

Ingenieurmathematik

Skript zur Vorlesung im Master-Studiengang Geodäsie

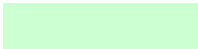
Prof. Dr. Martin Rumpf, Dr. Martin Lenz, Dr. Patrick Penzler

Version vom 16. Dezember 2013

Inhaltsverzeichnis

0	Einleitung	5
1	Finite Differenzen	7
1.1	Finite Differenzen in 1D	7
1.2	Konvergenztheorie für Differenzenverfahren	10
1.3	Ein zweidimensionales Beispiel	13
1.4	Totale Variation	17
2	Finite Elemente	29
2.1	Schwache Lösungen	29
2.2	Approximation durch Finite Elemente	33
2.3	Randwerte	37
2.4	Finite Elemente auf Dreiecksgittern	39
2.5	Assemblierung der Matrizen	41
2.6	Randwerte	45

Zur besseren Übersicht werden in diesem Skript folgende Farben verwendet:

Satz, Lemma, Folgerung	
Definition, Notation	
Schema	
Problemstellung	
Programmieraufgabe	

0 Einleitung

Thema: Numerische Lösungsverfahren für partielle Differentialgleichungen

Modellgleichung:

$$\frac{\partial^2}{\partial x^2}u(x, y) + \frac{\partial^2}{\partial y^2}u(x, y) = f(x, y)$$

Gesucht ist u , so dass die Differentialgleichung auf einem Gebiet Ω erfüllt ist für eine gegebene Funktion f und bestimmte Randwerte für u auf $\partial\Omega$. Diese sogenannte „Poisson-Gleichung“ tritt beispielsweise in der Elektrodynamik bei der Datenglättung oder der Modellierung eines Schwerfelds auf.

Notation 0.1. *Wir schreiben*

$$\frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} =: \partial_{xx} + \partial_{yy} =: \Delta.$$

Diskretisierungsmethoden: Folgende Diskretisierungsmethoden werden in der Vorlesung behandelt:

- Finite Differenzen
- Finite Elemente
- Randelemente

Algorithmische Umsetzung: MATLAB

Literatur:

- Dietrich Braess: „*Finite Elemente*“ (Kapitel I & II), Springer.
- Wolfgang Dahmen und Arnold Reusken: „*Numerik für Ingenieure und Naturwissenschaftler*“ (Kapitel 12), Springer.
- Lothar Gaul, Martin Kögl und Marcus Wagner: „*Boundary Element Methods for Engineers and Scientists*“ (Kapitel 4), Springer.

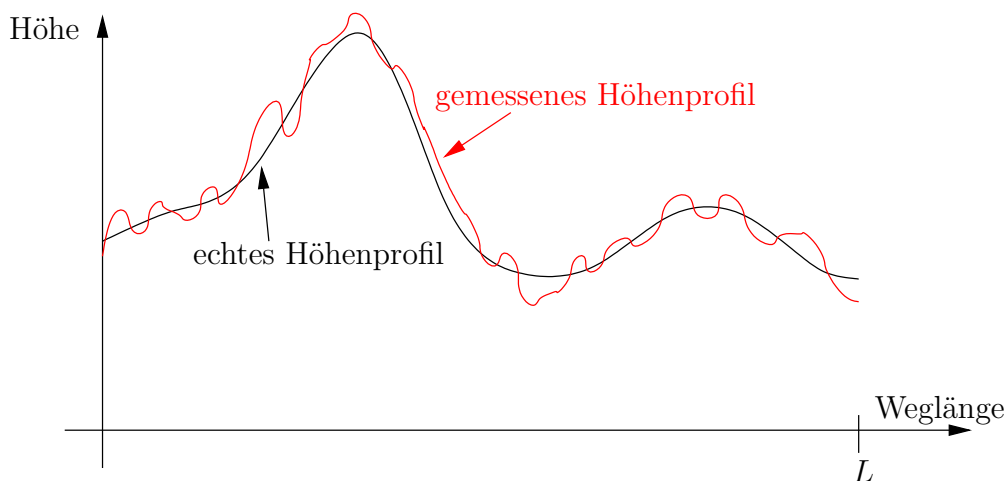
1 Finite Differenzen

1.1 Finite Differenzen in 1D

Modellproblem: Datenglättung für Funktionen $f : \mathbb{R} \rightarrow \mathbb{R}$

Als Beispiel wollen wir gemessene Höhendaten entlang eines Weges betrachten. Die Höhenwerte mögen dabei starkes Rauschen (z.B. durch schlechten GPS-Empfang) enthalten.

Wie erhält man daraus ein „vernünftiges“ Höhenprofil?



Einfacher Ansatz: Interpretiere die Höhendaten als *kontinuierliche* Funktion

$$f : [0, L] \rightarrow \mathbb{R}.$$

Gesucht ist nun eine Funktion

$$u : [0, L] \rightarrow \mathbb{R},$$

die die wesentlichen Eigenschaften von f widerspiegelt, aber weniger Rauschen enthält.

Ansatz: Was wollen wir nicht?

- Die Funktion u soll nicht weit weg von f sein, d.h. $|u(x) - f(x)|$ soll klein sein.
- u soll nicht verrauscht sein, d.h. $|u'(x)|$ soll klein sein.

Dies führt uns auf das folgende Funktional, das wir minimieren wollen:

$$E[u] := \int_0^L (u(x) - f(x))^2 + \beta (u'(x))^2 dx, \quad \beta > 0 \text{ konstant.}$$

Erinnerung: Hat eine differenzierbare Funktion ein Minimum, so ist die Ableitung dort 0.

Wie sieht nun die Ableitung von E nach u aus? Wir betrachten hierzu die Richtungsableitung $\left. \frac{d}{dt} E[u + tv] \right|_{t=0}$ von E in Richtung einer differenzierbaren Funktion $v : [0, L] \rightarrow \mathbb{R}$ für die $v(0) = v(L) = 0$ gilt. Dann gilt:

$$\begin{aligned}
 0 &= \left. \frac{d}{dt} E[u + tv] \right|_{t=0} \\
 &= \left. \frac{d}{dt} \int_0^L (u(x) + tv(x) - f(x))^2 + \beta(u'(x) + tv'(x))^2 dx \right|_{t=0} \\
 &= \left. \int_0^L \frac{d}{dt} (u(x) + tv(x) - f(x))^2 + \beta \frac{d}{dt} (u'(x) + tv'(x))^2 dx \right|_{t=0} \\
 &\stackrel{\text{Kettenregel}}{=} \left. \int_0^L 2(u(x) + tv(x) - f(x))v(x) + 2\beta(u'(x) + tv'(x))v'(x) dx \right|_{t=0} \\
 &= \int_0^L 2(u(x) - f(x))v(x) + 2\beta u'(x)v'(x) dx \\
 &\stackrel{\text{part. Int.}}{=} \left. \int_0^L 2(u(x) - f(x))v(x) dx - \int_0^L 2\beta u''(x)v(x) dx + u'(x)v(x) \right|_{x=0}^{x=L} \\
 &= 2 \int_0^L (u(x) - f(x) - \beta u''(x))v(x) dx
 \end{aligned}$$

Die Richtungsableitung muss im Minimum für alle Richtungen 0 ergeben, d.h. es muss gelten:

$$2 \int_0^L (u(x) - f(x) - \beta u''(x))v(x) dx = 0$$

für alle differenzierbaren $v : [0, L] \rightarrow \mathbb{R}$ mit $v(0) = v(L) = 0$. Daraus ergibt sich

$$u(x) - f(x) - \beta u''(x) = 0 \quad \text{für alle } x \in (0, L).$$

Was passiert am Rand?

Einfacher Fall: Wir kennen die Höhen am Start- und Zielpunkt sehr genau. Dann können wir diese einfach vorschreiben: $u(0) = f(0)$, $u(L) = f(L)$. Insgesamt erhalten wir:

Problemstellung 1.1. Gegeben sei eine stetige Funktion $f : [0, L] \rightarrow \mathbb{R}$ und $\beta > 0$. Wir suchen eine zweimal stetig differenzierbare Funktion $u : [0, L] \rightarrow \mathbb{R}$ mit

$$\begin{aligned}
 -\beta u''(x) + u(x) &= f(x), & x \in (0, L) \\
 u(0) &= f(0), \\
 u(L) &= f(L).
 \end{aligned} \tag{P}$$

Wie findet man u näherungsweise?

Wir approximieren u auf einem Gitter der Schrittweite $h = \frac{L}{N-1}$. Die Funktion $u : [0, L] \rightarrow \mathbb{R}$ wird dann durch Werte an den Knoten $x_i = (i-1) \cdot h$ approximiert. $U \in \mathbb{R}^N$ ist der Vektor mit den Knotenwerten, $U_i \approx u(x_i)$.

Approximation der Ableitung? Hierbei dürfen nur Werte an den Knoten verwendet werden! Idee: Differenzenquotienten (vgl. Ingenieur-Mathematik I, II Skript, Abschnitt 7.2)

Problemstellung 1.2. Gegeben sei eine stetige Funktion $f : [0, L] \rightarrow \mathbb{R}$ und $\beta > 0$. Sei $N \in \mathbb{N}$ und $h = L/(N-1)$. Wir suchen $U \in \mathbb{R}^N$ mit

$$\begin{aligned} -\frac{\beta}{h^2}(U_{i-1} - 2U_i + U_{i+1}) + U_i &= f(x_i), \quad i = 2, \dots, N-1, \\ U_1 &= f(x_1), \\ U_N &= f(x_N). \end{aligned} \quad (P_h)$$

Dies ist ein lineares Gleichungssystem für $U \in \mathbb{R}^N$. In Matrixschreibweise ergibt sich die folgende Darstellung.

Schema 1.3. (Matrixdarstellung von (P_h))

$$\left(-\frac{\beta}{h^2} \begin{pmatrix} 0 & & & 0 \\ 1 & -2 & 1 & \\ & \ddots & \ddots & \ddots \\ & & 1 & -2 & 1 \\ 0 & & & & 0 \end{pmatrix} + \begin{pmatrix} 1 & & & 0 \\ & 1 & & \\ & & \ddots & \\ & & & 1 \\ 0 & & & & 1 \end{pmatrix} \right) \begin{pmatrix} U_1 \\ U_2 \\ \vdots \\ U_{N-1} \\ U_N \end{pmatrix} = \begin{pmatrix} f(x_1) \\ f(x_2) \\ \vdots \\ f(x_{N-1}) \\ f(x_N) \end{pmatrix}.$$

Der zugehörige MATLAB-Code sieht folgendermaßen aus:

Schema 1.4. (MATLAB-Code zu (P_h))

```
e = ones (N, 1);
% Matrix for second derivative
A = spdiags ([e -2*e e], [-1 0 1], N, N);
% First and last line
A (1, :) = zeros (1, N);
A (N, :) = zeros (1, N);
% Identity matrix
B = spdiags (e, 0, N, N);
% Complete matrix
L = - beta / (h*h) * A + B;
```

Programmieraufgabe 1. Implementieren Sie das beschriebene Finite-Differenzen-Verfahren zur Datenglättung in **MATLAB**. Experimentieren Sie mit dem Parameter β .

1.2 Konvergenztheorie für Differenzenverfahren

(P_h) führt auf ein lineares Gleichungssystem. Wir wenden uns nun den folgenden Fragen zu:

1. Ist dieses stets *eindeutig lösbar*?
2. Konvergiert die Lösung von (P_h) für $N \rightarrow \infty$ (bzw. $h \rightarrow 0$) gegen die Lösung von (P) ? Wenn ja, wie schnell?

Dazu benötigen wir zunächst das folgende Lemma, das sowohl für die kontinuierliche Lösung von (P) als auch für die diskrete Lösung von (P_h) gilt. Um die Analogie herauszuarbeiten, beweisen wir beide Varianten.

Lemma 1.5 (diskretes Maximumprinzip).
Sei U Lösung von (P_h) . Dann ist

$$\begin{aligned}\max_{i=1,\dots,N} U_i &\leq \max_{i=1,\dots,N} f(x_i), \\ \min_{i=1,\dots,N} U_i &\geq \min_{i=1,\dots,N} f(x_i).\end{aligned}$$

Beweis.

Sei $j \in \{1, \dots, N\}$ Maximalstelle von U . Falls $j = 1$ oder $j = N$, so ist

$$\max_{i=1,\dots,N} U_i = U_j = f(x_j) \leq \max_{i=1,\dots,N} f(x_i).$$

Falls $j \in \{2, \dots, N-1\}$, so gilt (j ist Maximalstelle)

$$\begin{aligned}U_{j-1} &\leq U_j, \\ U_{j+1} &\leq U_j.\end{aligned}$$

Daraus folgt

$$\begin{aligned}U_{j-1} - 2U_j + U_{j+1} \leq 0 &\Rightarrow -\frac{\beta}{n^2}(U_{j-1} - 2U_j + U_{j+1}) \geq 0 \\ &\Rightarrow f(x_j) - U_j \geq 0 \\ &\Rightarrow U_j \leq f(x_j),\end{aligned}$$

wobei wie oben $f(x_j) \leq \max_{i=1,\dots,N} f(x_i)$.

Für das Minimum betrachte $\tilde{U} = -U$, $\tilde{f} = -f$. □

Folgerung 1.6. *Das lineare Gleichungssystem zu (P_h) ist stets eindeutig lösbar.*

Beweis. Das lineare Gleichungssystem ist eindeutig lösbar genau dann, wenn seine Matrix regulär ist. Die Matrix ist regulär, wenn das homogene Gleichungssystem (mit rechter Seite Null) nur die Lösung $U = 0$ hat.

Betrachten wir also das homogene Gleichungssystem d. h. für alle $i = 1, \dots, N$ ist $f(x_i) = 0$. Es gilt also

$$\max_{i=1, \dots, N} f(x_i) = \min_{i=1, \dots, N} f(x_i) = 0.$$

Nach dem Maximumprinzip folgt dann

$$0 \leq \min_{i=1, \dots, N} U_i \leq \max_{i=1, \dots, N} U_i \leq 0,$$

d.h. $U_i = 0$ für alle $i = 1, \dots, N$. □

Kommen wir nun zur zweiten Frage. Dazu benötigen wir vorweg einige Definitionen, die uns viel Schreibarbeit sparen werden.

Definition 1.7 (Diskreter Differentialoperator).

Sei $\Omega = (0, L)$ und $u : \Omega \rightarrow \mathbb{R}$. Wir schreiben

$$L^\beta u(x) := -\beta u''(x) + u(x).$$

Sei weiter $N \in \mathbb{N}$, $h = \frac{L}{N-1}$ und

$$\Omega_h := \{x_i = (i-1)h \mid i = 1, \dots, N\}$$

bezeichne die Menge der Gitterpunkte. Dann heißt $L_h^\beta u : \Omega_h \rightarrow \mathbb{R}$,

$$(L_h^\beta u)(x_i) := -\frac{\beta}{h^2}(u(x_{i-1}) - 2u(x_i) + u(x_{i+1})) + u(x_i),$$

diskreter Differentialoperator zu L^β . Analog schreiben wir für $U \in \mathbb{R}^N$

$$(L_h^\beta U)_i := -\frac{\beta}{h^2}(U_{i-1} - 2U_i + U_{i+1}) + U_i.$$

Mit diesen Bezeichnungen können wir nun den Begriff der *Konsistenz* definieren. Dies bezeichnet die Approximationsqualität der Ableitungen durch den gewählten Diskretisierungsansatz.

Satz 1.8 (Konsistenz). *Sei u auf Ω viermal stetig differenzierbar. Dann gilt*

$$\max_{i=2, \dots, N-1} |L^\beta u(x_i) - (L_h^\beta u)(x_i)| \leq Ch^2$$

(mit einer Konstanten C unabhängig von h). Ein solches Verfahren heißt konsistent der Ordnung 2.

1 Finite Differenzen

Beweis.

$$\begin{aligned} | -\beta u''(x_i) + u(x_i) - (L_h^\beta u)(x_i) | &= | \beta | - u''(x_i) + \frac{1}{h^2}(u(x_{i-1}) - 2u(x_i) + u(x_{i+1})) | \\ &\leq \beta Ch^2. \end{aligned}$$

Dies folgt aus den Überlegungen anhand des Satzes von Taylor, die wir in Abschnitt 1.1 angestellt haben. $\square$

Satz 1.9 (Stabilität, stetige Abhängigkeit von der rechten Seite).

Sei $U \in \mathbb{R}^N$ Lösung von (P_h) . Dann gilt

$$\max_{i=1,\dots,N} |U_i| \leq C \max_{i=1,\dots,N} |f(x_i)|$$

(mit einer Konstanten C unabhängig von h). Ein solches Verfahren heißt stabil.

Beweis. Aus dem Maximumprinzip folgt

$$\begin{aligned} \max_{i=1,\dots,N} U_i &\leq \max_{i=1,\dots,N} f(x_i), \\ \max_{i=1,\dots,N} -U_i = -\min_{i=1,\dots,N} U_i &\leq -\min_{i=1,\dots,N} f(x_i) = \max_{i=1,\dots,N} -f(x_i). \end{aligned}$$

Also erhalten wir

$$\max_{i=1,\dots,N} |U_i| \leq \max_{i=1,\dots,N} |f(x_i)| \quad (\text{hier } C = 1).$$

$\square$

Aus diesen beiden Eigenschaften können wir nun die Konvergenz des Verfahrens herleiten. Ein wichtiger **Merksatz** ist

Konsistenz + Stabilität $\Rightarrow$ Konvergenz

Satz 1.10 (Konvergenz). Sei u Lösung von (P) und viermal stetig differenzierbar, sei U Lösung von (P_h) . Dann gilt

$$\max_{i=1,\dots,N} |u(x_i) - U_i| \leq Ch^2.$$

Ein solches Verfahren heißt konvergent zweiter Ordnung.

Beweis. Zunächst zeigen wir, dass der Fehler $e : \Omega_h \rightarrow \mathbb{R}$, $e(x_i) := U_i - u(x_i)$ eine Variante des diskreten Problems (P_h) mit einer speziellen rechten Seite erfüllt:

$$\begin{aligned} (L_h^\beta e)(x_i) &= (L_h^\beta U)_i - (L_h^\beta u)(x_i) \\ &= f(x_i) - (L_h^\beta u)(x_i) \\ &= -\beta u''(x_i) + u(x_i) - (L_h^\beta u)(x_i) \\ &= (L_h^\beta u)(x_i) - (L_h^\beta u)(x_i) =: r(x_i) \quad \text{für } i = 2, \dots, N-1. \end{aligned}$$

Setzt man zusätzlich $r(x_1) := r(x_N) := 0$, so löst e tatsächlich das Problem (P_h) mit r anstelle von f . Da unser Verfahren *stabil* ist, wissen wir also

$$\max_{i=1,\dots,N} |e(x_i)| \leq C \max_{i=1,\dots,N} |r(x_i)|$$

und aus der *Konsistenz* erhalten wir für $i = 2, \dots, N-1$

$$|r(x_i)| \leq Ch^2.$$

An den Randpunkten gilt dies sowieso, da r dort gleich Null ist. Zusammen folgt somit

$$\max_{i=1,\dots,N} |e(x_i)| \leq Ch^2.$$

□

Dabei verwenden wir die Bezeichnung C für alle von h unabhängigen Konstanten und verzichten auf eine Unterscheidung.

1.3 Ein zweidimensionales Beispiel

Als Beispiel betrachten wir das Entrauschen von Luftbildern. Dabei interpretieren wir ein quadratisches Graustufen-Bild als kontinuierliche Funktion

$$f : \bar{\Omega} \rightarrow [0, 1] \quad \text{mit} \quad \Omega = (0, 1)^2, \bar{\Omega} = [0, 1]^2.$$

Dabei ordnen wir beispielsweise schwarz dem Wert 0 und weiß dem Wert 1 zu. Wir setzen wieder an

$$E[u] := \int_{\Omega} \beta |\nabla u(x, y)|^2 + (u(x, y) - f(x, y))^2 \, dx \, dy$$

und wollen wie im eindimensionalen

$$\left. \frac{d}{dt} E[u + tv] \right|_{t=0} = 0$$

für alle differenzierbaren v mit $v = 0$ auf $\partial\Omega$ berechnen.

Bemerkung: Im Folgenden verwenden wir stets $|\cdot|$ für die euklidische Norm im $\mathbb{R}^n$. Die Schreibweise $\|\cdot\|$ ist für Normen auf Funktionenräumen (vgl. Definition 2.3) reserviert. Zunächst erinnern wir uns an einige Definitionen und Resultate aus der Analysis. Dazu sei im Folgenden $u : \mathbb{R}^2 \rightarrow \mathbb{R}$ stetig differenzierbare Funktion und $F : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ stetig differenzierbares Vektorfeld.

Notation 1.11 (Gradient, Divergenz, Laplace). *Wir bezeichnen den Gradienten von u mit*

$$\begin{pmatrix} \partial_x u \\ \partial_y u \end{pmatrix} =: \nabla u = \text{grad } u.$$

Wir unterscheiden insbesondere nicht zwischen ∇u und $\text{grad } u$.

Die Divergenz von F ist gegeben durch

$$\partial_x F_1 + \partial_y F_2 =: \operatorname{div} F.$$

Schließlich schreiben wir

$$\Delta u := \operatorname{div} \nabla u = \partial_{xx} u + \partial_{yy} u.$$

Satz 1.12 (Satz von Gauß). Sei Ω eine offene und beschränkte Teilmenge von $\mathbb{R}^2$ mit stückweise glattem Rand. Dann gilt

$$\int_{\Omega} \operatorname{div} F = \int_{\partial\Omega} F \cdot \nu,$$

wobei ν äußere Normale an Ω ist.

Zur Erinnerung (vgl. Ingenieur-Mathematik I, II Skript, Abschnitt 11.7)

Produktregel Es ist $\operatorname{div}(uF) = \nabla u \cdot F + u(\operatorname{div} F)$, denn

$$\begin{aligned} \operatorname{div}(uF) &= \sum_{i=1}^2 \partial_{x_i}(uF)_i = \sum_{i=1}^2 \partial_{x_i}(uF_i) = \sum_{i=1}^2 (\partial_{x_i} u F_i + u \partial_{x_i} F_i) \\ &= \sum_{i=1}^2 F_i \partial_{x_i} u + u \sum_{i=1}^2 \partial_{x_i} F_i = F \cdot \nabla u + u(\operatorname{div} F). \end{aligned}$$

Nach Satz von Gauß folgt dann

$$\int_{\Omega} \nabla u \cdot F + \int_{\Omega} u(\operatorname{div} F) = \int_{\Omega} \operatorname{div}(uF) = \int_{\partial\Omega} uF \cdot \nu.$$

Ersetzen wir F durch ∇v , so erhalten wir

$$\int_{\Omega} \nabla u \cdot \nabla v = - \int_{\Omega} u(\Delta v) + \int_{\partial\Omega} u \nabla v \cdot \nu.$$

Nun können wir die “Richtungsableitungen” von E berechnen

$$\begin{aligned} \frac{d}{dt} E[u + tv] \Big|_{t=0} &= \frac{d}{dt} \int_{\Omega} \beta |\nabla u + t \nabla v|^2 + (u + tv - f)^2 \Big|_{t=0} \\ &= \int_{\Omega} 2\beta (\nabla u + t \nabla v) \cdot \nabla v + 2(u + tv - f)v \Big|_{t=0} \\ &= 2\beta \int_{\Omega} \nabla u \cdot \nabla v + 2 \int_{\Omega} (u - f)v \\ &= -2\beta \int_{\Omega} \Delta u v + 2\beta \int_{\partial\Omega} (\nabla u \cdot \nu) \underbrace{v}_{=0} - 2 \int_{\Omega} (f - u)v \\ &= 2 \int_{\Omega} (-\beta \Delta u + u - f)v \end{aligned}$$

Damit dieses Integral für alle v mit $v = 0$ auf $\partial\Omega$ verschwindet, muss gelten

$$-\beta\Delta u + u = f \quad \text{in } \Omega.$$

Wir sagen $-\beta\Delta u + u - f$ ist die *erste Variation* von E .

Analog zum 1D-Fall nehmen wir hier ebenfalls an, dass auf dem Rand die Werte von u und f übereinstimmen sollen:

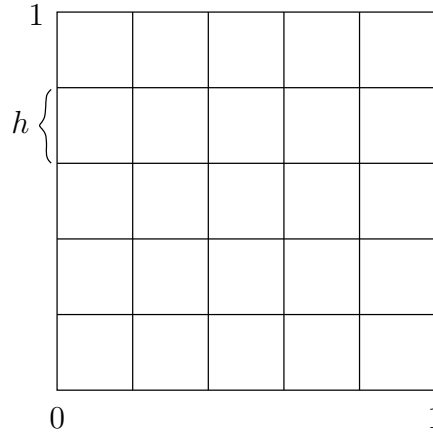
Problemstellung 1.13. Es seien $\Omega = (0, 1)^2$, $\bar{\Omega} = [0, 1]^2$ und $\partial\Omega = \text{Rand von } \Omega$. Ferner sei $f : \bar{\Omega} \rightarrow \mathbb{R}$ stetig. Gesucht ist eine zweimal stetig differenzierbare Funktion $u : \bar{\Omega} \rightarrow \mathbb{R}$ mit

$$\begin{aligned} -\beta\Delta u(x, y) + u(x, y) &= f(x, y) \quad \text{für } (x, y) \in \Omega, \\ u(x, y) &= f(x, y) \quad \text{für } (x, y) \in \partial\Omega. \end{aligned} \quad (P^2)$$

Bemerkung 1.14. Aus dem Maximumprinzip, welches analog in diesem Fall gilt, folgt

$$f : \bar{\Omega} \rightarrow [0, 1] \quad \Rightarrow \quad u : \bar{\Omega} \rightarrow [0, 1].$$

Nun approximieren wir die kontinuierliche Funktion u wiederum auf einem Gitter mit Gitterweite $h = \frac{1}{N-1}$, dabei soll jeder Gitterpunkt einem Pixel des ursprünglichen Bildes entsprechen.



Wir definieren die Knoten $(x_i, y_j) = ((i-1)h, (j-1)h)$ mit den Knotenwerten $u(x_i, y_j)$. Ein Bild bestehe aus $N \times N$ Pixeln und U_{ij} entspreche dem Grauwert des Pixels (i, j) . Wie approximieren wir nun $\Delta u(x, y)$? Im Kontinuierlichen ist

$$\Delta u(x, y) = \partial_{xx}u(x, y) + \partial_{yy}u(x, y),$$

also die Summe aus den zweiten Richtungsableitungen. Zur Approximation benutzen wir die bekannte eindimensionale Approximation der Richtungsableitungen:

$$\begin{aligned} \partial_{xx}u(x_i, y_j) &\approx \frac{1}{h^2} (u(x_{i-1}, y_j) - 2u(x_i, y_j) + u(x_{i+1}, y_j)), \\ \partial_{yy}u(x_i, y_j) &\approx \frac{1}{h^2} (u(x_i, y_{j-1}) - 2u(x_i, y_j) + u(x_i, y_{j+1})). \end{aligned}$$

Zusammen ergibt sich für den diskreten Laplace-Operator

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

Man spricht bei dieser Diskretisierung auch von dem *5-Punkte-Stern*

$$\frac{1}{h^2} \begin{bmatrix} & 1 & \\ 1 & -4 & 1 \\ & 1 & \end{bmatrix}.$$

Wie im Eindimensionalen erhält man Konsistenz der Ordnung 2.

Diese Diskretisierung ergibt wie in 1D ein lineares Gleichungssystem.

Schema 1.15. (Matrixdarstellung von (P^2))

$$\left(-\frac{\beta}{h^2} \begin{pmatrix} 0 & & & 0 \\ B & D & B & \\ & \ddots & \ddots & \ddots \\ & & B & D & B \\ 0 & & & & 0 \end{pmatrix} + \begin{pmatrix} 1 & & & 0 \\ & \ddots & & \\ & & \ddots & \\ 0 & & & \ddots & 1 \end{pmatrix} \right) \begin{pmatrix} U_1 \\ \vdots \\ U_{N^2} \end{pmatrix} = \begin{pmatrix} F_1 \\ \vdots \\ F_{N^2} \end{pmatrix}$$

$$\text{wobei } D = \begin{pmatrix} 0 & & & 0 \\ 1 & -4 & 1 & \\ & \ddots & \ddots & \ddots \\ & & 1 & -4 & 1 \\ 0 & & & & 0 \end{pmatrix} \in \mathbb{R}^{N \times N}, \quad B = \begin{pmatrix} 0 & & & \\ & 1 & & \\ & & \ddots & \\ & & & 1 \\ & & & & 0 \end{pmatrix} \in \mathbb{R}^{N \times N}$$

und $F_{N(j-1)+i} = f((i-1)h, (j-1)h)$.

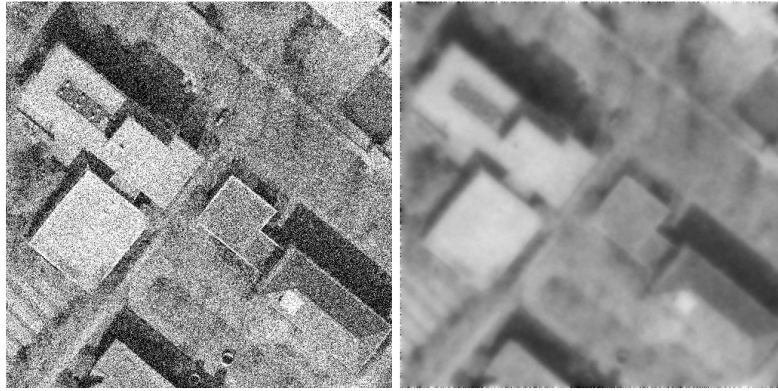
Man sieht dass in jeder Zeile der ersten Matrix (sofern sie zu einem inneren Punkt gehört) die Einträge 1, 1, -4, 1, 1 auftauchen, wobei die -4 auf der Diagonalen steht, die Einsen im Abstand von 1 bzw. N links bzw. rechts der Diagonalen.

Lediglich die Zeilen, die zu Randknoten gehören (d.h. $1, \dots, N$ (unten); $(N-1)N+1, \dots, N^2$ (oben); $1, N+1, 2N+1, \dots, (N-1)N+1$ (links) und $N, 2N, 3N, \dots, N^2$ (rechts)) sind leer.

Anhand dieser Überlegungen kann man die Matrix wie in 1D mit `spdiags` aufbauen und anschließend die entsprechenden Zeilen auf Null setzen.

Programmieraufgabe 2. Implementieren Sie das beschriebene Finite-Differenzen Verfahren zur Bildglättung in **MATLAB**.

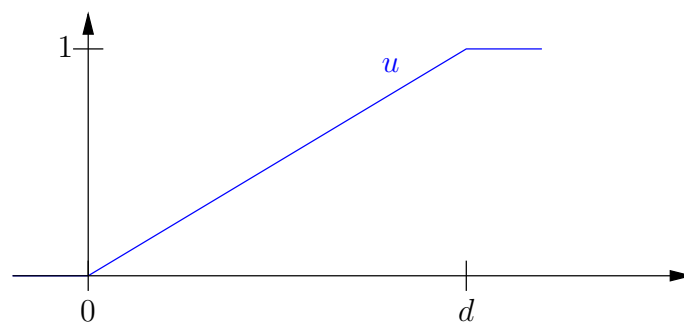
Als Ergebnis für $\beta = 10^{-4}$ erhalten wir



Beobachtung: Das Bild ist kaum noch verrauscht, aber “verschwommen”. Woran liegt das?

1.4 Totale Variation

Wir haben gesehen, dass die gewählte Energie dazu führt, dass Bilder stark geglättet werden. Um zu sehen, wie wir die Energie abändern müssen, um dies zu verhindern, schauen wir uns die Energie eines Übergangs von 0 nach 1 mit der Breite d (in 1D) an:



Die Funktion u ist also gegeben durch

$$u(x) = \begin{cases} 0 & x \leq 0 \\ \frac{1}{d}x & 0 \leq x \leq d \\ 1 & x \geq d \end{cases}.$$

Wir interessieren uns zunächst für den ersten Beitrag der Energie, also für $\int u'(x)^2 dx$. Da die Ableitung außerhalb der Intervalls $(0, d)$ verschwindet, betrachten wir nur dieses Intervall:

$$\int_0^d u'(x)^2 dx = \int_0^d \left(\frac{1}{d}\right)^2 dx = \frac{1}{(d)^2} \int_0^d 1 = \frac{1}{(d)^2} (d) = \frac{1}{d}.$$

Dieser Wert wird kleiner, wenn d groß ist. Wir sehen also, dass dieser Energiebeitrag, den wir ja zur Verringerung des Rauschens angesetzt hatten, lange Übergänge bevorzugt und damit die Unschärfe erzeugt.

Würden wir statt dessen die Energie $\int u'(x) dx$ (also ohne das Quadrat) betrachten, so erhielten wir

$$\int_0^d u'(x) dx = \int_0^d \frac{1}{d} dx = \frac{1}{d} \int_0^d 1 = \frac{1}{d}(d) = 1.$$

Die Energie des Übergangs ist also unabhängig von der Länge des Übergangs! Es sollten also auch steile Übergänge und somit scharfe Kanten möglich sein. Da die Energie nicht von der Richtung des Übergangs abhängen darf und nicht negativ werden soll, müssen wir noch den Betrag von u' nehmen. Täten wir das nicht, könnte man durch Einfügen von zusätzlichen Übergängen in Gegenrichtung die Energie verringern. Wir benutzen also die neue Energie

$$E[u] = \int_{\Omega} \beta |\nabla u| + (u - f)^2. \quad (1.1)$$

Den Term $\int_{\Omega} |\nabla u|$ bezeichnet man mit *Totale Variation* von u , kurz TV. Der zweite Term, $\int_{\Omega} (u - f)^2$, ist die L^2 -Norm von $u - f$. Daher bezeichnet man die Energie mit $TV\text{-}L^2$. Wir sind nun an einer minimierenden Funktion in der Klasse der Funktionen mit $\nabla u \cdot n = 0$ (n ist hier die äußere Normale an $\partial\Omega$) interessiert.

Bei der Minimierung solcher Energien können Funktionen als Minimierer auftreten mit unendlich steilen Flanken, d.h. mit Sprüngen. Wir sind später an einer Diskretisierung mittels finiter Differenzen interessiert, bei der dieses Problem nicht mehr auftritt. Im Folgenden argumentieren wir deshalb im kontinuierlichen Fall nur formal und nicht mathematisch rigoros. Entscheidend ist nun die folgende alternative (man spricht auch von dualer) Formulierung von $\int_{\Omega} |\nabla u(x)| dx$:

Satz 1.16. *Es gilt*

$$\int_{\Omega} |\nabla u(x)| dx = \max_p \left\{ \int_{\Omega} u(x) (\operatorname{div} p(x)) dx \right\},$$

wobei über alle stetig differenzierbaren Vektorfelder $p : \Omega \rightarrow \mathbb{R}^2$ maximiert wird, die punktweise maximal Länge eins haben, $|p(x)| \leq 1 \forall x \in \Omega$, und die in einer Umgebung des Randes $\partial\Omega$ verschwinden.

Beweis. Zunächst sehen wir mittels partieller Integration, dass

$$\int_{\Omega} u(\operatorname{div} p) = - \int_{\Omega} \nabla u \cdot p,$$

da p auf dem Rand verschwindet. Wir betrachten die Gleichung nun punktweise. Falls $\nabla u(x) = 0$, dann ist offenbar für alle p die Gleichung $|\nabla u| = \nabla u \cdot p$ erfüllt. Nehmen wir nun also an, dass $\nabla u(x) \neq 0$. Es ist $|\nabla u \cdot p| \leq |\nabla u| |p|$ und es gilt Gleichheit, falls ∇u und p parallel sind, d.h. $p = \lambda \nabla u$ mit $\lambda \in \mathbb{R}$. Da wir das Maximum von $-\nabla u \cdot p$ suchen, muss p in die ∇u entgegengesetzte Richtung zeigen, d.h. $\lambda < 0$ sein. Um die Beschränkung $|p| \leq 1$ einzuhalten, wählen wir

$$p = - \frac{\nabla u}{|\nabla u|}.$$

Damit erhalten wir

$$\max_p \int_{\Omega} u(\operatorname{div} p) = \max_p - \int_{\Omega} \nabla u \cdot p = \int_{\Omega} \nabla u \cdot \frac{\nabla u}{|\nabla u|} = \int_{\Omega} |\nabla u|.$$

□

Durch diesen Trick können wir die Energie (1.1) umschreiben in

$$E[u] = \max_{|p| \leq 1} \beta \int_{\Omega} u(\operatorname{div} p) + \int_{\Omega} (u - f)^2,$$

wobei die Bedingungen an p aus dem Satz mit $|p| \leq 1$ abgekürzt wurden. Die Minimierungsaufgabe ist nun also

$$\begin{aligned} \min_u E[u] &= \min_u \max_{|p| \leq 1} \beta \int_{\Omega} u(\operatorname{div} p) + \int_{\Omega} (u - f)^2 \, dx \\ &= \max_{|p| \leq 1} \min_u \beta \int_{\Omega} u(\operatorname{div} p) + \int_{\Omega} (u - f)^2 \, dx. \end{aligned}$$

Wir können nun die Minimierungsaufgabe (für p fest) wie üblich durch berechnen der ersten Variation lösen:

$$\begin{aligned} 0 &= \frac{d}{dt} \int_{\Omega} \beta(u + tv)(\operatorname{div} p) + (u + tv - f)^2 \Big|_{t=0} \\ &= \int_{\Omega} \beta v(\operatorname{div} p) + 2(u + tv - f)v \Big|_{t=0} \\ &= \int_{\Omega} (\beta(\operatorname{div} p) + 2(u - f)) v. \end{aligned}$$

Als Bedingung erhalten wir also

$$\beta(\operatorname{div} p) + 2(u - f) = 0 \iff u = -\frac{\beta}{2}(\operatorname{div} p) + f. \quad (1.2)$$

Einsetzen liefert

$$\begin{aligned} \min_u E[u] &= \max_{|p| \leq 1} \int_{\Omega} \beta \left(-\frac{\beta}{2}(\operatorname{div} p) + f \right) (\operatorname{div} p) + \int_{\Omega} \frac{\beta^2}{4} (\operatorname{div} p)^2, \\ &= \max_{|p| \leq 1} - \int_{\Omega} \frac{\beta^2}{2} (\operatorname{div} p)^2 + \int_{\Omega} \beta f(\operatorname{div} p) + \int_{\Omega} \frac{\beta^2}{4} (\operatorname{div} p)^2 \\ &= \max_{|p| \leq 1} \beta \int_{\Omega} f(\operatorname{div} p) - \int_{\Omega} \frac{\beta^2}{4} (\operatorname{div} p)^2. \end{aligned}$$

Statt der Minimierung eines nicht differenzierbaren Funktional haben wir also nun die Maximierung unter Nebenbedingungen eines differenzierbaren Funktional erhalten. Hat man das Maximum p^* gefunden, so erhält man die Funktion u^* durch (1.2).

Minimierung unter Ungleichungs-Nebenbedingungen

In diesem Abschnitt leiten wir Optimalitätsbedingungen für Minimierungen unter Ungleichungs-Nebenbedingungen (restringierte Optimierung) her. Wir besprechen die Theorie für Funktionen im $\mathbb{R}^n$ und übertragen die Ergebnisse dann auf den Fall unseres Energiefunktional.

Erinnern wir uns zunächst an Gleichungs-Nebenbedingungen: Zu dem Minimierungsproblem

$$\min_{p \in \mathbb{R}^n} f(p) \quad \text{unter der Nebenbedingung } g(p) = 0$$

heißt

$$L(p, \lambda) := f(p) + \lambda g(p)$$

Lagrange-Funktion und $\lambda \in \mathbb{R}$ heißt *Lagrange-Multiplikator*.

Zu einer Minimalstelle p^* gibt es dann stets ein λ^* , so dass

$$\begin{aligned} 0 &= \nabla_p L(p^*, \lambda^*) = \nabla f(p^*) + \lambda^* \nabla g(p^*), \\ 0 &= \nabla_\lambda L(p^*, \lambda^*) = g(p^*). \end{aligned}$$

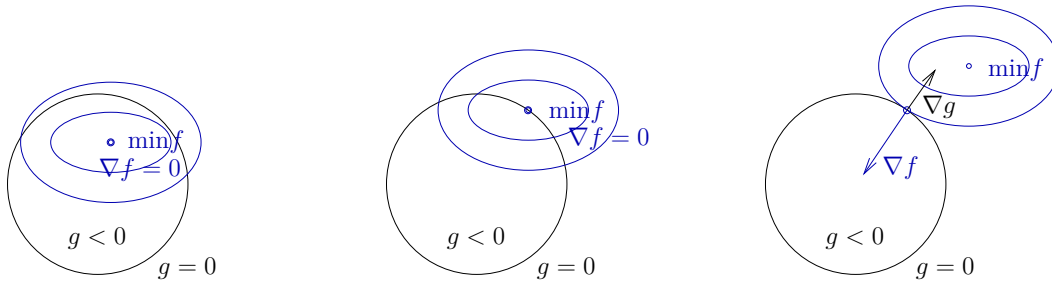
Die zweite Gleichung gilt, da die Nebenbedingung erfüllt sein muss. Die erste Gleichung besagt, dass $\nabla f(p^*)$ in die selbe Richtung wie $\nabla g(p^*)$ zeigt. Der Anteil von $\nabla f(p^*)$, der senkrecht zu $\nabla g(p^*)$ ist – also tangential zur durch $g(p) = 0$ beschriebenen Niveaumenge – muss nämlich Null sein.

Die selbe Funktion $L(p, \lambda)$ kann man auch im Fall von Ungleichungs-Nebenbedingungen definieren. Wir betrachten nun zur Veranschaulichung die verschiedenen Fälle, wo das unrestringierte Minimum von f in Bezug auf die zulässige Menge liegt:

Im Inneren

Auf dem Rand

Außen



$$\begin{array}{lll} \nabla f(p^*) = 0 = -\lambda^* \nabla g(p^*), & \nabla f(p^*) = 0 = -\lambda^* \nabla g(p^*), & \nabla f(p^*) = -\lambda^* \nabla g(p^*), \\ \lambda^* = 0, & \lambda^* = 0, & \lambda^* > 0, \\ g(p^*) < 0, & g(p^*) = 0, & g(p^*) = 0. \end{array}$$

Die Situation im (zweiten und) dritten Fall liegt analog zur Gleichungs-Nebenbedingung, da $g(p^*) = 0$. Zusammen erhalten wir also

$$\begin{aligned} \nabla f(p^*) &= -\lambda^* \nabla g(p^*), \\ \lambda^* &\geq 0, \\ g(p^*) &\leq 0, \\ \lambda^* g(p^*) &= 0. \end{aligned}$$

Für die Untersuchung solcher Minimierungsaufgaben unter Nebenbedingungen beschränken wir uns auf konvexe Funktionen und konvexe Nebenbedingungen zurück.

Definition 1.17 (Konvexe Funktion). *Eine Funktion $f : \mathbb{R}^n \rightarrow \mathbb{R}$ heißt konvex, wenn*

$$f(\lambda x + (1 - \lambda)y) \leq \lambda f(x) + (1 - \lambda)f(y)$$

für alle $x, y \in \mathbb{R}^n$ und alle $\lambda \in [0, 1]$ gilt.

Für eine differenzierbare und konvexe Funktion $f : \mathbb{R}^n \rightarrow \mathbb{R}$ gilt, dass

$$f(y) \geq f(x) + \nabla f(x) \cdot (y - x)$$

für alle $x, y \in \mathbb{R}^n$.

Das heisst, Tangenten ($t \rightarrow f(y) + t \nabla f(y) \cdot (x - y)$) liegen immer unterhalb des Graphen.

Dies sieht man wie folgt ein: Aus der Konvexität folgt

$$f(x + t(y - x)) \leq f(x)(1 - t) + f(y)t$$

für jedes $t \in [0, 1]$ und festes $x, y \in \mathbb{R}^n$. Daraus folgt aber direkt

$$\frac{f(x + t(y - x)) - f(x)}{t} \leq f(y) - f(x)$$

Für $t \rightarrow 0$ erhalten wir $\lim_{t \rightarrow 0} \frac{f(x + t(y - x)) - f(x)}{t} = \nabla f(x) \cdot (y - x)$ und somit

$$\nabla f(x) \cdot (y - x) \leq f(y) - f(x)$$

In 1d mit $h := x - y$ heißt das $f(y + h) \geq f(y) + f'(y)h$.

Nun kommen wir zu den Optimalitätsbedingungen.

Definition 1.18 (Optimalitätsbedingungen). *Seien $f : \mathbb{R}^n \rightarrow \mathbb{R}$ und $g : \mathbb{R}^n \rightarrow \mathbb{R}$ differenzierbar und konvex. Zu dem Minimierungsproblem*

$$\min_{p \in \mathbb{R}^n} f(p) \quad \text{unter der Nebenbedingung } g(p) \leq 0 \quad (1.3)$$

heißt

$$L(p, \lambda) := f(p) + \lambda g(p)$$

Lagrange-Funktion und $\lambda \in \mathbb{R}$ heißt Lagrange-Multiplikator. Ein Punkt $(p^, \lambda^*) \in \mathbb{R}^n \times \mathbb{R}$ heißt KKT-Punkt (nach Karush, Kuhn und Tucker) von (1.3), wenn er die Bedingungen*

$$\nabla_p L(p^*, \lambda^*) = 0$$

$$\lambda^* \geq 0$$

$$g(p^*) \leq 0$$

$$\lambda^* g(p^*) = 0$$

erfüllt.

Bemerkung 1.19. Die letzte Bedingung heißt, dass in einem KKT-Punkt immer $\lambda^* = 0$ oder $g(p^*) = 0$ gilt. Im ersten Fall ist p^* auch ein Minimum von f ohne Nebenbedingung, denn $L(p, 0) = f(p)$ und daher $0 = \nabla_p L(p^*, 0) = \nabla f$. Wir sagen in diesem Fall, dass die Nebenbedingung nicht aktiv ist. Im zweiten Fall, wenn $g(p^*) = 0$, liegt p^* also auf dem Rand des zulässigen Gebietes. Wir sagen, die Nebenbedingung ist aktiv.

Satz 1.20. Seien $f : \mathbb{R}^n \rightarrow \mathbb{R}$ und $g : \mathbb{R}^n \rightarrow \mathbb{R}$ konvex und (p^*, λ^*) sei KKT-Punkt von (1.3). Dann ist p^* Minimalpunkt von f unter der Nebenbedingung $g(p) \leq 0$.

Beweis. Nach Voraussetzung ist

$$\nabla f(p^*) + \lambda^* \nabla g(p^*) = 0 \iff \nabla f(p^*) = -\lambda^* \nabla g(p^*).$$

Sei nun p ein beliebiger zulässiger Vektor, d.h. $g(p) \leq 0$. Da f konvex ist, gilt

$$\begin{aligned} f(p) &\geq f(p^*) + \nabla f(p^*) \cdot (p - p^*) \\ &= f(p^*) - \lambda^* \nabla g(p^*) \cdot (p - p^*). \end{aligned}$$

Falls $\lambda^* = 0$ ist, dann gilt also $f(p) \geq f(p^*)$. Ist andererseits $\lambda^* > 0$, dann muss $g(p^*) = 0$ sein und da g konvex

$$\begin{aligned} g(p) &\geq \underbrace{g(p^*)}_{=0} + \nabla g(p^*) \cdot (p - p^*) \\ &\iff \nabla g(p^*) \cdot (p - p^*) \leq g(p) \leq 0. \end{aligned}$$

Damit erhalten wir auch im zweiten Fall $f(p) \geq f(p^*)$. Da dies für alle zulässigen p gilt, ist p^* ein Minimierer von f . $\square$

Beispiel (in 2D):

$$f(x, y) = \left| \begin{pmatrix} x \\ y \end{pmatrix} - \begin{pmatrix} x_0 \\ 0 \end{pmatrix} \right|^2 \quad \text{mit } x_0 \in \mathbb{R}.$$

Wir berechnen den KKT-Punkt zum Minimierungsproblem

$$\min_{x^2+y^2 \leq 1} f(x, y).$$

Dazu setzen wir $g(x, y) := x^2 + y^2 - 1$. Dann ist g konvex und

$$g(x, y) \leq 0 \iff x^2 + y^2 \leq 1.$$

Die Bedingungen für einen KKT-Punkt sind

$$\nabla f(x^*, y^*) + \lambda^* \nabla g(x^*, y^*) = 0 \iff \begin{aligned} 2(x^* - x_0) + 2\lambda^* x^* &= 0 \\ 2y^* + 2\lambda^* y^* &= 0 \end{aligned}$$

sowie

$$\lambda^* \geq 0, \quad g(x^*, y^*) \leq 0, \quad \text{und} \quad \lambda^* g(x^*, y^*) = 0.$$

Aus der zweiten Gleichung ergibt sich

$$(1 + \lambda^*)y^* = 0,$$

und da $\lambda^* \geq 0$ muss $y^* = 0$ sein. Die Nebenbedingung $g(x^*, y^*) \leq 0$ ist also äquivalent zu $|x^*| \leq 1$.

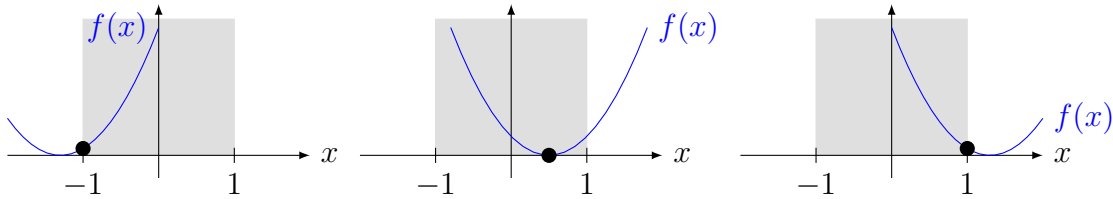
Falls die Nebenbedingung nicht aktiv ist, also $\lambda^* = 0$, dann folgt aus der ersten Gleichung $x^* = x_0$ und $((x_0, 0), 0)$ ist KKT-Punkt. Hierzu muss $|x_0| \leq 1$ sein. Wie erwartet ist dies auch das Minimum des unrestringierten Minimierungsproblems.

Falls die Nebenbedingung aktiv ist, also $\lambda^* > 0$, dann muss $g(x^*, y^*) = 0$ sein, also $x^* = \pm 1$. Wir unterscheiden

- $x^* = +1$. Dann $(1 + \lambda^*) = x_0 \iff \lambda^* = x_0 - 1$ und $((1, 0), x_0 - 1)$ ist KKT-Punkt falls $x_0 > 1$ (wegen $\lambda^* > 0$).
- $x^* = -1$. Dann $-(1 + \lambda^*) = x_0 \iff \lambda^* = -x_0 - 1$ und $((-1, 0), -x_0 - 1)$ ist KKT-Punkt falls $x_0 < -1$.

Zusammen erhalten wir also für den KKT-Punkt

$$((x^*, y^*), \lambda^*) = \begin{cases} ((-1, 0), -x_0 - 1) & x_0 < -1, \\ ((x_0, 0), 0) & -1 \leq x_0 \leq 1, \\ ((1, 0), x_0 - 1) & x_0 > 1. \end{cases}$$



Kommen wir nun zurück zu unserem Minimierungsproblem

$$\min_u E[u] = \max_{|p| \leq 1} F[p]$$

für $F[p] = \beta \int_{\Omega} f(\operatorname{div} p) - \int_{\Omega} \frac{\beta^2}{4} (\operatorname{div} p)^2$.

Wir können dies leicht in ein Minimierungsproblem umschreiben:

$$\min_u E[u] = - \min_{|p| \leq 1} (-F[p]).$$

Die Nebenbedingung $|p(x)| \leq 1$ beschreiben wir mittels der Funktion $g(p)$ die jedem x den Wert

$$g(p)(x) = |p(x)|^2 - 1$$

zuordnet. Die erste Variation von g ist dann gegeben als die Funktion $x \mapsto 2p(x)$. Um die KKT-Bedingungen herzuleiten, berechnen wir wieder die erste Variation des

Terms auf der rechten Seite.

$$\begin{aligned}
 \frac{d}{dt}(-F[p + tq])\Big|_{t=0} &= \frac{d}{dt} \int_{\Omega} -\beta f(\operatorname{div} p + t \operatorname{div} q) + \frac{\beta^2}{4} (\operatorname{div} p + t \operatorname{div} q)^2 \Big|_{t=0} \\
 &= -\beta \int_{\Omega} f(\operatorname{div} q) + \frac{\beta^2}{4} \int_{\Omega} 2(\operatorname{div} p + t \operatorname{div} q)(\operatorname{div} q) \Big|_{t=0} \\
 &= -\beta \int_{\Omega} f(\operatorname{div} q) + \frac{\beta^2}{2} \int_{\Omega} (\operatorname{div} p)(\operatorname{div} q) \\
 &= \beta \int_{\Omega} \nabla f \cdot q - \frac{\beta^2}{2} \int_{\Omega} \nabla \operatorname{div} p \cdot q \\
 &= \int_{\Omega} \left(-\frac{\beta^2}{2} \nabla \operatorname{div} p + \beta \nabla f \right) \cdot q.
 \end{aligned}$$

Dabei verwenden wir wieder Testfunktionen q , die auf dem Rand Null sind, so dass das Randintegral bei der partiellen Integration wegfällt. Die erste Variation von $-F$ ist also gegeben als die Funktion

$$x \mapsto -\beta \nabla \left(\frac{\beta}{2} \operatorname{div} p(x) - f(x) \right).$$

die jedem x einen lokalen Variationswert zuordnet. Wenn wir die erste Variation statt des Gradienten verwenden, erhalten wir als Bedingung für einen KKT-Punkt

$$-\beta \nabla \left(\frac{\beta}{2} \operatorname{div} p(x) - f(x) \right) + 2\lambda(x)p(x) = 0. \quad (1.4)$$

Anders als im Fall einer Minimierung in $\mathbb{R}^n$ ist der Lagrange-Multiplikator λ hier keine Zahl, sondern eine Funktion, die wiederum jedem x einen Lagrange-Multiplikator-Wert zuordnet.

Lemma 1.21. *Es ist $\lambda(x) = \frac{\beta}{2} \left| \nabla \left(\frac{\beta}{2} (\operatorname{div} p(x)) - f(x) \right) \right|$.*

Beweis. Wir zeigen die Aussage wieder punktweise. Falls $\lambda(x) > 0$, dann muss $|p(x)| = 1$ sein. Nehmen wir den Betrag von (1.4), so erhalten wir in x

$$\lambda(x) = |\lambda(x)| = |\lambda(x)p(x)| = \frac{\beta}{2} \left| \nabla \left(\frac{\beta}{2} (\operatorname{div} p(x)) - f(x) \right) \right|.$$

Falls $\lambda(x) = 0$, so ist nach (1.4) auch die rechte Seite Null. □

Einsetzen liefert

$$-\nabla \left(\frac{\beta}{2} (\operatorname{div} p(x)) - f(x) \right) + \left| \nabla \left(\frac{\beta}{2} (\operatorname{div} p(x)) - f(x) \right) \right| p(x) = 0. \quad (1.5)$$

Diese Gleichung wollen wir nun mit Hilfe einer Fixpunktiteration numerisch lösen.

Fixpunktiterationen

Allgemein ist eine Fixpunktgleichung von der Form $f(x^*) = x^*$. Gesucht ist also ein Fixpunkt x^* von f . Man kann nun eine Gleichung in eine Fixpunktgleichung umformulieren und dann die Fixpunktgleichung mittels einer Fixpunktiteration $x^{n+1} = f(x^n)$ zu gegebenen x^0 lösen.

Bevor wir die Fixpunktiteration auf die Gleichung (1.5) anwenden, teilen wir die Gleichung durch $\beta/2$ und erhalten nach weiterer Umformung

$$\begin{aligned} p(x) &= p(x) + \tau \nabla \left((\operatorname{div} p(x)) - \frac{2}{\beta} f(x) \right) - \tau \left| \nabla \left((\operatorname{div} p(x)) - \frac{2}{\beta} f(x) \right) \right| p(x) \Rightarrow \\ p(x) &= F(p)(x) := \frac{p(x) + \tau \nabla (\operatorname{div} p(x) - \frac{2}{\beta} f(x))}{1 + \tau \left| \nabla (\operatorname{div} p(x) - \frac{2}{\beta} f(x)) \right|}. \end{aligned}$$

Dann verwenden wir punktweise die daraus abgeleitete folgende Fixpunktiteration:

$$p^{n+1}(x) = \frac{p^n(x) + \tau \nabla (\operatorname{div} p^n(x) - \frac{2}{\beta} f(x))}{1 + \tau \left| \nabla (\operatorname{div} p^n(x) - \frac{2}{\beta} f(x)) \right|}.$$

Der Vorteil dieser Formulierung besteht darin, dass

- der neue Wert p^{n+1} einfach und ohne Lösen eines Gleichungssystems berechnet werden kann,
- die Konvergenz des Verfahrens sich durch die Wahl von τ steuern lässt und
- die Iteration für τ klein genug konvergiert.

Wir stoppen die Fixpunktiteration, wenn sich p nicht mehr oder kaum noch ändert, denn wenn $p^{n+1} \approx p^n$, dann ist p^n (fast) ein Fixpunkt.

Ortsdiskretisierung

Wir benutzen wie in Abschnitt 1.3 ein äquidistantes zweidimensionales Gitter mit Gitterweite h . Wie vorher diskretisieren wir die eindimensionalen Richtungsableitungen ∂_x und ∂_y durch Vorwärtsdifferenzen. Für das Enttauschen eines Bildes schien unsere Randbedingung $u = f$ auf $\partial\Omega$ nicht ideal zu sein, da sich so ein unschöner Rand bildete. Wir verwenden diesmal die Randbedingung

$$\nabla u \cdot \nu = 0,$$

wobei ν die äußere Normale ist. Bei einem rechteckigen Ω heißt dies, dass $\partial_x u = 0$ am linken und rechten Rand und $\partial_y u = 0$ am oberen und unteren Rand. Wir erhalten also für die diskreten, approximierten Gradienten ∂_x^h und ∂_y^h

$$\partial_x^h u(x_i, y_j) = \begin{cases} \frac{1}{h} (u(x_{i+1}, y_j) - u(x_i, y_j)) & \text{falls } i < N \\ 0 & \text{falls } i = N \end{cases}$$

und

$$\partial_y^h u(x_i, y_j) = \begin{cases} \frac{1}{h}(u(x_i, y_{j+1}) - u(x_i, y_j)) & \text{falls } j < N \\ 0 & \text{falls } j = N \end{cases}$$

Das Vektorfeld p speichern wir in zwei Komponenten $p^{(1)}$ und $p^{(2)}$. Die Divergenz von p erhält man dann durch Negieren und Transponieren der Matrixdarstellung des Gradienten:

$$\begin{aligned} (\operatorname{div}^h p)(x_i, y_j) = & \frac{1}{h} \begin{cases} p^{(1)}(x_i, y_j) & \text{falls } i = 1 \\ p^{(1)}(x_i, y_j) - p^{(1)}(x_{i-1}, y_j) & \text{falls } 1 < i < N \\ -p^{(1)}(x_{i-1}, y_j) & \text{falls } i = N \end{cases} \\ & + \frac{1}{h} \begin{cases} p^{(2)}(x_i, y_j) & \text{falls } j = 1 \\ p^{(2)}(x_i, y_j) - p^{(2)}(x_i, y_{j-1}) & \text{falls } 1 < j < N \\ -p^{(2)}(x_i, y_{j-1}) & \text{falls } j = N \end{cases} \end{aligned}$$

Das Aufstellen der Matrizen verläuft analog zu Abschnitt 1.3. Statt `spdiags` zu verwenden und die entsprechenden Zeilen und Spalten Null zu setzen, kann man auch die Funktion `sparse` benutzen. Der Aufruf ist

`S=sparse(i,j,s,n,n);`

mit Vektoren i, j, s , die für jeden nicht-Null-Eintrag der Matrix einen Eintrag haben. Die Matrix S ist dann eine $n \times n$ Matrix mit $S_{i(k),j(k)} = s(k)$ für alle k . Für den Gradienten in x-Richtung also z.B.

Schema 1.22 (MATLAB-Code für x-Ableitung).

```
rows=N*N;
nzs=2*N*(N-1);
row=zeros(nzs,1);
col=zeros(nzs,1);
val=zeros(nzs,1);
count=0;
for i=1:(N-1)
    for j=1:N
        count=count+1;
        row(count)=(i-1)*N+j;
        col(count)=(i-1)*N+j;
        val(count)=-1/h;
        count=count+1;
        row(count)=(i-1)*N+j;
        col(count)=i*N+j;
        val(count)=1/h;
    end
end
dx=sparse(row,col,val,rows,rows);
```

Achtung: Die Vektoren i, j und s in jedem Schritt um 1 zu verlängern wäre sehr langsam. Daher sollte unbedingt vorher mittels `zeros` die richtige Menge Speicher reserviert werden.

Als Abbruchkriterium wählen wir

$$\max_{ij} |p^{n+1}(x_i, y_j) - p^n(x_i, y_j)| < \theta$$

mit θ klein.

Bemerkung 1.23. *Mit der Energie aus Abschnitt 1.3 war es egal, ob die Helligkeit eines Bildpunktes durch Werte in $[0, 1]$ oder $[0, 255]$ (wie `imread` es liefert) gerechnet wurde, denn*

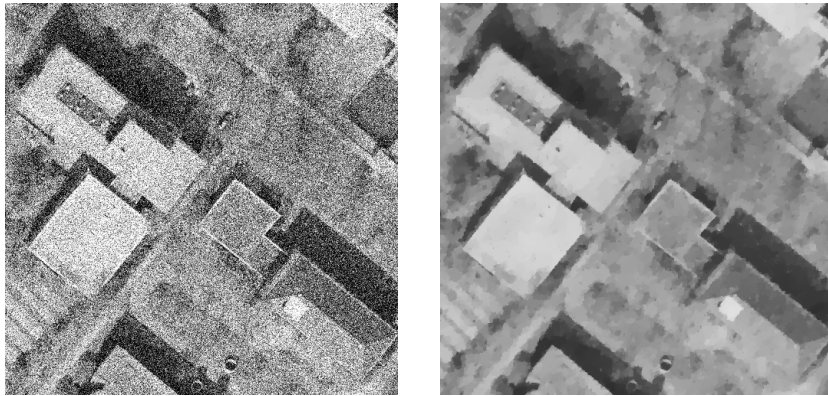
$$E[\lambda u] = \int_{\Omega} \beta |\lambda \nabla u|^2 + (\lambda u - \lambda f)^2 = \lambda^2 \int_{\Omega} \beta |\nabla u|^2 + (u - f)^2.$$

Dies ist nun anders, denn

$$E[\lambda u] = \int_{\Omega} \beta |\lambda \nabla u| + (\lambda u - \lambda f)^2 = \lambda^2 \int_{\Omega} \hat{\beta} |\nabla u| + (u - f)^2$$

mit $\hat{\beta} = \beta/\lambda$. Man muss $\hat{\beta}$ also je nach verwendeter Skalierung anpassen.

Als Ergebnis für $\tau = 0.2h^2$, $\beta = 0.4h$, $\theta = 10^{-3}$, $\Omega = [0, 1]^2$, $h = 1/511$ erhalten wir nach ca. 60 Sekunden



Wie wir sehen, ist das Ergebnis weit weniger verschwommen als zuvor. Ein Nachteil dieses Verfahrens ist allerdings, dass der Kontrast des Bildes verringert wird.

Bemerkung 1.24. *In der Rechnung wurde ein Bild mit den Dimensionen $\Omega = [0, 1]^2$ angenommen. Da das Bild 512×512 Pixel hat, ergibt sich ein Wert von $h = 1/511$. Hierbei ist zu beachten, dass die Schrittweite τ sich mit h^2 verändert. Im Programm wurde $\tau = 0.2h^2$ gesetzt.*

Als andere Möglichkeit würde man $h = 1$ setzen, also $\Omega = [1, 512]^2$. Dann ist τ zwar unabhängig von der Auflösung des Bildes, aber dafür muss man β anpassen, wenn sich die Auflösung ändert.

Der Vorteil der ersten Variante ist also, dass man das passende β für ein niedrig aufgelöstes Bild suchen kann und dann die höher aufgelöste Variante mit demselben β rechnen kann.

Programmieraufgabe 3.

1. Implementieren Sie das $TV-L^2$ -Verfahren zur Bildglättung in *MATLAB* als Funktion.
2. Für welche Werte von τ konvergiert die Iteration und wie schnell?

2 Finite Elemente

Finite Differenzen haben den Vorteil, einfach implementierbar und effizient zu sein. Allerdings

- sind sie unflexibel bei der Approximation des Gebiets Ω , falls dieses kein Rechteck bildet (z.B. bei Messpunkten auf einem Dreiecksgitter).
- besitzen sie nur dann eine gute Konvergenzordnung, wenn die kontinuierliche Lösung sehr glatt ist (viermal stetig differenzierbar!)

Abhilfe schafft die sogenannte Finite-Elemente-Methode.

2.1 Schwache Lösungen

Wir wollen nun einen Lösungsbegriff definieren, der auch dann gültig ist, wenn u weniger glatt ist.

Erinnern wir uns nochmal an die Herleitung der notwendigen Bedingung

$$-\beta \Delta u + u = f$$

in Abschnitt 1.3. Es war

$$E[u] = \int \beta |\nabla u|^2 + (u - f)^2$$

und es sollte gelten

$$0 = \frac{d}{dt} E[u + tv] \Big|_{t=0}$$

für alle stetig differenzierbaren $v : \bar{\Omega} \rightarrow \mathbb{R}$ mit $v = 0$ auf $\partial\Omega$, also

$$\begin{aligned} 0 &= \frac{d}{dt} \int_{\Omega} \beta |\nabla u + t \nabla v|^2 + (u + tv - f)^2 \Big|_{t=0} \\ &= \int_{\Omega} 2\beta (\nabla u + t \nabla v) \cdot \nabla v + (u + tv - f)^2 v \Big|_{t=0} \\ &= 2 \int_{\Omega} \beta \nabla u \cdot \nabla v + (u - f)v. \end{aligned}$$

Bisher hatten wir weiter gerechnet

$$\begin{aligned} &= 2 \int_{\Omega} -\beta \Delta u v + (u - f)v + 2 \int_{\partial\Omega} \nabla u \cdot \nu v \\ &= 2 \int_{\Omega} (-\beta \Delta u + (u - f))v. \end{aligned}$$

Stattdessen sagen wir, dass u eine *schwache Lösung* von $-\beta\Delta u + u = f$ ist, wenn für alle „Testfunktionen“ v gilt

$$\int_{\Omega} \beta \nabla u \cdot \nabla v + (u - f)v = 0.$$

Damit reicht es, wenn u einmal differenzierbar ist. Allerdings brauchen wir gar nicht, dass u stetig differenzierbar ist, sondern nur dass

$$\int_{\Omega} \nabla u \cdot \nabla v$$

existiert. Daher definieren wir

Definition 2.1. Sei $\Omega \subset \mathbb{R}^n$ beschränkt, der Rand von Ω sei stückweise glatt. Dann ist $C_0^\infty(\Omega)$ die Menge aller beliebig oft differenzierbarer Funktionen, die in einer Umgebung des Randes von Ω Null sind.

Wir nennen $w_i = \partial_i u$ schwache Ableitung bezüglich x_i von u auf Ω , falls für alle $\varphi \in C_0^\infty$ gilt, dass

$$\int_{\Omega} w_i \varphi = - \int_{\Omega} u \partial_i \varphi.$$

Bemerkung 2.2.

1. Schreiben wir $w := (w_1, \dots, w_n)$, so erhalten wir

$$\int_{\Omega} w \varphi = - \int_{\Omega} u \nabla \varphi \quad \forall \varphi \in C_0^\infty(\Omega)$$

oder

$$\int_{\Omega} w \cdot \Phi = - \int_{\Omega} u(\operatorname{div} \Phi) \quad \forall \Phi : \Omega \rightarrow \mathbb{R}^n \text{ mit } \Phi_i \in C_0^\infty(\Omega).$$

2. Falls u differenzierbar ist, so stimmen die übliche („starke“) Ableitung und die schwache Ableitung überein: mit $w := \nabla u$ erhalten wir mittels partieller Integration

$$\int_{\Omega} w \cdot \Phi = \int_{\Omega} \nabla u \cdot \Phi = - \int_{\Omega} u(\operatorname{div} \Phi) + \int_{\partial\Omega} u \underbrace{\Phi \cdot \nu}_{=0} = - \int_{\Omega} u(\operatorname{div} \Phi).$$

Allgemein ist für eine stetige Funktion $u : [a, c] \rightarrow \mathbb{R}$, die auf den Intervallen (a, b) und (b, c) differenzierbar ist (im Punkt b jedoch nicht), die schwache Ableitung gegeben durch

$$w(x) = \begin{cases} u'(x) & x \in (a, b) \\ 0 & x = b \\ u'(x) & x \in (b, c) \end{cases}$$

In der Tat ist

$$\begin{aligned}
\int_a^c w \varphi &= \int_a^b w \varphi + \int_b^c w \varphi \\
&= \int_a^b u' \varphi + \int_b^c u' \varphi \\
&= - \int_a^b u \varphi' - \int_b^c w \varphi' + u \varphi|_a^b + u \varphi|_b^c \\
&= - \int_a^c u \varphi' + u(c) \underbrace{\varphi(c)}_{=0} - \underbrace{u(b) \varphi(b) + u(b) \varphi(b)}_{=0} - u(a) \underbrace{\varphi(a)}_{=0} \\
&= - \int_a^c u \varphi'.
\end{aligned}$$

Dabei ist $u(b)$ einmal als linksseitiger Grenzwert, einmal als rechtsseitiger zu verstehen, die beiden Werte sind also nur gleich, wenn u stetig ist.

Also ist w schwache Ableitung von u , obwohl u im Punkt b nicht differenzierbar ist. Wie wir gesehen haben, „sieht“ das Integral einzelne Punkte nicht, denn

$$\int_a^c u = \int_a^b u + \int_b^c u$$

und somit ist die Festlegung $w(b) = 0$ beliebig und eigentlich auch überflüssig. Wie in obigem Beispiel schon geschehen, werden wir die schwache Ableitung in einzelnen Punkten nicht festlegen.

Definition 2.3 ($L^2(\Omega)$). Wir definieren durch

$$\langle u, w \rangle_{L^2(\Omega)} := \int_{\Omega} u \cdot w$$

das L^2 -Skalarprodukt. Man rechnet leicht nach, dass $\langle \cdot, \cdot \rangle$ symmetrisch, bilinear und positiv definit ist und damit ein Skalarprodukt. Mit diesem Skalarprodukt erhält man eine Norm, die L^2 -Norm

$$\|u\|_{L^2(\Omega)}^2 := \langle u, u \rangle_{L^2(\Omega)} = \int_{\Omega} u^2 \, dx.$$

Schließlich bezeichnen wir die Menge aller quadratintegrierbaren Funktionen mit $L^2(\Omega)$,

$$L^2(\Omega) = \left\{ u \mid \int_{\Omega} u^2 \, dx < \infty \right\} = \left\{ u \mid \|u\|_{L^2(\Omega)} < \infty \right\}.$$

$L^2(\Omega)$ ist ein Vektorraum.

Definition 2.4 (Sobolevraum). Sei $\Omega \subset \mathbb{R}^n$ offen mit einem stückweise glatten Rand. Wir bezeichnen mit $H^1(\Omega)$ die Menge aller Funktionen u aus $L^2(\Omega)$, die in alle Richtungen schwach differenzierbar sind deren schwache Ableitung in $L^2(\Omega)$ liegt (also $\int_{\Omega} (\partial_{x_i} u)^2 < \infty$ für alle $i = 1, \dots, N$).

Wir definieren eine Norm auf $H^1(\Omega)$ durch

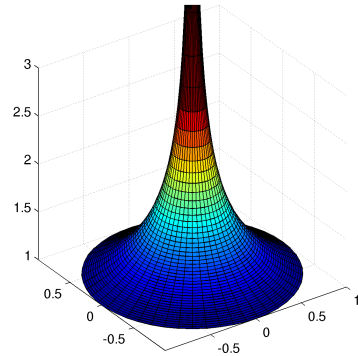
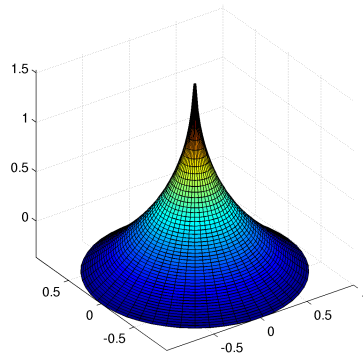
$$\begin{aligned}\|u\|_{H^1(\Omega)}^2 &:= \|u\|_{L^2(\Omega)}^2 + \|\nabla u\|_{L^2(\Omega)}^2 = \|u\|_{L^2(\Omega)}^2 + \sum_{i=1}^n \|\partial_{x_i} u\|_{L^2(\Omega)}^2 \\ &= \int_{\Omega} \left(u^2 + \sum_{i=1}^n (\partial_{x_i} u)^2 \right).\end{aligned}$$

Weiter sei $H_0^1(\Omega)$ die Menge der Funktionen $u \in H^1(\Omega)$ mit $u = 0$ auf $\partial\Omega$.

Notation 2.5. Wir schreiben kurz $\|u\|_{0,\Omega} := \|u\|_{L^2(\Omega)}$ (Null-mal schwach differenzierbar) und $\|u\|_{1,\Omega} := \|u\|_{H^1(\Omega)}$ (Ein-mal schwach differenzierbar).

Wir schauen uns nun einige Beispiele für Funktionen aus $H^1(\Omega)$ an. Das wichtigste Beispiel, nämlich stetige, stückweise differenzierbare Funktionen, haben wir schon gesehen. In 1d lassen alle Funktionen aus $H^1(\Omega)$ stetig ergänzen, indem man sie ggf. auf einer Ausnahmestelle A , die vom Integral „nicht gesehen“ wird ($\int_{\Omega \setminus A} 1 = \int_{\Omega} 1$), undefiniert.

In höheren Dimensionen müssen Funktionen aus $H^1(\Omega)$ nicht stetig ergänzbar sein. Betrachten wir z.B. die beiden Funktionen



$$f(x) = \log \left(\log \left(1 + \frac{1}{|x|} \right) \right),$$

$$g(x) = \frac{1}{\sqrt{|x|}},$$

die im Nullpunkt eine Singularität besitzen, sich also dort nicht stetig ergänzen lassen. Für beide Funktionen kann man zeigen, dass sie schwache Ableitungen besitzen (die außerhalb des Ursprungs ihren klassischen Ableitungen entsprechen), aber nur die schwache Ableitung von f liegt in $L^2(B_1(0))$. Damit ist $f \in H^1(B_1(0))$ und $g \notin H^1(B_1(0))$. Der Unterschied zwischen beiden Funktionen ist, wie schnell sie gegen Unendlich gehen.

Schließlich sei noch angemerkt, dass Funktionen mit einem Sprung nicht in $H^1(\Omega)$ liegen.

Für uns wird im weiteren Verlauf nur das Beispiel der stetigen, stückweise differenzierbaren Funktion relevant sein.

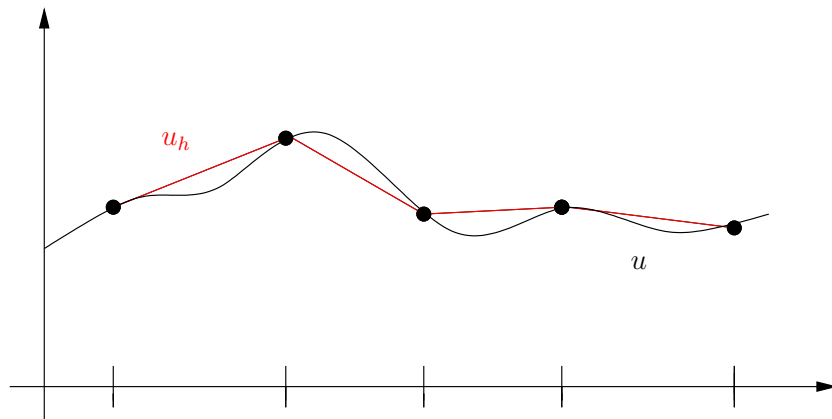
Mit diesen Definitionen können wir nun Problemstellung 1.13 wie folgt verallgemeinern.

Problemstellung 2.6 (Schwache Lösung). Sei $\Omega \subset \mathbb{R}^n$ beschränkt mit einem glatten Rand. Sei $f \in L^2(\Omega)$. Dann heißt $u \in H_0^1(\Omega)$ schwache Lösung von $-\beta \Delta u + u = f$ mit Nullrandwerten, falls für alle Testfunktionen $v \in H_0^1(\Omega)$ gilt, dass

$$\int_{\Omega} \beta \nabla u \cdot \nabla v + (u - f)v = 0. \quad (P^w)$$

2.2 Approximation durch Finite Elemente

Wir kommen nun zur Approximation von schwachen Lösungen. Zur besseren Veranschaulichung betrachten wir zunächst den eindimensionalen Fall ($n = 1$). Der Definitionsbereich einer Funktion u wird in Intervalle unterteilt, auf denen wir u durch affine Funktionen approximieren. Es ergibt sich eine stückweise affine, stetige Funktion u_h :

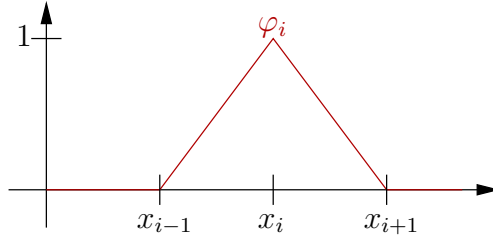


Wie können wir u_h einfach darstellen?

Definition 2.7 (Hütchenbasis in 1D). Sei $0 = x_1 < x_2 < \dots < x_N = L$. Dann bezeichnet $\varphi_i(x)$ diejenige stetige, stückweise affine Funktion, für die

$$\varphi_i(x_j) := \delta_{ij} = \begin{cases} 1 & \text{falls } i = j, \\ 0 & \text{sonst.} \end{cases}$$

gilt.



Satz 2.8. Jede stetige, in Bezug auf die Gitterpunkte $x_1 < \dots < x_N$ stückweise affine Funktion u_h kann als Linearkombination aus Hütchenfunktionen dargestellt werden:

$$u_h(x) = \sum_{i=1}^N U_i \varphi_i(x),$$

wobei $U_i = u_h(x_i)$ zu setzen ist.

Die Funktionen $\{\varphi_i\}$ bilden eine Basis des Vektorraums der stetigen stückweise affinen Funktionen. Diesen Raum bezeichnen wir mit V_h bzw. V_h^0 (mit Nullrandwerten, d.h. ohne φ_1, φ_N). Der Vektorraum V_h bzw. V_h^0 ist ein Untervektorraum von $H^1([0, L])$ bzw. $H_0^1([0, L])$, insbesondere sind Funktionen in V_h und V_h^0 schwach differenzierbar.

Nun lässt sich das Problem (P^w) einfach approximieren: Wir ersetzen lediglich $H_0^1(\Omega)$ durch V_h^0 !

Gesucht ist also $u_h \in V_h^0$, so dass

$$\beta \int_{\Omega} \nabla u_h \cdot \nabla v_h + \int_{\Omega} u_h v_h = \int_{\Omega} f_h v_h \quad \forall v_h \in V_h^0. \quad (2.1)$$

Die Gleichung (2.1) ist linear in Bezug auf v_h . Es genügt also, statt allen v_h nur $v_h = \varphi_i$ für $i = 2, \dots, N-1$ zu betrachten, also

$$\beta \int_{\Omega} \nabla u_h \cdot \nabla \varphi_i + \int_{\Omega} u_h \varphi_i = \int_{\Omega} f_h \varphi_i \quad i = 2, \dots, N-1.$$

Nun können wir noch u_h und f_h in der Basis darstellen

$$u_h(x) = \sum_{j=2}^{N-1} U_j \varphi_j(x) \quad \text{und} \quad f_h(x) = \sum_{j=2}^{N-1} F_j \varphi_j(x).$$

Damit muss für alle Basisfunktionen φ_i , $i = 2, \dots, N-1$ gelten:

$$\begin{aligned} \int_{\Omega} \beta \nabla \left(\sum_{j=2}^{N-1} U_j \varphi_j \right) \cdot \nabla \varphi_i + \int_{\Omega} \left(\sum_{j=2}^{N-1} U_j \varphi_j \right) \varphi_i &= \int_{\Omega} \left(\sum_{j=2}^{N-1} F_j \varphi_j \right) \varphi_i \\ \iff \sum_{j=2}^{N-1} U_j \underbrace{\beta \int_{\Omega} \nabla \varphi_i \cdot \nabla \varphi_j}_{=: L_{ij}} + \sum_{j=2}^{N-1} U_j \underbrace{\int_{\Omega} \varphi_i \varphi_j}_{=: M_{ij}} &= \sum_{j=2}^{N-1} F_j \int_{\Omega} \varphi_j \varphi_i \end{aligned}$$

Wir nennen die Matrix M *Massematrix* und die Matrix L *Steifigkeitsmatrix*. Daraus ergibt sich die Matrix-Vektor-Formulierung des Gleichungssystems:

$$(L + M)U = MF$$

Bemerkung 2.9. Wir mussten uns bei der Finite-Elemente-Diskretisierung keine Gedanken über die Art der Approximation des Gradienten machen. Nach Wahl der Approximationsfunktionen sind die Matrizen durch die Integrale für M_{ij} und L_{ij} gegeben, dabei können die schwachen Ableitungen der Basisfunktionen exakt berechnet werden. Die einzige Approximation findet durch die Wahl des Ansatzraumes V_h bzw. V_h^0 statt.

Berechnung der Matrizen M und L

Wir berechnen die Matrizen in 1d auf einem äquidistanten Gitter, also $x_i = (i-1)h$.

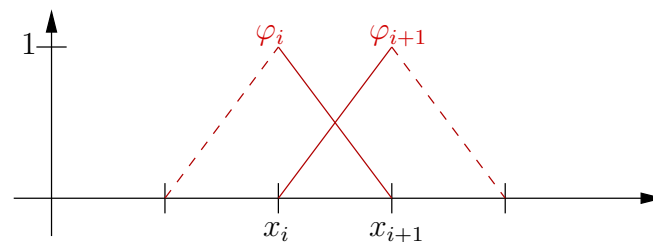
Für die Hütchenbasis gilt $\varphi_i(x_j) = \delta_{ij} = \begin{cases} 1 & \text{falls } i = j, \\ 0 & \text{sonst.} \end{cases}$

Dazwischen wird affin interpoliert.

Es ergeben sich damit folgende Definitionen für die Funktionen und ihre Ableitungen:

$$\varphi_i(x) = \begin{cases} 0 & \text{falls } x \leq x_{i-1} \text{ oder } x \geq x_{i+1}, \\ \frac{x - x_{i-1}}{h} & \text{falls } x \in (x_{i-1}, x_i], \\ \frac{x_{i+1} - x}{h} & \text{falls } x \in (x_i, x_{i+1}), \end{cases}$$

$$\varphi'_i(x) = \begin{cases} 0 & \text{falls } x \leq x_{i-1} \text{ oder } x \geq x_{i+1}, \\ \frac{1}{h} & \text{falls } x \in (x_{i-1}, x_i], \\ -\frac{1}{h} & \text{falls } x \in (x_i, x_{i+1}). \end{cases}$$



Da sich nur benachbarte Hütchenfunktionen überlappen, sind die meisten Einträge

der Matrizen Null. Die Nicht-Null-Einträge für die Massematrix sind

$$\begin{aligned}
 M_{ii} &= \int (\varphi_i)^2 = \frac{1}{h^2} \int_{x_{i-1}}^{x_i} (x - x_{i-1})^2 dx + \frac{1}{h^2} \int_{x_i}^{x_{i+1}} (x_{i+1} - x)^2 dx \\
 &= \frac{1}{h^2} \left[\frac{1}{3} (x - x_{i-1})^3 \right]_{x=x_{i-1}}^{x=x_i} - \frac{1}{h^2} \left[\frac{1}{3} (x_{i+1} - x)^3 \right]_{x=x_i}^{x=x_{i+1}} \\
 &= \frac{1}{h^2} \frac{1}{3} h^3 + \left(-\frac{1}{h^2} \left(-\frac{1}{3} h^3 \right) \right) = \frac{2}{3} h \\
 M_{i,i+1} &= \int \varphi_i \varphi_{i+1} = \frac{1}{h^2} \int_{x_i}^{x_{i+1}} \underbrace{(x_{i+1} - x)}_v \underbrace{(x - x_i)}_{u'} dx \\
 &\stackrel{\text{part. Int.}}{=} -\frac{1}{h^2} \int_{x_i}^{x_{i+1}} (-1) \frac{1}{2} (x - x_i)^2 dx + \underbrace{\frac{1}{h^2} \left[(x_{i+1} - x) \frac{1}{2} (x - x_i)^2 \right]_{x=x_i}^{x=x_{i+1}}}_{=0} \\
 &= \frac{1}{h^2} \int_{x_i}^{x_{i+1}} \frac{1}{2} (x - x_i)^2 dx = \frac{1}{h^2} \left[\frac{1}{6} (x - x_i)^3 \right]_{x=x_i}^{x=x_{i+1}} \\
 &= \frac{1}{h^2} \frac{1}{6} h^3 = \frac{1}{6} h = M_{i,i-1}
 \end{aligned}$$

also hat die $(N-2) \times (N-2)$ -Matrix M die Gestalt

$$M = \frac{h}{6} \begin{pmatrix} 4 & 1 & & & 0 \\ 1 & 4 & 1 & & \\ & \ddots & \ddots & \ddots & \\ & & 1 & 4 & 1 \\ 0 & & & 1 & 4 \end{pmatrix}.$$

Für die Steifigkeitsmatrix erhalten wir

$$\begin{aligned}
 L_{ii} &= \beta \int (\varphi'_i)^2 = \beta \int_{x_{i-1}}^{x_i} \frac{1}{h^2} + \beta \int_{x_i}^{x_{i+1}} \left(-\frac{1}{h^2} \right) \\
 &= \beta h \frac{1}{h^2} + \beta h \frac{1}{h^2} = \frac{2\beta}{h} \\
 L_{i,i+1} &= \beta \int \varphi'_i \varphi'_{i+1} = \beta \int_{x_i}^{x_{i+1}} \left(-\frac{1}{h} \right) \left(\frac{1}{h} \right) \\
 &= \beta h \left(-\frac{1}{h^2} \right) = -\frac{\beta}{h} = L_{i,i-1}
 \end{aligned}$$

somit hat die $(N-2) \times (N-2)$ -Matrix L die Gestalt

$$L = \frac{\beta}{h} \begin{pmatrix} 2 & -1 & & & 0 \\ -1 & 2 & -1 & & \\ & \ddots & \ddots & \ddots & \\ & & -1 & 2 & -1 \\ 0 & & & -1 & 2 \end{pmatrix}.$$

Bemerkung 2.10. Multipliziert man die Steifigkeitsmatrix mit $\frac{1}{h}$, so erhalten wir genau die selbe Matrix, die wir auch bei der Finite-Differenzen-Diskretisierung des

Laplace-Operators erhalten haben. Die Massmatrix unterscheidet sich jedoch (auch nach Multiplikation mit $\frac{1}{h}$) von der Matrix, die an dieser Stelle bei Finiten Differenzen auftrat – der Einheitsmatrix.

2.3 Randwerte

Bisher haben wir implizit Null-Randwerte durch Wahl der Räume $H_0^1(\Omega)$ bzw. V_h^0 vorgeschrieben, also

$$\begin{aligned} -\beta \Delta u + u &= f && \text{in } \Omega \\ u &= 0 && \text{auf } \partial\Omega \end{aligned}$$

Wir wollen aber (zunächst)

$$\begin{aligned} -\beta \Delta u + u &= f && \text{in } \Omega \\ u &= f && \text{auf } \partial\Omega \end{aligned}$$

Dazu führen wir das Problem auf ein Problem mit Nullrandwerten zurück: wir setzen $w := u - f$. Mit dieser Definition erhalten wir $w = 0$ auf $\partial\Omega$ und

$$\begin{aligned} -\beta \Delta w + w &= -\beta \Delta u + \beta \Delta f + u - f \\ &= \underbrace{-\beta \Delta u + u}_{=0} - f + \beta \Delta f. \end{aligned}$$

Wir erhalten also ein äquivalentes Problem mit Nullrandwerten:

$$\begin{aligned} -\beta \Delta w + w &= \beta \Delta f && \text{in } \Omega \\ w &= 0 && \text{auf } \partial\Omega. \end{aligned}$$

In der schwachen Formulierung erhalten wir (falls $f \in H^1(\Omega)$) aus

$$\int_{\Omega} \beta \nabla u \cdot \nabla v + (u - f)v = 0$$

mit $u = w + f$ die Gleichung

$$\int_{\Omega} \beta \nabla w \cdot \nabla v + wv + \beta \nabla f \cdot \nabla v = 0,$$

d.h. nach der Diskretisierung

$$LW + MW + LF = 0.$$

Die (Approximation an die) Funktion u erhält man nach Definition von w durch $u = w + f$.

Wir hatten aber gesehen, dass zumindest für die Bildglättung Nullrandwerte nicht unbedingt ideal sind. Daher betrachten wir wie in Abschnitt 1.4 die Randwerte

$$\nabla u \cdot \nu = 0.$$

2 Finite Elemente

Diese erhält man, indem man das Problem (P^w) mit $H^1(\Omega)$ statt $H_0^1(\Omega)$ betrachtet: Gesucht wird eine Funktion $u \in H^1(\Omega)$, so dass für alle $v \in H^1(\Omega)$

$$\int_{\Omega} \beta \nabla u \cdot \nabla v + (u - f)v = 0$$

gilt. Falls u zweimal differenzierbar ist, erhält man mit Hilfe des Satzes von Gauß

$$\int_{\Omega} -\beta \Delta u v + (u - f)v + \beta \int_{\partial\Omega} \nabla u \cdot \nu v = 0.$$

Verwendet man Testfunktionen v mit Nullrandwerten (diese sind zulässig da $H_0^1(\Omega) \subset H^1(\Omega)$), so ergibt sich wie bisher

$$-\beta \Delta u + u = f.$$

Daher ist

$$\int_{\partial\Omega} \nabla u \cdot \nu v = 0$$

für alle v , also

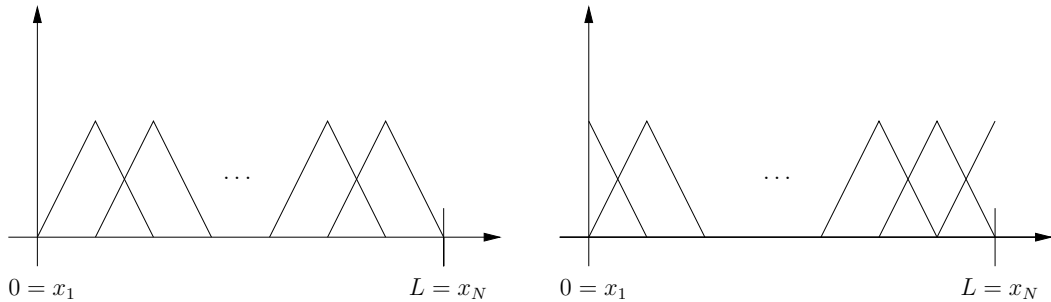
$$\nabla u \cdot \nu = 0.$$

Dies nennt man *natürliche Randwerte*. In der schwachen Formulierung unterscheiden sich natürliche von Null-Randwerten *nur* durch Verwendung von $H^1(\Omega)$ statt $H_0^1(\Omega)$. Analog verwendet man in der Diskretisierung V_h statt V_h^0 .

Im eindimensionalen Fall betrachten wir die folgenden Räume:

$$V_h^0 = \text{span}\{\varphi_i \mid i = 2, \dots, N-1\}$$

$$V_h = \text{span}\{\varphi_i \mid i = 1, \dots, N\}$$



Die rechte Abbildung enthält zusätzlich die „abgeschnittenen“ Basisfunktionen zu den Randknoten.

Die entstehenden Masse- und Steifigkeitsmatrizen sind $N \times N$ -Matrizen. Die Struktur entspricht den Matrizen für Nullrandwerte bis auf den linken oberen und den rechten unteren Eintrag. Bei den Integralen $\int_{\Omega} \varphi_1 \varphi_1$, $\int_{\Omega} \nabla \varphi_1 \cdot \nabla \varphi_1$ und den analogen Integralen von φ_N wird jeweils nur über „halbe“ Basisfunktionen integriert, so dass sich dort auch jeweils der halbe Wert ergibt. Wir erhalten also die $N \times N$ Matrizen

$$M = \frac{h}{6} \begin{pmatrix} 2 & 1 & & & 0 \\ 1 & 4 & 1 & & \\ & \ddots & \ddots & \ddots & \\ & & 1 & 4 & 1 \\ 0 & & & 1 & 2 \end{pmatrix}, \quad L = \frac{\beta}{h} \begin{pmatrix} 1 & -1 & & & 0 \\ -1 & 2 & -1 & & \\ & \ddots & \ddots & \ddots & \\ & & -1 & 2 & -1 \\ 0 & & & -1 & 1 \end{pmatrix}.$$

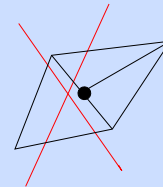
2.4 Finite Elemente auf Dreiecksgittern

Als Beispiel wollen wir die Generalisierung von DHM (digitalen Höhenmodellen) betrachten. Hier besteht die Aufgabe darin, die grobskaligen bzw. die „wichtigen“ Geländeeigenschaften zu finden. Dazu betrachten wir Dreiecksgitter in den Ebenen (z.B. Gauß-Krüger-Koordinaten) und der Funktionswert u stellt die Höhe des Gitterpunktes dar. Dies führt auf die bekannte partielle Differentialgleichung

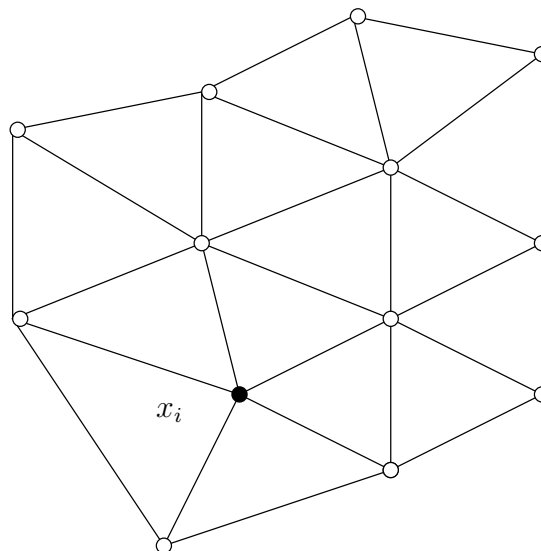
$$\begin{aligned} -\beta \Delta u + u &= f && \text{in } \Omega, \\ \nabla u \cdot \nu &= 0 && \text{auf } \partial\Omega. \end{aligned}$$

Definition 2.11. Eine Triangulierung eines polygonalen Gebietes Ω_h ist eine Unterteilung von Ω_h in Dreiecke, wobei gilt

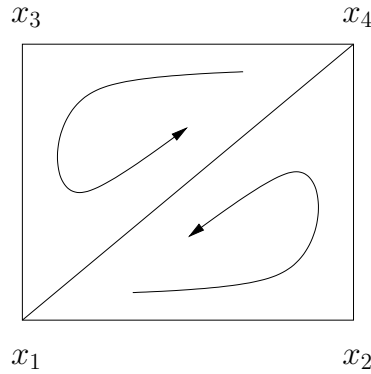
- ganz Ω_h wird überdeckt
- zwei Dreiecke schneiden sich entweder gar nicht, in einem Punkt, oder einer kompletten Kante, es gibt also keine „hängenden“ Knoten
- die Winkel sind nicht „zu spitz“



Eine derartige Triangulierung hat beispielsweise die Form:



Um eine Triangulierung im Rechner zu implementieren, speichert man ihre Punkte und Dreiecke ab. Nehmen wir als Beispiel das einfache Gitter



Liste von Punkten:

Liste von Dreiecken:

$$x_1 = (0, 0)$$

$$T_1 = (x_1, x_2, x_4)$$

$$x_2 = (1, 0)$$

$$T_2 = (x_4, x_3, x_1)$$

$$x_3 = (0, 1)$$

$$x_4 = (1, 1)$$

Dabei ist es sinnvoll, eine Konvention für die Orientierung der Dreiecke zu haben, z.B. gegen den Uhrzeigersinn.

In eine Datei schreibe man üblicherweise zunächst die Anzahl der Punkte und die Anzahl der Dreiecke, listet dann die Koordinaten der Punkte und schließlich für jedes Dreieck die Indizes der Punkte:

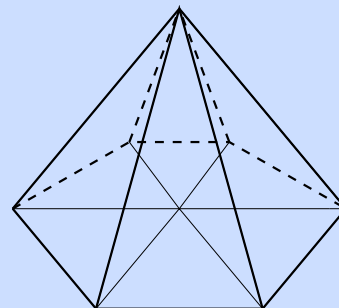
```

4      2
0.0    0.0
1.0    0.0
0.0    1.0
1.0    1.0
1      2      4
4      3      1
    
```

Auch in 2d definieren wir wieder eine Basis für stetige, stückweise affine Funktionen:

Definition 2.12. Zu einer Triangulierung wie oben bezeichnet φ_i die stetige stückweise affine Funktion, für die am Knoten x_j gilt:

$$\varphi_i(x_j) = \delta_{ij} = \begin{cases} 1 & \text{für } i = j \\ 0 & \text{sonst.} \end{cases}$$



Analog zum 1d-Fall bezeichnen wir mit V_h den von diesen Funktionen aufgespannten Vektorraum und setzen $V_h^0 := \text{span} \left\{ \varphi_i(x) \mid x_i \text{ liegt im Inneren von } \Omega_h \right\}$.

Wie in 1d gilt

$$u_h(x) = \sum_{i=1}^N U_i \varphi_i(x).$$

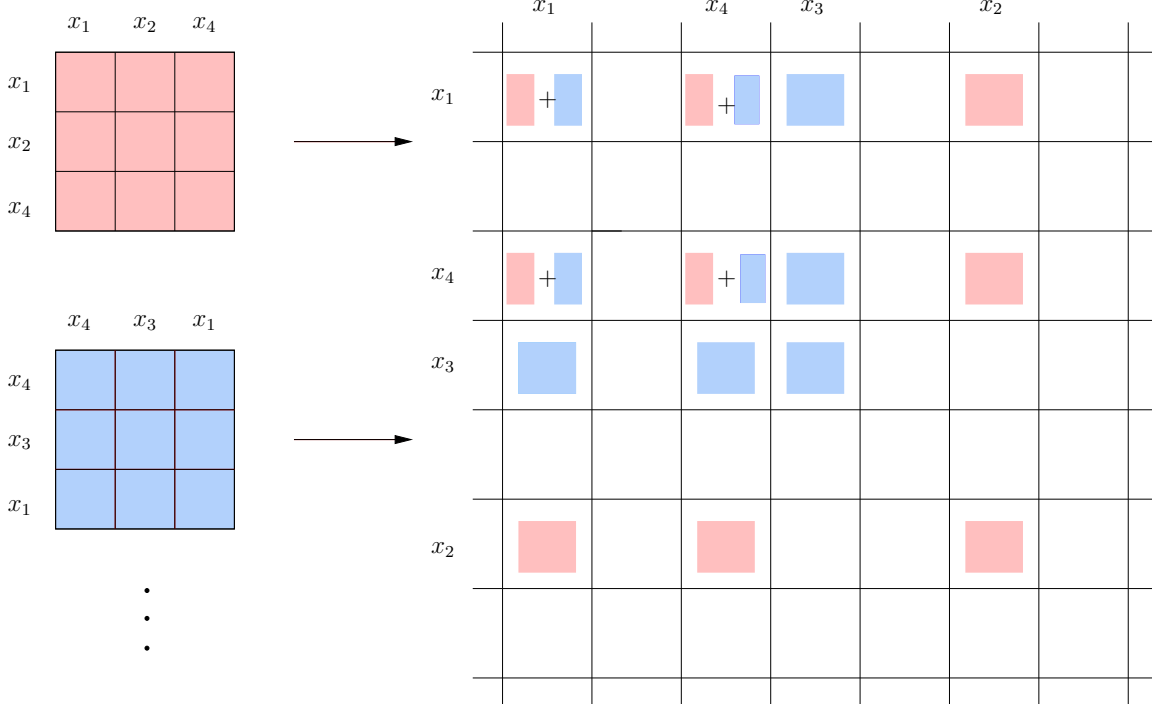
2.5 Assemblierung der Matrizen

Um den (i, j) -ten Eintrag der Massematrix (und analog der Steifigkeitsmatrix) zu ermitteln, berechnen wir $\int_{\Omega} \varphi_i \varphi_j$ nicht direkt für das gesamte Ω , sondern werten das Integral

$$\int_T \varphi_i \varphi_j$$

für jedes einzelne Dreieck T und für alle φ_i, φ_j , die auf diesem Dreieck nicht Null sind, aus.

Hierbei wissen wir genau, welche φ_i auf dem Dreieck T nicht Null sind: diejenigen, die zu einer Ecke des Dreiecks gehören. Für $T = (x_1, x_2, x_3)$ sind wären das also $i, j \in \{1, 2, 3\}$. Dies führt zu 3×3 -Kombinationen, die auf die entsprechenden Stellen der Matrix verteilt werden, dabei addiert man die Beiträge aus allen beteiligten Dreiecken auf.



Zunächst ein Wort zur Notation: Um die Menge an Bezeichnungen übersichtlich zu halten, sind wir im Folgenden etwas ungenau. Wir unterscheiden nicht explizit zwischen

- lokalen $(1, 2, 3)$ und globalen $(1, \dots, N)$ Kantennummern,
- der Punktmenge, aus der eine Seite besteht und dem zugehörigen Kantenvektor.

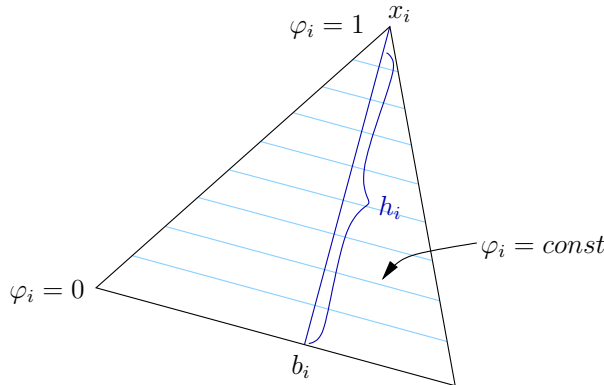
Die Bedeutung ergibt sich jeweils aus dem Kontext.

Steifigkeitsmatrix

Die Steifigkeitsmatrix ist einfacher, denn $\nabla\varphi_i$ und $\nabla\varphi_j$ sind konstante Vektoren auf T , also

$$\int_T \nabla\varphi_i \cdot \nabla\varphi_j = |T| \nabla\varphi_i \cdot \nabla\varphi_j.$$

Hierfür müssen wir die Gradienten $\nabla\varphi_i$ berechnen:



$$|\nabla\varphi_i| = \frac{|\varphi_i(x_i) - \varphi_i(b_i)|}{|x_i - b_i|} = \frac{1}{h_i}$$

Sei n_i die Normale an der Kante gegenüber von x_i . Dann gilt

$$\nabla\varphi_i = \frac{n_i}{h_i}$$

und wegen $n_i = D^{90} \frac{s_i}{|s_i|}$ auch

$$n_i \cdot n_j = \frac{s_i}{|s_i|} \cdot \frac{s_j}{|s_j|}.$$

Somit erhalten wir

$$\int_T \nabla\varphi_i \cdot \nabla\varphi_j = |T| \frac{n_i \cdot n_j}{h_i h_j} = \frac{|T|}{h_i h_j} \frac{s_i}{|s_i|} \cdot \frac{s_j}{|s_j|} = \frac{s_i \cdot s_j}{4 |T|}.$$

Hierbei haben wir

$$|T| = \frac{1}{2} |s_i| h_i \iff h_i = \frac{2 |T|}{|s_i|}.$$

verwendet. Insgesamt erhalten wir

$$\begin{aligned} \int_T |\nabla\varphi_i|^2 &= \frac{|s_i|^2}{4 |T|}, \\ \int_T \nabla\varphi_i \cdot \nabla\varphi_j &= \frac{s_i \cdot s_j}{4 |T|}. \end{aligned}$$

Zur Berechnung der Einträge brauchen wir noch die Fläche des Dreiecks. Dazu verwenden wir die Formel

$$|T| = \frac{1}{2} \left| \det \begin{pmatrix} -s_1 & - \\ -s_2 & - \end{pmatrix} \right|.$$

Massematrix

Die Einträge der Massematrix berechnen wir durch eine numerische Quadraturformel, die auf den vorkommenden Funktionen den exakten Wert des Integrals liefert.

Lemma 2.13. *Sei T ein Dreieck mit den Seitenmitten a, b, c und p sei quadratisches Polynom. Dann gilt*

$$\int_T p(x) \, dx = \frac{|T|}{3} (p(a) + p(b) + p(c)) \quad (2.2)$$

Beweis. Zunächst zeigen wir das Lemma für das Einheitsdreieck $\hat{T}$ mit den Ecken $(0, 0)$, $(0, 1)$ und $(1, 0)$. Es reicht, die Aussage für die Basis $\{1, x, y, xy, x^2, y^2\}$ der quadratischen Polynome zu zeigen. Aus Symmetriegründen genügt es, die Funktionen $\{1, x, xy, x^2\}$ zu betrachten. Es gilt

- $p(x, y) = 1$

$$\int_T 1 \, dx = \frac{1}{2}, \quad \frac{1}{3}(1 + 1 + 1) = \frac{1}{2} \quad \checkmark$$

- $p(x, y) = x$

$$\begin{aligned} \int_T x \, dx &= \int_0^1 \int_0^{1-y} x \, dx \, dy = \int_0^1 \frac{1}{2} (1-y)^2 \, dy = \frac{1}{2} \int_0^1 1 - 2y + y^2 \, dy \\ &= \frac{1}{2} \left[y - y^2 + \frac{1}{3} y^3 \right]_0^1 = \frac{1}{2} \left(1 - 1 + \frac{1}{3} \right) = \frac{1}{6}, \\ \frac{1}{3} \left(0 + \frac{1}{2} + \frac{1}{2} \right) &= \frac{1}{6} \quad \checkmark \end{aligned}$$

- $p(x, y) = xy$

$$\begin{aligned} \int_T xy \, dx &= \int_0^1 \int_0^{1-y} xy \, dx \, dy = \int_0^1 \frac{1}{2} (1-y)^2 y \, dy = \frac{1}{2} \int_0^1 y - 2y^2 + y^3 \, dy \\ &= \frac{1}{2} \left[\frac{1}{2} y^2 - \frac{2}{3} y^3 + \frac{1}{4} y^4 \right]_0^1 = \frac{1}{2} \left(\frac{1}{2} - \frac{2}{3} + \frac{1}{4} \right) = \frac{1}{2} \cdot \frac{6 - 8 + 3}{12} = \frac{1}{24}, \\ \frac{1}{3} \left(0 + 0 + \frac{1}{4} \right) &= \frac{1}{24} \quad \checkmark \end{aligned}$$

- $p(x, y) = x^2$

$$\begin{aligned} \int_T x^2 \, dx &= \int_0^1 \int_0^{1-y} x^2 \, dx \, dy = \int_0^1 \frac{1}{3} (1-y)^3 \, dy \\ &= \frac{1}{3} \int_0^1 1 - 3y + 3y^2 - y^3 \, dy = \frac{1}{3} \left[y - \frac{3}{2} y^2 + y^3 - \frac{1}{4} y^4 \right]_0^1 \\ &= \frac{1}{3} \left(1 - \frac{3}{2} + 1 - \frac{1}{4} \right) = \frac{1}{3} \cdot \frac{8 - 12 + 8 - 2}{8} = \frac{1}{12}, \\ \frac{1}{3} \left(0 + \frac{1}{4} + \frac{1}{4} \right) &= \frac{1}{12} \quad \checkmark \end{aligned}$$

Wir haben also gezeigt, dass

$$\int_{\hat{T}} p(x) \, dx = \frac{|\hat{T}|}{3} (p(a) + p(b) + p(c)).$$

Wir können ein beliebiges Dreieck T mit einer affinen Abbildung auf das Einheitsdreieck $\hat{T}$ abbilden. Aufgrund des Transformationssatzes wissen wir, dass das Integral dabei genau wie die rechte Seite mit der Flächenänderung skaliert. Da eine affine Abbildung Seitenmitten auf Seitenmitten abbildet und p nach der Transformation ein quadratisches Polynom bleibt, gilt die Formel auch auf T . $\square$

Wie sehen nun die Einträge der (lokalen) Massenmatrix aus? Die Basisfunktionen hat in der Mitte einer Seite den Wert 0, in den beiden anderen der Wert $\frac{1}{2}$. Mit Lemma 2.13 erhalten wir

$$\begin{aligned} \int_T \varphi_i \varphi_i &= \frac{|T|}{6}, \\ \int_T \varphi_i \varphi_j &= \frac{|T|}{12} \quad \text{für } i \neq j \end{aligned}$$

Für die lokale Massenmatrix erhalten wir also

$$M_T = \frac{|T|}{12} \begin{pmatrix} 2 & 1 & 1 \\ 1 & 2 & 1 \\ 1 & 1 & 2 \end{pmatrix}.$$

Bemerkung 2.14 (Mass lumping). *Einfacher, aber ungenauer ist die Quadraturformel*

$$\int_T p(x) \approx \frac{|T|}{3} (p(x_1) + p(x_2) + p(x_3)),$$

wobei x_i die Ecken des Dreiecks sind. Sie ist nur für affine p exakt. Der Vorteil ist, dass wegen $\varphi_i(x_j) = \delta_{ij}$

$$\begin{aligned} \int_T \varphi_i \varphi_i &\approx \frac{|T|}{3} (1 + 0 + 0) \\ \int_T \varphi_i \varphi_j &\approx \frac{|T|}{3} (0 + 0 + 0) \quad \text{falls } i \neq j. \end{aligned}$$

Damit ist die lokale Massematrix

$$M_T = \frac{|T|}{3} \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

und so ist auch die Massematrix M eine Diagonalmatrix (und damit sehr effizient in der Handhabung). Man bezeichnet diesen Ansatz als „lumped masses“.

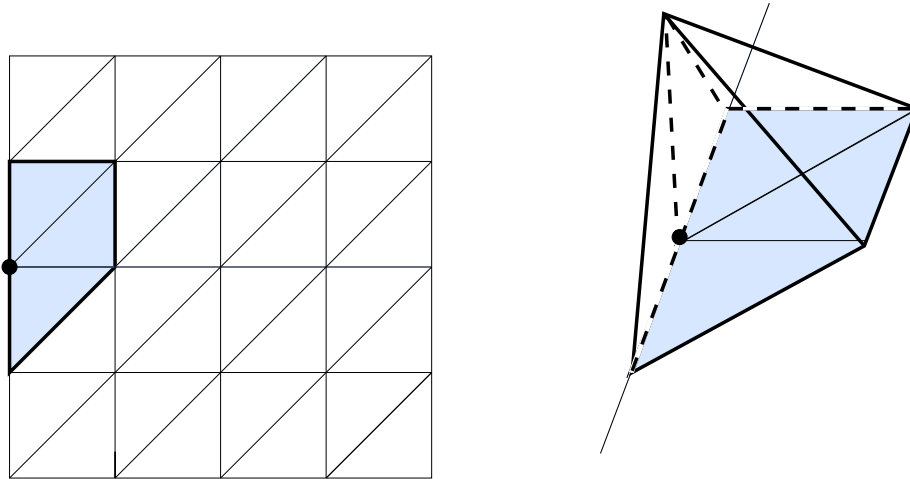
Assemblierung von Matrizen in MATLAB

In MATLAB kann man das eigentliche Assemblieren dem Aufruf `L = sparse ( I, J, V )`; überlassen, wenn man die Vektoren **I**, **J** und **V** nacheinander mit den 9 Beiträgen jedes einzelnen Dreiecks füllt. Beim Aufruf von `sparse` werden dann alle Werte, die zu mehrfach auftretenden Indexpaaren (i, j) gehören, automatisch aufaddiert.

2.6 Randwerte

Natürliche Randwerte

Wie schon in 1d wird für Randknoten nur der innere Teil der Hütchenfunktion betrachtet. Wenn man die Matrix wie oben beschrieben dreiecksweise assembliert, geschieht dies automatisch: Es werden für alle Dreiecke im Gebiet die auf dem jeweiligen Dreieck lebenden Teile der Basisfunktionen integriert.



Null-Randwerte

Um Nullrandwerte zu erhalten, betrachten wir wieder den Raum V_h^0 , d.h. nur diejenigen Ansatzfunktionen, die zu inneren Knoten gehören. Unbekannte sind also nur diejenigen Knoten, die nicht auf dem Rand liegen. Beim Assemblieren beider Matrizen auf einem Dreieck verwendet man nur die Nicht-Rand-Knoten.

Schwierigkeit: die Nummern der Randknoten sind über die Liste der Knoten verteilt. Um dies zu umgehen, gibt es zwei Möglichkeiten.

- eine Umnummerierung der Knoten
- Zu Randknoten werden nur Beiträge auf die entsprechende Diagonale in der Matrix addiert, alle anderen Assemblierungsbeiträge werden ignoriert. Ebenso werden keine Beiträge zur rechten Seite assembliert. Dies führt dazu, dass für einen Randknoten i alle Einträge der i -ten Zeile und i -ten Spalte 0 sind, ausgenommen dem Diagonalelement (welches z.B. 1 ist). Im Vektor f auf der rechten Seite ist der zugehörigen Eintrag ebenfalls 0. (Statt „1“ ist auf der Diagonalen

jeder nicht-Null-Eintrag möglich.)

$$i \rightarrow \begin{pmatrix} & & & i \\ & & & \downarrow \\ & & & 0 \\ & & & | \\ & & & 0 \\ 0 & \cdots & 0 & 1 & 0 & \cdots & 0 \\ & & & | \\ & & & 0 \end{pmatrix}$$

Dies ergibt dann $u_i = 0$ ohne jede Kopplung an andere Freiheitsgrade.

Alternativ: Eine normale Assemblierung (für natürliche Randwerte) durchführen, anschließend für Randknoten i die Zeile i und Spalte i auf Null, den Diagonaleintrag auf 1 setzen (ergibt dieselbe Matrix) und den Eintrag i der rechten Seite ebenfalls Null setzen. Wir haben allerdings schon bemerkt, dass das Nullsetzen von Zeilen langsam sein kann.

Nicht-Null-Randwerte

Genau wie in Abschnitt 2.3 wird dieser Fall auf den Fall von Nullrandwerten zurückgeführt.

Programmieraufgabe 4. Implementieren Sie das beschriebene Finite Elemente Verfahren zur Berechnung von

$$\begin{aligned} -\beta \Delta u + u &= f & \text{in } \Omega \\ \nabla u \cdot \nu &= 0 & \text{auf } \partial\Omega. \end{aligned}$$

zur Generalisierung digitaler Höhenmodelle auf Dreiecksgittern. Achten Sie auf eine effiziente Assemblierung der Matrix.

Als Ergebnis für $\beta = 5 \cdot 10^4$ (dieser Wert ist sicherlich zu groß, zerstört also zu viele Details) erhalten wir

