannahmen führt zu einer anderen Lösungsstrategie; Wir werden einige davon in den Übungen zu diesem Kapitel untersuchen.

3.1.3 Zusätzliche Beispiele

Eine wichtige Fähigkeit besteht darin, lineare Systeme „in freier Wildbahn“ identifizieren zu können. Hier zählen wir noch schnell ein paar weitere Beispiele auf.

Ausrichtung

Angenommen, wir machen zwei Fotos derselben Szene aus verschiedenen Positionen. Eine häufige Aufgabe bei Computer Vision und Grafik besteht darin, sie zusammenzufügen. Dazu kann der Benutzer (oder ein automatisches System) eine Anzahl von Punkten $x_k, y_k \in \mathbb{R}^2$ markieren, sodass x_k in Bild eins und y_k in Bild zwei entspricht. Natürlich wurden beim Abgleich dieser Punkte wahrscheinlich Fehler gemacht, daher möchten wir eine stabile Transformation zwischen den beiden Bildern finden, indem wir die Anzahl der erforderlichen Paare (x, y) überabtasten.

Vorausgesetzt, unsere Kamera verfügt über ein Standardobjektiv, sind die Kameraprojektionen linear, also ein vernünftiger Wert. Die Annahme ist, dass es ein $A \in \mathbb{R}^{2 \times 2}$ und einen Translationsvektor $b \in \mathbb{R}^2$ gibt, so dass

$$y_k = Ax_k + b.$$

¹ Es kann sinnvoll sein, an dieser Stelle noch einmal auf die Vorbemerkungen in Kapitel 0 zurückzukommen.

Unsere unbekannten Variablen sind hier A und b statt x_k und y_k .

In diesem Fall können wir die Transformation finden, indem wir Folgendes lösen:

$$\min_{A, b} \sum_{k=1}^P (Ax_k + b - y_k)^2.$$

Dieser Ausdruck ist wiederum eine Summe quadrierter linearer Ausdrücke in unseren Unbekannten A und b und kann durch eine ähnliche Ableitung wie bei unserer Diskussion des Problems der kleinsten Quadrate linear gelöst werden.

Dekonvolution

Oftmals machen wir aus Versehen Fotos, die etwas unscharf sind. Während ein Foto, das völlig unscharf ist, ein verlorener Fehler sein kann, können wir bei örtlicher oder kleiner Unschärfe möglicherweise mithilfe von Computertechniken ein schärferes Bild wiederherstellen. Eine einfache Strategie ist die Entfaltung, die unten erläutert wird.

Wir können uns ein Foto als einen Punkt in $\mathbb{R}^p$, wobei p die Anzahl der Pixel ist; natürlich, wenn die vorstellen. Wenn ein Foto in Farbe ist, benötigen wir möglicherweise drei Werte (RGB) pro Pixel, was zu einer ähnlichen Technik in $\mathbb{R}^{3p}$ führt. Unabhängig davon sind viele einfache Bildunschärfen linear, z. B. Gaußsche Faltung oder Operationen zur Mittelung von Pixeln mit ihren Nachbarn auf dem Bild. In der Bildverarbeitung haben diese linearen Operationen oft andere spezielle Eigenschaften wie Verschiebungsinvarianz, aber für unsere Zwecke können wir uns Unschärfe als einen linearen Operator $x \mapsto Gx$ vorstellen.

Angenommen, wir machen ein verschwommenes Foto $x_0 \in \mathbb{R}^p$. Dann könnten wir versuchen, das scharfe Bild wiederherzustellen $x \in \mathbb{R}^p$ durch Lösen des Problems der kleinsten Quadrate

$$\min_{x \in \mathbb{R}^p} \|x_0 - Gx\|^2.$$

Das heißt, wir verlangen, dass Sie das beobachtete Foto x_0 erhalten, wenn Sie x mit G verwischen. Natürlich können viele scharfe Bilder unter G das gleiche unscharfe Ergebnis liefern, daher fügen wir der obigen Minimierung oft zusätzliche Terme hinzu und verlangen, dass x_0 nicht stark variiert.

3.2 Besondere Eigenschaften linearer Systeme

Unsere Diskussion der Gaußschen Eliminierung und der LU-Faktorisierung führte zu einer völlig generischen Methode zur Lösung linearer Gleichungssysteme. Obwohl diese Strategie immer funktioniert, können wir manchmal Geschwindigkeit oder numerische Vorteile erzielen, indem wir das spezielle System untersuchen, das wir lösen. Hier diskutieren wir einige gängige Beispiele, bei denen das Wissen über das lineare System Lösungsstrategien vereinfachen kann.

3.2.1 Positiv-definite Matrizen und die Cholesky-Faktorisierung

Wie in Satz 3.1 gezeigt, liefert die Lösung eines Problems der kleinsten Quadrate $Ax \approx b$ eine Lösung x , die das quadratische lineare System $(A^T A)x = A^T b$ erfüllt. Unabhängig von A weist die Matrix $A^T A$ einige besondere Eigenschaften auf, die dieses System zu etwas Besonderem machen.

Erstens ist es leicht zu erkennen, dass $A^T A$ symmetrisch ist, da

$$(A^T A)^T = A^T (A^T)^T = A^T A.$$

Hier haben wir einfach die Identitäten $(AB) = BA$ und $(A^T)^T = A$ verwendet. Wir können diese Symmetrie indexweise ausdrücken, indem wir $(A^T A)_{ij} = (A A^T)_{ji}$ für alle Indizes i, j schreiben. Diese Eigenschaft impliziert, dass es ausreicht, nur die Werte von AA^T auf oder über der Diagonale zu speichern, da die restlichen Elemente durch Symmetrie erhalten werden können.

Darüber hinaus ist AA^T eine positive semidefinite Matrix, wie unten definiert:

Definition 3.1 (Positiv (Semi-)Definit). Eine Matrix $B \in \mathbb{R}^{n \times n}$ ist positiv semidefinit, wenn für alle $x \in \mathbb{R}^n$ $x^T B x \geq 0$. B ist positiv definit, wenn $x^T B x > 0$, wenn $x \neq 0$.

Es ist leicht zu zeigen, dass AA^T positiv semidefinit ist, denn:

$$x^T A A^T x = (A x)^T (A x) = (A x) \cdot (A x) = \|A x\|_2^2 \geq 0.$$

Wenn die Spalten von A tatsächlich linear unabhängig sind, dann ist AA^T positiv definit.

Nehmen wir allgemeiner an, wir möchten ein symmetrisches positiv definites System $Cx = d$ lösen. Wie wir bereits untersucht haben, könnten wir die Matrix C LU-faktorisieren, aber tatsächlich können wir es etwas besser machen. Wir schreiben $C \in \mathbb{R}^{n \times n}$ als Blockmatrix:

$$C = \begin{pmatrix} c_{11} & v \\ v^T & \tilde{C} \end{pmatrix}$$

wobei $v \in \mathbb{R}^{n-1}$ und $\tilde{C} \in \mathbb{R}^{(n-1) \times (n-1)}$. Dank der besonderen Struktur von C können wir folgende Beobachtung machen:

$$\begin{aligned} C e_1 &= \begin{pmatrix} 1 & 0 & \dots & 0 \end{pmatrix} e_1 = \begin{pmatrix} c_{11} & v \\ v^T & \tilde{C} \end{pmatrix} \begin{pmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix} \\ &= \begin{pmatrix} 1 & 0 & \dots & 0 \end{pmatrix} \begin{pmatrix} c_{11} \\ v \end{pmatrix} \\ &= c_{11} e_1 \\ &> 0, \text{ da } C \text{ positiv definit ist und } e_1 \neq 0. \end{aligned}$$

Dies zeigt, dass wir – abgesehen von numerischen Problemen – keine Pivotierung verwenden müssen, um sicherzustellen, dass $c_{11} \neq 0$ für den ersten Schritt der Gaußschen Eliminierung.

Wenn wir mit der Gaußschen Eliminierung fortfahren, können wir eine Vorwärtssubstitutionsmatrix E anwenden, die hat im Allgemeinen die Form

$$E = \begin{pmatrix} 1/\bar{c}_{11} & 0 \\ r & I_{(n-1) \times (n-1)} \end{pmatrix}.$$

Hier enthält der Vektor $r \in \mathbb{R}^{n-1}$ die Vielfachen von Zeile 1, um den Rest der ersten Spalte von C aufzuheben. Aus Gründen, die gleich klar werden, skalieren wir Zeile 1 auch um $1/\bar{c}_{11}$!

Nach der Vorwärtssubstitution kennen wir das Produkt konstruktionsbedingt

$$EC = \begin{pmatrix} \bar{c}_{11} & \bar{v} \\ 0 & D \end{pmatrix}$$

für ein $D \in \mathbb{R}^{(n-1) \times (n-1)}$.

Hier weichen wir von der Gaußschen Eliminierung ab: Anstatt mit der zweiten Zeile fortzufahren, können wir mit E nachmultiplizieren, um ein Produkt ECE zu erhalten:

$$\begin{aligned} ECE &= (EC)E\tilde{y} \\ &= \begin{pmatrix} c_{11} & \tilde{v}^T / \tilde{y}^T c_{11} \\ 0 & D \end{pmatrix} \begin{pmatrix} 1/\tilde{y}^T c_{11}^T \\ 0 \end{pmatrix} \begin{matrix} R \\ I(n-1) \times (n-1) \end{matrix} \\ &= \begin{pmatrix} 1 & 0 \\ 0 & D \end{pmatrix} \end{aligned}$$

Das heißt, wir haben die erste Zeile und die erste Spalte von C eliminiert! Darüber hinaus lässt sich leicht überprüfen, dass die Matrix D auch positiv definit ist.

Wir können diesen Vorgang wiederholen, um alle Zeilen und Spalten von C symmetrisch zu entfernen. Beachten dass wir sowohl Symmetrie als auch positive Bestimmtheit verwendet haben, um die Faktorisierung abzuleiten, da

- Die Symmetrie ermöglichte es uns, auf beiden Seiten dasselbe E anzuwenden, und
- Positive Bestimmtheit garantiert, dass $c_{11} > 0$, was impliziert, dass $\tilde{y}^T c_{11}$ existiert.

Am Ende erhalten wir nun ähnlich wie bei der LU-Faktorisierung eine Faktorisierung $C = LL^T$ für eine untere Dreiecksmatrix L. Dies ist als Cholesky-Faktorisierung von C bekannt. Wenn das Ziehen der Quadratwurzeln entlang der Diagonale zu numerischen Problemen führt, wird eine entsprechende LDL-Faktorisierung durchgeführt, wobei D eine Diagonalmatrix ist, vermeidet dieses Problem und lässt sich aus der obigen Diskussion einfach ableiten.

Die Cholesky-Faktorisierung ist aus mehreren Gründen wichtig. Besonders hervorzuheben ist, dass zum Speichern von L nur halb so viel Speicher benötigt wird wie für die LU-Faktorisierung von C oder sogar von C selbst, da die Elemente über der Diagonale Null sind und wie in LU das Lösen von $Cx = d$ so einfach ist wie eine Vorwärts- und Rückwärtssubstitution. Weitere Eigenschaften der Faktorisierung erforschen Sie in den Übungen.

Letztendlich kann der Code für die Cholesky-Faktorisierung sehr prägnant sein. Um eine besonders comp abzuleiten Nehmen wir an, wir wählen eine beliebige Zeile k aus und schreiben L in Blockform, um diese Zeile zu isolieren:

$$L = \begin{pmatrix} \tilde{y} & L_{11} & 0 & 0 & k & 0 \\ & k & & & & \\ & \tilde{y}\tilde{y}^T & L_{31} & k & L_{33} & \tilde{y}\tilde{y}^T \end{pmatrix}$$

Hier sind L_{11} und L_{33} beide untere dreieckige Quadratmatrizen. Dann ergibt die Durchführung eines Produkts:

$$\begin{aligned} LL^T &= \begin{pmatrix} \tilde{y} & L_{11} & 0 & 0 & k & 0 \\ & k & & & & \\ & L_{31} & k & L_{33} & \tilde{y}\tilde{y}^T & \end{pmatrix} \begin{pmatrix} L_{11} & k & L_{31}^T & 0 & 0 & L_{33}^T \\ 0 & k & & & & \\ 0 & 0 & L_{33} & k & \tilde{y}\tilde{y}^T & \end{pmatrix} \\ &= \begin{pmatrix} \tilde{y} & L_{11} & k & k & k & k \\ & L_{11} & k & k & k & k \\ & \tilde{y}\tilde{y}^T & L_{31} & k & L_{33} & \tilde{y}\tilde{y}^T \end{pmatrix} \begin{pmatrix} L_{11} & k & L_{31}^T & 0 & 0 & L_{33}^T \\ 0 & k & & & & \\ 0 & 0 & L_{33} & k & \tilde{y}\tilde{y}^T & \end{pmatrix} \end{aligned}$$

Werte des Produkts, die für unsere Ableitung nicht notwendig sind, lassen wir weg.

Am Ende wissen wir, dass wir $C = LL^T$ schreiben können. Das mittlere Element des Produkts zeigt:

$$kk = c_{kk} - \tilde{y}^T k$$

wobei $k \leq R_k - 1$ die Elemente der k -ten Zeile von L links von der Diagonale enthält. Außerdem, Das mittlere linke Element des Produkts wird angezeigt

$$L_{11}k = c_k$$

wobei c_k die Elemente von C an derselben Position enthält a_{kk} . Da L_{11} die untere Dreiecksform hat, ist dies System kann durch Vorwärtssubstitution gelöst werden!

Beachten Sie, dass unsere obige Diskussion einen Algorithmus zur Berechnung der Cholesky-Faktorisierung hervorbringt von oben nach unten, da L_{11} bereits berechnet ist, wenn wir Zeile k erreichen. Wir bieten an Pseudocode unten, angepasst von CITE:

```
// Nimmt als Eingabe eine n-by-n-Matrix A [i, j]
// Bearbeitet A an Ort und Stelle, um die Cholesky-Faktorisierung in seinem unteren Dreieck zu erhalten

für k von 1 bis n {
    // Zurück – ersetzen, um l_k zu finden
    für i von 1 bis k - 1 { // Element i von l_k
        Summe = 0;
        für j von 1 bis i - 1
            sum += A [i, j] * A [k, j];
        Ein [k, i] = ( A [k, i] - Summe ) / A [i, i];
    }

    // Wende die Formel für l_kk an
    normSquared = 0
    für j von 1 bis i - 1
        normSquared += A [k, j]^2;
    Ein [k, k] = sqrt ( A [k, k] - normSquared );
}
```

Wie bei der LU-Faktorisierung läuft dieser Algorithmus eindeutig in $O(n)^3$ Zeit.

3.2.2 Sparsität

Viele lineare Gleichungssysteme weisen von Natur aus Sparsity-Eigenschaften auf, was bedeutet, dass die meisten Einträge von A im System $Ax = b$ sind genau Null. Sparsity kann eine bestimmte Struktur in einem widerspiegeln gegebenes Problem, einschließlich der folgenden Anwendungsfälle:

- In der Bildverarbeitung drücken viele Systeme zur Fotobearbeitung und zum Verständnis Beziehungen zwischen den Werten von Pixeln und denen ihrer Nachbarn im Bildraster aus. Ein Bild kann ein Punkt in R^p für p Pixel sein, aber wenn man $Ax = b$ für ein neues Bild der Größe p auflöst, ist $A \in R^{p \times p}$ kann nur $O(p)$ statt $O(p^2)$ ungleich Null, da jede Zeile nur ein einzelnes Pixel umfasst und seine Nachbarn oben/unten/links/rechts.
- Beim maschinellen Lernen verwendet ein grafisches Modell eine Graphenstruktur $G = (V, E)$, um Wahrscheinlichkeitsverteilungen über mehrere Variablen auszudrücken. Jede Variable wird durch einen Knoten $v \in V$ von dargestellt des Graphen und die Kante $e \in E$ stellt eine probabilistische Abhängigkeit dar. Lineare Systeme entstehen in Dieser Kontext hat oft eine Zeile pro Scheitelpunkt $v \in V$ mit Nicht-Nulls nur in beteiligten Spalten v und seine Nachbarn.
- In der Computergeometrie werden Formen häufig durch Ansammlungen verknüpfter Dreiecke ausgedrückt zu einem Netz zusammenfügen. Gleichungen für die Oberflächenglättung und andere Aufgaben verknüpfen wiederum Positionen und andere Werte an einem bestimmten Scheitelpunkt mit denen an seinen Nachbarn im Netz.

Beispiel 3.5 (Harmonische Parametrisierung). Angenommen, wir möchten ein Bild verwenden, um ein Dreiecksnetz zu texturieren. Ein Netz kann als eine Ansammlung von Eckpunkten $V \subset \mathbb{R}^3$ dargestellt werden, die durch Kanten $E \subset V \times V$ miteinander verbunden sind, um Dreiecke zu bilden. Da sich die Eckpunkte der Geometrie im $\mathbb{R}^3$ befinden, müssen wir einen Weg finden, sie auf die Bildebene abzubilden, um die Textur als Bild zu speichern. Daher müssen wir jedem $v \in V$ Texturkoordinaten $t(v) \in \mathbb{R}^2$ auf der Bildebene zuweisen. Eine Veranschaulichung finden Sie in Abbildung NUMMER.

Eine Strategie zur Erstellung dieser Karte beinhaltet eine einzige lineare Lösung. Angenommen, das Netz hat eine Scheibentopologie, das heißt, es kann auf das Innere eines Kreises in der Ebene abgebildet werden. Für jeden Scheitelpunkt v_b an der Grenze des Netzes können wir die Position von v_b angeben, indem wir ihn auf einem Kreis platzieren. Im Inneren können wir verlangen, dass die Position der Texturkarte der Durchschnitt ihrer benachbarten Positionen ist:

$$t(v) = \frac{1}{|n(v)|} \sum_{w \in n(v)} t(w)$$

Hier ist $n(v) \subset V$ die Menge der Nachbarn von $v \in V$ auf dem Netz. Somit ist jedes $v \in V$ mit einer linearen Gleichung verbunden, die es entweder auf den Rand fixiert oder verlangt, dass seine Position dem Durchschnitt seiner benachbarten Positionen entspricht. Dieses $|V| \times |V|$ Gleichungssystem führt zu einer stabilen Parametrisierungsstrategie, die als harmonische Parametrisierung bekannt ist. Die Matrix des Systems hat nur $O(|V|)$ ungleich Nullen in Slots, die den Eckpunkten und ihren Nachbarn entsprechen.

Wenn $A \in \mathbb{R}^{n \times n}$ so dünnbesetzt ist, dass es $O(n)$ - statt $O(n^2)$ - Werte enthält, gibt es natürlich keinen Grund, A als $n \times n$ -Matrix zu speichern. Stattdessen speichern Speichertechniken für spärliche Matrizen nur die $O(n)$ ungleich Nullen in einer sinnvolleren Datenstruktur, z. B. einer Liste von Zeilen-/Spalten-/Werttripeln. Die Wahl einer Matrixdatenstruktur beinhaltet die Berücksichtigung der wahrscheinlichen Operationen, die auf der Matrix auftreten werden, möglicherweise einschließlich Multiplikation und Iteration über ungleich Nullen, oder Iteration über einzelne Zeilen oder Spalten.

Leider ist leicht zu erkennen, dass die LU-Faktorisierung eines dünn besetzten A möglicherweise nicht zu dünn besetzten L - und U -Matrizen führt. Dieser Strukturverlust schränkt die Anwendbarkeit dieser Methoden zur Lösung von $Ax = b$ erheblich ein, wenn A groß, aber dünn besetzt ist. Glücklicherweise gibt es viele direkte Sparse-Löser, die LU an dünn besetzte Matrizen anpassen und eine LU -ähnliche Faktorisierung erzeugen können, ohne viel Füllung oder zusätzliche Nicht-Nullen zu induzieren; Die Diskussion dieser Techniken liegt außerhalb des Rahmens dieses Textes. Alternativ wurden iterative Techniken verwendet, um Näherungslösungen für lineare Systeme zu erhalten; Wir werden die Diskussion dieser Methoden auf zukünftige Kapitel verschieben.

Bestimmte Matrizen sind nicht nur dünn besetzt, sondern auch strukturiert. Zum Beispiel ein tridiagonales System von Lineare Gleichungen haben das folgende Muster von Werten ungleich Null:

$$\begin{array}{ccccccc} & & x & x & & & \\ & & x & x & x & & \\ & & & x & x & x & \\ & & & & x & x & x \\ & & & & & x & x \\ & & & & & & x & x \end{array}$$

In den Übungen im Anschluss an dieses Kapitel werden Sie eine spezielle Version der Gaußschen Eliminierung für den Umgang mit dieser einfachen Bandstruktur herleiten.

In anderen Fällen sind Matrizen möglicherweise nicht dünn besetzt, lassen aber möglicherweise eine dünn besetzte Darstellung zu. Für die Prüfung Betrachten Sie beispielsweise die zyklische Matrix:

$$\begin{array}{cccccccc} & & & & & & & \\ & & & & & & & \\ & & & & & & & \\ & & & & & & & \\ & & & & & & & \\ & & & & & & & \\ & & & & & & & \\ & & & & & & & \end{array} \begin{array}{cccccccc} & & & & & & & \\ & & & & & & & \\ & & & & & & & \\ & & & & & & & \\ & & & & & & & \\ & & & & & & & \\ & & & & & & & \\ & & & & & & & \end{array}$$

abcddabcccdabbcda

Offensichtlich kann diese Matrix nur mit den Werten a, b, c, d gespeichert werden. Spezielle Techniken für diese und andere Matrizenklassen sind gut untersucht und oft effizienter als die generische Gaußsche Eliminierung.

3.3 Sensitivitätsanalyse

Wie wir gesehen haben, ist es wichtig, die Matrix eines linearen Systems zu untersuchen, um herauszufinden, ob sie spezielle Eigenschaften aufweist, die den Lösungsprozess vereinfachen können. Sparsität, positive Bestimmtheit, Symmetrie usw. können Hinweise auf den richtigen Löser für ein bestimmtes Problem geben.

Selbst wenn eine bestimmte Lösungsstrategie theoretisch funktionieren könnte, ist es jedoch ebenso wichtig zu verstehen, wie gut wir der Antwort eines bestimmten Löser auf ein lineares System vertrauen können. Beispielsweise könnte es aufgrund von Rundungen und anderen diskreten Effekten der Fall sein, dass eine Implementierung der Gaußschen Eliminierung zur Lösung von $Ax = b$ eine Lösung x_0 ergibt, so dass $0 < Ax_0 \leq 1$; mit anderen Worten, x_0 löst das System nur näherungsweise.

Eine Möglichkeit, die Wahrscheinlichkeit dieser Näherungseffekte zu verstehen, ist die Sensitivitätsanalyse. In diesem Ansatz fragen wir, was mit x passieren könnte, wenn wir, anstatt $Ax = b$ in Wirklichkeit zu lösen, ein gestörtes Gleichungssystem $(A + \tilde{A})x = b + \tilde{b}$ lösen. Es gibt zwei Möglichkeiten, die Schlussfolgerungen dieser Art von Analyse anzuzeigen:

1. Aufgrund von Rundungen und anderen Effekten machen wir wahrscheinlich Fehler bei der Darstellung von A und b . Diese Analyse zeigt dann die bestmögliche Genauigkeit, die wir für x angesichts der gemachten Fehler bei der Darstellung des Problems erwarten können.
2. Wenn unser Löser eine Näherung x_0 für die Lösung von $Ax = b$ generiert, ist es eine exakte Lösung für das System $Ax_0 = b_0$, wenn wir $b_0 \leq Ax_0$ definieren (stellen Sie sicher, dass Sie verstehen, warum dieser Satz keine Tautologie ist!). Wenn man versteht, wie sich Änderungen in x_0 auf Änderungen in b_0 auswirken, zeigt sich, wie empfindlich das System auf leicht falsche Antworten reagiert.

Beachten Sie, dass unsere Diskussion hier unseren Definitionen von Vorwärts- und Rückwärtsfehlern in den vorherigen Kapiteln ähnelt und tatsächlich durch diese motiviert ist.

3.3.1 Matrix- und Vektornormen

Bevor wir die Empfindlichkeit eines linearen Systems diskutieren können, müssen wir etwas sorgfältig definieren, was es bedeutet, dass eine Änderung $\tilde{x}$ „klein“ ist. Im Allgemeinen möchten wir die Länge oder Norm eines Vektors x messen. Wir sind bereits auf die Zweinorm eines Vektors gestoßen:

$$\|x\|_2 = \sqrt{x_1^2 + x_2^2 + \dots + x_n^2}$$

für $x \in \mathbb{R}^n$. Diese Norm ist aufgrund ihrer Verbindung zur euklidischen Geometrie beliebt, aber sie ist keineswegs die einzige Norm auf dem $\mathbb{R}^n$. Im Allgemeinen definieren wir eine Norm wie folgt:

Definition 3.2 (Vektornorm). Eine Vektornorm ist eine Funktion $\|\cdot\| : \mathbb{R}^n \rightarrow [0, \infty)$, die Folgendes erfüllt Bedingungen:

- $\|x\| = 0$ genau dann, wenn $x = 0$.

- $cx = |c|x$ für alle Skalare $c \in \mathbb{R}$ und Vektoren $x \in \mathbb{R}^n$
- $x + y \in \mathbb{R}^n$ für alle $x, y \in \mathbb{R}^n$

Während wir die beiden Indizes $\cdot 2$ verwenden, um die Zweinorm eines Vektors zu bezeichnen, verwenden wir, sofern wir nichts anderes vermerken, die Notation x , um die Zweinorm von x zu bezeichnen. Außer dieser Norm gibt es viele andere Beispiele:

- Die p -

Norm x_p für $p \geq 1$, gegeben durch:

$$x_p = (|x_1|^p + |x_2|^p + \dots + |x_n|^p)^{1/p}$$

Von besonderer Bedeutung ist die 1-Norm oder „Taxi“-Norm, gegeben durch

$$x_1 = \sum_{k=1}^n |x_k|.$$

Diese Norm erhält ihren Spitznamen, weil sie die Entfernung angibt, die ein Taxi zwischen zwei Punkten in einer Stadt zurücklegt, in der die Straßen nur in Nord-Süd- und Ost-West-Richtung verlaufen.

- Die ∞ -Norm x_∞ gegeben durch:

$$x_\infty = \max(|x_1|, |x_2|, \dots, |x_n|).$$

In gewissem Sinne sind viele Normen für $\mathbb{R}^n$ gleich. Nehmen wir insbesondere an, wir sagen, zwei Normen seien äquivalent, wenn sie die folgende Eigenschaft erfüllen: **Definition 3.3**

(Äquivalente Normen). Zwei Normen $\|\cdot\|$ und $\|\cdot\|'$ sind äquivalent, wenn es Konstanten c_{low} und c_{high} gibt, so dass

$$c_{\text{low}} \|x\|' \leq \|x\| \leq c_{\text{high}} \|x\|'$$

für alle $x \in \mathbb{R}^n$.

Diese Bedingung garantiert, dass bis auf einige konstante Faktoren alle Normen hinsichtlich der Vektoren übereinstimmen sind „klein“ und „groß“. Tatsächlich werden wir einen berühmten Satz aus der Analysis ohne Beweis angeben: **Satz 3.2**

(Äquivalenz von Normen auf $\mathbb{R}^n$). Alle Normen auf $\mathbb{R}^n$ sind äquivalent.

Dieses etwas überraschende Ergebnis impliziert, dass alle Vektornormen das gleiche grobe Verhalten aufweisen, die Wahl einer Norm zur Analyse oder Darstellung eines bestimmten Problems jedoch einen großen Unterschied machen kann. Beispielsweise geht die ∞ -Norm auf $\mathbb{R}^3$ davon aus, dass der Vektor $(1000, 1000, 1000)$ dieselbe Norm wie $(1000, 0, 0)$ hat, während die 2-Norm sicherlich von den zusätzlichen Werten ungleich Null beeinflusst wird.

Da wir nicht nur Vektoren, sondern auch Matrizen betrachten, müssen wir auch die Norm einer Matrix annehmen können. Natürlich ändert sich die grundlegende Definition einer Norm auf $\mathbb{R}^{n \times m}$ nicht. Aus diesem Grund können wir jede Matrix in $\mathbb{R}^{m \times n}$ zu einem Vektor in $\mathbb{R}^{nm}$ „abwickeln“, um jede Vektornorm in Matrizen zu übernehmen. Eine solche Norm ist die Frobenius-Norm, gegeben durch

$$\|A\|_{\text{Fro}} = \sqrt{\sum_{i,j} a_{ij}^2}.$$

Solche Anpassungen von Vektornormen sind jedoch nicht immer sehr aussagekräftig. Insbesondere ist die Priorität für das Verständnis der Struktur einer Matrix A häufig ihre Wirkung auf Vektoren, d. h. die wahrscheinlichen Ergebnisse, wenn A mit einem beliebigen x multipliziert wird. Mit dieser Motivation können wir die durch eine Vektornorm induzierte Norm wie folgt definieren:

Definition 3.4 (Induzierte Norm). Die durch eine Norm $\|\cdot\|$ auf $\mathbb{R}^n$ induzierte Norm auf $\mathbb{R}^{m \times n}$ ist gegeben durch

$$\|A\| = \max_{\|x\|=1} \|Ax\|.$$

Das heißt, die induzierte Norm ist die maximale Länge des Bildes eines Einheitsvektors multipliziert mit A .

Da Vektornormen $\|x\| = |c|x|$ erfüllen, ist leicht zu erkennen, dass diese Definition dem Erfordernis entspricht

$$\|A\| = \max_{x \in \mathbb{R}^n \setminus \{0\}} \frac{\|Ax\|}{\|x\|}.$$

Unter diesem Gesichtspunkt ist die durch $\|\cdot\|$ induzierte Norm von A das größte erreichbare Verhältnis der Norm von Ax relativ zu der der Eingabe x .

Diese allgemeine Definition macht es etwas schwierig, die Norm A bei gegebener Matrix A und einer Auswahl von $\|\cdot\|$ zu berechnen. Glücklicherweise können die von vielen gängigen Vektornormen abgeleiteten Matrixnormen vereinfacht werden. Wir geben einige solcher Ausdrücke ohne Beweis an:

- Die induzierte Einnorm von A ist die maximale Summe einer beliebigen Spalte von A :

$$\|A\|_1 = \max_{1 \leq j \leq n} \sum_{i=1}^m |a_{ij}|$$

- Die induzierte ∞ -Norm von A ist die maximale Summe einer beliebigen Zeile von A :

$$\|A\|_\infty = \max_{1 \leq i \leq m} \sum_{j=1}^n |a_{ij}|$$

- Die induzierte Zweinorm oder Spektralnorm von $A \in \mathbb{R}^{n \times n}$ ist die Quadratwurzel des größten Eigenwerts von $A^T A$. Das heißt,

$$\|A\|_2 = \sqrt{\lambda_{\max}(A^T A)} = \sqrt{\max\{\lambda : \text{es existiert } x \in \mathbb{R}^n \text{ mit } A^T A x = \lambda x\}}$$

Zumindest die ersten beiden Normen sind relativ einfach zu berechnen; Auf die dritte werden wir bei der Diskussion von Eigenwertproblemen zurückkommen.

3.3.2 Konditionsnummern

Da wir nun über Werkzeuge zum Messen der Wirkung einer Matrix verfügen, können wir die Bedingungsahl eines linearen Systems definieren, indem wir unsere generische Definition von Bedingungsahlen aus Kapitel 1 anpassen. Wir folgen der in CITE dargestellten Entwicklung.

Angenommen, wir haben eine Störung $\tilde{A}$ einer Matrix A und eine entsprechende Störung $\tilde{b}$. Für jede $\tilde{y} \neq 0$, unter Vernachlässigung technischer Invertierbarkeitsaspekte können wir einen Vektor $x(\tilde{y})$ als Lösung schreiben

$$(A + \tilde{y} \cdot \tilde{A})x(\tilde{y}) = b + \tilde{y} \cdot \tilde{b}.$$

Wenn wir beide Seiten nach $\tilde{y}$ differenzieren und die Produktregel anwenden, erhalten wir folgendes Ergebnis:

$$\tilde{y} A \cdot x + (A + \tilde{y} \cdot \tilde{A}) \frac{dx}{d\tilde{y}} = \tilde{y} \tilde{b}$$

Insbesondere wenn $\ddot{y} = 0$ gilt

$$\ddot{y}A \cdot x(0) + A = \ddot{y}b - d\ddot{y} \ddot{y}=0$$

oder gleichwertig,

$$\frac{dx}{d\tilde{y}} \Big|_{\tilde{y}=0} = A^{-1} (\tilde{y} b - \tilde{y} \tilde{y} A \cdot x(0)).$$

Mit der Taylor-Entwicklung können wir schreiben

$$x(\ddot{y}) = x + \ddot{y}x(0) + O(\ddot{y}^2),$$

wobei wir $x(0) = \text{gestörtes System definieren:}$ $\frac{dx}{d\tilde{y}}|_{\tilde{y}=0}$. Somit können wir den relativen Fehler erweitern, der durch die Lösung entsteht

$$\begin{aligned} \frac{x(\tilde{y}) - \tilde{y}x(0) - x(0)}{\tilde{y} - x(0)} &= \frac{\tilde{y}x(0) + O(\tilde{y}^2)}{\tilde{y} - x(0)} \quad \text{durch die Taylor-Entwicklung} \\ &= \frac{\tilde{y}A - \tilde{y}^1 (\tilde{y}b - \tilde{y}A \cdot x(0)) + O(\tilde{y}x(0)^2)}{\tilde{y} - x(0)} \quad \text{durch die von uns berechnete Ableitung} \\ \tilde{y} \frac{1}{\tilde{y} - x(0)} (A - \tilde{y}^1 \tilde{y}b + A \tilde{y}A \cdot x(0)) &+ O(\tilde{y}^2) \\ x(0) \text{ durch die Dreiecksungleichung } A + B &\tilde{y} A + B \\ \tilde{y} |\tilde{y}|A - \tilde{y}^1 \frac{\tilde{y}b}{x(0)} + \tilde{y}A + O(\tilde{y}^2) &\text{ durch die Identitt } AB \tilde{y} BA \\ = |\tilde{y}|A - \tilde{y}^1 A \frac{\tilde{y}b}{A x(0)} + \frac{\tilde{y}A}{A} + O(\tilde{y}^2) \\ \tilde{y} |\tilde{y}|A - \tilde{y}^1 A \frac{\tilde{y}b}{A x(0)} + \frac{\tilde{y}A}{A} + O(\tilde{y}^2) &\text{ da } Ax(0) \tilde{y} Ax(0) \\ = |\tilde{y}|A - \tilde{y}^1 A \frac{\tilde{y}b}{B} + \frac{\tilde{y}A}{A} + O(\tilde{y}^2) &\text{ da } Ax(0) \tilde{y} Ax(0) \end{aligned}$$

da per Definition $Ax(0) = b$ ist

Hier haben wir einige Eigenschaften der Matrixnorm angewendet, die sich aus entsprechenden Eigenschaften für Vektoren ergeben. Beachten Sie, dass die Summe $\| \tilde{b} \|_1 + \| \tilde{A} \|_1$ die relativen Störungen von A und b kodiert. Von diesem Standpunkt aus haben wir in erster Ordnung den relativen Fehler der Störung des Systems um $\tilde{y}$ unter Verwendung eines Faktors $\tilde{y} \| \tilde{A} \|_1$ begrenzt:

$$\frac{x(\ddot{y}) - \ddot{y}x(0) - x(0)}{\ddot{y} - \ddot{y} \cdot D \cdot \ddot{y}} + O(\ddot{y}^{-2})$$

Auf diese Weise begrenzt die Größe $\tilde{\gamma}$ die Konditionierung linearer Systeme mit A . Aus diesem Grund treffen wir die folgende Definition:

Definition 3.5 (Matrix-Bedingungsnummer). Die Bedingungsnummer von $A \in \mathbb{R}^{n \times n}$ für eine gegebene Matrixnorm $\|\cdot\|$ ist

$$\text{Bedingung } A = \frac{\|A\|}{\|A^{-1}\|}.$$

Wenn A nicht invertierbar ist, nehmen wir die Bedingung $A = \infty$.

Es ist leicht zu erkennen, dass $\text{cond } A \geq 1$ für alle A ist, dass die Skalierung von A keinen Einfluss auf seine Bedingungsnummer hat und dass die Bedingungsnummer der Identitätsmatrix 1 ist. Diese Eigenschaften stehen im Gegensatz zur Determinante, die nach oben und unten skaliert werden kann wenn Sie A erklimmen.

Wenn $\|\cdot\|$ durch eine Vektornorm induziert wird und A invertierbar ist, dann gilt

$$\begin{aligned} \|A\| &= \max_{\|x\|=1} \|Ax\| \quad \text{per Definition} \\ &= \max_{\|y\|=1} \frac{\|y\|}{\|A^{-1}y\|} \quad \text{durch Ersetzen von } y = Ax \\ &= \min_{\|y\|=1} \frac{1}{\|A^{-1}y\|} \quad \text{indem man den Kehrwert bildet} \end{aligned}$$

In diesem Fall ist die Konditionsnummer von A gegeben durch:

$$\text{Leitung } A = \frac{\|A\|}{\|A^{-1}\|} = \frac{\max_{\|x\|=1} \|Ax\|}{\min_{\|y\|=1} \|A^{-1}y\|}$$

Mit anderen Worten, $\text{cond } A$ misst die maximale bis minimal mögliche Ausdehnung eines Vektors x unter A .

Allgemeiner gesagt besteht eine wünschenswerte Stabilitätseigenschaft eines Systems $Ax = b$ darin, dass sich die Lösung x nicht wesentlich ändert, wenn eine Kugel gestört wird. Unsere Motivation für Bedingung A zeigt, dass bei einer kleinen Bedingungsnummer die Änderung von x im Verhältnis zur Änderung von A klein ist, wie in Abbildung NUMBER dargestellt. Andernfalls kann eine kleine Änderung der Parameter des linearen Systems zu großen Abweichungen in x führen; Diese Instabilität kann dazu führen, dass lineare Löser aufgrund von Rundungen und anderen Näherungen während des Lösungsprozesses große Fehler in x machen.

Die Norm $\|A\|$ machen. kann genauso schwierig sein wie die Berechnung $\|A^{-1}\|$. Eine Möglichkeit zum Absenken der vollständigen inversen A -Grenze. Die Bedingung Nummer besteht darin, die Identität $A^{-1}A = I$ anzuwenden. können wir A schreiben $A^{-1} = A^{-1}$. Somit gilt für jedes $x \neq 0$ $A^{-1}Ax = x$. Daher,

$$\text{Bedingung } A = \frac{\|A\|}{\|A^{-1}\|} = \frac{\|AA^{-1}x\|}{\|x\|}.$$

Wir können also die Bedingungsnummer begrenzen, indem wir $A^{-1}x$ für einige Vektoren x auflösen; Natürlich erzeugt die Notwendigkeit eines linearen Löser, $A^{-1}x$ zu finden, eine zirkuläre Abhängigkeit von der Bedingungsnummer, um die Qualität der Schätzung zu bewerten! Wenn $\|\cdot\|$ durch die Zweinorm induziert wird, werden wir in zukünftigen Kapiteln zuverlässigere Schätzungen bereitstellen.

3.4 Probleme

Etwas wie:

- Kernel-Regression als Beispiel für §3.1.1
- Lösung der kleinsten Norm für $Ax = b$, die Matrix der kleinsten Quadrate ist ansonsten invertierbar
- Variationsversionen der Tikhonov-Regularisierung/„Ridge“-Regression (nicht der übliche Ansatz).
dazu, aber was auch immer); Vervollständigung der unbestimmten Geschichte auf diese Weise
- 1 • L Ansätze zur Regularisierung für den Kontrast – zeichnen Sie ein Bild davon, warum mit Sparsity zu rechnen ist, zeichnen Sie
Einheitskreise, zeigen Sie, dass die p -Norm keine Norm für $p < 1$ ist, nehmen Sie den Grenzwert als $p \rightarrow 0$
- Mini-Riesz – Matrix für das innere Produkt ableiten, um zu zeigen, wie der Raum gedreht wird
- tridiagonale Lösung
- Eigenschaften der Konditionsnummer

Kapitel 4

Spaltenräume und QR

Eine Möglichkeit, das lineare Problem $Ax = b$ für x zu interpretieren, besteht darin, dass wir b als lineare Kombination der Spalten von A mit den in x angegebenen Gewichten schreiben möchten. Diese Perspektive ändert sich nicht, wenn wir zulassen, dass $A \in \mathbb{R}^{m \times n}$ nicht quadratisch ist, aber die Lösung existiert möglicherweise nicht oder ist abhängig von der Struktur des Spaltenraums nicht eindeutig. Aus diesen Gründen streben einige Techniken zum Faktorisieren von Matrizen und zum Analysieren linearer Systeme nach einfacheren Darstellungen des Spaltenraums, um die Lösbarkeit eindeutiger zu machen und eine explizitere Spannweite zu erreichen als zeilenbasierte Faktorisierungen wie LU.

4.1 Die Struktur der Normalgleichungen

Wie wir gezeigt haben, ist eine notwendige und hinreichende Bedingung dafür, dass x eine Lösung des Problems der kleinsten Quadrate $Ax \approx b$ ist, dass x die Normalgleichungen $(A^T A)x = A^T b$ erfüllt. Dieser Satz legt nahe, dass die Lösung der kleinsten Quadrate eine recht einfache Erweiterung linearer Techniken ist. Methoden wie die Cholesky-Faktorisierung zeigen auch, dass die spezielle Struktur von Problemen der kleinsten Quadrate zum Vorteil des Löfers genutzt werden kann.

Es gibt jedoch ein großes Problem, das die Verwendung dieses Ansatzes einschränkt. Nehmen wir zunächst einmal an, dass A quadratisch ist; dann können wir schreiben:

$$\begin{aligned} \text{cond } AA^T &= \frac{\lambda_{\max}(AA^T)}{\lambda_{\min}(AA^T)} = \frac{\lambda_{\max}(A) \lambda_{\max}(A^T)}{\lambda_{\min}(A) \lambda_{\min}(A^T)} = \left(\frac{\lambda_{\max}(A)}{\lambda_{\min}(A)} \right)^2 \\ &= \kappa(A)^2 \end{aligned}$$

abhängig von der Wahl des $\cdot$

Das heißt, die Konditionszahl von AA^T ist ungefähr das **Quadrat** der Konditionszahl von A ! Während also generische lineare Strategien bei AA^T funktionieren könnten, wenn das Problem der kleinsten Quadrate „einfach“ ist, erzeugen diese Strategien wahrscheinlich erhebliche Fehler, wenn die Spalten von A nahezu linear abhängig sind, da sie sich nicht direkt mit A befassen.

Intuitiv liegt ein Hauptgrund dafür, dass $\text{cond } A$ groß sein kann, darin, dass die Spalten von A möglicherweise „ähnlich“ aussehen. Stellen Sie sich jede Spalte von A als einen Vektor in $\mathbb{R}^m$ vor. Wenn zwei Spalten a_i und a_j $i \neq j$ erfüllen, würde die a_j -Restlänge der kleinsten Quadrate $b \approx Ax$ wahrscheinlich nicht viel darunter leiden, wenn wir Vielfache von a_i durch Vielfache von a_j ersetzen oder umgekehrt. Dieses breite Spektrum an nahezu – aber nicht vollständig – äquivalenten Lösungen führt zu einer schlechten Konditionierung. Der resultierende Vektor x ist zwar instabil, das Produkt jedoch

Axe bleibt aufgrund unserer Substitution nahezu unverändert. Wenn wir daher $Ax \approx b$ einfach lösen wollen, um b in den Spaltenraum von A zu schreiben, würde jede Lösung ausreichen.

Um solche schlecht konditionierten Probleme zu lösen, werden wir eine alternative Strategie anwenden, die stärker auf den Spaltenraum von A achtet, anstatt Zeilenoperationen wie bei der Gaußschen Eliminierung einzusetzen. Auf diese Weise können wir solche Beinahe-Abhängigkeiten explizit identifizieren und numerisch stabil damit umgehen.

4.2 Orthogonalität

Wir haben festgestellt, wann das Problem der kleinsten Quadrate schwierig ist, können uns aber auch fragen, wann es am einfachsten ist. Wenn wir ein System auf den einfachen Fall reduzieren können, ohne dabei Konditionierungsprobleme hervorzurufen, haben wir einen stabileren Weg gefunden, die in §4.1 erläuterten Probleme zu umgehen.

Offensichtlich ist das am einfachsten zu lösende lineare System $\mathbb{R}^n \times \mathbb{R}^n = \mathbb{R}^n$: Die Lösung ist einfach $x \approx b$! Es ist unwahrscheinlich, dass wir dieses spezielle lineare System explizit in unseren Löser eingeben, es kann jedoch sein, dass wir dies versehentlich beim Lösen der Methode der kleinsten Quadrate tun. Insbesondere selbst wenn $A = \mathbb{R}^n \times \mathbb{R}^n$ – A muss tatsächlich keine quadratische Matrix sein – können wir unter besonders glücklichen Umständen feststellen, dass die normale Matrix $A^T A = \mathbb{R}^n \times \mathbb{R}^n$ erfüllt. Um Verwechslungen mit dem allgemeinen Fall zu vermeiden, verwenden wir zur Darstellung einer solchen Matrix den Buchstaben Q .

Das einfache Beten, dass $Q^T Q = \mathbb{R}^n \times \mathbb{R}^n$ unwahrscheinlich ist, wird zu einer wünschenswerten Lösungsstrategie führen, aber wir können diesen Fall untersuchen, um zu sehen, wie er so günstig wird. Schreiben Sie die Spalten von Q als Vektoren $q_1, \dots, q_n \in \mathbb{R}^m$. Dann lässt sich leicht überprüfen, ob das Produkt $Q^T Q$ die folgende Struktur hat:

$$Q^T Q = \begin{pmatrix} q_1^T q_1 & q_1^T q_2 & \dots & q_1^T q_n \\ q_2^T q_1 & q_2^T q_2 & \dots & q_2^T q_n \\ \vdots & \vdots & \ddots & \vdots \\ q_n^T q_1 & q_n^T q_2 & \dots & q_n^T q_n \end{pmatrix} = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 \end{pmatrix}$$

Setzt man den Ausdruck auf der rechten Seite gleich $\mathbb{R}^n \times \mathbb{R}^n$, erhält man die folgende Beziehung:

$$q_i^T q_j = \begin{cases} 1 & \text{wenn } i = j \\ 0 & \text{wenn } i \neq j \end{cases}$$

Mit anderen Worten, die Spalten von Q haben eine Einheitslänge und sind orthogonal zueinander. Wir sagen, dass sie eine Orthonormalbasis für den Spaltenraum von Q bilden:

Definition 4.1 (Orthonormal; orthogonale Matrix). Eine Menge von Vektoren $\{v_1, \dots, v_k\}$ ist orthonormal, wenn $v_i^T v_i = 1$ für alle i und $v_i^T v_j = 0$ für alle $i \neq j$. Eine quadratische Matrix, deren Spalten orthonormal sind, wird orthogonale Matrix genannt.

Wir haben unsere Diskussion mit der Frage motiviert, wann wir mit $Q^T Q = \mathbb{R}^n \times \mathbb{R}^n$ rechnen können. Nun ist leicht zu erkennen, dass dies geschieht, wenn die Spalten von Q orthonormal sind. Wenn Q außerdem quadratisch und invertierbar ist mit $Q^T Q = \mathbb{R}^n \times \mathbb{R}^n$, dann erhalten wir durch einfaches Multiplizieren beider Seiten dieses Ausdrucks mit Q^{-1} $Q^{-1} Q^T Q = Q^{-1} \mathbb{R}^n \times \mathbb{R}^n$. Somit ist die Lösung von $Qx = b$ in diesem Fall so einfach wie das Multiplizieren beider Seiten durch das transponierte Q .

Orthonormalität hat auch eine starke geometrische Interpretation. Erinnern Sie sich an Kapitel 0, dass wir zwei orthogonale Vektoren a und b als senkrecht betrachten können. Also ein orthonormaler Satz von Vektoren

Ist einfach eine Menge senkrechter Vektoren mit Einheitslänge in $\mathbb{R}^n$, die $\cdot$. Wenn Q orthogonal ist, dann ist seine Wirkung orthogonal die Länge der Vektoren nicht beeinflussen:

$$\|Qx\|^2 = x^T Q^T Q x = x^T I x = x^T x = \|x\|^2$$

Ebenso kann Q den Winkel zwischen zwei Vektoren nicht beeinflussen, da:

$$(Qx) \cdot (Qy) = x^T Q^T Q y = x^T I y = x^T y = x \cdot y$$

Wenn Q orthogonal ist, stellt Q eine Isometrie von $\mathbb{R}^n$ dar, das heißt, es behält Längen und Winkel bei. Mit anderen Worten: Es kann Vektoren drehen oder spiegeln, sie jedoch nicht skalieren oder scheren.

Von einem hohen Niveau aus ist die lineare Algebra orthogonaler Matrizen einfacher, da ihre Wirkung die Geometrie des zugrunde liegenden Raums in keiner nichttrivialen Weise beeinflusst.

4.2.1 Strategie für nicht-orthogonale Matrizen

Außer unter besonderen Umständen sind die meisten unserer Matrizen A bei der Lösung von $Ax = b$ oder dem entsprechenden Problem der kleinsten Quadrate nicht orthogonal, sodass die Maschinerie von §4.2 nicht direkt anwendbar ist. Aus diesem Grund müssen wir einige zusätzliche Berechnungen durchführen, um den allgemeinen Fall mit dem orthogonalen zu verbinden.

Nehmen Sie eine Matrix $A \in \mathbb{R}^{m \times n}$, und bezeichnen seinen Spaltenraum als $\text{col } A$; Denken Sie daran, dass Spalte A die Spanne der Spalten von A darstellt. Nehmen wir nun an, dass eine Matrix $B \in \mathbb{R}^{n \times n}$ invertierbar ist. Wir können eine einfache Beobachtung über den Spaltenraum von AB relativ zu dem

von A machen: **Lemma 4.1** (Spaltenrauminvarianz). Für jedes $A \in \mathbb{R}^{m \times n}$ und das invertierbare $B \in \mathbb{R}^{n \times n}$

$$\text{Spalte } A = \text{Spalte } AB.$$

Nachweisen. Angenommen, $b \in \text{col } A$. Dann gibt es nach der Definition der Multiplikation mit A ein x mit $Ax = b$.

Dann ist $(AB) \cdot (B^{-1}x) = Ax = b$, so $b \in \text{col } AB$. Umgekehrt sei $c \in \text{col } AB$, also existiert y mit $(AB)y = c$. Dann ist $A \cdot (By) = c$, was zeigt, dass c in Spalte A steht. $\square$

Erinnern Sie sich an die „Eliminationsmatrix“-Beschreibung der Gaußschen Eliminierung: Wir begannen mit einer Matrix A und wendeten Zeilenoperationsmatrizen E_i an, so dass die Sequenz $A, E_1 A, E_2 E_1 A, \dots$ dargestellt sequentiell einfachere lineare Systeme. Das obige Lemma schlägt eine alternative Strategie für Situationen vor, in denen uns der Spaltenraum wichtig ist: Wenden Sie Spaltenoperationen durch Nachmultiplikation auf A an, bis die Spalten orthonormal sind. Das heißt, wir erhalten ein Produkt $Q = A E_1 E_2 \dots E_k$, so dass Q orthonormal ist. Solange die E_i invertierbar sind, zeigt das Lemma, dass $\text{col } Q = \text{col } A$. Das Invertieren dieser Operationen ergibt eine Faktorisierung $A = QR$ für $R = E_k^{-1} \dots E_1^{-1} A$. Wenn wir wie bei der LU-

Faktorisierung R sorgfältig entwerfen, ist die Lösung der kleinsten Quadratprobleme $Ax \approx b$ können vereinfacht werden. Insbesondere wenn $A = QR$ ist, können wir die Lösung von $Ax = b$ wie folgt schreiben:

$$\begin{aligned} x &= (A^T A)^{-1} A^T b \\ &= (R^T Q^T Q R)^{-1} R^T Q^T b, \text{ da } A = QR \\ &= (R^T R)^{-1} R^T Q^T b, \text{ da } Q \text{ orthogonal ist} \\ &= R^{-1} (R^T)^{-1} R^T Q^T b \text{ da } (AB)^{-1} = B^{-1} A^{-1} \\ &= R^{-1} Q^T b \end{aligned}$$

Oder äquivalent: $Rx = Q^T b$

Wenn wir also R als Dreiecksmatrix entwerfen, ist die Lösung des linearen Systems $Rx = Qb$ so einfach wie eine Rücksubstitution.

Unsere Aufgabe für den Rest des Kapitels besteht darin, Strategien für eine solche Faktorisierung zu entwerfen.

4.3 Gram-Schmidt-Orthogonalisierung

Unser erster Ansatz zum Finden von QR-Faktorisierungen ist am einfachsten zu beschreiben und zu implementieren, kann jedoch mit numerischen Problemen behaftet sein. Wir verwenden es hier als erste Strategie und werden es dann durch bessere Abläufe verbessern.

4.3.1 Projektionen

Angenommen, wir haben zwei Vektoren a und b . Dann könnten wir leicht fragen: „Welches Vielfache von a kommt am nächsten zu b ?“ Mathematisch entspricht diese Aufgabe der Minimierung von $\|ca - b\|^2$ über alle möglichen $c \in \mathbb{R}$. Wenn wir uns a und b als $n \times 1$ -Matrizen und c als 1×1 -Matrix vorstellen, dann ist dies nichts weiter als ein unkonventionelles Problem der kleinsten Quadrate $\cdot c$. In diesem Fall lauten die Normalgleichungen $a \cdot c = ab$, oder

$$c = \frac{a \cdot b}{a \cdot a} = \frac{a \cdot b}{\|a\|^2}.$$

Wir bezeichnen diese Projektion von b auf a als:

$$\text{proj}_A b = ca = \frac{a \cdot b}{a \cdot a} a = \frac{a \cdot b}{\|a\|^2} a$$

Offensichtlich ist $\text{proj}_A b$ parallel zu a . Was ist mit dem Rest $b - \text{proj}_A b$? Wir können ein einfaches machen
Berechnung um herauszufinden:

$$\begin{aligned} a \cdot (b - \text{proj}_A b) &= a \cdot b - a \cdot \frac{a \cdot b}{\|a\|^2} a \\ &= a \cdot b - \frac{a \cdot b}{\|a\|^2} (a \cdot a) \\ &= a \cdot b - a \cdot b \\ &= 0 \end{aligned}$$

Somit haben wir b in eine Komponente parallel zu a und eine weitere Komponente orthogonal zerlegt
zu einem.

Nehmen wir nun an, dass $a^1, a^2, \dots, a^k$ orthonormal sind; Wir werden Vektoren mit Einheitslänge mit Hüten überziehen. Dann können wir für jedes einzelne i Folgendes sehen:

$$\text{proj}_{a^i} b = (a^i \cdot b) a^i$$

Der Normterm kommt nicht vor, da per Definition $\hat{a}^i = 1$ ist. Wir könnten b auf die Spanne $\{\hat{a}^1, \dots, \hat{a}^k\}$ projizieren, indem wir die folgende Energie über $c_1, \dots, c_k \in \mathbb{R}$:

$$\begin{aligned} \|c_1 \hat{a}^1 + c_2 \hat{a}^2 + \dots + c_k \hat{a}^k - b\|^2 &= \sum_{i=1}^k \sum_{j=1}^k c_i c_j (\hat{a}^i \cdot \hat{a}^j) - 2b \cdot \sum_{i=1}^k c_i \hat{a}^i + b \cdot b \\ &\quad \text{durch Anwendung von } \|v\|^2 = v \cdot v \\ &= \sum_{i=1}^k c_i^2 - 2c_i b \cdot \hat{a}^i + b \cdot b \quad \text{da die } \hat{a}^i \text{ orthonormal sind} \end{aligned}$$

Beachten Sie, dass der zweite Schritt hier nur aufgrund der Orthonormalität gültig ist. Zumindest ist die Ableitung nach c_i für jedes c_i null, was Folgendes ergibt:

$$c_i = \hat{a}^i \cdot b$$

Wir haben also gezeigt, dass, wenn $\hat{a}^1, \dots, \hat{a}^k$ orthonormal sind, die folgende Beziehung gilt:

$$\text{projspan}\{\hat{a}^1, \dots, \hat{a}^k\} b = (\hat{a}^1 \cdot b) \hat{a}^1 + \dots + (\hat{a}^k \cdot b) \hat{a}^k$$

Dies ist lediglich eine Erweiterung unserer Projektionsformel, und anhand eines ähnlichen Beweises ist dies leicht zu erkennen

$$\hat{a}^i \cdot (b - \text{projspan}\{\hat{a}^1, \dots, \hat{a}^k\} b) = 0.$$

Das heißt, wir haben b in eine Komponente parallel zur Spanne der $\hat{a}^i$ und ein dazu senkrechtes Residuum aufgeteilt.

4.3.2 Gram-Schmidt-Orthogonalisierung

Unsere obigen Beobachtungen führen zu einem einfachen Algorithmus zur Orthogonalisierung oder zum Finden einer orthogonalen Basis $\{\hat{a}^1, \dots, \hat{a}^k\}$, deren Spanne dieselbe ist wie die einer Menge linear unabhängiger Eingabevektoren $\{v_1, \dots, v_k\}$:

1. Einstellen

$$\hat{a}^1 \leftarrow \frac{v_1}{\|v_1\|}.$$

Das heißt, wir nehmen an, dass $\hat{a}^1$ ein Einheitsvektor parallel zu v_1 ist.

2. Für i von 2 bis k ,

(a) Berechnen Sie die Projektion

$$p_i \leftarrow \text{projspan}\{\hat{a}^1, \dots, \hat{a}^{i-1}\} v_i.$$

Per Definition sind $\hat{a}^1, \dots, \hat{a}^{i-1}$ orthonormal, daher gilt unsere obige Formel.

(b) Definieren

$$\hat{a}^i \leftarrow \frac{v_i - p_i}{\|v_i - p_i\|}.$$

Diese als „Gram-Schmidt-Orthogonalisierung“ bekannte Technik ist eine einfache Anwendung unserer obigen Diskussion. Der Schlüssel zum Beweis dieser Technik besteht darin, zu beachten, dass $\text{span}\{v_1, \dots, v_i\} = \text{span}\{a^1, \dots, a^i\}$ für jedes $i \in \{1, \dots, k\}$. Schritt 1 macht dies eindeutig für $i = 1$ der Fall, und für $i > 1$ entfernt die Definition von a^i in Schritt 2b einfach die Projektion auf die Vektoren, die wir bereits gesehen haben.

Wenn wir mit einer Matrix A beginnen, deren Spalten $v_1, \dots, v_k$ sind, können wir Gram Schmidt als eine Reihe von Spaltenoperationen für A implementieren. Die Division der Spalte i von A durch ihre Norm entspricht der Nachmultiplikation von A mit $k \times k$ Diagonalmatrix. Ebenso ist das Subtrahieren der Projektion einer Spalte auf die orthonormalen Spalten links davon wie in Schritt 2 gleichbedeutend mit einer Nachmultiplikation mit einer oberen Dreiecksmatrix: Stellen Sie sicher, dass Sie verstehen, warum dies der Fall ist! Somit gilt unsere Diskussion in §4.2.1 und wir können Gram-Schmidt verwenden, um eine Faktorisierung $A = QR$ zu erhalten.

Leider kann der Gram-Schmidt-Algorithmus aufgrund des Subtraktionsschritts zu schwerwiegenden numerischen Instabilitäten führen. Nehmen wir zum Beispiel an, wir stellen die Vektoren $v_1 = (1, 1)$ und $v_2 = (1 + \tilde{y}, 1)$ als Eingabe für Gram-Schmidt für $0 < \tilde{y} < 1$ bereit. Beachten Sie, dass eine offensichtliche Basis für $\text{span}\{v_1, v_2\} = \{(1, 0), (0, 1)\}$. Wenn wir jedoch Gram-Schmidt anwenden, erhalten wir:

$$\begin{aligned} a^1 &= \frac{v_1}{\|v_1\|} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \\ p_2 &= \frac{v_2 \cdot a^1}{\|a^1\|} a^1 = \frac{1 + \tilde{y}}{2} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \\ v_2 - p_2 &= \begin{pmatrix} 1 + \tilde{y} \\ 1 \end{pmatrix} - \frac{1 + \tilde{y}}{2} \begin{pmatrix} 1 \\ 1 \end{pmatrix} = \begin{pmatrix} \frac{1 + \tilde{y}}{2} \\ \tilde{y} \end{pmatrix} \end{aligned}$$

Beachten Sie, dass $v_2 - p_2 = (\tilde{y}/2) \cdot \tilde{y}$, sodass die Berechnung von a^2 eine Division durch einen Skalar in der Größenordnung von $\tilde{y}$ erfordert. Die Division durch kleine Zahlen ist eine instabile numerische Operation, die wir vermeiden sollten.

4.4 Haushaltstransformationen

In §4.2.1 haben wir die Konstruktion der QR-Faktorisierung durch Postmultiplikation und Spaltenoperationen motiviert. Diese Konstruktion ist im Kontext der Analyse von Spaltenräumen sinnvoll, aber wie wir bei unserer Ableitung des Gram-Schmidt-Algorithmus gesehen haben, können die resultierenden numerischen Techniken instabil sein.

Anstatt mit A zu beginnen und nachträglich mit Spaltenoperationen zu multiplizieren, um $Q = A E_1 \cdots E_k$ zu erhalten, können wir jedoch unsere High-Level-Strategie vor der Gaußschen Eliminierung bewahren. Das heißt, wir können mit A beginnen und mit orthogonalen Matrizen Q_i vormultiplizieren, um $Q_k \cdots Q_1 A = R$ zu erhalten; Diese Q_i wirken wie Zeilenoperationen und eliminieren Elemente von A , bis das resultierende Produkt R die obere Dreiecksform hat. Dank der Orthogonalität der Q_i können wir dann $A = QR$ schreiben und so den QR-Faktor erhalten, da das Produkt orthogonaler Matrizen orthogonal ist.

Die Zeilenoperationsmatrizen, die wir bei der Gaußschen Eliminierung und LU verwendet haben, reichen für die QR-Faktorisierung nicht aus, da sie nicht orthogonal sind. Es wurde eine Reihe von Alternativen vorgeschlagen; Wir werden eine gemeinsame Strategie vorstellen, die 1958 von Alston Scott Householder eingeführt wurde.

Der Raum orthogonaler $n \times n$ -Matrizen ist sehr groß, daher müssen wir einen kleineren Raum von Q_i finden, mit dem wir einfacher arbeiten können. Aus unseren geometrischen Diskussionen in §4.2 wissen wir, dass orthogonale Matrizen Winkel und Längen beibehalten müssen, sodass sie intuitiv nur Vektoren drehen und reflektieren können.

Glücklicherweise lassen sich die Überlegungen leicht in Form von Projektionen beschreiben, wie in Abbildung NUMMER dargestellt. Angenommen, wir haben einen Vektor b , den wir über einen Vektor v spiegeln möchten. Wir haben gezeigt, dass das Residuum $r_v b$ steht senkrecht auf v . Wie in Abbildung NUMMER die Differenz $\tilde{y} b - \tilde{y}_v b$ reflektiert über v .

proj_v ist. Wir können unsere Reflexionsformel wie folgt erweitern:

$$\begin{aligned} 2\text{proj}_v \tilde{y} b &= 2 \frac{v \cdot b}{v \cdot v} v \tilde{y} b \text{ nach Definition der Projektion} \\ &= 2v \cdot \frac{vb}{V_v} \tilde{y} b \text{ unter Verwendung der Matrixschreibweise} \\ &= \frac{2vv}{V_v} \tilde{y} \text{In} \times n b \end{aligned}$$

$\tilde{y} \tilde{y} H_v b$, wobei das Negativ eingeführt wird, um es an andere Behandlungen anzupassen

Daher können wir uns die Spiegelung von b über v so vorstellen, als würden wir einen linearen Operator $\tilde{y} H_v$ auf b anwenden! Natürlich ist H_v ohne das Negativ immer noch orthogonal, daher werden wir es von nun an verwenden.

Angenommen, wir machen den ersten Schritt der Vorwärtssubstitution während der Gaußschen Eliminierung. Dann möchten wir A mit einer Matrix vormultiplizieren, die die erste Spalte von A , die wir mit a bezeichnen werden, zu einem Vielfachen des ersten Identitätsvektors e_1 führt. Mit anderen Worten, wir wollen für ein $c \in \mathbb{R}$:

$$\begin{aligned} ce_1 &= H_v a \\ &= \text{In} \times n \tilde{y} \frac{2vv}{V_v} A \\ &= a \tilde{y} 2v \frac{va}{V_v} \end{aligned}$$

Begriffe in Shows verschieben

$$v = (a \tilde{y} ce_1) \cdot \frac{V_v}{2va}$$

Mit anderen Worten, v muss parallel zur Differenz $a \tilde{y} ce_1$ sein. Tatsächlich hat die Skalierung von v keinen Einfluss auf die Formel für H_v , sodass wir $v = a \tilde{y} ce_1$ wählen können. Dann müssen wir es haben, damit unsere Beziehung hält

$$\begin{aligned} 1 &= \frac{V_v}{2va} \\ &= \frac{A^2 \tilde{y} 2ce_1 \cdot a + c^2}{2(a \cdot a \tilde{y} ce_1 \cdot a)} \\ \text{Oder } 0 &= a^2 \tilde{y} c^2 = \tilde{y} c = \pm a \end{aligned}$$

Mit dieser Wahl von c haben wir gezeigt:

$$H_v A = \begin{pmatrix} c & \times & \times & \times \\ \tilde{y} & 0 & \times & \times \\ & \vdots & \vdots & \vdots \\ & 0 & \times & \times \end{pmatrix} \tilde{y}$$

Wir haben gerade einen Schritt erreicht, der der Vorwärtseliminierung nur mit orthogonalen Matrizen ähnelt!

Weitergehend haben wir in der Notation von CITE während des k -ten Schritts der Triangularisierung einen Vektor a , das wir in zwei Komponenten aufteilen können:

$$a = \begin{pmatrix} a_1 \\ a_2 \end{pmatrix}$$

Hier gilt $a_1 \in \mathbb{R}^k$ und $a_2 \in \mathbb{R}^{m-k}$. Wir möchten v so finden

$$Hv = \begin{pmatrix} a_1 \\ 0 \\ \vdots \\ 0 \end{pmatrix}$$

Folgt man einer parallelen Ableitung zur obigen, lässt sich das leicht zeigen

$$v = \begin{pmatrix} 0 \\ a_2 \end{pmatrix} \quad \text{bis zu } c_k$$

führt genau diese Transformation durch, wenn $c = \pm a_2$; Normalerweise wählen wir das Vorzeichen von c , um eine Aufhebung zu vermeiden, indem wir dafür sorgen, dass es das Vorzeichen hat, das dem des k -ten Werts in a entgegengesetzt ist.

Der Algorithmus für Householder QR ist daher ziemlich einfach. Für jede Spalte von A berechnen wir v , indem wir die unteren Elemente der Spalte vernichten, und wenden Hv auf A an. Das Endergebnis ist eine obere Dreiecksmatrix $R = H_v \cdots H_1 A$. Die orthogonale Matrix Q ist durch das Produkt H which gegeben H kann implizit als Liste von Vektoren v gespeichert werden, die in das untere Dreieck wie oben gezeigt v_1^n passt.

4.5 Reduzierte QR-Faktorisierung

Wir schließen unsere Diskussion ab, indem wir zum allgemeinsten Fall $Ax = b$ zurückkehren, wenn $A \in \mathbb{R}^{m \times n}$ nicht quadratisch ist. Beachten Sie, dass beide in diesem Kapitel besprochenen Algorithmen nichtquadratische Matrizen A in Produkte QR faktorisieren können, die Ausgabe jedoch etwas anders ist:

- Bei der Anwendung von Gram-Schmidt führen wir Spaltenoperationen an A durch, um Q durch Orthogonalisierung zu erhalten. Aus diesem Grund ist die Dimension von A die von Q , was $Q \in \mathbb{R}^{m \times n}$ und $R \in \mathbb{R}^{n \times n}$ ergibt
- Bei Verwendung von Householder-Reflexionen erhalten wir Q als Produkt einer Zahl von $m \times m$ Reflexionsmatrizen, sodass $R \in \mathbb{R}^{m \times n}$ bleibt.

Angenommen, wir befinden uns im typischen Fall der kleinsten Quadrate, für den $m > n$ gilt. Aufgrund der numerischen Stabilität bevorzugen wir immer noch die Householder-Methode, aber jetzt ist die $m \times m$ -Matrix Q möglicherweise zu groß zum Speichern! Zum Glück wissen wir, dass R die obere Dreiecksform hat. Betrachten Sie beispielsweise die Struktur einer 5×3 -Matrix R :

$$R = \begin{pmatrix} \times & \times & \times \\ 0 & \times & \times \\ 0 & 0 & \times \\ \hline 0 & 0 & 0 & 0 \\ 0 & & & \end{pmatrix}$$

Es ist leicht zu erkennen, dass alles unterhalb des oberen $n \times n$ -Quadrats von R Null sein muss, was eine Vereinfachung ergibt. Kation:

$$A = QR = Q_1 Q_2 \begin{pmatrix} R_1 & \\ & 0 \end{pmatrix} = Q_1 R_1$$

Hier enthält $Q_1 \in \mathbb{R}^{m \times n}$ und $R_1 \in \mathbb{R}^{n \times n}$ noch das obere Dreieck von R . Dies wird als „reduziertes“ Dreieck bezeichnet. QR-Faktorisierung von A , da die Spalten von Q_1 eine Basis für den Spaltenraum von A und nicht für den gesamten $\mathbb{R}^m$ enthalten; es nimmt viel weniger Platz ein. Beachten Sie, dass die Diskussion in §4.2.1 weiterhin gilt, sodass die reduzierte QR-Faktorisierung auf ähnliche Weise für kleinste Quadrate verwendet werden kann.

4.6 Probleme

- Tridiagonalisierung mit Householder
- Gegebenheiten
- Unterbestimmter QR

Kapitel 5

Eigenvektoren

Wir wenden uns nun einem nichtlinearen Problem über Matrizen zu: der Ermittlung ihrer Eigenwerte und Eigenvektoren. Eigenvektoren x und ihre entsprechenden Eigenwerte λ einer quadratischen Matrix A werden durch die Gleichung $Ax = \lambda x$ bestimmt. Es gibt viele Möglichkeiten zu erkennen, dass dieses Problem nichtlinear ist. Zum Beispiel gibt es ein Produkt der Unbekannten λ und x , und um die triviale Lösung $x = 0$ zu vermeiden, beschränken wir $x = 1$; Diese Einschränkung ist eher kreisförmig als linear. Dank dieser Struktur werden sich unsere Methoden zum Finden von Eigenräumen erheblich von Techniken zum Lösen und Analysieren linearer Gleichungssysteme unterscheiden.

5.1 Motivation

Trotz der willkürlich wirkenden Form der Gleichung $Ax = \lambda x$ entsteht in vielen Fällen natürlicherweise das Problem, Eigenvektoren und Eigenwerte zu finden. Wir motivieren unsere Diskussion mit einigen Beispielen unten.

5.1.1 Statistik

Angenommen, wir verfügen über Maschinen zum Sammeln mehrerer statistischer Beobachtungen über eine Sammlung von Gegenständen. Beispielsweise können wir in einer medizinischen Studie Alter, Gewicht, Blutdruck und Herzfrequenz von 100 Patienten erfassen. Dann kann jeder Patient i durch einen Punkt x_i in $\mathbb{R}^4$ dargestellt werden, der diese vier Werte speichert.

Natürlich können solche Statistiken eine starke Korrelation aufweisen. Beispielsweise kann es bei Patienten mit höherem Blutdruck wahrscheinlich zu einem höheren Gewicht oder einer höheren Herzfrequenz kommen. Obwohl wir unsere Daten in $\mathbb{R}^4$ gesammelt haben, können sie aus diesem Grund in Wirklichkeit – bis zu einem gewissen Grad – in einem niedrigerdimensionalen Raum leben und die Beziehungen zwischen den verschiedenen Variablen besser erfassen.

Nehmen wir zunächst einmal an, dass es tatsächlich einen eindimensionalen Raum gibt, der unserem Datensatz nahekommt. Dann erwarten wir, dass alle Datenpunkte nahezu parallel zu einem Vektor v sind, sodass jeder als $x_i \approx c_i v$ für verschiedene $c_i \in \mathbb{R}$ geschrieben werden kann. Von zuvor wissen wir, dass die beste Näherung für x_i parallel zu v $\text{proj}_v x_i$ ist:

$$\begin{aligned} \text{proj}_v x_i &= \frac{x_i \cdot v}{v \cdot v} v \text{ per Definition} \\ &= (x_i \cdot \hat{v}) \hat{v}, \text{ da } \hat{v} \cdot \hat{v} = 1 \end{aligned}$$

Hier definieren wir $\hat{v} = v/v$. Natürlich spielt die Größe von v für das vorliegende Problem keine Rolle, daher ist es sinnvoll, den Raum der Einheitsvektoren $\hat{v}$ zu durchsuchen.

Nach dem Muster der kleinsten Quadrate haben wir ein neues Optimierungsproblem:

$$\begin{aligned} & \text{minimiere } \|\tilde{y} - \text{proj}_{\hat{v}} \tilde{x}\|^2 \\ & \text{so dass } \|\hat{v}\| = 1 \end{aligned}$$

Wir können unser Optimierungsziel etwas vereinfachen:

$$\begin{aligned} \|\tilde{y} - \text{proj}_{\hat{v}} \tilde{x}\|^2 &= \|\tilde{y} - (\tilde{x} \cdot \hat{v}) \hat{v}\|^2 \quad \text{per Definition von Projektion} \\ &= \|\tilde{y}\|^2 - 2 \tilde{y} \cdot (\tilde{x} \cdot \hat{v}) \hat{v} + \|\tilde{x} \cdot \hat{v}\|^2 \quad \text{da } \|\hat{v}\| = 1 \text{ und } w^2 = w \cdot w \\ &= \text{konst.} - 2 \tilde{y} \cdot (\tilde{x} \cdot \hat{v}) \end{aligned}$$

Diese Ableitung zeigt, dass wir ein äquivalentes Optimierungsproblem lösen können:

$$\begin{aligned} & \text{maximiere } \tilde{X} \hat{v} \\ & \text{so dass } \|\hat{v}\|^2 = 1, \end{aligned}$$

wobei die Spalten von X die Vektoren $\tilde{x}_i$ sind. Beachten Sie, dass $\tilde{X} \hat{v} = \tilde{X} \tilde{X}^T \hat{v}$, sodass nach Beispiel 0.27 der Vektor $\hat{v}$ dem Eigenvektor von $\tilde{X} \tilde{X}^T$ mit dem höchsten Eigenwert entspricht. Der Vektor $\hat{v}$ ist als erste Hauptkomponente des Datensatzes bekannt.

5.1.2 Differentialgleichungen

Viele physikalische Kräfte können als Funktionen der Position geschrieben werden. Beispielsweise kann die von einer Feder ausgeübte Kraft zwischen zwei Teilchen an den Positionen x und y in $\mathbb{R}^3$ nach dem Hookeschen Gesetz als $k(x - y)$ geschrieben werden; Solche Federkräfte werden in vielen Simulationssystemen zur Annäherung an die Kräfte verwendet, die Stoffe zusammenhalten. Obwohl diese Kräfte nicht unbedingt eine lineare Position haben, nähern wir sie uns oft linear an. Insbesondere in einem physikalischen System mit n Teilchen kodieren Sie die Positionen aller Teilchen gleichzeitig in einem Vektor $X \in \mathbb{R}^{3n}$. Wenn wir dann eine solche Näherung annehmen, können wir schreiben, dass die Kräfte im System für eine Matrix A ungefähr $F \approx AX$ sind.

Erinnern Sie sich an Newtons zweites Bewegungsgesetz $F = ma$, oder Kraft ist gleich Masse mal Beschleunigung. In unserem Kontext können wir eine diagonale Massenmatrix $M \in \mathbb{R}^{3n \times 3n}$ schreiben, die die Masse jedes Teilchens im System enthält. Dann wissen wir $F = MX$, wobei eine Primzahl die zeitliche Differenzierung bezeichnet. Natürlich ist $\dot{X} = (V)$, also haben wir am Ende ein Gleichungssystem erster Ordnung:

$$\frac{d}{dt} \begin{pmatrix} X \\ V \end{pmatrix} = \begin{pmatrix} 0 & I_{3n \times 3n} \\ M^{-1}A & 0 \end{pmatrix} \begin{pmatrix} X \\ V \end{pmatrix}$$

Hier berechnen wir gleichzeitig beide Positionen in $X \in \mathbb{R}^{3n}$ und Geschwindigkeiten $V \in \mathbb{R}^{3n}$ aller n Teilchen als Funktionen der Zeit.

Allgemeiner gesehen treten Differentialgleichungen der Form $\dot{x} = Ax$ in vielen Zusammenhängen auf, einschließlich der Simulation von Stoffen, Federn, Hitze, Wellen und anderen Phänomenen. Angenommen, wir kennen Eigenvektoren

$x_1, \dots, x_k$ von A , so dass $Ax_i = \tilde{y}_i x_i$.
Eigenvektoren schreiben, als

Wenn wir die Anfangsbedingung der Differentialgleichung anhand der

$$x(0) = c_1 x_1 + \dots + c_k x_k,$$

dann kann die Lösung der Gleichung in geschlossener Form geschrieben werden:

$$x(t) = c_1 e^{\tilde{y}_1 t} x_1 + \dots + c_k e^{\tilde{y}_k t} x_k,$$

Diese Lösung ist leicht von Hand zu überprüfen. Das heißt, wenn wir die Anfangsbedingungen dieser Differentialgleichung anhand der Eigenvektoren von A schreiben, dann kennen wir ihre Lösung für alle Zeiten $t \geq 0$ kostenlos.

Natürlich ist diese Formel nicht das Ende der Geschichte der Simulation: Das Finden des vollständigen Satzes von Eigenvektoren von A ist teuer, und A kann sich im Laufe der Zeit ändern.

5.2 Spektrale Einbettung

Angenommen, wir haben eine Sammlung von n Elementen in einem Datensatz und ein Maß $w_{ij} \geq 0$ dafür, wie ähnlich jedes Elementpaar i und j ist; Wir gehen davon aus, dass $w_{ij} = w_{ji}$ ist. Vielleicht erhalten wir zum Beispiel eine Sammlung von Fotos und verwenden w_{ij} , um die Ähnlichkeit ihrer Farbverteilungen zu vergleichen. Möglicherweise möchten wir die Fotos nach ihrer Ähnlichkeit sortieren, um die Betrachtung und Erkundung der Sammlung zu vereinfachen.

Ein Modell zum Ordnen der Sammlung könnte darin bestehen, jedem Element i eine Nummer x_i zuzuweisen und zu verlangen, dass ähnlichen Objekten ähnliche Nummern zugewiesen werden. Mithilfe der Energie können wir messen, wie gut eine Aufgabe ähnliche Objekte gruppiert

$$E(x) = \sum_{i,j} w_{ij} (x_i - x_j)^2.$$

Das heißt, $E(x)$ verlangt, dass die Elemente i und j mit hohen Ähnlichkeitswerten auf nahegelegene Werte abgebildet werden.

Natürlich ergibt die Minimierung von $E(x)$ ohne Einschränkungen ein offensichtliches Minimum: $x_i = \text{const.}$ für alle i . Durch das Hinzufügen einer Einschränkung $\sum x_i^2 = 1$ wird diese konstante Lösung nicht entfernt! Insbesondere ergibt die Annahme von $x_i = 1/\sqrt{n}$ für alle i auf uninteressante Weise $\sum x_i^2 = 1$ und $E(x) = 0$. Daher müssen wir auch diesen Fall entfernen:

$$\begin{aligned} &E(x) \text{ minimieren} \\ &\text{so dass } \sum x_i^2 = 1 \\ &\quad \sum x_i = 0 \end{aligned}$$

Beachten Sie, dass unsere zweite Einschränkung verlangt, dass die Summe von x Null ist.

Noch einmal können wir die Energie vereinfachen:

$$\begin{aligned} E(x) &= \sum_{i,j} w_{ij} (x_i - x_j)^2 \\ &= \sum_{i,j} w_{ij} (x_i^2 - 2x_i x_j + x_j^2) \\ &= \sum_i a_i x_i^2 - 2 \sum_{i,j} w_{ij} x_i x_j + \sum_j b_j x_j^2 \\ &\quad \text{für } a_i = \sum_j w_{ij} \text{ und } b_j = \sum_i w_{ij} \\ &= x^T (A - 2W + B)x \text{ wobei } \text{diag}(A) = a \text{ und } \text{diag}(B) = b = x^T (2A - 2W)x \text{ aufgrund} \\ &\quad \text{der Symmetrie von } W \end{aligned}$$

Es lässt sich leicht überprüfen, dass 1 ein Eigenvektor von $2A \tilde{y} 2W$ mit dem Eigenwert 0 ist. Interessanter ist, dass der Eigenvektor, der dem zweitkleinsten Eigenwert entspricht, der Lösung unseres obigen Minimierungsziels entspricht! (TODO: KKT-Beweis aus Vorlesung hinzufügen)

5.3 Eigenschaften von Eigenvektoren

Wir haben eine Vielzahl von Anwendungen etabliert, die eine Eigenraumberechnung benötigen. Bevor wir jedoch Algorithmen für diesen Zweck untersuchen können, werden wir die Struktur des Eigenwertproblems genauer untersuchen.

Wir können mit einigen Definitionen beginnen, die an dieser Stelle wahrscheinlich offensichtlich sind:

Definition 5.1 (Eigenwert und Eigenvektor). Ein Eigenvektor $x \neq 0$ einer Matrix $A \in \mathbb{R}^{n \times n}$ ist jeder Vektor, der $Ax = \tilde{y}x$ für ein $\tilde{y} \in \mathbb{R}$ erfüllt; das entsprechende $\tilde{y}$ wird als Eigenwert bezeichnet. Komplexe Eigenwerte und Eigenvektoren erfüllen die gleichen Beziehungen mit $\tilde{y} \in \mathbb{C}$ und $x \in \mathbb{C}^n$.

Definition 5.2 (Spektrum und Spektralradius). Das Spektrum von A ist die Menge der Eigenwerte von A . Der Spektralradius $\rho(A)$ ist der Eigenwert $\tilde{y}$, der $|\tilde{y}|$ maximiert.

Der Maßstab eines Eigenvektors ist nicht wichtig. Insbesondere ergibt die Skalierung eines Eigenvektors x mit c $A(cx) = cAx = c\tilde{y}x = \tilde{y}(cx)$, sodass cx ein Eigenvektor mit demselben Eigenwert ist. Wir schränken unsere Suche oft ein, indem wir eine Einschränkung $\|x\| = 1$ hinzufügen. Selbst diese Einschränkung beseitigt die Mehrdeutigkeit nicht vollständig, da nun $\pm x$ beide Eigenvektoren mit demselben Eigenwert sind.

Die algebraischen Eigenschaften von Eigenvektoren und Eigenwerten könnten leicht ein Buch füllen. Wir werden unsere Diskussion auf einige wichtige Theoreme beschränken, die den Entwurf numerischer Algorithmen beeinflussen; Wir werden die Entwicklung von CITE AXLER verfolgen. Damit unsere Suche nicht umsonst ist, sollten wir zunächst prüfen, ob jede Matrix mindestens einen Eigenvektor hat. Unsere übliche Strategie besteht darin, zu beachten, dass, wenn $\tilde{y}$ ein Eigenwert mit $Ax = \tilde{y}x$ ist, $(A - \tilde{y}I_n)x = 0$; somit ist $\tilde{y}$ genau dann ein Eigenwert, wenn die Matrix $A - \tilde{y}I_n$ nicht vollrangig ist.

Lemma 5.1 (CITE-Theorem 2.1). Jede Matrix $A \in \mathbb{R}^{n \times n}$ hat mindestens einen (komplexen) Eigenvektor.

Nachweisen. Nehmen Sie einen beliebigen Vektor $x \in \mathbb{R}^n \setminus \{0\}$. Die Menge $\{x, Ax, A^2x, \dots, A^nx\}$ muss linear abhängig sein, da sie $n + 1$ Vektoren in n Dimensionen enthält. Es gibt also Konstanten $c_0, \dots, c_n \in \mathbb{R}$ mit $c_n \neq 0$, so dass

$$0 = c_0x + c_1Ax + \dots + c_nA^nx.$$

Wir können ein Polynom aufschreiben

$$f(z) = c_0 + c_1z + \dots + c_nz^n.$$

Nach dem Fundamentalsatz der Algebra gibt es n Wurzeln $z_i \in \mathbb{C}$ mit $f(z) = c_n(z - z_1)$

$$(z - z_2) \cdots (z - z_n).$$

Dann haben wir:

$$\begin{aligned} 0 &= c_0x + c_1Ax + \dots + c_nA^nx \\ &= (c_0I_n + c_1A + \dots + c_nA^n)x \\ &= c_n(A - z_1I_n) \cdots (A - z_nI_n)x \text{ durch unsere Faktorisierung} \end{aligned}$$

Somit hat mindestens ein $A - z_iI_n$ einen Nullraum, was zeigt, dass es je nach Bedarf v mit $Av = z_iv$ gibt. □

Es gibt noch eine weitere Tatsache, die es wert ist, überprüft zu werden, um unsere Diskussion über Eigenvektorberechnungen zu motivieren
Stationierung:

Lemma 5.2 (CITE Proposition 2.2). Eigenvektoren, die unterschiedlichen Eigenwerten entsprechen, müssen linear unabhängig sein.

Nachweisen. Angenommen, dies ist nicht der Fall. Dann gibt es Eigenvektoren $x_1, \dots, x_k$ mit unterschiedlichen Eigenwerten $\tilde{y}_1, \dots, \tilde{y}_k$, die linear abhängig sind. Dies impliziert, dass es Koeffizienten $c_1, \dots, c_k$ nicht alle Null mit $0 = c_1 x_1 + \dots + c_k x_k$.

Wenn wir mit der Matrix $(A - \tilde{y}_1 I_n) \dots (A - \tilde{y}_k I_n)$ vormultiplizieren, finden wir:

$$0 = (A - \tilde{y}_1 I_n) \dots (A - \tilde{y}_k I_n)(c_1 x_1 + \dots + c_k x_k) = c_1 (\tilde{y}_1 - \tilde{y}_2) \dots (\tilde{y}_1 - \tilde{y}_k) x_1 \text{ da } Ax_i = \tilde{y}_i x_i$$

Da alle $\tilde{y}_i$ unterschiedlich sind, ergibt sich $c_1 = 0$. Ein ähnlicher Beweis zeigt, dass die restlichen c_i Null sein müssen, was der linearen Abhängigkeit widerspricht. $\square$

Dieses Lemma zeigt, dass eine $n \times n$ -Matrix höchstens n verschiedene Eigenwerte haben kann, da eine Menge von n Eigenwerten n linear unabhängige Vektoren ergibt. Die maximale Anzahl linear unabhängiger Eigenvektoren, die einem einzelnen Eigenwert $\tilde{y}$ entsprechen, wird als geometrische Vielfachheit von $\tilde{y}$ bezeichnet.

Es stimmt jedoch nicht, dass eine Matrix genau n linear unabhängige Eigenvektoren haben muss. Dies ist bei vielen Matrizen der Fall, die wir als nicht fehlerhaft bezeichnen:

Definition 5.3 (Nicht defekt). Eine Matrix $A \in \mathbb{R}^{n \times n}$ ist nichtdefekt oder diagonalisierbar, wenn ihre Eigenvektoren $\mathbb{R}^n$ aufspannen.

Wir nennen eine solche Matrix aus folgendem Grund diagonalisierbar: Wenn eine Matrix diagonalisierbar ist, dann hat sie n Eigenvektoren $x_1, \dots, x_n \in \mathbb{R}^n$ mit entsprechenden (ggf. nicht eindeutigen) Eigenwerten $\tilde{y}_1, \dots, \tilde{y}_n$.

Nehmen Sie die Spalten von X als die Vektoren x_i und definieren Sie D als die Diagonalmatrix mit den Eigenwerten $\tilde{y}_1, \dots, \tilde{y}_n$ entlang der Diagonale. Dann gilt per Definition der Eigenwerte $AX = XD$; Dies ist einfach eine „gestapelte“ Version von $Ax_i = \tilde{y}_i x_i$.

Mit anderen Worten,

$$D = X^{-1}AX,$$

bedeutet, dass A durch eine Ähnlichkeitstransformation $A \mapsto X^{-1}AX$ diagonalisiert wird: **Definition**

5.4 (Ähnliche Matrizen). Zwei Matrizen A und B sind ähnlich, wenn es T mit $B = T^{-1}AT$ gibt.

Ähnliche Matrizen haben die gleichen Eigenwerte, denn wenn $Bx = \tilde{y}x$, dann ist $T^{-1}ATx = \tilde{y}x$. Äquivalent ist $A(Tx) = \tilde{y}(Tx)$, was zeigt, dass Tx ein Eigenvektor mit dem Eigenwert $\tilde{y}$ ist.

5.3.1 Symmetrische und positiv definite Matrizen

Ansichts unserer besonderen Berücksichtigung normaler Matrizen AA^T überrascht es nicht, dass symmetrische und/oder positiv definite Matrizen eine besondere Eigenvektorstruktur aufweisen. Wenn wir eine dieser Eigenschaften überprüfen können, können spezielle Algorithmen verwendet werden, um ihre Eigenvektoren schneller zu extrahieren.

Erstens können wir eine Eigenschaft symmetrischer Matrizen beweisen, die die Notwendigkeit einer Komplexität überflüssig macht Arithmetik. Wir beginnen mit einer Verallgemeinerung symmetrischer Matrizen auf Matrizen in $\mathbb{C}^{n \times n}$:

Definition 5.5 (Komplexes Konjugat). Die komplexe Konjugierte einer Zahl $z = a + bi \in \mathbb{C}$ ist $\bar{z} = a - bi$.

Definition 5.6 (Konjugierte Transponierung). Die konjugierte Transponierte von $A \in \mathbb{C}^{m \times n}$ ist $A^H \in \mathbb{C}^{n \times m}$.

Definition 5.7 (Hermitesche Matrix). Eine Matrix $A \in \mathbb{C}^{n \times n}$ ist hermitesch, wenn $A = A^H$.

Beachten Sie, dass eine symmetrische Matrix $A \in \mathbb{R}^{n \times n}$ automatisch hermitesch ist, da sie keinen komplexen Teil hat. Mit dieser leichten Verallgemeinerung können wir eine Symmetrieeigenschaft für Eigenwerte beweisen.

Unser Beweis wird das Skalarprodukt der Vektoren in $\mathbb{C}^n$ verwenden, gegeben durch

$$\langle x, y \rangle = \sum_{i=1}^n x_i \overline{y_i},$$

wobei $x, y \in \mathbb{C}^n$. Beachten Sie, dass diese Definition wiederum mit $\langle x, y \rangle = \overline{\langle y, x \rangle}$ übereinstimmt, wenn $x, y \in \mathbb{R}^n$. Die Eigenschaften dieses inneren Produkts stimmen größtenteils mit denen des Skalarprodukts auf $\mathbb{R}^n$ überein, mit der Ausnahme, eine bemerkenswerte dass $\langle v, w \rangle = \overline{\langle w, v \rangle}$.

Lemma 5.3. Alle Eigenwerte hermitescher Matrizen sind reell.

Nachweisen. Angenommen, $A \in \mathbb{C}^{n \times n}$ ist hermitesch mit $Ax = \lambda x$. Durch Skalierung können wir $\|x\| = 1$ annehmen. Dann gilt:

$$\begin{aligned} \langle Ax, x \rangle &= \langle \lambda x, x \rangle = \lambda \langle x, x \rangle = \lambda, \text{ da } x \text{ die Norm 1 hat} \\ &= \langle x, Ax \rangle \text{ durch Linearität von } \langle \cdot, \cdot \rangle \\ &= \langle x, \lambda x \rangle = \overline{\lambda} \langle x, x \rangle = \overline{\lambda}, \text{ da } Ax = \lambda x \\ &= \overline{\langle Ax, x \rangle} \text{ nach Definition von } \langle \cdot, \cdot \rangle \\ &= \overline{\langle x, Ax \rangle} \text{ durch Entwicklung des Produkts und Verwendung von } \overline{ab} = \overline{a} \overline{b} = \\ &= \overline{\langle x, \lambda x \rangle} = \overline{\lambda \langle x, x \rangle} = \overline{\lambda} \overline{\langle x, x \rangle} = \overline{\lambda}, \text{ da } \langle x, x \rangle = 1 \\ &= \overline{\lambda}, \text{ da } A = A^H \text{ und } \langle x, Ax \rangle = \langle Ax, x \rangle \\ &= \overline{\lambda}, \text{ da } Ax = \lambda x, \text{ da } A = A^H \\ &= \overline{\lambda}, \text{ da } Ax = \lambda x, \text{ da } x \text{ die Norm 1 hat} \end{aligned}$$

Somit ist $\lambda = \overline{\lambda}$, was je nach Bedarf nur dann passieren kann, wenn $\lambda \in \mathbb{R}$. □

Symmetrische und hermitesche Matrizen verfügen außerdem über eine besondere Orthogonalitätseigenschaft für ihre Eigenvektoren:

Lemma 5.4. Eigenvektoren, die unterschiedlichen Eigenwerten hermitescher Matrizen entsprechen, müssen orthogonal sein.

Nachweisen. Angenommen, $A \in \mathbb{C}^{n \times n}$ sei hermitesch und sei $\lambda \neq \mu$ mit $Ax = \lambda x$ und $Ay = \mu y$. Durch das vorherige Lemma wissen wir $\lambda, \mu \in \mathbb{R}$. Dann ist $\langle Ax, y \rangle = \lambda \langle x, y \rangle$. Da A aber hermitesch ist, können wir auch $\langle Ax, y \rangle = \langle x, Ay \rangle = \mu \langle x, y \rangle$ schreiben. Somit ist $\lambda \langle x, y \rangle = \mu \langle x, y \rangle$. Da $\lambda \neq \mu$, müssen wir $\langle x, y \rangle = 0$ haben. □

Schließlich können wir ohne Beweis ein krönendes Ergebnis der linearen Algebra angeben, den Spektralsatz. Dieser Satz besagt, dass keine symmetrische oder hermitesche Matrix fehlerhaft sein kann, was bedeutet, dass eine $n \times n$ -Matrix, die diese Eigenschaft erfüllt, genau n orthogonale Eigenvektoren hat.

Satz 5.1 (Spektralsatz). Angenommen, $A \in \mathbb{C}^{n \times n}$ ist hermitesch (wenn $A \in \mathbb{R}^{n \times n}$, nehmen wir an, dass es symmetrisch ist). Dann hat A genau n orthonormale Eigenvektoren $x_1, \dots, x_n$ mit (ggf. wiederholten) Eigenwerten $\lambda_1, \dots, \lambda_n$. Mit anderen Worten: Es gibt eine Orthonormalmatrix X von Eigenvektoren und eine Diagonalmatrix D von Eigenwerten, sodass $D = X^H A X$ ist.

Dieser Satz impliziert, dass jeder Vektor $y \in \mathbb{R}^n$ in eine lineare Kombination der Eigenvektoren einer hermiteschen Matrix A zerlegt werden kann. Viele Berechnungen sind auf dieser Grundlage einfacher, wie unten gezeigt:

Beispiel 5.1 (Berechnung mit Eigenvektoren). Nimm $x_1, \dots, x_n \in \mathbb{R}^n$ sind die Einheitslängen-Eigenvektoren der symmetrischen Matrix $A \in \mathbb{R}^{n \times n}$. Angenommen, wir möchten $Ay = b$ lösen. Wir können schreiben

$$b = c_1 x_1 + \dots + c_n x_n,$$

wobei $c_i = b \cdot x_i$ durch Orthonormalität. Es ist leicht, die folgende Lösung zu erraten:

$$\frac{c_1}{x_1} x_1 + \dots + \frac{c_n}{x_n} x_n = y$$

Insbesondere finden wir:

$$\begin{aligned} Ay &= A \left(\frac{c_1}{x_1} x_1 + \dots + \frac{c_n}{x_n} x_n \right) \\ &= \frac{c_1}{x_1} Ax_1 + \dots + \frac{c_n}{x_n} Ax_n \\ &= c_1 x_1 + \dots + c_n x_n \\ &= b, \text{ wie gewünscht.} \end{aligned}$$

Die obige Berechnung ergibt sowohl ein positives als auch ein negatives Ergebnis. Es zeigt, dass Operationen wie die Inversion angesichts der Eigenvektoren von symmetrischem A unkompliziert sind. Auf der anderen Seite bedeutet dies, dass das Finden des vollständigen Satzes von Eigenvektoren einer symmetrischen Matrix A „mindestens“ so schwierig ist wie das Lösen von $Ax = b$.

Nach unserem Streifzug durch die komplexen Zahlen kehren wir zu den reellen Zahlen zurück, um eine letzte nützliche, wenn auch einfache Tatsache über positiv definite Matrizen zu beweisen:

Lemma 5.5. Alle Eigenwerte positiver definiter Matrizen sind nichtnegativ.

Nachweisen. Nehmen Sie $A \in \mathbb{R}^{n \times n}$ positiv definit und nehmen Sie $Ax = \lambda x$ mit $\|x\| = 1$ an. Durch positive Definitheit wissen wir, dass $\lambda > 0$ ist. Aber je nach Bedarf gilt $x \cdot Ax = x \cdot (\lambda x) = \lambda x \cdot x = \lambda$. □

5.3.2 Spezialisierte Eigenschaften 1

Charakteristisches Polynom

Denken Sie daran, dass die Determinante einer Matrix $\det A$ genau dann die Beziehung erfüllt, dass $\det A = 0$ ist, wenn A invertierbar ist. Daher besteht eine Möglichkeit, Eigenwerte einer Matrix zu finden, darin, Wurzeln des charakteristischen Polynoms zu finden

$$p_A(\lambda) = \det(A - \lambda I_n).$$

Wir werden in unserer Diskussion hier keine Determinanten definieren, aber die Vereinfachung von p_A zeigt, dass es sich um ein Polynom n -ten Grades in λ handelt. Dies liefert einen alternativen Grund, warum es höchstens n verschiedene Eigenwerte gibt, da es höchstens n Wurzeln dieser Funktion gibt.

Aus dieser Konstruktion können wir die algebraische Multiplizität eines Eigenwerts als seine Multiplizität als Wurzel von p_A definieren. Es ist leicht zu erkennen, dass die algebraische Multiplizität mindestens so groß ist wie die geometrische

¹ Dieser Abschnitt kann übersprungen werden, wenn den Lesern nicht genügend Hintergrundwissen vorliegt, er wird jedoch der Vollständigkeit halber eingefügt.

Vielzahl. Wenn die algebraische Multiplizität 1 ist, heißt die Wurzel einfach, weil sie einem einzelnen Eigenvektor entspricht, der von allen anderen linear abhängig ist. Eigenwerte, bei denen die algebraische und geometrische Multiplizität nicht gleich sind, werden als defekt bezeichnet.

In der numerischen Analyse vermeiden wir die Diskussion der Determinante einer Matrix. Obwohl es sich um eine praktische theoretische Konstruktion handelt, ist ihr praktischer Nutzen begrenzt. Determinanten sind schwer zu berechnen. Tatsächlich versuchen Eigenwertalgorithmen nicht, Wurzeln von p_A zu finden, da dies die Auswertung einer Determinante erfordern würde. Darüber hinaus hat die Determinante $\det A$ nichts mit der Konditionierung von A zu tun, sodass eine Determinante von $\det(A - \lambda I_{n \times n})$ nahe Null möglicherweise nicht zeigt, dass λ nahezu ein Eigenwert von A ist.

Jordan-Normalform

Wir können eine Matrix nur dann diagonalisieren, wenn sie einen vollständigen Eigenraum hat. Alle Matrizen ähneln jedoch einer Matrix in Jordan-Normalform, die folgende Form hat:

Werte ungleich Null befinden sich auf den Diagonaleinträgen a_{ii} und auf der „Superdiagonalen“ $a_{i(i+1)}$.

- Diagonalwerte sind Eigenwerte, die so oft wie ihre Multiplizität wiederholt werden; die Matrix ist Blockdiagonale um diese Cluster.
- Außerdiagonale Werte sind 1 oder 0.

Somit sieht die Form in etwa wie folgt aus

$$\begin{pmatrix} \lambda_1 & 1 & & \\ & \lambda_1 & 1 & \\ & & \ddots & \\ & & & \lambda_n \end{pmatrix}$$

Die Jordan-Normalform ist theoretisch attraktiv, weil sie immer existiert, aber die 1/0-Struktur ist diskret und unter numerischen Störungen instabil.

5.4 Berechnung von Eigenwerten

Die Berechnung und Schätzung der Eigenwerte einer Matrix ist ein gut untersuchtes Problem mit vielen möglichen Lösungen. Jede Lösung ist auf eine andere Situation abgestimmt und das Erreichen einer maximalen Konditionierung oder Geschwindigkeit erfordert das Experimentieren mit mehreren Techniken. Hier behandeln wir einige der beliebtesten und einfachsten Lösungen für das Eigenwertproblem, die in der Praxis häufig vorkommen.

5.4.1 Power-Iteration

Nehmen wir zunächst an, dass $A \in \mathbb{R}^{n \times n}$ symmetrisch ist. Dann können wir nach dem Spektralsatz die Eigenvektoren $x_1, \dots, x_n \in \mathbb{R}^n$; Wir sortieren sie so, dass ihre entsprechenden Eigenwerte $|\lambda_1| \geq |\lambda_2| \geq \dots \geq |\lambda_n|$ erfüllen.

Angenommen, wir nehmen einen beliebigen Vektor v . Da die Eigenvektoren von $A \in \mathbb{R}^n$ aufspannen, wir können schreiben:

$$v = c_1 x_1 + \dots + c_n x_n.$$

Dann,

$$\begin{aligned} Av &= c_1 A x_1 + \dots + c_n A x_n \\ &= c_1 \tilde{\gamma}_1 x_1 + \dots + c_n \tilde{\gamma}_n x_n, \text{ da } A x_i = \tilde{\gamma}_i x_i \\ &= \tilde{\gamma}_1 c_1 x_1 + \tilde{\gamma}_2 c_2 x_2 + \dots + \tilde{\gamma}_n c_n x_n \\ A^2 v &= \tilde{\gamma}_1^2 c_1 x_1 + \tilde{\gamma}_2^2 c_2 x_2 + \dots + \tilde{\gamma}_n^2 c_n x_n \\ &\vdots \\ A^k v &= \tilde{\gamma}_1^k c_1 x_1 + \tilde{\gamma}_2^k c_2 x_2 + \dots + \tilde{\gamma}_n^k c_n x_n \end{aligned}$$

Beachten Sie, dass für $k \rightarrow \infty$ das Verhältnis $(\tilde{\gamma}_i/\tilde{\gamma}_1)^k \rightarrow 0$ ist, es sei denn, $\tilde{\gamma}_i = \tilde{\gamma}_1$, da $\tilde{\gamma}_1$ per Definition den größten Betrag aller Eigenwerte hat. Wenn also x die Projektion von v auf den Raum der Eigenvektoren mit den Eigenwerten $\tilde{\gamma}_1$ ist, dann gilt für $k \rightarrow \infty$ die folgende Näherung immer genauer:

$$\tilde{\gamma}_1^k x = \frac{1}{\tilde{\gamma}_1^k} A^k v$$

Diese Beobachtung führt zu einem äußerst einfachen Algorithmus zur Berechnung eines Eigenvektors x von A entspricht dem größten Eigenwert $\tilde{\gamma}_1$:

1. Nehmen Sie $v_1 \in \mathbb{R}^n$ als einen beliebigen Vektor ungleich Null.
2. Iterieren Sie bis zur Konvergenz zur Erhöhung von k :

$$v_k = A v_{k-1}$$

Dieser als Power-Iteration bekannte Algorithmus erzeugt Vektoren v_k , die immer parallel zum gewünschten x_1 sind. Es ist garantiert, dass sie konvergiert, selbst wenn A asymmetrisch ist, obwohl der Beweis dieser Tatsache aufwändiger ist als die obige Herleitung. Das einzige Mal, dass diese Technik fehlschlagen kann, ist, wenn wir versehentlich v_1 so wählen, dass $c_1 = 0$, aber die Wahrscheinlichkeit, dass dies geschieht, ist gering bis gar nicht vorhanden.

Natürlich, wenn $|\tilde{\gamma}_1| > 1$, dann $\|v_k\| \rightarrow \infty$ als $k \rightarrow \infty$, eine unerwünschte Eigenschaft für Gleitkommaarithmetik. Denken Sie daran, dass uns nur die Richtung des Eigenvektors und nicht seine Größe wichtig ist, sodass die Skalierung keinen Einfluss auf die Qualität unserer Lösung hat. Um diese Divergenzsituation zu vermeiden, können wir einfach bei jedem Schritt normalisieren und so den normalisierten Potenziterationsalgorithmus erstellen:

1. Nehmen Sie $v_1 \in \mathbb{R}^n$ als einen beliebigen Vektor ungleich Null.
2. Iterieren Sie bis zur Konvergenz zur Erhöhung von k :

$$\begin{aligned} w_k &= A v_{k-1} \\ v_k &= \frac{w_k}{\|w_k\|} \end{aligned}$$

Beachten Sie, dass wir die Norm $\|\cdot\|$ nicht mit einem bestimmten Index versehen haben. Mathematisch gesehen reicht jede Norm aus, um das Divergenzproblem zu verhindern, da wir gezeigt haben, dass alle Normen auf $\mathbb{R}^n$ äquivalent sind. In der Praxis verwenden wir oft die Unendlichkeitsnorm $\|\cdot\|_\infty$; In diesem Fall lässt sich leicht überprüfen, dass $\|w_k - \tilde{y}\|_\infty \leq \|\tilde{y}\|_\infty$.

5.4.2 Inverse Iteration

Wir haben jetzt eine Strategie, um den Eigenwert $\tilde{\lambda}_1$ mit dem größten Betrag zu finden. Angenommen, A ist invertierbar, sodass wir $y = A^{-1}v$ auswerten können, indem wir $Ay = v$ mithilfe der in den vorherigen Kapiteln behandelten Techniken lösen.

Wenn $Ax = \tilde{\lambda}x$, dann ist $x = \tilde{\lambda}^{-1}Ax$ oder gleichwertig

$$A^{-1}x = \frac{1}{\tilde{\lambda}}x.$$

Damit haben wir gezeigt, dass $1/\tilde{\lambda}$ ein Eigenwert von A^{-1} und dann $|\tilde{\lambda}|^{-1}$ mit Eigenvektor x . Beachten Sie, dass wenn $|a| > |b|$ ist $\tilde{\lambda}^{-1}$ für jedes $a, b \in \mathbb{C}$, also ist der kleinste Eigenwert von A der größte. Diese Beobachtung ergibt eine Strategie zum Finden Betrageigenvektor von A , sogenannte $\tilde{\lambda}_1$ von $\tilde{\lambda}_n$ anstelle von $\tilde{\lambda}_1$ inverse Potenziteration:

1. Nehmen Sie $v_1 \in \mathbb{R}^n$ als einen beliebigen Vektor ungleich Null.

2. Iterieren Sie bis zur Konvergenz zur Erhöhung von k :

(a) Lösen Sie nach w_k auf: $Aw_k = v_k \tilde{\lambda}_1$

(b) Normalisieren: $v_k = \frac{w_k}{\|w_k\|}$

Wir lösen wiederholt Gleichungssysteme mit derselben Matrix A , was eine perfekte Anwendung der Faktorisierungstechniken aus früheren Kapiteln ist. Wenn wir beispielsweise $A = LU$ schreiben, könnten wir eine äquivalente, aber wesentlich effizientere Version der Umkehrpotenziteration formulieren:

1. Faktor $A = LU$

2. Nehmen Sie $v_1 \in \mathbb{R}^n$ als einen beliebigen Vektor ungleich Null.

3. Iterieren Sie bis zur Konvergenz zur Erhöhung von k :

(a) Lösen Sie nach y_k durch Vorwärtssubstitution auf: $Ly_k = v_k \tilde{\lambda}_1$ (b)

Lösen Sie nach w_k durch Rückwärtssubstitution auf: $Uw_k = y_k$

(c) Normalisieren Sie: $v_k = \frac{w_k}{\|w_k\|}$

5.4.3 Schalten

Angenommen, $\tilde{\lambda}_2$ ist der Eigenwert mit der zweitgrößten Größe von A . Angesichts unserer ursprünglichen Ableitung der Potenziteration ist leicht zu erkennen, dass die Potenziteration am schnellsten konvergiert, wenn $|\tilde{\lambda}_2/\tilde{\lambda}_1|$ ist klein, da in diesem Fall die Leistung $(\tilde{\lambda}_2/\tilde{\lambda}_1)^k$ schnell abfällt. Wenn dieses Verhältnis hingegen nahezu 1 beträgt, kann es viele Iterationen der Potenziteration erfordern, bis ein einzelner Eigenvektor isoliert wird.

Wenn die Eigenwerte von A λ_1 sind, $\dots$, λ_n , dann ist es leicht zu erkennen, dass die Eigenwerte von $A - \lambda_1 I_{n \times n}$ $\lambda_2 - \lambda_1, \dots, \lambda_n - \lambda_1$ sind. Dann besteht eine Strategie, um die Potenziteration schnell konvergieren zu lassen, darin, λ_1 so zu wählen, dass:

$$\frac{\lambda_2 - \lambda_1}{\lambda_1 - \lambda_1} < \frac{\lambda_2 - \lambda_1}{\lambda_1 - \lambda_1}.$$

Natürlich kann es eine Kunst sein, ein solches λ_1 zu erraten, da die Eigenwerte von A offensichtlich zunächst nicht bekannt sind. Wenn wir ebenfalls annehmen, dass λ_1 nahe einem Eigenwert von A liegt, dann hat $A - \lambda_1 I_{n \times n}$ einen Eigenwert nahe 0, den wir durch inverse Iteration ermitteln können.

Eine Strategie, die sich diese Beobachtung zunutze macht, ist als Rayleigh-Quotienten-Iteration bekannt. Wenn wir eine feste Schätzung für einen Eigenvektor x von A haben, dann ist durch NUMBER die Kleinste-Quadrate-Approximation des entsprechenden Eigenwerts λ gegeben

$$\lambda = \frac{x^T A x}{x^T x}.$$

Dieser Bruch wird als Rayleigh-Quotient bezeichnet. Daher können wir versuchen, die Konvergenz zu erhöhen, indem wir wie folgt iterieren:

1. Nehmen Sie an, dass $v_1 \in \mathbb{R}^n$ ein beliebiger Vektor ungleich Null oder eine anfängliche Schätzung eines Eigenvektors ist.
2. Iterieren Sie bis zur Konvergenz zur Erhöhung von k :

(a) Schreiben Sie die aktuelle Schätzung des Eigenwerts

$$\lambda_k = \frac{v_k^T A v_k}{v_k^T v_k}.$$

(b) Lösen Sie nach w auf: $(A - \lambda_k I_{n \times n})w = v_k$ (c) Normalisieren $w = \frac{v_k}{\|v_k\|}$

Sie: $v_k = \frac{w}{\|w\|}$

Bei einer guten Ausgangsschätzung konvergiert diese Strategie viel schneller, aber die Matrix $A - \lambda_k I_{n \times n}$ ist bei jeder Iteration unterschiedlich und kann nicht mit LU oder den meisten anderen Strategien vorkonfiguriert werden. Somit sind weniger Iterationen notwendig, aber jede Iteration nimmt mehr Zeit in Anspruch!

5.4.4 Finden mehrerer Eigenwerte

Bisher haben wir Techniken zum Finden eines einzelnen Eigenwert/Eigenvektor-Paares beschrieben: Potenziteration, um den größten Eigenwert zu finden, inverse Iteration, um den kleinsten zu finden, und Verschiebung zu Zielwerten dazwischen. Für viele Anwendungen reicht ein einzelner Eigenwert natürlich nicht aus. Zum Glück können wir unsere Strategien erweitern, um auch diesen Fall zu bearbeiten.

Deflation

Erinnern Sie sich an unsere Potenziterationsstrategie: Wählen Sie ein beliebiges v_1 und multiplizieren Sie es iterativ mit A , bis nur noch der größte Eigenwert λ_1 übrig bleibt. Nehmen Sie x_1 als den entsprechenden Eigenvektor.

Wir haben jedoch schnell einen unwahrscheinlichen Fehlermodus dieses Algorithmus verworfen, wenn $v_1 \cdot x_1 = 0$. In diesem Fall werden Sie niemals einen Vektor wiederherstellen, egal wie oft Sie mit A vormultiplizieren

parallel zu x_1 , da man eine Nullkomponente nicht verstärken kann. Die Wahrscheinlichkeit, eine solche v_1 zu wählen, ist genau null, sodass die Power-Iteration in allen Fällen, außer in den schädlichsten Fällen, sicher bleibt.

Wir können diesen Nachteil auf den Kopf stellen und eine Strategie formulieren, um mehr als einen Eigenwert zu finden, wenn A symmetrisch ist. Angenommen, wir finden x_1 und $\tilde{y}_1$ wie zuvor durch Potenziteration. Jetzt starten wir die Power-Iteration neu, jedoch bevor wir mit Projekt x_1 aus v_1 beginnen. Da die Eigenvektoren von A dann orthogonal sind, wird durch die Potenziteration der zweitgrößte Eigenwert wiederhergestellt!

Aufgrund numerischer Probleme kann es vorkommen, dass die Anwendung von A auf einen Vektor eine kleine Komponente parallel zu x_1 einführt. In der Praxis können wir diesen Effekt vermeiden, indem wir in jeder Iteration projizieren. Am Ende ergibt diese Strategie den folgenden Algorithmus zur Berechnung der Eigenwerte in der Reihenfolge absteigender Größe:

- Für jeden gewünschten Eigenwert = 1, 2, . . .

1. Nehmen Sie $v_1 \in \mathbb{R}^n$ als einen beliebigen Vektor ungleich Null.

2. Iterieren Sie bis zur Konvergenz zur Erhöhung von

k: (a) Projizieren Sie die Eigenvektoren, die wir bereits berechnet haben:

$$u_k = v_k - \tilde{y}_k \text{projspan}\{x_1, \dots, x_{k-1}\} = v_k - \tilde{y}_k$$

(b) Multiplizieren Sie $A u_k = w$

(c) Normalisieren Sie: $v_k = \frac{w}{\|w\|}$

3. Fügen Sie das Ergebnis der Iteration zur Menge von x_i hinzu

Die innere Schleife entspricht einer Potenziteration auf der Matrix $A P$, wobei P $x_1, \dots, x_{k-1}$ herausprojiziert.

Es ist leicht zu erkennen, dass $A P$ die gleichen Eigenvektoren wie A hat; seine Eigenwerte sind $\tilde{y}_1, \dots, \tilde{y}_n$, wobei die restlichen Eigenwerte auf Null gesetzt werden.

Allgemeiner gesagt besteht die Deflationsstrategie darin, die Matrix A so zu modifizieren, dass die Potenziteration einen Eigenvektor offenbart, den Sie noch nicht berechnet haben. Beispielsweise ist $A P$ eine Modifikation von A , sodass die großen Eigenwerte, die wir bereits berechnet haben, auf Null gesetzt werden.

Unsere Projektionsstrategie schlägt fehl, wenn A asymmetrisch ist, da seine Eigenvektoren in diesem Fall möglicherweise nicht orthogonal sind. Andere, weniger offensichtliche Deflationsstrategien können in diesem Fall funktionieren. Angenommen, $A x_1 = \tilde{y}_1 x_1$ mit $x_1 = 1$. Nehmen Sie H als die Householder-Matrix an, sodass $H x_1 = e_1$, der erste Standardbasisvektor. Ähnlichkeitstransformationen wirken sich wiederum nicht auf die Menge der Eigenvektoren aus, daher könnten wir versuchen, mit H zu konjugieren. Überlegen Sie, was passiert, wenn wir $H A H$ mit e_1 multiplizieren:

$$H A H e_1 = H A e_1, \text{ da } H \text{ symmetrisch ist}$$

$$= H A x_1, \text{ da } H x_1 = e_1 \text{ und } H = \tilde{y}_1 H x_1, \quad \tilde{y}_1^2 = I_{n \times n}$$

$$\text{da } A x_1 = \tilde{y}_1 x_1$$

$$= \tilde{y}_1 e_1 \text{ nach Definition von } H$$

Somit ist die erste Spalte von $H A H$ $\tilde{y}_1 e_1$, was zeigt, dass $H A H$ die folgende Struktur hat (CITE HEIDE):

$$H A H = \begin{pmatrix} \tilde{y}_1 & 0 \\ 0 & B \end{pmatrix}.$$

Die Matrix $B \in \mathbb{R}^{(n-1) \times (n-1)}$ hat Eigenwerte $\tilde{y}_2, \dots, \tilde{y}_n$. Daher besteht eine andere Strategie zur Deflation darin, immer kleinere B -Matrizen zu konstruieren, wobei jeder Eigenwert mithilfe einer Potenziteration berechnet wird.

QR-Iteration

Die Deflation hat den Nachteil, dass wir jeden Eigenvektor separat berechnen müssen, was langsam sein kann und zu Fehlern führen kann, wenn einzelne Eigenwerte nicht genau sind. Unsere verbleibenden Strategien versuchen, mehr als einen Eigenvektor gleichzeitig zu finden.

Denken Sie daran, dass ähnliche Matrizen A und $B = T^{-1}AT$ dieselben Eigenwerte haben müssen. Somit kann ein Algorithmus, der versucht, die Eigenwerte von A zu finden, Ähnlichkeitstransformationen frei auf A anwenden. Natürlich kann die Anwendung von T im Allgemeinen schwierig sein, da sie effektiv eine Invertierung von T erfordern würde. Daher suchen wir nach T -Matrizen, deren Umkehrungen leicht zu ermitteln sind anwenden.

Einer unserer Beweggründe für die Ableitung der QR-Faktorisierung war, dass die Matrix Q orthogonal ist und $Q^{-1} = Q^T$ erfüllt. Daher sind Q und Q^T gleichermaßen einfach anzuwenden, was orthogonale Matrizen zu einer guten Wahl für Ähnlichkeitstransformationen macht.

Aber welche orthogonale Matrix Q sollten wir wählen? Idealerweise sollte Q die Struktur von A beinhalten und gleichzeitig einfach zu berechnen sein. Es ist unklar, wie man einfache Transformationen wie Householder-Matrizen strategisch anwendet, um mehrere Eigenwerte aufzudecken,² aber wir wissen, wie man ein solches Q einfach durch Faktorisieren von $A = QR$ erzeugt. Dann könnten wir A mit Q konjugieren und finden: $Q^{-1}AQ = Q^T A Q = Q^T (QR)Q = (Q^T Q)RQ = RQ$

Erstaunlicherweise ist die Konjugation von $A = QR$ durch die orthogonale Matrix Q identisch mit dem Schreiben des Produkts RQ !

Auf der Grundlage dieser Überlegungen stellten in den 1950er Jahren mehrere Gruppen europäischer Mathematiker Hypothesen auf bemessen den gleichen eleganten iterativen Algorithmus zum Finden der Eigenwerte einer Matrix A :

1. Nehmen Sie $A_1 = A$.

2. Für $k = 1, 2, \dots$

(a) Faktor $A_k = Q_k R_k$. (b)

Schreiben Sie $A_{k+1} = R_k Q_k$.

Nach unserer obigen Ableitung haben die Matrizen A_k alle die gleichen Eigenwerte wie A . Nehmen wir außerdem an, dass die A_k -Werte gegen ein A_∞ konvergieren. Dann können wir $A_\infty = Q_\infty R_\infty$ faktorisieren und durch Konvergenz wissen wir $A_\infty = Q_\infty R_\infty = R_\infty Q_\infty$. Nach ZAHL sind die Eigenwerte von R_∞ einfach die Werte entlang der Diagonale von R_∞ , und nach ZAHL muss das Produkt $R_\infty Q_\infty = A_\infty$ wiederum die gleichen Eigenwerte haben. Schließlich hat A_∞ konstruktionsbedingt die gleichen Eigenwerte wie A . Wir haben also gezeigt, dass die QR-Iteration bei Konvergenz die Eigenwerte von A auf einfache Weise offenbart.

Natürlich geht die obige Ableitung davon aus, dass es A_∞ mit A_k als $k \rightarrow \infty$ gibt. Tatsächlich ist die QR-Iteration eine stabile Methode, die in vielen wichtigen Situationen garantiert konvergiert, und die Konvergenz kann sogar durch eine Änderung der Strategie verbessert werden. Wir werden hier keine genauen Bedingungen herleiten, sondern können stattdessen eine Vorstellung davon geben, warum diese scheinbar willkürliche Strategie konvergieren sollte. Im Folgenden geben wir einige Hinweise für den symmetrischen Fall $A = A^T$, der dank der Orthogonalität der Eigenvektoren in diesem Fall einfacher zu analysieren ist.

Angenommen, die Spalten von A sind durch $a_1, \dots, a_n$ gegeben, und betrachten Sie die Matrix A^k für großes k . Wir kann schreiben:

$$A^k = A \cdot A^{k-1} = \begin{pmatrix} \tilde{y} & | & \tilde{y} & | & \tilde{y} \\ A^{k-1}a_1 & A^{k-1}a_2 & \dots & A^{k-1}a_n \end{pmatrix} \begin{pmatrix} \tilde{y} \\ | \\ \tilde{y} \end{pmatrix}$$

²Fortgeschrittenere Techniken bewirken jedoch genau dies!

Nach unserer Ableitung der Potenziteration ist die erste Spalte von A^k im Allgemeinen parallel zum Eigenvektor x_1 mit dem größten Betrag $|y_1|$ da wir einen Vektor a_1 genommen und ihn viele Male mit A multipliziert haben.

Wenn wir unsere Intuition aus der Deflation anwenden, nehmen wir an, dass wir a_1 aus der zweiten Spalte von A^k . Das projizieren. Ein Vektor muss nahezu parallel zu x_2 sein, da es sich um den zweitdominantesten Eigenwert handelt! Weiter Induktiv würde die Faktorisierung von $A = QR$ einen Satz von Eigenvektoren als Spalten von Q in der Reihenfolge abnehmender Eigenwertgröße mit den entsprechenden Eigenwerten entlang der Diagonale von R ergeben.

Natürlich nimmt die Berechnung von A^k für große k die Bedingungszahl von A zur k -ten Potenz, sodass QR auf der resultierenden Matrix wahrscheinlich fehlschlägt; Dies ist deutlich zu erkennen, da alle Spalten von A^k für große k wie x_1 aussehen sollten. Wir können jedoch folgende Beobachtung machen:

$$\begin{aligned} A &= Q_1 R_1 \\ A^2 &= (Q_1 R_1)(Q_1 R_1) \\ &= Q_1 (R_1 Q_1) R_1 \\ &= Q_1 Q_2 R_2 R_1 \text{ unter Verwendung der obigen Notation der QR-Iteration, da } A^2 = R_1 Q_1 \\ &\vdots \\ A^k &= Q_1 Q_2 \dots Q_k R_k R_{k-1} \dots R_1 \end{aligned}$$

Die separate Gruppierung der Q_i -Variablen und der R_i -Variablen liefert eine QR-Faktorisierung von A . Wir erwarten, dass die Spalten von $Q_1 \dots Q_k$ zu den Eigenvektoren von A konvergieren. ^k. Daher,

Durch ein ähnliches Argument können wir finden

$$\begin{aligned} A &= Q_1 R_1 \\ A^2 &= R_1 Q_1 \text{ durch unsere Konstruktion der QR-Iteration} \\ &= Q_2 R_2 \text{ nach Definition der Faktorisierung} \\ A^3 &= R_2 Q_2 \text{ aus QR-Iteration} \\ &= Q_2 A^2 Q_2, \text{ da } A^2 = Q_2 R_2 \\ &= Q_2 R_1 Q_1 Q_2, \text{ da } A^2 = R_1 Q_1 \\ &= Q_2 Q_1 A Q_1 Q_2, \text{ da } A = Q_1 R_1 \\ &\vdots \\ A_{k+1} &= Q_k \dots Q_1 A Q_1 \dots Q_k \text{ induktiv} \\ &= (Q_1 \dots Q_k) A (Q_1 \dots Q_k) \end{aligned}$$

wobei A_k die k -te Matrix aus der QR-Iteration ist. Somit ist A_{k+1} einfach die durch das Produkt $Q^{-k} = Q_1 \dots Q_k$ konjugierte Matrix A . Wir haben zuvor argumentiert, dass die Spalten von Q^{-k} gegen die Eigenvektoren von A konvergieren. Da die Konjugation durch die Matrix der Eigenvektoren eine diagonale Matrix von Eigenwerten ergibt, wissen wir also $A_{k+1} = Q^{-k} A Q^k$ wird entlang seiner Diagonale ungefähre Eigenwerte von A haben, wenn $k \rightarrow \infty$.

Krylov-Subraummethoden

Unsere Begründung der QR-Iteration umfasste die Analyse der Spalten von A der Power- k als $k \cdot \tilde{y}$ als Erweiterung Iteration. Allgemeiner ausgedrückt gilt für einen Vektor $\tilde{y} \in \mathbb{R}^n$, Wir können die sogenannte Krylov-Matrix untersuchen

$$K_k = \begin{bmatrix} \tilde{y} & A\tilde{y} & A^2\tilde{y} & \dots & A^{k-1}\tilde{y} \end{bmatrix}$$

Methoden zur Analyse von K_k zur Ermittlung von Eigenvektoren und Eigenwerten werden im Allgemeinen als Krylov-Unterraummethoden bezeichnet. Beispielsweise verwendet der Arnoldi-Iterationsalgorithmus die Gram-Schmidt-Orthogonalisierung, um eine orthogonale Basis $\{q_1, \dots, q_k\}$ für den Spaltenraum von K_k :

1. Beginnen Sie damit, dass q_1 ein beliebiger Einheitsnormvektor ist

2. Für $k = 2, 3, \dots$

(a) $\tilde{b}_k = A q_{k-1}$ (b)

Projizieren Sie die $\tilde{b}_k$, die Sie bereits berechnet haben:

$$\tilde{b}_k = \tilde{b}_k - \text{proj}_{\text{span}\{q_1, \dots, q_{k-1}\}} \tilde{b}_k$$

(c) Renormieren Sie, um das nächste $q_k = \tilde{b}_k / \|\tilde{b}_k\|$ zu finden.

Die Matrix Q_k , deren Spalten die oben gefundenen Vektoren sind, ist eine orthogonale Matrix mit demselben Spaltenraum wie K_k , und Eigenwertschätzungen können aus der Struktur von $Q^T A Q_k$ wiederhergestellt werden.

Die Verwendung von Gram-Schmidt macht diese Technik instabil und das Timing wird mit zunehmendem k immer schlechter. Es sind jedoch viele Erweiterungen erforderlich, um sie durchführbar zu machen. Eine Strategie besteht beispielsweise darin, einige Arnoldi-Iterationen auszuführen, die Ausgabe zu verwenden, um eine bessere Schätzung für das anfängliche q_1 zu generieren, und dann neu zu starten. Methoden dieser Klasse eignen sich für Probleme, die mehrere Eigenvektoren an einem Ende des Spektrums erfordern, ohne den vollständigen Satz zu berechnen.

5.5 Sensibilität und Konditionierung

Wie gewarnt, haben wir nur einige Eigenwerttechniken aus einer umfangreichen und langjährigen Literatur skizziert. Mit nahezu jeder algorithmischen Technik wurde zum Auffinden von Spektren experimentiert, von iterativen Methoden über die Wurzelfindung auf dem charakteristischen Polynom bis hin zu Methoden, die Matrizen zur parallelen Verarbeitung in Blöcke unterteilen.

Genau wie bei linearen Lösern können wir die Konditionierung eines Eigenwertproblems unabhängig von der Lösungstechnik bewerten. Diese Analyse kann helfen zu verstehen, ob ein vereinfachtes iteratives Schema zum Finden der Eigenvektoren einer gegebenen Matrix erfolgreich ist oder ob komplexere Methoden erforderlich sind. Es ist wichtig zu beachten, dass die Konditionierung eines Eigenwertproblems nicht mit der Bedingungszahl der Matrix zur Lösung von Systemen identisch ist, da es sich um separate Probleme handelt.

Angenommen, eine Matrix A hat einen Eigenvektor x mit dem Eigenwert $\tilde{\lambda}$. Die Analyse der Konditionierung des Eigenwertproblems beinhaltet die Analyse der Stabilität von x und $\tilde{\lambda}$ gegenüber Störungen in A . Zu diesem Zweck könnten wir A durch eine kleine Matrix $\tilde{A}$ stören und so den Satz der Eigenvektoren ändern. Insbesondere können wir Eigenvektoren von $A + \tilde{A}$ als Störungen von Eigenvektoren von A schreiben, indem wir das Problem lösen

$$(A + \tilde{A})(x + \tilde{y}) = (\tilde{\lambda} + \tilde{\mu})(x + \tilde{y})$$

Die Erweiterung beider Seiten ergibt:

$$Ax + A\tilde{y}x + \tilde{y}A \cdot x + \tilde{y}A \cdot \tilde{y}x = \tilde{y}x + \tilde{y}\tilde{y}x + \tilde{y}\tilde{y} \cdot x + \tilde{y}\tilde{y} \cdot \tilde{y}x$$

Unter der Annahme, dass $\tilde{y}A$ klein ist, gehen wir davon aus³, dass $\tilde{y}x$ und $\tilde{y}\tilde{y}$ ebenfalls klein sind. Produkte zwischen diesen Variablen sind dann vernachlässigbar, was die folgende Näherung ergibt:

$$Ax + A\tilde{y}x + \tilde{y}A \cdot x \approx \tilde{y}x + \tilde{y}\tilde{y}x + \tilde{y}\tilde{y} \cdot x$$

Da $Ax = \tilde{y}x$, können wir diesen Wert von beiden Seiten subtrahieren, um Folgendes zu finden:

$$A\tilde{y}x + \tilde{y}A \cdot x \approx \tilde{y}\tilde{y}x + \tilde{y}\tilde{y} \cdot x$$

Wir wenden nun einen analytischen Trick an, um unsere Ableitung zu vervollständigen. Da $Ax = \tilde{y}x$, wissen wir $(A - \tilde{y}\tilde{y})x = 0$, also ist $A - \tilde{y}\tilde{y}$ nicht vollrangig. Die Transponierte einer Matrix ist nur dann vollrangig, wenn die Matrix vollrangig ist. Wir wissen also, dass $(A - \tilde{y}\tilde{y})^T = A^T - \tilde{y}\tilde{y}^T$ auch einen Nullraumvektor y hat. Also $Ay = \tilde{y}y$; wir können y den linken Eigenvektor nennen, der x entspricht. Wir können unsere obige Störungsschätzung mit y links multiplizieren:

$$y(A\tilde{y}x + \tilde{y}A \cdot x) \approx y(\tilde{y}\tilde{y}x + \tilde{y}\tilde{y} \cdot x)$$

Da $Ay = \tilde{y}y$ ist, können wir Folgendes vereinfachen:

$$y\tilde{y}A \cdot x \approx y\tilde{y}\tilde{y}y \cdot x$$

Erträge neu ordnen:

$$\tilde{y} \approx \frac{y(\tilde{y}A)x \cdot \tilde{y}\tilde{y}}{yx}$$

Angenommen, $x = 1$ und $y = 1$. Wenn wir dann Normen auf beiden Seiten nehmen, finden wir:

$$|\tilde{y}\tilde{y}| \approx \frac{|\tilde{y}A|^2}{y \cdot x}$$

Im Allgemeinen hängt die Konditionierung des Eigenwertproblems also – wie erwartet – von der Größe der Störung $\tilde{y}A$ und dem Winkel zwischen den linken und rechten Eigenvektoren x und y ab. Wir können $1/(x \cdot y)$ als ungefähre Bedingungsanzahl verwenden. Beachten Sie, dass $x = y$, wenn A symmetrisch ist, was eine Bedingungsanzahl von 1 ergibt; Dies spiegelt die Tatsache wider, dass die Eigenvektoren symmetrischer Matrizen orthogonal und daher maximal getrennt sind.

5.6 Probleme

³Diese Annahme sollte in einer strengeren Behandlung überprüft werden!

Kapitel 6

Einzelwertzerlegung

In Kapitel 5 haben wir eine Reihe von Algorithmen zur Berechnung der Eigenwerte und Eigenvektoren von Matrizen $A \in \mathbb{R}^{n \times n}$ abgeleitet. Nachdem wir diese Maschinerie entwickelt haben, schließen wir unsere anfängliche Diskussion der numerischen linearen Algebra ab, indem wir eine letzte Matrixfaktorisierung ableiten und verwenden, die für jede Matrix $A \in \mathbb{R}^{m \times n}$ existiert: die Singularwertzerlegung (SVD).

6.1 Ableitung der SVD

Für $A \in \mathbb{R}^{m \times n}$, Wir können uns die Funktion $x \mapsto Ax$ als eine Karte vorstellen, die Punkte in $\mathbb{R}^n$ zu Punkten in $\mathbb{R}^m$ führt. Aus dieser Perspektive könnten wir fragen, was dabei mit der Geometrie von $\mathbb{R}^n$ passiert und insbesondere welche Auswirkung A auf die Längen und Winkel zwischen Vektoren hat.

Wenn wir unseren üblichen Ausgangspunkt für Eigenwertprobleme anwenden, können wir nach der Auswirkung fragen, die A auf hat die Längen von Vektoren durch Untersuchung kritischer Punkte des Verhältnisses

$$R(x) = \frac{Ax^t}{x}$$

über verschiedene Werte von x . Die Skalierung von x spielt keine Rolle, da

$$R(\tilde{x}) = \frac{A \cdot \tilde{x}}{\tilde{x}} = \frac{|\tilde{y}|}{|\tilde{y}|} \cdot \frac{Ax^t}{x} = \frac{Ax^t}{x} = R(x).$$

Somit können wir unsere Suche auf x mit $|x| = 1$ beschränken. Da $R(x) \in \mathbb{R}$ ist, können wir außerdem stattdessen $[R(x)]^2 = Ax^t Ax$ betrachten. Wie wir jedoch in den vorherigen Kapiteln gezeigt haben, sind kritische Punkte von $x^t Ax$ unter der Bedingung $|x| = 1$ genau die Eigenvektoren x_i , die $Ax_i = \lambda_i x_i$ erfüllen; Beachten Sie $\lambda_i \in \mathbb{R}$ und $x_i \cdot x_j = 0$, wenn $i \neq j$, da AA^t symmetrisch und positiv semidefinit ist.

Basierend auf unserer Verwendung der Funktion R ist die $\{x_i\}$ -Basis eine sinnvolle Basis für die Untersuchung der geometrischen Effekte von A . Zurück zu diesem ursprünglichen Ziel: Definieren Sie $y_i \in \mathbb{R}^m$ durch $y_i = Ax_i$. Wir können eine zusätzliche Beobachtung zu y_i machen, die eine noch stärkere Eigenwertstruktur offenbart:

$$\begin{aligned} \tilde{y}_i y_i &= \tilde{y}_i \cdot Ax_i \text{ nach Definition von } y_i \\ &= A(\tilde{y}_i x_i) \\ &= A(Ax_i), \text{ da } x_i \text{ ein Eigenvektor von } AA^t \text{ ist} \\ &= (AA^t)(Ax_i) \text{ durch Assoziativität} \\ &= (AA^t)y_i \end{aligned}$$

Wir haben also zwei Fälle:

1. Wenn $\tilde{y}_i = 0$, dann ist $y_i = 0$. In diesem Fall ist x_i ein Eigenvektor von AA und $y_i = Ax_i$ ist $AAx_i = \frac{1}{2} x_i$ ein entsprechender Eigenvektor von AA mit $y_i = Ax_i = \frac{1}{2} x_i$.

2. Wenn $\tilde{y}_i = 0$, ist $y_i = 0$.

Ein identischer Beweis zeigt, dass, wenn y ein Eigenvektor von AA ist, $x \in A y$ entweder Null oder ein Eigenvektor von AA mit demselben Eigenwert ist.

Nehmen Sie k als die Anzahl der oben diskutierten streng positiven Eigenwerte $\tilde{y}_i > 0$ an. Mit unserer obigen Konstruktion können wir $x_1, \dots, x_k \in \mathbb{R}^n$ als Eigenvektoren von AA und entsprechende Eigenvektoren $y_1, \dots, y_k \in \mathbb{R}^m$ von AA so dass

$$A x_i = \tilde{y}_i x_i$$

$$A A y_i = \tilde{y}_i y_i$$

für Eigenwerte $\tilde{y}_i > 0$; hier normalisieren wir so, dass $x_i = y_i = 1$ für alle i . Der traditionellen Notation folgend, können wir Matrizen $V \in \mathbb{R}^{n \times k}$ und $U \in \mathbb{R}^{m \times k}$ definieren, deren Spalten x_i bzw. y_i sind.

Wir können die Auswirkung dieser neuen Basismatrizen auf A untersuchen. Nehmen Sie e_i als den i -ten Standardbasisvektor. Dann,

$$\begin{aligned} U^T A V e_i &= U^T A x_i \text{ nach Definition von } V \\ &= \frac{1}{\tilde{y}_i} U^T A (\tilde{y}_i x_i), \text{ da wir } \tilde{y}_i > 0 \text{ } \tilde{y}_i = 1 \text{ angenommen} \\ &\quad \text{haben} \\ &= \frac{1}{\tilde{y}_i} U^T A (A x_i), \text{ da } x_i \text{ ein Eigenvektor von } AA \text{ } \tilde{y}_i = 1 \text{ ist} \\ &= \frac{1}{\tilde{y}_i} U^T (AA) x_i \text{ durch Assoziativität } \tilde{y}_i = 1 \\ &= \frac{1}{\tilde{y}_i} U^T (AA) y_i, \text{ da wir neu skaliert haben, sodass } y_i = 1 \text{ } \tilde{y}_i = 1 \\ &= \tilde{y}_i U^T y_i, \text{ da } A A y_i = \tilde{y}_i y_i \\ &= \tilde{y}_i e_i \end{aligned}$$

Nehmen Sie $\tilde{Y} = \text{diag}(\tilde{y}_1, \dots, \tilde{y}_k)$. Dann zeigt die obige Ableitung, dass $U^T A V = \tilde{Y}$.

Vervollständigen Sie die Spalten von U und V zu $U \in \mathbb{R}^{m \times m}$ und $V \in \mathbb{R}^{n \times n}$, indem Sie Orthonormalvektoren x_i und y_i mit $A x_i = 0$ bzw. $A A y_i = 0$ hinzufügen. In diesem Fall ist es einfach, $U A V e_i = 0$ und/oder $U A V = 0$ anzuzeigen. Also, wenn wir nehmen

$$\tilde{y}_{ij} = \begin{cases} \tilde{y}_i & i = j \\ 0 & i \neq j \end{cases} \text{ und } i \neq k \text{ sonst}$$

Dann können wir unsere vorherige Beziehung erweitern, um $U A V = \tilde{Y}$ oder eine gleichwertige Gleichung zu zeigen

$$A = U \tilde{Y} V.$$

Diese Faktorisierung ist genau die Singularwertzerlegung (SVD) von A . Die Spalten von U überspannen den Spaltenraum von A und werden dessen linke Singulärvektoren genannt; Die Spalten von V überspannen seinen Zeilenraum und sind die rechten singulären Vektoren. Die Diagonalelemente $\tilde{\gamma}_i$ von $\tilde{\gamma}$ sind die Singulärwerte von A ; Normalerweise sind sie so sortiert, dass $\tilde{\gamma}_1 \geq \tilde{\gamma}_2 \geq \dots \geq 0$. Sowohl U als auch V sind orthogonale Matrizen.

Die SVD liefert eine vollständige geometrische Charakterisierung der Wirkung von A . Da U und V orthogonal sind, können sie als Rotationsmatrizen betrachtet werden; Als Diagonalmatrix skaliert $\tilde{\gamma}$ einfach einzelne Koordinaten. Somit sind alle Matrizen $A \in \mathbb{R}^{m \times n}$ eine Zusammensetzung aus einer Rotation, einer Skala und einer zweiten Rotation.

6.1.1 Berechnung der SVD

Denken Sie daran, dass die Spalten von V einfach die Eigenvektoren von AA^T sind und daher mit den im vorherigen Kapitel beschriebenen Techniken berechnet werden können. Da $A = U\tilde{\gamma}V^T$ ist, wissen wir $AV = U\tilde{\gamma}$. Somit sind die Spalten von U , die singulären Werten ungleich Null in $\tilde{\gamma}$ entsprechen, einfach normalisierte Spalten von AV ; Die verbleibenden Spalten erfüllen $AA^T u_i = 0$, was mithilfe der LU-Faktorisierung gelöst werden kann.

Diese Strategie ist keineswegs der effizienteste oder stabilste Ansatz zur Berechnung der SVD, funktioniert aber für viele Anwendungen einigermaßen gut. Wir werden speziellere Ansätze zur Ermittlung der SVD weglassen, beachten jedoch, dass es sich bei vielen um einfache Erweiterungen der Potenziteration und anderer Strategien handelt, die wir bereits behandelt haben und die ohne explizite Bildung von AA^T oder $A^T A$ funktionieren.

6.2 Anwendungen des SVD

Den Rest dieses Kapitels widmen wir der Vorstellung vieler Anwendungen der SVD. Die SVD taucht sowohl in der Theorie als auch in der Praxis der numerischen linearen Algebra unzählige Male auf, und ihre Bedeutung kann kaum überbewertet werden.

6.2.1 Lösung linearer Systeme und der Pseudoinversen

Im Spezialfall, in dem $A \in \mathbb{R}^{n \times n}$ quadratisch und invertierbar ist, ist es wichtig zu beachten, dass die SVD zur Lösung des linearen Problems $Ax = b$ verwendet werden kann. Insbesondere gilt $U\tilde{\gamma}V^T x = b$, oder

$$x = V\tilde{\gamma}^{-1}U^T b.$$

In diesem Fall ist $\tilde{\gamma}$ eine quadratische Diagonalmatrix, also $\tilde{\gamma}^{-1}$ ist einfach die Matrix, deren diagonale Einträge sind $1/\tilde{\gamma}_i$.

Die Berechnung der SVD ist weitaus aufwendiger als die meisten linearen Lösungstechniken, die wir in Kapitel 2 vorgestellt haben, daher ist diese erste Beobachtung hauptsächlich von theoretischem Interesse. Nehmen wir allgemeiner an, wir möchten eine Lösung der kleinsten Quadrate für $Ax \approx b$ finden, wobei $A \in \mathbb{R}^{m \times n}$ nicht unbedingt quadratisch ist. Aus unserer Diskussion der Normalgleichungen wissen wir, dass $x^T A^T A x = A^T b$ erfüllen muss. Bisher haben wir den Fall, dass A „kurz“ oder „unterbestimmt“ ist, also wenn A mehr Spalten als Zeilen hat, größtenteils außer Acht gelassen. In diesem Fall ist die Lösung der Normalgleichungen nicht eindeutig.

Um alle drei Fälle abzudecken, können wir ein Optimierungsproblem der folgenden Form lösen:

$$\begin{aligned} & \text{x minimieren} \quad \|Ax - b\|^2 \\ & \text{so dass } Ax = b \end{aligned}$$

In Worten ausgedrückt verlangt diese Optimierung, dass x die Normalgleichungen mit der kleinstmöglichen Norm erfüllt. Schreiben wir nun $A = U\tilde{V}$. Dann,

$$\begin{aligned} AA^T &= (U\tilde{V})^T (U\tilde{V}) \\ &= \tilde{V}^T U^T U \tilde{V} \text{ da } (AB)^T = B^T A^T \\ &= \tilde{V}^T \tilde{V}, \text{ da } U \text{ orthogonal ist} \end{aligned}$$

Daher ist die Frage, dass $Ax = b$ ist, dasselbe wie die Frage

$$\tilde{V}^T \tilde{V} x = \tilde{V}^T U b$$

$$\text{Oder äquivalent: } \tilde{y} = d$$

wenn wir $d = U b$ und $\tilde{y} = x$ annehmen. Beachten Sie, dass $y = x$, da U orthogonal ist, sodass unsere Optimierung zu Folgendem führt:

$$\begin{aligned} &\text{Minimiere } y^T y, \\ &\text{dass } \tilde{y} = d \end{aligned}$$

Da $\tilde{y}$ jedoch diagonal ist, besagt die Bedingung $\tilde{y} = d$ einfach $\tilde{y}_i = d_i$; Wenn also $\tilde{y}_i = 0$, müssen wir $y_i = d_i / \tilde{y}_i$ haben. Wenn $\tilde{y}_i = 0$, gibt es keine Einschränkung für y_i . Da wir also y minimieren, können wir genauso gut $y_i = 0$ annehmen. Mit anderen Worten, die Lösung für diese Optimierung ist $y = \tilde{y} + d$, wobei $\tilde{y} \in \mathbb{R}^{n \times m}$ hat die folgende Form:

$$+ \tilde{y}_{ij} \frac{1}{\tilde{y}_i} \text{ i = j, } \tilde{y}_i = 0 \text{ und sonst i \neq j}$$

Diese Form wiederum ergibt $x = Vy = V\tilde{y} + d = V\tilde{y} + Ub$.

Aus dieser Motivation heraus treffen wir folgende Definition:

Definition 6.1 (Pseudoinvers). Die Pseudoinverse von $A = U\tilde{V} \in \mathbb{R}^{m \times n}$ ist $A^+ = \tilde{V} \tilde{y} + U \in \mathbb{R}^{n \times m}$.

Unsere obige Ableitung zeigt, dass die Pseudoinverse von A die folgenden Eigenschaften aufweist:

- Wenn A quadratisch und invertierbar ist, ist $A^+ = A^{-1}$.
- Wenn A überbestimmt ist, ergibt $A^+ b$ die Lösung der kleinsten Quadrate für $Ax = b$.
- Wenn A unterbestimmt ist, ergibt $A^+ b$ die Lösung der kleinsten Quadrate für $Ax = b$ mit minimalem (Euklidische) Norm.

Auf diese Weise können wir endlich die unterbestimmten, vollständig bestimmten und überbestimmten Fälle von $Ax = b$ vereinheitlichen.

6.2.2 Zerlegung in äußere Produkte und Approximationen mit niedrigem Rang

Entwickeln wir das Produkt $A = U\tilde{V}$

, Es ist leicht zu zeigen, dass diese Beziehung Folgendes impliziert:

$$A = \sum_{i=1} \tilde{y}_i u_i v_i^T$$

wobei $\tilde{y} = \min\{m, n\}$ und u_i und v_i die i -ten Spalten von U bzw. V sind. Unsere Summe geht nur zu $\min\{m, n\}$, da wir wissen, dass die verbleibenden Spalten von U oder V durch $\tilde{y}$ auf Null gesetzt werden.

Dieser Ausdruck zeigt, dass jede Matrix als Summe der äußeren Produkte von zerlegt werden kann Vektoren:

Definition 6.2 (Äußeres Produkt). Das äußere Produkt von $u \in \mathbb{R}^m$ und $v \in \mathbb{R}^n$ ist die Matrix $u \tilde{y} v^T \in \mathbb{R}^{m \times n}$.

Angenommen, wir möchten das Produkt Ax schreiben. Dann könnten wir stattdessen schreiben:

$$\begin{aligned} Ax &= \sum_{i=1}^{\tilde{y}} u_i v_i^T x \\ &= \sum_{i=1}^{\tilde{y}} \tilde{y}_i u_i (v_i^T x) \\ &= \sum_{i=1}^{\tilde{y}} \tilde{y}_i (v_i^T x) u_i, \text{ da } x^T y = xy \end{aligned}$$

Das Anwenden von A auf x ist also dasselbe wie das lineare Kombinieren der u_i -Vektoren mit Gewichten $\tilde{y}_i(v_i^T x)$. Diese Strategie zur Berechnung von Ax kann zu erheblichen Einsparungen führen, wenn die Anzahl der $\tilde{y}_i$ -Werte ungleich Null relativ klein ist. Darüber hinaus können wir kleine Werte von $\tilde{y}_i$ ignorieren und diese Summe effektiv kürzen, um Ax mit weniger Aufwand anzunähern.

In ähnlicher Weise können wir aus §6.2.1 die Pseudoinverse von A wie folgt schreiben:

$$A^+ = \sum_{\tilde{y}_i > 0} \frac{v_i u_i^T}{\tilde{y}_i}.$$

Offensichtlich können wir den gleichen Trick anwenden, um A^+x auszuwerten, und tatsächlich können wir A^+x annähern, indem wir nur diejenigen Terme in der Summe auswerten, für die $\tilde{y}_i$ relativ klein ist. In der Praxis berechnen wir die Singulärwerte $\tilde{y}_i$ als Quadratwurzeln der Eigenwerte von AA^T oder $A^T A$, und Methoden wie Potenziteration können verwendet werden, um einen Teilsatz anstelle eines vollständigen Satzes von Eigenwerten aufzudecken.

Wenn wir also eine Reihe von Problemen der kleinsten Quadrate $Ax_i \approx b_i$ für verschiedene b_i lösen müssen und mit einer Näherung für x_i zufrieden sind, kann es sinnvoll sein, zunächst die kleinsten $\tilde{y}_i$ -Werte zu berechnen und die obige Näherung zu verwenden. Diese Strategie vermeidet außerdem, jemals die vollständige A^+ -Matrix berechnen oder speichern zu müssen, und kann genau sein, wenn A einen großen Bereich singulärer Werte aufweist.

Wenn wir zu unserer ursprünglichen Notation $A = U\tilde{Y}V^T$ zurückkehren, zeigt unser obiges Argument effektiv, dass eine potenziell nützliche Näherung für $A^+ \approx \tilde{Y}^+ U^T V^T$ ist, wobei $\tilde{Y}^+$ kleine Werte von $\tilde{Y}$ auf Null rundet.

Es lässt sich leicht überprüfen, dass der Spaltenraum von A^+ eine Dimension hat, die der Anzahl der Werte ungleich Null auf der Diagonale von $\tilde{Y}^+$ entspricht. Tatsächlich handelt es sich bei dieser Näherung nicht um eine Ad-hoc-Schätzung, sondern um die Lösung eines schwierigen Optimierungsproblems, das von dem folgenden berühmten Autor veröffentlicht wurde Satz (ohne Beweis angegeben):

Satz 6.1 (Eckart-Young, 1936). Angenommen, A wird aus $A = U\tilde{Y}V^T$ erhalten, indem alle außer $A \tilde{Y} A^T$ und $A \tilde{Y} A^T$ größten Singulärwerten $\tilde{y}_i$ von A auf Null gekürzt werden, abhängig von den k Einschränkungen, dass der Spaltenraum von A höchstens die Dimension k hat.

6.2.3 Matrixnormen

Die Konstruktion der SVD ermöglicht es uns auch, zu unserer Diskussion der Matrixnormen aus §3.3.1 zurückzukehren. Erinnern Sie sich beispielsweise daran, dass wir die Frobenius-Norm von A als definiert haben

$$\|A\|_{\text{Her}}^2 = \sum_{ij} a_{ij}^2.$$

Wenn wir $A = U\tilde{y}V$ schreiben, Wir können diesen Ausdruck vereinfachen:

$$\begin{aligned} \|A\|_{\text{Her}}^2 &= \sum_j \|\tilde{y} A e_j\|^2 \text{ da dieses Produkt die } j\text{-te Spalte von } A \text{ ist} \\ &= \sum_j \|\tilde{y} U \tilde{y} V e_j\|^2, \text{ ersetzt die SVD} \\ &= \sum_j \|e_j V \tilde{y}^T V e_j\|^2 \text{ seit } x^T = x x^T \text{ und } U \text{ ist orthogonal} \\ &= \sum_j \|V \tilde{y}\|_{\text{Her}}^2 \text{ Hin und her nach der gleichen Logik} \\ &= \sum_j \|V \tilde{y}\|_{\text{Her}}^2 \text{ da eine Matrix und ihre Transponierte dieselbe Frobenius-Norm haben} \\ &= \sum_j \|V \tilde{y} e_j\|^2 = \sum_j \tilde{y}^T V e_j V^T e_j \tilde{y} \text{ durch Diagonalität von } \tilde{y} \\ &= \sum_j \tilde{y}_j^2 \text{ da } V \text{ orthogonal ist} \end{aligned}$$

Somit ist die Frobenius-Norm von $A \in \mathbb{R}^{m \times n}$ die Summe der Quadrate seiner Singulärwerte.

Dieses Ergebnis ist von theoretischem Interesse, aber praktisch gesehen ist die grundlegende Definition der Frobenius-Norm bereits einfach zu bewerten. Noch interessanter ist, dass die induzierte Zwei-Norm von A gegeben ist durch

$$\|A\|_2 = \max\{\|\tilde{y}\| : \text{es existiert } x \in \mathbb{R}^n \text{ mit } A x = \tilde{y} x\}.$$

Nachdem wir nun Eigenwertprobleme untersucht haben, erkennen wir, dass dieser Wert die Quadratwurzel des größten Eigenwerts von AA^T oder ein Äquivalent ist

$$\|A\|_2 = \sqrt{\lambda_{\max}(AA^T)}.$$

Mit anderen Worten: Wir können die Zweinorm von A direkt aus seinen Eigenwerten ablesen.

Denken Sie auch daran, dass die Bedingungsanzahl von A durch $\text{cond } A = \|A\|_2 \|A^{-1}\|_2$ gegeben ist. Bei uns müssen Ableitung von A^{-1} , die Singulärwerte von A^{-1} die Kehrwerte der singulären Werte von A sein. Kombinieren Sie dies mit unserer Vereinfachung der $\|A\|_2$ -Erträge:

$$\text{cond } A = \frac{\tilde{y}_{\max}}{\tilde{y}_{\min}}.$$

Dieser Ausdruck liefert eine Strategie zur Bewertung der Konditionierung von A . Natürlich erfordert die Berechnung von $\tilde{y}_{\min}$ die Lösung von Systemen $Ax = b$, ein Prozess, der an sich unter einer schlechten Konditionierung von A leiden kann; Wenn dies ein Problem darstellt, kann die Konditionierung mithilfe verschiedener Näherungen der Singulärwerte von A begrenzt und angenähert werden.

6.2.4 Das Procrustes-Problem und die Ausrichtung

Viele Techniken in der Computer Vision beinhalten die Ausrichtung dreidimensionaler Formen. Angenommen, wir haben einen dreidimensionalen Scanner, der zwei Punktwolken desselben starren Objekts aus verschiedenen Ansichten sammelt. Eine typische Aufgabe könnte darin bestehen, diese beiden Punktwolken in einem einzigen Koordinatenrahmen auszurichten.

Da das Objekt starr ist, erwarten wir, dass es eine Rotationsmatrix R und eine Translation $t \in \mathbb{R}^3$ gibt, so dass durch Drehen der ersten Punktwolke um R und anschließendes Verschieben um t die beiden Datensätze ausgerichtet werden. Unsere Aufgabe ist es, diese beiden Objekte abzuschätzen.

Wenn sich die beiden Scans überschneiden, kann der Benutzer oder ein automatisiertes System n entsprechende Punkte markieren, die zwischen den beiden Scans übereinstimmen; wir können diese in zwei Matrizen $X_1, X_2 \in \mathbb{R}^{3 \times n}$ speichern. Dann erwarten wir für jede Spalte x_{1i} von X_1 und x_{2i} von X_2 $Rx_{1i} + t = x_{2i}$. Wir können eine Energiefunktion schreiben, die misst, inwieweit diese Beziehung zutrifft:

$$E \approx \sum_{i=1}^n \|Rx_{1i} + t - x_{2i}\|^2.$$

Wenn wir R fixieren und bezüglich t minimieren, wird die Optimierung von E offensichtlich zu einem Problem der kleinsten Quadrate. Nehmen wir nun an, wir optimieren für R mit festem Wert. Dies ist dasselbe wie das Minimieren $\|RX_1 - X_2\|_{\text{Her}}^2$, wobei die $\| \cdot \|_{\text{Her}}$ sind die von X_2 übersetzten Bytes, vorausgesetzt, dass R eine 3×3 -Rotationsmatrix ist, Spalten von X , also $RR^T = I_{3 \times 3}$, sind. Dies ist als orthogonales Procrustes-Problem bekannt.

Um dieses Problem zu lösen, führen wir die Spur einer quadratischen Matrix wie folgt ein:

Definition 6.3 (Trace). Die Spur von $A \in \mathbb{R}^{n \times n}$ ist die Summe ihrer Diagonalen:

$$\text{tr}(A) = \sum_{i=1}^n a_{ii}.$$

Es ist einfach zu überprüfen, ob $\|RX_1 - X_2\|_{\text{Her}}^2 = \text{tr}(A)$. Somit können wir E wie folgt vereinfachen:

$$\begin{aligned} \|RX_1 - X_2\|_{\text{Her}}^2 &= \text{tr}((RX_1 - X_2)(RX_1 - X_2)^T) \\ &= \text{tr}(RX_1X_1^T - RX_1X_2^T - X_2X_1^T + X_2X_2^T) \\ &= \text{konst.} - 2\text{tr}(X_2^T RX_1) \end{aligned}$$

da $\text{tr}(A + B) = \text{tr} A + \text{tr} B$ und $\text{tr}(A^T) = \text{tr}(A)$

Daher möchten wir $\text{tr}(X_2^T RX_1)$ mit $RR^T = I_{3 \times 3}$ maximieren. In den Übungen beweisen Sie, dass $\text{tr}(AB) = \text{tr}(BA)$. Somit lässt sich unser Ziel leicht zu $\text{tr}(RC)$ mit $C = X_1X_2^T$ vereinfachen. Wenn wir die SVD anwenden und $C = U\tilde{V}$ zerlegen, können wir noch mehr vereinfachen:

$$\begin{aligned} \text{tr}(RC) &= \text{tr}(RU\tilde{V}) \text{ per Definition} = \text{tr}(\tilde{V}^T RU) \\ &\text{da } \text{tr}(AB) = \text{tr}(BA) = \text{tr}(\tilde{R}^T) \text{ wenn wir } \tilde{R} = \\ &\tilde{V}^T RU \text{ definieren, was ebenfalls orthogonal ist } \tilde{R}^T \tilde{R} = I, \text{ da } \tilde{V} \text{ diagonal ist} \end{aligned}$$

Da $\tilde{R}$ orthogonal ist, haben alle seine Spalten eine Einheitslänge. Dies impliziert, dass $\tilde{r}_{ii} \leq 1$, da sonst die Norm der Spalte i zu groß wäre. Da $\tilde{r}_{ii} \geq 0$ für alle i ist, zeigt dieses Argument, dass wir $\text{tr}(RC)$ maximieren können, indem wir $\tilde{R} = I_{3 \times 3}$ annehmen. Wenn wir unsere Ersetzungen rückgängig machen, ergibt sich $R = V\tilde{R}U^T = VU^T$.

Allgemeiner gesagt haben wir Folgendes gezeigt:

Satz 6.2 (Orthogonaler Prokrustes). Die orthogonale Matrix R minimiert $\|RX - Y\|_F$, wobei SVD auf den Faktor $XY = U\tilde{Y}V$ angewendet wird. ² ist gegeben durch

Zurück zum Ausrichtungsproblem: Eine typische Strategie ist ein alternierender Ansatz:

1. Fixiere R und minimiere E bezüglich U .
2. Fixieren Sie das Resultierende und minimieren Sie E in Bezug auf R unter der Voraussetzung $RR^T = I_{3 \times 3}$.
3. Kehren Sie zu Schritt 1 zurück.

Die Energie E nimmt mit jedem Schritt ab und konvergiert somit zu einem lokalen Minimum. Da wir nie gleichzeitig optimieren und R , können wir nicht garantieren, dass das Ergebnis der kleinstmögliche Wert von E ist, aber in der Praxis funktioniert diese Methode gut.

6.2.5 Hauptkomponentenanalyse (PCA)

Erinnern Sie sich an den Aufbau aus §5.1.1: Wir möchten eine niedrigdimensionale Näherung für eine Menge von Datenpunkten finden, die wir in einer Matrix $X \in \mathbb{R}^{n \times k}$ für k Beobachtungen in n Dimensionen speichern können. Zuvor haben wir gezeigt, dass, wenn uns nur eine einzige Dimension erlaubt ist, die bestmögliche Richtung durch den dominanten Eigenvektor von XX^T gegeben ist.

Angenommen, wir dürfen stattdessen auf die Spanne von d Vektoren mit $d \leq \min\{k, n\}$ projizieren und möchten diese Vektoren optimal auswählen. Wir könnten sie in einer $n \times d$ -Matrix C schreiben; Da wir Gram-Schmidt auf jede Menge von Vektoren anwenden können, können wir annehmen, dass die Spalten von C orthonormal sind, was $CC^T = I_{d \times d}$ zeigt. Da C orthonormale Spalten hat, ist die Projektion von X auf den Spaltenraum von C durch die Normalengleichungen durch $CC^T X$ gegeben.

In diesem Aufbau möchten wir $\|X - CC^T X\|_F^2$ unter der Bedingung $CC^T = I_{d \times d}$ minimieren. Wir können vereinfachen. Unser Problem einigermassen:

$$\begin{aligned} \|X - CC^T X\|_F^2 &= \text{tr}((X - CC^T X)(X - CC^T X)^T) \text{ da } A = \text{tr}(A^T A) \\ &= \text{tr}(X X^T - 2X C C^T X + X C C^T C C^T X) = \text{const.} - 2 \text{tr}(X C C^T X) \\ &= \text{tr}(X C C^T X) \text{ da } CC^T = I_{d \times d} \\ &= \text{tr}(C^T X X^T C) \end{aligned}$$

Wir können also gleichbedeutend mit $C^T X X^T C$ den Mittelwert Null maximieren, sodass wir die Varianz der Projektion $C^T X$ maximieren möchten.

Nehmen wir nun an, wir faktorisieren $X = U\tilde{Y}V$. Dann möchten wir $\|C U\tilde{Y}V\|_F^2 = \|C\tilde{Y}\|_F^2$ durch Orthogonalität von V maximieren, wenn wir $C^T = CU$ annehmen. Wenn die Elemente von C^T c_{ij} sind, dann ergibt die Erweiterung dieser Norm

$$\|C\tilde{Y}\|_F^2 = \sum_{i,j} \tilde{y}_{ij}^2 c_{ij}^2$$

Aufgrund der Orthogonalität der Spalten von C^T wissen wir, dass $\sum_j c_{ij}^2 = 1$ für alle i und da C^T weniger als n Spalten haben kann, ist $\sum_i c_{ij}^2 \leq 1$. Somit ist der Koeffizient neben $\tilde{y}_{ij}^2$ in der obigen Summe höchstens 1, wenn wir also klar sind. Das Maximum wird erreicht, indem unsere Koordinatenänderung, man die Spalten von C^T so sortiert, dass $\tilde{y}_1 \geq \tilde{y}_2 \geq \dots \geq \tilde{y}_d$, ...ed. Wenn wir rückgängig machen, sehen wir, dass unsere Wahl von C die ersten d Spalten von U sein sollte.

Wir haben gezeigt, dass die SVD von X zur Lösung eines solchen Hauptkomponentenanalyseproblems (PCA) verwendet werden kann. In der Praxis werden die Zeilen von X normalerweise so verschoben, dass sie vor der SVD den Mittelwert Null haben; Wie in Abbildung NUMBER gezeigt, zentriert dies den Datensatz um den Ursprung und liefert aussagekräftigere PCA-Vektoren u_i .

6.3 Probleme

Teil III

Nichtlineare Techniken

Kapitel 7

Nichtlineare Systeme

So sehr wir uns auch bemühen, es ist einfach nicht möglich, alle Gleichungssysteme in dem linearen Rahmen auszudrücken, den wir in den letzten Kapiteln entwickelt haben. Es ist kaum notwendig, die Verwendung von Logarithmen, Exponentialfunktionen, trigonometrischen Funktionen, Absolutwerten, Polynomen usw. in praktischen Problemen zu motivieren, aber außer in einigen Sonderfällen ist keine dieser Funktionen linear. Wenn diese Funktionen auftreten, müssen wir einen allgemeineren, wenn auch weniger effizienten Maschinensatz einsetzen.

7.1 Probleme mit einer Variablen

Wir beginnen unsere Diskussion mit der Betrachtung von Problemen einer einzelnen Skalarvariablen. Insbesondere möchten wir für eine gegebene Funktion $f(x) : \mathbb{R} \rightarrow \mathbb{R}$ Strategien entwickeln, um Punkte $x \in \mathbb{R}$ zu finden, so dass $f(x) = 0$; wir nennen x eine Wurzel von f . Probleme mit einer Variablen in der linearen Algebra sind nicht besonders = b/ interessant; Schließlich können wir die Gleichung $ax + b = 0$ in geschlossener Form als $x = -b/a$ a. Die Lösung nichtlineare Gleichung lösen, $x^2 + e^x \cos y - 3 = 0$ ist jedoch weitaus weniger offensichtlich (übrigens die Lösung). z. B. y ist $y = \pm 1,30246 \dots$).

7.1.1 Charakterisierung von Problemen

Wir können nicht mehr davon ausgehen, dass f linear ist, aber ohne Annahmen über seine Struktur werden wir bei der Lösung von Systemen mit einer Variablen wahrscheinlich keine Fortschritte machen. Beispielsweise kann ein Löser garantiert nicht die Nullstellen von $f(x)$ finden, die durch gegeben sind

$$f(x) = \begin{cases} 1 & \text{für } x \in [0, 1] \\ 0 & \text{sonst} \end{cases}$$

Oder schlimmer:

$$f(x) = \begin{cases} 1 & \text{für } x \in \mathbb{Q} \\ 0 & \text{sonst} \end{cases}$$

Diese Beispiele sind in dem Sinne trivial, dass ein rationaler Kunde einer Root-Findungssoftware in diesem Fall wahrscheinlich nicht erwarten würde, dass dies gelingt, aber weit weniger offensichtliche Fälle sind nicht viel schwieriger konstruieren.

Aus diesem Grund müssen wir einige „regulierende“ Annahmen darüber hinzufügen, dass f einen Halt in die Möglichkeit bietet, Techniken zur Wurzelfindung zu entwerfen. Typische Annahmen dieser Art sind unten aufgeführt, in aufsteigender Reihenfolge ihrer Stärke:

- **Stetigkeit:** Eine Funktion f ist stetig, wenn sie gezeichnet werden kann, ohne einen Stift anzuheben; mehr formal ist f stetig, wenn die Differenz $f(x) - f(y)$ für $x \rightarrow y$ verschwindet.
 - **Lipschitz:** Eine Funktion f ist Lipschitz-stetig, wenn es eine Konstante C mit $|f(x) - f(y)| \leq C|x - y|$ gibt; Lipschitz-Funktionen müssen nicht differenzierbar sein, sind aber in ihren Änderungsraten begrenzt.
 - **Differenzierbarkeit:** Eine Funktion f ist differenzierbar, wenn ihre Ableitung f' für alle x existiert.
- dass f k -mal differenzierbar ist und jede dieser k Ableitungen stetig ist; • C : Eine Funktion ist C^k und zeigt an, dass f existieren und stetig sind.

Indem wir immer stärkere Annahmen über f hinzufügen, können wir effektivere Algorithmen entwerfen, um $f(x) = 0$ zu lösen. Wir werden diesen Effekt durch die Betrachtung einiger nachfolgender Algorithmen veranschaulichen.

7.1.2 Kontinuität und Halbierung

Nehmen wir an, alles, was wir über f wissen, ist, dass es stetig ist. In diesem Fall können wir einen intuitiven Satz aus der Standardrechnung mit einzelnen Variablen formulieren:

Satz 7.1 (Zwischenwertsatz). Angenommen, $f : [a, b] \rightarrow \mathbb{R}$ ist stetig. Angenommen, $f(x) < u < f(y)$. Dann existiert z zwischen x und y , so dass $f(z) = u$.

Mit anderen Worten: Die Funktion f muss jeden Wert zwischen $f(x)$ und $f(y)$ erreichen.

Angenommen, wir erhalten als Eingabe die Funktion f sowie zwei Werte u und r , so dass $f(u) \cdot f(r) < 0$; Beachten Sie, dass dies bedeutet, dass $f(u)$ und $f(r)$ entgegengesetzte Vorzeichen haben. Dann wissen wir durch den Zwischenwertsatz, dass es irgendwo zwischen u und r eine Wurzel von f gibt! Dies bietet eine offensichtliche Halbierungsstrategie zum Finden von x :

$\bar{y}$:

1. Berechnen Sie $c = (u+r)/2$.
2. Wenn $f(c) = 0$, gib x zurück $\bar{y} = c$.
3. Wenn $f(u) \cdot f(c) < 0$, nehme $r \leftarrow c$. Ansonsten nimm $\bar{y} \leftarrow c$.
4. Wenn $|r - \bar{y}| < \epsilon$, gib x zurück $\bar{y} \leftarrow c$.
5. Gehen Sie zurück zu Schritt 1

Diese Strategie teilt einfach das Intervall $[u, r]$ iterativ in zwei Hälften, wobei jedes Mal die Seite beibehalten wird, auf der bekanntermaßen eine Wurzel existiert. Nach dem Zwischenwertsatz konvergiert es eindeutig bedingungslos, in dem Sinne, dass, solange $f(u) \cdot f(r) < 0$, schließlich beide u und r garantiert zu einer gültigen Wurzel x konvergieren.

7.1.3 Analyse der Wurzelfindung

Die Bisektion ist die einfachste, aber nicht unbedingt die effektivste Technik zur Wurzelfindung. Wie bei den meisten Eigenwertmethoden ist die Halbierung von Natur aus iterativ und liefert möglicherweise nie eine exakte Lösung. Wir können x^k jedoch fragen, wie nahe der Wert x^k von c in der k -ten Iteration an der Wurzel x liegt, die wir berechnen möchten. Diese Analyse bietet eine Grundlage für den Vergleich mit anderen Methoden.

Nehmen wir im Allgemeinen an, dass wir eine Fehlergrenze E_k festlegen können, so dass die Schätzung x_k der x -Wurzel $\tilde{y}$ während der k -ten Iteration einer Wurzelfindungsmethode $|x_k - \tilde{y}|$ erfüllt $< E_k$. Offensichtlich stellt jeder Algorithmus mit $E_k \rightarrow 0$ ein konvergentes Schema dar; Die Konvergenzgeschwindigkeit kann jedoch durch die Geschwindigkeit charakterisiert werden, mit der sich $E_k \rightarrow 0$ nähert.

Beispielsweise in der Halbierung, da sowohl der x_k - als auch der x -Fehler durch E_k gegeben sind $|x_k - \tilde{y}|$. Da wir das Intervall bei jeder Iteration halbieren, wissen wir $E_{k+1} = 1/2 E_k$. Da E_{k+1} in E_k linear ist, sagen wir, dass die Halbierung eine lineare Konvergenz aufweist.

7.1.4 Fixpunktiteration

Es ist garantiert, dass die Halbierung für jede stetige Funktion f gegen eine Wurzel konvergiert. Wenn wir jedoch mehr über f wissen, können wir Algorithmen formulieren, die schneller konvergieren können.

Nehmen wir als Beispiel an, wir möchten x finden mit $g(x) = x$; Natürlich ist dieser Aufbau äquivalent zum Wurzelfindungsproblem, da das Lösen von $f(x) = 0$ dasselbe ist wie das Lösen von $f(x) + x = x$. Als zusätzliche Information könnten wir jedoch auch wissen, dass g Lipschitz mit der Konstante $C < 1$ ist.

Das System $g(x) = x$ legt eine mögliche Strategie nahe, die wir annehmen könnten:

1. Nehmen Sie x_0 als eine anfängliche Schätzung einer Wurzel.

2. Iteriere $x_k = g(x_{k-1})$.

Wenn diese Strategie konvergiert, ist das Ergebnis eindeutig ein Fixpunkt von g , der die oben genannten Kriterien erfüllt.

Glücklicherweise stellt die Lipschitz-Eigenschaft sicher, dass diese Strategie zu einer Wurzel konvergiert, sofern eine solche existiert. Wenn wir $E_k = |x_k - \tilde{y}|$ annehmen, dann haben wir die folgende Eigenschaft:

$$\begin{aligned} E_k &= |x_k - \tilde{y}| \\ &= |g(x_{k-1}) - g(\tilde{y})| \text{ durch Design des iterativen Schemas und Definition von } x_k \\ &\leq C |x_{k-1} - \tilde{y}| \text{ da } g \text{ Lipschitz ist} \\ &= C E_{k-1} \end{aligned}$$

Die induktive Anwendung dieser Aussage zeigt, dass $E_k \leq C^k |E_0|$ als $k \rightarrow \infty$. Also Festkomma-Iteration zum gewünschten x konvergiert!

Wenn g tatsächlich Lipschitz mit der Konstante $C < 1$ in einer Umgebung $[x - \delta, x + \delta]$, dann konvergiert die Fixpunktiteration, solange x_0 in diesem Intervall gewählt wird. Dies ist wahr, da unser obiger Ausdruck für E_k zeigt, dass er bei jeder Iteration kleiner wird. und $|g(x) - g(\tilde{y})| < 1$. Aufgrund

Ein wichtiger Fall tritt ein, wenn g C ist und man δ der Stetigkeit von g in diesem Fall $\tilde{y}$ mit $|g(x) - g(\tilde{y})| < 1$ für jedes $x \in N$, weiß, dass es eine Umgebung $N = [x - \delta, x + \delta]$,

$$|g(x) - g(y)| = |g(\tilde{y}) - g(y)| \cdot |x - \tilde{y}| \text{ nach dem Mittelwertsatz der Grundrechnung für ein } \tilde{y} \in [x, y] < (1 - C) |x - y|$$

Dies zeigt, dass g Lipschitz mit der Konstante $1 - C < 1$ in N ist. Wenn g also stetig differenzierbar ist und $|g'(x)| < 1$, konvergiert die Fixpunktiteration gegen $\tilde{y}$, wenn die anfängliche Schätzung x_0 nahe bei $\tilde{y}$ liegt.

¹ Diese Aussage ist schwer zu verstehen: Stellen Sie sicher, dass Sie sie verstehen!

Bisher haben wir kaum einen Grund, die Fixpunktiteration zu verwenden: Wir haben gezeigt, dass die Konvergenz nur dann garantiert ist, wenn g Lipschitz ist, und unser Argument über die E_k zeigt eine lineare Konvergenz wie eine Halbierung. Es gibt jedoch einen Fall, in dem die Festkomma-Iteration einen Vorteil bietet.

Angenommen, g ist differenzierbar mit $g'(x) = 0$. Dann verschwindet der Term erster Ordnung in der Taylor-Reihe für g und hinterlässt:

$$g(x_k) = g(x) + g'(x)(x_k - x) + \frac{1}{2} g''(\xi)(x_k - x)^2 + O((x_k - x)^3).$$

Somit haben wir in diesem Fall:

$$\begin{aligned} E_k &= |x_k - x| = |g(x_k) - g(x)| \text{ wie vorher} \\ &= \frac{1}{2} |g''(\xi)(x_k - x)^2| + O((x_k - x)^3) \text{ aus dem Taylor-Argument} \\ &= \frac{1}{2} (|g''(\xi)| + \tilde{\gamma})(x_k - x)^2 \text{ für einige } \tilde{\gamma}, \text{ solange } x_k \text{ nahe bei } x \text{ liegt} \\ &= \frac{1}{2} (|g''(\xi)| + \tilde{\gamma}) E_{k-1}^2 \end{aligned}$$

In diesem Fall ist E_k also quadratisch in E_{k-1} , wir sagen also, dass eine Fixpunktiteration quadratische Konvergenz haben kann; Beachten Sie, dass dieser Beweis der quadratischen Konvergenz nur gilt, weil wir $E_k \rightarrow 0$ bereits aus unserem allgemeineren Konvergenzbeweis kennen. Dies impliziert, dass $E_k \rightarrow 0$ viel schneller ist, sodass wir weniger Iterationen benötigen, um eine vernünftige Nullstelle zu erreichen.

Beispiel 7.1 (Konvergenz der Festkomma-Iteration).

7.1.5 Newtons Methode

Wir verschärfen unsere Funktionsklasse noch einmal, um eine Methode abzuleiten, die eine konsistentere quadratische Konvergenz aufweist. Nehmen wir nun erneut an, wir möchten $f(x) = 0$ lösen, gehen aber nun davon aus, dass f eine etwas engere Bedingung als Lipschitz ist.

An einem Punkt $x_k \in \mathbb{R}$ können wir ihn, da f jetzt differenzierbar ist, mit einer Tangente approximieren: $f(x) \approx$

$$f(x_k) + f'(x_k)(x - x_k)$$

Das Auflösen dieser Näherung nach $f(x) = 0$ ergibt eine Wurzel

$$x_{k+1} = x_k - \frac{f(x_k)}{f'(x_k)}.$$

Die Iteration dieser Formel ist als Newtons Methode zur Wurzelfindung bekannt und läuft darauf hinaus, eine lineare Näherung des nichtlinearen Problems iterativ zu lösen.

Beachten Sie das, wenn wir definieren

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

dann läuft Newtons Methode auf eine Fixpunktiteration auf g hinaus. Durch Differenzieren finden wir:

$$\begin{aligned} g'(x) &= 1 - \frac{f'(x)^2 f(x) + f(x)^2 f''(x)}{f(x)^2} \text{ nach der Quotientenregel } f(x) \\ &= \frac{f(x)^2 - f'(x)^2 f(x) - f(x)^2 f''(x)}{f(x)^2} \end{aligned}$$

Fixpunktiteration ist eine einfache Wurzel, was $f(x) = 0$ bedeutet. Dann ist $g(x) = 0$, und durch unsere Ableitung der obigen „Angenommen x “ wissen wir, dass Newtons Methode für einen hinreichend nahen Punkt quadratisch gegen x konvergiert erste Vermutung. Wenn also f mit einer einfachen Wurzel differenzierbar ist, liefert das Newton-Verfahren eine Festkomma-Iterationsformel, die garantiert quadratisch konvergiert; Wenn x jedoch nicht einfach ist, kann die Konvergenz linear oder schlechter sein. ⁹ Ist

Die Ableitung der Newton-Methode legt andere Methoden nahe, die durch die Verwendung weiterer Terme in der Taylor-Reihe abgeleitet werden. Beispielsweise fügt die „Halley-Methode“ den Iterationen Terme hinzu, an denen f beteiligt ist, und eine Klasse von „Haushaltsmethoden“ akzeptiert eine beliebige Anzahl von Ableitungen. Diese Techniken bieten eine noch höhere Konvergenz, allerdings auf Kosten der Auswertung komplexerer Iterationen und der Möglichkeit exotischerer Fehlermodi. Andere Methoden ersetzen Taylor-Reihen durch andere Grundformen; Beispielsweise verwendet die lineare Bruchinterpolation rationale Funktionen, um Funktionen mit Asymptotenstruktur besser anzunähern.

7.1.6 Sekantenmethode

Ein Effizienzproblem, auf das wir noch nicht eingegangen sind, sind die Kosten für die Bewertung von f und seinen Derivaten. Wenn f eine sehr komplizierte Funktion ist, möchten wir möglicherweise die Häufigkeit, mit der wir f berechnen müssen, oder noch schlimmer, f minimieren. Höhere Konvergenzordnungen helfen bei diesem Problem, aber wir können auch numerische Methoden entwerfen, die die Auswertung kostspieliger Ableitungen vermeiden.

Beispiel 7.2 (Design). Angenommen, wir entwerfen eine Rakete und möchten wissen, wie viel Treibstoff wir dem Motor hinzufügen müssen. Für eine gegebene Gallonenzahl x können wir eine Funktion $f(x)$ schreiben, die die maximale Höhe der Rakete angibt; Unsere Ingenieure haben angegeben, dass die Rakete eine Höhe h erreichen soll, also müssen wir $f(x) = h$ lösen. Die Bewertung von $f(x)$ erfordert die Simulation einer Rakete beim Start und die Überwachung ihres Treibstoffverbrauchs, was eine teure Angelegenheit ist, und obwohl wir vermuten könnten, dass f differenzierbar ist, können wir f möglicherweise nicht innerhalb einer praktikablen Zeitspanne bewerten.

Eine Strategie zum Entwerfen von Methoden mit geringeren Auswirkungen besteht darin, Daten so weit wie möglich wiederzuverwenden. Für Beispielsweise könnten wir leicht Folgendes annähern:

$$\frac{f(x_k) - f(x_{k-1})}{x_k - x_{k-1}}.$$

Das heißt, da wir $f(x_{k-1})$ in der vorherigen Iteration berechnen mussten, verwenden wir einfach die Steigung zu $f(x_k)$, um die Ableitung anzunähern. Sicherlich funktioniert diese Näherung gut, insbesondere wenn x_k nahe der Konvergenz liegt.

Wenn wir unsere Näherung in die Newton-Methode integrieren, ergibt sich ein neues iteratives Schema:

$$x_{k+1} = x_k - \frac{f(x_k)(x_k - x_{k-1}) - f(x_k)}{f(x_{k-1})}$$

Beachten Sie, dass der Benutzer zwei anfängliche Schätzungen x_0 und x_1 angeben muss, um dieses Schema zu starten, oder eine einzelne Iteration von Newton ausführen kann, um es zu starten.

Die Analyse der Sekantenmethode ist etwas komplizierter als die der anderen von uns betrachteten Methoden, da sie sowohl $f(x_k)$ als auch $f(x_{k-1})$ verwendet; Der Beweis seiner Konvergenz liegt außerhalb des Rahmens unserer Diskussion. Interessanterweise zeigt die Fehleranalyse, dass der Fehler zwischen linear und quadratisch mit einer Rate von $1 + \sqrt{5}/2$ (dem „Goldenen Schnitt“) abnimmt; Da die Konvergenz der des Newton-Verfahrens nahe kommt, ohne dass f ausgewertet werden muss, kann das Sekantenverfahren eine starke Alternative darstellen.

7.1.7 Hybridtechniken

Es kann zusätzliches Engineering durchgeführt werden, um zu versuchen, die Vorteile verschiedener Wurzelfindungsalgorithmen zu kombinieren. Beispielsweise könnten wir die folgenden Beobachtungen zu zwei Methoden machen, die wir besprochen haben:

- Die Halbierung konvergiert garantiert bedingungslos, jedoch nur mit einer linearen Rate.
- Die Sekantenmethode konvergiert schneller, wenn sie eine Wurzel erreicht, in manchen Fällen jedoch nicht konvergieren.

Angenommen, wir haben eine Wurzel von $f(x)$ in einem Intervall $[k, rk]$ wie bei einer Halbierung eingeklammert. Wir können das sagen, aktuellen Schätzung von x_k und ϵ ist gegeben durch $x_k = k$ wenn $|f(k)| < |f(rk)|$ andernfalls $x_k = rk$. Wenn wir bei unserer x_{k+1} im Auge behalten, könnten wir x_{k+1} als die nächste durch die Sekantenmethode gegebene Schätzung der Wurzel annehmen. Wenn x_k jedoch außerhalb des Intervalls $[k, rk]$ liegt, können wir es durch $k+rk/2$ ersetzen. Diese Korrektur garantiert, dass $x_{k+1} \in [k, rk]$, und unabhängig von der Wahl können wir durch Untersuchung des Vorzeichens von $f(x_{k+1})$ eine gültige Klammer $[k+1, rk+1]$ wie in der Halbierung erhalten. Dieser Algorithmus ist als „Dekker-Methode“ bekannt.

Die obige Strategie versucht, die unbedingte Konvergenz der Halbierung mit den stärkeren Wurzelschätzungen der Sekantenmethode zu kombinieren. In vielen Fällen ist es erfolgreich, aber seine Konvergenzrate ist etwas schwierig zu analysieren; Spezielle Fehlermodi können diese Methode auf lineare Konvergenz oder Schlimmeres reduzieren – tatsächlich kann Halbierung in manchen Fällen überraschenderweise schneller konvergieren! Andere Techniken, z. B. die „Brent-Methode“, führen häufiger Halbierungsschritte durch, um diesen Fall zu vermeiden, und können auf Kosten einer etwas komplexeren Implementierung ein garantiertes Verhalten zeigen.

7.1.8 Einzelvariablenfall: Zusammenfassung

Wir haben nun eine Reihe von Methoden zur Lösung von $f(x) = 0$ im Fall einer einzelnen Variablen vorgestellt und analysiert. An diesem Punkt ist es wahrscheinlich offensichtlich, dass wir nur an der Oberfläche solcher Techniken gekratzt haben; Es gibt viele iterative Schemata zur Wurzelfindung, alle mit unterschiedlichen Garantien, Konvergenzraten und Einschränkungen. Unabhängig davon können wir aufgrund unserer Erfahrungen eine Reihe von Beobachtungen machen:

- Aufgrund der möglichen generischen Form von f ist es unwahrscheinlich, dass wir Wurzeln x exakt finden können und uns stattdessen mit iterativen Verfahren zufrieden geben.
- Wir wünschen uns, dass die Folge x_k der Wurzelschätzungen x erreicht ϵ schnellstens. Wenn E_k eine Fehlergrenze ist, können wir eine Reihe von Konvergenzsituationen charakterisieren, indem wir $E_k \rightarrow 0$ als $k \rightarrow \infty$ annehmen. Eine vollständige Liste der Bedingungen, die gelten müssen, wenn k groß genug ist, finden Sie unten:
 1. Lineare Konvergenz: $E_{k+1} \leq C E_k$ für einige $C < 1$
 2. Superlineare Konvergenz: $E_{k+1} \leq C E_k^r$ für $r > 1$ (jetzt benötigen wir keine $C < 1$, da sich die r -Potenz aufheben kann, wenn E_k klein genug ist die Auswirkungen von C)
 3. Quadratische Konvergenz: $E_{k+1} \leq C E_k^2$
 4. Kubische Konvergenz: $E_{k+1} \leq C E_k^3$ (und so weiter)
- Eine Methode konvergiert möglicherweise schneller, erfordert aber bei jeder einzelnen Iteration zusätzliche Berechnungen. Aus diesem Grund kann es vorzuziehen sein, mehr Iterationen einer einfacheren Methode durchzuführen als weniger Iterationen einer komplexeren Methode.

7.2 Multivariable Probleme

Einige Anwendungen erfordern möglicherweise die Lösung eines allgemeineren Problems $f(x) = 0$ für eine Funktion $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$. Wir haben bereits einen Fall dieses Problems bei der Lösung von $Ax = b$ gesehen, was dem Finden von Wurzeln von $f(x) = Ax - b$ entspricht, aber der allgemeine Fall ist wesentlich schwieriger. Insbesondere Strategien wie die Halbierung lassen sich nur schwer erweitern, da wir nun weitgehend garantieren, dass m verschiedene Werte gleichzeitig alle Null sind.

7.2.1 Newtons Methode

Glücklicherweise lässt sich eine unserer Strategien auf unkomplizierte Weise erweitern. Denken Sie daran, dass wir für $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$ die Jacobi-Matrix schreiben können, die die Ableitung jeder Komponente von f in jeder Koordinatenrichtung angibt:

$$(Df)_{ij} = \frac{df_i}{dx_j}$$

Wir können die Jacobi-Funktion von f verwenden, um unsere Ableitung der Newton-Methode auf mehrere Dimensionen zu erweitern. Insbesondere ist die Näherung erster Ordnung von f gegeben durch:

$$f(x) \approx f(x_k) + Df(x_k) \cdot (x - x_k).$$

Das Einsetzen des gewünschten $f(x) = 0$ ergibt das folgende lineare System für die nächste Iteration x_{k+1} :

$$Df(x_k) \cdot (x_{k+1} - x_k) = -f(x_k)$$

Diese Gleichung kann mit der Pseudoinversen gelöst werden, wenn $m < n$; Wenn $m > n$, kann man die Methode der kleinsten Quadrate ausprobieren, aber sowohl die Existenz einer Wurzel als auch die Konvergenz dieser Technik sind unwahrscheinlich. Wenn Df jedoch quadratisch ist, entspricht die , wir erhalten die typische Iteration für Newton Methode $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$:

$$x_{k+1} = x_k - [Df(x_k)]^{-1} f(x_k), \text{ wobei}$$

wir wie immer die Matrix $[Df(x_k)]^{-1}$ nicht explizit berechnen, sondern sie vielmehr verwenden, um die Lösung eines linearen Systems zu signalisieren.

Die Konvergenz von Festkommamethoden wie der Newton-Methode, die $x_{k+1} = g(x_k)$ iteriert, erfordert, dass der Eigenwert der maximalen Größe des Jacobi-Dg kleiner als 1 ist. Nach Überprüfung dieser Annahme zeigt sich ein Argument, das dem eindimensionalen Fall ähnelt dass Newtons Methode $\approx$ kann, für die $Df(x)$ nichtsingulär ist. haben Wurzeln x quadratische Konvergenz in der Nähe von

7.2.2 Newton schneller machen: Quasi-Newton und Broyen

Wenn m und n zunehmen, wird die Newton-Methode sehr teuer. Für jede Iteration muss eine andere Matrix $Df(x_k)$ invertiert werden; Da es sich so oft ändert, hilft die Vorfaktorisierung von $Df(x_k) = L_k U_k$ nicht.

Einige Quasi-Newton-Strategien versuchen, unterschiedliche Approximationsstrategien anzuwenden, um einzelne Iterationen zu vereinfachen. Ein einfacher Ansatz könnte beispielsweise Df aus früheren Iterationen wiederverwenden und gleichzeitig $f(x_k)$ unter der Annahme neu berechnen, dass sich die Ableitung nicht sehr schnell ändert. Wir werden auf diese Strategien zurückkommen, wenn wir die Anwendung der Newton-Methode auf die Optimierung diskutieren.

Eine andere Möglichkeit besteht darin, zu versuchen, eine Parallele zu unserer Ableitung der Sekantenmethode herzustellen. So wie die Sekantenmethode immer noch eine Division enthält, machen solche Näherungen nicht unbedingt die Notwendigkeit, eine Matrix zu invertieren, zu verringern, sie ermöglichen jedoch die Durchführung einer Optimierung ohne explizite Berechnung des Jacobian Df . Solche Erweiterungen sind nicht ganz offensichtlich, da geteilte Differenzen keine vollständige approximative Jacobi-Matrix ergeben.

Bedenken Sie jedoch, dass die Richtungsableitung von f in Richtung v durch $D_v f = Df \cdot v$ gegeben ist. Wie bei der Sekantenmethode können wir diese Beobachtung zu unserem Vorteil nutzen, indem wir verlangen, dass unsere Näherung J einer Jacobi-Methode erfüllt

$$J \cdot (x_k - x_{k-1}) \approx f(x_k) - f(x_{k-1}).$$

Broydens Methode ist eine solche Erweiterung der Sekantenmethode, die nicht nur eine Schätzung x_k von x verfolgt, sondern auch eine Matrix J_k , die den Jacobi-Wert schätzt; Es müssen beide Anfangsschätzungen J_0 und x_0 angegeben werden. Angenommen, wir haben eine vorherige Schätzung J_{k-1} des Jacobi-Werts aus der vorherigen Iteration. Wir haben jetzt einen neuen Datenpunkt x_k , an dem wir $f(x_k)$ ausgewertet haben, daher möchten wir J_{k-1} unter Berücksichtigung dieser neuen Beobachtung auf einen neuen Jacobi- J_k aktualisieren. Ein vernünftiges Modell besteht darin, zu verlangen, dass die neue Näherung der alten Näherung bis auf die Richtung $x_k - x_{k-1}$ möglichst ähnlich ist :

$$\begin{aligned} &\text{minimiere } \|J_k - J_{k-1}\|_{\text{Her}}^2, \\ &\text{so dass } J_k \cdot (x_k - x_{k-1}) = f(x_k) - f(x_{k-1}) \end{aligned}$$

Um dieses Problem zu lösen, definieren Sie $\tilde{y} = J_k - J_{k-1}$ und $\tilde{x} = x_k - x_{k-1}$. Die Substitution ergibt die folgende Form:

$$\begin{aligned} &\text{minimiere } \|\tilde{y}\|_{\text{Her}}^2 \\ &\text{so, dass } \tilde{y} \cdot \tilde{x} = d \end{aligned}$$

Wenn wir $\tilde{y}$ als Lagrange-Multiplikator betrachten, ist diese Minimierung gleichbedeutend mit der Suche nach kritischen Punkten der Lagrange-Funktion $\tilde{y}$:

$$\tilde{y} = \tilde{y}_J + \tilde{\lambda} (\tilde{y} \cdot \tilde{x} - d)$$

Differenzieren nach $(\tilde{y}_J)_i$ zeigt:

$$0 = \frac{\partial \tilde{y}}{\partial (\tilde{y}_J)_i} = 2(\tilde{y}_J)_i + \tilde{\lambda} (\tilde{y}_J)_i = \tilde{y}_J = \tilde{y}(\tilde{y}_J)_i = \tilde{y}(\tilde{y}_J)_i^2$$

Das Einsetzen in $\tilde{y} \cdot \tilde{x} = d$ ergibt $\tilde{y}(\tilde{y}_J) \cdot \tilde{x} = \tilde{y}^2 d$, oder äquivalent $\tilde{y} = \tilde{y}^2 d / \tilde{y}_J^2$. Schließlich können wir Folgendes ersetzen:

$$\tilde{y}_J = \tilde{y}^2 - \tilde{y}(\tilde{y}_J) = \tilde{y}_J^2 \frac{d(\tilde{y}_J)}{d(\tilde{y}_J)}$$

Die Erweiterung unserer Vertretung zeigt:

$$\begin{aligned} J_k &= J_{k-1} + \tilde{y}_J \\ &= J_{k-1} + \frac{d(\tilde{y}_J)}{d(\tilde{y}_J)} \\ &= J_{k-1} + \frac{\tilde{y} f(x_{k-1}) - J_{k-1} \cdot \tilde{x}}{(x_k - x_{k-1})^2} \end{aligned}$$

Daher wechselt Broydens Methode einfach zwischen dieser Aktualisierung und dem entsprechenden Newton $f(x_k)$. In Schritt $x_{k+1} = x_k - J_k^{-1} f(x_k)$ manchen Fällen lässt sich zusätzliche Effizienz erzielen, indem man den Überblick behält die Matrix J_k^{-1} explizit anstelle der Matrix J_k , die mit einer ähnlichen Formel aktualisiert werden kann.

7.3 Konditionierung

Wir haben bereits in Beispiel 1.7 gezeigt, dass die Bedingungsanzahl der Wurzelfindung in einer einzelnen Variablen ist:

$$\text{cond}_x f = \left| \frac{f'(x)}{f(x)} \right|$$

Wie in Abbildung NUMBER dargestellt, zeigt diese Bedingungsnummer, dass die bestmögliche Situation für die Wurzelfindung dann auftritt, wenn sich f in der Nähe von x schnell ändert, da in diesem Fall eine Störung von x zu $f(x)$ führt, dass f Werte weit von 0 annimmt.

Die Anwendung eines identischen Arguments, wenn f mehrdimensional ist, ergibt eine Bedingungsanzahl von $\|Df(x)\|^{-1}$. Beachten Sie, dass die Bedingungsanzahl unendlich ist, wenn Df nicht invertierbar ist. Diese Kuriosität tritt auf, weil Störungen erster Ordnung x zu $f(x)$ bewahrt $f(x) = 0$, und eine solche Bedingung kann dies tatsächlich anspruchsvollen Wurzelfindungsfällen führen, wie sie in Abbildung NUMBER dargestellt sind.

7.4 Probleme

Viele Möglichkeiten, darunter:

- Viele mögliche Festkomma-Iterationsschemata für ein gegebenes Wurzelfindungsproblem, grafische Version der Festkomma-Iteration
- Mittlere Felditeration in ML
- Müllers Methode – komplexe Wurzeln
- Iterative Methoden höherer Ordnung – Householder-Methoden
- Interpretation von Eigensachen als Wurzelfindung
- Konvergenz der Sekantenmethode
- Wurzeln von Polynomen
- Newton-Fourier-Methode (!)
- „Modifizierte Newton-Methode bei nichtquadratischer Konvergenz“
- Konvergenz – Spektralradius für mehrdimensionales Newton; quadratische Konvergenz
- Sherman-Morrison-Update für Broyden

Kapitel 8

Uneingeschränkte Optimierung

In den vorherigen Kapiteln haben wir uns für einen weitgehend variierenden Ansatz zur Ableitung von Standardalgorithmen für die rechnerische lineare Algebra entschieden. Das heißt, wir definieren eine Zielfunktion, möglicherweise mit Einschränkungen, und stellen unsere Algorithmen als Minimierungs- oder Maximierungsproblem dar. Eine Auswahl aus unserer vorherigen Diskussion ist unten aufgeführt:

Problem	Zielsetzung	Einschränkungen
Kleinsten Quadrate	$E(x) = \ Ax - b\ ^2$	Keiner
Projektb aufa	$E(c) = \ ca - b\ ^2$	Keiner
Eigenvektoren der symmetrischen Matrix	$E(x) = x^T A x$	$x = 1$
Pseudoinvers	$E(x) = \ x\ ^2$	$A x = b$
Hauptkomponentenanalyse	$E(C) = \ X - C C^T X\ _{\text{Fro}}^2$	$C C^T = I_{d \times d}$
Broyden-Schritt	$E(J_k) = \ J_k - J_k^1\ _{\text{Her}}^2$	$J_k \cdot (x_k - x_{k-1}) = f(x_k) - f(x_{k-1})$

Offensichtlich ist die Formulierung von Problemen auf diese Weise ein wirkungsvoller und allgemeiner Ansatz. Aus diesem Grund ist es wertvoll, Algorithmen zu entwerfen, die ohne eine spezielle Form für die Energie E funktionieren, genauso wie wir Strategien entwickelt haben, um Wurzeln von f zu finden, ohne die Form von vornherein zu kennen.

8.1 Uneingeschränkte Optimierung: Motivation

In diesem Kapitel werden wir unbeschränkte Probleme betrachten, das heißt Probleme, die als Minimierung oder Maximierung einer Funktion $f: \mathbb{R}^n \rightarrow \mathbb{R}$ ohne Anforderungen an die Eingabe gestellt werden können. Es ist nicht schwer, in der Praxis auf solche Probleme zu stoßen; Nachfolgend listen wir einige Beispiele auf.

Beispiel 8.1 (Nichtlineare kleinste Quadrate). Angenommen, wir erhalten eine Anzahl von Paaren (x_i, y_i) , so dass $f(x_i) \approx y_i$, und wir möchten den besten Näherungswert für f innerhalb einer bestimmten Klasse finden. Beispielsweise können wir erwarten, dass f exponentiell ist. In diesem Fall sollten wir $f(x) = ce^{ax}$ für einige c und einige a schreiben können; Unsere Aufgabe ist es, diese Parameter zu finden. Eine einfache Strategie könnte darin bestehen, zu versuchen, die folgende Energie zu minimieren:

$$E(a, c) = \sum (y_i - ce^{ax_i})^2.$$

Diese Form für E ist nicht quadratisch, daher sind unsere linearen Methoden der kleinsten Quadrate nicht anwendbar.

Beispiel 8.2 (Maximum-Likelihood-Schätzung). Beim maschinellen Lernen besteht das Problem der Parameterschätzung darin, die Ergebnisse eines randomisierten Experiments zu untersuchen und zu versuchen, sie mithilfe einer Wahrscheinlichkeitsverteilung einer bestimmten Form zusammenzufassen. Beispielsweise könnten wir die Körpergröße jedes Schülers in einer Klasse messen und so eine Liste der Körpergrößen h_i für jeden Schüler i erhalten. Wenn wir viele Schüler haben, können wir die Verteilung der Schülergrößen mithilfe einer Normalverteilung modellieren:

$$g(h; \mu, \tilde{y}) = \frac{1}{\tilde{y} \sqrt{2\pi}} e^{-\frac{(h-\mu)^2}{2\tilde{y}^2}},$$

Dabei ist μ der Mittelwert der Verteilung und $\tilde{y}$ die Standardabweichung.

Unter dieser Normalverteilung ist die Wahrscheinlichkeit, dass wir die Körpergröße h_i für Schüler i beobachten, durch $g(h_i; \mu, \tilde{y})$ gegeben, und unter der (vernünftigen) Annahme, dass die Körpergröße von Schüler i wahrscheinlich unabhängig von der von Schüler j ist, Die Wahrscheinlichkeit, den gesamten Satz beobachteter Höhen zu beobachten, ist durch das Produkt gegeben

$$P(\{h_1, \dots, h_n\}; \mu, \tilde{y}) = \tilde{y}^{-n} \prod_{i=1}^n g(h_i; \mu, \tilde{y}).$$

Eine übliche Methode zum Schätzen der Parameter μ und $\tilde{y}$ von g besteht darin, P , betrachtet als Funktion von μ und $\tilde{y}$, mit festem $\{h_i\}$ zu maximieren; Dies wird als Maximum-Likelihood-Schätzung von μ und $\tilde{y}$ bezeichnet. In der Praxis optimieren wir normalerweise die logarithmische Wahrscheinlichkeit $(\mu, \tilde{y}) \mapsto \log P(\{h_1, \dots, h_n\}; \mu, \tilde{y})$; Diese Funktion hat die gleichen Maxima, verfügt aber über bessere numerische und mathematische Eigenschaften.

Beispiel 8.3 (Geometrische Probleme). Viele Geometrieprobleme, die in Grafik und Vision auftreten, lassen sich nicht auf Energien der kleinsten Quadrate reduzieren. Angenommen, wir haben eine Anzahl von Punkten $x_1, \dots, x_k \in \mathbb{R}^3$. Wenn wir diese Punkte gruppieren möchten, möchten wir sie möglicherweise mit einem einzigen x zusammenfassen, indem wir Folgendes minimieren:

$$E(x) = \sum_{i=1}^k \|x - x_i\|^2.$$

Das $x \in \mathbb{R}^3$, das E minimiert, ist als geometrischer Median von $\{x_1, \dots, x_k\}$. Beachten Sie, dass die Norm der Differenz $\|x - x_i\|$ in E nicht quadratisch ist, sodass die Energie in den Komponenten von x nicht mehr quadratisch ist.

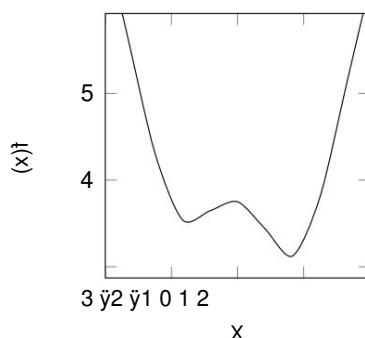
Beispiel 8.4 (Physikalische Gleichgewichte, angepasst von CITE). Angenommen, wir befestigen einen Gegenstand an einem Satz Federn. Jede Feder ist am Punkt $x_i \in \mathbb{R}^3$ verankert und hat die natürliche Länge L_i und die Konstante k_i . In Abwesenheit der Schwerkraft, wenn sich unser Objekt an der Position $p \in \mathbb{R}^3$ befindet, Das Quellennetz verfügt über potentielle Energie

$$E(p) = \frac{1}{2} \sum_{i=1}^n k_i (\|p - x_i\| - L_i)^2.$$

Gleichgewichte dieses Systems sind durch Minima von E gegeben und spiegeln Punkte p wider, an denen die Federkräfte alle ausgeglichen sind. Solche Gleichungssysteme werden verwendet, um Graphen $G = (V, E)$ zu visualisieren, indem für jedes Paar in E Eckpunkte in V mit Federn verbunden werden.

8.2 Optimalität

Bevor wir diskutieren, wie eine Funktion minimiert oder maximiert werden kann, sollten wir uns darüber im Klaren sein, wonach wir suchen. Beachten Sie, dass das Maximieren von f dasselbe ist wie das Minimieren von $-f$, sodass das Minimierungsproblem für unsere Betrachtung ausreichend ist. Für ein bestimmtes $f: \mathbb{R}^n \rightarrow \mathbb{R}$ und $x \in \mathbb{R}^n$ müssen wir $f(x)$ hat. Optimalitätsbedingungen, die ableiten, dass $\tilde{y}$ den kleinstmöglichen Wert verifizieren, dass x Natürlich möchten wir im Idealfall globale Optima von f finden:

Abbildung 8.1: Eine Funktion $f(x)$ mit mehreren Optima.

Definition 8.1 (Globales Minimum). Der Punkt $x^* \in \mathbb{R}^n$ ist ein globales Minimum von $f : \mathbb{R}^n \rightarrow \mathbb{R}$ wenn $f(x^*) \leq f(x)$ für alle $x \in \mathbb{R}^n$.

Das Finden eines globalen Minimums von f ohne Informationen über die Struktur von f erfordert effektiv eine Suche im Dunkeln. Angenommen, ein Optimierungsalgorithmus identifiziert das lokale Minimum in der Nähe von $x = -1$ in der Funktion in Abbildung 8.1. Es ist fast unmöglich, einfach durch Erraten der x -Werte zu erkennen, dass es ein zweites, niedrigeres Minimum in der Nähe von $x = 1$ gibt – soweit wir wissen, könnte es bei $x = 1000$ ein drittes, noch niedrigeres Minimum von f geben!

Daher begnügen wir uns in vielen Fällen damit, ein lokales Minimum zu finden:

Definition 8.2 (Lokales Minimum). Der Punkt $x^* \in \mathbb{R}^n$ ist ein lokales Minimum von $f : \mathbb{R}^n \rightarrow \mathbb{R}$, wenn $f(x^*) \leq f(x)$ für alle $x \in \mathbb{R}^n$, die $\|x - x^*\| \leq \delta$ erfüllen für ein $\delta > 0$.

Definition erfordert, dass der δ -Radius erreicht den kleinsten Wert in einer Umgebung, die durch definiert ist. Diese δ ist. Beachten Sie, dass lokale Optimierungsalgorithmen eine schwerwiegende Einschränkung haben, da sie nicht garantieren können, dass sie den niedrigstmöglichen Wert von f liefern, wie in Abbildung 8.1, wenn das linke lokale Minimum erreicht wird; Viele heuristische und andere Strategien werden angewendet, um die Landschaft möglicher x -Werte zu erkunden und so die Gewissheit zu gewinnen, dass ein lokales Minimum den bestmöglichen Wert hat.

8.2.1 Differenzielle Optimalität

Eine bekannte Geschichte aus der Ein- und Mehrvariablenrechnung besagt, dass das Finden potenzieller Minima und Maxima einer Funktion $f : \mathbb{R}^n \rightarrow \mathbb{R}$ einfacher ist, wenn f differenzierbar ist. Denken Sie daran, dass der Gradientenvektor $\nabla f = (\frac{\partial f}{\partial x_1}, \dots, \frac{\partial f}{\partial x_n})$ in die Richtung zeigt, in der f am stärksten zunimmt; der Vektor $-\nabla f$ zeigt in Richtung der größten Abnahme. Eine Möglichkeit, dies zu erkennen, besteht darin, sich daran zu erinnern, dass f in der Nähe von x_0 aussieht wie der lineare Funktionspunkt $x_0 \in \mathbb{R}^n$

$$f(x) \approx f(x_0) + \nabla f(x_0) \cdot (x - x_0).$$

Wenn wir $x - x_0 = \delta \nabla f(x_0)$ nehmen, dann finden wir:

$$f(x_0 + \delta \nabla f(x_0)) \approx f(x_0) + \delta \nabla f(x_0) \cdot \nabla f(x_0)$$

Wenn $\delta \nabla f(x_0) > 0$, bestimmt das Vorzeichen von δ , ob f zunimmt oder abnimmt.

Es ist nicht schwer, das obige Argument zu formalisieren, um zu zeigen, dass, wenn x_0 ein lokales Minimum ist, $\nabla f(x_0) = 0$ gelten muss. Beachten Sie, dass diese Bedingung notwendig, aber nicht ausreichend ist: Maxima und Sattel

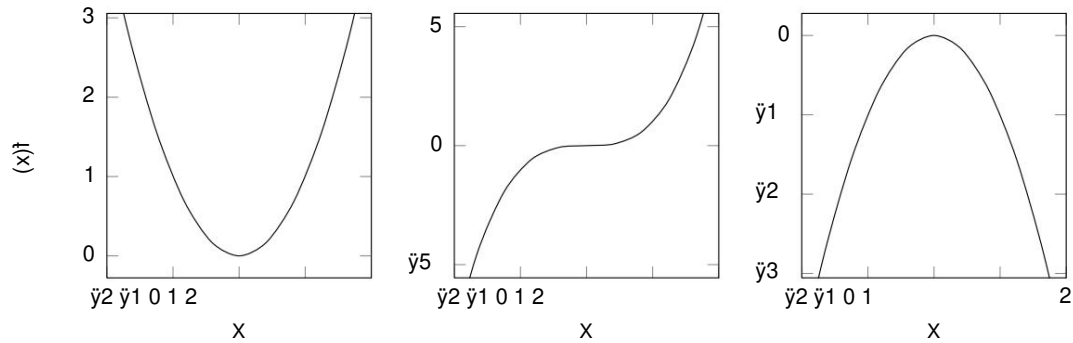


Abbildung 8.2: Kritische Punkte können viele Formen annehmen; Hier zeigen wir ein lokales Minimum, einen Sattelpunkt und ein lokales Maximum.

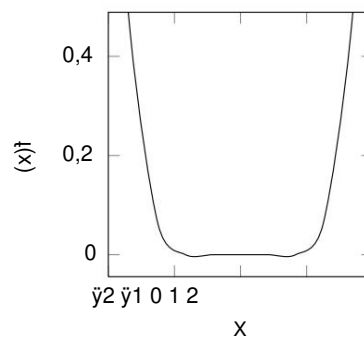


Abbildung 8.3: Eine Funktion mit vielen stationären Punkten.

Punkte haben auch $\nabla f(x) = 0$, wie in Abbildung 8.2 dargestellt. Dennoch liefert diese Beobachtung über Minima differenzierbarer Funktionen eine gemeinsame Strategie zur Wurzelfindung:

1. Finden Sie Punkte x_i , die $\nabla f(x_i) = 0$ erfüllen.
2. Prüfen Sie, welcher dieser Punkte ein lokales Minimum und kein Maximum oder Sattelpunkt ist.

Aufgrund ihrer wichtigen Rolle in dieser Strategie geben wir den von uns gesuchten Punkten einen besonderen Namen:

Definition 8.3 (Stationärer Punkt). Ein stationärer Punkt von $f : \mathbb{R}^n \rightarrow \mathbb{R}$ ist ein Punkt $x \in \mathbb{R}^n$, der $\nabla f(x) = 0$ erfüllt.

Das heißt, unsere Minimierungsstrategie kann darin bestehen, stationäre Punkte von f zu finden und dann diejenigen zu eliminieren, die keine Minima sind.

Es ist wichtig, im Auge zu behalten, wann wir mit dem Erfolg unserer Minimierungsstrategien rechnen können. In den meisten Fällen, wie beispielsweise in den in Abbildung 8.2 gezeigten, sind die stationären Punkte von f isoliert, was bedeutet, dass wir sie in eine diskrete Liste $\{x_0, x_1, \dots\}$ schreiben können. Ein degenerierter Fall ist jedoch in Abbildung 8.3 dargestellt; Hier besteht das gesamte Intervall $[-1/2, 1/2]$ aus stationären Punkten, so dass es unmöglich ist, sie einzeln zu betrachten. Wir werden Probleme wie degenerierte Fälle größtenteils ignorieren, werden aber auf sie zurückkommen, wenn wir die Konditionierung des Minimierungsproblems betrachten.

Angenommen, wir identifizieren einen Punkt $x \in \mathbb{R}^n$ als stationären Punkt von f und möchten nun prüfen, ob es sich um ein lokales Minimum handelt. Wenn f zweimal differenzierbar ist, besteht eine Strategie, die wir anwenden können, darin, seine Hessesche Funktion zu schreiben

Matrix:

$$Hf(x) = \begin{pmatrix} \frac{\partial^2 f}{\partial x_1^2} & \frac{\partial^2 f}{\partial x_1 \partial x_2} & \dots & \frac{\partial^2 f}{\partial x_1 \partial x_n} \\ \frac{\partial^2 f}{\partial x_2 \partial x_1} & \frac{\partial^2 f}{\partial x_2^2} & \dots & \frac{\partial^2 f}{\partial x_2 \partial x_n} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial^2 f}{\partial x_n \partial x_1} & \frac{\partial^2 f}{\partial x_n \partial x_2} & \dots & \frac{\partial^2 f}{\partial x_n^2} \end{pmatrix}$$

Wir können unserer Taylor-Entwicklung von f einen weiteren Term hinzufügen, um die Rolle von Hf zu sehen:

$$f(x) \approx f(x_0) + \nabla f(x_0) \cdot (x - x_0) + \frac{1}{2} (x - x_0)^T Hf(x_0) (x - x_0)$$

Wenn wir einen stationären Punkt x^* ersetzen, dann wissen wir per Definition:

$$f(x) \approx f(x^*) + \frac{1}{2} (x - x^*)^T Hf(x^*) (x - x^*)$$

Wenn Hf positiv definit ist, dann zeigt dieser Ausdruck $f(x) \approx f(x^*) + \frac{1}{2} (x - x^*)^T Hf(x^*) (x - x^*)$ und somit x^* ist ein lokales Minimum. Im Allgemeinen kann eine der folgenden Situationen auftreten:

- Wenn Hf positiv definit ist, dann x^* ist ein lokales Minimum von f .
- Wenn Hf negativ definit ist, dann x^* ist ein lokales Maximum von f .
- Wenn Hf unbestimmt ist, dann x^* ist ein Sattelpunkt von f .
- Wenn Hf nicht invertierbar ist, können Kuriositäten wie die Funktion in Abbildung 8.3 auftreten.

Um zu überprüfen, ob eine Matrix positiv definit ist, kann man prüfen, ob ihre Cholesky-Faktorisierung existiert, oder – langsamer – indem man prüft, ob alle ihre Eigenwerte positiv sind. Wenn also die Hesse-Funktion von f bekannt ist, können wir mithilfe der obigen Liste stationäre Punkte auf Optimalität prüfen. Viele Optimierungsalgorithmen, einschließlich derjenigen, die wir besprechen werden, ignorieren einfach den letzten Fall und benachrichtigen den Benutzer, da dies relativ unwahrscheinlich ist.

8.2.2 Optimalität über Funktionseigenschaften

Gelegentlich können wir, wenn wir mehr Informationen über $f : \mathbb{R}^n \rightarrow \mathbb{R}$ kennen, Optimalitätsbedingungen angeben, die stärker oder einfacher zu überprüfen sind als die oben genannten.

Eine Eigenschaft von f , die starke Auswirkungen auf die Optimierung hat, ist die Konvexität, dargestellt in Abbildung NUMMER:

Definition 8.4 (Konvex). Eine Funktion $f : \mathbb{R}^n \rightarrow \mathbb{R}$ ist konvex, wenn für alle $x, y \in \mathbb{R}^n$ und $\lambda \in (0, 1)$ die folgende Beziehung gilt:

$$f((1 - \lambda)x + \lambda y) \leq (1 - \lambda)f(x) + \lambda f(y).$$

Wenn die Ungleichung streng ist, ist die Funktion streng konvex.

Konvexität bedeutet, dass, wenn man im $\mathbb{R}^n$ zwei Punkte mit einer Linie verbindet, die Werte von f entlang der Linie kleiner oder gleich denen sind, die man durch lineare Interpolation erhalten würde.

Konvexe Funktionen verfügen über viele starke Eigenschaften, von denen die grundlegendste die folgende ist:

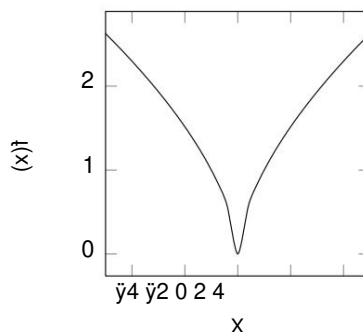


Abbildung 8.4: Eine quasikonvexe Funktion.

Vorschlag 8.1. Ein lokales Minimum einer konvexen Funktion $f : \mathbb{R}^n \rightarrow \mathbb{R}$ ist notwendigerweise ein globales Minimum.

Nachweisen. Nehmen Sie x als ein solches lokales Minimum und nehmen Sie an, dass es $x' = x$ mit $f(x') < f(x)$ gibt. Dann ist für $\gamma \in (0, 1)$

$$f(x + \gamma(x' - x)) \leq (1 - \gamma)f(x) + \gamma f(x') < f(x) \text{ durch Konvexität}$$

$$< f(x) \text{ da } f(x') < f(x)$$

Aber wenn man $\gamma \rightarrow 0$ annimmt, zeigt sich, dass x unmöglich ein lokales Minimum sein kann. □

Dieser Satz und damit verbundene Beobachtungen zeigen, dass es möglich ist, zu überprüfen, ob Sie ein globales Minimum einer konvexen Funktion erreicht haben, indem Sie einfach die Optimalität erster Ordnung anwenden. Daher ist es wertvoll, von Hand zu prüfen, ob eine zu optimierende Funktion zufällig konvex ist, eine Situation, die im wissenschaftlichen Rechnen überraschend häufig vorkommt; Eine hinreichende Bedingung, die einfacher zu überprüfen ist, wenn f zweimal differenzierbar ist, besteht darin, dass Hf überall positiv definit ist.

Andere Optimierungstechniken haben Garantien unter anderen Annahmen über f . Für die Prüfung beispielsweise ist eine schwächere Version der Konvexität die Quasikonvexität, die erreicht wird, wenn

$$f((1 - \gamma)x + \gamma y) \leq \max(f(x), f(y)).$$

Ein Beispiel für eine quasikonvexe Funktion ist in Abbildung 8.4 dargestellt; Obwohl sie nicht die charakteristische „Schalenform“ einer konvexen Funktion aufweist, verfügt sie über ein einzigartiges Optimum.

8.3 Eindimensionale Strategien

Wie im letzten Kapitel beginnen wir mit der eindimensionalen Optimierung von $f : \mathbb{R} \rightarrow \mathbb{R}$ und erweitern dann unsere Strategien auf allgemeinere Funktionen $f : \mathbb{R}^n \rightarrow \mathbb{R}$.

8.3.1 Newtons Methode

Unsere Hauptstrategie zur Minimierung differenzierbarer Funktionen $f : \mathbb{R}^n \rightarrow \mathbb{R}$ besteht darin, x zu finden, das $\nabla f(x)$ Maxima, Minima $\nabla f(x) = 0$ erfüllt. Angenommen, wir können überprüfen, ob stationäre Punkte Tionspunkte sind x oder Sattelpunkte als Post- In diesem Verarbeitungsschritt konzentrieren wir uns auf das Problem, die stationären Punkte x zu finden.

Nehmen wir dazu an, dass $f: \mathbb{R} \rightarrow \mathbb{R}$ differenzierbar ist. Dann wie in unserer Ableitung von Newton Methode zur Wurzelfindung können wir Folgendes annähern:

$$f(x) \approx f(x_k) + f'(x_k)(x - x_k) + \frac{f''(x_k)}{2}(x - x_k)^2.$$

Die Näherung auf der rechten Seite ist eine Parabel, deren Scheitelpunkt bei $x_k - f(x_k)/f'(x_k)$ liegt.

Natürlich ist f in Wirklichkeit nicht unbedingt eine Parabel, daher iteriert Newtons Methode einfach die Formel

$$x_{k+1} = x_k - \frac{f(x_k)}{f'(x_k)}.$$

Diese Technik lässt sich angesichts der Arbeit, die wir bereits im vorherigen Kapitel in das Verständnis der Newton-Methode zur Wurzelfindung gesteckt haben, leicht analysieren. Eine alternative Möglichkeit, die obige Formel abzuleiten, besteht insbesondere darin, die Wurzel auf $f(x)$ zu finden, da stationäre Punkte $f'(x) = 0$ erfüllen.

Daher weist Newtons Optimierungsverfahren in den meisten Fällen quadratische Konvergenz auf, vorausgesetzt, dass die anfängliche Schätzung x_0 hinreichend nahe an x liegt.

Eine natürliche Frage ist, ob die Sekantenmethode auf analoge Weise angewendet werden kann.

Unsere obige Ableitung der Newtonschen Methode findet Wurzeln von f' , sodass die Sekantenmethode verwendet werden könnte, um die Auswertung von f' , aber nicht von f zu eliminieren; Situationen, in denen wir f' , aber nicht f kennen, sind relativ selten. Eine geeignetere Parallele besteht darin, die zur Annäherung von f in der Sekantenmethode verwendeten Liniensegmente durch Parabeln zu ersetzen. Diese Strategie, bekannt als sukzessive parabolische Interpolation, minimiert auch eine quadratische Näherung von f bei jeder Iteration, aber anstatt $f(x_k)$, $f'(x_k)$ und $f''(x_k)$ zum Konstruieren der Näherung zu verwenden, verwendet sie $f(x_k)$, $f(x_{k+1})$ und $f(x_{k+2})$. Die Ableitung dieser Technik ist relativ einfach und sie konvergiert superlinear.

8.3.2 Suche im Goldenen Schnitt

Wir haben die Halbierung in unserer Parallele zu Techniken zur Wurzelfindung mit einer Variablen übersprungen. Für dieses Versäumnis gibt es viele Gründe. Unsere Motivation für die Halbierung war, dass sie nur die schwächste Annahme über f verwendete, die zum Finden von Wurzeln erforderlich war: Kontinuität. Der Zwischenwertsatz lässt sich jedoch nicht intuitiv auf Minima anwenden, sodass es den Anschein hat, dass es einen derart einfachen Ansatz nicht gibt.

Es ist jedoch wertvoll, mindestens eine Minimierungsstrategie zur Verfügung zu haben, die keine Differenzierbarkeit von f als zugrunde liegende Annahme erfordert; Schließlich gibt es nicht differenzierbare Funktionen, die eindeutige Minima haben, wie $f(x) = |x|$ bei $x = 0$. Zu diesem Zweck könnte eine alternative Annahme sein, dass f unimodular ist:

Definition 8.5 (Unimodular). Eine Funktion $f: [a, b] \rightarrow \mathbb{R}$ ist unimodular, wenn es $\bar{x} \in [a, b]$ gibt, sodass f für $x \in [a, \bar{x}]$ abnehmend und für $x \in [\bar{x}, b]$ steigend ist.

Mit anderen Worten: Eine unimodulare Funktion nimmt für einige Zeit ab und beginnt dann zuzunehmen; Es sind keine lokalisierten Minima zulässig. Beachten Sie, dass dies wie $|x|$ funktioniert sind nicht differenzierbar, aber dennoch unimodular.

Angenommen, wir haben zwei Werte x_0 und x_1 , so dass $a < x_0 < x_1 < b$. Wir können zwei Beobachtungen machen, die uns bei der Formulierung einer Optimierungstechnik helfen werden:

- Wenn $f(x_0) \geq f(x_1)$, dann wissen wir, dass $f(x) \geq f(x_1)$ für alle $x \in [a, x_0]$. Somit ist das Intervall $[a, x_0]$ kann bei unserer Suche nach einem Minimum von f verworfen werden.

- Wenn $f(x_1) \leq f(x_0)$, dann wissen wir, dass $f(x) \leq f(x_0)$ für alle $x \in [x_1, b]$, und können daher verwerfen $[x_1, b]$.

Diese Struktur legt eine mögliche Strategie zur Minimierung nahe, die mit dem Intervall $[a, b]$ beginnt und Teile gemäß den oben genannten Regeln iterativ entfernt.

Ein wichtiges Detail bleibt jedoch bestehen. Unsere Konvergenzgarantie für den Halbierungsalgorithmus beruhte auf der Tatsache, dass wir in jeder Iteration die Hälfte des betreffenden Intervalls entfernen konnten.

Wir könnten auf ähnliche Weise vorgehen und jedes Mal ein Drittel des Intervalls entfernen; Dies erfordert zwei Auswertungen von f während jeder Iteration an neuen x_0 - und x_1 -Positionen. Wenn die Auswertung von f jedoch teuer ist, möchten wir möglicherweise Informationen aus früheren Iterationen wiederverwenden, um mindestens eine dieser beiden Auswertungen zu vermeiden.

Vorerst gilt $a = 0$ und $b = 1$; Die Strategien, die wir unten ableiten, funktionieren allgemeiner durch Verschiebung und Skalierung. In Ermangelung weiterer Informationen über f könnten wir genauso gut eine symmetrische Wahl $x_0 = \bar{y}$ und $x_1 = 1 - \bar{y}$ für ein $\bar{y} \in (0, 1/2)$ treffen. Angenommen, unsere Iteration entfernt das Intervall ganz rechts $[x_1, b]$. Dann wird das Suchintervall zu $[0, 1 - \bar{y}]$ und wir kennen $f(\bar{y})$ aus der vorherigen Iteration. Die nächste Iteration wird $[0, 1 - \bar{y}]$ so dividieren, dass $x_0 = \bar{y}(1 - \bar{y})$ und $x_1 = (1 - \bar{y})$ ². Wenn wir

Wenn Sie $f(\bar{y})$ aus der vorherigen Iteration wiederverwenden möchten, können Sie $(1 - \bar{y})$ festlegen. ² = $\bar{y}$, ergibt:

$$\bar{y} = \frac{1}{2} (3 - \bar{y} - 5)$$

$$1 - \bar{y} = 2 - (\bar{y} - 5 - 1)$$

Der Wert von $1 - \bar{y}$ oben ist der Goldene Schnitt! Es ermöglicht die Wiederverwendung einer der Funktionsauswertungen aus den vorherigen Iterationen; Ein symmetrisches Argument zeigt, dass die gleiche Wahl von $\bar{y}$ funktioniert, wenn wir das linke Intervall anstelle des rechten entfernt hätten.

Der Golden-Section-Suchalgorithmus nutzt diese Konstruktion (CITE):

1. Nehmen Sie $\bar{y} = \frac{1}{2}(3 - \sqrt{5})$ und initialisieren Sie a und b so, dass f unimodal auf $[a, b]$ ist.
2. Machen Sie eine anfängliche Unterteilung $x_0 = a + (1 - \bar{y})(b - a)$ und $x_1 = a + \bar{y}(b - a)$.
3. Initialisieren Sie $f_0 = f(x_0)$ und $f_1 = f(x_1)$.
4. Iterieren Sie, bis $b - a$ ausreichend klein ist:

(a) Wenn $f_0 \leq f_1$, dann entfernen Sie das Intervall $[a, x_0]$ wie folgt:

- Nach links verschieben: $a \leftarrow x_0$
- Vorherige Iteration wiederverwenden: $x_0 \leftarrow x_1$, $f_0 \leftarrow f_1$
- Neue Stichprobe generieren: $x_1 \leftarrow a + \bar{y}(b - a)$, $f_1 \leftarrow f(x_1)$

(b) Wenn $f_1 > f_0$, dann entfernen Sie das Intervall $[x_1, b]$ wie folgt:

- Nach rechts verschieben: $b \leftarrow x_1$
- Vorherige Iteration wiederverwenden: $x_1 \leftarrow x_0$, $f_1 \leftarrow f_0$
- Neue Stichprobe generieren: $x_0 \leftarrow a + (1 - \bar{y})(b - a)$, $f_0 \leftarrow f(x_0)$

Dieser Algorithmus konvergiert offensichtlich bedingungslos und linear. Wenn f nicht global unimodal ist, kann es schwierig sein, $[a, b]$ zu finden, sodass f in diesem Intervall unimodal ist, was die Anwendungen dieser Technik etwas einschränkt; Im Allgemeinen wird $[a, b]$ geschätzt, indem versucht wird, ein lokales Minimum von f einzuklammern.

8.4 Multivariable Strategien

Wir setzen unsere Parallele zu unserer Diskussion über die Wurzelfindung fort, indem wir unsere Diskussion auf multivariable Probleme erweitern. Wie bei der Wurzelfindung sind multivariable Probleme erheblich schwieriger als Probleme mit einer einzelnen Variablen, kommen aber in der Praxis so oft vor, dass sie eine sorgfältige Betrachtung wert sind.

Wir betrachten hier nur den Fall, dass $f: \mathbb{R}^n \rightarrow \mathbb{R}$ differenzierbar ist. Optimierungsmethoden, die eher der Suche nach nicht differenzierbaren Funktionen im Goldenen Schnitt ähneln, haben begrenzte Anwendungsmöglichkeiten und sind schwer zu formulieren.

8.4.1 Gradientenabstieg

Erinnern Sie sich an unsere vorherige Diskussion, dass $\nabla f(x)$ in die Richtung des „steilsten Anstiegs“ von f bei x zeigt; In ähnlicher Weise ist der Vektor $-\nabla f(x)$ die Richtung des „steilsten Abstiegs“. Wenn nichts anderes garantiert, garantiert diese Definition, dass wir für kleine $\epsilon > 0$ haben müssen, wenn $\nabla f(x) = 0$

$$f(x - \epsilon \nabla f(x)) \leq f(x).$$

Angenommen, unsere aktuelle Schätzung des Ortes des Minimums von f ist x_k . Dann möchten wir vielleicht x_{k+1} wählen, sodass $f(x_{k+1}) < f(x_k)$ für eine iterative Minimierungsstrategie ist. Eine Möglichkeit, die Suche nach x_{k+1} zu vereinfachen, wäre die Verwendung eines unserer eindimensionalen Algorithmen aus §8.3 auf ein einfacheres Problem. Betrachten Sie insbesondere die Funktion $g_k(t) = f(x_k - t \nabla f(x_k))$, die f auf die Gerade durch x_k parallel zu $-\nabla f(x_k)$ beschränkt. Dank unserer Diskussion des Gradienten wissen wir, dass ein kleines t zu einer Verringerung von f führt.

Der Gradientenabstiegsalgorithmus löst diese eindimensionalen Probleme iterativ, um unsere Schätzung von x_k zu verbessern :

1. Wählen Sie eine erste Schätzung x_0

2. Iteriere bis zur Konvergenz von x_k :

(a) Nehmen Sie $g_k(t) = f(x_k - t \nabla f(x_k))$ (b)

Verwenden Sie einen eindimensionalen Algorithmus, um t zu $\min_{t \geq 0} g_k(t)$ (c) Nehmen Sie $x_{k+1} = x_k - t \nabla f(x_k)$

Jede Iteration des Gradientenabstiegs verringert $f(x_k)$, sodass die Zielwerte konvergieren. Der Algorithmus endet nur, wenn $\nabla f(x_k) = 0$, was zeigt, dass der Gradientenabfall mindestens ein lokales Minimum erreichen muss; Die Konvergenz ist jedoch für die meisten Funktionen f langsam. Der Liniensuchprozess kann durch eine Methode ersetzt werden, die das Ziel einfach um einen nicht vernachlässigbaren, wenn auch suboptimalen Betrag verringert, obwohl es in diesem Fall schwieriger ist, Konvergenz zu gewährleisten.

8.4.2 Newtons Methode

Parallel zu unserer Ableitung des Einzelvariablenfalls können wir eine Taylor-Reihennäherung von $f: \mathbb{R}^n \rightarrow \mathbb{R}$ unter Verwendung seines hessischen H_f schreiben :

$$f(x) \approx f(x_k) + \nabla f(x_k) \cdot (x - x_k) + \frac{1}{2} (x - x_k)^T H_f(x_k) (x - x_k)$$

Wenn man nach x differenziert und das Ergebnis gleich Null setzt, erhält man das folgende iterative Schema:

$$x_{k+1} = x_k - [Hf(x_k)]^{-1} \nabla f(x_k)$$

Es lässt sich leicht überprüfen, dass dieser Ausdruck eine Verallgemeinerung des Ausdrucks in §8.3.1 ist, und er konvergiert erneut quadratisch, wenn x_0 nahe einem Minimum liegt.

Abhängig vom Optimierungsziel f kann die Newton-Methode effizienter sein als der Gradientenabstieg. Denken Sie daran, dass jede Iteration des Gradientenabstiegs möglicherweise viele Auswertungen von f während der Liniensuchprozedur erfordert. Andererseits müssen wir den Hf -Wert bei jeder Iteration der Newton-Methode auswerten und invertieren. Beachten Sie, dass diese Faktoren keinen Einfluss auf die Anzahl der Iterationen, wohl aber auf die Laufzeit haben: Dies ist ein Kompromiss, der bei herkömmlichen Analysen möglicherweise nicht offensichtlich ist.

Es ist intuitiv, warum die Newton-Methode schnell konvergiert, wenn sie sich einem Optimum nähert. Insbesondere hat der Gradientenabstieg keine Kenntnis von Hf ; Es verhält sich analog zum Bergabgehen, indem man nur auf die Füße schaut. Durch die Verwendung von Hf erhält Newtons Methode ein umfassenderes Bild der Form von f in der Nähe.

Wenn Hf jedoch nicht positiv definit ist, sieht das Ziel lokal eher wie ein Sattel oder eine Spitze als wie eine Schüssel aus. In diesem Fall ist es möglicherweise nicht sinnvoll, zu einem ungefähren stationären Punkt zu springen.

Daher könnten adaptive Techniken prüfen, ob Hf positiv definit ist, bevor ein Newton-Schritt angewendet wird. Wenn es nicht positiv definit ist, können die Methoden auf den Gradientenabstieg zurückgreifen, um eine bessere Annäherung an das Minimum zu finden. Alternativ können sie Hf modifizieren, indem sie beispielsweise auf die nächste positiv definite Matrix projizieren.

8.4.3 Optimierung ohne Derivate: BFGS

Es kann schwierig sein, die Newton-Methode auf komplizierte Funktionen $f: \mathbb{R}^n \rightarrow \mathbb{R}$ anzuwenden. Die zweite Ableitung von f ist möglicherweise wesentlich komplizierter als die Form von f , und Hf ändert sich mit jeder Iteration, was es schwierig macht, Arbeiten aus früheren Iterationen wiederzuverwenden. Darüber hinaus hat Hf die Größe $n \times n$, sodass die Speicherung von Hf $O(n^2)$ Platz, der inakzeptabel sein kann.

Wie in unserer Erörterung der Wurzelfindung werden Minimierungstechniken, die die Newton-Methode imitieren, aber Näherungsableitungen verwenden, als Quasi-Newton-Methoden bezeichnet. Oftmals können sie ähnlich starke Konvergenzeigenschaften aufweisen, ohne dass bei jeder Iteration eine explizite Neubewertung und sogar Faktorisierung des Hesse-Werts erforderlich ist. In unserer Diskussion werden wir die Entwicklung von (CITE NO CEDAL AND WRIGHT) verfolgen.

Angenommen, wir möchten $f: \mathbb{R}^n \rightarrow \mathbb{R}$ mithilfe eines iterativen Schemas minimieren. Nahe der aktuellen Schätzung x_k der Wurzel, könnten wir f mit einem quadratischen Modell schätzen:

$$f(x_k + \tilde{y}) \approx f(x_k) + \nabla f(x_k) \cdot \tilde{y} + \frac{1}{2} (\tilde{y})^T B_k(\tilde{y}).$$

Beachten Sie, dass wir verlangt haben, dass unsere Näherung mit f erster Ordnung bei x_k übereinstimmt; Wie bei Broydens Methode zur Wurzelfindung lassen wir jedoch zu, dass unsere Schätzung des hessischen B_k variiert.

Dieses quadratische Modell wird minimiert, indem $\tilde{y} = \tilde{y}_k$ angenommen wird, $\tilde{y}_k^1 \nabla f(x_k)$. Falls $\tilde{y}_k^2$ groß ist und wir keinen so großen Schritt machen möchten, erlauben wir uns, diesen Unterschied um eine Schrittweite $\tilde{y}_k$ zu skalieren, was ergibt

$$x_{k+1} = x_k - \tilde{y}_k B_k^{-1} \nabla f(x_k).$$

Unser Ziel ist es, durch Aktualisierung von B_k eine vernünftige Schätzung B_{k+1} zu finden, damit wir diesen Vorgang wiederholen können.

Die Hesse-Funktion von f ist nichts anderes als die Ableitung von $\tilde{y} f$, daher können wir eine Bedingung im Sekantenstil für B_{k+1} schreiben:

$$B_{k+1}(x_{k+1} - y_k) = \tilde{y} f(x_{k+1}) - \tilde{y} f(x_k).$$

Wir ersetzen $\tilde{y} x_{k+1} - y_k$ und $y_k - \tilde{y} f(x_{k+1}) - \tilde{y} f(x_k)$, was eine äquivalente Bedingung $B_{k+1} s_k = y_k$ ergibt.

Angeichts der vorliegenden Optimierung wünschen wir uns, dass B_k zwei Eigenschaften hat:

- B_k sollte eine symmetrische Matrix sein, wie die hessische H_f .
- B_k sollte positiv (semi-)definit sein, sodass wir eher nach Minima als nach Maxima suchen Sattelpunkte.

Die Symmetriebedingung reicht aus, um die Möglichkeit der Verwendung der Broyden-Schätzung auszuschließen, die wir im vorherigen Kapitel entwickelt haben.

Die positiv definite Einschränkung stellt implizit eine Bedingung für die Beziehung zwischen s_k und y_k . Insbesondere wird die Beziehung $B_{k+1} s_k = y_k$ mit $s_k^T B_{k+1} s_k = s_k^T y_k$ vormultipliziert. Für dich s_k^T zeigt s_k

Damit B_{k+1} positiv definit ist, müssen wir dann $s_k^T y_k > 0$ haben. Diese Beobachtung kann unsere Wahl von $\tilde{y}_k$ leiten; es ist leicht zu erkennen, dass dies für hinreichend kleine $\tilde{y}_k > 0$ gilt.

Nehmen Sie an, dass s_k und y_k unsere Kompatibilitätsbedingung erfüllen. Wenn das erledigt ist, können wir schreiben eine Optimierung im Broyden-Stil, die zu einer möglichen Näherung B_{k+1} führt:

$$\begin{aligned} &\text{minimieren } B_{k+1} B_{k+1}^{-1} \tilde{y} B_k, \text{ so} \\ &\text{dass } B_{k+1} s_k = y_k \end{aligned}$$

Zur geeigneten Auswahl von Normen: Iteratives. Diese Optimierung ergibt das bekannte DFP (Davidon) Schema von Fletcher-Powell.

Anstatt die Details des DFP-Schemas auszuarbeiten, gehen wir zu einer populäreren Methode über, die als BFGS-Formel (Broyden-Fletcher-Goldfarb-Shanno) bekannt ist und in vielen modernen Systemen vorkommt. Beachten Sie, dass – wenn wir vorerst unsere Wahl von $\tilde{y}_k$ ignorieren – unsere Näherung zweiter Ordnung minimiert wurde, indem $\tilde{y}_x = \tilde{y} B$ angenommen wurde

wird durch die inverse Matrix B_k vorgegeben. Die Annahme, dass $B_{k+1} \tilde{y} B_k$ klein ist, kann dennoch relativ große Unterschiede zwischen der Wirkung von B_k implizieren $\tilde{y}_k^1$ und das von $B_{k+1} \tilde{y}_k^1$. Somit ist am Ende das Verhalten unseres iterativen Schemas

Unter Berücksichtigung dieser Beobachtung nimmt das BFGS-Schema eine kleine Änderung an der obigen Ableitung vor. Anstatt B_k bei jeder Iteration zu berechnen, können wir seinen Kehrwert $H_k \tilde{y} B$ direkt berechnen. Jetzt wird unsere Bedingung $B_{k+1} s_k = y_k$ umgekehrt zu $s_k = H_{k+1} y_k$; Die Bedingung, dass B_k symmetrisch ist, ist dasselbe wie die Forderung, dass H_k symmetrisch ist. Wir lösen eine Optimierung

$$\begin{aligned} &\text{minimieren } H_{k+1} H_{k+1}^{-1} \tilde{y} H_k, \text{ so} \\ &\text{dass } H_{k+1} s_k = H_{k+1} y_k \end{aligned}$$

Diese Konstruktion hat den schönen Nebenvorteil, dass zur Berechnung von $\tilde{y}_x = \tilde{y} H_k \tilde{y} f(x_k)$ keine Matrixinversion erforderlich ist.

Um eine Formel für H_{k+1} abzuleiten, müssen wir uns für eine Matrixnorm entscheiden. Wie bei unserer vorherigen Diskussion kommt die Frobenius-Norm der Optimierung der kleinsten Quadrate am nächsten, was es wahrscheinlich macht, dass wir einen Ausdruck in geschlossener Form für H_{k+1} generieren können, anstatt die obige Minimierung als Unterprogramm der BFGS-Optimierung lösen zu müssen.

Die Frobenius-Norm hat jedoch einen gravierenden Nachteil für hessische Matrizen. Denken Sie daran, dass die Hesse-Matrix Einträge $(Hf)_{ij} = \ddot{y} f_i / \ddot{y} x_j$ hat. Oftmals können die Größen x_i für verschiedene i unterschiedliche Einheiten haben; Erwägen Sie beispielsweise die Maximierung des Gewinns (in Dollar), der durch den Verkauf eines Cheeseburgers mit dem Radius r (in Zoll) und dem Preis p (in Dollar) erzielt wird, was zu f : (Zoll, Dollar) $\ddot{y}$ Dollar führt. Es macht keinen Sinn, diese unterschiedlichen Größen zu quadrieren und zu addieren.

Angenommen, wir finden eine symmetrische positiv definite Matrix W , so dass $W s_k = y_k$; Wir werden in den Übungen prüfen, ob eine solche Matrix existiert. Eine solche Matrix nimmt die Einheiten von $s_k = x_{k+1} - x_k$ zu denen von $y_k = \ddot{y} f(x_{k+1}) - \ddot{y} f(x_k)$ auf. Inspiriert von unserem Ausdruck $A = \text{Tr}_{\text{Her}}^2(A, A)$ können wir eine gewichtete Frobenius-Norm einer Matrix A definieren als

$$A^2_W \ddot{y} \text{Tr}(A W A W)$$

Es ist einfach zu überprüfen, ob dieser Ausdruck konsistente Einheiten hat, wenn er auf unsere Optimierung für H_{k+1} angewendet wird. Wenn sowohl W als auch A symmetrisch zu den Spalten w_i bzw. a_i sind, ergibt die Erweiterung des obigen Ausdrucks:

$$A^2_W = \ddot{y} \sum_{ij} (w_i \cdot a_j)(w_j \cdot a_i).$$

Diese Wahl der Norm in Kombination mit der Wahl von W ergibt eine besonders saubere Formel für H_{k+1} und y_k : gegeben H_k, s_k ,

$$H_{k+1} = (\ln \times n \ddot{y} \ddot{y} s_k s_k^T) H_k (\ln \times n \ddot{y} \ddot{y} y_k y_k^T) + \ddot{y} s_k s_k^T,$$

wobei $\ddot{y} k \ddot{y} 1/y$ -s. Wir zeigen im Anhang zu diesem Kapitel, wie man diese Formel herleitet.

Der BFGS-Algorithmus vermeidet die Notwendigkeit, eine Hesse-Matrix für f zu berechnen und zu invertieren, erfordert aber dennoch $O(n^2)$ Speicher für H_k . Eine Variante namens L-BFGS („Limited-Memory BFGS“) vermeidet dieses Problem, indem sie eine begrenzte Historie der Vektoren y_k führt und Anwenden von H_k durch rekursive Erweiterung seiner Formel. Dieser Ansatz kann trotz seiner kompakten Raumnutzung tatsächlich bessere numerische Eigenschaften haben; insbesondere sind alte Vektoren y_k und s_k möglicherweise nicht mehr relevant und sollten ignoriert werden.

8.5 Probleme

Ideenliste:

- Leiten Sie Gauss-Newton ab
- Stochastische Methoden, AdaGrad
- VSCG-Algorithmus
- Wolfe-Bedingungen für den Gradientenabstieg; An BFGS anschließen
- Sherman-Morrison-Woodbury-Formel für B_k für BFGS
- Beweisen Sie, dass BFGS konvergiert; zeigen Existenz einer Matrix W
- (Verallgemeinerter) reduzierter Gradientenalgorithmus
- Konditionsnummer zur Optimierung

Anhang: Ableitung von BFGS Update¹

Unsere Optimierung für H_{k+1} hat den folgenden Lagrange-Multiplikatorausdruck (zur Vereinfachung der Notation nehmen wir $H_{k+1} \approx H$ und $H_k = H$):

$$\begin{aligned} & \sum_{ij} \tilde{y}_i \tilde{y}_j (w_i \cdot (h_j \tilde{y}_h - \tilde{y}_j)) (w_j \cdot (h_i \tilde{y}_h - \tilde{y}_i)) \tilde{y}_i \tilde{y}_j (H_{ij} - H_{ji}) \tilde{y}_i \tilde{y}_j (H_{yk} - \tilde{y}_{sk}) \\ & = \sum_{ij} \tilde{y}_i (w_i \cdot (h_j \tilde{y}_h - \tilde{y}_j)) (w_j \cdot (h_i \tilde{y}_h - \tilde{y}_i)) \tilde{y}_i \tilde{y}_j H_{ij} \tilde{y}_i \tilde{y}_j (H_{yk} - \tilde{y}_{sk}), \text{ wenn wir } \tilde{y}_{ij} = \tilde{y}_i \tilde{y}_j \text{ annehmen} \end{aligned}$$

Die Bildung von Ableitungen zum Finden kritischer Punkte zeigt (für $y \approx y_k, s \approx s_k$):

$$\begin{aligned} 0 &= \frac{\partial \tilde{y}}{\partial \tilde{y}} = \tilde{y} 2w_i (w_j \cdot (h_j \tilde{y}_h - \tilde{y}_j)) \tilde{y}_i \tilde{y}_j \tilde{y}_i \tilde{y}_j H_{ij} = \\ & w_i (W(H \tilde{y} H \tilde{y}))_j \tilde{y}_i \tilde{y}_j \tilde{y}_i \tilde{y}_j \\ &= 2 \tilde{y} (W(H \tilde{y} H \tilde{y}))_{ji} \tilde{y}_i \tilde{y}_j \tilde{y}_i \tilde{y}_j \text{ aufgrund der Symmetrie von } W \\ &= 2 (W(H \tilde{y} H \tilde{y}) W)_{ji} \tilde{y}_i \tilde{y}_j \tilde{y}_i \tilde{y}_j \\ &= 2 (W(H \tilde{y} H \tilde{y}) W)_{ij} \tilde{y}_i \tilde{y}_j \tilde{y}_i \tilde{y}_j \text{ aufgrund der Symmetrie von } W \text{ und } H \end{aligned}$$

In Matrixform haben wir also die folgende Liste von Fakten:

$$\begin{aligned} 0 &= 2W(H \tilde{y} H \tilde{y}) W \tilde{y} A \tilde{y} \tilde{y}, \text{ wobei } A_{ij} = \tilde{y}_{ij} \\ A &= \tilde{y} A, W = W, H = H, (H \tilde{y}) = H \\ H y &= s, W s = y \end{aligned}$$

Durch Transposition in Kombination mit der Symmetrie von H und W und der Asymmetrie von A können wir ein Beziehungspaar erreichen:

$$\begin{aligned} 0 &= 2W(H \tilde{y} H \tilde{y}) W \tilde{y} A \tilde{y} \tilde{y} \\ 0 &= 2W(H \tilde{y} H \tilde{y}) W + A \tilde{y} \tilde{y} \tilde{y} = \tilde{y} 0 = \\ & 4W(H \tilde{y} H \tilde{y}) W \tilde{y} \tilde{y} \tilde{y} \end{aligned}$$

Die nachträgliche Multiplikation dieser Beziehung mit s zeigt:

$$0 = 4(y \tilde{y} W H \tilde{y} y) \tilde{y} \tilde{y} (y \cdot s) \tilde{y} y (\tilde{y} \cdot s)$$

Nehmen Sie nun das Skalarprodukt mit:

$$0 = 4(y \cdot s) \tilde{y} 4(y H \tilde{y} y) \tilde{y} 2(y \cdot s) (\tilde{y} \cdot s)$$

Das zeigt:

$$\tilde{y} \cdot s = 2 \tilde{y} y (s \tilde{y} H \tilde{y} y), \text{ für } \tilde{y} \approx 1/y \cdot s$$

¹Besonderer Dank geht an Tao Du für das Debuggen mehrerer Teile dieser Ableitung.

Nun setzen wir dies in unsere Vektorgleichung ein:

$$\begin{aligned} 0 &= 4(y - WHy) - y(y - s) - y(y - s) \text{ von vorher} = \\ &= 4(y - WHy) - y(y - s) - y[2y - (s - Hy)] \text{ aus unserer Vereinfachung} \\ &= y - y = 4y - 4WHy - y - 2y - 2y(s - Hy) \end{aligned}$$

Die nachträgliche Multiplikation mit y ergibt:

$$yy = 4y(y - WHy) - y - 2y - 2y(s - Hy)yy$$

Nehmen Sie die Transponierung,

$$y\tilde{y} = 4\tilde{y}y(y - Hy - W) - \tilde{y} - 2\tilde{y} - 2\tilde{y}(s - Hy)y$$

Kombinieren Sie diese Ergebnisse und dividieren Sie sie durch vier Shows:

$$\frac{1}{4}(\tilde{y}y + y\tilde{y}) = \tilde{y}(2yy - WHy - yy - yHy - W) - \tilde{y} - 2\tilde{y}(s - Hy)y$$

Jetzt führen wir eine Vor- und Nachmultiplikation mit $W\tilde{y}1$ durch. Da $Ws = y$, können wir äquivalent schreiben: $= W\tilde{y}1y$; außerdem wissen wir dann aufgrund der Symmetrie von W , dass $yW\tilde{y}1 = s$ ist. Die Anwendung dieser Identitäten auf den obigen Ausdruck zeigt:

$$\begin{aligned} \frac{1}{4}W\tilde{y}1(\tilde{y}y + y\tilde{y})W\tilde{y}1 &= 2\tilde{y}ss - \tilde{y}Hy - ysy - y - \tilde{y} - 2(y - Hy)ss + \tilde{y} - 2(y - Hy)yss \\ &= 2\tilde{y}ss - \tilde{y}Hy - ysy - y - \tilde{y} - \tilde{y}ss + s\tilde{y} - 2(y - Hy)s \text{ nach Definition von } \tilde{y} \\ &= \tilde{y}ss - \tilde{y}Hy - ysy - y - \tilde{y} + s\tilde{y} - 2(y - Hy)s \end{aligned}$$

Abschließend können wir unsere Ableitung des BFGS-Schritts wie folgt abschließen:

$$\begin{aligned} 0 &= 4W(H - Hy - W) - y - y - y \text{ von vorher} \\ = \tilde{y}H &= 4 \left[-W\tilde{y}1(\tilde{y}y + y\tilde{y})W\tilde{y}1 + H - ysy - y - \tilde{y} - 2(y - Hy)s + H - y \right] \text{ aus dem letzten Absatz} \\ &= H - \tilde{y}(I - \tilde{y}y) + \tilde{y}ss - \tilde{y}sy - H\tilde{y}(I - \tilde{y}y) + (\tilde{y}sy)H\tilde{y}(\tilde{y}ys) \\ &= H - \tilde{y}ys + \tilde{y}ss - \tilde{y}sy - H\tilde{y}(I - \tilde{y}y) = \tilde{y}ss + (I - \tilde{y}y) \\ &\quad \tilde{y}sy)H\tilde{y}(I - \tilde{y}y) \end{aligned}$$

Dieser letzte Ausdruck ist genau der im Kapitel vorgestellte BFGS-Schritt.

Kapitel 9

Optimierungsprobleme

Wir setzen unsere Betrachtung von Optimierungsproblemen fort, indem wir den eingeschränkten Fall untersuchen. Diese Probleme haben die folgende allgemeine Form:

$$\begin{aligned} &\text{Minimiere } f(x) \text{ so,} \\ &\text{dass } g(x) = 0 \text{ } h(x) \leq 0 \end{aligned}$$

Hier gilt $f : \mathbb{R}^n \rightarrow \mathbb{R}$, $g : \mathbb{R}^n \rightarrow \mathbb{R}^m$ und $h : \mathbb{R}^n \rightarrow \mathbb{R}^p$. Offensichtlich ist diese Form äußerst generisch, sodass es nicht schwer vorherzusagen ist, dass Algorithmen zur Lösung solcher Probleme ohne zusätzliche Annahmen zu f , g oder h schwer zu formulieren sein können und Entartungen wie lokalen Minima und dem Fehlen von solchen unterliegen Konvergenz. Tatsächlich kodiert diese Optimierung andere Probleme, die wir bereits berücksichtigt haben; Wenn wir $f(x) \leq 0$ annehmen, wird diese eingeschränkte Optimierung zur Wurzelfindung auf g , während sie sich, wenn wir $g(x) = h(x) \leq 0$ annehmen, auf eine uneingeschränkte Optimierung auf f reduziert.

Trotz dieser etwas düsteren Aussichten können Optimierungen für allgemeine eingeschränkte Fälle wertvoll sein, wenn f , g und h keine nützliche Struktur haben oder zu spezialisiert sind, um eine spezielle Behandlung zu verdienen. Wenn f ohnehin heuristisch ist, ist es darüber hinaus wertvoll, einfach ein zulässiges x zu finden, für das $f(x) < f(x_0)$ für eine anfängliche Schätzung x_0 gilt. Eine einfache Anwendung in diesem Bereich wäre ein Wirtschaftssystem, in dem f die Kosten misst; Natürlich möchten wir die Kosten minimieren, aber wenn x_0 die aktuelle Konfiguration darstellt, ist jedes x , das f verringert, eine wertvolle Ausgabe.

9.1 Motivation

In der Praxis ist es nicht schwierig, auf eingeschränkte Optimierungsprobleme zu stoßen. Tatsächlich haben wir viele Anwendungen dieser Probleme bereits aufgelistet, als wir Eigenvektoren und Eigenwerte besprochen haben, da dieses Problem so gestellt werden kann, dass kritische Punkte von $x^T A x$ unter der Bedingung $x^T x = 1$ gefunden werden; Natürlich gibt es für den besonderen Fall der Eigenwertberechnung spezielle Algorithmen, die das Problem einfacher machen.

Hier listen wir andere Optimierungen auf, die nicht die Struktur von Eigenwertproblemen aufweisen:

Beispiel 9.1 (Geometrische Projektion). Viele Flächen S im $\mathbb{R}^3$ können für ein g implizit in der Form $g(x) = 0$ geschrieben werden. Beispielsweise ergibt sich die Einheitskugel aus der Annahme $g(x) = x^T x - 1$, während dies bei einem Würfel der Fall ist

kann durch die Annahme von $g(x) = x^T y$ konstruiert werden. Tatsächlich erlauben es einige 3D-Modellierungsumgebungen Benutzern, „blobby“-Objekte, wie in Abbildung NUMBER, als Summen anzugeben

$$g(x) = c + \sum_{i=1}^n \alpha_i \left(\frac{x - x_i}{r_i} \right)^2.$$

Angenommen, wir erhalten einen Punkt $y \in \mathbb{R}^3$ und möchten den Punkt auf S finden, der y am nächsten liegt. Dieses Problem wird durch die folgende eingeschränkte Minimierung gelöst:

$$\begin{aligned} &\text{minimiere } \|x - y\|^2, \text{ so} \\ &\text{dass } g(x) = 0 \end{aligned}$$

Beispiel 9.2 (Fertigung). Angenommen, Sie haben m verschiedene Materialien; Sie haben S_i Einheiten von jedem Material auf Lager. Sie können k verschiedene Produkte herstellen; Produkt j bringt Ihnen Gewinn p_j und verwendet c_{ij} des Materials i für die Herstellung. Um den Gewinn zu maximieren, können Sie die folgende Optimierung für die Gesamtmenge x_j durchführen, die Sie von jedem Artikel j herstellen sollten:

$$\begin{aligned} &\text{maximiere } \sum_{j=1}^k p_j x_j \\ &\text{so dass } x_j \geq 0 \quad j \in \{1, \dots, k\} \\ &\quad \sum_{j=1}^k c_{ij} x_j \leq S_i \quad i \in \{1, \dots, m\} \end{aligned}$$

Die erste Einschränkung stellt sicher, dass Sie bei keinem Produkt negative Mengen herstellen, und die zweite stellt sicher, dass Sie von jedem Material nicht mehr als Ihren Lagerbestand verbrauchen.

Beispiel 9.3 (Nichtnegative kleinste Quadrate). Wir haben bereits zahlreiche Beispiele für Kleinste-Quadrate-Probleme gesehen, aber manchmal sind negative Werte im Lösungsvektor möglicherweise nicht sinnvoll. In der Computergrafik könnte ein animiertes Modell beispielsweise als sich verformende Knochenstruktur plus vernetzter „Haut“ ausgedrückt werden; Für jeden Punkt auf der Haut kann eine Liste von Gewichten berechnet werden, um den Einfluss der Positionen der Knochengelenke auf die Position der Hautsteuerelemente zu approximieren (CITE). Solche Gewichte sollten auf nichtnegative Werte beschränkt werden, um ein degeneriertes Verhalten bei der Verformung der Oberfläche zu vermeiden. In einem solchen Fall können wir das Problem der „nichtnegativen kleinsten Quadrate“ lösen:

$$\begin{aligned} &\text{minimiere } \|Ax - b\|^2, \text{ so} \\ &\text{dass } x_i \geq 0 \quad i \end{aligned}$$

Neuere Forschungen umfassen die Charakterisierung der Sparsität von nichtnegativen Lösungen der kleinsten Quadrate, die oft mehrere Werte x_i haben, die $x_i = 0$ genau erfüllen (CITE).

Beispiel 9.4 (Bündelanpassung). Nehmen wir beim Computer Vision an, dass wir ein Objekt aus mehreren Winkeln fotografieren. Eine natürliche Aufgabe besteht darin, die dreidimensionale Form des Objekts zu rekonstruieren. Dazu könnten wir auf jedem Bild eine entsprechende Menge von Punkten markieren; Insbesondere können wir $x_{ij} \in \mathbb{R}^2$ als die Position des Merkmalspunkts j auf dem Bild i annehmen. In Wirklichkeit hat jeder Merkmalspunkt eine Position $y_j \in \mathbb{R}^3$ im Raum, die wir berechnen möchten. Darüber hinaus müssen wir die Positionen der Kameras selbst finden, die wir als darstellen können

unbekannte Projektionsmatrizen P_i . Dieses als Bündelanpassung bekannte Problem kann mit einer Optimierungsstrategie angegangen werden:

$$\min_{ij} \sum_{ij} P_{ij} \tilde{y} \tilde{y}^T P_{ij} \tilde{x} \tilde{x}^T$$

so dass P_i orthogonal $\tilde{y}_i$ ist

Die Orthogonalitätsbeschränkung stellt sicher, dass die Kameratransformationen sinnvoll sind.

9.2 Theorie der eingeschränkten Optimierung

In unserer Diskussion gehen wir davon aus, dass f , g und h differenzierbar sind. Es gibt einige Methoden, die nur schwache Kontinuitäts- oder Lipschitz-Annahmen zugrunde legen, aber diese Techniken sind ziemlich spezialisiert und erfordern fortgeschrittene analytische Überlegungen.

Obwohl wir noch keine Algorithmen für die allgemeine Optimierung unter Nebenbedingungen entwickelt haben, haben wir bei der Betrachtung von Eigenwertmethoden implizit die Theorie solcher Probleme genutzt. Erinnern Sie sich insbesondere an die Methode der Lagrange-Multiplikatoren, die in Satz 0.1 eingeführt wurde. Bei dieser Technik werden kritische Punkte $f(x)$, die $g(x)$ unterliegen, als kritische Punkte der unbeschränkten La-Grange-Multiplikatorfunktion $\tilde{y}(x, \tilde{y}) = f(x) - \tilde{y}^T \cdot g(x)$ in Bezug auf beide $\tilde{y}$ charakterisiert und x gleichzeitig.

Dieser Satz ermöglichte es uns, Variationsinterpretationen von Eigenwertproblemen bereitzustellen; Allgemeiner gesagt gibt es ein alternatives (notwendiges, aber nicht ausreichendes) Kriterium dafür, dass x ein kritischer Punkt einer gleichheitsbeschränkten Optimierung ist.

Es kann jedoch eine erhebliche Herausforderung sein, einfach ein x zu finden, das die Einschränkungen erfüllt. Wir können diese Probleme trennen, indem wir einige Definitionen vornehmen:

Definition 9.1 (Zulässiger Punkt und zulässige Menge). Ein zulässiger Punkt eines eingeschränkten Optimierungsproblems ist jeder Punkt x , der $g(x) = 0$ und $h(x) \leq 0$ erfüllt. Die zulässige Menge ist die Menge aller Punkte x , die diese Einschränkungen erfüllen.

Definition 9.2 (Kritischer Punkt der eingeschränkten Optimierung). Ein kritischer Punkt einer eingeschränkten Optimierung ist einer, der die Einschränkungen erfüllt und gleichzeitig ein lokales Maximum, Minimum oder Sattelpunkt von f innerhalb der zulässigen Menge ist.

Beschränkte Optimierungen sind schwierig, weil sie gleichzeitig Wurzelfindungsprobleme (die $g(x) = 0$ -Einschränkung), Erfüllbarkeitsprobleme (die $h(x) \leq 0$ -Einschränkung) und Minimierung (die Funktion f) lösen. Abgesehen davon müssen wir, um unsere Differentialtechniken vollständig allgemeingültig zu machen, einen Weg finden, Ungleichheitsbeschränkungen zum Lagrange-Multiplikatorsystem hinzuzufügen. Angenommen, wir haben das Minimum der Optimierung gefunden, das mit x bezeichnet wird. Für jede Ungleichheitsbedingung $h_i(x) \leq 0$ haben wir zwei Möglichkeiten:

- $h_i(x) = 0$: Eine solche Einschränkung ist aktiv, was wahrscheinlich darauf hindeutet, dass sich das Optimum ändern könnte, wenn die Einschränkung entfernt würde.
- $h_i(x) < 0$: Eine solche Einschränkung ist inaktiv, was bedeutet, dass wir in einer Umgebung von x auch dann das gleiche Minimum erreicht hätten, wenn wir diese Einschränkung entfernt hätten.

Natürlich wissen wir erst nach der Berechnung, welche Einschränkungen bei x aktiv oder inaktiv sein werden.

Wenn alle unsere Einschränkungen aktiv wären, könnten wir unsere Einschränkung $h(x) \leq 0$ in eine Gleichheit ändern, ohne das Minimum zu beeinflussen. Dies könnte dazu motivieren, das folgende Lagrange-Multiplikatorsystem zu studieren:

$$\tilde{y}(x, \tilde{y}, \mu) = f(x) - \sum_i \tilde{y}_i g_i(x) - \sum_j \mu_j h_j(x)$$

Wir können jedoch nicht mehr sagen, dass x ein kritischer Punkt von $\tilde{y}$ ist, da inaktive Einschränkungen die oben genannten Terme entfernen würden. Wenn wir diese (wichtige!) Frage vorerst ignorieren, könnten wir blind vorgehen und nach kritischen Punkten dieses neuen $\tilde{y}$ in Bezug auf x fragen, die Folgendes erfüllen:

$$0 = \tilde{y} f(x) - \sum_i \tilde{y}_i g_i(x) - \sum_j \mu_j h_j(x)$$

Hier haben wir die einzelnen Komponenten von g und h herausgetrennt und als Skalarfunktionen behandelt, um eine komplexe Notation zu vermeiden.

Ein cleverer Trick kann diese Optimalitätsbedingung auf ungleichheitsbeschränkte Systeme erweitern. Beachten Sie, dass, wenn wir $\mu_j = 0$ angenommen hätten, wann immer h_j inaktiv ist, dies die irrelevanten Terme aus den Optimalitätsbedingungen entfernt. Mit anderen Worten, wir können den Lagrange-Multiplikatoren eine Einschränkung hinzufügen:

$$\mu_j h_j(x) = 0.$$

Mit dieser Einschränkung wissen wir, dass mindestens eines von μ_j und $h_j(x)$ Null sein muss, und daher gilt unsere Optimalitätsbedingung erster Ordnung immer noch!

Bisher hat unsere Konstruktion nicht zwischen der Einschränkung $h_j(x) \leq 0$ und der Einschränkung $h_j(x) = 0$ unterschieden. Wenn die Einschränkung inaktiv wäre, hätte sie weggelassen werden können, ohne das Ergebnis der Optimierung lokal zu beeinflussen, so wir Betrachten Sie den Fall, dass die Einschränkung aktiv ist. Intuitiv¹ gehen wir in diesem Fall davon aus, dass es eine Möglichkeit gibt, f durch Verletzung der Einschränkung zu verringern. Lokal gesehen ist die Richtung, in der f abnimmt, $-\tilde{y} f(x)$ und die Richtung, in der h_j abnimmt, ist $-\tilde{y} h_j(x)$. Beginnend bei x können wir f noch weiter verringern, indem wir die Einschränkung $h_j(x) \leq 0$ verletzen, wenn $-\tilde{y} f(x) \cdot -\tilde{y} h_j(x) > 0$.

Natürlich ist es schwierig, mit Produkten der Gradienten von f und h_j zu arbeiten. Bedenken Sie jedoch, dass unsere $\tilde{y} \times$ Optimalitätsbedingung erster Ordnung Folgendes besagt:

$$\tilde{y} f(x) = \sum_i \tilde{y}_i g_i(x) + \sum_{j \text{ aktiv}} \mu_j \tilde{y} h_j(x)$$

Die inaktiven μ_j -Werte sind Null und können entfernt werden. Tatsächlich können wir die Einschränkungen für $g(x) = 0$ entfernen, indem wir die Ungleichheitsbeschränkungen $g(x) \leq 0$ und $g(x) = 0$ zu h hinzufügen; Dies ist eher eine mathematische Erleichterung beim Verfassen eines Beweises als ein numerisches Manöver. Dann ergibt die Bildung von Skalarprodukten mit $\tilde{y} h_k$ für jedes feste k :

$$\sum_{j \text{ aktiv}} \tilde{y}_j \tilde{y} h_j(x) \cdot \tilde{y} h_k(x) = \tilde{y} f(x) \cdot \tilde{y} h_k(x) \leq 0$$

Die Vektorisierung dieses Ausdrucks zeigt $Dh(x) \tilde{y} Dh(x) \tilde{y} \mu$ initiale $\tilde{y} \geq 0$. Da $Dh(x) \tilde{y} Dh(x) \tilde{y}$ positiv semidef ist, dies impliziert $\mu \tilde{y} \geq 0$. Somit manifestiert sich die Beobachtung $\tilde{y} f(x) \cdot \tilde{y} h_j(x) \leq 0$ einfach durch Tatsache, dass $\mu_j \tilde{y} \geq 0$.

Unsere Beobachtungen können formalisiert werden, um eine Optimalitätsbedingung erster Ordnung für ungleichheitsbeschränkte Optimierungen zu beweisen:

¹Sie sollten unsere Diskussion nicht als formalen Beweis betrachten, da wir nicht viele Grenzfälle betrachten.

Satz 9.1 (Karush-Kuhn-Tucker (KKT)-Bedingungen). Der Vektor $x \in \mathbb{R}^n$ ist ein kritischer Punkt für die Minimierung von f , sofern $g(x) = 0$ und $h(x) \leq 0$ gilt, wenn $\tilde{y} \in \mathbb{R}^m$ und $\mu \in \mathbb{R}^p$ existieren, so dass:

$$\bullet \quad 0 = \tilde{y} f(x) - \tilde{y} \tilde{y}^T \tilde{y} g(x) - \sum_j \mu_j h_j(x) \quad (\text{„Stationarität“})$$

$$\bullet \quad g(x) = 0 \text{ und } h(x) \leq 0 \quad (\text{„primale Machbarkeit“}) \bullet \quad \mu_j h_j(x) = 0$$

$$0 \text{ für alle } j \quad (\text{„komplementäre Slackness“}) \bullet \quad \mu_j \geq 0 \text{ für alle } j \quad (\text{„duale$$

$$\text{Machbarkeit“})$$

Beachten Sie, dass sich dieser Satz auf das Lagrange-Multiplikator-kriterium reduziert, wenn h entfernt wird.

Beispiel 9.5 (Einfache Optimierung²). Angenommen, wir möchten eine Lösung finden

Maximiere xy ,

$$\text{sodass } x + y \leq 2 \text{ und } x, y \geq 0$$

$$0$$

In diesem Fall haben wir keine $\tilde{y}$ s und drei μ s. Wir nehmen $f(x, y) = xy$, $h_1(x, y) = 2 - x - y$ und $h_3(x, y) = y$. Die KKT-Bedingungen sind:

$$\text{Stationarität: } 0 = \tilde{y}y + \mu_1 \tilde{y}$$

$$\mu_2 0 = \tilde{y}x + 2\mu_1 y - \mu_3$$

$$\text{Ursprüngliche Machbarkeit: } x^2 + y^2$$

$$+ yx, y \geq 0$$

$$\text{Komplementäres Spiel: } \mu_1(2 - x - y) + \mu_2 x = 0, \mu_3 y = 0$$

$$0$$

$$\text{Duale Machbarkeit: } \mu_1, \mu_2, \mu_3 \geq 0$$

Beispiel 9.6 (Lineare Programmierung). Betrachten Sie die Optimierung:

$$\text{minimiere } b \cdot x$$

$$\text{so dass } Ax \leq c$$

Beachten Sie, dass Beispiel 9.2 auf diese Weise geschrieben werden kann. Die KKT-Bedingungen für dieses Problem sind:

$$\text{Stationarität: } A\mu = b$$

$$\text{Ursprüngliche Machbarkeit: } Ax \leq c$$

$$\text{Komplementärer Spielraum: } \mu_i(a_i \cdot x - c_i) = 0 \quad \forall i, \text{ wobei } a_i \text{ ist Zeile } i \text{ von } A$$

$$\text{Doppelte Machbarkeit: } \mu \geq 0$$

Wie im Fall der Lagrange-Multiplikatoren können wir nicht davon ausgehen, dass jedes $x \in \mathbb{R}^n$, das die KKT-Bedingungen erfüllt, f vorbehaltlich der Einschränkungen automatisch minimiert, auch nicht lokal. Eine Möglichkeit, die lokale Optimalität zu überprüfen, besteht darin, die Hesse-Funktion von f zu untersuchen, die auf den Unterraum von $\mathbb{R}^n$ beschränkt ist, in dem sich x bewegen kann, ohne die Einschränkungen zu verletzen; Wenn dieser „reduzierte“ Hessesche Wert positiv definit ist, hat die Optimierung ein lokales Minimum erreicht.

²Von <http://www.math.ubc.ca/~israel/m340/kkt2.pdf>

9.3 Optimierungsalgorithmen

Eine sorgfältige Betrachtung von Algorithmen zur eingeschränkten Optimierung liegt außerhalb des Rahmens unserer Diskussion; Glücklicherweise gibt es viele stabile Implementierungen dieser Techniken und vieles kann als „Client“ dieser Software erstellt werden, anstatt sie von Grund auf neu zu schreiben. Dennoch ist es sinnvoll, einige mögliche Ansätze zu skizzieren, um einen Einblick in die Funktionsweise dieser Bibliotheken zu gewinnen.

9.3.1 Sequentielle quadratische Programmierung (SQP)

Ähnlich wie bei BFGS und anderen Methoden, die wir in unserer Diskussion der uneingeschränkten Optimierung betrachtet haben, besteht eine typische Strategie für die eingeschränkte Optimierung darin, f , g und h mit einfacheren Funktionen zu approximieren, die angenäherte Optimierung zu lösen und zu iterieren.

Angenommen, wir haben eine Schätzung x_k der Lösung des eingeschränkten Optimierungsproblems. Wir könnten eine Taylor-Entwicklung zweiter Ordnung auf f und eine Näherung erster Ordnung auf g und h anwenden, um eine nächste Iteration wie folgt zu definieren:

$$x_{k+1} \approx x_k + \arg \min_d \frac{1}{2} d^T Hf(x_k) d + \tilde{y}^T f(x_k) \cdot d + f(x_k)$$

$$\text{so dass } g_i(x_k) + \tilde{y}_i g_i(x_k) \cdot d = 0 \quad h_i(x_k) + \tilde{y}_i h_i(x_k) \cdot d \approx 0$$

Die Optimierung zum Finden von d hat ein quadratisches Ziel mit linearen Einschränkungen, für die die Optimierung mit einer von vielen Strategien erheblich einfacher sein kann. Es ist als quadratisches Programm bekannt.

Natürlich funktioniert diese Taylor-Näherung nur in der Umgebung des optimalen Punktes. Wenn eine gute anfängliche Schätzung x_0 nicht verfügbar ist, werden diese Strategien wahrscheinlich scheitern.

Gleichheitsbeschränkungen Wenn die einzigen Beschränkungen Gleichheiten sind und h entfernt wird, weist das quadratische Programm für d die Lagrange-Multiplikator-Optimalitätsbedingungen auf, die wie folgt abgeleitet werden:

$$\frac{1}{2} \tilde{y}^T (d, \tilde{y}) \tilde{y} \approx \frac{1}{2} d^T Hf(x_k) d + \tilde{y}^T f(x_k) \cdot d + f(x_k) + \tilde{y}^T (g(x_k) + Dg(x_k)d)$$

$$\Rightarrow 0 = \tilde{y}^T d = \tilde{y}^T d = Hf(x_k) d + \tilde{y}^T f(x_k) + [Dg(x_k)]^T \tilde{y}$$

Kombiniert man dies mit der Gleichheitsbedingung, erhält man ein lineares System:

$$\begin{bmatrix} Hf(x_k) & [Dg(x_k)]^T \\ Dg(x_k) & 0 \end{bmatrix} \begin{bmatrix} d \\ \tilde{y} \end{bmatrix} = \begin{bmatrix} -\tilde{y}^T f(x_k) \\ \tilde{y}^T g(x_k) \end{bmatrix}$$

Somit kann jede Iteration der sequentiellen quadratischen Programmierung bei Vorhandensein nur von Gleichheitsbeschränkungen erreicht werden, indem dieses lineare System bei jeder Iteration gelöst wird, um $x_{k+1} \approx x_k + d$ zu erhalten. Es ist wichtig zu beachten, dass das obige lineare System nicht positiv definit ist und daher im großen Maßstab schwierig zu lösen sein kann.

Erweiterungen dieser Strategie funktionieren als BFGS und ähnliche Näherungen funktionieren für eine uneingeschränkte Optimierung, indem Näherungen des Hessischen Hf eingeführt werden. Stabilität kann auch durch die Begrenzung der in einer einzelnen Iteration zurückgelegten Distanz erreicht werden.

Ungleichungsbeschränkungen Es gibt spezielle Algorithmen zum Lösen quadratischer Programme anstelle allgemeiner nichtlinearer Programme, und diese können zum Generieren von SQP-Schritten verwendet werden. Eine bemerkenswerte Strategie besteht darin, einen „aktiven Satz“ von Einschränkungen beizubehalten, die in Bezug auf d auf dem Minimum aktiv sind; Dann können die oben genannten gleichheitsbeschränkten Methoden angewendet werden, indem inaktive Einschränkungen ignoriert werden. Iterationen der Active-Set-Optimierung aktualisieren den aktiven Constraint-Satz, indem verletzte Constraints zum aktiven Set hinzugefügt und diejenigen Ungleichheits-Constraints h_j entfernt werden, für die $\bar{y}^T f - \bar{y}^T h_j > 0$ ist, wie in unserer Diskussion der KKT-Bedingungen.

9.3.2 Barrieremethoden

Eine weitere Möglichkeit zur Minimierung bei Vorliegen von Einschränkungen besteht darin, die Einschränkungen in Energieterme umzuwandeln. Im Fall der Gleichheitsbeschränkung könnten wir beispielsweise ein „erweitertes“ Ziel wie folgt minimieren:

$$f_{\bar{y}}(x) = f(x) + \frac{\bar{y}^2}{2} g(x)$$

Beachten Sie, dass die Annahme von $\bar{y} \rightarrow 0$ dazu führt, dass $g(x)$ so klein wie möglich ist, sodass wir schließlich $g(x) \rightarrow 0$ erreichen. Daher wendet die Barrieremethode der eingeschränkten Optimierung iterative uneingeschränkte Optimierungstechniken auf $f_{\bar{y}}$ an und prüft, wie gut die Einschränkungen erfüllt sind; Wenn sie nicht innerhalb einer bestimmten Toleranz liegen, wird $\bar{y}$ erhöht und die Optimierung wird unter Verwendung der vorherigen Iteration als Ausgangspunkt fortgesetzt.

Barrieremethoden sind einfach zu implementieren und zu verwenden, können jedoch einige schädliche Fehlermodi aufweisen. Insbesondere nimmt mit zunehmendem $\bar{y}$ der Einfluss von f auf die Zielfunktion ab und die Hesse-Funktion von $f_{\bar{y}}$ wird immer schlechter konditioniert.

Barrieremethoden können auch auf Ungleichheitsbeschränkungen angewendet werden. Hier müssen wir sicherstellen, dass $h_i(x) > 0$ für alle i ; Typische Optionen für Barrierefunktionen könnten $1/h_i(x)$ (die „inverse Barriere“) oder $\bar{y} \log h_i(x)$ (die „logarithmische Barriere“) sein.

9.4 Konvexe Programmierung

Im Allgemeinen bieten Methoden wie die, die wir für die eingeschränkte Optimierung beschrieben haben, nur wenige oder gar keine Garantien für die Qualität der Ausgabe. Sicherlich sind diese Methoden nicht in der Lage, globale Minima ohne eine gute anfängliche Schätzung x_0 zu erhalten, und in bestimmten Fällen, z. B. wenn die Hesse-Methode verwendet wird

in der Nähe von x^* nicht positiv definit ist, konvergieren sie möglicherweise überhaupt nicht.

Es gibt eine bemerkenswerte Ausnahme von dieser Regel, die in vielen wichtigen Optimierungen vorkommt: die konvexe Programmierung. Die Idee dabei ist, dass die Optimierung ein eindeutiges Minimum besitzt, wenn f eine konvexe Funktion ist und die zulässige Menge selbst konvex ist. Wir haben bereits eine konvexe Funktion definiert, müssen aber verstehen, was es für eine Reihe von Einschränkungen bedeutet

konvex:

Definition 9.3 (Konvexe Menge). Eine Menge $S \subseteq \mathbb{R}^n$ ist konvex, wenn für jedes $x, y \in S$ der Punkt $tx + (1-t)y$ auch für jedes $t \in [0, 1]$ in S liegt.

Wie in Abbildung NUMMER gezeigt, ist eine Menge intuitiv konvex, wenn sich ihre Grenzform nicht sowohl nach innen als auch nach außen biegen kann.

Beispiel 9.7 (Kreise). Die Scheibe $\{x \in \mathbb{R}^n : x_2 \leq 1\}$ ist konvex, der Einheitskreis $\{x \in \mathbb{R}^n : x_2 = 1\}$ hingegen nicht.

Es ist leicht zu erkennen, dass eine konvexe Funktion ein eindeutiges Minimum hat, selbst wenn diese Funktion auf einen konvexen Bereich beschränkt ist. Insbesondere wenn die Funktion zwei lokale Minima hätte, dann darf die Punktlinie zwischen diesen Minima Werte von f ergeben, die nicht größer sind als die an den Endpunkten.

Für konvexe Optimierungen stehen starke Konvergenzgarantien zur Verfügung, die das Finden des globalen Minimums garantieren, solange f konvex ist und die Einschränkungen für g und h eine konvexe zulässige Menge ergeben. Daher ist es eine wertvolle Übung für fast jedes Optimierungsproblem, zu prüfen, ob es konvex ist, da eine solche Beobachtung das Vertrauen in die Lösungsqualität und die Erfolgsaussichten um einen großen Faktor erhöhen kann.

Ein neues Gebiet namens disziplinierte konvexe Programmierung versucht, einfache Regeln zur Konvexität zu verketteten, um konvexe Optimierungen zu generieren (CITE CVX), wodurch der Endbenutzer einfache konvexe Energieterme und Einschränkungen kombinieren kann, solange sie Kriterien erfüllen, die die endgültige Optimierung konvex machen. Zu den nützlichen Aussagen zur Konvexität in diesem Bereich gehören die folgenden:

- Der Durchschnitt konvexer Mengen ist konvex; Daher ist das Hinzufügen mehrerer konvexer Einschränkungen eine zulässige Operation.
- Die Summe konvexer Funktionen ist konvex.
- Wenn f und g konvex sind, gilt auch $h(x) \leq \max\{f(x), g(x)\}$.
- Wenn f eine konvexe Funktion ist, ist die Menge $\{x : f(x) \leq c\}$ konvex.

Tools wie die CVX-Bibliothek helfen dabei, die Implementierung verschiedener konvexer Ziele von deren Minimierung zu trennen.

Beispiel 9.8 (Konvexe Programmierung).

- Das nichtnegative Kleinste-Quadrate-Problem in Beispiel 9.3 ist konvex, weil $Ax \leq b$ konvex ist Funktion von x und die Menge $x \geq 0$ ist konvex.
- Das lineare Programmierproblem in Beispiel 9.6 ist konvex, weil es ein lineares Ziel hat und linear ist Einschränkungen.
- Wir können x_1 in ein konvexes Optimierungsziel einbeziehen, indem wir eine Variable y einführen. Dazu fügen wir die Einschränkungen $y_i \leq x_i$ und $y_i \leq -x_i$ für jedes i und ein Ziel $\sum y_i$ hinzu. Diese Summe hat Terme, die mindestens so groß sind wie $|x_i|$ und dass die Energie und die Beschränkungen konvex sind. Im Minimum müssen wir $y_i = |x_i|$ haben da wir $y_i \leq |x_i|$ eingeschränkt haben und wir möchten den Energieverbrauch minimieren. „Disziplinierte“ konvexe Bibliotheken können solche Operationen im Hintergrund durchführen, ohne dass solche Ersetzungen dem Endbenutzer offengelegt werden.

Ein besonders wichtiges Beispiel für eine konvexe Optimierung ist die lineare Programmierung aus Beispiel 9.6.

Der berühmte Simplex-Algorithmus verfolgt die aktiven Einschränkungen, löst das resultierende x mithilfe eines linearen Systems auf und prüft, ob die aktive Menge aktualisiert werden muss. Es sind keine Taylor-Näherungen erforderlich, da die objektive und zulässige Menge durch lineare Maschinen gegeben sind. Für diese Probleme sind auch interne punktlinare Programmierstrategien wie die Barrieremethode erfolgreich. Aus diesem Grund können lineare Programme in großem Maßstab gelöst werden – bis zu Millionen oder Milliarden von Variablen! – und tauchen häufig bei Problemen wie Terminplanung oder Preisgestaltung auf.

9.5 Probleme

- Simplex ableiten?
- Dualität der linearen Programmierung

Kapitel 10

Iterative lineare Löser

In den beiden vorherigen Kapiteln haben wir Strategien zur Lösung einer neuen Klasse von Problemen entwickelt, bei denen es um die Minimierung einer Funktion $f(x)$ mit oder ohne Einschränkungen für x geht. Dabei haben wir unseren Standpunkt von der numerischen linearen Algebra und insbesondere der Gaußschen Eliminierung, dass wir eine exakte Lösung für ein Gleichungssystem finden müssen, gelockert und uns stattdessen iterativen Schemata zugewandt, die das Minimum einer Funktion garantiert immer besser annähern. Iteriere immer mehr. Selbst wenn wir das Minimum nie genau finden, wissen wir, dass wir irgendwann ein x_0 mit $f(x_0) \approx 0$ mit beliebiger Ebenenqualität finden werden, abhängig von der Anzahl der von uns ausgeführten Iterationen.

Wir haben jetzt eine Wiederholung unseres Lieblingsproblems aus der numerischen linearen Algebra, nämlich die Lösung von $Ax = b$ für x , wenden jedoch einen iterativen Ansatz an, anstatt zu erwarten, eine Lösung in geschlossener Form zu finden. Diese Strategie offenbart eine neue Klasse linearer Systemlöser, die in erstaunlich wenigen Iterationen zuverlässige Näherungen für x finden können. Wir haben bereits in unserer Diskussion der linearen Algebra vorgeschlagen, wie wir dies angehen können, indem wir vorgeschlagen haben, dass Lösungen für lineare Systeme unter anderem Minima der Energie $Ax \approx b$ sind.

22,

Warum sollte man sich die Mühe machen, eine weitere Klasse linearer Systemlöser abzuleiten? Bisher erfordern die meisten unserer direkten Ansätze, dass wir A als vollständige $n \times n$ -Matrix darstellen, und Algorithmen wie LU, QR oder Cholesky-Faktorisierung erfordern alle etwa $O(n^3)$ Zeit. Bei Potenzen sind zwei Fälle zu beachten:

1. Wenn A dünn besetzt ist, neigen Methoden wie die Gaußsche Eliminierung dazu, eine Füllung zu induzieren, was bedeutet, dass sogar wenn A $O(n)$ Werte ungleich Null enthält, können Zwischenschritte der Eliminierung $O(n^2)$ Werte ungleich Null einführen. Diese Eigenschaft kann schnell dazu führen, dass lineare Algebrasysteme nicht mehr über genügend Speicher verfügen. Im Gegensatz dazu erfordern die Algorithmen in diesem Kapitel nur, dass Sie sie anwenden können A in Vektoren, was zeitlich proportional zur Anzahl der Werte ungleich Null in einer Matrix erfolgen kann.
2. Wir möchten vielleicht das $O(n^3)$ besiegen. Laufzeit von Standardtechniken zur Matrixfaktorisierung. Insbesondere wenn ein iteratives Schema in wenigen Iterationen eine ziemlich genaue Lösung für $Ax = b$ finden kann, können die Laufzeiten erheblich verkürzt werden.

Beachten Sie außerdem, dass viele der von uns besprochenen nichtlinearen Optimierungsmethoden, insbesondere diejenigen, die auf einem Newton-ähnlichen Schritt basieren, die Lösung eines linearen Systems in jeder Iteration erfordern! Daher kann die Formulierung des schnellstmöglichen Lösern einen erheblichen Unterschied machen, wenn umfangreiche Optimierungsmethoden implementiert werden, die eine oder mehrere lineare Lösungen pro Iteration erfordern. Tatsächlich könnte in diesem Fall eine ungenaue, aber schnelle Lösung eines linearen Systems akzeptabel sein, da sie ohnehin in eine größere iterative Technik einfließt.

Bitte beachten Sie, dass ein Großteil unserer Diskussion auf CITE zurückzuführen ist, auch wenn unsere Entwicklung dies sein kann angesichts der Entwicklung in den vorherigen Kapiteln etwas kürzer.

10.1 Gefälleabstieg

Wir werden unsere Diskussion auf die Lösung von $Ax = b$ konzentrieren, wobei A drei Eigenschaften hat:

1. $A \in \mathbb{R}^{n \times n}$ ist quadratisch
2. A ist symmetrisch, d. h. $A = A^T$
3. A ist positiv definit, d. h. für alle $x \neq 0$ ist $x^T A x > 0$

Gegen Ende dieses Kapitels werden wir diese Annahmen lockern. Beachten Sie zunächst, dass wir $Ax = b$ durch die Normalgleichungen $A^T A x = A^T b$ ersetzen können, um diese Kriterien zu erfüllen, obwohl diese Ersetzung, wie wir besprochen haben, zu numerischen Konditionierungsproblemen führen kann.

10.1.1 Ableitung des iterativen Schemas

In diesem Fall lässt sich leicht überprüfen, dass Lösungen von $Ax = b$ Minima der durch die quadratische Form gegebenen Funktion $f(x)$ sind

$$f(x) = \frac{1}{2} x^T A x - b^T x + c \quad \text{für jedes } c$$

in $\mathbb{R}^n$. Insbesondere zeigt die Ableitung von f

$$\nabla f(x) = Ax - b,$$

und das Setzen von $\nabla f(x) = 0$ liefert das gewünschte Ergebnis.

Anstatt $\nabla f(x) = 0$ direkt zu lösen, wie wir es in der Vergangenheit getan haben, nehmen wir an, dass wir Folgendes anwenden: Gradientenabstiegsstrategie zu dieser Minimierung. Erinnern Sie sich an den grundlegenden Gradientenabstiegsalgorithmus:

1. Berechnen Sie die Suchrichtung $d_k \in \mathbb{R}^n$ mit $\nabla f(x_k) = -A x_k + b$.
2. Definieren Sie $x_{k+1} = x_k + \gamma_k d_k$, wobei γ_k so gewählt ist, dass $f(x_k) > f(x_{k+1})$

Für eine generische Funktion f kann die Entscheidung über den Wert von γ_k ein schwieriges eindimensionales „Liniensuch“-Problem sein, das darauf hinausläuft, $f(x_k + \gamma d_k)$ als Funktion einer einzelnen Variablen $\gamma \in \mathbb{R}$ zu minimieren. Für $f(x) = \frac{1}{2} x^T A x - b^T x + c$ können wir jedoch eine Liniensuche in geschlossener Form durchführen. Insbesondere definieren

$$\begin{aligned} g(\gamma) &= f(x_k + \gamma d_k) \\ &= \frac{1}{2} (x_k + \gamma d_k)^T A (x_k + \gamma d_k) - b^T (x_k + \gamma d_k) + c \\ &= \frac{1}{2} (x_k^T A x_k + 2\gamma x_k^T A d_k + \gamma^2 d_k^T A d_k) - b^T x_k - \gamma b^T d_k + c \quad \text{aufgrund der Symmetrie von } A \\ &= \frac{1}{2} d_k^T A d_k \gamma^2 + (x_k^T A d_k - b^T d_k) \gamma + \text{const.} \\ \Rightarrow \frac{dg}{d\gamma}(\gamma) &= \gamma d_k^T A d_k + d_k^T (A x_k - b) \end{aligned}$$

Wenn wir also g bezüglich $\tilde{y}$ minimieren wollen, wählen wir einfach

$$\tilde{y} = \frac{d(b - \tilde{y} Ax) dAd}{dAd}$$

Insbesondere für den Gradientenabstieg haben wir d gewählt $k = b - \tilde{y} Ax_k$, also nimmt $\tilde{y}_k$ tatsächlich eine schöne Form an:

$$\tilde{y}_k = \frac{d_k d_k}{D_k Ad_k}$$

Am Ende ergibt unsere Formel für die Liniensuche das folgende iterative Gradientenabstiegsschema zum Lösen von $Ax = b$ im symmetrischen positiv definiten Fall:

$$d_k = b - Ax_{k-1}$$

$$\tilde{y}_k = \frac{d_k d_k}{d_k Ad_k}$$

$$x_k = x_{k-1} + \tilde{y}_k d_k$$

10.1.2 Konvergenz

Konstruktionsbedingt verringert unsere Strategie für den Gradientenabstieg $f(x_k)$, wenn $k \rightarrow \infty$. Dennoch haben wir nicht gezeigt, dass der Algorithmus den minimal möglichen Wert von f erreicht, und wir konnten nicht charakterisieren, wie viele Iterationen wir ausführen sollten, um ein angemessenes Maß an Sicherheit zu erreichen, dass $Ax_k \rightarrow b$.

Eine einfache Strategie zum Verständnis der Konvergenz des Gradientenabstiegsalgorithmus für unsere Wahl von f besteht darin, die Änderung des Rückwärtsfehlers von Iteration zu Iteration zu untersuchen.¹ Angenommen, es handelt sich um die gesuchte Lösung, nämlich $\tilde{y} \times Ax = b$. Dann können wir das Rückwärtsverhältnis untersuchen Fehler von Iteration zu Iteration:

$$R_k = \frac{f(x_k) - f(\tilde{y})}{f(x_{k-1}) - f(\tilde{y})}$$

Offensichtlich zeigt die Begrenzung von $R_k < 1$ für einige $\tilde{y}$, dass der Gradientenabfall konvergiert.

Der Einfachheit halber können wir $f(x_k)$ erweitern:

$$\begin{aligned} f(x_k) &= f(x_{k-1} + \tilde{y}_k d_k) \text{ nach unserem iterativen Schema} \\ &= \frac{1}{2} (x_{k-1} + \tilde{y}_k d_k)^T A (x_{k-1} + \tilde{y}_k d_k) - b^T (x_{k-1} + \tilde{y}_k d_k) + c \\ &= f(x_{k-1}) + \tilde{y}_k d_k^T A x_{k-1} + \tilde{y}_k^2 d_k^T A d_k - b^T \tilde{y}_k d_k \text{ nach Definition von } f \\ &= f(x_{k-1}) + \tilde{y}_k d_k^T A x_{k-1} + \tilde{y}_k^2 d_k^T A d_k - b^T \tilde{y}_k d_k \text{ seit } d_k = b - Ax_{k-1} \\ &= f(x_{k-1}) + \tilde{y}_k d_k^T A x_{k-1} + \tilde{y}_k^2 d_k^T A d_k - b^T \tilde{y}_k d_k \end{aligned}$$

¹Dieses Argument wird z. B. in <http://www-personal.umich.edu/~mepelman/teaching/IOE511/Handouts/präsentiert.511notes07-7.pdf>.

$$\begin{aligned}
&= f(x_k) - \frac{1}{\|D_k\|} \frac{D_k^T}{\|D_k\|} \frac{D_k}{\|D_k\|} f(x_k) - \frac{1}{\|D_k\|} \frac{D_k^T}{\|D_k\|} \frac{D_k}{\|D_k\|} f(x_k) \text{ nach Definition von } \tilde{y}_k \\
&= f(x_k) - \frac{1}{\|D_k\|} \frac{(D_k^T D_k)^2}{2} = f(x_k) - \frac{1}{\|D_k\|} \frac{(D_k^T D_k)^2}{2}
\end{aligned}$$

Somit können wir zu unserem Bruch zurückkehren:

$$\begin{aligned}
R_k &= \frac{f(x_k) - \frac{1}{\|D_k\|} \frac{(D_k^T D_k)^2}{2} f(x_k)}{f(x_k) - f(x_k)} \text{ nach unserer Formel für } f(x_k) \\
&= 1 - \frac{(D_k^T D_k)^2}{2 \|D_k\| f(x_k) - f(x_k)}
\end{aligned}$$

Beachten Sie, dass $Ax = b$, also können wir schreiben:

$$\begin{aligned}
f(x_k) - f(x_k) &= \frac{1}{2} x_k^T A x_k - \frac{1}{2} b^T x_k + \frac{1}{2} c \\
&= \frac{1}{2} x_k^T A x_k - \frac{1}{2} b^T x_k + \frac{1}{2} c \\
&= -\frac{1}{2} (A x_k - b)^T (A x_k - b) \text{ durch Symmetrie von } A \\
&= -\frac{1}{2} \|A x_k - b\|^2 \text{ nach Definition von } d_k
\end{aligned}$$

Daher,

$$\begin{aligned}
R_k &= 1 - \frac{(D_k^T D_k)^2}{2 \|D_k\| f(x_k) - f(x_k)} \\
&= 1 - \frac{(D_k^T D_k)^2}{2 \|D_k\| f(x_k) - f(x_k)} \text{ durch unsere neueste Vereinfachung} \\
&= 1 - \frac{D_k^T D_k}{\|D_k\|} \cdot \frac{D_k^T D_k}{\|D_k\|} \\
&= 1 - \min_{d=1} \frac{1}{\|d\|} \min_{d=1} \frac{1}{\|d\|} \text{ da dadurch der zweite Term kleiner wird} \\
&= 1 - \max_{d=1} \frac{1}{\|d\|} = \max_{d=1} \frac{1}{\|d\|} \\
&= 1 - \frac{\tilde{y}_{\min}}{\tilde{y}_{\max}} \text{ wobei } \tilde{y}_{\min} \text{ und } \tilde{y}_{\max} \text{ die minimalen und maximalen Singulärwerte von } A \text{ sind} \\
&= 1 - \frac{1}{\text{cond } A}
\end{aligned}$$

Es erfordert einiges an Algebra, aber wir haben eine wichtige Tatsache bewiesen:

Die Konvergenz des Gradientenabstiegs auf f hängt von der Konditionierung von A ab.

Das heißt, je besser A konditioniert ist, desto schneller wird der Gradientenabfall konvergieren. Da $\text{cond } A \gg 1$ ist, wissen wir außerdem, dass unsere obige Gradientenabstiegsstrategie bedingungslos gegen x konvergiert, obwohl die Konvergenz langsam sein kann, wenn A schlecht konditioniert ist.

Abbildung NUMBER veranschaulicht das Verhalten des Gradientenabstiegs für gut und schlecht konditionierte Matrizen. Wie Sie sehen, kann es beim Gradientenabstieg schwierig sein, das Minimum unserer quadratischen Funktion f zu finden, wenn die Eigenwerte von A weit gestreut sind.

10.2 Konjugierte Farbverläufe

Denken Sie daran, dass das Lösen von $Ax = b$ für $A \in \mathbb{R}^{n \times n}$ die $O(n^3)$ Zeit. Erneute Untersuchung des Gradientenabstiegs obige $O(n)$ -Strategie verwendet hat. Wir sehen, dass jede $O(n^2)$ Zeit, da wir den Matrixvektor berechnen müssen Iteration $O(n)$ Produkte zwischen A , x_k und d_k erfordert. Wenn also der Gradientenabstieg mehr als n Iterationen erfordert, müssen wir hätte genauso gut die Gaußsche Eliminierung anwenden können, die die exakte Lösung in der gleichen Zeitspanne wiederherstellt. Leider können wir nicht zeigen, dass der Gradientenabstieg eine endliche Anzahl von Iterationen erfordern muss, und in schlecht konditionierten Fällen tatsächlich auch Es kann eine große Anzahl von Iterationen erfordern, um das Minimum zu finden.

Aus diesem Grund werden wir einen Algorithmus entwerfen, der garantiert in höchstens n Schritten Unter Beibehaltung des $O(n)$ konvergiert (Worst-Case-Timing für die Lösung linearer Systeme). Unterwegs werden wir finden $O(n)$ weist dieser Algorithmus tatsächlich insgesamt bessere Konvergenzeigenschaften auf, was ihn zu einer vernünftigen Wahl macht, auch wenn wir ihn nicht vollständig ausführen.

10.2.1 Motivation

Unsere Ableitung des Algorithmus für konjugierte Gradienten basiert auf einer ziemlich einfachen Beobachtung. Angenommen, wir kennen die Lösung x auf $Ax = b$. Dann können wir unsere quadratische Form f in Ax andere Weise:

$$\begin{aligned} f(x) &= \frac{1}{2} x^T A x - b^T x + c \quad \text{per Definition 2.1} \\ &= \frac{1}{2} (x - \tilde{x})^T A (x - \tilde{x}) + \frac{1}{2} \tilde{x}^T A \tilde{x} - b^T \tilde{x} + c \\ &\quad \text{durch Addition und Subtraktion derselben Terme} \\ &= \frac{1}{2} (x - \tilde{x})^T A (x - \tilde{x}) + \frac{1}{2} \tilde{x}^T A \tilde{x} - b^T \tilde{x} + c \\ &\quad \text{da } Ax = b \text{ seit } Ax = b \text{ } \\ &= \frac{1}{2} (x - \tilde{x})^T A (x - \tilde{x}) + \text{const. da die } b\text{-Bedingungen ungültig werden} \end{aligned}$$

zu einer konstanten Verschiebung, dass f das gleiche ist wie das $\frac{1}{2} (x - \tilde{x})^T A (x - \tilde{x})$. Natürlich wissen wir bis Produkt x , aber diese Beobachtung zeigt uns die Natur von f ; es misst einfach den Abstand $\tilde{x}$ von x nach $\tilde{x}$ in Bezug auf die „A-Norm“ $\sqrt{\tilde{x}^T A \tilde{x}}$. Da A

symmetrisch und positiv definit ist, wissen wir tatsächlich, dass es mithilfe der Cholesky-Strategie als faktorisiert werden kann, auch wenn die Umsetzung in der Praxis langsam sein könnte $A = LL^T$. Mit dieser Faktorisierung nimmt f eine noch schönere Form an:

$$f(x) = \frac{1}{2} \|L(x - \tilde{x})\|_2^2 + \text{Konst.}$$

Da L eine invertierbare Matrix ist, ist diese Norm tatsächlich ein Abstandsmaß zwischen x und x .

Definieren Sie $y = Lx$ und $y = Lx$. Dann minimieren wir von diesem neuen Standpunkt aus 2 . Natürlich, $f(y) = \frac{1}{2} y^T y$ wenn wir wirklich durch Cholesky-Faktorisierung und Optimierung an diesen Punkt gelangen könnten f wäre außerordentlich einfach, aber um ein Schema für diese Minimierung ohne L abzuleiten, betrachten wir die Möglichkeit, f zu minimieren, indem wir nur die in §10.1.1 abgeleiteten Zeilensuchen verwenden.

Wir machen eine einfache Beobachtung zur Minimierung unserer vereinfachten Funktion f mithilfe einer solchen Strategie egy , dargestellt in Abbildung NUMMER:

Vorschlag 10.1. Angenommen $\{w_1, \dots, w_n\}$ sind orthogonal im $\mathbb{R}^n$. Dann wird f in höchstens n Schritten durch Liniensuche in Richtung w_1 , dann in Richtung w_2 usw. minimiert.

Nachweisen. Nehmen Sie an, dass die Spalten von $Q \in \mathbb{R}^{n \times n}$ die Vektoren w_i sind; Q ist eine orthogonale Matrix. Da $Q^T f(y) = \frac{1}{2} (Qy)^T (Qy)$ können wir $z = Qy$ schreiben; Mit anderen Worten, wir können so, dass $Qy = z$ der erste Standardbasisvektor ist, w_2 der zweite und so weiter. Wenn wir $z = Qy$ schreiben und nach der zweiten Iteration $z = Qy$ und $z = Qy$, Dann müssen wir nach der ersten Iteration eindeutig $z_1 = z$ haben, $z_2 = z$, so weiter. Nach n Schritten erreichen wir $z_n = z$, was das gewünschte Ergebnis liefert. $\square$

Die Optimierung von f kann also immer durch n -zeilige Suchvorgänge erreicht werden, solange diese Suchvorgänge in orthogonalen Richtungen erfolgen.

Um von f nach f zu gelangen, haben wir lediglich die Koordinaten mit L gedreht. So eine lineare Transformation mit geraden Linien zu geraden Linien, also führt eine Liniensuche zu einer f entlang eines Vektors w ist gleichbedeutend Liniensuche entlang $(L^{-1})^T w$ auf unserer ursprünglichen quadratischen Funktion f . Wenn wir umgekehrt n Liniensuchen auf f in Richtungen v_i durchführen, so dass $L v_i = w_i$ orthogonal sind, dann müssen wir nach Proposition 10.1 x gefunden haben. Beachten Sie, dass die Frage $w_i \cdot w_j = 0$ dasselbe ist wie die Frage j

$$0 = w_i \cdot w_j = (L v_i)^T (L v_j) = v_i^T (L^T L) v_j = v_i^T A v_j.$$

Wir haben gerade eine wichtige Folgerung zu Proposition 10.1 dargelegt. Definieren Sie konjugierte Vektoren wie folgt:

Definition 10.1 (A-konjugierte Vektoren). Zwei Vektoren v, w sind A-konjugiert, wenn $v^T A w = 0$.

Basierend auf unserer Diskussion haben wir dann gezeigt:

Vorschlag 10.2. Angenommen, $\{v_1, \dots, v_n\}$ sind A-konjugiert. Dann wird f in höchstens n Schritten durch Liniensuche in Richtung v_1 , dann in Richtung v_2 usw. minimiert.

Auf einer hohen Ebene wendet der Algorithmus für konjugierte Gradienten diesen Vorschlag einfach an, indem er entlang A-konjugierter Richtungen generiert und sucht, anstatt sich entlang ∇f zu bewegen. Beachten Sie, dass dieses Ergebnis möglicherweise etwas kontraintuitiv erscheint: Wir bewegen uns nicht unbedingt entlang der steilsten Abstiegsrichtung, sondern verlangen vielmehr, dass unsere Suchrichtungen ein globales Kriterium erfüllen, um sicherzustellen, dass wir die Arbeit nicht wiederholen. Dieser Aufbau garantiert Konvergenz in einer endlichen Anzahl von Iterationen und berücksichtigt die oben diskutierte Struktur von f in Bezug auf f .

Denken Sie daran, dass wir A-konjugierte Richtungen motiviert haben, indem wir festgestellt haben, dass sie orthogonal sind, nachdem L aus der Faktorisierung $A = LL^T$ angewendet wurde. Von diesem Standpunkt aus haben wir es mit zwei Skalarprodukten zu tun: $x_i \cdot x_j$ und $y_i \cdot y_j = (L x_i)^T (L x_j) = x_i^T L^T L x_j = x_i^T A x_j$. Diese beiden Produkte werden in gleicher Menge in unsere nachfolgende Diskussion einfließen, daher bezeichnen wir das „A-innere Produkt“ als

$$u, v_A = (L u)^T (L v) = u^T A v.$$

10.2.2 Suboptimalität des Gradientenabstiegs

Bisher wissen wir, dass wir $Ax = b$ in n Schritten über Liniensuchen entlang dieser Richtungen lösen können, wenn wir n A -konjugierte Suchrichtungen finden können. Was bleibt, ist, eine Strategie zu finden, um diese Richtungen möglichst effizient zu finden. Dazu werden wir eine weitere Eigenschaft des Gradientenabstiegsalgorithmus untersuchen, die zu einem verfeinerten Ansatz inspirieren wird.

Angenommen, wir befinden uns während einer iterativen Liniensuchmethode auf $f(x)$ bei x_k ; Wir nennen die Richtung des steilsten Abfalls von f bei x_k das Residuum $r_k \tilde{=} b \tilde{-} Ax_k$. Wir entscheiden uns möglicherweise nicht für eine Liniensuche entlang r_k wie beim Gradientenabstieg, da die Gradientenrichtungen nicht unbedingt A -konjugiert sind. Wenn wir also etwas verallgemeinern, werden wir x_{k+1} über eine Liniensuche entlang einer noch unbestimmten Richtung v_{k+1} finden.

Aus unserer Ableitung des Gradientenabstiegs sollten wir $x_{k+1} = x_k + \tilde{\gamma}_{k+1} v_{k+1}$ wählen, wobei $\tilde{\gamma}_{k+1}$ gegeben ist durch

$$\tilde{\gamma}_{k+1} = \frac{v_{k+1}^T r_k}{v_{k+1}^T A v_{k+1}}.$$

Wenn wir diese Erweiterung von x_{k+1} anwenden, können wir eine alternative Aktualisierungsformel für das Residuum schreiben:

$$\begin{aligned} r_{k+1} &= b - Ax_{k+1} = b - A(x_k + \tilde{\gamma}_{k+1} v_{k+1}) \text{ per Definition von } x_{k+1} = (b - Ax_k) - \tilde{\gamma}_{k+1} A v_{k+1} \\ &= r_k - \tilde{\gamma}_{k+1} A v_{k+1} \text{ per Definition von } r_k \end{aligned}$$

Diese Formel gilt unabhängig von unserer Wahl von v_{k+1} und kann auf jede iterative Liniensuchmethode angewendet werden.

Im Fall des Gradientenabstiegs haben wir jedoch $v_{k+1} \tilde{=} r_k$ gewählt. Diese Auswahl ergibt eine Wiederholungsbeziehung:

$$r_{k+1} = r_k - \tilde{\gamma}_{k+1} A r_k.$$

Diese einfache Formel führt zu einer lehrreichen Aussage:

Vorschlag 10.3. Wenn ein Gradientenabstieg auf f durchgeführt wird, $\text{span}\{r_0, \dots, r_k\} = \text{span}\{r_0, A r_0, \dots, A^k r_0\}$. Ein $k r_0$.

Nachweisen. Diese Aussage folgt induktiv aus unserer obigen Formel für r_{k+1} . □

Die Struktur, die wir entdecken, ähnelt stark den in Kapitel 5 erwähnten Krylov-Unterraummethoden: Das ist kein Fehler!

Der Gradientenabstieg erreicht x_k , indem man sich entlang r_0 , dann r_1 usw. über r_k bewegt. Somit wissen wir am Ende, dass die Iteration x_k des Gradientenabstiegs auf f irgendwo in der Ebene $x_0 + A^k \tilde{\gamma}_1 r_0$ liegt, nach Satz 10.3. Leider ist es span die Iteration x_k in diesem Unterraum optimal ist, wenn wir $\{r_0, r_1, \dots, r_k\} = x_0 + \text{span}\{r_0, A r_0, \dots, A^k r_0\}$. Es stimmt nicht, dass einen Gradientenabstieg ausführen. Mit anderen Worten: Im Allgemeinen kann es so sein

$$\tilde{\gamma} x_0 = v \tilde{\gamma} \text{span} \min_{\{r_0, A r_0, \dots, A^k r_0\}} x_k \quad f(x_0 + v) \arg$$

Im Idealfall würde die Umwandlung dieser Ungleichung in eine Gleichheit sicherstellen, dass die Generierung von x_{k+1} aus x_k keine Arbeit „zunichthemacht“, die während der Iterationen 1 bis $k \tilde{-} 1$ geleistet wurde.

Wenn wir unseren Beweis von Proposition 10.1 unter Berücksichtigung dieser Tatsache erneut prüfen, können wir eine Beobachtung machen, die darauf hindeutet, wie wir Konjugation nutzen könnten, um den Gradientenabstieg zu verbessern. Insbesondere wenn z zu z wechselt, ändert sich der Wert in einer zukünftigen Iteration nie mehr. Mit anderen Worten, nach der Drehung von z zu i , x gilt folgender Satz:

Vorschlag 10.4. Nehmen Sie x_k als die k -te Iteration des Prozesses aus Proposition 10.1 nach der Suche entlang v_k an. Dann,

$$x_k \approx x_0 = \arg \min_{v \in \text{span}\{v_1, \dots, v_k\}} f(x_0 + v)$$

In der besten aller möglichen Welten könnten wir also bei dem Versuch, den Gradientenabstieg zu übertreffen, hoffen, $k \geq 1$ A -konjugierte Richtungen $\{v_1, \dots, v_n\}$, so dass die Spanne $\{v_1, \dots, v_k\} = \text{span}\{r_0, Ar_0, \dots, A^{k-1}r_0\}$ für jedes k ; dass A Dann ist gewährleistet, dass unser iteratives Schema während einer bestimmten Iteration nicht schlechter abschneidet als der Gradientenabstieg. Aber wir möchten dies unbedingt tun, ohne Orthogonalisierung oder mehr als eine endliche Anzahl von Vektoren gleichzeitig zu speichern.

10.2.3 A-konjugierte Richtungen erzeugen

Natürlich können wir jede gegebene Richtungsmenge mit einer Methode wie der Gram-Schmidt-Orthogonalisierung A -orthogonal machen. Leider ist die Orthogonalisierung von $\{r_0, Ar_0, \dots\}$ das Finden der Suchrichtungen ist kostspielig und würde erfordern, dass wir eine vollständige Liste der Richtungen führen v_k ; Diese Konstruktion würde wahrscheinlich sogar den Zeit- und Speicherbedarf einer Gaußschen Eliminierung übersteigen.

Wir werden eine letzte Beobachtung über Gram-Schmidt offenbaren, die konjugierte Gradienten durch die Erzeugung konjugierter Richtungen ohne einen teuren Orthogonalisierungsprozess handhabbar macht.

Wenn wir diese Probleme ignorieren, könnten wir eine „Methode konjugierter Richtungen“ wie folgt schreiben:

$$\text{Suchrichtung aktualisieren (schlechter Gram-Schmidt-Schritt): } v_k = A^{k-1} r_0 - \sum_{i < k} \frac{A^{k-1} r_0, v_i}{v_i, v_i} v_i$$

$$\text{Zeilensuche: } \tilde{y}_k = \frac{v_k^T A v_k}{v_k^T A v_k}$$

$$\text{Schätzung aktualisieren: } x_k = x_{k-1} + \tilde{y}_k v_k$$

$$\text{Residuum aktualisieren: } r_k = r_{k-1} - \tilde{y}_k A v_k$$

Hier berechnen wir die k -te Suchrichtung v_k einfach durch Projizieren von $v_1, \dots, v_{k-1}$ aus dem Vektor $A^{k-1}r_0$. Dieser Algorithmus hat offensichtlich die Eigenschaft $\text{span}\{v_1, \dots, v_k\} = \text{span}\{r_0, Ar_0, \dots, A^{k-1}r_0\}$ vorgeschlagen in §10.2.2, hat aber zwei Probleme:

1. Ähnlich wie bei der Potenziteration für Eigenvektoren sieht die Potenz $A^{k-1}r_0$ wahrscheinlich größtenteils wie der erste Eigenvektor von A aus, wodurch die Projektion immer schlechter konditioniert wird
2. Wir müssen v_1 behalten $\dots, v_{k-1}$ um v_k zu berechnen; Daher benötigt jede Iteration dieses Algorithmus mehr Speicher und Zeit als die letzte.

Das erste Problem können wir relativ unkompliziert beheben. Insbesondere projizieren wir derzeit die vorherigen Suchrichtungen aus $A^{k-1}r_0$ heraus, aber in Wirklichkeit können wir vorherige Richtungen aus jedem Vektor w projizieren, solange

$$w \notin \text{span}\{r_0, Ar_0, \dots, A^{k-1}r_0\} \setminus \text{span}\{r_0, Ar_0, \dots, A^{k-2}r_0\},$$

das heißt, solange w eine Komponente im neuen Teil des Raums hat.

Eine alternative Wahl von w mit dieser Eigenschaft ist das Residuum r_k . Diese Eigenschaft folgt aus der Restaktualisierung $r_k = r_{k-1} - \tilde{y}_k A v_k$; In diesem Ausdruck multiplizieren wir v_k mit A und führen so die neue Potenz von A ein, die wir benötigen. Diese Wahl ahmt auch den Gradientenabstiegsalgorithmus besser nach, der $v_k = r_{k-1}$ annimmt. Somit können wir unseren Algorithmus etwas aktualisieren:

Suchrichtung aktualisieren (schlechtes Gram-Schmidt auf Residuum): $v_k = r_{k-1} - \sum_{i < k} \frac{r_{k-1}^T v_i}{v_i^T v_i} v_i$

Zeilensuche: $\tilde{y}_k = \frac{v_k^T r_{k-1}}{v_k^T A v_k}$

Schätzung aktualisieren: $x_k = x_{k-1} + \tilde{y}_k v_k$

Residuum aktualisieren: $r_k = r_{k-1} - \tilde{y}_k A v_k$

Jetzt führen wir keine Arithmetik mit dem schlecht konditionierten Vektor $A \tilde{r}_0$ durch, haben aber immer noch das obige „Speicherproblem“.

Tatsächlich ist die überraschende Beobachtung zum obigen Orthogonalisierungsschritt, dass die meisten Terme in der Summe genau Null sind! Diese erstaunliche Beobachtung ermöglicht es, jede Iteration konjugierter Gradienten durchzuführen, ohne den Speicherverbrauch zu erhöhen. Wir halten dieses Ergebnis in einem Satz fest:

Vorschlag 10.5. In der oben genannten Methode der „konjugierten Richtung“ ist $r_k^T v_i = 0$ für alle $i < k$.

Nachweisen. Wir gehen induktiv vor. Für den Basisfall $k = 1$ gibt es nichts zu beweisen. Nehmen Sie also an, dass $k > 1$ ist und dass das Ergebnis für alle $k < k$ gilt. Durch die Restaktualisierungsformel wissen wir:

$$r_k^T v_i = r_{k-1}^T v_i - \tilde{y}_k A v_k^T v_i, \quad r_k^T v_i = r_{k-1}^T v_i - \tilde{y}_k v_k^T A v_i, \quad A v_i$$

wobei die zweite Gleichheit aus der Symmetrie von A folgt.

Nehmen wir zunächst an, dass $i < k - 1$. Dann ist der erste Term der obigen Differenz durch Induktion Null. Darüber hinaus ist $A v_i \in \text{span}\{v_1, \dots, v_{i+1}\}$, da wir also unsere Suchrichtungen A -konjugiert konstruiert haben, wissen wir, dass der zweite Term ebenfalls Null sein muss.

Um den Beweis abzuschließen, betrachten wir den Fall $i = k - 1$. Mithilfe der Restaktualisierungsformel wissen wir:

$$A v_{k-1} = \frac{1}{\tilde{y}_{k-1}} (r_{k-2} - \tilde{y}_{k-1} r_{k-1})$$

Die Vormultiplikation von r_k zeigt:

$$r_k^T A v_{k-1} = \frac{1}{\tilde{y}_{k-1}} r_k^T (r_{k-2} - \tilde{y}_{k-1} r_{k-1})$$

Die Differenz $r_{k-2} - \tilde{y}_{k-1} r_{k-1}$ lebt im Bereich $\{r_0, A r_0, \dots, A^{k-1} r_0\}$, nach der Restaktualisierungsformel. zeigt, dass x_k in diesem Unterraum optimal ist. Da $r_k = -\tilde{y}_k f(x_k)$, impliziert dies $A^{k-1} r_0$, da es sonst eine Richtung gäbe, von x_k zu einer Verringerung von f zu gelangen. Insbesondere zeigt dies wie gewünscht das innere Produkt über $r_k, A v_{k-1} = 0$. □

Somit zeigt unser obiger Beweis, dass wir eine neue Richtung v_k wie folgt finden können:

$$v_k = r_{k-1} - \sum_{i=1}^{k-1} \frac{r_{k-1}, v_i}{v_i, v_i} v_i \quad \text{vi nach der Gram-Schmidt-Formel}$$

$$= r_{k-1} - \frac{r_{k-1}, v_{k-1}}{v_{k-1}, v_{k-1}} v_{k-1} \quad \text{weil die restlichen Terme } v_{k-1}, v_{k-1} \text{A verschwinden}$$

Da die Summierung über i verschwindet, hängen die Kosten für die Berechnung von v_k nicht von k ab.

10.2.4 Formulieren des Algorithmus für konjugierte Gradienten

Da wir nun über eine Strategie verfügen, die A-konjugierte Suchrichtungen mit relativ geringem Rechenaufwand liefert, wenden wir diese Strategie einfach an, um den Algorithmus für konjugierte Gradienten zu formulieren.

Nehmen wir insbesondere an, dass x_0 eine anfängliche Schätzung der Lösung von $Ax = b$ ist und dass $r_0 = b - Ax_0$ gilt.

Nehmen Sie der Einfachheit halber $v_0 = 0$. Dann aktualisieren wir x_{k-1} iterativ auf x_k , indem wir eine Reihe von Schritten für $k = 1, 2, \dots$ verwenden. . . :

$$\text{Suchrichtung aktualisieren: } v_k = r_{k-1} - \frac{r_{k-1}, v_{k-1}}{v_{k-1}, v_{k-1}} v_{k-1}$$

$$\text{Zeilensuche: } \gamma_k = \frac{v_k, r_{k-1}}{v_k, A v_k}$$

$$\text{Schätzung aktualisieren: } x_k = x_{k-1} + \gamma_k v_k$$

$$\text{Residuum aktualisieren: } r_k = r_{k-1} - \gamma_k A v_k$$

Dieses iterative Schema stellt nur eine geringfügige Anpassung des Gradientenabstiegsalgorithmus dar, weist jedoch konstruktionsbedingt viele wünschenswerte Eigenschaften auf:

- $f(x_k)$ wird nach oben durch die k -te Iteration des Gradientenabstiegs begrenzt
- Der Algorithmus konvergiert gegen x^* in n Schritten
- Bei jedem Schritt ist die Iteration x_k in dem von den ersten k Suchrichtungen aufgespannten Unterraum optimal

Um die maximale numerische Qualität aus konjugierten Gradienten herauszuholen, können wir versuchen, die Zahlen der obigen Ausdrücke zu vereinfachen. Wenn wir beispielsweise die Aktualisierung der Suchrichtung in die Formel für γ_k einbauen, können wir durch Orthogonalität schreiben:

$$\gamma_k = \frac{r_{k-1}, r_{k-1}}{v_k, A v_k}$$

Der Zähler dieses Bruchs ist nun ohne numerische Genauigkeit garantiert nicht negativ Probleme.

Ebenso können wir eine Konstante γ_k definieren, um die Aktualisierung der Suchrichtung in zwei Schritte aufzuteilen:

$$\tilde{\gamma}_k = \frac{r_{k-1}, v_{k-1}}{v_{k-1}, v_{k-1}}$$

$$v_k = r_{k-1} + \tilde{\gamma}_k v_{k-1}$$

Wir können unsere Formel für $\tilde{y}_k$ vereinfachen :

$$\begin{aligned}
 \tilde{y}_k &= \frac{r_{k\tilde{y}1}, v_{k\tilde{y}1}A}{v_{k\tilde{y}1}, v_{k\tilde{y}1}A} \\
 &= \frac{r_{k\tilde{y}1}Av_{k\tilde{y}1}}{v_{k\tilde{y}1}Av_{k\tilde{y}1}} \quad \text{nach Definition von } \tilde{y}_k, A \\
 &= \frac{r_{k\tilde{y}1} - (r_{k\tilde{y}2} \tilde{y}_{rk\tilde{y}1})}{\tilde{y}_{k\tilde{y}1}v_{k\tilde{y}1}Av_{k\tilde{y}1}} \quad \text{dark} = r_{k\tilde{y}1} - \tilde{y}_{rk\tilde{y}1}Av_k \\
 &= \frac{R_{k\tilde{y}1} r_{k\tilde{y}1}}{v_{k\tilde{y}1}Av_{k\tilde{y}1}} \quad \text{durch eine Berechnung unten } \tilde{y}_{k\tilde{y}1}v \\
 &= \frac{R_{k\tilde{y}1} r_{k\tilde{y}1}}{R_{k\tilde{y}2} r_{k\tilde{y}2}} \quad \text{nach unserer letzten Formel für } \tilde{y}_k
 \end{aligned}$$

Dieser Ausdruck zeigt, dass $\tilde{y}_k \neq 0$ ist, eine Eigenschaft, die nach Problemen mit der numerischen Genauigkeit möglicherweise nicht gilt. Wir haben unten noch eine verbleibende Berechnung: $(r_{k\tilde{y}2} \tilde{y}_{rk\tilde{y}1})$

$$\begin{aligned}
 R_{k\tilde{y}2} r_{k\tilde{y}1} &= r_{k\tilde{y}2} - \tilde{y}_{k\tilde{y}1}Av_{k\tilde{y}1} \quad \text{durch unsere Restaktualisierungsformel} \\
 &= r_{k\tilde{y}2} - \tilde{y}_{rk\tilde{y}1}v_{k\tilde{y}1} \frac{R_{k\tilde{y}2} r_{k\tilde{y}2}}{v_{k\tilde{y}1}Av_{k\tilde{y}1}} \quad \text{nach unserer Formel für } \tilde{y}_k \tilde{y}_{rk\tilde{y}1} \\
 &= r_{k\tilde{y}2} - \tilde{y}_{rk\tilde{y}1}v_{k\tilde{y}1} \frac{R_{k\tilde{y}2} r_{k\tilde{y}2}}{v_{k\tilde{y}1}Av_{k\tilde{y}1}} \quad \text{durch die Aktualisierung für vk und A-Konjugation der vk 's } \tilde{y}_{rk\tilde{y}1} \\
 &= 0, \text{ je nach Bedarf.}
 \end{aligned}$$

Mit diesen Vereinfachungen haben wir eine alternative Version konjugierter Gradienten:

$$\begin{aligned}
 \text{Suchrichtung aktualisieren: } \tilde{y}_k &= \frac{R_{k\tilde{y}1} r_{k\tilde{y}1}}{R_{k\tilde{y}2} r_{k\tilde{y}2}} \\
 v_k &= r_{k\tilde{y}1} + \tilde{y}_k v_{k\tilde{y}1} \\
 \text{Zeilensuche: } \tilde{y}_k &= \frac{R_{k\tilde{y}1} r_{k\tilde{y}1}}{v_k^T A v_k} \\
 \text{Schätzung aktualisieren: } x_k &= x_{k\tilde{y}1} + \tilde{y}_k v_k \\
 \text{Residuum aktualisieren: } r_k &= r_{k\tilde{y}1} - \tilde{y}_k A v_k
 \end{aligned}$$

Aus numerischen Gründen ist es gelegentlich ratsam, anstelle der Aktualisierungsformel für r_k die Restformel $r_k = b - A x_k$ zu verwenden. Diese Formel erfordert eine zusätzliche Matrix-Vektor-Multiplikation, repariert jedoch die numerische „Drift“, die durch das Runden mit endlicher Genauigkeit verursacht wird. Beachten Sie, dass es nicht erforderlich ist, eine lange Liste vorheriger Residuen oder Suchrichtungen zu speichern: Konjugierte Gradienten nehmen von Iteration zu Iteration eine konstante Menge an Platz ein.

10.2.5 Konvergenz- und Stoppbedingungen

Durch die Konstruktion wird garantiert, dass der Algorithmus für konjugierte Gradienten (CG) nicht langsamer konvergiert als der Gradientenabstieg auf f , während er nicht schwieriger zu implementieren ist und eine Reihe anderer aufweist

positive Eigenschaften. Eine detaillierte Diskussion der CG-Konvergenz würde den Rahmen unserer Diskussion sprengen, aber im Allgemeinen verhält sich der Algorithmus am besten bei Matrizen mit gleichmäßig verteilten Eigenwerten über einen kleinen Bereich. Eine grobe Schätzung parallel zu unserer Schätzung in §10.1.2 zeigt, dass der CG-Algorithmus Folgendes erfüllt:

$$\frac{f(x_k) - f(x^*)}{f(x_0) - f(x^*)} \leq 2 \frac{\sqrt{\kappa} - 1}{\sqrt{\kappa} + 1} \sqrt{\kappa}^k$$

wobei $\sqrt{\kappa} = \sqrt{\lambda_{\max}/\lambda_{\min}}$ cond A. Allgemeiner kann die Anzahl der Iterationen, die für den konjugierten Gradienten erforderlich sind, um einen gegebenen Fehlerwert zu erreichen, normalerweise durch eine Funktion von $\sqrt{\kappa}$ begrenzt werden, wohingegen die Grenzen für die Konvergenz des Gradientenabfalls proportional zu $\sqrt{\kappa}$ sind.

Wir wissen, dass konjugierte Gradienten garantiert genau in n Schritten gegen x^* konvergieren, aber wenn n groß ist, kann es besser sein, früher aufzuhören. Tatsächlich dividiert die Formel für $\|r_k\|$ durch Null, wenn das Residuum sehr kurz wird, was zu Problemen mit der numerischen Genauigkeit in der Nähe des Minimums von f führen kann. Daher wird in der Praxis der Schwerpunkt normalerweise dann angehalten, wenn das Verhältnis $\|r_k\|/\|r_0\|$ ausreichend klein ist.

10.3 Vorkonditionierung

Wir verfügen nun über zwei leistungsstarke iterative Schemata, um Lösungen für $Ax = b$ zu finden, wenn A symmetrisch und positiv definit ist: Gradientenabstieg und konjugierte Gradienten. Beide Strategien konvergieren bedingungslos, was bedeutet, dass wir unabhängig von der anfänglichen Schätzung x_0 mit genügend Iterationen in einer endlichen Anzahl von Iterationen genau in einer endlichen Anzahl von Iterationen $\sqrt{\kappa}$; Tatsächlich garantieren konjugierte Gradienten, dass wir x^* beliebig nahe an die wahre Lösung x herankommen. Natürlich ist die Zeit, die für beide Methoden benötigt wird, um eine Lösung von $Ax = b$ zu erreichen, direkt proportional zur Anzahl der Iterationen, die erforderlich sind, um x innerhalb einer akzeptablen Toleranz zu erreichen. Daher ist es sinnvoll, eine Strategie so weit wie möglich abzustimmen, um die Anzahl der Iterationen für die Konvergenz zu minimieren.

Zu diesem Zweck stellen wir fest, dass wir in der Lage sind, die Konvergenzraten beider Algorithmen und vieler weiterer verwandter iterativer Techniken anhand der Bedingungszahl cond A zu charakterisieren. Das heißt, je kleiner der Wert von cond A ist, desto weniger Zeit sollte dafür benötigt werden um $Ax = b$ zu lösen. Beachten Sie, dass diese Situation bei der Gaußschen Eliminierung etwas anders ist, da diese unabhängig von A die gleiche Anzahl an Schritten erfordert. Mit anderen Worten: Die Konditionierung von A beeinflusst nicht nur die Qualität der Ausgabe iterativer Methoden, sondern auch die Geschwindigkeit, mit der x $\sqrt{\kappa}$ wird angefahren.

Natürlich gilt für jede invertierbare Matrix P, dass die Lösung von $PAx = Pb$ äquivalent zur Lösung von $Ax = b$ ist. Der Trick besteht jedoch darin, dass die Konditionsnummer von PA nicht mit der von A identisch sein muss; Im (unerreichbaren) Extrem würden wir natürlich, wenn wir $P = A$ nehmen würden, Konditionierungsprobleme vollständig beseitigen! Allgemeiner angenommen, $P \approx A$ cond A, und daher kann es ratsam sein, P vor der Lösung des $\sqrt{\kappa}$. Dann erwarten wir cond PA linearen Systems anzuwenden. In diesem Fall nennen wir P einen Vorkonditionierer.

Obwohl die Idee der Vorkonditionierung attraktiv erscheint, bleiben zwei Probleme bestehen:

1. Während A symmetrisch und positiv definit sein kann, wird das Produkt PA im Allgemeinen nicht gefallen diese Eigenschaften.
2. Wir müssen $P \approx A$ finden $\sqrt{\kappa}$ das ist einfacher zu berechnen als A $\sqrt{\kappa}$ selbst.

Auf diese Probleme gehen wir in den folgenden Abschnitten ein.

10.3.1 CG mit Vorkonditionierung

Wir werden unsere Diskussion auf konjugierte Gradienten konzentrieren, da diese über bessere Konvergenzeigenschaften verfügen, obwohl sich die meisten unserer Konstruktionen auch relativ leicht auf den Gradientenabstieg anwenden lassen. Wenn wir von diesem Standpunkt aus unsere Konstruktionen in §10.2.1 betrachten, wird klar, dass unsere Konstruktion von CG stark sowohl von der Symmetrie als auch der positiven Bestimmtheit von A abhängt, sodass die Ausführung von CG auf PA normalerweise nicht sofort konvergiert.

Nehmen wir jedoch an, dass der Vorkonditionierer P selbst symmetrisch und positiv definit ist. Dies ist eine vernünftige Annahme, da A diese Eigenschaften erfüllen muss. Dann können wir wieder $a = EE^T$ schreiben. Wir machen folgende Beobachtung: Cholesky-Faktorisierung des inversen P - **Satz 10.6.**

Die Konditionszahl von PA ist die gleiche wie die von $E^T A E$.

Nachweisen. Wir zeigen, dass PA und $E^T A E$ die gleichen Singulärwerte haben; Die Bedingungszahl ist das Verhältnis des maximalen zum minimalen Singulärwert, daher ist diese Aussage mehr als ausreichend. Insbesondere ist klar, dass $E^T A E$ symmetrisch und positiv definit ist, sodass seine Eigenvektoren seine Singulärwerte sind. Nehmen wir also an, dass $E^T A E \tilde{y} = \tilde{\lambda} \tilde{y}$ ist. Wir kennen P . Wenn wir also beide Seiten unseres Eigenvektorausdrucks vorab mit $E \tilde{y}$ multiplizieren, finden wir $PAE \tilde{y} = E \tilde{\lambda} \tilde{y}$.

Die Definition von $\tilde{y} \in \mathbb{R}^n$ zeigt $PAy = \tilde{\lambda} y$. Somit haben PA und $E^T A E$ beide vollständige Eigenräume und identische Eigenwerte. $\square$

Diese Aussage impliziert, dass wir, wenn wir CG auf der symmetrischen positiv definiten Matrix $E^T A E$ durchführen, die gleichen Konditionierungsvorteile erhalten, die wir hätten, wenn wir auf PA iterieren könnten. Wie in unserem Beweis von Proposition 10.6 könnten wir unsere neue Lösung für $y = E x$ in zwei Schritten durchführen:

1. Lösen Sie $E^T A E \tilde{y} = E^T b$.
2. Lösen Sie $x = E \tilde{y}$.

Das Finden von E wäre ein wesentlicher Bestandteil dieser Strategie, ist aber wahrscheinlich schwierig, aber wir werden in Kürze beweisen, dass es unnötig ist.

Wenn wir die Berechnung von E ignorieren, könnten wir Schritt 1 mit CG wie folgt ausführen:

$$\begin{aligned} \text{Suchrichtung aktualisieren: } \tilde{y}_k &= \frac{r_k \tilde{y}_1}{r_k^T r_k} \\ v_k &= r_k \tilde{y}_1 + \tilde{y}_k v_k \\ \text{Zeilensuche: } \tilde{y}_k &= \frac{r_k \tilde{y}_1}{v_k^T E^T A E v_k} \\ \text{Schätzung aktualisieren: } y_k &= y_k \tilde{y}_1 + \tilde{y}_k v_k \\ \text{Residuum aktualisieren: } r_k &= r_k \tilde{y}_1 - \tilde{y}_k E^T A E v_k \end{aligned}$$

Dieses iterative Schema wird entsprechend der Konditionierung unserer Matrix $E^T A E$ konvergieren.

Definieren Sie $\tilde{r}_k \in \mathbb{R}^n$, $\tilde{v}_k \in \mathbb{R}^n$ und $\tilde{x}_k \in \mathbb{R}^n$. Wenn wir uns an die Beziehung $P = E \tilde{y} E^T$ erinnern, $\tilde{y}_1$, wir können schreiben wir unsere vorkonditionierte konjugierte Gradienteniteration unter Verwendung dieser neuen Variablen neu:

$$\begin{aligned} \text{Suchrichtung aktualisieren: } \tilde{y}_k &= \frac{\tilde{r}_k \tilde{y}_1}{\tilde{r}_k^T \tilde{y}_1} \\ \tilde{r}_k &= \tilde{r}_k \tilde{y}_1 - \tilde{y}_k \tilde{v}_k^T \tilde{y}_1 \end{aligned}$$

$$\tilde{v}^k = P^T \tilde{r}^{k-1} + \tilde{y}^k v^k$$

$$\text{Zeilensuche: } \tilde{y}^k = \frac{r^T_{k-1} P^T \tilde{r}^{k-1}}{\tilde{v}^k A \tilde{v}^k}$$

$$\text{Schätzung aktualisieren: } x^k = x^{k-1} + \tilde{y}^k v^k$$

$$\text{Residuum aktualisieren: } \tilde{r}^k = \tilde{r}^{k-1} - \tilde{y}^k A v^k$$

Diese Iteration hängt nicht von der Cholesky-Faktorisierung von P ab, die ausschließlich $\tilde{y}^1$ überhaupt, sondern kann unter Verwendung von P und A durchgeführt werden. Es ist leicht zu erkennen, dass das $x^k \tilde{y}^k$ -Schema $\tilde{y}^1$, tatsächlich so sein die Vorteile der Vorkonditionierung genießt, ohne dass der Vorkonditionierer faktorisiert werden muss.

Nebenbei bemerkt: Eine noch effektivere Vorkonditionierung kann durchgeführt werden, indem A durch PAQ für eine zweite Matrix Q ersetzt wird, obwohl für die Anwendung dieser zweiten Matrix zusätzliche Berechnungen erforderlich sind. Dieses Beispiel stellt einen häufigen Kompromiss dar: Wenn die Anwendung eines Vorkonditionierers selbst in einer einzelnen Iteration von CG oder einer anderen Methode zu lange dauert, lohnt sich die reduzierte Anzahl von Iterationen möglicherweise nicht.

10.3.2 Gemeinsame Vorkonditionierer

In der Praxis gute Vorkonditionierer zu finden, ist ebenso eine Kunst wie eine Wissenschaft. Das Finden des Besten hängt Näherung P von A usw. $\tilde{y}^1$ von der Struktur von A , der jeweiligen Anwendung usw. ab

Selbst grobe Näherungen können die Konvergenz jedoch erheblich verbessern, so dass CG-Anwendungen, die keinen Vorkonditionierer verwenden, selten vorkommen.

Die beste Strategie zur Formulierung von P ist oft anwendungsspezifisch, und ein interessantes technisches Näherungsproblem besteht darin, verschiedene P s für den besten Vorkonditionierer zu entwerfen und zu testen.

Nachfolgend finden Sie zwei gängige Strategien:

- Ein diagonaler (oder „Jacobi“) Vorkonditionierer nimmt einfach P als die Matrix an, die man durch Invertieren diagonalen Elemente von A erhält; das heißt, P ist die Diagonalmatrix mit den Einträgen $1/a_{ii}$. Diese Strategie kann eine ungleichmäßige Skalierung von Zeile zu Zeile verringern, die eine häufige Ursache für schlechte Konditionierung ist.
- Der spärliche approximative inverse Vorkonditionierer wird durch Lösen eines Teilproblems $\min_{P \in S} \|AP - I\|_F$ formuliert, wobei P darauf beschränkt ist, in einer Menge S von Matrizen zu sein, über die es weniger schwierig ist, ein solches Ziel zu optimieren. Eine häufige Einschränkung besteht beispielsweise darin, ein Sparsity-Muster für P vorzuschreiben, z. B. nur auf der Diagonale ungleich Nullen oder wenn A ungleich Nullen hat.
- Die unvollständigen Cholesky-Vorbedingungsfaktoren $A \tilde{y}^1 L \tilde{y}^1 L$ und nähert sich dann A an $\tilde{y}^1$ durch Lösen der entsprechenden Vorwärts- und Rücksubstitutionsprobleme. Eine beliebte Strategie besteht beispielsweise darin, die Schritte der Cholesky-Faktorisierung durchzuführen, die Ausgabe jedoch nur an den Positionen (i, j) zu speichern, an denen $a_{ij} = 0$ ist.
- Die Werte ungleich Null in A können als Diagramm betrachtet werden, und das Entfernen von Kanten im Diagramm oder das Gruppieren von Knoten kann dazu führen, dass verschiedene Komponenten getrennt werden. Das resultierende System ist nach dem Permutieren von Zeilen und Spalten blockdiagonal und kann daher mithilfe einer Folge kleinerer Lösungen gelöst werden. Eine solche Domänenzerlegungsstrategie kann für lineare Systeme effektiv sein, die aus Differentialgleichungen wie den in Kapitel NUMBER betrachteten entstehen.

Einige Vorkonditionierer verfügen über Grenzen, die Änderungen an der Konditionierung von A nach dem Ersetzen durch PA beschreiben. In den meisten Fällen handelt es sich jedoch um heuristische Strategien, die getestet und verfeinert werden sollten.

10.4 Andere iterative Schemata

Die Algorithmen, die wir in diesem Kapitel ausführlich entwickelt haben, gelten für die Lösung von $Ax = b$, wenn A quadratisch, symmetrisch und positiv definit ist. Wir haben uns auf diesen Fall konzentriert, weil er in der Praxis so häufig vorkommt, aber es gibt Fälle, in denen A asymmetrisch, unbestimmt oder sogar rechteckig ist. Es liegt außerhalb des Rahmens unserer Diskussion, in jedem Fall iterative Algorithmen abzuleiten, da viele eine spezielle Analyse oder fortgeschrittene Entwicklung erfordern. Wir fassen hier jedoch einige Techniken auf hoher Ebene zusammen (JEDEN zitieren):

- Aufteilungsmethoden zerlegen $A = M - N$ und beachten, dass $Ax = b$ äquivalent zu $Mx = Nx + b$ ist. Wenn M leicht zu invertieren ist, kann ein Festkommaschema abgeleitet werden, indem man schreibt: $Mx_k = Nx_{k-1} + b$ (CITE); Diese Techniken sind einfach zu implementieren, weisen jedoch eine Konvergenz auf, die vom Spektrum der Matrix $G = M^{-1}N$ abhängt, und können insbesondere divergieren, wenn der Spektralradius von G größer als eins ist. Eine beliebte Wahl für M ist die Diagonale von A . Methoden wie die sukzessive Überrelaxation (SOR) gewichten diese beiden Terme für eine bessere Konvergenz.
- Die Methode der konjugierten Gradientennormalgleichung (CGNR) wendet einfach den CG-Algorithmus auf die Normalengleichungen $A^T A x = A^T b$ an. Diese Methode ist einfach zu implementieren und konvergiert garantiert, solange A den vollen Rang hat. Die Konvergenz kann jedoch aufgrund der schlechten Konditionierung von $A^T A$ langsam sein, wie in Kapitel NUMBER erläutert.
- Die Methode des konjugierten Gradientennormalgleichungsfehlers (CGNE) löst auf ähnliche Weise $A A^T y = b$; Dann die Lösung von $Ax = b$ ist einfach $A^T y$.
- Methoden wie MINRES und SYMMLQ gelten für symmetrische, aber nicht unbedingt positiv definite Matrizen A , indem wir unsere quadratische Form $f(x)$ durch $g(x) = \frac{1}{2} x^T A x$ ersetzen; Diese Funktion g wird bei Lösungen für $Ax = b$ minimiert, unabhängig von der Bestimmtheit von A .
- Angesichts der schlechten Konditionierung von CGNR und CGNE minimieren die Algorithmen LSQR und LSMR auch $g(x)$ mit weniger Annahmen zu A , was insbesondere die Lösung von Systemen der kleinsten Quadrate ermöglicht.
- Verallgemeinerte Methoden einschließlich GMRES, QMR, BiCG, CGS und BiCGStab lösen $Ax = b$ mit der einzigen Einschränkung, dass A quadratisch und invertierbar ist. Sie optimieren ähnliche Energien, müssen jedoch häufig mehr Informationen über frühere Iterationen speichern und möglicherweise Zwischenmatrizen faktorisieren, um die Konvergenz mit dieser Allgemeingültigkeit zu gewährleisten.
- Schließlich kehren die Fletcher-Reeves-, Polak-Ribière- und andere Methoden zum allgemeineren Problem der Minimierung einer nichtquadratischen Funktion f zurück, indem sie konjugierte Gradientenschritte anwenden, um neue Liniensuchrichtungen zu finden. Funktionen f , die durch Quadrate gut approximiert werden, können mit diesen Strategien sehr effektiv minimiert werden, obwohl sie nicht unbedingt die Hesse-Funktion nutzen; Beispielsweise ersetzt die Fletcher-Reeves-Methode einfach das Residuum in CG-Iterationen durch den negativen Gradienten $-\nabla f$. Es ist möglich, die Konvergenz dieser Methoden zu charakterisieren, wenn sie von ausreichend effektiven Liniensuchstrategien begleitet werden.

Viele dieser Algorithmen sind fast so einfach zu implementieren wie CG oder Gradient Descent, und es gibt viele Implementierungen, die lediglich die Eingabe von A und b erfordern. Viele der oben aufgeführten Algorithmen erfordern die Anwendung von A und A , was in manchen Fällen eine technische Herausforderung darstellen kann. Als Faustregel gilt: Je allgemeiner eine Methode ist – d. h. je weniger Annahmen sie über die Struktur der Matrix A trifft –, desto mehr Iterationen ist wahrscheinlich erforderlich, um diesen Mangel an Annahmen auszugleichen. Allerdings gibt es keine festen Regeln, wenn man sich einfach nur das erfolgreichste iterative Schema anschaut, obwohl es nur begrenzte theoretische Diskussionen zum Vergleich der Vor- und Nachteile jeder dieser Methoden gibt (CITE).

10.5 Probleme

- Leiten Sie CGNR und/oder CGNE ab
- MINRES ableiten
- Leiten Sie Fletcher-Reeves ab
- Folie 13 von http://math.ntnu.edu.tw/~min/matrix_computation/Ch4_Slide4_CG_2011.pdf

Teil IV

Funktionen, Ableitungen und Integrale

Kapitel 11

Interpolation

Bisher haben wir Methoden zur Analyse von Funktionen f abgeleitet, z. B. zum Finden ihrer Minima und Wurzeln. Die Auswertung von $f(x)$ für ein bestimmtes $x \in \mathbb{R}^n$ könnte teuer sein, aber eine Grundannahme der Methoden, die wir in den vorherigen Kapiteln entwickelt haben, ist, dass wir $f(x)$ unabhängig von x erhalten können, wenn wir es wollen.

Es gibt viele Kontexte, in denen diese Annahme nicht realistisch ist. Wenn wir beispielsweise ein Foto mit einer Digitalkamera aufnehmen, erhalten wir ein $n \times m$ -Raster aus Pixelfarbwerten, das das Kontinuum des Lichts abtastet, das in ein Kameraobjektiv einfällt. Wir könnten uns ein Foto als eine kontinuierliche Funktion von der Bildposition (x, y) bis zur Farbe (r, g, b) vorstellen, aber in Wirklichkeit kennen wir den Bildwert nur an nm getrennten Orten auf der Bildebene. In ähnlicher Weise erhalten wir beim maschinellen Lernen und in der Statistik oft nur Stichproben einer Funktion an den Punkten, an denen wir Daten gesammelt haben, und wir müssen sie interpolieren, um an anderer Stelle Werte zu haben; In einem medizinischen Umfeld überwachen wir möglicherweise die Reaktion eines Patienten auf unterschiedliche Dosierungen eines Arzneimittels, können jedoch nur vorhersagen, was bei einer Dosierung passieren wird, die wir nicht explizit ausprobiert haben.

Bevor wir in diesen Fällen eine Funktion minimieren, ihre Wurzeln finden oder sogar Werte $f(x)$ an beliebigen Orten x berechnen können, benötigen wir ein Modell für die Interpolation von $f(x)$ auf den gesamten $\mathbb{R}^n$ (oder eine Teilmenge davon) bei gegebenem a Sammlung von Proben $f(x_i)$. Natürlich sind Techniken zur Lösung dieses Interpolationsproblems von Natur aus Näherungsmethoden, da wir die wahren Werte von f nicht kennen. Stattdessen streben wir danach, dass die interpolierte Funktion glatt ist und als „vernünftige“ Vorhersage der Funktionswerte dient.

In diesem Kapitel gehen wir davon aus, dass die Werte $f(x_i)$ mit völliger Sicherheit bekannt sind; In diesem Fall könnten wir uns das Problem auch so vorstellen, dass f auf den Rest der Domäne ausgedehnt wird, ohne den Wert an einer der Eingabestellen zu stören. In Kapitel NUMBER (SCHREIBEN SIE MICH IN 2014) werden wir das Regressionsproblem betrachten, bei dem der Wert $f(x_i)$ mit einer gewissen Unsicherheit bekannt ist. In diesem Fall können wir auf die Anpassung von $f(x_i)$ vollständig verzichten und stattdessen f glatter machen.

11.1 Interpolation in einer einzelnen Variablen

Bevor wir den allgemeinsten Fall betrachten, werden wir Methoden zur Interpolation von Funktionen einer einzelnen Variablen $f: \mathbb{R} \rightarrow \mathbb{R}$ entwerfen. Als Eingabe nehmen wir eine Menge von k Paaren (x_i, y_i) mit der Annahme $f(x_i) = y_i$; Unsere Aufgabe ist es, $f(x)$ für $x \in \{x_1, \dots, x_k\}$.

Unsere Strategie in diesem und anderen Abschnitten lässt sich von der linearen Algebra inspirieren, indem wir $f(x)$ in eine Basis schreiben. Das heißt, die Menge aller möglichen Funktionen $f: \mathbb{R} \rightarrow \mathbb{R}$ ist viel zu groß, um damit zu arbeiten, und enthält viele Funktionen, die in einer rechnerischen Umgebung nicht praktikabel sind. Daher vereinfachen wir den Suchraum, indem wir erzwingen, dass f als Linearkombination einfacherer Bausteine geschrieben wird

Basisfunktionen. Diese Strategie ist bereits aus der Grundrechnung bekannt: Die Taylor-Entwicklung schreibt Funktionen auf der Basis von Polynomen, während Fourier-Reihen Sinus und Cosinus verwenden.

11.1.1 Polynominterpolation

Der vielleicht einfachste Interpolant besteht darin, anzunehmen, dass $f(x)$ in $R[x]$, der Menge der Polynome, liegt. Polynome sind glatt und es ist einfach, ein Polynom vom Grad $k-1$ durch k Abtastpunkte zu finden.

Tatsächlich werden in Beispiel 3.3 bereits die Details einer solchen Interpolationstechnik erläutert. Als ein Zur Erinnerung: Angenommen, wir möchten $f(x) \approx a_0 + a_1x + a_2x^2 + \dots + a_{k-1}x^{k-1}$ finden; hier sind unsere Unbekannten die Werte $a_0, \dots, a_{k-1}$. Das Einsetzen des Ausdrucks $y_i = f(x_i)$ für jedes i zeigt, dass der Vektora das $k \times k$ Vandermonde-System erfüllt:

$$\begin{pmatrix} 1 & x_1 & x_1^2 & \dots & x_1^{k-1} \\ 1 & x_2 & x_2^2 & \dots & x_2^{k-1} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_{k-1} & x_{k-1}^2 & \dots & x_{k-1}^{k-1} \end{pmatrix} \begin{pmatrix} a_0 \\ a_1 \\ a_2 \\ \vdots \\ a_{k-1} \end{pmatrix} = \begin{pmatrix} y_0 \\ y_1 \\ \vdots \\ y_{k-1} \end{pmatrix}$$

Somit kann die Durchführung einer Grad- k -Polynominterpolation mithilfe der linearen $k \times k$ -Lösung unter Anwendung unserer generischen Strategien aus den vorherigen Kapiteln durchgeführt werden, aber tatsächlich können wir es besser machen.

Eine Möglichkeit, über unsere Form für $f(x)$ nachzudenken, besteht darin, dass sie in einer Basis geschrieben ist. Genauso wie eine Basis für R^n eine Menge von n linear unabhängigen Vektoren $v_1, \dots, v_n$, hier der Raum der Polynome für $k-1$ sein, die in der Spanne der Monome $\{1, x, x^2, \dots, x^{k-1}\}$. Es mag die offensichtlichste Basis geschrieben wird, aber unsere aktuelle Wahl hat nur wenige Eigenschaften, die sie für das Interpolationsproblem nützlich machen. Eine Möglichkeit, dieses Problem zu erkennen, besteht darin, die Folge der Funktionen $1, x, x^2, \dots, x^{k-1}$ für $x \in [0, 1]$; in diesem Intervall ist es leicht zu erkennen, dass die Funktionen x größer werden, je größer x ist, und alle ähnlich

Wenn wir weiterhin unsere Intuition aus der linearen Algebra anwenden, können wir uns dafür entscheiden, unser Polynom auf einer Basis zu schreiben, die für das vorliegende Problem besser geeignet ist. Denken Sie dieses Mal daran, dass wir k Paare (x_1, y_1) erhalten $\dots, (x_k, y_k)$. Wir werden diese (Fix-)Punkte verwenden, um die Lagrange-Interpolationsbasis $\tilde{y}_1, \dots, \tilde{y}_k$ indem man schreibt:

$$\tilde{y}_i(x) = \frac{\prod_{j \neq i} (x - x_j)}{\prod_{j \neq i} (x_i - x_j)}$$

Obwohl es nicht in der Basis $1, x, x^2, \dots, x^{k-1}$ geschrieben ist, Es ist leicht zu erkennen, dass jedes $\tilde{y}_i$ immer noch ein Poly ist x Nominal vom Grad $k-1$. Darüber hinaus hat die Lagrange-Basis die folgende wünschenswerte Eigenschaft:

$$\tilde{y}_i(x) = \begin{cases} 1 & \text{wenn } x = x_i \\ 0 & \text{sonst.} \end{cases}$$

Daher ist es auf der Lagrange-Basis einfach, das eindeutige Polynom vom Grad $k-1$ zu finden, das zu unseren (x_i, y_i) -Paaren passt:

$$f(x) \approx \sum_{i=1}^k y_i \tilde{y}_i(x)$$

Insbesondere wenn wir $x = x_j$ einsetzen, finden wir: $f(x_j)$

$$= \sum_{i=1}^k y_i \tilde{y}_i(x_j)$$

$$= y_j \text{ durch unseren obigen Ausdruck für } \tilde{y}_i(x).$$

Somit können wir in der Lagrange-Basis eine geschlossene Formel für $f(x)$ schreiben, die keine Lösung des Vandermonde-Systems erfordert. Der Nachteil besteht jedoch darin, dass jedes $\tilde{y}_i(x)$ $O(k^2)$ Zeit benötigt, um unter Verwendung der obigen Formel für ein gegebenes x ausgewertet zu werden. Für die Ermittlung von $f(x)$ $O(k^2)$ Zeit; wenn wir die Koeffizienten finden wird also explizit $O(n^2)$ aus dem Vandermonde-System benötigt, jedoch die Auswertung Die Zeit kann auf $O(n)$ reduziert werden.

Abgesehen von der Rechenzeit hat die Lagrange-Basis einen zusätzlichen numerischen Nachteil. Beachten Sie, dass der Nenner das Produkt mehrerer Terme ist. Wenn die x_i nahe beieinander liegen, enthält das Produkt möglicherweise viele Terme nahe Null, sodass wir durch eine potenziell kleine Zahl dividieren. Wie wir gesehen haben, kann dieser Vorgang zu numerischen Problemen führen, die wir vermeiden möchten.

Eine Basis für Polynome vom Grad $k \geq 1$, die versucht, einen Kompromiss zwischen der numerischen Qualität der Monome und der Effizienz der Lagrange-Basis zu finden, ist die Newton-Basis, die wie folgt definiert ist:

$$\tilde{y}_i(x) = \prod_{j=1}^{i-1} (x - x_j)$$

Wir definieren $\tilde{y}_1(x) \equiv 1$. Beachten Sie, dass $\tilde{y}_i(x)$ ein Polynom vom Grad $i-1$ ist. Durch die Definition von $\tilde{y}_i$ ist klar, dass $\tilde{y}_i(x) = 0$ für alle $x = x_j$ mit $j < i$. Wenn wir $f(x) = \sum_{i=1}^k c_i \tilde{y}_i(x)$ schreiben und diese Beobachtung expliziter formulieren wollen, finden wir:

$$\begin{aligned} f(x_1) &= c_1 \tilde{y}_1(x_1) + c_2 \tilde{y}_2(x_1) + \dots + c_k \tilde{y}_k(x_1) \\ &= c_1 \tilde{y}_1(x_2) + c_2 \tilde{y}_2(x_2) + \dots + c_k \tilde{y}_k(x_2) \\ &\vdots \\ &= c_1 \tilde{y}_1(x_k) + c_2 \tilde{y}_2(x_k) + \dots + c_k \tilde{y}_k(x_k) \end{aligned}$$

Mit anderen Worten, wir können die folgende untere Dreieckssystemkraft lösen:

$$\begin{pmatrix} \tilde{y}_1(x_1) & 0 & \dots & 0 \\ \tilde{y}_2(x_1) & \tilde{y}_2(x_2) & \dots & 0 \\ \tilde{y}_3(x_1) & \tilde{y}_3(x_2) & \tilde{y}_3(x_3) & \dots & 0 \\ \vdots & \vdots & \vdots & \dots & \vdots \\ \tilde{y}_k(x_1) & \tilde{y}_k(x_2) & \tilde{y}_k(x_3) & \dots & \tilde{y}_k(x_k) \end{pmatrix} \begin{pmatrix} c_1 \\ c_2 \\ c_3 \\ \vdots \\ c_k \end{pmatrix} = \begin{pmatrix} y_1 \\ y_2 \\ y_3 \\ \vdots \\ y_k \end{pmatrix}$$

Dieses System kann in $O(n^2)$ Zeit unter Verwendung der Vorwärtssubstitution anstelle der $O(n^3)$ Zeit, die zur Lösung des Vandermonde-Systems benötigt werden.

Wir haben jetzt drei Strategien zum Interpolieren von k Datenpunkten unter Verwendung eines Polynoms vom Grad $k-1$, indem wir es auf der Monom-, Lagrange- und Newton-Basis schreiben. Alle drei stellen unterschiedliche Kompromisse zwischen numerischer Qualität und Geschwindigkeit dar. Eine wichtige Eigenschaft ist jedoch, dass die resultierende interpolierte Funktion $f(x)$ in jedem Fall dieselbe ist. Genauer gesagt gibt es genau ein Polynom vom Grad $k-1$, das durch eine Menge von k Punkten geht. Da also alle unsere Interpolanten vom Grad $k-1$ sind, müssen sie die gleiche Ausgabe haben.

11.1.2 Alternative Basen

Obwohl sich Polynomfunktionen besonders gut für die mathematische Analyse eignen, gibt es keinen grundsätzlichen Grund, warum unsere Interpolationsbasis nicht aus verschiedenen Arten von Funktionen bestehen kann. Ein krönendes Ergebnis der Fourier-Analyse impliziert beispielsweise, dass eine große Klasse von Funktionen durch Summen der trigonometrischen Funktionen $\cos(kx)$ und $\sin(kx)$ für $k \in \mathbb{N}$ gut angenähert werden kann. Eine Konstruktion

wie das Vandermonde-System gilt in diesem Fall immer noch, und tatsächlich zeigt der Fast-Fourier-Transformationsalgorithmus (der eine ausführlichere Diskussion verdient) wie man eine solche Interpolation noch schneller durchführen kann.

Eine kleinere Erweiterung der Entwicklung in §11.1.1 betrifft rationale Funktionen der Form:

$$f(x) = \frac{p_0 + p_1x + p_2x^2 + \dots + p_mx^m}{q_0 + q_1x + q_2x^2 + \dots + q_nx^n}$$

Beachten Sie, dass wir bei gegebenen k Paaren (x_i, y_i) $m + n + 1 = k$ benötigen, damit diese Funktion wohldefiniert ist. Es muss ein zusätzlicher Freiheitsgrad festgelegt werden, um der Tatsache Rechnung zu tragen, dass dieselbe rationale Funktion durch identische Skalierung von Zähler und Nenner auf mehrere Arten ausgedrückt werden kann.

Rationale Funktionen können Asymptoten und andere Muster aufweisen, die mit nur Polynomen nicht erreichbar sind, daher können sie wünschenswerte Interpolanten für Funktionen sein, die sich schnell ändern oder Pole haben. Sobald m und n festgelegt sind, können die Koeffizienten p_i und q_i tatsächlich immer noch mithilfe linearer Techniken ermittelt werden, indem beide Seiten mit dem Nenner multipliziert werden:

$$y_i(q_0 + q_1x_i + q_2x_i^2 + \dots + q_nx_i^n) = p_0 + p_1x_i + p_2x_i^2 + \dots + p_mx_i^m$$

Auch hier sind die Unbekannten in diesem Ausdruck die p s und q s.

Die Flexibilität rationaler Funktionen kann jedoch einige Probleme verursachen. Betrachten Sie zum Beispiel das folgende Beispiel:

Beispiel 11.1 (Fehler der rationalen Interpolation, Bulirsch-Stoer §2.2). Angenommen, wir möchten $f(x)$ mit den folgenden Datenpunkten finden: $(0, 1)$, $(1, 2)$, $(2, 2)$. Wir könnten $m = n = 1$ wählen. Dann werden unsere linearen Bedingungen zu:

$$\begin{aligned} q_0 &= p_0 \\ 2(q_0 + q_1) &= p_0 + p_1 \\ 2q_1 &= p_0 + 2p_1 \end{aligned}$$

Eine nicht triviale Lösung für dieses System ist:

$$\begin{aligned} p_0 &= 0 \\ p_1 &= 2 \\ q_0 &= 0 \\ q_1 &= 1 \end{aligned}$$

Dies impliziert die folgende Form für $f(x)$:

$$f(x) = \frac{2x}{x}$$

Diese Funktion weist bei $x = 0$ eine Entartung auf, und tatsächlich führt die Streichung des x im Zähler und Nenner nicht zu $f(0) = 1$, wie wir es uns wünschen.

Dieses Beispiel veranschaulicht ein größeres Phänomen. Unser lineares System zum Finden der p - und q -Werte kann auf Probleme stoßen, wenn der resultierende Nenner $\sum p_i x^i$ eine Wurzel bei einem der festen x_i hat. Es kann gezeigt werden, dass in diesem Fall keine rationale Funktion mit der festen Wahl von m und n existiert, die die gegebenen Werte interpolieren. Eine typische Teillösung für diesen Fall wird in (CITE) dargestellt, bei der m und n abwechselnd erhöht werden, bis eine nichttriviale Lösung vorliegt. Aus praktischer Sicht ist die Spezialisierung dieser Methoden jedoch ein guter Indikator dafür, dass alternative Interpolationsstrategien vorzuziehen sein könnten, wenn die grundlegenden rationalen Methoden versagen.

11.1.3 Stückweise Interpolation

Bisher haben wir unsere Interpolationsstrategien durch die Kombination einfacher Funktionen für ganz $\mathbb{R}$ konstruiert.

Wenn die Anzahl k der Datenpunkte jedoch hoch wird, werden viele Entartungen sichtbar.

Abbildung NUMBER zeigt beispielsweise Beispiele, in denen die Anpassung von Polynomen höheren Grades an Eingabedaten zu unerwarteten Ergebnissen führen kann. Darüber hinaus zeigt Abbildung NUMBER, dass diese Strategien nichtlokal sind, was bedeutet, dass die Änderung eines einzelnen Werts y_i in den Eingabedaten das Verhalten von f für alle x ändern kann, auch für diejenigen, die weit vom entsprechenden x_i entfernt sind. Irgendwie ist diese Eigenschaft unrealistisch: Wir erwarten, dass nur die Eingabedaten in der Nähe eines bestimmten x den Wert von $f(x)$ beeinflussen, insbesondere wenn es eine große Wolke von Eingabepunkten gibt.

Wenn wir aus diesen Gründen einen Satz von Basisfunktionen $\tilde{y}_1, \dots, \tilde{y}_k$, eine wünschenswerte Eigenschaft ist nicht nur, dass sie einfach zu bearbeiten sind, sondern auch, dass sie eine kompakte Unterstützung haben:

Definition 11.1 (Kompakte Unterstützung). Eine Funktion $g(x)$ hat kompakte Unterstützung, wenn es $C \in \mathbb{R}$ gibt, so dass $g(x) = 0$ für jedes x mit $|x| > C$.

Das heißt, kompakt unterstützte Funktionen haben nur einen endlichen Bereich von Punkten, in denen sie Werte ungleich Null annehmen können.

Eine übliche Strategie zur Konstruktion interpolierender Basen mit kompakter Unterstützung besteht darin, dies stückweise zu tun. Insbesondere ein Großteil der Literatur zur Computergrafik basiert auf der Konstruktion stückweiser Polynome, die definiert werden, indem $\mathbb{R}$ in eine Reihe von Intervallen aufgeteilt und in jedes Intervall ein anderes Polynom geschrieben wird. Dazu ordnen wir unsere Datenpunkte so an, dass $x_1 < x_2 < \dots < x_k$. Dann sind zwei einfache Beispiele für stückweise Interpolanten die folgenden:

- **Stückweise Konstante** (Abbildung NUMBER): Finden Sie für ein gegebenes x den Datenpunkt x_i , der $|x - x_i|$ minimiert und definiere $f(x) = y_i$.
- **Stückweise linear** (Abbildung NUMBER): Wenn $x < x_1$, ist $f(x) = y_1$, und wenn $x > x_k$, ist $f(x) = y_k$. Andernfalls suchen Sie ein Intervall mit $x \in [x_i, x_{i+1}]$ und definieren Sie es

$$f(x) = y_{i+1} \cdot \frac{x - x_i}{x_{i+1} - x_i} + y_i \cdot \frac{x_{i+1} - x}{x_{i+1} - x_i}.$$

Allgemeiner gesagt können wir in jedem Intervall $[x_i, x_{i+1}]$ ein anderes Polynom schreiben. Beachten Sie unser bisheriges Muster: Stückweise konstante Polynome sind diskontinuierlich, während stückweise lineare Funktionen stetig sind. Es ist leicht zu erkennen, dass stückweise Quadrate C usw. sein können. Diese erhöhte Kontinuität und Differenzierbarkeit tritt auf, obwohl jedes y_i lokale Unterstützung hat; Diese Theorie wird im Detail bei der Konstruktion von „Splines“ oder Kurven, die zwischen Punkten mit gegebenen Funktionswerten und Tangenten interpolieren, ausgearbeitet.

Diese erhöhte Kontinuität hat jedoch ihre eigenen Nachteile. Mit jedem zusätzlichen Differenzierbarkeitsgrad setzen wir eine stärkere Glätteannahme für f . Diese Annahme kann unrealistisch sein: Viele physikalische Phänomene sind tatsächlich verrauscht oder diskontinuierlich, und diese erhöhte Glätte kann sich negativ auf die Interpolationsergebnisse auswirken. Ein Bereich, in dem dieser Effekt besonders deutlich wird, ist die Verwendung der Interpolation in Verbindung mit physikalischen Simulationswerkzeugen. Durch die Simulation turbulenter Flüssigkeitsströme mit überglätteten Funktionen können diskontinuierliche Phänomene wie Stoßwellen, die als Ausgabe wünschenswert sind, entfernt werden.

Abgesehen von diesen Problemen können stückweise Polynome immer noch als lineare Kombinationen von Basisfunktionen geschrieben werden. Als Basis für die stückweise konstanten Funktionen dienen beispielsweise folgende Funktionen:

$$\tilde{y}_i(x) = \begin{cases} 1 & \text{wenn } x_i \leq x < x_{i+1} \\ 0 & \text{sonst} \end{cases}$$

Diese Basis setzt einfach die Konstante 1 in die Nähe von x_i und 0 anderswo; Die stückweise konstante Interpolation einer Menge von Punkten (x_i, y_i) wird geschrieben als $f(x) = \sum y_i \tilde{y}_i(x)$. In ähnlicher Weise umfasst die in Abbildung NUMBER gezeigte sogenannte „Hut“-Basis die Menge stückweise linearer Funktionen mit scharfen Kanten an unseren Datenpunkten x_i :

$$\tilde{y}_i(x) = \begin{cases} \frac{x - x_{i-1}}{x_i - x_{i-1}} & \text{wenn } x_{i-1} < x < x_i \\ \frac{x_{i+1} - x}{x_{i+1} - x_i} & \text{wenn } x_i < x < x_{i+1} \\ 0 & \text{sonst} \end{cases}$$

Auch hier ist die stückweise lineare Interpolation der gegebenen Datenpunkte konstruktionsbedingt $f(x) = \sum y_i \tilde{y}_i(x)$.

11.1.4 Gaußsche Prozesse und Kriging

Nicht abgedeckt in CS 205A, Herbst 2013.

11.2 Multivariable Interpolation

Es gibt viele Erweiterungen der oben genannten Strategien zum Interpolieren einer Funktion mit gegebenen Datenpunkten (x_i, y_i) , bei denen $x_i \in \mathbb{R}^n$ nun mehrdimensional sein kann. Allerdings sind die Interpolationsstrategien in diesem Fall nicht ganz so klar, da es weniger offensichtlich ist, $\mathbb{R}^n$ in eine kleine Anzahl von Regionen um x_i herum zu unterteilen. Aus diesem Grund besteht ein gängiges Muster darin, mit Funktionen relativ niedriger Ordnung zu interpolieren, d. h. einfache und effiziente Interpolationsstrategien gegenüber solchen zu bevorzugen, die C^2 -Funktionen ausgeben.

Wenn alles, was uns gegeben wird, die Menge der Ein- und Ausgänge (x_i, y_i) ist, dann ist eine stückweise konstante Strategie für die Interpolation die Verwendung der Interpolation des nächsten Nachbarn. In diesem Fall nimmt $f(x)$ einfach den Wert y_i an, der x_i entspricht und x x_i minimiert; Einfache Implementierungen iterieren über alle i , um diesen Wert zu finden, obwohl Datenstrukturen wie KD-Bäume die nächsten Nachbarn schneller finden können. So wie stückweise konstante Interpolationen $\mathbb{R}^n$ in Intervalle um die Datenpunkte x_i aufteilen, teilt die Nächste-Nachbarn-Strategie $\mathbb{R}^n$ in eine Menge von Voronoi-Zellen:

Definition 11.2 (Voronoi-Zelle). Gegeben sei eine Menge von Punkten $S = \{x_1, x_2, \dots, x_k\} \subset \mathbb{R}^n$, die Voronoi-Zelle entsprechend einem bestimmten x_i ist die Menge $V_i = \{x : \|x - x_i\| < \|x - x_j\| \text{ für alle } j \neq i\}$. Das heißt, es ist die Menge der Punkte, die näher an x_i liegen als an jedem anderen x_j in S .

Abbildung NUMBER zeigt ein Beispiel dafür, dass die Voronoi-Zellen über eine Reihe von Datenpunkten in $\mathbb{R}^2$. Diese Zellen viele günstige Eigenschaften haben; Beispielsweise handelt es sich um konvexe Polygone, die um jedes x_i herum lokalisiert sind. Tatsächlich ist die Konnektivität von Voronoi-Zellen ein gut untersuchtes Problem der Computergeometrie, das zur Konstruktion der berühmten Delaunay-Triangulation führte.

Es gibt viele Optionen für die kontinuierliche Interpolation von Funktionen auf $\mathbb{R}^n$, jede mit ihren eigenen Vor- und Nachteilen. Wenn wir beispielsweise unsere obige Nächste-Nachbarn-Strategie erweitern möchten, könnten wir mehrere nächste Nachbarn von x berechnen und $f(x)$ basierend auf x interpolieren

x_i^2 für jeden nächsten Nachbarn x_i . Bestimmte „k-nächste Nachbar“-Datenstrukturen können Abfragen beschleunigen, bei denen Sie mehrere Punkte in einem Datensatz finden möchten, die einem bestimmten x am nächsten liegen.

Eine weitere Strategie, die in der Literatur zur Computergrafik häufig vorkommt, ist die baryzentrische Interpolation. Angenommen, wir haben genau $n + 1$ Abtastpunkte $(x_1, y_1), \dots, (x_{n+1}, y_{n+1})$, wobei $x_i \in \mathbb{R}^n$ und wir wie immer die y -Werte auf das gesamte $\mathbb{R}^n$ interpolieren möchten; Auf der Ebene würden uns beispielsweise drei Werte gegeben, die den Eckpunkten eines Dreiecks zugeordnet sind. Jeder Punkt $x \in \mathbb{R}^n$ kann eindeutig als lineare Kombination $x = \sum_{i=1}^{n+1} a_i x_i$ mit der zusätzlichen Einschränkung geschrieben werden, dass $\sum_{i=1}^{n+1} a_i = 1$; mit anderen Worten, wir schreiben x als gewichteten Durchschnitt der Punkte x_i . Die baryzentrische Interpolation schreibt in diesem Fall einfach $f(x) = \sum_{i=1}^{n+1} a_i(x) y_i$.

Auf der Ebene $\mathbb{R}^2$ weist die baryzentrische Interpolation eine direkte geometrische Interpretation in umlaufenden Dreiecksflächen auf, wie in Abbildung NUMBER dargestellt. Darüber hinaus lässt sich leicht überprüfen, ob die resultierende interpolierte Funktion $f(x)$ affin ist, was bedeutet, dass sie für ein $c \in \mathbb{R}$ und $d \in \mathbb{R}^n$ als $f(x) = c + d \cdot x$ geschrieben werden kann.

Im Allgemeinen ist das Gleichungssystem, das wir für eine baryzentrische Interpolation lösen möchten, bei einigen $x \in \mathbb{R}^n$ ist:

$$\sum_{i=1}^{n+1} a_i x_i = x$$

$$\sum_{i=1}^{n+1} a_i = 1$$

In Abwesenheit von Entartungen ist dieses System invertierbar, wenn es $n + 1$ Punkte x_i gibt. Wenn jedoch mehr x_i vorhanden sind, wird das System für a unterbestimmt. Das bedeutet, dass es mehrere Möglichkeiten gibt, ein gegebenes x als gewichteten Durchschnitt der x_i zu schreiben.

Eine Lösung dieses Problems besteht darin, weitere Bedingungen für den Vektor der Durchschnittsgewichte hinzuzufügen. Diese Strategie führt zu verallgemeinerten baryzentrischen Koordinaten, einem Forschungsthema in der modernen Mathematik und Technik. Typische Einschränkungen erfordern, dass es als Funktion auf $\mathbb{R}^n$ glatt und im Inneren der Menge von x_i nicht negativ ist, wenn diese Punkte ein Polygon oder Polyeder definieren. Abbildung NUMBER zeigt ein Beispiel für verallgemeinerte Schwerpunktkoordinaten, die aus Datenpunkten auf einem Polygon mit mehr als $n + 1$ Punkten berechnet wurden.

Eine alternative Lösung des unterbestimmten Problems für baryzentrische Koordinaten bezieht sich auf die Idee, stückweise Funktionen für die Interpolation zu verwenden; Der Einfachheit halber beschränken wir unsere Diskussion hier auf $x \in \mathbb{R}^2$, obwohl die Erweiterungen auf höhere Dimensionen relativ offensichtlich sind.

Oftmals erhalten wir nicht nur die Menge der Punkte x_i , sondern auch eine Zerlegung des Bereichs, um den wir uns kümmern (in diesem Fall eine Teilmenge von $\mathbb{R}^2$), in $n + 1$ -dimensionale Objekte, wobei diese Punkte als Eckpunkte verwendet werden. Abbildung NUMBER zeigt beispielsweise eine solche Tessellation eines Teils von $\mathbb{R}^2$ in Dreiecke.

Die Interpolation ist in diesem Fall unkompliziert: Das Innere jedes Dreiecks wird mithilfe baryzentrischer Koordinaten interpoliert.

Beispiel 11.2 (Schattierung). In der Computergrafik ist eine der häufigsten Darstellungen einer Form eine Reihe von Dreiecken in einem Netz. Im Schattierungsmodell pro Scheitelpunkt wird eine Farbe für jeden Scheitelpunkt auf einem Netz berechnet. Um das Bild dann auf dem Bildschirm darzustellen, werden diese Werte pro Scheitelpunkt mithilfe einer baryzentrischen Interpolation in die Innenräume der Dreiecke interpoliert. Ähnliche Strategien werden für die Texturierung und andere häufige Aufgaben verwendet. Abbildung NUMBER zeigt ein Beispiel für dieses einfache Schattierungsmodell. Nebenbei bemerkt ist das Zusammenspiel von Perspektiventransformationen und Interpolationsstrategien ein spezifisches Problem der Computergrafik. Die baryzentrische Interpolation von Farben auf einer 3D-Oberfläche und die anschließende Projektion dieser Farbe auf die Bildebene ist nicht dasselbe wie die Projektion von Dreiecken auf die Bildebene und die anschließende Interpolation von Farben in das Innere des Dreiecks. Daher müssen Algorithmen in diesem Bereich eine perspektivische Korrektur anwenden, um diesen Fehler zu berücksichtigen.

Bei einer Menge von Punkten in $\mathbb{R}^2$ ist das Problem der Triangulation alles andere als trivial, und Algorithmen für diese Art von Berechnung lassen sich oft nur schlecht auf $\mathbb{R}^n$ übertragen. Daher werden in höheren Dimensionen Next-Neighbor- oder Regressionsstrategien vorzuziehen (siehe Kapitel NUMMER).

Die baryzentrische Interpolation führt zu einer Verallgemeinerung der stückweise linearen Hutfunktionen aus §11.1.3, dargestellt in Abbildung NUMBER. Denken Sie daran, dass unsere interpolatorische Ausgabe vollständig durch die Werte y_i an den Eckpunkten der Dreiecke bestimmt wird. Tatsächlich können wir uns $f(x)$ als eine lineare Kombination $\sum y_i \tilde{y}_i(x)$ vorstellen, wobei jedes $\tilde{y}_i(x)$ die stückweise Schwerpunktfunktion ist, die man erhält, indem man auf x_i eine 1 und überall sonst eine 0 setzt, wie in Abbildung NUMBER. Diese dreieckigen Hutfunktionen bilden die Grundlage der „Finite-Elemente-Methode erster Ordnung“, die wir in zukünftigen Kapiteln untersuchen werden; Spezielle Konstruktionen, die Polynome höherer Ordnung verwenden, werden als „Elemente höherer Ordnung“ bezeichnet und können verwendet werden, um die Differenzierbarkeit entlang der Dreiecksanten zu gewährleisten.

Eine alternative und ebenso wichtige Zerlegung des Bereichs von f erfolgt, wenn die Punkte x_i auf einem regelmäßigen Gitter in $\mathbb{R}^n$ liegen. Die folgenden Beispiele veranschaulichen Situationen, in denen dies der Fall ist:

Beispiel 11.3 (Bildverarbeitung). Wie in der Einleitung erwähnt, wird ein typisches digitales Foto als $m \times n$ -Raster aus roten, grünen und blauen Farbintensitäten dargestellt. Wir können uns diese Werte so vorstellen, als ob sie auf einem Gitter in $\mathbb{Z} \times \mathbb{Z}$ leben. Angenommen, wir möchten das Bild um einen Winkel drehen, der jedoch kein Vielfaches von 90° ist. Dann müssen wir, wie in Abbildung NUMBER dargestellt, Bildwerte an möglicherweise nicht ganzzahligen Positionen nachschlagen, was die Interpolation der Farben auf $\mathbb{R} \times \mathbb{R}$ erfordert.

Beispiel 11.4 (Medizinische Bildgebung). Die typische Ausgabe eines Magnetresonanztomographiegeräts (MRT) ist ein Werteraster von einem $n \times n \times p$, das die Gewebedichte an verschiedenen Punkten darstellt. Theoretisch ist das typische Modell für diese Funktion $f: \mathbb{R}^3 \rightarrow \mathbb{R}$. Wir können die äußere Oberfläche eines bestimmten Organs extrahieren, wie in Abbildung NUMBER gezeigt, indem wir den Ebenensatz $\{x: f(x) = c\}$ für ein c ermitteln. Um diesen Ebenensatz zu finden, müssen wir f auf das gesamte Voxelgitter erweitern, um genau zu finden, wo es c kreuzt.

Gitterbasierte Interpolationsstrategien wenden typischerweise die eindimensionalen Formeln aus §11.1.3 Dimension für Dimension an. Beispielsweise interpolieren bilineare Interpolationsschemata in $\mathbb{R}^2$ jeweils eine Dimension linear, um den Ausgabewert zu erhalten:

Beispiel 11.5 (Bilineare Interpolation). Angenommen, f nimmt die folgenden Werte an:

- $f(0, 0) = 1$
- $f(0, 1) = \frac{3}{4}$
- $f(1, 0) = 5$
- $f(1, 1) = \frac{11}{4}$

und dass dazwischen f durch bilineare Interpolation erhalten wird. Um $f(\frac{1}{4}, \frac{1}{2})$ zu finden, interpolieren wir zunächst in x_1 , um Folgendes zu finden:

$$F = \frac{1}{4}, \quad \frac{3}{4} f(0, 0) + \frac{1}{4} f(1, 0) = \frac{3}{4} \cdot 1 + \frac{1}{4} \cdot 5 = \frac{8}{4} = 2$$

$$F = \frac{1}{4}, \quad \frac{3}{4} f(0, 1) + \frac{1}{4} f(1, 1) = \frac{3}{4} \cdot \frac{3}{4} + \frac{1}{4} \cdot \frac{11}{4} = \frac{9}{16} + \frac{11}{16} = \frac{20}{16} = \frac{5}{4}$$

Als nächstes interpolieren wir in x_2 :

$$F\left(\frac{1}{4}, \frac{1}{2}\right) = \frac{1}{2} f\left(\frac{1}{4}, 0\right) + \frac{1}{2} f\left(\frac{1}{4}, \frac{3}{4}\right) = \frac{1}{2} - \frac{1}{4} = \frac{1}{4}$$

Eine wichtige Eigenschaft der bilinearen Interpolation besteht darin, dass wir die gleiche Ausgabe erhalten, wenn wir zuerst in x_2 und dann in x_1 interpolieren.

Methoden höherer Ordnung wie die bikubische Interpolation und die Lanczos-Interpolation verwenden wiederum mehr Polynomterme, sind jedoch langsamer zu berechnen. Insbesondere bei der Interpolation von Bildern erfordern bikubische Strategien mehr Datenpunkte als das Quadrat der Funktionswerte, die einem Punkt x am nächsten liegen; Dieser zusätzliche Aufwand kann grafische Tools verlangsamen, für die jede Suche im Speicher zusätzliche Rechenzeit verursacht.

11.3 Theorie der Interpolation

Bisher war unsere Behandlung der Interpolation ziemlich heuristisch. Während wir uns größtenteils auf unsere Intuition verlassen, was eine „vernünftige“ Interpolation für eine Reihe von Funktionswerten eine akzeptable Strategie ist, können bei verschiedenen Interpolationsmethoden subtile Probleme auftreten, die es zu berücksichtigen gilt.

11.3.1 Lineare Algebra von Funktionen

Wir begannen unsere Diskussion damit, verschiedene Interpolationsstrategien als unterschiedliche Grundlagen für die Menge der Funktionen $f: \mathbf{R} \rightarrow \mathbf{R}$ vorzustellen. Diese Analogie zu Vektorräumen erstreckt sich auf eine vollständige geometrische Funktionstheorie und im Wesentlichen auf frühe Arbeiten auf dem Gebiet der Funktionalanalyse erweitert die Geometrie von $\mathbf{R}^n$ auf Mengen von Funktionen. Hier werden wir Funktionen einer Variablen diskutieren, obwohl viele Aspekte der Erweiterung auf allgemeinere Funktionen einfach durchzuführen sind.

So wie wir für Funktionen die Begriffe Spanne und Linearkombination definieren können, können wir für feste $a, b \in \mathbf{R}$ ein inneres Produkt der Funktionen $f(x)$ und $g(x)$ wie folgt definieren:

$$(f, g) = \int_a^b f(x)g(x) dx.$$

So wie das A -innere Produkt von Vektoren uns bei der Ableitung des Algorithmus für konjugierte Gradienten half und viel mit dem Skalarprodukt gemeinsam hatte, kann das funktionale innere Produkt verwendet werden, um Methoden der linearen Algebra für den Umgang mit Funktionsräumen und das Verständnis ihrer Spanne zu definieren. Wir definieren eine Norm einer Funktion auch als $\|f\| = \sqrt{(f, f)}$.

Beispiel 11.6. Funktion Inneres Produkt Nehmen wir $p_n(x) = x^n$ das n -te Monom sein. Dann gilt für $a = 0$ und $b = 1$, wir haben:

$$\begin{aligned} (p_n, p_m) &= \int_0^1 x^n \cdot x^m dx \\ &= \int_0^1 x^{n+m} dx \\ &= \frac{1}{n+m+1} \end{aligned}$$

Beachten Sie, dass dies Folgendes zeigt:

$$\begin{aligned} \frac{p_n}{p_n}, \frac{Uhr}{Uhr} &= \frac{p_n}{p_m} \\ &= \frac{p_n p_m (2n+1)(2m+1)}{n+m+1} \end{aligned}$$

Dieser Wert beträgt ungefähr 1, wenn $n \approx m$, aber $n = m$, was unsere frühere Behauptung untermauert, dass sich die Monome auf $[0, 1]$ erheblich „überlappen“.

Angesichts dieses inneren Produkts können wir den Gram-Schmidt-Algorithmus anwenden, um eine Orthonormalbasis für die Menge der Polynome zu finden. Wenn wir $a = -1$ und $b = 1$ nehmen, erhalten wir die Legendre-Polynome, dargestellt in Abbildung NUMBER:

$$\begin{aligned} P_0(x) &= 1 \\ P_1(x) &= x \\ P_2(x) &= \frac{1}{2}(3x^2 - 1) \\ P_3(x) &= \frac{1}{2}(5x^3 - 3x) \\ P_4(x) &= \frac{1}{8}(35x^4 - 30x^2 + 3) \\ &\vdots \end{aligned}$$

Diese Polynome haben aufgrund ihrer Orthogonalität viele nützliche Eigenschaften. Angenommen, wir möchten $f(x)$ mit einer Summe $\sum a_i P_i(x)$ approximieren. Wenn wir $f \approx \sum a_i P_i$ in der Funktionsnorm minimieren wollen, ist dies ein Problem der kleinsten Quadrate! Aufgrund der Orthogonalität der Legendre-Basis für $R[x]$ zeigt eine einfache Erweiterung unserer Projektionsmethoden:

$$a_i = \frac{f, P_i}{P_i, P_i}$$

Somit kann die Approximation von f mithilfe von Polynomen einfach durch Integration von f gegen die Mitglieder der Legendre-Basis erreicht werden; Im nächsten Kapitel werden wir erfahren, wie dieses Integral näherungsweise ausgeführt werden könnte.

Bei einer gegebenen positiven Funktion $w(x)$ können wir durch Schreiben ein allgemeineres inneres Produkt $\cdot_w$ definieren

$$f, g_w = \int_A^B w(x) f(x) g(x) dx.$$

Wenn wir $w(x) = \frac{1}{\sqrt{1-x^2}}$ mit $a = -1$ und $b = 1$, dann ergibt die Anwendung von Gram-Schmidt den Tschebyscheff Polynome nehmen:

$$\begin{aligned} T_0(x) &= 1 \\ T_1(x) &= x \\ T_2(x) &= 2x^2 - 1 \\ T_3(x) &= 4x^3 - 3x \\ T_4(x) &= 8x^4 - 8x^2 + 1 \\ &\vdots \end{aligned}$$

Tatsächlich gilt für diese Polynome eine überraschende Identität:

$$T_k(x) = \cos(k \arccos(x)).$$

Diese Formel kann überprüft werden, indem man sie explizit auf T_0 und T_1 überprüft und dann die Beobachtung induktiv anwendet:

$$\begin{aligned} T_{k+1}(x) &= \cos((k+1) \arccos(x)) \\ &= 2x \cos(k \arccos(x)) - \cos((k-1) \arccos(x)) \text{ durch die} \\ &\quad \text{Identität } \cos((k+1)\tilde{y}) = 2 \cos(k\tilde{y}) \cos(\tilde{y}) - \cos((k-1)\tilde{y}) \\ &= 2xT_k(x) - T_{k-1}(x) \end{aligned}$$

Diese „Drei-Term-Rekurrenz“-Formel bietet auch eine einfache Möglichkeit, die Tschebyscheff-Polynome zu generieren.

Wie in Abbildung NUMBER dargestellt, ist dank der trigonometrischen Formel für die Tschebyscheff-Polynome leicht zu erkennen, dass die Minima und Maxima von T_k zwischen $+1$ und -1 schwanken. Darüber hinaus liegen diese Extrema bei $\cos(i\tilde{y}/k)$ (die sogenannten „Tschebyscheff-Punkte“) für i von 0 bis k ; Diese schöne Verteilung der Extrema vermeidet Oszillationsphänomene wie die in Abbildung NUMBER gezeigten, wenn eine endliche Anzahl von Polynomtermen zur Approximation einer Funktion verwendet wird. Tatsächlich empfehlen technischere Behandlungen der Polynominterpolation, x_i für die Interpolation in der Nähe von Tschebyscheff-Punkten zu platzieren, um eine gleichmäßige Ausgabe zu erhalten.

11.3.2 Approximation über stückweise Polynome

Angenommen, wir möchten eine Funktion $f(x)$ mit einem Polynom vom Grad n auf einem Intervall $[a, b]$ approximieren. Definieren Sie $\tilde{y}_x$ als den Abstand $b - x$. Ein Maß für den Fehler einer Näherung ist die Funktion von $\tilde{y}_x$, die verschwinden sollte, wenn $\tilde{y}_x \rightarrow 0$. Wenn wir dann f mit stückweisen Polynomen approximieren, sagt uns diese Art der Analyse, wie weit wir die Polynome voneinander entfernen sollten, um dies zu erreichen ein gewünschtes Maß an Annäherung.

Angenommen, wir approximieren f mit einer konstanten $c = f(\frac{a+b}{2})$, wie in stückweise konstant. Wenn wir $|f(x) - c| < M$ für alle $x \in [a, b]$ annehmen, haben wir:

$$\max_{x \in [a, b]} |f(x) - c| \leq M \quad \text{M nach dem Mittelwertsatz } x \in [a, b]$$

Daher erwarten wir einen $O(\tilde{y}_x)$ -Fehler, wenn wir die stückweise konstante Interpolation verwenden.

Angenommen, wir approximieren stattdessen f durch stückweise lineare Interpolation, also durch Nehmen

$$f(x) = f(a) + \frac{b-x}{b-a} (f(b) - f(a)).$$

Durch den Mittelwertsatz wissen wir $f(x) = f(\tilde{y})$ für ein $\tilde{y} \in [a, b]$. Das Schreiben der Taylor-Entwicklung über $\tilde{y}$ zeigt $f(x) = f(\tilde{y}) + f'(\tilde{y})(x - \tilde{y}) + O((x - \tilde{y})^2)$ auf $[a, b]$, während wir unsere lineare $f(x) = f(\tilde{y}) + f'(\tilde{y})(x - \tilde{y})$ umschreiben können. Die Subtraktion dieser beiden Ausdrücke zeigt also, dass

Näherungsfehler von f auf $O(\tilde{y}_x)$ mit einem Polynom vom Grad n abnimmt, ist $O(\tilde{y}_x)$ für stückweise lineare Approximationen für die meisten Anwendungen ausreichend.

11.4 Probleme

Ideen:

- Horner's Methode zur Auswertung von Polynomen
- Rekursive Strategie für Newton-Polynomkoeffizienten.
- Splines, deCasteljeau
- Prüfen Sie die Dreieckflächeninterpolation der baryzentrischen Interpolation

Kapitel 12

Numerische Integration und Differenzierung

Im vorherigen Kapitel haben wir Werkzeuge entwickelt, um sinnvolle Werte einer Funktion $f(x)$ anhand einer Stichprobe von Werten $(x_i, f(x_i))$ im Bereich von f einzugeben. Offensichtlich ist dieses Interpolationsproblem an sich nützlich, um Funktionen zu vervollständigen, von denen bekannt ist, dass sie stetig oder differenzierbar sind, deren Werte jedoch nur an einer Reihe isolierter Punkte bekannt sind. In einigen Fällen möchten wir jedoch die Eigenschaften dieser Funktionen untersuchen. Insbesondere wenn wir Werkzeuge aus der Analysis auf f anwenden möchten, müssen wir in der Lage sein, seine Integrale und Ableitungen zu approximieren.

Tatsächlich gibt es viele Anwendungen, bei denen numerische Integration und Differentiation eine Schlüsselrolle bei der Berechnung spielen. Im einfachsten Fall werden einige bekannte Funktionen als Integrale definiert. Beispielsweise wird die „Fehlerfunktion“, die als kumulative Verteilung einer Gauß- oder Glockenkurve verwendet wird, wie folgt geschrieben:

$$\operatorname{erf}(x) = \frac{2}{\sqrt{\pi}} \int_0^x e^{-t^2} dt$$

Näherungen von $\operatorname{erf}(x)$ werden in vielen statistischen Zusammenhängen benötigt, und ein sinnvoller Ansatz zur Ermittlung dieser Werte besteht darin, das obige Integral numerisch auszuführen.

In anderen Fällen sind numerische Näherungen von Ableitungen und Integralen Teil eines größeren Systems. Beispielsweise werden Methoden, die wir in zukünftigen Kapiteln zur Approximation von Lösungen für Differentialgleichungen entwickeln, stark von diesen Approximationen abhängen. In ähnlicher Weise erscheinen in der rechnerischen Elektrodynamik Integralgleichungen, die nach einer unbekannten Funktion $\tilde{y}$ bei gegebenem Kern K und Ausgabe f aufgelöst werden, in der Beziehung:

$$f(x) = \int_{\mathbb{R}^n} K(x,y) \tilde{y}(y) dy.$$

Diese Arten von Gleichungen müssen gelöst werden, um elektrische und magnetische Felder abzuschätzen. Wenn $\tilde{y}$ und K jedoch nicht sehr speziell sind, können wir nicht hoffen, ein solches Integral in geschlossener Form zu finden, und lösen diese Gleichung allein für die unbekannte Funktion $\tilde{y}$.

In diesem Kapitel werden wir verschiedene Methoden zur numerischen Integration und Differenzierung anhand einer Stichprobe von Funktionswerten entwickeln. Bei diesen Algorithmen handelt es sich in der Regel um relativ einfache Annäherungen. Um sie zu vergleichen, werden wir auch einige Strategien entwickeln, die bewerten, wie gut wir die Leistung verschiedener Methoden erwarten.

12.1 Motivation

Es ist nicht schwer, einfache Anwendungen der numerischen Integration und Differentiation zu formulieren, wenn man bedenkt, wie oft die Werkzeuge der Analysis in den Grundformeln und Techniken der Physik, Statistik und anderen Bereichen vorkommen. Hier schlagen wir einige weniger offensichtliche Orte vor, an denen Integration und Differenzierung auftreten.

Beispiel 12.1 (Stichprobe aus einer Verteilung). Angenommen, wir erhalten eine Wahrscheinlichkeitsverteilung $p(t)$ für das Intervall $[0, 1]$; Das heißt, wenn wir zufällig Werte gemäß dieser Verteilung abtasten, erwarten wir, dass $p(t)$ proportional zu der Häufigkeit ist, mit der wir einen Wert in der Nähe von t zeichnen. Eine häufige Aufgabe besteht darin, Zufallszahlen zu generieren, die wie $p(t)$ verteilt sind.

Anstatt jedes Mal, wenn wir ein neues $p(t)$ erhalten, eine spezielle Methode zu entwickeln, ist es möglich, eine nützliche Beobachtung zu machen. Wir definieren die kumulative Verteilungsfunktion von p als

$$F(t) = \int_0^t p(x) dx.$$

Wenn X dann eine gleichmäßig in $[0, 1]$ verteilte Zufallszahl ist, kann man zeigen, dass $F \circ X$ wie p verteilt ist, wobei $F \circ X$ die Umkehrung von F ist. $F \circ X$ können wir Zufallszahlen gemäß einer beliebigen Verteilung p generieren; Diese Näherung läuft auf die Integration von p hinaus, was möglicherweise numerisch erfolgen muss, wenn die Integrale nicht in geschlossener Form bekannt sind.

Beispiel 12.2 (Optimierung). Denken Sie daran, dass die meisten unserer Methoden zum Minimieren und Finden von Wurzeln einer Funktion f davon abhängen, nicht nur Werte $f(x)$, sondern auch ihren Gradienten $\nabla f(x)$ und sogar Hessian H_f zu haben. Wir haben gesehen, dass Algorithmen wie BFGS und die Broyden-Methode während des Optimierungsprozesses grobe Näherungen der Ableitungen von f aufbauen. Wenn f jedoch hohe Frequenzen aufweist, ist es möglicherweise besser, ∇f in der Nähe der aktuellen Iteration x_k zu approximieren, anstatt Werte von möglicherweise weit entfernten Punkten x für $|x - x_k|$ zu verwenden.

Beispiel 12.3 (Rendering). Die Rendering-Gleichung von Raytracing und anderen Algorithmen für hochwertiges Rendering ist ein Integral, das besagt, dass das eine Oberfläche verlassende Licht gleich dem Integral des Lichts ist, das über alle möglichen Einfallsrichtungen in die Oberfläche eindringt, nachdem es reflektiert und gestreut wurde; Im Wesentlichen heißt es, dass die Lichtenergie vor und nach der Wechselwirkung von Licht mit einem Objekt erhalten bleiben muss. Algorithmen zum Rendern müssen dieses Integral annähern, um die Lichtmenge zu berechnen, die von einer Oberfläche emittiert wird, die Licht in einer Szene reflektiert.

Beispiel 12.4 (Bildverarbeitung). Angenommen, wir stellen uns ein Bild als Funktion zweier Variablen $I(x, y)$ vor. Viele Filter, einschließlich Gaußscher Unschärfen, können als Faltungen betrachtet werden, gegeben durch

$$(I \circ g)(x, y) = \int I(u, v) g(x - u, y - v) du dv.$$

Um beispielsweise ein Bild unscharf zu machen, könnten wir g als Gaußsche Funktion betrachten; in diesem Fall kann man sich $(I \circ g)(x, y)$ als gewichteten Durchschnitt der Farben von I in der Nähe des Punktes (x, y) vorstellen. In der Praxis handelt es sich bei Bildern um diskrete Pixelgitter, daher muss dieses Integral angenähert werden.

Beispiel 12.5 (Bayes-Regel). Angenommen, X und Y sind Zufallsvariablen mit kontinuierlichem Wert; Wir können $P(X)$ und $P(Y)$ verwenden, um die Wahrscheinlichkeiten auszudrücken, dass X und Y bestimmte Werte annehmen. Manchmal kann die Kenntnis von X unser Wissen über Y beeinflussen. Wenn X beispielsweise der Blutdruck eines Patienten und Y das Gewicht eines Patienten ist,

Dann kann das Wissen, dass ein Patient ein hohes Gewicht hat, darauf hindeuten, dass er auch einen hohen Blutdruck hat. Wir können daher auch bedingte Wahrscheinlichkeitsverteilungen $P(X|Y)$ (sprich „die Wahrscheinlichkeit von X bei gegebenem Y“) schreiben, die solche Beziehungen ausdrücken.

Eine Grundlage der modernen Wahrscheinlichkeitstheorie besagt, dass $P(X|Y)$ und $P(Y|X)$ wie folgt

$$P(X|Y) = \frac{\text{zusammenhängen: } P(Y|X)P(X)}{P(Y|X)P(X) dY}$$

Das Schätzen des Integrals im Nenner kann bei maschinellen Lernalgorithmen, bei denen die Wahrscheinlichkeitsverteilungen komplexe Formen annehmen, ein ernstes Problem darstellen. Daher werden für Algorithmen zur Parameterauswahl, die diesen Wert als Teil einer größeren Optimierungstechnik verwenden, approximative und oft randomisierte Integrations schemata benötigt.

12.2 Quadratur

Wir beginnen mit der Betrachtung des Problems der numerischen Integration oder Quadratur. Dieses Problem – in einer einzelnen Variablen – kann ausgedrückt werden als: „Gegeben eine Stichprobe von n Punkten aus einer Funktion $f(x)$, $f(x)$ Finden Sie eine Annäherung an $\int_A^B dx$.“ Im vorherigen Abschnitt haben wir mehrere Situationen vorgestellt genau diese Technik.

Es gibt einige Variationen des Problems, die eine leicht unterschiedliche Behandlung oder Anpassung erfordern tion:

- Die Endpunkte a und b können fest sein, oder wir möchten möglicherweise ein Quadratschema finden, das Integrale für viele (a, b) -Paare effizient approximieren kann.
- Möglicherweise können wir $f(x)$ an jedem x abfragen, möchten aber das Integral mit relativ wenigen Stichproben annähern, oder wir erhalten möglicherweise eine Liste vorberechneter Paare $(x_i, f(x_i))$ und sind auf die Verwendung dieser Daten beschränkt Punkte in unserer Näherung.

Diese Überlegungen sollten im Hinterkopf behalten werden, wenn wir verschiedene Algorithmen für das Quadraturproblem entwerfen.

12.2.1 Interpolatorische Quadratur

Viele der im vorherigen Kapitel entwickelten Interpolationsstrategien können mithilfe einer sehr einfachen Beobachtung auf Methoden für die Quadratur erweitert werden. Angenommen, wir schreiben eine Funktion $f(x)$ anhand einer Menge von Basisfunktionen $\tilde{y}_i(x)$:

$$f(x) = \sum_i \tilde{y}_i(x) c_i.$$

Dann können wir das Integral von f wie folgt finden:

$$\begin{aligned} \int_A^B f(x) dx &= \int_A^B \sum_i \tilde{y}_i(x) c_i dx \text{ nach Definition von } f \\ &= \sum_i c_i \int_A^B \tilde{y}_i(x) dx \\ &= \sum_i c_i a_i, \text{ wenn wir die Definition } a_i = \int_A^B \tilde{y}_i(x) dx \end{aligned}$$

Mit anderen Worten: Bei der Integration von f werden einfach die Integrale der Basisfunktionen, aus denen f besteht, linear kombiniert.

Beispiel 12.6 (Monome). Angenommen, wir schreiben $f(x) = \sum_k a_k x^k$. Wir wissen

$$\int_0^1 x^k dx = \frac{1}{k+1},$$

Wenn wir also die obige Ableitung anwenden, wissen wir es

$$\int_0^1 f(x) dx = \sum_k a_k \frac{1}{k+1}.$$

Mit anderen Worten, in unserer obigen Notation haben wir $c_k = \frac{1}{k+1}$ definiert.

Schemata, bei denen wir eine Funktion durch Interpolation von Abtastwerten und Integration der interpolierten Funktion integrieren, werden als interpolatorische Quadraturregeln bezeichnet. Fast alle Methoden, die wir im Folgenden vorstellen, können auf diese Weise geschrieben werden. Natürlich kann es zu einem Henne-Ei-Problem kommen, wenn das Integral $\int_0^1 f(x) dx$ selbst nicht in geschlossener Form bekannt ist. Bestimmte Methoden in finiten Elementen höherer Ordnung lösen dieses Problem, indem sie zusätzliche Rechenzeit in die Erstellung einer qualitativ hochwertigen numerischen Näherung des Integrals eines einzelnen f_i investieren und dann, da alle f_i eine ähnliche Form haben, Koordinatenänderungsformeln anwenden. Schreiben Sie Integrale für die verbleibenden Basisfunktionen. Dieses kanonische Integral kann mithilfe eines hochpräzisen Schemas offline angenähert und dann wiederverwendet werden.

12.2.2 Quadraturregeln

Wenn wir eine Menge von $(x_i, f(x_i))$ -Paaren erhalten, schlägt unsere obige Diskussion die folgende Form für eine Quadraturregel zur Approximation des Integrals von f in einem bestimmten Intervall vor:

$$Q[f] \approx \sum w_i f(x_i).$$

Unterschiedliche Gewichte ergeben unterschiedliche Näherungen des Integrals, von denen wir hoffen, dass sie immer ähnlicher werden, je dichter wir die x_i abtasten.

Tatsächlich legt sogar die klassische Integrationstheorie nahe, dass diese Formel ein vernünftiger Ausgangspunkt ist. Beispielsweise hat das Riemann-Integral, das in vielen Einführungskursen in die Analysis vorgestellt wird, die Form:

$$\int_a^b f(x) dx = \lim_{n \rightarrow \infty} \sum_{k=1}^n f(\tilde{x}_k) (x_{k+1} - x_k)$$

Hier wird das Intervall $[a, b]$ in Teile $a = x_1 < x_2 < \dots < x_n = b$ unterteilt, wobei $\tilde{x}_k = x_{k+1} - x_k$ und $\tilde{x}_k$ ein beliebiger Punkt in $[x_k, x_{k+1}]$ ist. Für eine feste Menge von x_k vor der Grenzziehung kann dieses Integral eindeutig in der obigen $Q[f]$ -Form geschrieben werden.

Aus dieser Perspektive bestimmen die Entscheidungen von $\{x_i\}$ und $\{w_i\}$ vollständig eine Strategie für die Quadratur. Es gibt viele Möglichkeiten, diese Werte zu bestimmen, wie wir im nächsten Abschnitt sehen werden und wie wir bereits für die interpolatorische Quadratur gesehen haben.

Beispiel 12.7 (Methode unbestimmter Koeffizienten). Angenommen, wir reparieren $x_1, \dots, x_n$ und möchten einen sinnvollen Satz begleitender Gewichte w_i finden, sodass $\sum w_i f(x_i)$ eine geeignete Näherung für das Integral ist

aus . Eine Alternative zur oben aufgeführten Basisfunktionsstrategie ist die Verwendung der Methode unbestimmter Koeffizienten. In dieser Strategie wählen wir n Funktionen $f_1(x), \dots, f_n(x)$, deren Integrale bekannt sind, und verlangen, dass unsere Quadraturregel die Integrale dieser Funktionen genau wiederherstellt:

$$\begin{aligned} \int_A^B f_1(x) dx &= w_1 f_1(x_1) + w_2 f_1(x_2) + \dots + w_n f_1(x_n) \\ \int_A^B f_2(x) dx &= w_1 f_2(x_1) + w_2 f_2(x_2) + \dots + w_n f_2(x_n) \\ &\vdots \\ \int_A^B f_n(x) dx &= w_1 f_n(x_1) + w_2 f_n(x_2) + \dots + w_n f_n(x_n) \end{aligned}$$

Dadurch entsteht ein $n \times n$ lineares Gleichungssystem für die w_i .

Eine häufige Wahl besteht darin, $f_k(x) = x^{k-1}$ als Integrale, das heißt, um sicherzustellen, dass das Quadraturschema wiederhergestellt wird von Polynomen niedriger Ordnung zu verwenden. Wir wissen

$$\int_A^B x^{k-1} dx = \frac{b^k - a^k}{k}.$$

Somit erhalten wir das folgende lineare Gleichungssystem für die w_i :

$$\begin{aligned} w_1 + w_2 + \dots + w_n &= \frac{b^2 - a^2}{2} \\ b x_1 w_1 + x_2 w_2 + \dots + x_n w_n &= \frac{b^3 - a^3}{2} \\ x_1^2 w_1 + x_2^2 w_2 + \dots + x_n^2 w_n &= \frac{b^4 - a^4}{2} \\ &\vdots \\ x_1^{n-1} w_1 + x_2^{n-1} w_2 + \dots + x_n^{n-1} w_n &= \frac{b^n - a^n}{2} \end{aligned}$$

Dieses System ist genau das in §11.1.1 diskutierte Vandermonde-System.

12.2.3 Newton-Cotes-Quadratur

Quadratur regiert, wenn das x regiert. x_i s, die gleichmäßig in $[a, b]$ verteilt sind, werden als Newton-Cotes-Quadratur bezeichnet. Wie in Abbildung NUMMER dargestellt, gibt es zwei sinnvolle Möglichkeiten für gleichmäßig verteilte Proben:

- Geschlossene Newton-Cotes-Quadratur platziert x_i an a und b . Insbesondere für $k \in \{1, \dots, n\}$ wir nehmen

$$x_k = a + \frac{(k-1)(b-a)}{n-1}.$$

- Die offene Newton-Cotes-Quadratur platziert kein x_i an a oder b :

$$x_k = a + \frac{k(b-a)}{n+1}.$$

Nachdem diese Wahl getroffen wurde, wenden die Newton-Cotes-Formeln einfach eine Polynominterpolation an, um das Integral von a nach b anzunähern; Der Grad des Polynoms muss offensichtlich $n \geq 1$ sein, um die Quadraturregel wohldefiniert zu halten.

Im Allgemeinen werden wir n relativ klein halten. Auf diese Weise vermeiden wir Oszillations- und Rauschphänomene, die beim Anpassen von Polynomen höheren Grades an einen Satz von Datenpunkten auftreten. Wie bei der stückweisen Polynominterpolation verketten wir dann kleine Teile zu zusammengesetzten Regeln, wenn wir über ein großes Intervall $[a, b]$ integrieren.

Geschlossene Regeln. Geschlossene Newton-Cotes-Quadraturstrategien erfordern $n \geq 2$, um eine Division durch Null zu vermeiden. In der Praxis tauchen häufig zwei Strategien auf:

- Die Trapezregel erhält man für $n = 2$ (also $x_1 = a$ und $x_2 = b$) durch lineare Interpolation von $f(a)$ nach $f(b)$. Es sagt, dass

$$\int_a^b f(x) dx \approx (b-a) \frac{f(a) + f(b)}{2}.$$

- Die Simpson-Regel ergibt sich aus der Annahme von $n = 3$, also gilt jetzt

$$\begin{aligned} x_1 &= a \\ x_2 &= \frac{a+b}{2} \\ x_3 &= b \end{aligned}$$

Integriert man die Parabel, die durch diese drei Punkte geht, erhält man

$$\int_a^b f(x) dx \approx \frac{b-a}{6} \left[f(a) + 4f\left(\frac{a+b}{2}\right) + f(b) \right].$$

Offene Regeln. Offene Regeln für die Quadratur ermöglichen die Möglichkeit von $n = 1$, was die vereinfachte Mittelpunktsregel ergibt:

$$\int_a^b f(x) dx \approx (b-a) f\left(\frac{a+b}{2}\right).$$

Größere Werte von n ergeben Regeln ähnlich der Simpson-Regel und der Trapezregel.

Zusammengesetzte Integration. Im Allgemeinen möchten wir möglicherweise $f(x)$ mit mehr als einem, zwei oder drei Werten x_i integrieren. Es ist offensichtlich, wie man eine zusammengesetzte Regel aus den obigen Mittelpunkt- oder Trapezregeln konstruieren kann, wie in Abbildung NUMMER dargestellt; Summieren Sie einfach die Werte entlang jedes Intervalls. Wenn wir beispielsweise $[a, b]$ in k Intervalle unterteilen, können wir $x_i = a + i\frac{b-a}{k}$ annehmen. Dann lautet die zusammengesetzte Mittelpunktsregel:

$$\int_a^b f(x) dx \approx \frac{b-a}{k} \sum_{i=1}^k f\left(\frac{x_{i-1} + x_i}{2}\right).$$

Ebenso lautet die zusammengesetzte Trapezregel:

$$\int_A^B f(x) dx \approx \sum_{i=1}^k \frac{f(x_i) + f(x_{i+1})}{2} \Delta x$$

$$= \Delta x \left[\frac{1}{2} f(a) + f(x_1) + f(x_2) + \dots + f(x_{k-1}) + \frac{1}{2} f(b) \right]$$

durch Trennung der beiden gemittelten Werte von f in der ersten Zeile und Neuindizierung

Eine alternative Behandlung der zusammengesetzten Mittelpunktsregel besteht darin, die interpolatorische Quadraturformel aus §12.2.1 auf die stückweise lineare Interpolation anzuwenden; In ähnlicher Weise beruht die zusammengesetzte Version der Trapezregel auf einer stückweisen linearen Interpolation.

Die zusammengesetzte Version der Simpson-Regel, dargestellt in Abbildung NUMMER, verkettet jeweils drei Punkte, um parabolische Näherungen vorzunehmen. Benachbarte Parabeln treffen sich an geraden x_i -Achsen und dürfen keine gemeinsamen Tangenten haben. Diese Summation, die nur existiert, wenn n gerade ist, wird zu:

$$\int_A^B f(x) dx \approx \frac{\Delta x}{3} \left[f(a) + 2 \sum_{i=1}^{n/2-1} f(x_{2i}) + 4 \sum_{i=1}^{n/2} f(x_{2i-1}) + f(b) \right]$$

$$= \frac{\Delta x}{3} [f(a) + 4f(x_1) + 2f(x_2) + 4f(x_3) + 2f(x_4) + \dots + 4f(x_{n-1}) + f(b)]$$

Genauigkeit. Bisher haben wir eine Reihe von Quadraturregeln entwickelt, die denselben Satz von $f(x_i)$ auf unterschiedliche Weise effektiv kombinieren, um unterschiedliche Näherungen des Integrals von f zu erhalten. Jede Näherung basiert auf einer anderen technischen Annahme, daher ist unklar, ob eine dieser Regeln besser ist als jede andere. Daher müssen wir Fehlerschätzungen entwickeln, die ihr jeweiliges Verhalten charakterisieren. Wir werden unsere oben genannten Newton-Cotes-Integratoren verwenden, um zu zeigen, wie solche Vergleiche durchgeführt werden könnten, wie in CITE dargestellt.

Betrachten Sie zunächst die Mittelpunktsquadraturregel für ein einzelnes Intervall $[a, b]$. Definieren Sie $c = \frac{1}{2}(a + b)$. Der Taylorreihe von f über c ist:

$$f(x) = f(c) + f'(c)(x - c) + \frac{f''(c)}{2!}(x - c)^2 + \frac{f'''(c)}{3!}(x - c)^3 + \frac{f^{(4)}(c)}{4!}(x - c)^4 + \dots$$

Aufgrund der Symmetrie um c fallen also die ungeraden Terme weg:

$$\int_A^B f(x) dx = (b - a)f(c) + \frac{1}{24}f''(c)(b - a)^3 + \frac{1}{1920}f^{(4)}(c)(b - a)^5 + \dots$$

Beachten Sie, dass der erste Term dieser Summe genau die Schätzung der $\int_A^B f(x) dx$ vom Mittelpunkt bereitgestellt Regel ist, sodass diese Regel bis zu $O(\Delta x^3)$ genau ist).

Wenn man nun a und b in unsere Taylor-Reihe für f über c einfügt, zeigt sich:

$$f(a) = f(c) + f'(c)(a - c) + \frac{f''(c)}{2!}(a - c)^2 + \frac{f'''(c)}{3!}(a - c)^3 + \dots$$

$$f(b) = f(c) + f'(c)(b - c) + \frac{f''(c)}{2!}(b - c)^2 + \frac{f'''(c)}{3!}(b - c)^3 + \dots$$

Addiert man diese und multipliziert beide Seiten mit $(b - a)/2$, erhält man:

$$(b - a) \left(\frac{f(a) + f(b)}{2} - f(c) \right) = \frac{(b - a)^3}{24} f''(c) + \frac{(b - a)^5}{1920} f^{(4)}(c) + \dots$$

Der $f(c)$ -Term verschwindet per Definition von c . Beachten Sie, dass die linke Seite die Integralschätzung der Trapezregel ist und die rechte Seite bis auf den kubischen Term mit unserer Taylor-Reihe für $f(x)$ dx übereinstimmt. Mit anderen Worten, die Trapezregel lautet auch $O(\tilde{y}x^3)$ in einem einzigen Intervall genau.

Wir machen hier eine Pause, um ein zunächst überraschendes Ergebnis zu bemerken: Die Trapez- und Mittelpunktregele haben die gleiche Genauigkeitsreihenfolge! Tatsächlich zeigt die Untersuchung des Termes dritter Ordnung, dass die Mittelpunktsregel etwa zweimal genauer ist als die Trapezregel. Dieses Ergebnis scheint kontraintuitiv zu sein, da die Trapezregel eine lineare Näherung verwendet, während die Mittelpunktsregel konstant ist.

Wie in Abbildung NUMBER dargestellt, stellt die Mittelpunktsregel jedoch tatsächlich das Integral linearer Funktionen wieder her, was ihren zusätzlichen Grad an Genauigkeit erklärt.

Ein ähnliches Argument gilt für die Suche nach einer Fehlerschätzung für die Simpson-Regel. [ERKLÄRUNG SCHREIBEN⁵]. NATION HIER; AUS 205A WENDEN]. Am Ende stellen wir fest, dass die Simpson-Regel einen Fehler wie $O(\tilde{y}x^5)$ hat

Für diese Art der Analyse gilt ein wichtiger Vorbehalt. Im Allgemeinen gilt der Satz von Taylor nur, wenn $\tilde{y}x$ ausreichend klein ist. Wenn die Proben weit voneinander entfernt sind, gelten die Nachteile der Polynominterpolation, und die in Abschnitt NUMBER erläuterten Oszillationsphänomene können bei Integrationsschemata höherer Ordnung zu instabilen Ergebnissen führen.

Kehren wir also zu dem Fall zurück, in dem a und b weit voneinander entfernt sind, unterteilen wir nun $[a, b]$ in Intervalle der Breite $\tilde{y}x$ und wenden innerhalb dieser Intervalle eine unserer Quadraturregeln an. Beachten Sie, dass unsere Gesamtzahl an Intervallen $b/a/\tilde{y}x$ beträgt, daher müssen wir in diesem Fall unsere Fehlerschätzungen mit $1/\tilde{y}x$ multiplizieren. Insbesondere gelten folgende Genauigkeitsordnungen:

- Zusammengesetzter Mittelpunkt: $O(\tilde{y}x^2)$
- Zusammengesetztes Trapez: $O(\tilde{y}x^2)$
- Zusammengesetzter Simpson: $O(\tilde{y}x^4)$

12.2.4 Gaußsche Quadratur

In einigen Anwendungen können wir die Orte x_i auswählen, an denen f abgetastet wird. In diesem Fall können wir nicht nur die Gewichte für die Quadraturregel, sondern auch die Orte x_i optimieren, um die höchste Qualität zu erzielen. Diese Beobachtung führt zu anspruchsvollen, aber theoretisch ansprechenden Quadraturregeln.

Die Einzelheiten dieser Technik liegen außerhalb des Rahmens unserer Diskussion, wir bieten jedoch einen einfachen Weg zu ihrer Ableitung. Nehmen wir insbesondere wie in Beispiel 12.7 an, dass wir x_1 optimieren möchten $\dots, x_n$ und $w_1, \dots, w_n$ gleichzeitig die Ordnung eines Integrationsschemas erhöhen. Jetzt haben wir $2n$ statt n Bekannte, sodass wir Gleichheit für $2n$ Beispiele erzwingen können:

$$\begin{array}{l} \int_A^B f_1(x) dx = w_1 f_1(x_1) + w_2 f_1(x_2) + \dots + w_n f_1(x_n) \\ \int_A^B f_2(x) dx = w_1 f_2(x_1) + w_2 f_2(x_2) + \dots + w_n f_2(x_n) \\ \vdots \\ \int_A^B f_{2n}(x) dx = w_1 f_{2n}(x_1) + w_2 f_{2n}(x_2) + \dots + w_n f_{2n}(x_n) \end{array}$$

Jetzt sind sowohl die x_i als auch die w_i unbekannt, sodass dieses Gleichungssystem nicht mehr linear ist. Wenn wir diese Werte beispielsweise für Polynome im Intervall $[y_1, 1]$ optimieren möchten, würden wir dies tun

müssen das folgende Polynomsystem (CITE) lösen:

$$\begin{aligned} w_1 + w_2 &= \int_{\tilde{y}_1}^1 1 \, dx = 2 \\ w_1 x_1 + w_2 x_2 &= \int_{\tilde{y}_1}^1 x \, dx = 0 \\ 2 w_1 x_1^2 + w_2 x_2^2 &= \int_{\tilde{y}_1}^1 x^2 \, dx = \frac{1}{3} \\ 3 w_1 x_1^3 + w_2 x_2^3 &= \int_{\tilde{y}_1}^1 x^3 \, dx = 0 \end{aligned}$$

Es kann vorkommen, dass Systeme wie dieses mehrere Wurzeln und andere Entartungen haben, die nicht nur von der Wahl der f_i (typischerweise Polynome) abhängen, sondern auch vom Intervall, über das wir ein Integral approximieren. Darüber hinaus sind diese Regeln nicht progressiv in dem Sinne, dass die Menge von x_i für n Datenpunkte nichts mit denen für k Datenpunkte gemeinsam hat, wenn $k = n$, sodass es schwierig ist, Daten wiederzuverwenden, um eine bessere Schätzung zu erzielen. Wenn sie jedoch anwendbar sind, hat die Gaußsche Quadratur den höchstmöglichen Grad für festes n . Die Kronrod-Quadraturregeln versuchen, dieses Problem zu vermeiden, indem sie die Quadratur mit $2n + 1$ Punkten optimieren und gleichzeitig die Gaußschen Punkte wiederverwenden.

12.2.5 Adaptive Quadratur

Wie wir bereits gezeigt haben, gibt es bestimmte Funktionen f , deren Integrale besser mit einer gegebenen Quadraturregel angenähert werden können als andere; Beispielsweise integrieren die Mittelpunkt- und Trapezregeln lineare Funktionen mit voller Genauigkeit, während es zu Abtastproblemen und anderen Problemen kommen kann, wenn f schnell oszilliert.

Denken Sie daran, dass die Gaußsche Quadraturregel darauf hindeutet, dass die Platzierung der x_i einen Einfluss auf die Qualität eines Quadraturschemas haben kann. Eine Information haben wir allerdings noch nicht genutzt: die Werte $f(x_i)$. Schließlich bestimmen diese die Qualität unseres Quadraturschemas.

Vor diesem Hintergrund untersuchen adaptive Quadraturstrategien die aktuelle Schätzung und generieren neue x_i , wenn der Integrand komplizierter ist. Strategien zur adaptiven Integration vergleichen oft die Ausgabe mehrerer Quadraturtechniken, z. B. Trapez und Mittelpunkt, mit der Annahme, dass sie dort übereinstimmen, wo die Abtastung von f ausreichend ist (siehe Abbildung NUMBER). Wenn sie mit einer bestimmten Toleranz in einem bestimmten Intervall nicht übereinstimmen, wird ein zusätzlicher Stichprobenpunkt generiert und die Integralschätzungen werden aktualisiert.

WEITERE DETAILS ODER EIN BEISPIEL HINZUFÜGEN; Diskutieren Sie den rekursiven Algorithmus; GÄNSERICH UND GAUTSCHI

12.2.6 Mehrere Variablen

Oft möchten

wir Funktionen $f(x)$ integrieren, wobei $x \in \mathbb{R}^n$ ist. Wenn beispielsweise $n = 2$ ist, könnten wir durch Berechnen über ein Rechteck integrieren

$$\int_A^B \int_C^D f(x, y) \, dx \, dy.$$

Allgemeiner ausgedrückt möchten wir, wie in Abbildung NUMBER, dargestellt, möglicherweise ein Integral $\int_{\tilde{y}} f(x) \, dx$, finden, bei dem $\tilde{y}$ eine Teilmenge von $\mathbb{R}^n$ ist

Ein „Fluch der Dimensionalität“ erschwert die Integration mit zunehmender Dimension exponentiell. Insbesondere nimmt die Anzahl der Abtastwerte von f , die erforderlich sind, um eine vergleichbare Quadraturgenauigkeit für ein Integral in R_k zu erreichen, wie $O(n^k)$ zu. Diese Beobachtung mag entmutigend sein, ist aber einigermaßen vernünftig: Je mehr Eingabedimensionen für f vorhanden sind, desto mehr Stichproben werden benötigt, um sein Verhalten in allen Dimensionen zu verstehen.

Die einfachste Strategie zur Integration in R_k ist das integrierte Integral. Wenn f beispielsweise eine Funktion ist. Nehmen wir an, wir möchten $\int_a^b \int_c^d f(x, y) dx dy$ finden, um zwei Variablen zu betrachten. Für festes y können wir das innere Integral über x mithilfe einer eindimensionalen Quadraturregel annähern; Anschließend integrieren wir diese Werte über y mithilfe einer anderen Quadraturregel. Offensichtlich verursachen beide Integrationsschemata einen gewissen Fehler, sodass wir möglicherweise n^2 dichter als in einer Dimension abtasten müssen, um die gewünschte Ausgabequalität zu erzielen.

Alternativ können wir, genau wie wir $[a, b]$ in Intervalle unterteilt haben, $\tilde{y}$ in Dreiecke und Rechtecke in 2D, Polyeder oder Kästchen in 3D usw. unterteilen und in jedem Teil einfache interpolatorische Quadraturregeln verwenden. Eine beliebte Option besteht beispielsweise darin, die Ausgabe der baryzentrischen Interpolation innerhalb von Polyedern zu integrieren, da dieses Integral in geschlossener Form bekannt ist.

Wenn n jedoch hoch ist, ist es nicht praktikabel, die Domäne wie vorgeschlagen zu teilen. In diesem Fall können wir die randomisierte Monte-Carlo-Methode verwenden. In diesem Fall generieren wir einfach k Zufallspunkte $x_i \in \tilde{y}$ mit beispielsweise gleichmäßiger Wahrscheinlichkeit. Die Mittelung der Werte $f(x_i)$ ergibt eine Näherung von $\int_{\tilde{y}} f(x) dx$, die wie $1/\tilde{y} \cdot k$ konvergiert – unabhängig von der Dimension von $\tilde{y}$! Also in großen Dimensionen Die Monte-Carlo-Schätzung ist den oben genannten deterministischen Quadraturmethoden vorzuziehen.

WEITERE DETAILS ZUR MONTE-CARLO-KONVERGENZ UND DER AUSWAHL DER VERTEILUNGEN
ÜBER $\tilde{y}$

12.2.7 Konditionierung

Bisher haben wir die Qualität einer Quadraturmethode unter Verwendung der Genauigkeitswerte $O(\tilde{y}^k)$ betrachtet. Aufgrund dieser Metrik ist offensichtlich ein Satz Quadraturgewichte mit großem k vorzuziehen.

Ein anderes Maß gleicht jedoch die mithilfe von Taylor-Argumenten ermittelten Genauigkeitsmessungen aus. Erinnern Sie sich insbesondere daran, dass wir unsere Quadraturregel als $Q[f] = \sum_{i=1}^n w_i f(x_i)$ geschrieben haben. Angenommen, wir stören f in ein anderes $\tilde{f}$. Definieren Sie $\kappa[f, \tilde{f}] = \max_{x \in [a, b]} |f(x) - \tilde{f}(x)|$. Dann,

$$\frac{|Q[f] - Q[\tilde{f}]|}{\|f - \tilde{f}\|} = \frac{|\sum_{i=1}^n w_i (f(x_i) - \tilde{f}(x_i))|}{\|f - \tilde{f}\|} \leq \sum_{i=1}^n |w_i| \frac{|f(x_i) - \tilde{f}(x_i)|}{\|f - \tilde{f}\|} \leq \sum_{i=1}^n |w_i| \kappa[f, \tilde{f}]$$

die Dreiecksungleichung $\|f - \tilde{f}\| \geq |f(x_i) - \tilde{f}(x_i)|$ durch

da $\|f - \tilde{f}\| = \max_{x \in [a, b]} |f(x) - \tilde{f}(x)|$ per Definition.

Somit hängt die Stabilität oder Konditionierung einer Quadraturregel von der Norm des Gewichtssatzes ab w .

Im Allgemeinen lässt sich leicht überprüfen, dass die Konditionierung w schlechter wird, wenn wir die Ordnung der Quadraturgenauigkeit erhöhen, da die w_i große negative Werte annehmen; Dies steht im Gegensatz zum vollständig positiven Fall, bei dem die Konditionierung durch $b - a$ begrenzt ist, da $\sum_{i=1}^n w_i = b - a$ für Polynome in Terpolationsschemata gilt und die meisten Methoden niedriger Ordnung nur positive Koeffizienten haben (CHECK). Diese Tatsache spiegelt die gleiche Intuition wider, dass wir Funktionen nicht mit Polynomen höherer Ordnung interpolieren sollten. Daher bevorzugen wir in der Praxis in der Regel die zusammengesetzte Quadratur gegenüber Methoden höherer Ordnung, die möglicherweise bessere Schätzungen liefern, bei numerischen Störungen jedoch instabil sein können.

12.3 Differenzierung

Die numerische Integration ist ein relativ stabiles Problem, dass der Einfluss eines einzelnen Werts $f(x)$ auf $\int_a^b f(x) dx$ auf Null schrumpft, wenn a und b weit auseinander liegen. Die Approximation der Ableitung einer Funktion $f(x)$ hingegen weist keine solche Stabilitätseigenschaft auf. Aus der Perspektive der Fourier-Analyse kann man zeigen, dass das Integral $f(x)$ im Allgemeinen niedrigere Frequenzen als f hat, während die Differenzierung zur Erzeugung von f die hohen Frequenzen von f verstärkt, was Stichprobenbeschränkungen, Konditionierung und Stabilität für die Approximation von f besonders schwierig macht.

Trotz der schwierigen Umstände sind Approximationen von Ableitungen in der Regel relativ einfach zu berechnen und können abhängig von der jeweiligen Funktion stabil sein. Tatsächlich haben wir bei der Entwicklung der Sekantenregel, der Broyden-Methode usw. einfache Näherungen von Ableitungen und Gradienten verwendet, um Optimierungsroutinen zu leiten.

Hier konzentrieren wir uns auf die Approximation von f für $f: \mathbf{R} \rightarrow \mathbf{R}$. Das Finden von Gradienten und Jacobi-Funktionen erfolgt häufig durch Differenzieren in jeweils einer Dimension, wodurch wir uns effektiv auf das eindimensionale Problem reduzieren, das wir hier betrachten.

12.3.1 Basisfunktionen differenzieren

Der einfachste Fall der Differenzierung liegt bei Funktionen vor, die mithilfe von Interpolationsroutinen erstellt werden. Wenn wir wie in §12.2.1 $f(x) = \sum_i a_i \tilde{y}_i(x)$ schreiben können, dann wissen wir es durch Linearität

$$f'(x) = \sum_i a_i \tilde{y}_i'(x)$$

Mit anderen Worten: Wir können uns die Funktionen $\tilde{y}_i$ als Basis für Ableitungen von Funktionen vorstellen, die in der $\tilde{y}_i$ -Basis geschrieben sind!

Ein Beispiel für dieses Verfahren ist in Abbildung NUMBER dargestellt. Dieses Phänomen verbindet häufig verschiedene Interpolationsschemata. Beispielsweise haben stückweise lineare Funktionen stückweise konstante Ableitungen, Polynomfunktionen haben Polynomableitungen niedrigeren Grades und so weiter; Wir werden auf diese Struktur zurückkommen, wenn wir Diskretisierungen partieller Differentialgleichungen betrachten. In diesem Fall ist es jedoch wertvoll zu wissen, dass f mit voller Sicherheit bekannt ist, obwohl seine Ableitungen, wie in Abbildung NUMMER, unerwünschte Diskontinuitäten aufweisen können.

12.3.2 Endliche Differenzen

Ein häufigerer Fall ist, dass wir eine Funktion $f(x)$ haben, die wir abfragen können, deren Ableitungen jedoch unbekannt sind. Dies geschieht häufig, wenn f eine komplexe Form annimmt oder wenn ein Benutzer $f(x)$ als Unterprogramm ohne analytische Informationen über seine Struktur bereitstellt.

Die Definition des Derivats legt einen sinnvollen Ansatz nahe:

$$\lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h}$$

Wie zu erwarten ist, gilt für ein endliches $h > 0$ mit kleinem $|h|$. Der Ausdruck im Grenzwert liefert einen möglichen Wert, der $f'(x)$ annähert.

Um diese Intuition zu untermauern, können wir Taylor-Reihen verwenden, um zu schreiben:

$$f(x+h) = f(x) + f'(x)h + \frac{f''(x)}{2}h^2 + \dots$$

Das Umordnen dieses Ausdrucks zeigt:

$$= \frac{f(x+h) - f(x)}{h} - \frac{f(x) - f(x-h)}{h} + O(h)$$

Somit weist die folgende Vorwärtsdifferenznäherung von f' eine lineare Konvergenz auf:

$$\frac{f(x+h) - f(x)}{h}$$

In ähnlicher Weise zeigt das Umdrehen des Vorzeichens von h , dass auch Rückwärtsdifferenzen eine lineare Konvergenz aufweisen:

$$\frac{f(x) - f(x-h)}{h}$$

Mit einem Trick können wir tatsächlich die Konvergenz unserer Näherung verbessern. Von Taylor Satz können wir schreiben:

$$\begin{aligned} f(x+h) &= f(x) + f'(x)h + \frac{1}{2}f''(x)h^2 + \frac{1}{6}f'''(x)h^3 + \dots \\ f(x-h) &= f(x) - f'(x)h + \frac{1}{2}f''(x)h^2 - \frac{1}{6}f'''(x)h^3 + \dots \\ \frac{f(x+h) - f(x-h)}{2h} &= f'(x) + \frac{1}{6}f'''(x)h^2 + \dots \end{aligned}$$

Somit ergibt diese zentrierte Differenz eine Näherung von $f'(x)$ mit quadratischer Konvergenz; Dies ist die höchste Konvergenzordnung, die wir mit einer geteilten Differenz erwarten können. Wir können jedoch eine höhere Genauigkeit erreichen, indem wir f an anderen Punkten auswerten, z. B. $x + 2h$, obwohl diese Näherung in der Praxis nicht oft zugunsten einer einfachen Verringerung von h verwendet wird.

Die Erstellung von Schätzungen für Ableitungen höherer Ordnung kann durch ähnliche Konstruktionen erfolgen. Für Wenn wir beispielsweise die Taylor-Entwicklungen von $f(x+h)$ und $f(x-h)$ addieren, sehen wir

$$\begin{aligned} f(x+h) + f(x-h) &= 2f(x) + \frac{1}{3}f'''(x)h^3 + O(h^5) \\ \frac{f(x+h) + f(x-h) - 2f(x)}{h^2} &= \frac{1}{3}f'''(x) + O(h^2) \end{aligned}$$

Um ähnliche Kombinationen für höhere Ableitungen vorherzusagen, besteht ein Trick darin, unseren zweiten zu beachten Die Ableitungsformel kann unterschiedlich faktorisiert werden:

$$\frac{f(x+h) - 2f(x) + f(x-h))}{h^2} = \frac{f(x+h)-f(x)}{h} - \frac{f(x)-f(x-h)}{h}$$

Das heißt, unsere Näherung der zweiten Ableitung ist eine „endliche Differenz endlicher Differenzen“.

Eine Möglichkeit, diese Formel zu interpretieren, ist in Abbildung NUMBER dargestellt. Wenn wir die Vorwärtsdifferenznäherung von f' zwischen x und $x+h$ berechnen, können wir uns vorstellen, dass diese Steigung bei $x+h/2$ liegt; In ähnlicher Weise können wir Rückwärtsdifferenzen verwenden, um eine Steigung bei $x-h/2$ zu platzieren. Durch Ermitteln der Steigung zwischen diesen Werten wird die Näherung auf x zurückgesetzt.

Eine Strategie, die die Konvergenz der oben genannten Näherungen verbessern kann, ist die Richardson-Extrapolation. Als Beispiel für ein allgemeineres Muster nehmen wir an, wir möchten Vorwärtsdifferenzen verwenden, um f anzunähern. Definieren

$$D(h) \approx h \frac{f(x+h) - f(x)}{h}.$$

Offensichtlich nähert sich $D(h) \approx f'(x)$, wenn $h \rightarrow 0$. Genauer gesagt wissen wir jedoch aus unserer Diskussion in §12.3.2, dass $D(h)$ die Form annimmt:

$$D(h) = f'(x) + \frac{1}{2}f''(x)h + O(h^2)$$

Angenommen, wir kennen $D(h)$ und $D(\tilde{y}h)$ für $0 < \tilde{y} < 1$. Wir wissen:

$$D(\tilde{y}h) = f'(x) + \frac{1}{2}f''(x)\tilde{y}h + O(h^2)$$

Wir können diese beiden Beziehungen in einer Matrix schreiben:

$$\begin{pmatrix} 1 & \frac{1}{2}h \\ 1 & \frac{1}{2}\tilde{y}h \end{pmatrix} \begin{pmatrix} f(x) \\ f(x) \end{pmatrix} = \begin{pmatrix} D(h) \\ D(\tilde{y}h) \end{pmatrix} + O(h^2)$$

Oder gleichwertig,

$$\begin{pmatrix} f(x) \\ f(x) \end{pmatrix} = \begin{pmatrix} \frac{1}{2}h & \tilde{y} \\ 1 & \frac{1}{2}\tilde{y}h \end{pmatrix}^{-1} \begin{pmatrix} D(h) \\ D(\tilde{y}h) \end{pmatrix} + O(h^2)$$

Das heißt, wir haben eine $O(h)$ -Näherung von $f'(x)$ unter Verwendung von $D(h)$ genommen und $O(h^2)$ ca. daraus eine $O(h)$ -Imitation gemacht! Diese clevere Technik ist eine Methode zur Sequenzbeschleunigung, da sie die Konvergenzordnung der verbesserten Näherung $D(h)$ verbessert. Der gleiche Trick lässt sich allgemeiner auf viele andere Probleme anwenden, indem man eine Näherung $D(h) = a + bhn + O(h^m)$ schreibt, wobei $m > n$ ist, wobei a die Größe ist, die wir schätzen möchten und b ist der nächste Term in der Taylor-Entwicklung. Tatsächlich kann die Richardson-Extrapolation sogar rekursiv angewendet werden, um Näherungen immer höherer Ordnung zu erstellen.

12.3.3 Auswahl der Schrittgröße

Im Gegensatz zur Quadratur hat die numerische Differentiation eine merkwürdige Eigenschaft. Es scheint, dass jede Methode, die wir wählen, beliebig genau sein kann, indem wir einfach ein ausreichend kleines h wählen. Diese Beobachtung ist aus der Perspektive interessant, dass wir ohne zusätzliche Rechenzeit qualitativ hochwertigere Näherungen erzielen können. Der Haken ist jedoch, dass wir durch h dividieren und immer mehr ähnliche Werte $f(x)$ und $f(x+h)$ vergleichen müssen; In der Arithmetik mit endlicher Genauigkeit führt das Addieren und/oder Dividieren durch Werte nahe Null zu numerischen Problemen und Instabilitäten. Somit gibt es einen Bereich von h -Werten, die nicht groß genug sind, um einen signifikanten Diskretisierungsfehler hervorzurufen, und nicht klein genug, um numerische Probleme zu verursachen; Abbildung NUMBER zeigt ein Beispiel für die Differenzierung einer einfachen Funktion in der IEEE-Gleitkomma-Arithmetik.

12.3.4 Integrierte Größen

Nicht abgedeckt in CS 205A, Herbst 2013.

12.4 Probleme

- Gaußsche Quadratur – enthält immer Mittelpunkte, Strategie unter Verwendung orthogonaler Polynome
- Adaptive Quadratur
- Anwendungen der Richardson-Extrapolation an anderer Stelle

Kapitel 13

Gewöhnliche Differentialgleichungen

Wir haben das Problem der Interpolation in Kapitel 11 motiviert, indem wir vom Analysieren zum Finden von Funktionen übergegangen sind. Das heißt, bei Problemen wie Interpolation und Regression ist die Unbekannte eine Funktion f und die Aufgabe des Algorithmus besteht darin, fehlende Daten zu ergänzen.

Wir setzen diese Diskussion fort, indem wir ähnliche Probleme beim Ausfüllen von Funktionswerten betrachten. Hier ist unsere Unbekannte weiterhin eine Funktion f , aber anstatt einfach fehlende Werte zu erraten, möchten wir komplexere Designprobleme lösen. Betrachten Sie beispielsweise die folgenden Probleme:

- Finden Sie f , das eine andere Funktion f_0 annähert, aber zusätzliche Kriterien erfüllt (Glattheit, Kontinuität, Niederfrequenz usw.).
- Simulieren Sie eine dynamische oder physikalische Beziehung als $f(t)$, wobei t die Zeit ist.
- Finden Sie f mit ähnlichen Werten wie f_0 , aber bestimmten Eigenschaften, die mit einer anderen Funktion gemeinsam sind g_0 .

In jedem dieser Fälle ist unsere Unbekannte eine Funktion f , aber unser Erfolgskriterium ist komplexer als „passt zu einem bestimmten Satz von Datenpunkten“.

Die Theorien der gewöhnlichen Differentialgleichungen (ODEs) und der partiellen Differentialgleichungen (PDEs) untersuchen den Fall, dass wir eine Funktion $f(x)$ basierend auf Informationen über oder Beziehungen zwischen ihren Ableitungen finden möchten. Beachten Sie, dass wir in unserer Diskussion der Quadratur bereits eine einfache Version dieses Problems gelöst haben: Bei gegebenem $f(t)$ bieten Methoden für die Quadratur Möglichkeiten, $f(t)$ durch Integration zu approximieren.

In diesem Kapitel betrachten wir den Fall gewöhnlicher Differentialgleichungen und insbesondere Anfangswertprobleme. Hier ist die Unbekannte eine Funktion $f(t) : \mathbf{R} \rightarrow \mathbf{R}_n$ und wir haben eine Gleichung gegeben, die durch f und seine Ableitungen sowie $f(0)$ erfüllt ist; Unser Ziel ist es, $f(t)$ für $t > 0$ vorherzusagen. Wir werden mehrere Beispiele für ODEs liefern, die in der Informatikliteratur vorkommen, und dann mit der Beschreibung gängiger Lösungstechniken fortfahren.

Als Hinweis verwenden wir die Notation f' , um die Ableitung $d f/dt$ von $f : [0, \infty) \rightarrow \mathbf{R}_n$ zu bezeichnen. Unser Ziel wird es sein, $f(t)$ gegebene Beziehungen zwischen t , $f(t)$, $f'(t)$, $f''(t)$ usw. zu finden.

13.1 Motivation

ODEs tauchen in nahezu jedem Teil wissenschaftlicher Beispiele auf, und es ist nicht schwer, auf praktische Situationen zu stoßen, die eine Lösung erfordern. Die Grundgesetze der physikalischen Bewegung werden beispielsweise durch eine ODE angegeben:

Beispiel 13.1 (Newtons zweites Gesetz). Erinnern Sie sich in Fortsetzung von §5.1.2 daran, dass Newtons zweites Bewegungsgesetz besagt, dass $F = ma$, das heißt, die Gesamtkraft auf ein Objekt ist gleich seiner Masse mal seiner Beschleunigung. Wenn wir n Teilchen gleichzeitig simulieren, können wir uns vorstellen, alle ihre Positionen in einem Vektor $x \in \mathbb{R}^{3n}$ zusammenzufassen. In ähnlicher Weise können wir eine Funktion $F(t, x, \dot{x}) \in \mathbb{R}^{3n}$ schreiben, die die Zeit, die Positionen der Teilchen und ihre Geschwindigkeiten annimmt und die Gesamtkraft auf jedes Teilchen zurückgibt. Diese Funktion kann Wechselbeziehungen zwischen Partikeln (z. B. Gravitationskräfte oder Federn), externe Effekte wie Windwiderstand (der von x abhängt), externe Kräfte, die sich mit der Zeit t ändern, usw. berücksichtigen.

Um dann die Positionen aller Teilchen als Funktionen der Zeit zu ermitteln, möchten wir die Gleichung $\ddot{x} = F(t, x, \dot{x})/m$ lösen. Als Ausgangsbedingung werden uns üblicherweise die Positionen und Geschwindigkeiten aller Teilchen zum Zeitpunkt $t = 0$ gegeben.

Beispiel 13.2 (Proteinfaltung). Im kleineren Maßstab sind die Gleichungen, die die Bewegung von Molekülen regeln, auch gewöhnliche Differentialgleichungen. Ein besonders anspruchsvoller Fall ist die Proteinfaltung, bei der die geometrische Struktur eines Proteins durch die Simulation intermolekularer Kräfte über die Zeit vorhergesagt wird. Diese Kräfte nehmen viele oft nichtlineare Formen an, die Forscher in der Computerbiologie weiterhin vor Herausforderungen stellen.

Beispiel 13.3 (Gradientenabstieg). Angenommen, wir möchten eine Energiefunktion $E(x)$ über alle x minimieren. Wir haben in Kapitel 8 gelernt, dass $\nabla E(x)$ in die Richtung zeigt, in der E bei einem gegebenen x am stärksten abnimmt, also haben wir eine Liniensuche entlang dieser Richtung von x durchgeführt, um E lokal zu minimieren. Eine alternative Option, die in bestimmten Theorien beliebt ist, besteht darin, eine ODE der Form $\dot{x} = -\nabla E(x)$ zu lösen; Mit anderen Worten: Stellen Sie sich x als eine Funktion der Zeit $x(t)$ vor, die versucht, E zu verringern, indem sie bergab geht.

Nehmen wir zum Beispiel an, wir möchten $Ax = b$ für das symmetrische positiv definite A auflösen. Wir wissen aus §10.1.1, dass dies äquivalent zur Minimierung von $E(x) = \frac{1}{2}Ax \cdot \tilde{A}x + c$. Somit könnten wir versuchen, die ODE zu lösen $\ddot{y} f(x) = b - \tilde{A}Ax$ ist. Für $t \rightarrow \infty$ erwarten wir, dass $x(t)$ das lineare System immer besser erfüllt.

Beispiel 13.4 (Massensimulation). Angenommen, wir schreiben Videospielsoftware, die eine realistische Simulation virtueller Menschenmengen, Tiere, Raumschiffe und dergleichen erfordert. Eine Strategie zur Erzeugung plausibler Bewegungen, dargestellt in Abbildung NUMMER, ist die Verwendung von Differentialgleichungen. Hierbei wird die Geschwindigkeit eines Mitglieds der Menschenmenge als Funktion seiner Umgebung bestimmt; In Menschenmengen können beispielsweise die Nähe anderer Menschen, die Entfernung zu Hindernissen usw. die Bewegungsrichtung eines bestimmten Agenten beeinflussen. Diese Regeln können einfach sein, aber insgesamt ist ihr Zusammenspiel komplex. Dieser Maschinerie liegen stabile Integratoren für Differentialgleichungen zugrunde, da wir kein merklich unrealistisches oder unphysikalisches Verhalten wünschen.

13.2 Theorie der ODEs

Eine vollständige Behandlung der Theorie gewöhnlicher Differentialgleichungen liegt außerhalb des Rahmens unserer Diskussion und wir verweisen den Leser für weitere Einzelheiten auf CITE. Abgesehen davon erwähnen wir hier einige Highlights, die für unsere Entwicklung in zukünftigen Abschnitten relevant sein werden.

13.2.1 Grundbegriffe

Das allgemeinste ODE-Anfangswertproblem hat die folgende Form:

$$\begin{aligned} &\text{Finden Sie } f(t) : \mathbb{R} \rightarrow \mathbb{R}^n \\ &\mathbf{R}, \text{ das } F[t, f(t), f'(t), f''(t), \dots, f^{(k)}(t)] = 0 \text{ Gegeben} \\ &f(0), f'(0), f''(0), \dots, f^{(k-1)}(0) \end{aligned}$$

Hier ist F eine Beziehung zwischen f und allen seinen Ableitungen; Wir verwenden $f^{(k)}$, um die k -te Ableitung von f zu bezeichnen. Wir können uns ODEs als bestimmende Entwicklung von f über die Zeit t vorstellen; Wir kennen f und seine Ableitungen zum Zeitpunkt Null und möchten es für die Zukunft vorhersagen.

ODEs nehmen selbst in einer einzigen Variablen viele Formen an. Bezeichnen Sie zum Beispiel $y = f(t)$ und nehmen Sie an, dass $y \in \mathbb{R}^1$. Dann sind Beispiele für ODEs:

- $y' = 1 + \cos t$: Diese ODE kann durch Integration beider Seiten, z. B. mittels Quadratur, gelöst werden
- $y' = ay$: Diese ODE ist linear in y
- $y' = ay + e^t$: Diese ODE ist zeit- und ortsabhängig
- $y'' + 3y' + y = t$: Diese ODE beinhaltet mehrere Ableitungen von y
- $y \sin y = e^{-t}$: Diese ODE ist in y und t nichtlinear.

Offensichtlich kann es schwierig sein, die allgemeinsten ODEs zu lösen. Wir werden den Großteil unserer Diskussion auf den Fall expliziter ODEs beschränken, bei denen die Ableitung höchster Ordnung isoliert werden kann:

Definition 13.1 (Explizite ODE). Eine ODE ist explizit, wenn sie in der Form geschrieben werden kann

$$f^{(k)}(t) = F[t, f(t), f'(t), f''(t), \dots, f^{(k-1)}(t)].$$

Eine explizite Form des zweiten Newtonschen Gesetzes ist beispielsweise $x''(t) = \frac{1}{m} a(t, x(t), x'(t))$.

Überraschenderweise lässt sich bei Verallgemeinerung des Tricks in §5.1.2 tatsächlich jede explizite ODE in eine Gleichung erster Ordnung $f'(t) = F[t, f(t)]$ umwandeln, wobei f eine mehrdimensionale Ausgabe hat. Diese Beobachtung impliziert, dass wir bei unserer Behandlung von ODE-Algorithmen nicht mehr als eine Ableitung benötigen. Um diese Beziehung zu sehen, erinnern wir uns einfach daran, dass $d^2y/dt^2 = d/dt(dy/dt)$. Somit können wir eine Zwischenvariable $z = dy/dt$ definieren und d^2y/dt^2 als dz/dt mit der Einschränkung $z = dy/dt$ verstehen. Allgemeiner gesagt, wenn wir das explizite Problem lösen wollen

$$f^{(k)}(t) = F[t, f(t), f'(t), f''(t), \dots, f^{(k-1)}(t)],$$

wobei $f : \mathbb{R} \rightarrow \mathbb{R}^n$, dann definieren wir $g(t) : \mathbb{R} \rightarrow \mathbb{R}^{kn}$ unter Verwendung der ODE erster Ordnung:

$$\begin{array}{ccccccc} & \ddot{y} & g_1(t) & \ddot{y} & \ddot{y} & g_2(t) & \ddot{y} \\ \frac{d}{dt} & & g_2(t) & & & g_3(t) & \\ & & \vdots & & & \vdots & \\ & & g_{k-1}(t) & & & g_k(t) & \\ & & g_k(t) & & & & \\ & & & & & F[t, g_1(t), g_2(t), \dots, g_{k-1}(t)] & \end{array}$$

Hier bezeichnen wir $g(t) : \mathbb{R} \rightarrow \mathbb{R}^n$, um n Komponenten von g zu enthalten. Dann erfüllt $g_1(t)$ die ursprüngliche ODE. Um dies zu sehen, überprüfen wir einfach, ob unsere obige Gleichung $g_2(t) = g_1'(t)$, $g_3(t) = g_2'(t) = g_1''(t)$ usw. impliziert. Somit zeigt die Durchführung dieser Ersetzungen, dass die letzte Zeile die ursprüngliche ODE kodiert.

Der obige Trick wird unsere Notation vereinfachen, es sollte jedoch darauf geachtet werden, dass dieser Ansatz die Berechnungen nicht trivialisiert. Insbesondere wird unsere Funktion $f(t)$ in vielen Fällen nur eine einzige Ausgabe haben, die ODE jedoch in mehreren Ableitungen vorliegen. Wir ersetzen diesen Fall durch eine Ableitung und mehrere Ausgaben.

Beispiel 13.5 (ODE-Erweiterung). Angenommen, wir möchten $y' = 3y - 2y + y$ lösen, wobei $y(t) : \mathbb{R} \rightarrow \mathbb{R}$. Diese Gleichung entspricht:

$$\frac{D}{dt} \begin{pmatrix} y \\ y' \\ y'' \end{pmatrix} = \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ -2 & 3 & 1 \end{pmatrix} \begin{pmatrix} y \\ y' \\ y'' \end{pmatrix}$$

So wie unser obiger Trick es uns ermöglicht, nur ODEs erster Ordnung zu berücksichtigen, können wir unsere Notation noch stärker auf autonome ODEs beschränken. Diese Gleichungen haben die Form $f(t) = F[f(t)]$, das heißt, F hängt nicht mehr von t ab. Dazu könnten wir definieren

$$g(t) = \begin{pmatrix} f(t) \\ g^-(t) \end{pmatrix}.$$

Dann können wir stattdessen die folgende ODE für g lösen:

$$g'(t) = \begin{pmatrix} f(t) \\ g^-(t) \end{pmatrix} = F \begin{pmatrix} f(t) \\ g^-(t) \end{pmatrix}.$$

Insbesondere gilt $g^-(t) = t$ unter der Annahme, dass $g^-(0) = 0$ ist.

Es gibt viele Möglichkeiten, das Verhalten von ODEs zu visualisieren, wie in Abbildung NUMBER dargestellt. Wenn beispielsweise die Unbekannte $f(t)$ eine Funktion einer einzelnen Variablen ist, können wir uns $F[f(t)]$ als die Steigung von $f(t)$ vorstellen, wie in Abbildung NUMBER dargestellt. Alternativ können wir, wenn $f(t)$ eine Ausgabe in $\mathbb{R}^2$ hat, die Abhängigkeit von der Zeit t nicht mehr visualisieren, aber wir können einen Phasenraum zeichnen, der den Tangens von $f(t)$ an jedem $(x, y) \in \mathbb{R}^2$ zeigt

13.2.2 Existenz und Einzigartigkeit

Bevor wir mit der Diskretisierung des Anfangswertproblems fortfahren, sollten wir kurz anerkennen, dass nicht alle Differentialgleichungsprobleme lösbar sind. Darüber hinaus lassen einige Differentialgleichungen mehrere Lösungen zu.

Beispiel 13.6 (Unlösbare ODE). Betrachten Sie die Gleichung $y' = 2y/t$, wobei $y(0) = 0$ gegeben ist; Beachten Sie, dass wir nicht durch Null dividieren, da $y(0)$ vorgeschrieben ist. Umschreiben als

$$\frac{1}{y} \frac{dy}{dt} = \frac{2}{t}$$

und die Integration bezüglich t auf beiden Seiten zeigt:

$$\ln |y| = 2 \ln t + c$$

Oder äquivalent: $y = Ct^2$ für ein $C \in \mathbb{R}$. Beachten Sie, dass $y(0) = 0$ in diesem Ausdruck ist, was unseren Anfangsbedingungen widerspricht. Somit hat diese ODE keine Lösung mit den gegebenen Anfangsbedingungen.

Beispiel 13.7 (Nicht eindeutige Lösungen). Betrachten Sie nun dieselbe ODE mit $y(0) = 0$. Betrachten Sie $y(t)$, gegeben durch $y(t) = Ct^2$ für jedes $C \in \mathbb{R}$. Dann ist $y'(t) = 2Ct$. Daher,

$$\frac{2Ct^2}{t} = \frac{2Ct^2}{t} = 2Ct = y'(t),$$

Dies zeigt, dass die ODE unabhängig von C durch diese Funktion gelöst wird. Daher sind Lösungen dieses Problems nicht eindeutig.

Zum Glück gibt es eine umfangreiche Theorie, die das Verhalten und die Stabilität von Lösungen von Differentialgleichungen charakterisiert. Unsere Entwicklung im nächsten Kapitel wird einen stärkeren Satz von Bedingungen haben, die für die Existenz einer Lösung erforderlich sind, aber tatsächlich ist es unter schwachen Bedingungen für f möglich zu zeigen, dass eine ODE $f'(t) = F(f(t))$ eine Lösung hat. Ein solcher Satz garantiert beispielsweise die lokale Existenz einer Lösung:

Satz 13.1 (Lokale Existenz und Einzigartigkeit). Angenommen, F ist stetig und Lipschitz, das heißt $|F(y) - F(x)| \leq L|y - x|$ für ein L . Dann lässt die ODE $f'(t) = F(f(t))$ genau eine Lösung für alle $t \in [0, \delta]$ unabhängig von den Anfangsbedingungen.

In unserer weiteren Entwicklung gehen wir davon aus, dass die ODE, die wir zu lösen versuchen, die Bedingungen eines solchen Theorems erfüllt; Diese Annahme ist insofern ziemlich realistisch, als es zumindest lokal ein ziemlich degeneriertes Verhalten geben müsste, um solche schwachen Annahmen zu durchbrechen.

13.2.3 Modellgleichungen

Eine Möglichkeit, das Verhalten von ODEs zu verstehen, besteht darin, das Verhalten von Lösungen einiger einfacher Modellgleichungen zu untersuchen, die in geschlossener Form gelöst werden können. Diese Gleichungen stellen Linearisierungen praktischerer Gleichungen dar und modellieren somit lokal die Art von Verhalten, die wir erwarten können.

Wir beginnen mit ODEs in einer einzelnen Variablen. Angesichts unserer Vereinfachungen in §13.2.1 wäre die einfachste Gleichung, mit der wir arbeiten könnten, $y' = F(y)$, wobei $y(t) : \mathbb{R} \rightarrow \mathbb{R}$. Eine lineare Näherung würde Gleichungen vom Typ $y' = ay + b$ ergeben. Das Einsetzen von $y' = y + b/a$ ergibt: $y' = y + b/a = a(y - b/a) + b/a$. Somit induziert die Konstante b in unseren Modellgleichungen einfach eine Verschiebung, und für unsere phänomenologische Untersuchung in diesem Abschnitt können wir $b = 0$ annehmen.

Mit dem obigen Argument können wir das Verhalten von $y' = F(y)$ lokal verstehen, indem wir die Linearität untersuchen Gleichung $y' = ay$. Tatsächlich zeigt die Anwendung von Standardargumenten aus der Infinitesimalrechnung dies

$$y(t) = Ce^{at}.$$

Offensichtlich gibt es drei Fälle, die in Abbildung NUMMER dargestellt sind:

1. $a > 0$: In diesem Fall werden die Lösungen immer größer; in der Tat, wenn $y(t)$ und $y'(t)$ beide die erfüllen ODE mit leicht unterschiedlichen Startbedingungen, da sie bei $t \rightarrow \infty$ divergieren.
2. $a = 0$: Das System wird in diesem Fall durch Konstanten gelöst; Lösungen mit unterschiedlichen Startpunkten bleiben im gleichen Abstand voneinander.
3. $a < 0$: Dann nähern sich alle Lösungen der ODE 0, wenn $t \rightarrow \infty$.

Wir sagen, dass die Fälle 2 und 3 stabil sind, in dem Sinne, dass eine leichte Störung von $y(0)$ zu Lösungen führt, die sich mit der Zeit immer weiter annähern; Fall 1 ist instabil, da ein kleiner Fehler bei der Angabe des Eingabeparameters $y(0)$ mit fortschreitender Zeit t verstärkt wird. Instabile ODEs erzeugen schlecht gestellte Rechenprobleme; Ohne sorgfältige Überlegung können wir in diesem Fall nicht erwarten, dass numerische Methoden brauchbare Lösungen liefern, da selbst theoretische Ausgaben sehr empfindlich auf Störungen der Eingabe reagieren. Andererseits sind stabile Probleme gut gestellt, da kleine Fehler in $y(0)$ mit der Zeit abnehmen.

Wenn wir zu mehreren Dimensionen vordringen, könnten wir die linearisierte Gleichung untersuchen

$$y = J a.$$

Wie in §5.1.2 erläutert, sind $y_1, \dots, y_k$ Eigenvektoren von A mit den Eigenwerten $\tilde{y}_1, \dots, \tilde{y}_k$ und $y(0) = c_1 y_1 + \dots + c_k y_k$,

$$y(t) = c_1 e^{\tilde{y}_1 t} y_1 + \dots + c_k e^{\tilde{y}_k t} y_k.$$

Mit anderen Worten, die Eigenwerte von A treten in unserem eindimensionalen Beispiel an die Stelle von a . Aus diesem Ergebnis lässt sich leicht ableiten, dass ein multivariates System genau dann stabil ist, wenn sein Spektralradius kleiner als eins ist.

In Wirklichkeit möchten wir $y = F[y]$ für allgemeine Funktionen F lösen. Unter der Annahme, dass F differenzierbar ist, können wir $F[y] \approx F[y_0] + JF(y_0)(y - y_0)$ schreiben, was die obige Modellgleichung ergibt nach einer Schicht. Daher erwarten wir für kurze Zeiträume ein ähnliches Verhalten wie die Modellgleichung. Darüber hinaus können die Bedingungen in Satz 13.1 als eine Grenze für das Verhalten von JF angesehen werden, was eine Verbindung zu weniger lokalisierten ODE-Theorien herstellt.

13.3 Zeitschrittschemata

Wir beschreiben nun mehrere Methoden zur Lösung der nichtlinearen ODE $y = F[y]$ für potenziell nichtlineare Funktionen F . Im Allgemeinen werden unsere Methoden bei einem gegebenen „Zeitschritt“ h verwendet, um Schätzungen von $y(t + h)$ zu generieren, gegeben $y(t)$. Durch die Anwendung dieser Methoden werden iterativ Schätzungen von $y_0 \approx y(t)$, $y_1 \approx y(t + h)$, $y_2 \approx y(t + 2h)$, $y_3 \approx y(t + 3h)$ usw. generiert. Beachten Sie, dass der Mechanismus zur Generierung jedes zusätzlichen Schritts derselbe ist wie der erste, da F keine t -Abhängigkeit aufweist, sodass wir größtenteils nur einen einzelnen Schritt dieser Methoden beschreiben müssen. Wir nennen Methoden zum Erzeugen von Näherungen für $y(t)$ Integratoren, was die Tatsache widerspiegelt, dass sie die Ableitungen in der Eingabegleichung herausintegrieren.

Von zentraler Bedeutung für unsere Überlegungen ist der Gedanke der Stabilität. So wie ODEs stabil oder instabil sein können, können auch Diskretisierungen stabil oder instabil sein. Wenn h beispielsweise zu groß ist, häufen sich bei einigen Schemata Fehler exponentiell an; Im Gegensatz dazu sind andere Methoden insofern stabil, als die Lösungen auch dann begrenzt bleiben, wenn h groß ist. Stabilität kann jedoch mit Genauigkeit konkurrieren; Oft sind zeitstabile Schemata schlechte Annäherungen an $y(t)$, auch wenn garantiert ist, dass sie kein wildes Verhalten aufweisen.

13.3.1 Vorwärts-Euler

Unsere erste ODE-Strategie ergibt sich aus unserer Konstruktion des Vorwärtsdifferenzierungsschemas in §12.3.2:

$$F[y_k] = y(t) = \frac{y_{k+1} - y_k}{h} + O(h) h$$

Das Auflösen dieser Beziehung für y_{k+1} zeigt

$$y_{k+1} = y_k + hF[y_k] + O(h^2) \approx y_k + hF[y_k].$$

Daher wendet das Vorwärts-Euler-Schema die Formel rechts an, um y_{k+1} zu schätzen. Dies ist eine der effizientesten Strategien für die Zeitmessung, da sie einfach F auswertet und ein Vielfaches des Ergebnisses zu y_k addiert. Aus diesem Grund nennen wir es eine explizite Methode, das heißt, es gibt eine explizite Formel für y_{k+1} in Bezug auf y_k und F .

Die Analyse der Genauigkeit dieser Methode ist ziemlich einfach. Beachten Sie, dass unsere Näherung y_{k+1} ist $O(h^2)$ Schritt einen quadratischen Fehler induziert. Wir nennen diesen Fehler lokalisierten Kürzungsfehler, da es sich um den Fehler handelt, der durch einen einzelnen Schritt verursacht wird. Das Wort „Trunkierung“ bezieht sich auf die Tatsache, dass wir eine Taylor-Reihe gekürzt haben, um diese Formel zu erhalten. Natürlich kann es sein, dass unser Iterationsk aufgrund angesamelter Kürzungsfehler aus früheren Iterationen bereits ungenau ist. Wenn wir von t_0 bis t mit $O(1/h)$ Schritten integrieren, dann sieht unser Gesamtfehler wie folgt aus: $O(h)$; Diese Schätzung stellt einen globalen Kürzungsfehler dar, und daher schreiben wir normalerweise, dass das Vorwärts-Euler-Schema „genau erster Ordnung“ ist.

Die Stabilität dieser Methode erfordert etwas mehr Überlegung. In unserer Diskussion werden wir die Stabilität von Methoden im Fall $y' = ay$ mit einer Variablen ausarbeiten, mit der Intuition, dass ähnliche Aussagen auf mehrdimensionale Gleichungen übertragen werden können, indem wir a durch den Spektralradius ersetzen. In diesem Fall wissen wir es

$$y_{k+1} = y_k + ah y_k = (1 + ah)y_k.$$

Mit anderen Worten: $y_k = (1 + ah)^k y_0$. Somit ist der Integrator stabil, wenn $|1 + ah| \leq 1$, da sonst $|y_k|$ exponentiell. Unter der Annahme, dass $a < 0$ ist (andernfalls ist das Problem schlecht gestellt), können wir Folgendes vereinfachen:

$$|1 + ah| \leq 1 \iff 1 - |a|h \leq 1 + ah \leq 1 + |a|h \iff 0 \leq |a|h \leq \frac{2}{|a|}, \text{ da } a < 0$$

Somit lässt Vorwärts-Euler eine Zeitschrittbeschränkung für die Stabilität zu, die durch unsere letzte Bedingung für h gegeben ist. Mit anderen Worten: Die Ausgabe von Vorwärts-Euler kann explodieren, selbst wenn $y' = ay$ stabil ist, wenn h nicht klein genug ist. Abbildung NUMMER zeigt, was passiert, wenn diese Bedingung befolgt oder verletzt wird. In mehreren Dimensionen können wir diese Einschränkung durch eine analoge ersetzen, indem wir den Spektralradius von A verwenden. Für nichtlineare ODEs gibt diese Formel einen Leitfaden für die Stabilität zumindest lokal in der Zeit; Global muss h möglicherweise angepasst werden, wenn die Jacobi-Funktion von F schlechter konditioniert wird.

13.3.2 Rückwärts-Euler

In ähnlicher Weise hätten wir das Rückwärtsdifferenzierungsschema bei y_{k+1} anwenden können, um einen ODE-Integrator zu entwerfen:

$$F[y_{k+1}] = y'(t_k) = \frac{y_{k+1} - y_k}{h} + O(h)$$

Somit lösen wir das folgende potenziell nichtlineare Gleichungssystem für y_{k+1} :

$$y_k = y_{k+1} - hF[y_{k+1}].$$

Da wir diese Gleichung nach y_{k+1} lösen müssen, ist der Rückwärts-Euler ein impliziter Integrator.

Diese Methode ist durch einen identischen Beweis erster Ordnung genau wie Forward Euler. Die Stabilität dieser Methode steht jedoch im erheblichen Gegensatz zur Vorwärts-Euler-Methode. Unter erneuter Betrachtung der Modellgleichung $y' = ay$ schreiben wir:

$$y_k = y_{k+1} - h a y_{k+1} \Rightarrow y_{k+1} = \frac{y_k}{1 - h a}$$

Parallel zu unserem vorherigen Argument ist der Rückwärts-Euler unter der folgenden Bedingung stabil:

$$\frac{1}{1 - h a} \approx 1 + h a + \frac{1}{2} (h a)^2 + \dots$$

$$|1 - h a| > 1 \quad \text{oder} \quad |1 - h a| > 1$$

$$2 h a > h a \quad \text{oder} \quad h a > 0, \text{ für } a < 0$$

Offensichtlich nehmen wir immer $h > 0$ an, sodass Rückwärts-Euler bedingungslos stabil ist.

Selbst wenn der Rückwärts-Euler stabil ist, ist er natürlich nicht unbedingt genau. Wenn h zu groß ist, geht y_k viel zu schnell gegen Null. Bei der Simulation von Stoffen und anderen physikalischen Materialien, die viele Hochfrequenzdetails erfordern, um realistisch zu sein, ist die Rückwärts-Euler-Methode möglicherweise keine effektive Wahl.

Darüber hinaus müssen wir $F[\cdot]$ invertieren, um nach y_{k+1} aufzulösen.

Beispiel 13.8 (Rückwärts-Euler). Angenommen, wir möchten $y' = Ay$ für $A \in \mathbb{R}^{n \times n}$ lösen. Um dann y_{k+1} zu finden, lösen wir das folgende System:

$$y_k = y_{k+1} - h A y_{k+1} \Rightarrow y_{k+1} = (I_n - h A)^{-1} y_k$$

13.3.3 Trapezmethode

Angenommen, y_k ist zum Zeitpunkt t_k bekannt und y_{k+1} stellt den Wert zum Zeitpunkt $t_{k+1} = t_k + h$ dar. Angenommen, wir kennen auch $y_{k+1/2}$ auf halbem Weg zwischen diesen beiden Schritten. Dann wissen wir durch unsere Ableitung der zentrierten Differenzierung:

$$y_{k+1} - y_k = h F[y_{k+1/2}] + O(h^3)$$

Aus unserer Ableitung der Trapezregel:

$$\frac{F[y_{k+1}] + F[y_k]}{2} = F[y_{k+1/2}] + O(h^2)$$

Das Ersetzen dieser Beziehung ergibt unser erstes Integrationsschema zweiter Ordnung, die Trapezmethode zur Integration von ODEs:

$$y_{k+1} = y_k + h \frac{F[y_{k+1}] + F[y_k]}{2}$$

Wie die Rückwärts-Euler-Methode ist diese Methode implizit, da wir diese Gleichung nach y_{k+1} lösen müssen.

Wenn wir erneut eine Stabilitätsanalyse für $y' = ay$ durchführen, finden wir in diesem Fall Zeitschritte der Trapezmethode

$$y_{k+1} = y_k + \frac{1}{2} h a (y_{k+1} + y_k)$$

Mit anderen Worten,

$$y_k = \frac{1 + \frac{1}{2} h a}{1 - \frac{1}{2} h a} y_0$$

Die Methode ist somit stabil, wenn

$$\frac{h a + 1}{1 - \frac{h^2 a^2}{2}} < 1.$$

Es ist leicht zu erkennen, dass diese Ungleichung immer dann gilt, wenn $a < 0$ und $h > 0$, was zeigt, dass die Trapezmethode bedingungslos stabil ist.

Trotz ihrer höheren Genauigkeit bei gleichbleibender Stabilität weist die Trapezmethode jedoch einige Nachteile auf, die sie weniger beliebt machen als die Rückwärts-Euler-Methode. Berücksichtigen Sie insbesondere die

Verhältnis

$$R \approx \frac{y_{k+1}}{y_k} = \frac{1 + \frac{1}{2} h a}{1 - \frac{1}{2} h a}$$

Wenn $a < 0$ ist und h groß genug ist, wird dieses Verhältnis schließlich negativ; Tatsächlich gilt für $h \approx R \approx -1$. Wie in Abbildung NUMBER dargestellt, neigt die trapezförmige Integrationsmethode bei zu großen Zeitschritten dazu, ein unerwünschtes Oszillationsverhalten zu zeigen, das überhaupt nicht dem entspricht, was wir für Lösungen von $y' = ay$ erwarten würden.

13.3.4 Runge-Kutta-Methoden

Eine Klasse von Integratoren kann durch die folgende Beobachtung abgeleitet werden:

$$\begin{aligned} y_{k+1} &= y_k + \int_{t_k}^{t_{k+1}} y(t) dt \text{ nach dem Fundamentalsatz der Analysis} \\ &= y_k + \int_{t_k}^{t_{k+1}} F[y(t)] dt \end{aligned}$$

Natürlich funktioniert die direkte Verwendung dieser Formel nicht zur Formulierung einer Methode zur Zeitschrittmessung, da wir $y(t)$ nicht kennen, aber eine sorgfältige Anwendung unserer Quadraturformeln aus dem vorherigen Kapitel kann praktikable Strategien hervorbringen.

Angenommen, wir wenden zur Integration die Trapezmethode an. Dann finden wir:

$$h y_{k+1} = y_k + \frac{h}{2} (F[y_k] + F[y_{k+1}]) + O(h^3)$$

Dies ist die Formel, die wir für die Trapezmethode in §13.3.3 geschrieben haben.

Wenn wir jedoch nicht implizit nach y_{k+1} auflösen wollen, müssen wir einen Ausdruck für $\text{approx } F[y_{k+1}]$ finden. Herstellung Kumpel $F[y_{k+1}]$. Mit der Euler-Methode wissen wir jedoch, dass $y_{k+1} = y_k + hF[y_k] + O(h^2)$. Diese Substitution für y_{k+1} hat keinen Einfluss auf die Reihenfolge der Approximation des trapezförmigen Zeitschritts oben, sodass wir schreiben können:

$$= y_k + \frac{h}{2} (F[y_k] + F[y_k + hF[y_k]]) + O(h^3)$$

Ignorieren des $O(h^3)$ Begriffe ergeben eine neue Integrationsstrategie, die als Heun-Methode bekannt ist zweiter Ordnung genau und explizit.

Wenn wir das Stabilitätsverhalten der Heun-Methode für $y' = ay$ für $a < 0$ untersuchen, wissen wir:

$$y_{k+1} = y_k + h(a y_k + a(y_k + h a y_k)) \\ + \frac{h^2}{2} a^2 (2 + h a) y_k = 1 \\ 1 = 1 + h a + h^2 \frac{a^2}{2} (2 + h a) \quad \text{Ja}$$

Somit ist die Methode stabil, wenn

$$1 - h a + h^2 \frac{a^2}{2} (2 + h a) \geq 0$$

Die Ungleichung auf der rechten Seite zeigt $h |a| < 2$ und die auf der linken Seite gilt immer für $h > 0$ und $|a| > 0$, daher ist die Stabilitätsbedingung $h |a| < 2$. Die Heun-

Methode ist ein Beispiel für eine Runge-Kutta-Methode, die durch die Anwendung von Quadraturmethoden auf das obige Integral und das Einsetzen von Euler-Schritten in $F[\cdot]$ abgeleitet wird. Forward Euler ist eine genaue Runge-Kutta-Methode erster Ordnung, und Heuns Methode ist zweiter Ordnung. Eine beliebte Runge-Kutta-Methode vierter Ordnung (abgekürzt „RK4“) ist gegeben durch:

$$y_{k+1} = y_k + h \frac{1}{6} (k_1 + 4k_2 + k_3) y_{k+1}$$

wobei $k_1 = F[y_k]$

$$k_2 = F[y_k + \frac{h}{2} k_1]$$

$$k_3 = F[y_k + \frac{h}{2} k_2]$$

$$k_4 = F[y_k + h k_3]$$

Diese Methode kann durch Anwendung der Simpson-Regel für die Quadratur abgeleitet werden.

Runge-Kutta-Methoden sind beliebt, weil sie explizit und daher leicht auszuwerten sind und gleichzeitig ein hohes Maß an Genauigkeit bieten. Der Preis dieser Genauigkeit besteht jedoch darin, dass $F[\cdot]$ mehrmals ausgewertet werden muss. Darüber hinaus können Runge-Kutta-Strategien auf implizite Methoden erweitert werden, die steife Gleichungen lösen können.

13.3.5 Exponentielle Integratoren Eine Klasse

von Integratoren, die eine hohe Genauigkeit erreichen, wenn $F[\cdot]$ annähernd linear ist, besteht darin, unsere Lösung der Modellgleichung explizit zu verwenden. Insbesondere wenn wir die ODE $y' = Ay$ unter Verwendung von Eigenvektoren von A (oder einer anderen Methode) lösen würden, könnten wir eine explizite Lösung $y(t)$ finden, wie in §13.2.3 erläutert. Normalerweise schreiben wir $y_{k+1} = e^{Ah} y_k$, kodiert unsere Potenzierung der Eigenwerte λ (tatsächlich können wir jetzt eine Matrix e^{Ah} finden, wenn wir schreiben e^{Ah} aus diesem Ausdruck, der die ODE zur Zeit h löst).

$$y' = Ay + G[y],$$

Wenn G eine nichtlineare, aber kleine Funktion ist, können wir eine ziemlich hohe Genauigkeit erreichen, indem wir den A -Teil explizit integrieren und dann den nichtlinearen G -Teil separat approximieren. Beispielsweise wendet der Exponentialintegrator erster Ordnung Vorwärts-Euler auf den nichtlinearen G -Term an:

$$A h y_{k+1} = e^{y_k \tilde{y} A} (1 - e^{A h}) G[y_k]$$

Die Analyse, die die Vorteile dieser Methode aufzeigt, ist komplexer als das, was wir geschrieben haben, aber intuitiv ist klar, dass sich diese Methoden besonders gut verhalten, wenn G klein ist.

13.4 Mehrwertige Methoden

Die Transformationen in §13.2.1 ermöglichten es uns, die Notation im vorherigen Abschnitt erheblich zu vereinfachen, indem wir alle expliziten ODEs auf die Form $y = F[y]$ reduzierten. Obwohl alle expliziten ODEs auf diese Weise geschrieben werden können, ist nicht klar, ob dies immer der Fall sein sollte.

Insbesondere als wir ODEs k -ter Ordnung auf ODEs erster Ordnung reduzierten, führten wir eine Reihe von Variablen ein, die die ersten bis $k-1$ -ten Ableitungen der gewünschten Ausgabe darstellen. Tatsächlich kümmern wir uns in unserer endgültigen Lösung nur um die nullte Ableitung, also um die Funktion selbst, sodass Genauigkeitsreihenfolgen für die temporären Variablen weniger wichtig sind.

Betrachten Sie aus dieser Perspektive die Taylor-Reihe

$$y(t_k + h) = y(t_k) + h y'(t_k) + \frac{h^2}{2} y''(t_k) + O(h^3)$$

Wenn wir y nur bis $O(h^2)$, hat dies keinen Einfluss auf unsere Näherung, da y mit h multipliziert wird. Wenn wir y nur bis zu $O(h)$ kennen, hat diese Näherung ebenfalls keinen Einfluss auf die obigen Taylor-Reihenterme, da sie mit $h^2/2$ multipliziert werden. Daher betrachten wir nun „mehrwertiges“ $Meth(k)(t) = F[t, y(t), y'(t), \dots, y^{(k-1)}(t)]$ mit Genauigkeit Funktion y zu integrieren. unterschiedlicher Ordnung für ods , entworfen, um y verschiedene Ableitungen der

Angesichts der Bedeutung von Newtons zweitem Gesetz $F = ma$ beschränken wir uns auf den Fall $y = F[t, y, y']$; Für den weniger verbreiteten Fall k -ter Ordnung gibt es viele Erweiterungen. Wir führen einen „Geschwindigkeits“-Vektor $v(t) = y'(t)$ und einen „Beschleunigungs“-Vektora ein. Mit unserer bisherigen Reduktion wollen wir das folgende System erster Ordnung lösen:

$$\begin{aligned} y'(t) &= v(t) \\ v'(t) &= a(t) \\ a(t) &= F[t, y(t), v(t)] \end{aligned}$$

Unser Ziel ist es, einen speziell auf dieses System zugeschnittenen Integrator abzuleiten.

13.4.1 Newmark-Systeme

Wir beginnen mit der Ableitung der berühmten Klasse der Newmark-Integratoren.¹ Bezeichnen Sie y_k , v_k und a_k als Positions-, Geschwindigkeits- und Beschleunigungsvektoren zum Zeitpunkt t_k ; Unser Ziel ist es, zur Zeit $t_{k+1} = t_k + h$ vorzurücken.

¹Wir verfolgen die Entwicklung unter http://www.stanford.edu/group/frg/course_work/AA242B/CA-AA242B-Ch7.pdf.

Verwenden Sie $y(t)$, $v(t)$ und $a(t)$, um die Funktionen der Zeit zu bezeichnen, vorausgesetzt, wir beginnen bei t_k . Dann natürlich wir können schreiben

$$v_{k+1} = v_k + \int_{t_k}^{t_{k+1}} a(t) dt$$

Wir können y_{k+1} auch als Integral mit $a(t)$ schreiben, indem wir ein paar Schritte befolgen:

$$\begin{aligned} y_{k+1} &= y_k + \int_{t_k}^{t_{k+1}} v(t) dt \\ &= y_k + [tv(t)]_{t_k}^{t_{k+1}} - \int_{t_k}^{t_{k+1}} t a(t) dt \text{ nach partieller Integration} \\ &= y_k + t_{k+1} v_{k+1} - t_k v_k - \int_{t_k}^{t_{k+1}} t a(t) dt \text{ durch Erweiterung des Differenzterms} \\ &= y_k + h v_k + t_{k+1} v_{k+1} - t_{k+1} v_k - \int_{t_k}^{t_{k+1}} t a(t) dt \text{ durch Addieren und Subtrahieren von } h v_k \\ &= y_k + h v_k + t_{k+1} (v_{k+1} - v_k) - \int_{t_k}^{t_{k+1}} t a(t) dt \text{ nach Faktorisierung} \\ &= y_k + h v_k + t_{k+1} \int_{t_k}^{t_{k+1}} a(t) dt - \int_{t_k}^{t_{k+1}} t a(t) dt \text{ da } v(t) = a(t) \\ &= y_k + h v_k + \int_{t_k}^{t_{k+1}} (t_{k+1} - t) a(t) dt \end{aligned}$$

Angenommen, wir wählen $\tilde{y} \in [t_k, t_{k+1}]$. Dann können wir Ausdrücke für a_{k+1} unter Verwendung der Taylor-Reihe über $\tilde{y}$ schreiben:

$$\begin{aligned} a_k &= a(\tilde{y}) + a'(\tilde{y})(t_k - \tilde{y}) + O(h^2) \\ a_{k+1} &= a(\tilde{y}) + a'(\tilde{y})(t_{k+1} - \tilde{y}) + O(h^2) \end{aligned}$$

Wenn wir für jede Konstante $\tilde{y} \in \mathbb{R}$ die erste Gleichung mit $1 - \tilde{y}$ und die zweite mit $\tilde{y}$ skalieren und die Ergebnisse summieren, finden wir:

$$\begin{aligned} a(\tilde{y}) &= (1 - \tilde{y}) a_k + \tilde{y} a_{k+1} + a'(\tilde{y})((1 - \tilde{y})(t_k - \tilde{y}) + \tilde{y}(t_{k+1} - \tilde{y})) + O(h^2) \\ &= (1 - \tilde{y}) a_k + \tilde{y} a_{k+1} + a'(\tilde{y})(\tilde{y} - h\tilde{y} - t_k) + O(h^2) \text{ nach Einsetzen von } t_{k+1} = t_k + h \end{aligned}$$

Am Ende möchten wir von t_k bis t_{k+1} integrieren, um die Geschwindigkeitsänderung zu erhalten. Somit berechnen wir:

$$\begin{aligned} \int_{t_k}^{t_{k+1}} a(\tilde{y}) d\tilde{y} &= (1 - \tilde{y}) h a_k + \tilde{y} h a_{k+1} + \int_{t_k}^{t_{k+1}} a'(\tilde{y})(\tilde{y} - h\tilde{y} - t_k) d\tilde{y} + O(h^3) \\ &= (1 - \tilde{y}) h a_k + \tilde{y} h a_{k+1} + O(h^2), \end{aligned}$$

wobei der zweite Schritt gilt, weil der Integrand $O(h)$ ist und das Integrationsintervall die Breite h hat. Mit anderen Worten, wir wissen jetzt:

$$v_{k+1} = v_k + (1 - \tilde{y}) h a_k + \tilde{y} h a_{k+1} + O(h^2)$$

Um eine ähnliche Näherung für y_{k+1} zu erreichen, können wir schreiben

$$\begin{aligned} \int_{t_k}^{t_{k+1}} (t_{k+1} - \tilde{y}) a(t) dt &= \int_{t_k}^{t_{k+1}} (t_{k+1} - \tilde{y}) ((1 - \tilde{y}) a_k + \tilde{y} a_{k+1} + a'(\tilde{y})(\tilde{y} - h\tilde{y} - t_k)) d\tilde{y} + O(h^3) \\ &= \frac{1}{2} (1 - \tilde{y}) h^2 a_k + \frac{1}{2} \tilde{y} h^2 a_{k+1} + O(h^2) \text{ mit einem ähnlichen Argument.} \end{aligned}$$

Somit können wir unsere frühere Beziehung verwenden, um zu zeigen:

$$y_{k+1} = y_k + hv_k + \int_{t_k}^{t_{k+1}} (\ddot{y} + \ddot{y}_k) a(t) dt \text{ von vorher}$$

$$= y_k + hv_k + \ddot{y} \ddot{y}_k h^2 + \ddot{y}_k h^2 a_{k+1} + O(h^2)$$

Hier verwenden wir $\ddot{y}$ anstelle von $\ddot{y}$ (und absorbieren dabei einen Faktor zwei), da das $\ddot{y}$, das wir für die Approximation von y_{k+1} wählen, nicht dasselbe sein muss wie das, das wir für die Approximation von v_{k+1} wählen.

Nach all dieser Integration haben wir die Klasse der Newmark-Schemata mit zwei Eingaben abgeleitet. Parameter $\ddot{y}$ und $\ddot{y}$, die durch den obigen Beweis eine Genauigkeit erster Ordnung haben:

$$y_{k+1} = y_k + hv_k + \ddot{y} \ddot{y}_k h^2 + \ddot{y}_k h^2 a_{k+1}$$

$$v_{k+1} = v_k + (1 - \ddot{y}) h a_k + \ddot{y} h a_{k+1}$$

$$a_k = F[t_k, y_k, v_k]$$

Unterschiedliche Wahlen von $\ddot{y}$ und $\ddot{y}$ führen zu unterschiedlichen Schemata. Betrachten Sie beispielsweise die folgenden Beispiele:

- $\ddot{y} = \ddot{y} = 0$ ergibt den Konstantbeschleunigungsintegrator:

$$y_{k+1} = y_k + hv_k + h^2 a_k$$

$$v_{k+1} = v_k + h a_k$$

Dieser Integrator ist explizit und gilt genau dann, wenn die Beschleunigung eine konstante Funktion ist.

- $\ddot{y} = 1/2, \ddot{y} = 1$ ergibt den konstanten impliziten Beschleunigungsintegrator:

$$y_{k+1} = y_k + hv_k + h^2 a_{k+1}/2$$

$$v_{k+1} = v_k + h a_{k+1}$$

Hier wird die Geschwindigkeit implizit mithilfe der Rückwärts-Euler-Methode schrittweise erhöht, was eine Genauigkeit erster Ordnung ergibt. Das Update von y kann jedoch geschrieben werden

$$y_{k+1} = y_k + 2h(v_k + v_{k+1}),$$

zeigt, dass es lokal durch die Mittelpunkregel aktualisiert wird; Dies ist unser erstes Beispiel für ein Schema, bei dem die Geschwindigkeits- und Positionsaktualisierungen unterschiedliche Genauigkeitsordnungen haben. Dennoch lässt sich zeigen, dass diese Technik in y global genau erster Ordnung ist.

- $\ddot{y} = 1/4, \ddot{y} = 1/2$ ergibt nach etwas Algebra das folgende Trapezschaema zweiter Ordnung:

$$x_{k+1} = x_k + 2h(v_k + v_{k+1})$$

$$v_{k+1} = v_k + h(a_k + a_{k+1})$$

- $\tilde{\gamma} = 0$, $\tilde{\gamma} = 1/2$ ergibt ein genaues zentrales Differenzierungsschema zweiter Ordnung. Im Kanonischen Form, wir haben

$$x_{k+1} = x_k + h v_k + h^2 a_k \frac{1}{2}$$

$$v_{k+1} = v_k + h (a_k + a_{k+1}) \frac{1}{2}$$

Die Methode trägt ihren Namen, weil die Vereinfachung der obigen Gleichungen zu der alternativen Form führt:

$$v_{k+1} = \frac{y_{k+2} - y_k}{2h}$$

$$2y_{k+1} + y_k a_{k+1} = h^2$$

Es lässt sich zeigen, dass Newmarks Methoden bedingungslos stabil sind, wenn $4\tilde{\gamma} > 2\tilde{\gamma} > 1$, und dass Genauigkeit zweiter Ordnung genau dann auftritt, wenn $\tilde{\gamma} = 1/2$ (CHECK).

13.4.2 Versetztes Raster

Eine andere Möglichkeit, eine Genauigkeit zweiter Ordnung in y zu erreichen, besteht darin, zentrierte Differenzen über die Zeit $t_{k+1/2} = t_k + h/2$ zu verwenden:

$$y_{k+1} = y_k + h v_{k+1/2}$$

Anstatt Taylor-Argumente zu verwenden, um zu versuchen, $v_{k+1/2}$ zu bewegen, können wir Geschwindigkeiten v einfach an halben Punkten auf dem Gitter der Zeitschritte speichern.

Dann können wir ein ähnliches Update verwenden, um die Geschwindigkeiten zu erhöhen:

$$v_{k+3/2} = v_{k+1/2} + h a_{k+1}$$

Beachten Sie, dass diese Aktualisierung tatsächlich auch für x zweiter Ordnung genau ist, denn wenn wir unsere Ausdrücke für $v_{k+1/2}$ und $v_{k+3/2}$ ersetzen, können wir schreiben:

$$a_{k+1} = \frac{1}{h^2} (y_{k+2} - 2y_{k+1} + y_k)$$

Schließlich genügt für den Beschleunigungsterm eine einfache Näherung, da es sich um einen Term höherer Ordnung handelt:

$$a_{k+1} = F(t_{k+1}, x_{k+1}, 2(v_{k+1/2} + v_{k+3/2}))$$

Dieser Ausdruck kann in den Ausdruck für $v_{k+3/2}$ eingesetzt werden.

Wenn $F[\cdot]$ keine Abhängigkeit von v hat, ist die Methode vollständig explizit:

$$y_{k+1} = y_k + h v_{k+1/2}$$

$$= F(t_{k+1}, y_{k+1}) v_{k+3/2} =$$

$$v_{k+1/2} + h a_{k+1}$$

Aufgrund des gestaffelten Zeitrasters und der Tatsache, dass jeder Mittelpunkt zur Aktualisierung der nächsten Geschwindigkeit oder Position verwendet wird, ist dies als Leapfrog-Integrationsmethode bekannt.

Andernfalls wird die Methode implizit, wenn die Geschwindigkeitsaktualisierung von v abhängt. Oftmals ist die Abhängigkeit von der Geschwindigkeit symmetrisch; Beispielsweise ändert der Windwiderstand einfach das Vorzeichen, wenn Sie die Bewegungsrichtung umkehren. Diese Eigenschaft kann im impliziten Schritt zur Geschwindigkeitsaktualisierung zu symmetrischen Matrizen führen, wodurch es möglich wird, konjugierte Gradienten und verwandte schnelle iterative Methoden zur Lösung zu verwenden.

13.5 Zu erledigen

- Definieren Sie steife ODE
- Geben Sie eine Tabelle mit Zeitschrittmethoden für $F[t;y]$ an.
- Verwenden Sie die y -Notation konsequenter

13.6 Probleme

- TVD RK
- Mehrschritt-/Mehrwertmethoden a la Heath
- Verlet-Integration
- Symplektische Integratoren

Kapitel 14

Partielle Differentialgleichungen

Unsere Intuition für gewöhnliche Differentialgleichungen beruht im Allgemeinen auf der zeitlichen Entwicklung physikalischer Systeme. Gleichungen wie das zweite Newtonsche Gesetz, das die Bewegung physikalischer Objekte über die Zeit bestimmt, dominieren die Literatur zu solchen Anfangswertproblemen; Weitere Beispiele sind chemische Konzentrationen, die im Laufe der Zeit reagieren, Populationen von Raubtieren und Beutetieren, die von Saison zu Saison interagieren, und so weiter. In jedem Fall ist die Anfangskonfiguration – z. B. die Positionen und Geschwindigkeiten von Teilchen in einem System zum Zeitpunkt Null – bekannt, und die Aufgabe besteht darin, das Verhalten im Laufe der Zeit vorherzusagen.

In diesem Kapitel beschäftigen wir uns jedoch mit der Möglichkeit, Beziehungen zwischen verschiedenen Ableitungen einer Funktion zu koppeln. Es ist nicht schwer, Beispiele zu finden, bei denen diese Kopplung notwendig ist. Wenn Sie beispielsweise Rauch- oder Gasgrößen wie „Druckgradienten“, die Ableitung des Drucks eines Gases im Raum, simulieren, können Sie ermitteln, wie sich das Gas im Laufe der Zeit bewegt. Diese Struktur ist sinnvoll, da Gas auf natürliche Weise von Hochdruckregionen in Tiefdruckregionen diffundiert. Bei der Bildverarbeitung koppeln Ableitungen noch natürlicher, da Messungen an Bildern tendenziell gleichzeitig in x- und y-Richtung erfolgen.

Gleichungen, die Ableitungen von Funktionen miteinander verknüpfen, werden als partielle Differentialgleichungen bezeichnet. Sie sind Gegenstand einer reichhaltigen, aber stark nuancierten Theorie, die einer umfassenderen Behandlung würdig ist. Daher wird es hier unser Ziel sein, Schlüsselideen zusammenzufassen und ausreichend Material zur Lösung von Problemen bereitzustellen, die in der Praxis häufig auftreten.

14.1 Motivation

Partielle Differentialgleichungen (PDEs) entstehen, wenn die Unbekannte eine Funktion $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$ ist. Wir erhalten eine oder mehrere Beziehungen zwischen den partiellen Ableitungen von f und das Ziel besteht darin, ein f zu finden, das die Kriterien erfüllt. PDEs kommen in fast allen Bereichen der angewandten Mathematik vor, und wir listen nachfolgend nur einige auf.

Abgesehen davon sollten wir vor der Einführung spezifischer PDEs einige Notationen einführen. Insbesondere gibt es einige Kombinationen partieller Ableitungen, die in der Welt der PDEs häufig vorkommen. Wenn $f : \mathbb{R}^3 \rightarrow \mathbb{R}$ und $v : \mathbb{R}^3 \rightarrow \mathbb{R}^3$, Dann sollten Sie sich die folgenden Operatoren merken:

$$\text{Steigung: } \nabla f = \left(\frac{\partial f}{\partial x_1}, \frac{\partial f}{\partial x_2}, \frac{\partial f}{\partial x_3} \right)$$

$$\text{Divergenz: } \vec{v} \cdot \vec{v} = v_1 + v_2 + v_3$$

$$\text{Lorenz: } \vec{v} \times \vec{v} = \begin{pmatrix} v_2 v_3 - v_3 v_2 \\ v_3 v_1 - v_1 v_3 \\ v_1 v_2 - v_2 v_1 \end{pmatrix}$$

$$\text{Laplace-Operator: } \Delta f = \frac{\partial^2 f}{\partial x_1^2} + \frac{\partial^2 f}{\partial x_2^2} + \frac{\partial^2 f}{\partial x_3^2}$$

Beispiel 14.1 (Flüssigkeitssimulation). Der Fluss von Flüssigkeiten wie Wasser und Rauch wird durch die Navier-Stokes-Gleichungen bestimmt, ein System von PDEs mit vielen Variablen. Nehmen wir insbesondere an, dass sich eine Flüssigkeit in einem Bereich $\Omega \subset \mathbb{R}^3$ bewegt. Wir definieren die folgenden Variablen, dargestellt in Abbildung NUMBER:

- $t \in [0, \infty)$: Zeit
- $\mathbf{v}(t) : \Omega \rightarrow \mathbb{R}^3$: Die Geschwindigkeit der Flüssigkeit
- $\rho(t) : \Omega \rightarrow \mathbb{R}$: Die Dichte der Flüssigkeit
- $p(t) : \Omega \rightarrow \mathbb{R}$: Der Druck der Flüssigkeit
- $\mathbf{f}(t) : \Omega \rightarrow \mathbb{R}^3$: Äußere Kräfte wie Schwerkraft auf die Flüssigkeit

Wenn die Flüssigkeit die Viskosität μ hat und wir davon ausgehen, dass sie inkompressibel ist, lauten die Navier-Stokes-Gleichungen:

$$\rho \frac{d\mathbf{v}}{dt} = \rho \mathbf{f} - \nabla p + \mu \Delta \mathbf{v}$$

Hier ist $\frac{d}{dt} = \frac{\partial}{\partial t} + \mathbf{v} \cdot \nabla$; Wir stellen uns den Gradienten ∇ als einen Gradienten im Raum vor und nicht als $\frac{\partial}{\partial x_1}, \frac{\partial}{\partial x_2}, \frac{\partial}{\partial x_3}$, Zeit, also $\mathbf{f} = (f_1, f_2, f_3)$. Dieses Gleichungssystem bestimmt die Zeitdynamik der Flüssigkeitsbewegung und tatsächlich konstruiert werden, indem man Newtons zweites Gesetz auf die Verfolgung von „Teilchen“ einer Flüssigkeit anwendet. Seine Aussage beinhaltet jedoch nicht nur Ableitungen in der Zeit $\frac{\partial}{\partial t}$, aber auch Ableitungen im Raum ∇ , was es zu einer PDE macht.

Beispiel 14.2 (Maxwell-Gleichungen). Maxwells Gleichungen bestimmen die Wechselwirkung von elektrischen Feldern $\mathbf{E}$ und magnetischen Feldern $\mathbf{B}$ über die Zeit. Wie bei den Navier-Stokes-Gleichungen stellen wir uns Gradient, Divergenz und Curl als partielle Ableitungen im Raum (und nicht in der Zeit t) vor. Dann kann das Maxwell-System (in „starker“ Form) geschrieben werden:

$$\begin{aligned} \text{Gaußsches Gesetz für elektrische Felder: } \nabla \cdot \mathbf{E} &= \frac{\rho}{\epsilon_0} \\ \text{Gaußsches Gesetz für Magnetismus: } \nabla \cdot \mathbf{B} &= 0 \\ \text{Faradaysches Gesetz: } \nabla \times \mathbf{E} &= -\frac{\partial \mathbf{B}}{\partial t} \\ \text{Ampèresches Gesetz: } \nabla \times \mathbf{B} &= \mu_0 \mathbf{J} + \mu_0 \frac{\partial \mathbf{E}}{\partial t} \end{aligned}$$

Hier sind ϵ_0 und μ_0 physikalische Konstanten und $\mathbf{J}$ kodiert die Dichte des elektrischen Stroms. Genau wie die Navier-Stokes-Gleichungen verknüpfen die Maxwell-Gleichungen die Ableitungen physikalischer Größen in der Zeit t mit ihren Ableitungen im Raum (gegeben durch Curl- und Divergenzterme).

Beispiel 14.3 (Laplace-Gleichung). Angenommen, $\tilde{\gamma}$ ist ein Bereich in $\mathbb{R}^2$ mit der Grenze $\tilde{\gamma}$ und wir erhalten eine Funktion $g: \tilde{\gamma} \rightarrow \mathbb{R}$, dargestellt in Abbildung NUMBER. Möglicherweise möchten wir g in das Innere von $\tilde{\gamma}$ interpolieren. Wenn $\tilde{\gamma}$ jedoch eine unregelmäßige Form hat, können unsere Interpolationsstrategien aus Kapitel 11 scheitern.

Angenommen, wir definieren $f(x) : \tilde{\gamma} \rightarrow \mathbb{R}$ als Interpolationsfunktion. Dann besteht eine von unserem Ansatz zur Methode der kleinsten Quadrate inspirierte Strategie darin, ein Energiefunktional zu definieren:

$$E[f] = \int_{\tilde{\gamma}} |\nabla f(x)|^2 dx$$

Das heißt, $E[f]$ misst die „Gesamtableitung“ von f , gemessen durch Nehmen der Norm seines Gradienten und Integrieren dieser Größe über ganz $\tilde{\gamma}$. Stark schwankende Funktionen f werden hohe Werte von $E[f]$ haben, da die Steigung ∇f an vielen Stellen groß sein wird; glatte und niederfrequente Funktionen f haben dagegen ein kleines $E[f]$, da ihre Steigung überall klein ist.¹ Dann könnten wir verlangen, dass f g interpoliert und dabei im Inneren von $\tilde{\gamma}$ möglichst glatt ist folgende Optimierung:

$$\begin{aligned} &\text{minimiere } E[f] \\ &\text{mit } f(x) = g(x) \text{ für } x \in \tilde{\gamma} \end{aligned}$$

Dieser Aufbau sieht aus wie Optimierungen, die wir in früheren Beispielen gelöst haben, aber jetzt ist unsere Unbekannte eine Funktion f und kein Punkt in $\mathbb{R}^n$!

Wenn E minimiert, dann ist $E[f + h] \geq E[f]$ für alle Funktionen $h(x)$. Diese Aussage gilt auch für kleine Störungen $E[f + \tilde{\gamma}h]$, wenn $\tilde{\gamma} \rightarrow 0$. Wenn wir durch $\tilde{\gamma}$ dividieren und den Grenzwert als $\tilde{\gamma} \rightarrow 0$ annehmen, müssen wir $E[f + \tilde{\gamma}h]|_{\tilde{\gamma}=0} = 0$ haben; dies ist so, als würde man Richtungsableitungen einer Funktion gleich Null setzen, um deren Minima zu ermitteln. Wir können Folgendes vereinfachen:

$$\begin{aligned} E[f + \tilde{\gamma}h] &= \int_{\tilde{\gamma}} |\nabla f(x) + \tilde{\gamma}\nabla h(x)|^2 dx \\ &= \int_{\tilde{\gamma}} (|\nabla f(x)|^2 + 2\tilde{\gamma}\nabla f(x) \cdot \nabla h(x) + \tilde{\gamma}^2 |\nabla h(x)|^2) dx \end{aligned}$$

Differenzierende Shows:

$$\begin{aligned} \frac{D}{d\tilde{\gamma}} E[f + \tilde{\gamma}h] &= \int_{\tilde{\gamma}} (2\nabla f(x) \cdot \nabla h(x) + 2\tilde{\gamma}|\nabla h(x)|^2) dx \\ \frac{D}{d\tilde{\gamma}} E[f + \tilde{\gamma}h]|_{\tilde{\gamma}=0} &= 2 \int_{\tilde{\gamma}} [\nabla f(x) \cdot \nabla h(x)] dx \end{aligned}$$

Diese Ableitung muss für alle h gleich Null sein, daher können wir insbesondere $h(x) = 0$ für alle $x \in \tilde{\gamma}$ wählen. Wenn wir dann die partielle Integration anwenden, erhalten wir:

$$\frac{D}{d\tilde{\gamma}} E[f + \tilde{\gamma}h]|_{\tilde{\gamma}=0} = 2 \int_{\tilde{\gamma}} \nabla h(x) \cdot \nabla f(x) dx$$

Dieser Ausdruck muss für alle (alle!) Störungen h gleich Null sein, also müssen wir $\nabla f(x) = 0$ für alle $x \in \tilde{\gamma}$ haben (ein formaler Beweis liegt außerhalb des Rahmens unserer Diskussion). Das heißt, das obige Interpolationsproblem

¹Die hier verwendete Notation $E[\cdot]$ steht nicht für „Erwartung“, wie es in der Wahrscheinlichkeitstheorie der Fall wäre, sondern ist einfach so ein „Energie“-Funktional; Es ist die Standardnotation in Bereichen der Funktionsanalyse.

kann mit der folgenden PDE gelöst werden: $\Delta^2 u = 0$

$$f(x) = 0 \quad f(x) = g(x) \quad \Delta^2 u = 0$$

Diese Gleichung ist als Laplace-Gleichung bekannt und kann mit dünn besetzten positiv definiten linearen Methoden gelöst werden, wie wir sie in Kapitel 10 behandelt haben. Wie wir gesehen haben, kann sie auf Interpolationsprobleme für unregelmäßige Domänen Ω angewendet werden; Darüber hinaus kann $E[f]$ erweitert werden, um andere Eigenschaften von f zu messen, z. B. wie gut f eine verrauschte Funktion f_0 annähert, um verwandte PDEs durch Parallelisierung des obigen Arguments abzuleiten.

Beispiel 14.4 (Eikonal-Gleichung). Angenommen, $\Omega \subset \mathbb{R}^n$ sei ein geschlossener Raumbereich. Dann könnten wir $d(x)$ als eine Funktion betrachten, die den Abstand von einem Punkt x_0 zu x vollständig innerhalb von Ω misst. Wenn Ω konvex ist, können wir d in geschlossener Form schreiben:

$$d(x) = |x - x_0|.$$

Wie in Abbildung NUMBER dargestellt, wird es jedoch schwieriger, Abstände zu berechnen, wenn Ω nicht konvex ist oder ein komplizierterer Bereich wie eine Oberfläche ist. In diesem Fall erfüllen die Distanzfunktionen d die lokalisierte Bedingung, die als Eikonalgleichung bekannt ist:

$$|\nabla d|^2 = 1.$$

Wenn wir es berechnen können, kann d für Aufgaben wie die Planung von Roboterpfaden verwendet werden, indem die Distanz, die sie zurücklegen müssen, minimiert wird, mit der Einschränkung, dass sie sich nur in Ω bewegen können.

Spezielle Algorithmen, sogenannte Fast-Marching-Methoden, werden verwendet, um Schätzungen von d bei gegebenem x_0 und Ω durch Integration der Eikonalgleichung zu finden. Diese Gleichung ist in der Ableitung $|\nabla d|^2$ nichtlinear, daher sind die Integrationsmethoden für diese Gleichung etwas spezialisiert und der Nachweis ihrer Wirksamkeit ist komplex. Interessanterweise, aber nicht überraschend, haben viele Algorithmen zur Lösung der Eikonalgleichung eine ähnliche Struktur wie Dijkstras Algorithmus zur Berechnung kürzester Pfade entlang von Graphen.

Beispiel 14.5 (Harmonische Analyse). Verschiedene Objekte reagieren unterschiedlich auf Vibrationen, und diese Reaktionen hängen zum großen Teil von der Geometrie der Objekte ab. Beispielsweise können Celli und Klaviere die gleiche Note spielen, aber selbst ein unerfahrener Musiker kann die von ihnen erzeugten Klänge leicht unterscheiden. Aus mathematischer Sicht können wir $\Omega \subset \mathbb{R}^3$ als eine Form annehmen, die entweder als Oberfläche oder als Volumen dargestellt wird.

Wenn wir die Kanten der Form einklemmen, ist ihr Frequenzspektrum durch Lösungen des folgenden Differentialeigenwertproblems gegeben:

$$\begin{aligned} \Delta^2 u &= \lambda u \quad f(x) \\ &= 0 \quad \text{für } x \in \partial\Omega, \text{ wobei } \Delta^2 \end{aligned}$$

der Laplace-Operator von Ω und $\partial\Omega$ der Rand von Ω ist. Abbildung NUMBER zeigt Beispiele dieser Funktionen auf verschiedenen Domänen Ω .

Es lässt sich leicht überprüfen, dass $\sin kx$ dieses Problem löst, wenn Ω das Intervall $[0, 2\pi]$ für $k \in \mathbb{Z}$ ist. Insbesondere ist der Laplace-Operator in einer Dimension $\Delta^2 u = -u''$ und so können wir überprüfen:

$$\begin{aligned} \frac{\Delta^2 u}{-u''} &= \sin kx = k^2 \cos kx \quad \Delta^2 u \\ &= -k^2 \sin kx \\ \sin k \cdot 0 &= 0 \\ \sin k \cdot 2\pi &= 0 \end{aligned}$$

Somit sind die Eigenfunktionen $\sin kx$ mit Eigenwerten $-k^2$.

14.2 Grundlegende Definitionen

Unter Verwendung der CITE-Notation gehen wir davon aus, dass unsere Unbekannte eine Funktion $f: \mathbb{R}^n \rightarrow \mathbb{R}$ ist. Für Gleichungen mit bis zu drei Variablen verwenden wir die Indexnotation, um partielle Ableitungen zu bezeichnen:

$$\begin{aligned} f_{x_i} &= \frac{\partial f}{\partial x_i}, \\ f_{x_i x_j} &= \frac{\partial^2 f}{\partial x_i \partial x_j}, \\ f_{x_i x_j x_k} &= \frac{\partial^3 f}{\partial x_i \partial x_j \partial x_k}, \end{aligned}$$

und so weiter.

Partielle Ableitungen werden im Folgenden üblicherweise als Beziehungen zwischen zwei oder mehr Ableitungen von f angegeben: als

- Linear, homogen: $f_{xx} + f_{xy} = 0$

Linear: $f_{xx} + f_{yy} + f = 0$

- Nichtlinear: $f_{xx} = f_{xy}$

Im Allgemeinen möchten wir $f: \mathbb{R}^n \rightarrow \mathbb{R}$ für ein $\mathbb{R}^n$ finden. So wie ODEs als Anfangswertprobleme angegeben wurden, werden wir die meisten PDEs als Randwertprobleme bezeichnen. Das heißt, unsere Aufgabe besteht darin, f im Inneren von $\mathbb{R}^n$ mit gegebenen Werten an seinem Rand auszufüllen. Tatsächlich können wir uns das ODE-Anfangswertproblem folgendermaßen vorstellen: Die Domäne ist $\mathbb{R} = [0, \infty)$ mit der Grenze $\partial\mathbb{R} = \{0\}$, wo wir Eingabedaten bereitstellen. Abbildung NUMBER veranschaulicht komplexere Beispiele. Randbedingungen für diese Probleme können viele Formen annehmen:

- Dirichlet-Bedingungen geben einfach den Wert von $f(x)$ auf $\partial\mathbb{R}^n$ an
- Neumann-Bedingungen geben die Ableitungen von $f(x)$ auf $\partial\mathbb{R}^n$ an
- Gemischte oder Robin-Bedingungen kombinieren diese beiden

14.3 Modellgleichungen

Erinnern Sie sich an das vorherige Kapitel, dass wir viele Eigenschaften von ODEs verstehen konnten, indem wir eine Modellgleichung $y' = ay$ untersuchten. Wir können versuchen, einen ähnlichen Ansatz für PDEs zu verfolgen, werden jedoch feststellen, dass die Geschichte nuancierter ist, wenn Derivate miteinander verknüpft werden.

Wie bei der Modellgleichung für ODEs werden wir den linearen Fall mit einer Variablen untersuchen. Wir beschränken uns auch auf Systeme zweiter Ordnung, also Systeme, die höchstens die zweite Ableitung von u enthalten; Die Modell-ODE war erster Ordnung, aber hier benötigen wir mindestens zwei Ordnungen, um zu untersuchen, wie Derivate auf nichttriviale Weise interagieren.

Eine lineare PDE zweiter Ordnung hat die folgende allgemeine Form:

$$= 0 \quad a_{ij} \frac{\partial^2 f}{\partial x_i \partial x_j} + b_i \frac{\partial f}{\partial x_i} + c$$

Formal könnten wir den „Gradientenoperator“ wie folgt definieren:

$$\vec{\nabla} = \left(\frac{\partial}{\partial x_1}, \frac{\partial}{\partial x_2}, \dots, \frac{\partial}{\partial x_n} \right)$$

Sie sollten überprüfen, ob dieser Operator eine sinnvolle Schreibweise hat, da Ausdrücke wie $\vec{\nabla} f$, $\vec{\nabla} \cdot \mathbf{v}$ und $\vec{\nabla} \times \mathbf{v}$ alle die richtigen Ausdrücke liefern. In dieser Notation kann man sich die PDE als eine Matrixform vorstellen:

$$(\vec{\nabla} A \vec{\nabla} + \vec{\nabla} \cdot \mathbf{b} + c)f = 0.$$

Diese Form hat viel mit unserer Untersuchung quadratischer Formen in konjugierten Gradienten gemeinsam, und tatsächlich charakterisieren wir PDEs normalerweise durch die Struktur von A:

- Wenn A positiv oder negativ definit ist, ist das System elliptisch.
- Wenn A positiv oder negativ semidefinit ist, ist das System parabolisch.
- Wenn A nur einen Eigenwert mit unterschiedlichem Vorzeichen vom Rest hat, ist das System hyperbolisch.
- Wenn A keines der Kriterien erfüllt, ist das System ultrahyperbolisch.

Diese Kriterien sind ungefähr in der Reihenfolge des Schwierigkeitsgrads der Lösung jedes Gleichungstyps aufgeführt. Im Folgenden betrachten wir die ersten drei Fälle und geben Beispiele für entsprechendes Verhalten; Ultrahyperbolische Gleichungen kommen in der Praxis nicht so häufig vor und erfordern hochspezialisierte Techniken zu ihrer Lösung.

TODO: Reduktion auf kanonische Form über Eigenmaterial von A (nicht in 205A)

14.3.1 Elliptische PDEs

So wie positiv definite Matrizen spezielle Algorithmen wie die Cholesky-Zerlegung und konjugierte Gradienten ermöglichen, die ihre Inversion vereinfachen, haben elliptische PDEs eine besonders starke Struktur, die zu effektiven Lösungstechniken führt.

Die elliptische Modell-PDE ist die Laplace-Gleichung, gegeben durch $\nabla^2 f = g$ für eine gegebene Funktion g als im Beispiel 14.3. Beispielsweise ergibt sich für zwei Variablen die Laplace-Gleichung

$$f_{xx} + f_{yy} = g.$$

Abbildung NUMBER zeigt einige Lösungen der Laplace-Gleichung für verschiedene Auswahlmöglichkeiten von u und f in einem zweidimensionalen Bereich.

Elliptische Gleichungen haben viele wichtige Eigenschaften. Von besonderer theoretischer und praktischer Bedeutung ist die Idee der elliptischen Regularität, dass Lösungen elliptischer PDGs automatisch glatte Funktionen in $C^\infty(\bar{\Omega})$ sind. Diese Eigenschaft ist nicht sofort offensichtlich: Eine PDE zweiter Ordnung in f erfordert nur, dass f zweimal differenzierbar ist, um einen Sinn zu ergeben, aber tatsächlich sind sie unter schwachen Bedingungen automatisch unendlich differenzierbar. Diese Eigenschaft unterstützt die physikalische Intuition, dass elliptische Gleichungen stabile physikalische Gleichgewichte darstellen, wie die Ruhelage einer ausgestreckten Gummiplatte.

Auch elliptische Gleichungen zweiter Ordnung in der obigen Form lassen im Gegensatz zu PDEs in einigen anderen Formen garantiert Lösungen zu.

Beispiel 14.6 (Poisson in einer Variablen). Die Laplace-Gleichung mit $g = 0$ erhält den speziellen Namen Poisson-Gleichung. In einer Variablen kann $f(x) = 0$ geschrieben werden, was trivialerweise durch $f(x) = \frac{1}{2}x^2 + \frac{1}{2}x$ gelöst wird. Diese Gleichung reicht aus, um mögliche Randbedingungen auf $[a, b]$ zu untersuchen:

- Dirichlet-Bedingungen für diese Gleichung geben einfach $f(a)$ und $f(b)$ an; Es gibt offensichtlich eine eindeutige Linie, die durch $(a, f(a))$ und $(b, f(b))$ geht und die Lösung der Gleichung liefert.
- Neumann-Bedingungen würden $f'(a)$ und $f'(b)$ spezifizieren. Aber $f'(a) = f'(b) = \tilde{y}$ für $f(x) = \tilde{y}x + \tilde{y}$. Auf diese Weise können Grenzwerte für Neumann-Probleme Kompatibilitätsbedingungen unterliegen, die für eine Lösung erforderlich sind. Darüber hinaus hat die Wahl von $\tilde{y}$ keinen Einfluss auf die Randbedingungen, sodass die Lösung nicht eindeutig ist, wenn diese erfüllt sind.

14.3.2 Parabolische PDEs

In Anlehnung an die Struktur der linearen Algebra sind positiv semidefinite Gleichungssysteme nur geringfügig schwieriger zu handhaben als positiv definite. Insbesondere positive semidefinite Matrizen lassen einen Nullraum zu, mit dem sorgfältig umgegangen werden muss, aber in den übrigen Richtungen verhalten sich die Matrizen genauso wie im definiten Fall.

Die Wärmegleichung ist die modellparabolische PDE. Angenommen, $f(0; x, y)$ ist eine Wärmeverteilung in einem Bereich $\tilde{y} \in \mathbb{R}^2$ zum Zeitpunkt $t = 0$. Dann bestimmt die Wärmegleichung, wie die Wärme über die Zeit t als Funktion $f(t; x, y)$ diffundiert. :

$$\frac{\partial f}{\partial t} = \tilde{y}^2 f, \quad \tilde{y}^2$$

wobei $\tilde{y} > 0$ und wir uns $\tilde{y}^2$ wieder als den Laplace-Operator in den Raumvariablen x und y vorstellen, also $\tilde{y}^2 = \frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2}$. Diese Gleichung muss parabolisch sein, da vor f_{xx} und f_{yy} derselbe Koeffizient $\tilde{y}$ steht, $\frac{\partial f}{\partial t}$ jedoch nicht in die Gleichung einfließt.

Abbildung NUMBER veranschaulicht eine phänomenologische Interpretation der Wärmegleichung. Wir können uns $\tilde{y}^2 f$ als Maß für die Konvexität von f vorstellen, wie in Abbildung NUMBER(a). Somit nimmt die Wärmegleichung u mit der Zeit zu, wenn ihr Wert nach oben „erhöht“ wird, andernfalls nimmt sie ab. Diese negative Rückkopplung ist stabil und führt zu einem Gleichgewicht, wenn $t \rightarrow \infty$.

Für die Wärmegleichung sind zwei Randbedingungen erforderlich, die beide mit einhergehen einfache physikalische Interpretationen:

- Die Wärmeverteilung $f(0; x, y)$ zum Zeitpunkt $t = 0$ an allen Punkten $(x, y) \in \tilde{y}$
- Verhalten von f , wenn $t > 0$ an Punkten $(x, y) \in \tilde{y}$. Diese Randbedingungen beschreiben das Verhalten an der Grenze der Domäne. Dirichlet-Bedingungen liefern hier $f(t; x, y)$ für alle $t \geq 0$ und $(x, y) \in \tilde{y}$, entsprechend der Situation, in der ein äußerer Agent die Temperaturen an der Grenze der Domäne festlegt. Diese Bedingungen können auftreten, wenn $\tilde{y}$ ein Stück Folie ist, das neben einer Wärmequelle liegt, deren Temperatur durch die Folie nicht wesentlich beeinflusst wird, wie z. B. ein großer Kühlschrank oder Ofen. Im Gegensatz dazu geben Neumann-Bedingungen die Ableitung von f in der Richtung normal zur Grenze $\tilde{y}$ an, wie in Abbildung NUMBER; Sie entsprechen der Festlegung des Wärmeflusses aus $\tilde{y}$, der durch verschiedene Arten der Isolierung verursacht wird.

14.3.3 Hyperbolische PDEs

Die endgültige Modellgleichung ist die Wellengleichung, die dem Fall der unbestimmten Matrix entspricht:

$$\frac{\partial^2 f}{\partial t^2} - \tilde{y}^2 f = 0$$

Die Wellengleichung ist hyperbolisch, da die zweite zeitliche Ableitung das entgegengesetzte Vorzeichen der beiden räumlichen Ableitungen hat. Diese Gleichung bestimmt die Bewegung von Wellen über ein elastisches Medium wie eine Gummipatte; Es kann beispielsweise abgeleitet werden, indem das zweite Newtonsche Gesetz auf Punkte auf einem Gummistück angewendet wird, wobei x und y Positionen auf dem Blatt sind und $f(t; x, y)$ die Höhe des Gummistücks zum Zeitpunkt t ist.

Abbildung NUMBER zeigt eine eindimensionale Lösung der Wellengleichung. Das Wellenverhalten steht in erheblichem Gegensatz zur Wärmediffusion, da bei $t \rightarrow \infty$ die Energie möglicherweise nicht diffundiert. Insbesondere können Wellen innerhalb einer Domäne unbegrenzt hin und her springen. Aus diesem Grund werden wir sehen, dass implizite Integrationsstrategien möglicherweise nicht für die Integration hyperbolischer PDEs geeignet sind, da sie dazu neigen, Bewegungen zu dämpfen.

Die Randbedingungen für die Wellengleichung ähneln denen der Wärmeleitung, jedoch jetzt wir müssen sowohl $f(0; x, y)$ als auch $f_t(0; x, y)$ zum Zeitpunkt Null angeben:

- Die Bedingungen bei $t = 0$ geben die Position und Geschwindigkeit der Welle zum Anfangszeitpunkt an.
- Randbedingungen an $\bar{y}$ bestimmen, was an den Enden des Materials passiert. Dirichlet-Bedingungen entsprechen dem Fixieren der Seiten der Welle, z. B. dem Zupfen einer Cellosaite, die an ihren beiden Enden flach am Instrument gehalten wird. Neumann-Bedingungen entsprechen dem Unberührtlassen der Enden der Welle wie dem Ende einer Peitsche.

14.4 Derivate als Operatoren

In PDEs und anderswo können wir uns Ableitungen als Operatoren vorstellen, die auf Funktionen genauso einwirken wie Matrizen auf Vektoren. Unsere Wahl der Notation spiegelt oft diese Parallele wider: Die Ableitung d/dx sieht aus wie das Produkt eines Operators d/dx und einer Funktion f . Tatsächlich ist die Differentiation ein linearer Operator, genau wie die Matrixmultiplikation, da für alle $f, g: \mathbf{R} \rightarrow \mathbf{R}$ und $a, b \in \mathbf{R}$ gilt

$$\frac{d}{dx}(a f(x) + b g(x)) = a \frac{d}{dx} f(x) + b \frac{d}{dx} g(x)$$

Wenn wir PDEs für numerische Lösungen diskretisieren, können wir diese Analogie tatsächlich vollständig ausführen. Betrachten Sie beispielsweise eine Funktion f auf $[0, 1]$, die mit $n + 1$ gleichmäßig verteilten Stichproben diskretisiert wurde, wie in Abbildung NUMBER. Denken Sie daran, dass der Abstand zwischen zwei Stichproben $h = 1/n$ beträgt. In Kapitel 12 haben wir eine Näherung für die zweite Ableitung $f''(x)$ entwickelt:

$$\frac{f(x+h) - 2f(x) + f(x-h))}{h^2} = f''(x) + O(h)$$

Angenommen, unsere n Stichproben von $f(x)$ auf $[0, 1]$ sind $y_0 = f(0)$, $y_1 = f(h)$, $y_2 = f(2h)$, $\dots$, $y_n = f(nh)$.

Dann ergibt die Anwendung unserer Formel oben eine Strategie zur Annäherung von f'' an jedem Gitterpunkt:

$$f''_{jk} \approx \frac{y_{k+1} - 2y_k + y_{k-1}}{h^2}$$

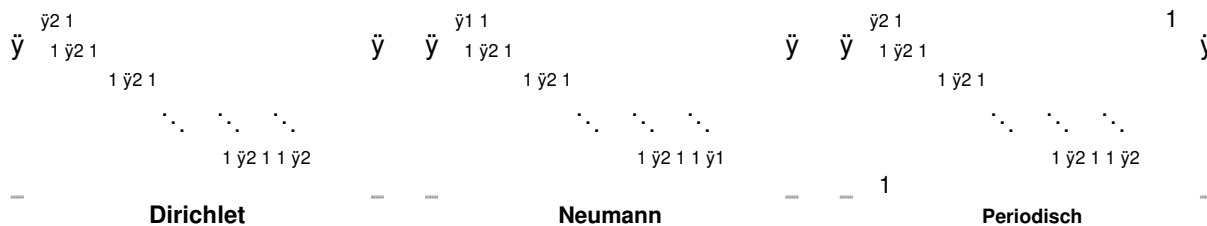
Das heißt, die zweite Ableitung einer Funktion auf einem Punktgitter kann mit der in Abbildung NUMBER(a) dargestellten 1—2-Schablone berechnet werden.

Eine Feinheit, auf die wir nicht eingegangen sind, ist, was bei y_0 und y_n passiert, da die obige Formel y_{j+1} und y_{j-1} erfordern würde. Tatsächlich kodiert diese Entscheidung die in §14.2 eingeführten Randbedingungen.

Unter Berücksichtigung der Tatsache, dass $y_0 = f(0)$ und $y_n = f(1)$, sind Beispiele für mögliche Randbedingungen für f :

- Dirichlet-Randbedingungen: $y_0 = y_{n+1} = 0$, das heißt, legen Sie einfach den Wert von y über dem fest Endpunkte
- Neumann-Randbedingungen: $y_0 = y_1$ und $y_n = y_{n+1}$, Kodierung der Randbedingung $f'(0) = f'(1) = 0$.
- Periodische Randbedingungen: $y_0 = y_n$ und $y_{n+1} = y_1$, was die Identifikation $f(0) = f(1)$ ergibt

Angenommen, wir stapeln die Stichproben y_k in einen Vektor $y \in \mathbb{R}^{n+1}$ und den Stichproben y_k in eine Sekunde Vektor $w \in \mathbb{R}^{n+1}$. Dann ist aus unserer obigen Konstruktion leicht zu erkennen, dass $h^2 w = L_1 y$, wobei L_1 eine der folgenden Optionen ist:



Das heißt, die Matrix L kann als diskretisierte Version des Operators $y \in \mathbb{R}^{n+1}$ und nicht als Funktion $\frac{d}{dx^2}$ Einwirken auf $f : [0, 1] \rightarrow \mathbb{R}$ betrachtet werden.

Wir können eine ähnliche Näherung für Δf schreiben, wenn wir $f : [0, 1] \times [0, 1] \rightarrow \mathbb{R}$ mit einem Wertegitter abtasten, wie in Abbildung NUMBER. Denken Sie insbesondere daran, dass in diesem Fall $\Delta f = f_{xx} + f_{yy}$ ist, sodass wir insbesondere die zweiten Ableitungen von x und y aufsummieren können, wie wir es im obigen eindimensionalen Beispiel getan haben. Dies führt zu einer verdoppelten 1-2-1-Schablone, wie in Abbildung NUMBER. Wenn wir unsere Stichproben als $y_k, y(kh, h)$ nummerieren, lautet unsere Formel für den Laplace-Operator von f in diesem Fall:

$$(\Delta y)_k = \frac{1}{h^2} (y(k-1, h) + y(k+1, h) + y(k, h-1) + y(k, h+1) - 4y_k)$$

Wenn wir unsere Stichproben von y und y noch einmal zu y und w kombinieren, können wir mit einer ähnlichen Konstruktion und Wahl der Randbedingungen wieder $h^2 w = L_2 y$ schreiben. Dieser zweidimensionale Gitter-Laplace-Operator L_2 erscheint in vielen Bildverarbeitungsanwendungen, in denen (k, j) zum Indizieren von Pixeln in einem Bild verwendet wird.

Nach der obigen Diskussion stellt sich natürlich die Frage, warum wir in unserer obigen Diskussion zur zweiten Ableitung des Laplace-Operators gesprungen sind, anstatt die erste Ableitung $f'(x)$ zu diskretisieren. Im Prinzip gibt es keinen Grund, warum wir nicht ähnliche Matrizen D erstellen könnten, die Vorwärts-, Rückwärts- oder symmetrische Differenznäherungen von f implementieren. Einige technische Details machen diese Aufgabe jedoch etwas schwieriger, wie weiter unten beschrieben.

Am wichtigsten ist, dass wir entscheiden müssen, welche Näherung der ersten Ableitung wir verwenden möchten. 1. als Vorwärtsdifferenz $h^{-1}(y_{k+1} - y_k)$ – Wenn wir beispielsweise y_k ($y_{k+1} - y_k$) schreiben, befinden wir uns in der unnatürlich asymmetrischen Situation, dass wir bei y_n eine Randbedingung benötigen, bei y_0 jedoch nicht. Im Gegensatz dazu könnten wir die symmetrische Differenz $(y_{k+1} - y_{k-1})/2$ verwenden, aber diese Diskretisierung leidet unter einem subtileren Zaunpfahl Problem, das in Abbildung NUMBER dargestellt ist. Insbesondere ignoriert diese Version von y_k den Wert von y_k und betrachtet nur seine Nachbarn y_{k-1} und y_{k+1} , was zu Artefakten führen kann, da jede Zeile von D nur y_k für entweder gerades oder ungerades k , aber nicht beides beinhaltet.

Wenn wir Vorwärts- oder Rückwärtsableitungen verwenden, um die Zaunpfahlprobleme zu vermeiden, verlieren wir eine Größenordnung an Genauigkeit und leiden außerdem unter den oben beschriebenen Asymmetrien. Wie bei der Leapfrog-Integration

Algorithmus in §13.4.2 besteht eine Möglichkeit, diese Probleme zu vermeiden, darin, sich die Ableitungen als auf halben Gitterpunkten lebend vorzustellen, wie in Abbildung NUMMER dargestellt. Im eindimensionalen Fall entspricht diese Änderung der Differenz an Scheitelpunkten, $\frac{1}{2\Delta x} (y_{k+1} - y_k)$ als $y_{k+1/2}$. Diese Technik, bei der verschiedene Ableitungen zur Kennzeichnung Kanten und Mittelpunkten von Gitterzellen platziert werden, kommt besonders häufig in der Flüssigkeitssimulation zum Einsatz, bei der Drücke, Flüssigkeitgeschwindigkeiten usw. an Stellen beibehalten werden, die die Berechnungen vereinfachen.

Abgesehen von diesen Feinheiten ist unsere wichtigste Schlussfolgerung aus dieser Diskussion, dass, wenn wir eine Funktion $f(x)$ diskretisieren, indem wir die Stichproben (x_i, y_i) verfolgen, die meisten vernünftigen Näherungen von Ableitungen von f als Produkt Lx für eine Matrix berechenbar sind L . Diese Beobachtung vervollständigt die Analogie: „Ableitungen wirken auf Funktionen wie Matrizen auf Vektoren.“ Oder in standardisierter Prüfungsnotation: Ableitungen: Funktionen:: Matrizen: Vektoren

14.5 PDEs numerisch lösen

Über die Theorie der PDEs bleibt noch viel zu sagen. Fragen der Existenz und Einzigartigkeit sowie der Möglichkeit, Lösungen für verschiedene PDEs zu charakterisieren, führen zu differenzierten Diskussionen unter Verwendung fortgeschrittener Aspekte der realen Analyse. Während ein vollständiges Verständnis dieser Eigenschaften erforderlich ist, um die Wirksamkeit von PDE-Diskretisierungen rigoros nachzuweisen, verfügen wir bereits über genügend Kenntnisse, um einige Techniken vorzuschlagen, die in der Praxis eingesetzt werden.

14.5.1 Elliptische Gleichungen lösen

Die meiste Arbeit zur Lösung elliptischer PDEs haben wir bereits in §14.4 erledigt. Nehmen wir insbesondere an, wir möchten eine lineare elliptische PDE der Form $Lf = g$ lösen. Hier ist L ein Differentialoperator; Um beispielsweise die Laplace-Gleichung zu lösen, würden wir $L = \nabla^2$ als Laplace-Operator nehmen. Dann haben wir in §14.4 gezeigt, dass, wenn wir f diskretisieren, indem wir eine Menge von Stichproben in einem Vektor y mit $y_i = f(x_i)$ nehmen, dann eine entsprechende Näherung von Lf für eine Matrix L als Ly geschrieben werden kann. Wenn wir auch g diskretisieren, Unter Verwendung von Stichproben in einem Vektor b wird die Lösung der elliptischen PDE $Lf = g$ durch die Lösung des linearen Systems $Ly = b$ angenähert.

Beispiel 14.7 (Elliptische PDE-Diskretisierung). Angenommen, wir möchten Lösungen für $f(x) = g(x)$ auf $[0, 1]$ mit Randbedingungen $f(0) = f(1) = 0$ approximieren. Wir werden $f(x)$ mit einem Vektor $y \in \mathbb{R}^n$ approximieren. Probenahme f wie folgt:

$$\begin{array}{ccc} y_1 & & f(h) \\ y_2 & & f(2h) \\ \vdots & & \vdots \\ y_n & & f(nh) \end{array}$$

wobei $h = 1/(n+1)$. Wir fügen keine Stichproben bei $x = 0$ oder $x = 1$ hinzu, da dort die Randbedingungen die Werte bestimmen. Wir werden b verwenden, um einen analogen Satz von Werten für $g(x)$ zu halten.

Unter Berücksichtigung unserer Randbedingungen diskretisieren wir $f(x)$ als $\frac{1}{h} \sum_{j=1}^n L y_j$, wobei L gegeben ist durch:

$$L \ddot{y} = \begin{pmatrix} \ddot{y}_1 & \ddot{y}_2 & \ddot{y}_3 & \ddot{y}_4 & \ddot{y}_5 \\ \ddot{y}_1 & 2\ddot{y}_2 & \ddot{y}_3 & \ddot{y}_4 & \ddot{y}_5 \\ \ddot{y}_2 & \ddot{y}_3 & 2\ddot{y}_4 & \ddot{y}_5 & \ddot{y}_6 \\ \ddot{y}_3 & \ddot{y}_4 & \ddot{y}_5 & 2\ddot{y}_6 & \ddot{y}_7 \\ \ddot{y}_4 & \ddot{y}_5 & \ddot{y}_6 & \ddot{y}_7 & 2\ddot{y}_8 \\ \ddot{y}_5 & \ddot{y}_6 & \ddot{y}_7 & \ddot{y}_8 & \ddot{y}_9 \end{pmatrix} \ddot{y}$$

Somit ist unsere Näherungslösung der PDE gegeben durch $y = h^{-1} L^{-1} b$.

So wie elliptische PDEs die am einfachsten zu lösenden PDEs sind, sind ihre Diskretisierungen mithilfe von Matrizen wie im obigen Beispiel am einfachsten zu lösen. Tatsächlich ist die Diskretisierung eines elliptischen Operators L im Allgemeinen eine positiv definite und dünn besetzte Matrix, die sich perfekt für die in Kapitel 10 abgeleiteten Lösungstechniken eignet.

Beispiel 14.8 (Elliptische Operatoren als Matrizen). Betrachten Sie die Matrix L aus Beispiel 14.7. Wir können zeigen, dass L negativ definit ist (und daher das positiv definite System $\ddot{y}L y = \ddot{y}h^{-1}b$ unter Verwendung konjugierter Gradienten gelöst werden kann), indem wir beachten, dass $\ddot{y}L = DD^T$ für die Matrix $D \in \mathbb{R}^{(n+1) \times n}$ gegeben durch:

$$D = \begin{pmatrix} \ddot{y}_1 & \ddot{y}_2 & \ddot{y}_3 & \ddot{y}_4 & \ddot{y}_5 \\ \ddot{y}_1 & \ddot{y}_2 & \ddot{y}_3 & \ddot{y}_4 & \ddot{y}_5 \\ \ddot{y}_2 & \ddot{y}_3 & \ddot{y}_4 & \ddot{y}_5 & \ddot{y}_6 \\ \ddot{y}_3 & \ddot{y}_4 & \ddot{y}_5 & \ddot{y}_6 & \ddot{y}_7 \\ \ddot{y}_4 & \ddot{y}_5 & \ddot{y}_6 & \ddot{y}_7 & \ddot{y}_8 \\ \ddot{y}_5 & \ddot{y}_6 & \ddot{y}_7 & \ddot{y}_8 & \ddot{y}_9 \end{pmatrix}$$

Diese Matrix ist nichts anderes als die endliche differenzierte erste Ableitung, daher entspricht diese Beobachtung der Tatsache, dass $\frac{d^2 f}{dx^2} = \frac{d}{dx} \left(\frac{df}{dx} \right)$. Somit ist $x^T L x = \ddot{y} x^T D D^T x = \ddot{y} D x^T D^T x = \ddot{y} \int_0^1 \left(\frac{df}{dx} \right)^2 dx$, was zeigt, dass L negativ semidefinit ist. Es ist leicht zu erkennen, dass $Dx = 0$ genau dann ist, wenn $x = 0$, was den Beweis vervollständigt, dass L tatsächlich negativ definit ist.

14.5.2 Parabolische und hyperbolische Gleichungen lösen

Parabolische und hyperbolische Gleichungen führen im Allgemeinen eine Zeitvariable in die Formulierung ein, die ebenfalls differenziert ist, jedoch möglicherweise in niedrigerer Ordnung. Da Lösungen parabolischer Gleichungen viele Stabilitätseigenschaften zulassen, sind numerische Techniken zum Umgang mit dieser Zeitvariablen oft stabil und gut konditioniert; Im Gegensatz dazu muss mehr Sorgfalt darauf verwendet werden, hyperbolisches Verhalten zu behandeln und eine Dämpfung der Bewegung im Laufe der Zeit zu verhindern.

Semidiskrete Methoden Der wahrscheinlich einfachste Ansatz zur Lösung einfacher zeitabhängiger Gleichungen ist die Verwendung einer semidiskreten Methode. Hier diskretisieren wir den Bereich, aber nicht die Zeitvariable, was zu einer ODE führt, die mit den Methoden aus Kapitel 13 gelöst werden kann.

Beispiel 14.9 (Semidiskrete Wärme Gleichung). Betrachten Sie die Wärme Gleichung in einer Variablen, gegeben durch $f_t = f_{xx}$, wobei $f(t; x)$ die Wärme an Position x und Zeit t darstellt. Als Grenzdaten stellt der Benutzer Folgendes bereit:

Funktion $f_0(x)$ mit $f(0; x) \approx f_0(x)$; wir verbinden auch den Rand $x \in \{0, 1\}$ mit einem Kühschrank und erzwingen damit $f(t; 0) = f(t; 1) = 0$.

Angenommen, wir diskretisieren die x -Variable, indem wir Folgendes definieren:

$$\begin{aligned} f_1(t) &\approx f(h; t) \\ f_2(t) &\approx f(2h; t) \\ &\vdots \\ f_n(t) &\approx f(nh; t), \end{aligned}$$

wobei wir wie in Beispiel 14.7 $h = 1/n+1$ nehmen und Stichproben bei $x \in \{0, 1\}$ weglassen, da sie durch die Randbedingungen bereitgestellt werden.

Durch die Kombination dieser f_i können wir $f(t) : \mathbb{R} \rightarrow \mathbb{R}^n$ als die semidiskrete Version von f definieren, bei der wir im Raum, aber nicht in der Zeit abgetastet haben. Nach unserer Konstruktion ist die semidiskrete PDE-Näherung die durch $f'(t) = L f(t)$ gegebene ODE.

Das vorherige Beispiel zeigt ein Beispiel eines sehr allgemeinen Musters für parabolische Gleichungen. Wenn wir kontinuierliche Phänomene simulieren, wie etwa Wärme, die sich über eine Domäne bewegt, oder Chemikalien, die durch eine Membran diffundieren, gibt es normalerweise eine Zeitvariable und dann mehrere räumliche Variablen, die auf elliptische Weise differenziert sind. Wenn wir dieses System semidiskret diskretisieren, können wir ODE-Integrationsstrategien für ihre Lösung verwenden. Genauso wie die Matrix, die zur Lösung einer linearen elliptischen Gleichung wie in §14.5.1 verwendet wird, im Allgemeinen positiv oder negativ definit ist, ist die Matrix L normalerweise negativ definit, wenn wir eine semidiskrete parabolische PDE $f' = L f$ schreiben. Diese Beobachtung impliziert, dass die Lösung dieser kontinuierlichen ODE durch f bedingungslos stabil ist, da negative Eigenwerte mit der Zeit gedämpft werden.

Wie im vorherigen Kapitel dargelegt, haben wir aufgrund einer räumlichen Diskretisierung viele Möglichkeiten, die ODE zeitlich zu lösen. Wenn die Zeitschritte klein und begrenzt sind, können explizite Methoden akzeptabel sein. Implizite Löser werden häufig zur Lösung parabolischer PDGs eingesetzt; Das diffusive Verhalten des impliziten Eulers kann zu Ungenauigkeiten führen, aber verhaltenstechnisch gesehen ähnelt es der durch die Wärme Gleichung bereitgestellten Diffusion und kann selbst bei relativ großen Zeitschritten akzeptabel sein. Hyperbolische PDEs erfordern möglicherweise implizite Schritte für die Stabilität, aber fortgeschrittene Integratoren wie „symplektische Integratoren“ können eine durch diese Art von Schritten verursachte Überglättung verhindern.

Ein kontrastierender Ansatz besteht darin, Lösungen semidiskreter Systeme $f' = L f$ in Form von Eigenvektoren von L zu schreiben. Angenommen, $v_1, \dots, v_n$ sind Eigenvektoren von L mit den Eigenwerten $\gamma_1, \dots, \gamma_n$ und dass wir $f(0) = c_1 v_1 + \dots + c_n v_n$ wissen. Denken Sie dann daran, dass die Lösung von $f' = L f$ gegeben ist durch:

$$f(t) = \sum_{i=1}^n c_i v_i e^{\gamma_i t}$$

Diese Formel ist nichts Neues über §5.1.2 hinaus, den wir während unserer Diskussion über Eigenvektoren und Eigenwerte eingeführt haben. Die Eigenvektoren von L können jedoch im Fall einer semidiskreten PDE eine physikalische Bedeutung haben, wie in Beispiel 14.5, das zeigte, dass Eigenvektoren der Laplace-Operatoren L verschiedenen Resonanzschwingungen der Domäne entsprechen. Somit kann dieser Eigenvektoransatz angewendet werden, um beispielsweise „Niederfrequenznäherungen“ der Anfangswertdaten zu entwickeln, indem die obige Summe über i gekürzt wird, mit dem Vorteil, dass die t -Abhängigkeit ohne Zeitschritt genau bekannt ist.

Beispiel 14.10 (Eigenfunktionen des Laplace-Operators). Abbildung NUMBER zeigt Eigenvektoren der Matrix L aus Beispiel 14.7. Eigenvektoren mit niedrigen Eigenwerten entsprechen Niederfrequenzfunktionen auf $[0, 1]$ mit an den Endpunkten festgelegten Werten und können gute Annäherungen an $f(x)$ sein, wenn es relativ glatt ist.

Vollständig diskrete Methoden Alternativ könnten wir die Raum- und Zeitvariablen demokratischer behandeln und beide gleichzeitig diskretisieren. Diese Strategie ergibt ein zu lösendes Gleichungssystem, das eher §14.5.1 ähnelt. Diese Methode lässt sich leicht formulieren, indem man den elliptischen Fall parallelisiert, aber die resultierenden linearen Gleichungssysteme können groß sein, wenn die Abhängigkeit zwischen Zeitschritten eine globale Reichweite hat.

Beispiel 14.11 (Vollständig diskrete Wärmediffusion). Explizit, implizit, Crank-Nicolson. Nicht in CS 205A abgedeckt.

Es ist wichtig zu beachten, dass letztendlich sogar semidiskrete Methoden als vollständig diskret betrachtet werden können, da die zeitschrittbasierte ODE-Methode immer noch die t -Variable diskretisiert; Der Unterschied besteht hauptsächlich in der Klassifizierung der Art und Weise, wie Methoden abgeleitet wurden. Ein Vorteil semidiskreter Techniken besteht jedoch darin, dass sie den Zeitschritt für t abhängig von der aktuellen Iteration anpassen können. Wenn sich Objekte beispielsweise in einer physikalischen Simulation schnell bewegen, kann es sinnvoll sein, mehr Zeitschritte zu verwenden, um diese Bewegung aufzulösen. Einige Methoden passen sogar die Diskretisierung des Bereichs der x -Werte an, falls in der Nähe lokaler Diskontinuitäten oder anderer Artefakte eine höhere Auflösung erforderlich ist.

14.6 Methode der Finiten Elemente

Nicht in 205A abgedeckt.

14.7 Beispiele aus der Praxis

Anstelle einer strengen Behandlung aller häufig verwendeten PDE-Techniken stellen wir in diesem Abschnitt Beispiele dafür vor, wo sie in der Informatik in der Praxis vorkommen.

14.7.1 Bildverarbeitung im Gradientenbereich

14.7.2 Kantenerhaltende Filterung

14.7.3 Gitterbasierte Flüssigkeiten

14.8 Zu erledigen

- Mehr zum Thema Existenz/Einzigartigkeit
- CFL-Bedingungen
- Laxer Äquivalenzsatz
- Beständigkeit, Stabilität und Freunde

14.9 Probleme

- Zeigen Sie, dass der 1d-Laplace-Operator als DD^T für die erste Ableitungsmatrix D faktorisiert werden kann
- Lösen Sie die PDE erster Ordnung

Danksagungen

[Diskussion geht hier]

Ich schulde den Studenten des Stanford CS 205A, Herbst 2013, viel Dank dafür, dass sie bei der Entwicklung dieses Buches zahlreiche Tippfehler und Fehler entdeckt haben. Das Folgende ist zweifellos eine unvollständige Liste der Studenten, die zu diesem Projekt beigetragen haben: Tao Du, Lennart Jansson, Miles Johnson, Luke Knepper, Minjae Lee, Nisha Masharani, John Reyna, William Song, Ben-Han Sung, Martina Troesch, Ozhan Turgut, Patrick Ward, Joongyeub Yeo und Yang Zhao.

Besonderer Dank geht an Jan Heiland und Tao Du für ihre Hilfe bei der Klärung der Ableitung des BFGS-Algorithmus.

[Weitere Diskussion hier]