

Analysis IV

INHALTSVERZEICHNIS

Einleitung	4
1. Topologische Räume I	5
1.1. Mengentheoretische Begriffe	5
1.2. Metrische Räume	5
1.3. Topologische Räume	7
1.4. Die Teilraumtopologie	10
1.5. Abgeschlossene Teilmengen	11
1.6. Stetige Abbildungen	12
1.7. Homöomorphismen	12
1.8. Kompakte topologische Räume	13
2. Definition und Beispiele von Untermannigfaltigkeiten	17
2.1. Die Definition von Untermannigfaltigkeiten	17
2.2. Untermannigfaltigkeiten und Parametrisierungen	21
2.3. Untermannigfaltigkeiten und Graphen	22
2.4. Der Satz vom regulären Wert	25
2.5. Differenzierbare Funktionen auf Untermannigfaltigkeiten	26
3. Integration und Untermannigfaltigkeiten	28
3.1. Erinnerung: Maß- und Integrationstheorie	28
3.2. Motivation für Integrale auf Untermannigfaltigkeiten	33
3.3. Einschub: Lineare Algebra	34
3.4. Integration auf Untermannigfaltigkeiten I	39
3.5. Integration auf Untermannigfaltigkeiten II	41
3.6. Beispiele	43
3.7. Integration auf Untermannigfaltigkeiten III	44
3.8. Das k -dimensionale Volumen	46
4. Untermannigfaltigkeiten mit Rand	48
5. Zusammenhängende topologische Räume	55
6. Der Gaußsche Integralsatz	58
6.1. Tangentialvektoren zu Untermannigfaltigkeiten	58
6.2. Vektorfelder auf Untermannigfaltigkeiten	60
6.3. Formulierung des Gaußschen Integralsatzes	65
6.4. Beweis des Gaußschen Integralsatzes I: Quader	68
6.5. Beweis des Gaußschen Integralsatzes II: Stempel	69
6.6. Beweis des Gaußschen Integralsatzes III: Der allgemeine Fall	73
6.7. Der Gaußsche Integralsatzes für Untermannigfaltigkeiten mit Kanten	76
6.8. Ausblick	79
7. Kategorien und Funktoren	80

7.1.	Induzierte Abbildungen auf Tangentialräumen	80
7.2.	Duale Vektorräume	83
7.3.	Kategorien	84
7.4.	Funktoren	86
8.	Differentialformen erster Ordnung	90
8.1.	Definition von Differentialformen erster Ordnung	90
8.2.	Integrierbarkeit von 1-Formen	97
8.3.	Integration und Parametertransformationen	100
8.4.	Orientierte eindimensionale Untermannigfaltigkeiten	102
9.	Differentialformen höherer Ordnung	105
9.1.	Motivation	105
9.2.	Alternierende Multilinearformen I: Definition	106
9.3.	Alternierende Multilinearformen II: Das Dachprodukt von Linearformen	109
9.4.	Differentialformen auf Untermannigfaltigkeiten	113
9.5.	Induzierte Abbildungen	114
9.6.	Glatte Differentialformen auf Untermannigfaltigkeiten	115
9.7.	Integration von Differentialformen I	117
9.8.	Orientierte Vektorräume und Untermannigfaltigkeiten	120
9.9.	Integration von Differentialformen II	126
9.10.	Integration von Differentialformen III	127
9.11.	Null-dimensionale Untermannigfaltigkeiten	132
10.	Vorbereitungen für den Satz von Stokes	133
10.1.	Orientierte Untermannigfaltigkeiten mit Rand	133
10.2.	Das Differential von Differentialformen	136
11.	Der Satz von Stokes	145
11.1.	Formulierung des Satz von Stokes	145
11.2.	Beweis vom Satz von Stokes	146
11.3.	Der klassische Satz von Stokes und der Satz von Green	148
11.4.	Anwendung auf holomorphe Funktionen	150
12.	Länge von Wegen und Geodäten	153
13.	Riemannsche Untermannigfaltigkeiten	157
13.1.	Definition von Riemannschen Untermannigfaltigkeiten	157
13.2.	Geodäten auf der Sphäre	167
14.	Hyperbolische Geometrie	171
14.1.	Der hyperbolische Raum	171
14.2.	Das Poincaré-Model für den hyperbolischen Raum	181
15.	Volumen und Krümmung auf Riemannschen Untermannigfaltigkeiten	185
15.1.	Volumen auf Riemannschen Untermannigfaltigkeiten	185
15.2.	Winkel und Dreiecke auf Riemannschen Untermannigfaltigkeiten	187
15.3.	Krümmung auf Riemannschen Untermannigfaltigkeiten	194
16.	Topologische Räume II	197
16.1.	Definition eines topologischen Raums und Beispiele	197

16.2.	Basis einer Topologie	198
16.3.	Das Produkt von topologischen Räumen	202
17.	Quotientenräume	205
17.1.	Äquivalenzrelationen	205
17.2.	Die Quotiententopologie	206
17.3.	Beispiele: Zylinder, Torus, das Möbius Band und die Kleinsche Flasche	209
17.4.	Beispiele: Flächen von höherem Geschlecht	213
18.	Topologische Mannigfaltigkeiten	217
18.1.	Zweitabzählbare topologische Räume	217
18.2.	Definition von topologischen Mannigfaltigkeiten	218
18.3.	Gruppenoperationen	220
18.4.	Stetige Operationen	222
18.5.	Gruppenoperationen auf topologischen Mannigfaltigkeiten	228
19.	Mannigfaltigkeiten	231
19.1.	Definition und Beispiele von Mannigfaltigkeiten	231
19.2.	Untermannigfaltigkeiten und Produkte von Mannigfaltigkeiten	235
19.3.	Die Fläche von Geschlecht 2 als Mannigfaltigkeit	237
20.	Riemannsche Mannigfaltigkeiten	241
20.1.	Der Tangentialraum einer Mannigfaltigkeit	241
20.2.	Riemannsche Mannigfaltigkeiten und Integration von Funktionen	245
20.3.	Mannigfaltigkeiten mit Rand	246
20.4.	Differentialformen auf Mannigfaltigkeiten	247
21.	Krümmung auf Mannigfaltigkeiten	252
21.1.	Krümmung des Raums	252
21.2.	Krümmung von 2-dimensionalen Mannigfaltigkeiten	253
21.3.	Flächen von höherem Geschlecht	257
21.4.	Reguläre n -Ecke und hyperbolische Geometrie	258
21.5.	Hyperbolische Metriken auf Flächen von Geschlecht 2	259
21.6.	Die Klassifizierung von geschlossenen Flächen	264
21.7.	Die Geometrie von 3-dimensionalen Mannigfaltigkeiten	266
22.	Die de Rham-Kohomologie von Mannigfaltigkeiten	268
22.1.	Quotientenvektorräume	268
22.2.	Die Definition der de Rham-Kohomologie	269
22.3.	Berechnungen der de Rham-Kohomologie	269
22.4.	Die Funktorialität der de Rham-Kohomologie	271
22.5.	Das Dachprodukt auf de Rham-Kohomologie	273
22.6.	Die Sphäre und der Torus sind nicht diffeomorph	275
23.	Ausblick	278

EINLEITUNG

In der Analysis II hatten wir schon den Begriff einer k -dimensionalen Untermannigfaltigkeit von $\mathbb{R}^n$ eingeführt. Vereinfacht gesprochen handelt es sich hierbei um eine “ k -dimensionale Teilmenge von $\mathbb{R}^n$ ”. Beispielsweise ist die Sphäre S^2 eine 2-dimensionale Untermannigfaltigkeit von $\mathbb{R}^3$. In der Analysis IV Vorlesung wollen wir folgende Themen behandeln:

- (1) Das Integral von reellwertigen Funktionen auf Untermannigfaltigkeiten.
- (2) Die klassischen Sätze der Vektoranalysis, z.B. der Satz von Gauß, der Satz von Stokes, und die Verallgemeinerungen auf höher dimensionale Untermannigfaltigkeiten.
- (3) Den Begriff der Krümmung einer Untermannigfaltigkeit.
- (4) Die Erweiterung des Begriffs von Untermannigfaltigkeiten von $\mathbb{R}^n$ zu Mannigfaltigkeiten.
- (5) Wir wollen beweisen, dass die Sphäre S^2 nicht diffeomorph zu einem Torus ist.

1. TOPOLOGISCHE RÄUME I

In diesem Kapitel führen wir den Begriff des topologischen Raumes ein. Diese Räume abstrahieren den Begriff von “offenen” und “abgeschlossenen” Teilmengen. Insbesondere übertragen sich die meisten Definitionen und Aussagen aus der Analysis II über metrische Räume, welche nur mithilfe von offenen Teilmengen eingeführt wurden, auf topologische Räume. Insbesondere macht es Sinn von kompakten topologischen Räumen und von stetigen Abbildungen zwischen topologischen Räumen zu sprechen.

1.1. Mengentheoretische Begriffe. Bevor wir die topologischen Räume einführen erinnern wir an ein paar grundlegende Definitionen und Aussagen aus der Mengentheorie:

- (1) Wir bezeichnen mit $\emptyset$ die leere Menge.
- (2) Es sei X eine Menge. Wir bezeichnen mit $\mathcal{P}(X)$ die *Potenzmenge von X* , d.h. die Menge aller Teilmengen von X .
- (3) Es seien X und Y Mengen, dann bezeichnen wir mit

$$X \times Y := \{(x, y) \mid x \in X, y \in Y\}$$

die Menge aller angeordneter Paare.

- (4) Es sei X eine Menge. Eine *Familie $A_i, i \in I$ von Teilmengen von X* besteht aus einer Indexmenge I (z.B. $I = \{1, \dots, k\}$ oder $I = \mathbb{Z}$), und einer Teilmenge $A_i \subset X$ für jedes $i \in I$.
- (5) Es sei nun $A_i, i \in I$ eine Familie von Teilmengen von X , dann gilt

$$X \setminus \bigcap_{i \in I} A_i = \bigcup_{i \in I} (X \setminus A_i) \quad \text{und} \quad X \setminus \bigcup_{i \in I} A_i = \bigcap_{i \in I} (X \setminus A_i).$$

Diese Aussagen werden manchmal die “de Morgansche Gesetze” genannt.

- (6) Es seien $A, B \subset X$ zwei Teilmengen. Wir sagen A und B sind *disjunkt*, wenn sich A und B nicht schneiden, d.h. wenn $A \cap B = \emptyset$.
- (7) Es sei $f: X \rightarrow Y$ eine Abbildung, dann gilt:

$$\begin{aligned} U &\supset f(f^{-1}(U)) && \text{für jede Teilmenge } U \subset Y \\ U &= f(f^{-1}(U)) && \text{für jede Teilmenge } U \subset Y, \text{ wenn } f \text{ surjektiv} \\ U &\subset f^{-1}(f(U)) && \text{für jede Teilmenge } U \subset X \\ f^{-1}\left(\bigcup_{i \in I} U_i\right) &= \bigcup_{i \in I} f^{-1}(U_i) && \text{für Teilmengen } U_i \in Y, i \in I \\ f\left(\bigcup_{i \in I} U_i\right) &= \bigcup_{i \in I} f(U_i) && \text{für Teilmengen } U_i \in X, i \in I. \end{aligned}$$

1.2. Metrische Räume. Wir hatten die metrischen Räume schon in Analysis II eingeführt und behandelt. Wir erinnern noch einmal kurz an die Definition:

Definition. Ein *metrischer Raum* ist ein Paar (X, d) bestehend aus einer Menge X und einer *Metrik* d auf X , d.h. einer Abbildung

$$\begin{aligned} X \times X &\rightarrow \mathbb{R}_{\geq 0} := \{x \in \mathbb{R} \mid x \geq 0\}, \\ (x, y) &\mapsto d(x, y), \end{aligned}$$

mit folgenden Eigenschaften:

- (1) $d(x, y) = 0$ genau dann, wenn $x = y$,
- (2) für alle $x, y \in X$ gilt

$$d(x, y) = d(y, x) \quad (\text{Symmetrie})$$

- (3) für alle $x, y, z \in X$ gilt

$$d(x, z) \leq d(x, y) + d(y, z) \quad (\text{Dreiecksungleichung}).$$

Wir nennen $d(x, y)$ manchmal den *Abstand von x zu y* .

Das vielleicht wichtigste Beispiel eines metrischen Raumes ist gegeben durch $X = \mathbb{R}^n$ mit der euklidischen Metrik, d.h. mit der Metrik, welche definiert ist durch

$$d(x, y) = \|x - y\| = \sqrt{\sum_{i=1}^n (x_i - y_i)^2}.$$

↑
wie üblich bezeichnen wir für $a \in \mathbb{R}^n$
die Koordinaten mit $a_1, \dots, a_n$

Wenn wir nichts anderes schreiben, dann betrachten wir $\mathbb{R}^n$ als metrischen Raum bezüglich der euklidischen Metrik.

Definition. Es sei (X, d) ein metrischer Raum.

- (1) Es sei $x \in X$ und $r \in \mathbb{R}$. Wir bezeichnen

$$B_r(x) := \{y \in X \mid d(x, y) < r\}$$

als die *offene r -Kugel um x* .

- (2) Eine Teilmenge $U \subset X$ heißt *offen*, wenn es zu jedem $x \in U$ ein $r > 0$ gibt, so dass $B_r(x)$ noch in U enthalten ist.
- (3) Wir sagen $A \subset X$ ist *abgeschlossen*, wenn $X \setminus A$ offen ist.

Mithilfe der Definitionen kann man leicht folgendes Lemma beweisen, welches wir schon als Lemma 1.7 aus der Analysis II kennen.

Lemma 1.1. *Es sei X ein metrischer Raum. Dann gilt*

- (1) $\emptyset$ und X sind offen,
- (2) der Schnitt endlich vieler offener Mengen ist wiederum offen, d.h. es gilt

$$U_1, \dots, U_k \text{ offen} \implies U_1 \cap \dots \cap U_k \text{ ist offen},$$

- (3) die Vereinigung beliebig vieler offener Mengen ist wiederum offen, d.h. es gilt

$$U_i, i \in I \text{ offen} \implies \bigcup_{i \in I} U_i \text{ ist offen}.$$

Definition. Es sei $f: X \rightarrow Y$ eine Abbildung zwischen metrischen Räumen. Wir sagen f ist *stetig*, wenn gilt¹

$$\forall_{x \in X} \quad \forall_{\epsilon > 0} \quad \exists_{\delta > 0} \quad \forall_{y \in B_\delta(x)} \quad d(f(x), f(y)) < \epsilon.$$

Diese Definition der Stetigkeit ist eine Verallgemeinerung der “ $\epsilon - \delta$ ”-Definition von Stetigkeit für Funktionen auf $\mathbb{R}$.

Mit der Definition von offenen Mengen können wir nun folgenden Satz formulieren:

Satz 1.2. *Es sei $f: X \rightarrow Y$ eine Abbildung zwischen zwei metrischen Räumen. Dann gilt:*²

$$\begin{aligned} & f \text{ ist stetig} \\ \iff & \text{für jede offene Menge } U \text{ in } Y \text{ ist } f^{-1}(U) \text{ offen in } X \\ \iff & \text{für jede abgeschlossene Menge } A \text{ in } Y \text{ ist } f^{-1}(A) \text{ abgeschlossen in } X. \end{aligned}$$

Dieser Satz wurde als Satz 2.8 in Analysis II formuliert. Der Beweis war eine Übungsaufgabe in Analysis II, und ist jetzt wiederum eine (weiterhin sehr sinnvolle) Übungsaufgabe.

Die Formulierung von Stetigkeit mit offenen Mengen kann man sich nicht nur leichter merken als die “ $\epsilon - \delta$ ”-Definition, diese Umformulierung erleichtert auch das Beweisen vieler Aussagen. Beispielsweise läßt sich folgendes Lemma jetzt ganz einfach beweisen:

Lemma 1.3. *Es seien $f: X \rightarrow Y$ und $g: Y \rightarrow Z$ stetige Abbildungen zwischen metrischen Räumen. Dann ist die Verknüpfung $g \circ f: X \rightarrow Z$ ebenfalls stetig.*

Beweis. Aus Satz 1.2 folgt, dass es genügt zu zeigen, dass für jede offene Teilmenge $U \subset Z$ auch das Urbild $(g \circ f)^{-1}(U) \subset X$ offen ist. Es sei also $U \subset Z$ offen. Dann gilt

$$\begin{aligned} (g \circ f)^{-1}(U) &= f^{-1}(g^{-1}(U)) = f^{-1}(\text{der offenen Menge } g^{-1}(U)) = \text{offen.} \\ &\quad \uparrow \qquad \qquad \qquad \uparrow \\ &\quad \text{denn } g \text{ ist stetig und } U \text{ ist offen} \qquad \qquad \text{denn } f \text{ ist stetig.} \quad \square \end{aligned}$$

1.3. Topologische Räume. Wir betrachten jetzt die sogenannten topologischen Räume. Diese kann man als eine Verallgemeinerung des Begriffes des metrischen Raumes auffassen.

Definition. Ein *topologischer Raum* ist ein Paar $(X, \mathcal{T})$, wobei X eine Menge ist und $\mathcal{T}$ eine *Topologie auf X* , d.h. $\mathcal{T}$ ist eine Menge von Teilmengen von X mit folgenden Eigenschaften:

- (1) die leere Menge und die ganze Menge X sind in $\mathcal{T}$ enthalten,
- (2) der Durchschnitt von *endlich vielen* Mengen in $\mathcal{T}$ ist wiederum eine Menge in $\mathcal{T}$,
- (3) die Vereinigung von *beliebig vielen* Mengen in $\mathcal{T}$ ist wiederum eine Menge in $\mathcal{T}$.

Die Mengen in $\mathcal{T}$ werden *offen* genannt.

Beispiel.

¹Wenn nichts anderes angegeben ist, dann wird die Metrik immer mit d bezeichnet, d.h. “Es sei X ein metrischer Raum” ist kurz für “Es sei (X, d) ein metrischer Raum”.

²Man kann die Aussage noch etwas knackiger formulieren: f ist stetig genau dann, wenn das Urbild jeder offenen Menge wiederum offen ist.

- (A) Wenn (X, d) ein metrischer Raum ist, dann folgt aus Lemma 1.1, dass

$$\mathcal{T} := \text{alle offenen Teilmengen von } (X, d)$$

eine Topologie auf X ist. Insbesondere betrachten wir den $\mathbb{R}^n$ mit der üblichen Topologie, außer wir geben explizit eine andere Topologie an.³

- (B) Es sei X eine Menge, dann ist $\mathcal{T} = \{\emptyset, X\}$ eine Topologie auf X . Diese Topologie wird manchmal die *triviale Topologie auf X* genannt.
- (C) Es sei X eine Menge und $\mathcal{T}$ die Menge aller Teilmengen von X . Dann ist $\mathcal{T}$ eine Topologie auf X . Die Topologie wird oft auch als die *diskrete Topologie auf X* bezeichnet.
- (D) Es sei $X = \mathbb{R}$ und es sei $\mathcal{T}$ wie folgt definiert:

$$U \in \mathcal{T} \iff \begin{array}{l} \text{entweder ist } U = \emptyset, \text{ oder} \\ U \text{ ist das Komplement von endlich vielen Punkten in } \mathbb{R}. \end{array}$$

Beispielsweise liegen $\emptyset, \mathbb{R} \setminus \{\pi\}, \mathbb{R} \setminus \{-1, \sqrt{2}\}$ aber auch $\mathbb{R}$ (als Komplement von null Punkten) in $\mathcal{T}$. Es folgt nun leicht aus den “Rechenregeln” für die Mengenlehre, dass $\mathcal{T}$ eine Topologie auf $X = \mathbb{R}$ ist.

- (E) Wir betrachten die Menge

$$X := \mathbb{R} \cup \{\infty\},$$

d.h. X besteht aus $\mathbb{R}$ und einem weiteren Punkt ∞ . Wir sagen $U \subset X$ ist offen⁴, wenn die beiden folgenden Bedingungen erfüllt sind:

- (a) zu jedem Punkt $x \in U \cap \mathbb{R}$ existiert ein $\epsilon > 0$, so dass $(x - \epsilon, x + \epsilon) \subset U$,
- (b) wenn $\infty \in U$, dann gibt es ein $C > 0$, so dass $(-\infty, -C) \cup (C, \infty) \subset U$.

Wir bezeichnen diesen topologischen Raum als die “Gerade mit einem Punkt im Unendlichen”. In Übungsblatt 1 werden wir wiederum sehen, dass dies in der Tat eine Topologie auf X definiert.

- (F) Wir betrachten die Menge

$$X := \mathbb{R} \cup \{*\},$$

d.h. X besteht aus $\mathbb{R}$ und einem weiteren Punkt $*$. Wir sagen $U \subset X$ ist offen, wenn gilt:

- (a) zu jedem Punkt $x \in U \cap \mathbb{R}$ existiert ein $\epsilon > 0$, so dass $(x - \epsilon, x + \epsilon) \subset U$,
- (b) wenn $*$ $\in U$, dann gibt es ein $\epsilon > 0$, so dass $(-\epsilon, 0) \cup (0, \epsilon) \subset U$.

In Übungsblatt 1 werden wir sehen, dass dies in der Tat eine Topologie auf X definiert. Wir bezeichnen diesen topologischen Raum als “Gerade mit zwei Nullen”.⁵

³Dieses Beispiel ist wohl das ursprüngliche Beispiel für einen topologischen Raum und erklärt auch, warum Mengen in einer Topologie “offen” genannt werden.

⁴Wenn wir eine Topologie spezifizieren wollen, müssen wir also festlegen, welche Teilmengen “offen” genannt werden sollen.

⁵Der etwas eigenwillige Name “Gerade mit zwei Nullen” rührt daher, dass 0 und $*$ die gleiche Rolle spielen. Formaler ausgedrückt, die Abbildung, welche 0 und $*$ vertauscht, ist ein Homöomorphismus.

- (G) Es sei $\mathbb{K}$ ein Körper (z.B. $\mathbb{K} = \mathbb{R}$, $\mathbb{K} = \mathbb{C}$ oder $\mathbb{K} = \mathbb{F}_p$, d.h. der Körper mit p Elementen, wobei p eine Primzahl ist). Wir betrachten

$$\mathbb{K}^2 = \{(x, y) \mid x, y \in \mathbb{K}\}.$$

Für eine Menge $\mathcal{F} = \{f_1(x, y), \dots, f_k(x, y)\}$ von Polynomen definieren wir die *Verschwindungsmenge*

$$V(f_1, \dots, f_k) := \{(x, y) \in \mathbb{K}^2 \mid f_1(x, y) = \dots = f_k(x, y) = 0\}.$$

Wir definieren nun eine Topologie of $\mathbb{K}^2$ wie folgt:

$$U \subset \mathbb{K}^2 \text{ ist offen} \quad :\Longleftrightarrow \quad \begin{array}{l} \text{es gibt Polynome } f_1(x, y), \dots, f_k(x, y), \text{ so} \\ \text{dass } U = \mathbb{K}^2 \setminus V(f_1(x, y), \dots, f_k(x, y)). \end{array}$$

Dann ist $\mathcal{T}$ eine Topologie auf $X = \mathbb{K}^2$. Wir überlassen den Beweis dieser Aussage als freiwillige Übungsaufgabe.⁶ Diese Topologie wird die *Zariski-Topologie*⁷ auf $\mathbb{K}^2$ genannt, und ist ein zentrales Objekt der “algebraischen Geometrie”. Die Mengen in der Zariski-Topologie werden oft *Zariski-offen* genannt.

Für $\mathbb{K} = \mathbb{R}$ ist jede Zariski-offene Menge auch offen im üblichen Sinne, aber die Umkehrung gilt nicht. Für $\mathbb{K} = \mathbb{F}_p$ erhalten wir nun einen Begriff von Offenheit, welcher wenig mit Anschauung zu tun hat, aber trotzdem sehr hilfreich ist.

Definition. Es sei $(X, \mathcal{T})$ ein topologischer Raum und es sei $A \subset X$ eine Teilmenge. Wir sagen $U \subset X$ ist eine *Umgebung von A*, falls es eine offene Menge V gibt, so dass $A \subset V \subset U$. Wir sagen U ist eine *offene Umgebung von A*, wenn U zudem offen ist.

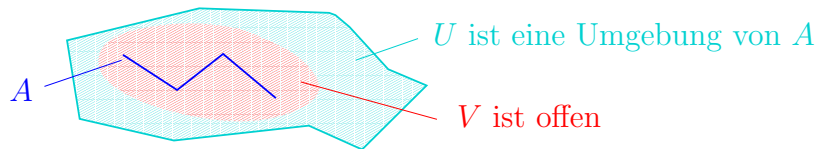


ABBILDUNG 1.

Beispiele.

- (1) Wir betrachten die Teilmenge $A = [-1, 1] \times \{0\}$ von dem topologischen Raum $\mathbb{R}^2$. Eine offene Umgebung von X ist dann beispielsweise gegeben durch die offene Scheibe $B_2(0)$.
- (2) Es sei X die Gerade mit einem Punkt im Unendlichen. Dann ist für jedes $C > 0$ die Menge $\{\infty\} \cup (-\infty, -C) \cup (C, \infty)$ eine offene Umgebung von ∞ .

⁶In dem Beweis benötigt man, dass der Ring $\mathbb{K}[x, y]$ noethersch ist, d.h. dass jedes Ideal in $\mathbb{K}[x, y]$ schon von endlich vielen Polynomen erzeugt wird. Diese Aussage wird in der (kommutativen) Algebra-Vorlesung bewiesen.

⁷Oskar Zariski (1899-1986) war ein amerikanischer Mathematiker.

Definition. Ein topologischer Raum heißt *Hausdorff*, wenn es zu zwei verschiedenen Punkten $x \neq y$ disjunkte offene Umgebungen U von x und V von y gibt.

Beispiel. Wir wollen nun überprüfen, ob die obigen Beispiele Hausdorff sind:

- (A) Wir hatten schon in Analysis II Satz 1.8 gesehen, dass der einem metrischen Raum (X, d) zugeordnete topologische Raum Hausdorff ist. In der Tat, denn für $x \neq y$ setzen wir $r = d(x, y)$, dann sind $B_{\frac{r}{2}}(x)$ und $B_{\frac{r}{2}}(y)$ disjunkte offene Umgebungen von x und y .
- (B) Wenn X mindestens zwei Elemente besitzt, dann ist $(X, \mathcal{T})$ nicht Hausdorff.
- (C) Diese Topologie ist Hausdorff, denn, für gegebene $x \neq y$ können wir $U = \{x\}$ und $V = \{y\}$ wählen.
- (D) Diese Topologie ist nicht Hausdorff, denn die Schnittmenge von zwei nichtleeren offenen Mengen in $(X, \mathcal{T})$ ist nichtleer, es gibt also keine disjunkten nichtleeren offenen Teilmengen in $(X, \mathcal{T})$.
- (E) Dies ist eine Übungsaufgabe in Übungsblatt 1.
- (F) Dies ist ebenfalls eine Übungsaufgabe in Übungsblatt 1.
- (G) Wenn $\mathbb{K} = \mathbb{R}$, dann ist der topologische Raum nicht Hausdorff. Wir werden diese Aussage nicht verwenden und skizzieren daher nur einen Beweis. Für jedes Polynom $f \neq 0$ ist $V(f)$ eine Nullmenge in $\mathbb{R}^2$. Es folgt nun leicht, dass jede nichtleere Zariski-offene Teilmenge $U \subset \mathbb{R}^2$ die Eigenschaft besitzt, dass $\mathbb{R}^2 \setminus U$ eine Nullmenge ist. Insbesondere schneiden sich je zwei nichtleere Zariski-offene Teilmengen von $\mathbb{R}^2$. Insbesondere kann $\mathbb{R}^2$ mit der Zariski-Topologie kein Hausdorff-Raum sein.

1.4. Die Teilraumtopologie. Es sei $(X, \mathcal{T})$ ein topologischer Raum und $Y \subset X$ eine Teilmenge. Wir betrachten dann

$$SS := \{Y \cap U \mid U \in \mathcal{T}\}.$$

Es ist völlig elementar sich davon zu überzeugen, dass SS in der Tat eine Topologie auf Y definiert. Wir nennen SS die *Teilraumtopologie* auf Y . Die Definition besagt also, dass eine Teilmenge U offen in Y ist, genau dann, wenn es eine Menge $V \subset X$ gibt, welche offen in X ist, so dass $U = Y \cap V$.

Konvention. Wir betrachten jede Teilmenge Y von $\mathbb{R}^n$ immer als topologischen Raum bezüglich der durch $\mathbb{R}^n$ induzierten Teilraumtopologie.

Beispiel. Es sei

$$Y := S^1 = \{z \in \mathbb{C} = \mathbb{R}^2 \mid |z| = 1\}.$$

Es seien $\alpha, \beta \in \mathbb{R}$. Dann ist das “offene Kreisstück”

$$U := \{e^{i\varphi} \in \mathbb{C} = \mathbb{R}^2 \mid \varphi \in (\alpha, \beta)\}$$

offen in $Y = S^1$, denn

$$U = \underbrace{\{re^{i\varphi} \in \mathbb{C} = \mathbb{R}^2 \mid r \in (\frac{1}{2}, \frac{3}{2}) \text{ und } \varphi \in (\alpha, \beta)\}}_{\text{offen in } \mathbb{C}=\mathbb{R}^2} \cap Y.$$

Dieses Beispiel ist in Abbildung 2 skizziert.

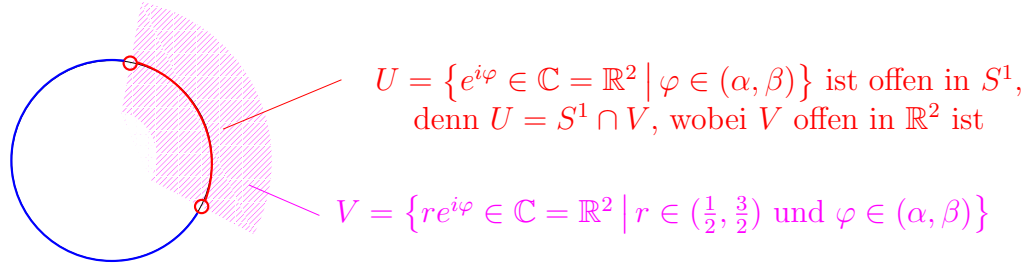


ABBILDUNG 2. Definition der Teilraumtopologie.

1.5. Abgeschlossene Teilmengen.

Definition. Es sei X ein topologischer Raum. Wir sagen $A \subset X$ ist *abgeschlossen*, wenn $X \setminus A$ offen ist.

Beispiel. Wir betrachten den topologischen Raum $X = (-1, 1)$. Dann ist $A = (-1, 0]$ eine abgeschlossene Teilmenge, denn $X \setminus A = (0, 1)$ ist eine offene Teilmenge von X .

Der Beweis von folgendem Lemma ist wort-wörtlich der gleiche wie der Beweis von Analysis II Lemma 1.7. Der Beweis folgt leicht aus den Axiomen von offenen Mengen und den de Morganschen Gesetzen.

Lemma 1.4. *Es sei X ein topologischer Raum. Dann gilt:*

- (1) $\emptyset$ und X sind abgeschlossen,
- (2) die Vereinigung von endlich vielen abgeschlossen Teilmengen ist abgeschlossen,
- (3) der Durchschnitt von beliebig vielen abgeschlossenen Teilmengen ist wiederum abgeschlossen.

Definition. Es sei X ein topologischer Raum und $A \subset X$ eine Teilmenge.

- (1) Der *Abschluss* $\overline{A}$ von A ist die Schnittmenge aller abgeschlossenen Mengen in X , die A enthalten.
- (2) Das *Innere* $\overset{\circ}{A}$ ist die Vereinigung aller offenen Mengen von X , die in A enthalten sind.

Es folgt sofort aus Lemma 1.4 (3), dass der Abschluss $\overline{A}$ eine abgeschlossene Menge in X ist, und es folgt aus Axiom (T3) einer Topologie, dass das Innere $\overset{\circ}{A}$ eine offene Menge in X ist.

In Übungsblatt 1 werden wir folgenden Satz beweisen.

Satz 1.5. *Es sei X ein topologischer Raum und $B \subset X$ eine Teilmenge. Dann gilt*

- (1) $x \in \overline{B} \iff$ für jede offene Umgebung U von x gilt $U \cap B \neq \emptyset$
- (2) $x \in \overset{\circ}{B} \iff$ es gibt eine offene Umgebung U von x mit $U \subset B$.

1.6. Stetige Abbildungen.

Definition. Es seien X und Y topologische Räume⁸ und es sei $f: X \rightarrow Y$ eine Abbildung. Wir definieren

$$f \text{ ist stetig} \quad :\Longleftrightarrow \quad \text{für jede offene Menge } U \text{ in } Y \text{ ist } f^{-1}(U) \text{ offen in } X.$$

Bemerkung.

- (1) Es sei $f: X \rightarrow Y$ eine Abbildung zwischen zwei metrischen Räumen. Dann folgt aus Satz 1.2, dass f stetig ist als Abbildung zwischen metrischen Räumen genau dann, wenn f stetig ist als Abbildung zwischen den entsprechenden topologischen Räumen.
- (2) Das gleiche Argument wie im Beweis von Satz 1.2 zeigt, dass wir in der Definition von Stetigkeit auch “offen” durch “abgeschlossen” ersetzen können.

Wir können nun folgenden Satz formulieren. Der Beweis hierbei ist wort-wörtlich der gleiche, wie der Beweis von Lemma 1.3.

Satz 1.6. *Es seien $f: X \rightarrow Y$ und $g: Y \rightarrow Z$ stetige Abbildungen zwischen topologischen Räumen, dann ist auch $g \circ f: X \rightarrow Z$ stetig.*

Wir beschließen das Kapitel mit folgendem Lemma.

Lemma 1.7. *Es sei X ein topologischer Raum und A eine Teilmenge von X . Dann gilt:*

- (1) *Die Inklusionsabbildung $i: A \rightarrow X$ ist stetig.*
- (2) *Die Einschränkung einer stetigen Abbildung $f: X \rightarrow Y$ auf A ist ebenfalls stetig.*

Beweis.

- (1) Wir betrachten die Inklusionsabbildung $i: A \rightarrow X$. Es sei $U \subset X$ offen. Dann ist $i^{-1}(U) = U \cap A$, aber $U \cap A$ ist per Definition offen in A . Also ist i stetig.
- (2) Es sei $f: X \rightarrow Y$ eine stetige Abbildung. Die Einschränkung auf A ist die Verknüpfung der stetigen Abbildungen $i: A \rightarrow X$ und $f: X \rightarrow Y$. Also ist nach Satz 1.6 die Einschränkung von f auf A ebenfalls stetig. □

1.7. Homöomorphismen. Folgende Definition hatten wir in Analysis II schon für metrische Räume eingeführt.

Definition. Es seien X und Y topologischen Räume und es sei $f: X \rightarrow Y$ eine Abbildung. Wir sagen f ist ein *Homöomorphismus*, wenn f folgende drei Eigenschaften besitzt:

- (1) f ist stetig,
- (2) f ist bijektiv,
- (3) $f^{-1}: Y \rightarrow X$ ist stetig.

⁸Wenn wir schreiben “es sei X ein topologischer Raum”, dann meinen wir das X eine Menge ist mit einer Topologie. Im Normalfall erwähnen wir die Topologie nicht in der Notation. Diese Konvention ist gang und gäbe in der Mathematik, beispielsweise schreiben wir in Linearer Algebra, “sei V ein Vektorraum”, wenn wir wirklich schreiben müssten, “sei $(V, +, \cdot)$ ein Vektorraum”.

Es seien X und Y topologische Räume. Wir sagen X und Y sind *homöomorph*, wenn es einen Homöomorphismus $f: X \rightarrow Y$ gibt.

Beispiele.

- (1) Die Abbildung

$$\begin{aligned}\mathbb{R} &\rightarrow (-1, 1) \\ x &\mapsto \frac{2}{\pi} \arctan(x)\end{aligned}$$

ist ein Homöomorphismus.

- (2) Die Abbildung

$$\begin{aligned}[0, 2\pi) &\rightarrow S^1 \\ x &\mapsto (\cos(x), \sin(x))\end{aligned}$$

ist eine stetige Bijektion, aber kein Homöomorphismus, nachdem die Umkehrabbildung im Punkt $(1, 0)$ nicht stetig ist. In der Tat kann es keinen Homöomorphismus zwischen $[0, 2\pi)$ und S^1 geben, denn $[0, 2\pi)$ ist nicht kompakt, aber S^1 ist kompakt. Es ist aber leicht zu sehen, dass wenn $f: X \rightarrow Y$ ein Homöomorphismus ist, dann ist X kompakt genau dann, wenn Y kompakt ist.

- (3) In Übungsblatt 1 werden wir zeigen, dass die Gerade mit einem Punkt im Unendlichen homöomorph zu S^1 ist.

1.8. Kompakte topologische Räume. In Analysis II Kapitel 3.1 hatten wir schon den Begriff von kompakten Teilmengen für metrische Räume eingeführt. Die Definition kann man problemlos auf topologische Räume verallgemeinern. Viele von den Sätzen, welche wir in Analysis II schon bewiesen hatten, gelten mit minimalen Abänderungen der Aussagen und Beweise auch für kompakte topologische Räume.

Definition. Es sei X ein topologischer Raum.

- (1) Eine *offene Überdeckung* von X ist eine Familie $\{U_i\}_{i \in I}$ von offenen Teilmengen von X , so dass

$$X = \bigcup_{i \in I} U_i.$$

- (2) Wir sagen X ist *kompakt*, wenn es zu jeder offenen Überdeckung $\{U_i\}_{i \in I}$ von X endlich viele Indizes $i_1, \dots, i_k \in I$ gibt, so dass

$$X = U_{i_1} \cup U_{i_2} \cup \dots \cup U_{i_k}.$$

- (3) Wir sagen eine Teilmenge $A \subset X$ ist kompakt, wenn A bezüglich der Teilraumtopologie kompakt ist.

Beispiele. Wir betrachten wiederum die obigen Beispiele.

- (A) Wenn (X, d) ein metrischer Raum ist, dann erhalten wir den gleichen Kompaktheitsbegriff wie in Analysis II.
 (B) Es sei X eine beliebige Menge mit der trivialen Topologie. Dann ist X offensichtlich kompakt.

- (C) Es sei X eine beliebige Menge mit der diskreten Topologie. Wenn X endlich ist, dann ist X offensichtlich kompakt. Wenn X unendlich ist, dann ist X nicht kompakt, denn $\{x\}_{x \in X}$ ist nun eine offene Überdeckung von X , aber X wird nicht von endlich vielen dieser offenen Menge überdeckt.
- (D) Dieser topologische Raum ist kompakt.
- (E) und (F) Wir werden in Übungsblatt 1 der Frage nachgehen, ob diese topologischen Räume kompakt sind.

Der folgende Satz von Heine–Borel gibt viele Beispiele von kompakten Mengen.

Satz 1.8. (Satz von Heine–Borel) *Es sei A eine Teilmenge von $\mathbb{R}^n$. Dann gilt:*

$$A \text{ ist kompakt} \iff A \text{ ist beschränkt und abgeschlossen.}$$

Der Beweis der folgenden Sätze ist fast identisch zu den Beweisen der analogen Aussagen aus der Analysis II.

Satz 1.9. *Es sei X ein kompakter topologischer Raum und $A \subset X$ eine abgeschlossene Teilmenge. Dann ist auch A kompakt.*

Satz 1.10. *Es sei $f: X \rightarrow Y$ eine stetige Abbildung zwischen topologischen Räumen. Wenn X kompakt ist, dann ist auch $f(X)$ kompakt.*

Satz 1.11. *Es sei $f: X \rightarrow \mathbb{R}$ eine stetige Abbildung auf einem kompakten topologischen Raum. Dann ist f beschränkt, und f nimmt ein Minimum und ein Maximum an, d.h. es existieren $x_0, x_1 \in X$, so dass*

$$f(x_0) \leq f(x) \leq f(x_1)$$

für alle $x \in X$.

Der Satz von Heine-Borel besagt insbesondere, dass eine kompakte Teilmenge von $\mathbb{R}^n$ abgeschlossen in $\mathbb{R}^n$ ist. Diese Aussage gilt im allgemeinen nicht für topologische Räume, d.h. eine kompakte Teilmenge eines topologischen Raumes X ist nicht notwendigerweise abgeschlossen in X .

Es sei beispielsweise X eine beliebige Menge X mit mindestens zwei Elementen und $\mathcal{T}$ die triviale Topologie auf X , d.h. $\mathcal{T} = \{\emptyset, X\}$. Es sei nun $A \subset X$ eine beliebige Teilmenge mit $A \neq \emptyset$ und $A \neq X$. Dann ist die Teilraumtopologie von A gegeben durch $\{\emptyset, A\}$, d.h. die Teilraumtopologie auf A ist die triviale Topologie. Insbesondere ist A kompakt. Andererseits ist A keine abgeschlossene Teilmenge von X .

Der folgende Satz besagt, dass im Gegenzug kompakte Teilmengen von *Hausdorff-Räumen* abgeschlossen sind.

Satz 1.12. *Es sei X ein topologischer Raum und $A \subset X$ eine kompakte Teilmenge. Wenn X ein Hausdorff-Raum ist, dann ist A abgeschlossen in X .*

Der Beweis des Satzes ähnelt dem Beweis von Satz 3.5 aus der Analysis II.

Beweis. Es sei X ein topologischer Raum, welcher Hausdorff ist. Es sei $A \subset X$ eine kompakte Teilmenge. Wir wollen zeigen, dass $X \setminus A$ offen ist. Es genügt folgende Behauptung zu beweisen:

Behauptung. Es sei $x \in X \setminus A$. Dann existiert eine offene Umgebung V von x , welche in $X \setminus A$ enthalten ist.

Wir wenden die Hausdorff-Eigenschaft auf x und jedes $y \in A$ an. Für jedes $y \in A$ erhalten wir offene disjunkte Umgebungen U_y von y und V_y von x . Offensichtlich ist

$$A = \bigcup_{y \in A} \{y\} \subset \bigcup_{y \in A} (U_y \cap A) \subset A.$$

Es folgt, also dass $\{U_y \cap A\}_{y \in A}$ eine offene Überdeckung von A ist. Nachdem A kompakt ist, gibt es $y_1, \dots, y_k$, so dass

$$A = \bigcup_{i=1}^k (U_{y_i} \cap A).$$

Wir betrachten nun

$$V := \bigcap_{i=1}^k V_{y_i}.$$

Als Durchschnitt von endlich vielen offenen Mengen ist V eine offene Umgebung von x . Zudem ist V von allen $U_{y_1}, \dots, U_{y_k}$ disjunkt, also auch von $A \subset U_{y_1} \cup \dots \cup U_{y_k}$. Wir haben damit die Behauptung bewiesen. $\square$

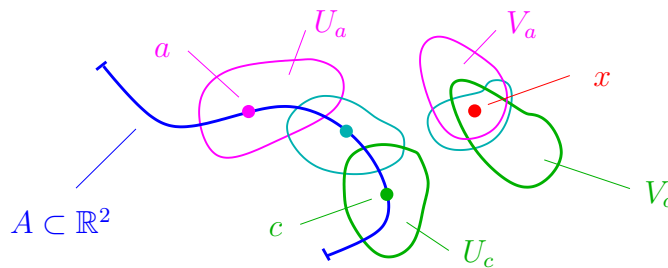


ABBILDUNG 3. Skizze zum Beweis von Satz 1.12.

Wir hatten schon auf Seite 13 gesehen, dass eine bijektive stetige Abbildung $f: X \rightarrow Y$ zwischen topologischen Räumen nicht notwendigerweise ein Homöomorphismus ist. Mithilfe des vorherigen Satzes können wir jedoch jetzt folgenden Satz beweisen, welcher uns später das Leben mehrmals erleichtern wird.

Satz 1.13. *Es sei $f: X \rightarrow Y$ eine bijektive stetige Abbildung zwischen topologischen Räumen. Wenn X kompakt ist, und wenn Y Hausdorff ist, dann ist f ein Homöomorphismus.*

Bemerkung. Man kann sich leicht davon überzeugen, dass die zusätzliche Voraussetzung, dass Y Hausdorff ist, notwendig ist.

Der Beweis von Satz 1.13 ist fast wort-wörtlich der gleiche wie von Satz 3.9 aus der Analysis II.

Beweis. Wir müssen also zeigen, dass die Umkehrabbildung $g := f^{-1}: Y \rightarrow X$ stetig ist. Wie wir auf Seite 12 angemerkt hatten genügt es zu zeigen, dass für alle abgeschlossenen Teilmengen $A \subset X$ auch das Urbild $g^{-1}(A)$ abgeschlossen ist.

In der Tat gilt

$$\begin{array}{ccccc}
 A \subset X \text{ abgeschlossen} & \Rightarrow & A \text{ ist kompakt} & \Rightarrow & f(A) \text{ kompakt} \\
 & \uparrow & & \uparrow & \\
 & \text{folgt aus Satz 1.9, da } X \text{ kompakt} & & \text{Satz 1.10} & \\
 & & \Rightarrow & f(A) = g^{-1}(A) \text{ ist abgeschlossen.} & \\
 & \uparrow & & & \\
 & \text{folgt aus Satz 1.12, da } Y \text{ Hausdorff} & & &
 \end{array}$$

□

Nach dieser kurzen Einführung in die topologischen Räume werden wir uns die nächsten Wochen nur auf Teilmengen von $\mathbb{R}^n$ mit der Teilraumtopologie konzentrieren. Die allgemeineren topologischen Räume werden erst gegen Ende der Vorlesung wieder eine wichtige Rolle spielen.

2. DEFINITION UND BEISPIELE VON UNTERMANNIGFALTIGKEITEN

2.1. Die Definition von Untermannigfaltigkeiten. In Analysis II hatten wir k -dimensionale Untermannigfaltigkeiten über Parametrisierungen definiert, und wir hatten dann in Satz 11.1 bewiesen, dass man k -dimensionale Untermannigfaltigkeiten auch über “Karten” definieren kann. In dieser Vorlesung werden wir durchgehend den zweiten Gesichtspunkt verwenden, insbesondere werden wir im Folgenden k -dimensionale Untermannigfaltigkeiten mithilfe von Karten definieren.

Wir erinnern zuerst an folgende Definition aus der Analysis II.

Definition.

- (1) Es sei $U \subset \mathbb{R}^n$ offen und es sei

$$f = (f_1, \dots, f_m): U \rightarrow \mathbb{R}^m$$

eine Abbildung. Wir nennen f eine C^∞ -Abbildung oder auch *glatt*, wenn die Koordinatenfunktionen $f_1, \dots, f_m$ beliebig oft stetig partiell differenzierbar sind.

- (2) Eine Abbildung $f: U \rightarrow V$ zwischen zwei offenen Teilmengen von $\mathbb{R}^n$ heißt *Diffeomorphismus*, wenn gilt:

- (a) f ist glatt,
- (b) f ist bijektiv, und
- (c) f^{-1} ist ebenfalls glatt.

Notation. Für $k \leq n$ bezeichnen wir im Folgenden

$$E_k := \{(x_1, \dots, x_k, 0, \dots, 0) \in \mathbb{R}^n \mid x_1, \dots, x_k \in \mathbb{R}\}$$

als die k -dimensionale Ebene in $\mathbb{R}^n$.

Definition. Eine Teilmenge $M \subset \mathbb{R}^n$ heißt k -dimensionale Untermannigfaltigkeit, wenn es zu jedem Punkt $P \in M$ eine Karte um P gibt, d.h. einen Diffeomorphismus $\Phi: U \rightarrow V$ zwischen einer offenen Umgebung von P und einer offenen Teilmenge $V \subset \mathbb{R}^n$, so dass⁹

$$\Phi(M \cap U) = E_k \cap V.$$

Ein *Atlas* für M ist eine Familie $\{\Phi_i: U_i \rightarrow V_i\}_{i \in I}$ von Karten mit $M \subset \bigcup_{i \in I} U_i$.

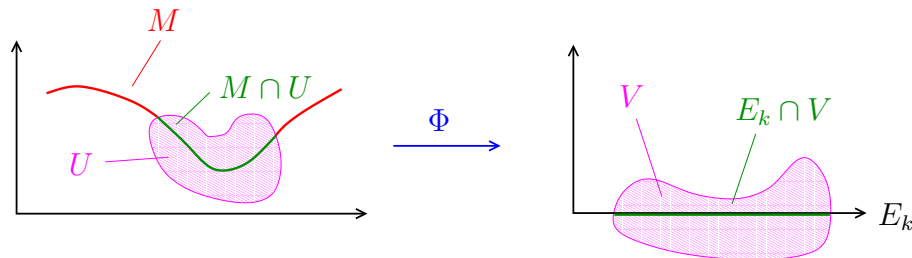


ABBILDUNG 4. Schematisches Bild einer Karte.

⁹Die leere Menge ist nun per Definition für jedes k eine k -dimensionale Untermannigfaltigkeit.

Bemerkung. Nach Satz 11.1 aus Analysis II ist diese Definition (fast) äquivalent zur Definition von Untermannigfaltigkeiten, welche wir in Analysis II gegeben hatten. Der einzige Unterschied ist, dass wir von nun an verlangen, dass die Karten Diffeomorphismen sind, und nicht nur C^1 -invertierbare Abbildungen. In der Praxis macht das allerdings normalerweise keinen großen Unterschied.

Wir wollen nun zeigen, dass der Kreis S^1 und die Sphäre S^2 Untermannigfaltigkeiten sind. In diesen und in vielen anderen Fällen ist es allerdings leichter, die Inversen von Karten explizit anzugeben. Wir verwenden den Satz über die Umkehrabbildung, welchen wir schon als Satz 9.5 und Korollar 9.8 in Analysis II bewiesen hatten, um zu zeigen, dass eine Abbildung in der Tat ein Diffeomorphismus ist.

Satz 2.1. (Satz über die Umkehrabbildung) *Es sei $V \subset \mathbb{R}^n$ offen, $\Psi: V \rightarrow \mathbb{R}^n$ eine glatte Abbildung und $Q \in V$. Wenn $D\Psi(Q)$ invertierbar ist, dann existieren offene Umgebungen $Y \subset V$ von Q und X von $\Psi(Q)$, so dass $\Psi: Y \rightarrow X$ ein Diffeomorphismus ist.*

Wir erhalten folgendes hilfreiche Korollar, welches es erlaubt zu zeigen, dass eine Abbildung ein Diffeomorphismus ist, ohne die Umkehrabbildung zu bestimmen.

Korollar 2.2. *Es sei $V \subset \mathbb{R}^n$ offen und es sei $\Psi: V \rightarrow \mathbb{R}^n$ eine glatte Abbildung. Wenn Ψ bijektiv ist und wenn für alle $Q \in V$ das Differential $D\Psi(Q)$ invertierbar ist, dann ist die Abbildung $\Psi: V \rightarrow \Psi(V)$ ein Diffeomorphismus.*

Beweis. Wir setzen $U := \Psi(V)$. Wir müssen nur noch zeigen, dass $\Psi^{-1}: U \rightarrow V$ ebenfalls eine glatte Abbildung ist. Es sei also $P \in U$. Wir müssen zeigen, dass es eine offene Umgebung Y von P gibt, so dass Ψ^{-1} auf Y glatt ist. Wir setzen $Q := \Psi^{-1}(P)$. Nach dem $D\Psi(Q)$ invertierbar ist können wir den Satz über die Umkehrabbildung anwenden. Wir erhalten insbesondere eine offene Umgebung X von P , so dass Ψ^{-1} auf X glatt ist. $\square$

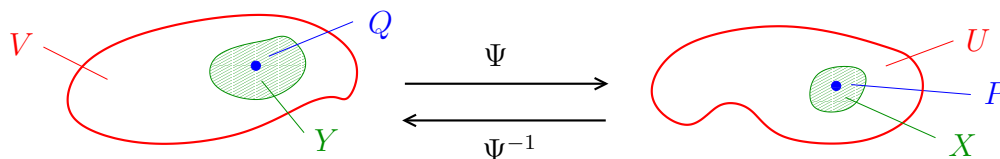


ABBILDUNG 5. Skizze zum Beweis von Korollar 2.2.

Wir wenden uns nun den Beispielen zu.

Beispiel. Wir wollen zeigen, dass der Kreis

$$S^1 = \{(x, y) \in \mathbb{R}^2 \mid x^2 + y^2 = 1\}$$

eine eindimensionale Untermannigfaltigkeit von $\mathbb{R}^2$ ist. Wir betrachten die Abbildung

$$\begin{aligned}\Psi: V := (-\pi, \pi) \times (-1, 1) &\rightarrow U := \Psi(V) \subset \mathbb{R}^2 \\ (\varphi, s) &\mapsto ((1+s) \cos \varphi, (1+s) \sin \varphi).\end{aligned}$$

Dies ist eine glatte bijektive Abbildung mit $\Psi(V \cap E_1) = S^1 \cap U$. Es folgt aus Korollar 2.2, dass dies sogar ein Diffeomorphismus ist. Also definiert

$$\Phi := \Psi^{-1}: U \rightarrow V$$

eine Karte. Diese Karte deckt alle Punkte in $S^1 \setminus \{(-1, 0)\}$ ab. Wir benötigen nun noch eine Karte um $(-1, 0)$. Diese ist beispielsweise gegeben durch die Umkehrabbildung Φ' von

$$\begin{aligned}\Psi': V' := (0, 2\pi) \times (-1, 1) &\rightarrow U' := \Psi'(V') \subset \mathbb{R}^2 \\ (\varphi, s) &\mapsto ((1+s) \cos \varphi, (1+s) \sin \varphi).\end{aligned}$$

Die beiden Karten Φ und Φ' bilden also einen Atlas von S^1 . Die beiden Karten werden in Abbildung 6 illustriert.

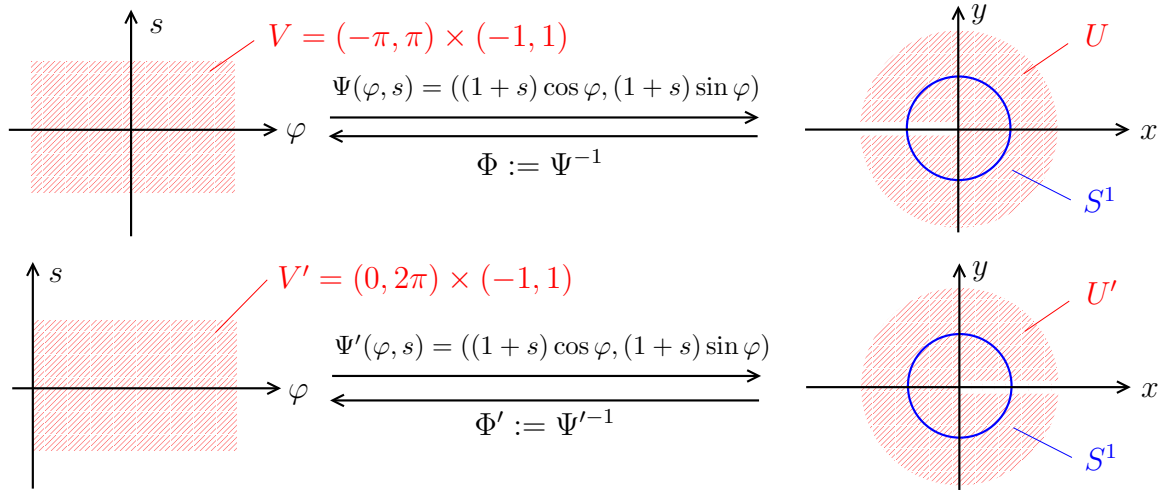


ABBILDUNG 6. Atlas für die eindimensionale Untermannigfaltigkeit S^1 .

Beispiel. Für jedes $r > 0$ und $C > 0$ betrachten wir nun

$$Z := \{(x, y, z) \in \mathbb{R}^3 \mid x^2 + y^2 = r^2 \text{ und } z \in (-C, C)\},$$

anschaulich gesprochen ist Z ein Zylinder mit Radius r . Ganz analog zum vorherigen Beispiel kann man zeigen, dass Z eine 2-dimensionale Untermannigfaltigkeit ist.

Beispiel. Wir wollen nun zeigen, dass die Kugel

$$S^2 = \{(x, y, z) \in \mathbb{R}^3 \mid x^2 + y^2 + z^2 = 1\}$$

eine 2-dimensionale Untermannigfaltigkeit von $\mathbb{R}^3$ ist. Der Beweis dieser Aussage läuft ganz analog zum vorherigen Beispiel. Wir betrachten zuerst die Abbildung

$$\begin{aligned} \Psi: V := (-\pi, \pi) \times (0, \pi) \times (-1, 1) &\rightarrow U := \Psi(V) \subset \mathbb{R}^3 \\ (\varphi, \theta, s) &\mapsto \begin{pmatrix} (1+s) \cos \varphi \cdot \sin \theta \\ (1+s) \sin \varphi \cdot \sin \theta \\ (1+s) \cos \theta \end{pmatrix}. \end{aligned}$$

Mithilfe von Korollar 2.2 zeigt man leicht, dass dies ein Diffeomorphismus ist. Zudem gilt $\Psi(V \cap E_2) = S^2 \cap U$. Also definiert

$$\Phi := \Psi^{-1}: U \rightarrow V$$

eine Karte für alle Punkte in $U \cap S^2$. Mit anderen Worten Φ deckt alle Punkte in S^2 außerhalb des Halbkreises $K := \{(x, y, z) \in S^2 \mid x \leq 0 \text{ und } y = 0\}$ ab. Wir brauchen nun noch eine Karte, welche die fehlenden Punkte abdeckt. Wir verwenden dazu die vorherige Abbildung Ψ , vertauschen die Rollen der y - und z -Koordinaten und wir spiegeln entlang der yz -Ebene. Genauer gesagt, wir betrachten die Abbildung

$$\begin{aligned} \Psi': V' := (-\pi, \pi) \times (0, \pi) \times (-1, 1) &\rightarrow U' := \Psi'(V') \subset \mathbb{R}^3 \\ (\varphi, \theta, s) &\mapsto \begin{pmatrix} -(1+s) \cos \varphi \cdot \sin \theta \\ (1+s) \cos \theta \\ (1+s) \sin \varphi \cdot \sin \theta \end{pmatrix}. \end{aligned}$$

Mithilfe von Korollar 2.2 zeigt man wiederum, dass ein Diffeomorphismus ist. Das Inverse

$$\Phi' := \Psi'^{-1}: U' \rightarrow V'$$

ist eine Karte, welche alle Punkte außerhalb von $L = \{(x, y, z) \in S^2 \mid x \geq 0 \text{ und } z = 0\}$ abdeckt. Nachdem es auf S^2 keinen Punkt gibt, welcher in den beiden Halbkreisen K und L liegt, decken die beiden Karten S^2 ab. Wir haben also insbesondere gezeigt, dass S^2 einen Atlas besitzt, welcher aus zwei Karten besteht.

Beispiel. Wir betrachten noch einmal die 2-dimensionale Sphäre

$$S^2 = \{(x, y, z) \mid x^2 + y^2 + z^2 = 1\}$$

mit den Punkten

$$N = (0, 0, 1) \quad \text{“Nordpol”} \quad \text{und} \quad S = (0, 0, -1) \quad \text{“Südpol”}.$$

Wir geben nun einen anderen Atlas für S^2 an, bei dem wir die Karten explizit hinschreiben können. Genauer gesagt, wir betrachten die Abbildung

$$\begin{aligned} \Phi: \{(x, y, z) \mid z < 1\} &\rightarrow \mathbb{R}^3 \\ (x, y, z) &\mapsto \left(\frac{x}{1-z}, \frac{y}{1-z}, 1 - x^2 - y^2 - z^2 \right). \end{aligned}$$

Eingeschränkt auf $S^2 \setminus N$ ist dies die stereographische Projektion bezüglich des Nordpols N . Wir betrachten auch die Abbildung

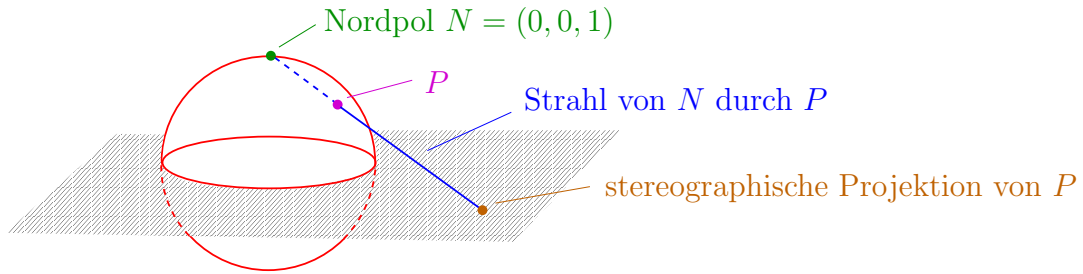


ABBILDUNG 7.

$$\begin{aligned} \Psi: \{(x, y, z) \mid z > -1\} &\rightarrow \mathbb{R}^3 \\ (x, y, z) &\mapsto \left(\frac{x}{1+z}, \frac{y}{1+z}, 1 - x^2 - y^2 - z^2 \right). \end{aligned}$$

Eingeschränkt auf $S^2 \setminus S$ ist dies die stereographische Projektion bezüglich des Südpols S . Man kann nun leicht überprüfen, dass Einschränkungen von Φ und Ψ auf geeignet gewählte offene Mengen Karten für S^2 sind, welche S^2 überdecken. Nachdem die Definitionsbereiche ganz S^2 abdecken, bilden diese Einschränkungen von Φ und Ψ einen Atlas für S^2 .

Beispiel. Wir bezeichnen nun mit T die Menge aller Punkte in $\mathbb{R}^3$, welche geschrieben werden können als Produkt

$$\underbrace{\begin{pmatrix} \cos \varphi & -\sin \varphi & 0 \\ \sin \varphi & \cos \varphi & 0 \\ 0 & 0 & 1 \end{pmatrix}}_{\text{Drehung um die } z\text{-Achse}} \cdot \underbrace{\begin{pmatrix} 2 + \sin \theta \\ 0 \\ \cos \theta \end{pmatrix}}_{\text{beschreibt Kreis in } xz\text{-Ebene}},$$

wobei $\varphi, \theta \in \mathbb{R}$. Mit anderen Worten, T ist die Menge

$$T := \{((2 + \sin \theta) \cos \varphi, (2 + \sin \theta) \sin \varphi, \cos \theta) \mid \theta, \varphi \in \mathbb{R}\}.$$

Diese Menge wird in Abbildung 8 skizziert. Mit ähnlichen Argumenten wie in den vorherigen Beispielen kann man zeigen, dass der Torus T eine 2-dimensionale Untermannigfaltigkeit ist.

2.2. Untermannigfaltigkeiten und Parametrisierungen. Wir erinnern noch an folgende Definition aus der Analysis II:

Definition. Eine k -dimensionale Parametrisierung einer Teilmenge M von $\mathbb{R}^n$ um den Punkt P ist eine Homöomorphismus $\varphi: T \rightarrow U$, wobei gilt:

- (1) T ist eine offene Teilmenge von $\mathbb{R}^k$,
- (2) U ist eine offene Umgebung von P in M ,

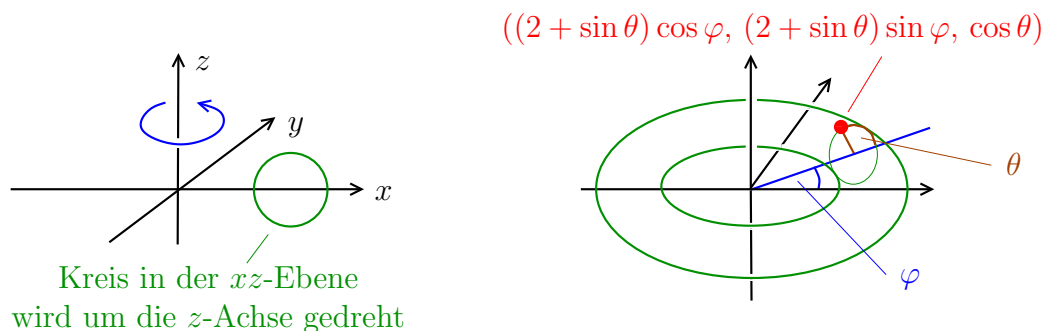


ABBILDUNG 8.

- (3) φ ist zudem glatt und für jeden Punkt P von T ist der Rang¹⁰ des Differentials $D\varphi(P)$ gleich k .

In Satz 11.1 aus der Analysis II hatten wir gesehen, dass wir Untermannigfaltigkeiten entweder über Karten oder über Parametrisierungen definieren können. Nachdem wir dieses Mal die Untermannigfaltigkeiten über Karten eingeführt hatten, lautet nun Satz 11.1 aus der Analysis II wie folgt.¹¹

Satz 2.3. Eine Teilmenge $M \subset \mathbb{R}^n$ ist eine k -dimensionale Untermannigfaltigkeit genau dann, wenn es um jedem Punkt $P \in M$ eine k -dimensionale Parametrisierung gibt. Genauer gesagt gelten folgende beide Aussagen:

- (1) Wenn $\Phi: U \rightarrow V$ eine Karte ist, dann ist $\Phi^{-1}|_{U \cap E_k} \rightarrow M$ eine Parametrisierung um den Punkt P ,
- (2) wenn $\varphi: T \rightarrow M$ eine Parametrisierung um den Punkt P ist, dann gibt es eine Karte $\Phi: U \rightarrow V$ um P , so dass $U \cap M \subset T$, und so dass $\varphi = \Phi^{-1}|_{U \cap M}$.

2.3. Untermannigfaltigkeiten und Graphen. In Übungsblatt 1 werden wir folgendes Lemma beweisen.

Lemma 2.4. Es sei $U \subset \mathbb{R}^n$ eine offene Teilmenge und es sei $f: U \rightarrow \mathbb{R}^k$ eine glatte Abbildung. Dann ist der Graph

$$\text{Graph}(f) = \{(x, f(x)) \mid x \in U\} \subset \mathbb{R}^{n+k}$$

eine n -dimensionale Untermannigfaltigkeit, welche einen Atlas mit nur einer Karte besitzt.

¹⁰Zur Erinnerung, der Rang einer $k \times n$ -Matrix A ist definiert als die Dimension des von den Spalten von A aufgespannten Teilraums von $\mathbb{R}^k$.

¹¹Wenn man's genau nimmt, wurde der Satz nur für stetig differenzierbare Karten und Parametrisierungen formuliert und bewiesen. Der Beweis überträgt sich aber wort-wörtlich auf den Fall von glatten Karten und glatten Parametrisierungen.

Beispiel. Die offene Halbsphäre

$$K := \{(x, y, 2 + \sqrt{1 - x^2 - y^2}) \mid x^2 + y^2 < 1\}$$

ist der Graph der Funktion $f(x, y) = 2 + \sqrt{1 - x^2 - y^2}$ auf der offenen Scheibe $D_1(0)$, also eine zweidimensionale Untermannigfaltigkeit. Zudem ist die unendliche Helix

$$H := \{(x, \cos(x), \sin(x)) \mid x \in \mathbb{R}\}$$

der Graph der Funktion $f(x) = (\cos(x), \sin(x))$ auf $\mathbb{R}$, also eine eindimensionale Untermannigfaltigkeit von $\mathbb{R}^3$.

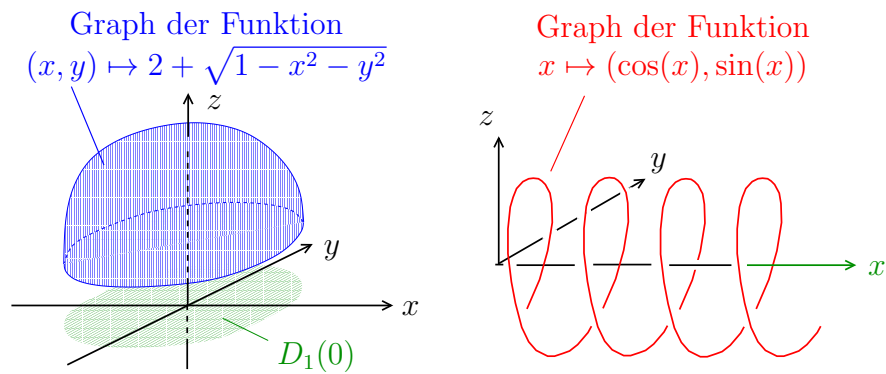


ABBILDUNG 9.

Der folgende Satz besagt nun, dass jede Untermannigfaltigkeit “lokal”, bis auf eine Rotation, ein Graph ist.

Satz 2.5. *Es sei $M \subset \mathbb{R}^{n+k}$ eine n -dimensionale Untermannigfaltigkeit und es sei $P \in M$. Dann existiert eine Permutationsmatrix A , eine offene Umgebung Y von P in M und eine offene Teilmenge $Z \subset \mathbb{R}^n$ sowie eine glatte Funktion $f: Z \rightarrow \mathbb{R}^k$, so dass*

$$A \cdot Y = \{(z, f(z)) \mid z \in Z\}.$$

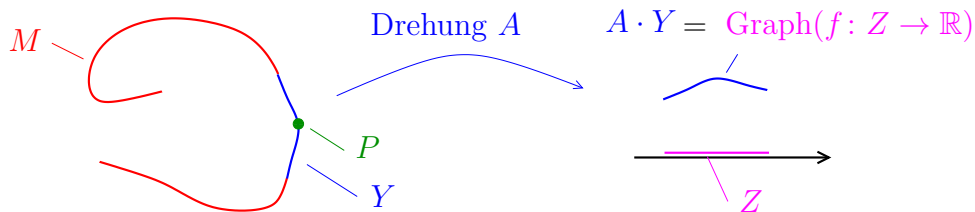


ABBILDUNG 10. Untermannigfaltigkeiten sind lokal ein Graph.

Beweis.

Es genügt zu zeigen, dass M , gegebenenfalls nach Anwendung einer Permutationsmatrix, um den Punkt P eine Parametrisierung $\varphi: Z \rightarrow \mathbb{R}^{n+k}$ besitzt, so dass φ von der Form $\varphi(z) = \begin{pmatrix} z \\ \psi(z) \end{pmatrix}$ ist.

Es sei $M \subset \mathbb{R}^{n+k}$ eine n -dimensionale Untermannigfaltigkeit und es sei $P \in M$. Nach Satz 2.3 existiert eine Parametrisierung $\varphi: T \rightarrow \mathbb{R}^{n+k}$ von M um den Punkt P . Es sei $Q \in T$ der eindeutig bestimmte Punkt mit $\varphi(Q) = P$. Wir schreiben nun $\varphi = \begin{pmatrix} \alpha \\ \beta \end{pmatrix}$, wobei $\alpha: T \rightarrow \mathbb{R}^n$ und $\beta: T \rightarrow \mathbb{R}^k$.

Nachdem φ eine Parametrisierung ist besitzt die $(n+k) \times n$ -Matrix $D\varphi(Q)$ insbesondere n linear unabhängige Zeilen. Wir nehmen zuerst an, dass die ersten n Zeilen von $D\varphi(Q)$ linear unabhängig sind. Das Differential $D\alpha(Q)$ ist gerade durch die ersten n Zeilen von $D\varphi(Q)$ gegeben. Also ist $D\alpha(Q)$ invertierbar. Nach dem Satz 2.1¹² über die Umkehrabbildung existiert nun eine offene Umgebung $T' \subset T$ von Q und eine offene Umgebung S' von $\alpha(Q)$, so dass $\alpha: T' \rightarrow S'$ ein Diffeomorphismus ist. Die Abbildung

$$\begin{aligned} S' &\rightarrow \mathbb{R}^{n+k} \\ z &\mapsto \varphi(\alpha^{-1}(z)) = \begin{pmatrix} z \\ \beta(\alpha^{-1}(z)) \end{pmatrix} \end{aligned}$$

ist dann eine Parametrisierung. Es ist nun offensichtlich, dass $M \cap \varphi(S')$ der Graph der Funktion $z \mapsto \beta(\alpha^{-1}(z))$ ist.

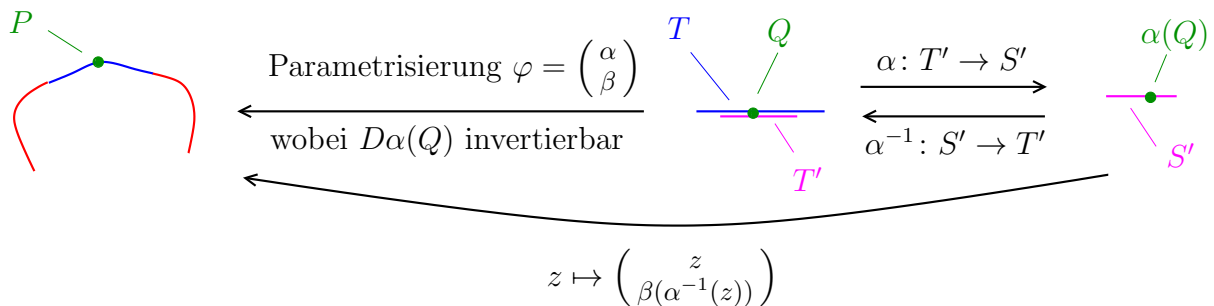


ABBILDUNG 11. Skizze zum Beweis von Satz 2.5.

Wir betrachten nun noch den Fall, dass die ersten n Zeilen von $D\varphi(Q)$ nicht linear unabhängig sind. In diesem Fall vertauschen wir die Koordinaten von M (d.h. wir wenden eine Permutationsmatrix auf M an), so dass die ersten n Zeilen von $D\varphi(Q)$ linear unabhängig werden. Wir fahren dann mit dem obigen Argument fort. $\square$

¹²In unserer Notation kann dieser wie folgt formuliert werden: Es sei $T \subset \mathbb{R}^n$ offen, $\alpha: T \rightarrow \mathbb{R}^n$ eine glatte Abbildung und $Q \in T$. Wenn $D\alpha(Q)$ invertierbar ist, dann existieren offene Umgebungen $T' \subset T$ von Q und S' von $\alpha(Q)$, so dass $\alpha: T' \rightarrow S'$ ein Diffeomorphismus ist.

2.4. Der Satz vom regulären Wert. Wir erinnern an folgende Definition aus Analysis II.

Definition. Es sei $U \subset \mathbb{R}^n$ offen und $f: U \rightarrow \mathbb{R}^k$ eine stetig differenzierbare Abbildung.

- (1) Wir sagen $x \in U$ ist ein *regulärer Punkt von f* , wenn das Differential $Df(x)$ den Rang k besitzt.
- (2) Wir sagen $z \in \mathbb{R}^k$ ist ein *regulärer Wert von f* , wenn alle Punkte in $f^{-1}(z)$ regulär sind.

Ein Punkt (bzw. Wert), welcher nicht regulär ist, wird *singulär* genannt.¹³

Bemerkung. Der Fall $k = 1$ ist ein wichtiger Spezialfall. Eine $1 \times n$ -Matrix besitzt Rang 1 genau dann, wenn die Matrix ungleich Null ist. Für $k = 1$ entspricht das Differential gerade dem Gradienten. Wir sehen also, dass ein Punkt $x \in U$ ein regulärer Punkt einer Abbildung $f: U \rightarrow \mathbb{R}$ ist, genau dann, wenn $\text{Grad } f(x) \neq 0$.

Wir können jetzt folgenden wichtigen Satz formulieren:

Satz 2.6. (Satz vom regulären Wert) *Es sei $X \subset \mathbb{R}^n$ offen und $f: X \rightarrow \mathbb{R}^k$ eine glatte Abbildung. Wenn $z \in \mathbb{R}^k$ ein regulärer Wert ist, dann ist $f^{-1}(z) \subset \mathbb{R}^n$ eine $(n - k)$ -dimensionale Untermannigfaltigkeit.*¹⁴

Beweis. Wir hatten den Satz vom regulären Wert in der Analysis II unter der etwas schwächeren Voraussetzung, dass f stetig differenzierbar ist, und der etwas schwächeren Aussage, dass man C^1 -Karten finden kann, in Analysis II bewiesen. Der gleiche Beweis gibt nun wort-wörtlich auch die jetzige Formulierung vom Satz vom regulären Wert. $\square$

Beispiel. Es seien $a_1, \dots, a_n$ reelle Zahlen, welche alle von null verschieden sind. Wir betrachten

$$\begin{aligned} f: \mathbb{R}^n &\rightarrow \mathbb{R} \\ (x_1, \dots, x_n) &\mapsto a_1 x_1^2 + \dots + a_n x_n^2, \end{aligned}$$

dann gilt für jeden Punkt $x = (x_1, \dots, x_n)$ in $\mathbb{R}^n$, dass

$$\text{Grad } f(x_1, \dots, x_n) = (2a_1 x_1, \dots, 2a_n x_n),$$

insbesondere ist der Gradient ungleich null für jeden Punkt $x \neq 0$. Anders ausgedrückt, die Menge der regulären Punkte von f ist $\mathbb{R}^n \setminus \{0\}$. Nachdem $f(x) = 0$ genau dann, wenn $x = 0$ folgt, dass alle Werte in $\mathbb{R} \setminus \{0\}$ regulär sind. Es folgt insbesondere aus dem Satz 2.6 vom regulären Wert, dass

$$M := \{(x_1, \dots, x_n) \in \mathbb{R}^n \mid a_1 x_1^2 + \dots + a_n x_n^2 = 1\}$$

eine $(n - 1)$ -dimensionale Untermannigfaltigkeit von $\mathbb{R}^n$ ist.

¹³Anders ausgedrückt, ein Wert $w \in \mathbb{R}^k$ ist regulär, wenn dieser nicht das Bild von einem singulären Punkt in U ist.

¹⁴Wenn z nicht im Bild von f liegt, dann ist $f^{-1}(z) = \emptyset$ die leere Menge, insbesondere eine (nicht besonders interessante) Untermannigfaltigkeit beliebiger Dimension.

2.5. Differenzierbare Funktionen auf Untermannigfaltigkeiten. Es sei $M \subset \mathbb{R}^n$ eine k -dimensionale Untermannigfaltigkeit. Wie immer betrachten wir M als topologischen Raum bezüglich der Teilraumtopologie, d.h. $U \subset M$ ist offen genau dann, wenn es eine offene Menge $V \subset \mathbb{R}^n$ gibt, mit $U = M \cap V$. Es sei nun $f: M \rightarrow \mathbb{R}$ eine Funktion. Nachdem wir M als topologischen Raum auffassen, haben wir schon einen Stetigkeitsbegriff für Funktionen auf M . In der Tat, eine Funktion $f: M \rightarrow \mathbb{R}$ heißt stetig, wenn für jede offene Menge $X \subset \mathbb{R}$, das Urbild $f^{-1}(X)$ eine offene Teilmenge von M ist.

Andererseits ist es a priori nicht klar, was es heißen soll, dass eine Funktion $f: M \rightarrow \mathbb{R}$ auf einer Untermannigfaltigkeit differenzierbar ist, denn wir können keine partiellen Ableitungen bilden, da für $P \in M$ die Punkte $P + he_i$ im Allgemeinen nicht mehr in M , d.h. im Definitionsbereich von f liegen.

Definition. Es sei M eine k -dimensionale Untermannigfaltigkeit und es sei $f: M \rightarrow \mathbb{R}$ eine Funktion. Wir sagen f ist *stetig differenzierbar* (beziehungsweise *glatt*), wenn es einen Atlas $\{\Phi_i: U_i \rightarrow V_i\}_{i \in I}$ gibt, so dass für alle $i \in I$ die Funktion¹⁵

$$V_i \cap E_k \xrightarrow{\Phi_i^{-1}} U_i \cap M \xrightarrow{f} \mathbb{R}$$

stetig differenzierbar (beziehungsweise glatt) ist.

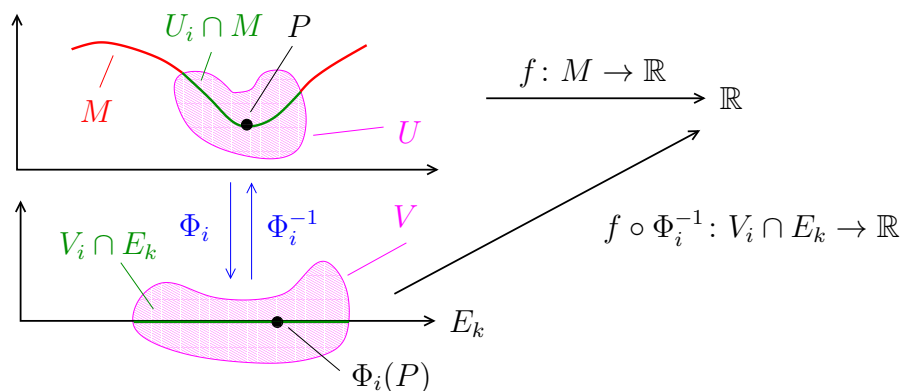


ABBILDUNG 12. Schematisches Bild für die Definition der Differenzierbarkeit einer Funktion auf einer Untermannigfaltigkeit.

Beispiel. Es sei $f: \mathbb{R}^n \rightarrow \mathbb{R}$ eine stetig differenzierbare Funktion. Dann ist die Einschränkung von f auf jede Untermannigfaltigkeit M von $\mathbb{R}^n$ ebenfalls stetig differenzierbar. In der Tat,

¹⁵ Wir fassen hierbei $V \cap E_k$ auf offensichtliche Weise als Teilmenge des $\mathbb{R}^k$ auf. Die Definition der stetigen Differenzierbarkeit macht nun in der Tat Sinn, nachdem

$$V \cap E_k \xrightarrow{\Phi^{-1}} U \cap M \xrightarrow{f} \mathbb{R}$$

eine Funktion von einer offenen Teilmenge in $\mathbb{R}^k$, nämlich $V \cap E_k$, nach $\mathbb{R}$ ist. Wir können darauf nun die Definition der stetigen Differenzierbarkeit aus der Analysis II anwenden.

denn für eine Karte $\Phi: U \rightarrow V$ ist dann $f \circ \Phi^{-1}: V \rightarrow \mathbb{R}$ die Verknüpfung von zwei stetig differenzierbaren Abbildungen, also ebenfalls stetig differenzierbar. Dann ist aber auch die Einschränkung von $f \circ \Phi^{-1}$ auf $V \cap E_k$ stetig differenzierbar.

Satz 2.7. *Es sei $f: M \rightarrow \mathbb{R}$ eine Funktion auf einer k -dimensionalen Untermannigfaltigkeit M . Dann gilt*

$$f \text{ ist glatt} \iff \text{für jede Karte } \Phi: U \rightarrow V \text{ ist } f \circ \Phi^{-1}: V \cap E_k \rightarrow \mathbb{R} \text{ glatt.}$$

Die analoge Aussage gilt auch für “stetig differenzierbar” anstatt “glatt”.

Beweis. Es sei M eine k -dimensionale Untermannigfaltigkeit und es sei $f: M \rightarrow \mathbb{R}$ eine Funktion. Wir beweisen die Aussage für “glatt”, die Aussage für “stetig differenzierbar” wird ganz analog bewiesen.

Die Aussage “ $\Leftarrow$ ” ist offensichtlich. Wir müssen nun noch die Aussage “ $\Rightarrow$ ” beweisen. Es sei also $\{\Phi_i: U_i \rightarrow V_i\}_{i \in I}$ ein Atlas, so dass für alle $i \in I$ die Funktion

$$V_i \cap E_k \xrightarrow{\Phi_i^{-1}} U_i \cap M \xrightarrow{f} \mathbb{R}$$

glatt ist. Es sei nun $\Phi: U \rightarrow V$ eine weitere Karte.

Es genügt zu zeigen, dass es für jedes $Q \in V \cap E_k$ eine offene Umgebung Y gibt, so dass $f \circ \Phi^{-1}: Y \cap E_k \rightarrow \mathbb{R}$ glatt ist. Es sei also $Q \in V \cap E_k$. Wir setzen $P = \Phi^{-1}(Q)$. Wir wählen ein $i \in I$, so dass $P \subset U_i$. Wir setzen $X = U \cap U_i$ und $Y = \Phi(X)$. Wir müssen zeigen, dass die Funktion

$$Y \cap E_k \xrightarrow{\Phi^{-1}} X \cap M \xrightarrow{f} \mathbb{R}$$

glatt ist. Dies ist der Fall, denn es ist

$$f \circ \Phi^{-1} = \underbrace{f \circ \Phi_i^{-1}}_{\text{glatt nach Voraussetzung}} \circ \underbrace{\Phi_i \circ \Phi^{-1}}_{\text{glatt, da } \Phi \text{ und } \Phi_i \text{ Diffeomorphismen sind}}.$$

□

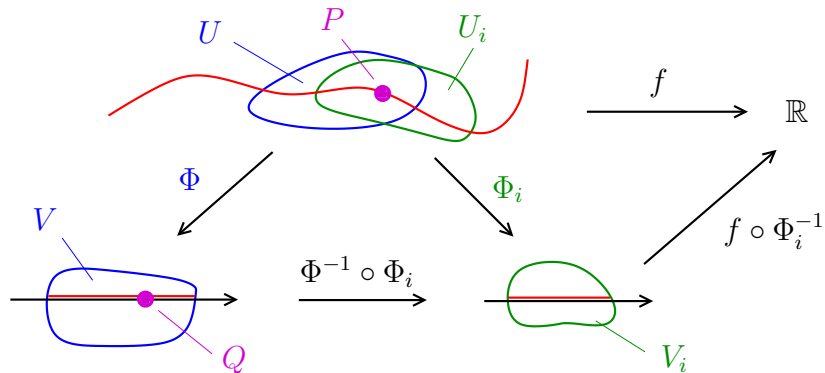


ABBILDUNG 13. Skizze zum Beweis von Satz 2.7.

3. INTEGRATION UND UNTERMANNIGFALTIGKEITEN

3.1. Erinnerung: Maß- und Integrationstheorie. In Analysis III hatten wir das Lebesgue-Maß und das Lebesgue-Integral eingeführt. Die genauen Definitionen interessieren uns jetzt nicht weiter. Wir erinnern in diesem kurzen Kapitel jedoch an die wichtigsten Eigenschaften vom Lebesgue-Maß und vom Lebesgue-Integral.

Folgenden Satz hatten wir in Kapitel 6 von Analysis III bewiesen.

Satz 3.1. *Jede kompakte Teilmenge und jede beschränkte offene Teilmenge von $\mathbb{R}^n$ ist messbar und besitzt ein endliches Volumen.*

Wir erinnern auch an Satz 7.4 aus der Analysis III.

Satz 3.2. *Es sei $T \in \text{GL}(n, \mathbb{R})$ eine invertierbare Matrix. Dann gilt für alle messbaren Teilmengen B von $\mathbb{R}^n$, dass*

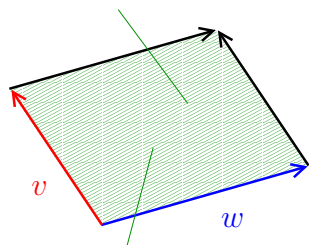
$$\text{Vol}_n(T \cdot B) = |\det(T)| \cdot \text{Vol}_n(B).$$

Für Vektoren $v_1, \dots, v_k$ in $\mathbb{R}^n$ bezeichnen wir mit

$$P(v_1, \dots, v_k) := \left\{ \sum_{i=1}^k \lambda_i v_i \mid \lambda_1, \dots, \lambda_k \in [0, 1] \right\}$$

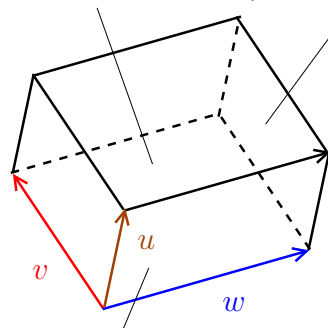
das durch $v_1, \dots, v_k$ in $\mathbb{R}^n$ aufgespannte *Parallelotop*. Wenn die Vektoren linear unabhängig sind, dann nennen wir $P(v_1, \dots, v_k)$ ein *k-dimensionales Parallelotop*. Für $k = n = 2$ ist dies insbesondere das durch v_1 und v_2 aufgespannte Parallelogramm. Für $k = n = 3$ ist dies der durch v_1, v_2 und v_3 aufgespannte Spat.

$P(v, w) = \{ \lambda v + \mu w \mid \lambda, \mu \in [0, 1] \}$
ist das von den Vektoren v und w
in $\mathbb{R}^2$ aufgespannte Parallelogramm



der Flächeninhalt beträgt $|\det(v \ w)|$

das Parallelotop $P(u, v, w)$ in $\mathbb{R}^3$



das Volumen beträgt $|\det(u \ v \ w)|$

ABBILDUNG 14.

Der folgende Satz 7.3 aus der Analysis III gibt nun insbesondere eine geometrische Interpretation des Absolutbetrags der Determinante einer quadratischen Matrix. Den Satz kann man auch leicht dadurch beweisen, dass man Satz 3.2 auf den Einheitswürfel $[0, 1]^n \subset \mathbb{R}^n$ anwendet.

Satz 3.3. *Es seien $v_1, \dots, v_n \in \mathbb{R}^n$. Dann gilt*

$$\text{Vol}_n(P(v_1, \dots, v_n)) = |\det(v_1 \dots v_n)|.$$

In Analysis III hatten wir den Begriff der Lebesgue-Integrierbarkeit und des Lebesgue-Integrals eingeführt. Im Folgenden zitieren wir Satz 9.4 aus Analysis III, dieser verbindet die Begriffe von Lebesgue-Integral und Riemann-Integral. Zudem erlaubt es uns dieser Satz in den meisten Fällen, dass Integral einer Funktion im reellen mithilfe der Methoden aus der Analysis I zu bestimmen.

Satz 3.4. *Jede Riemann-integrierbare Funktion $f: [a, b] \rightarrow \mathbb{R}$ ist auch Lebesgue-integrierbar und es gilt*

$$\underbrace{\int_a^b f(x) dx}_{\text{Riemann-Integral}} = \underbrace{\int_{[a,b]} f dv}_{\text{Lebesgue-Integral}}.$$

Im weiteren Verlauf der Vorlesung sagen wir einfach “integrierbar” anstatt “Lebesgue-integrierbar” und wir sagen “Integral” anstatt “Lebesgue-Integral”. Für eine integrierbare Funktion $f: U \rightarrow \mathbb{R}$, welche auf einer Teilmenge U von $\mathbb{R}^n$ definiert ist, verwenden wir, je nach Zusammenhang, unter anderem folgende Notationen für das Integral:

$$\int_U f dv = \int_U f(x) dx = \int_U f(y) dy.$$

Der nächste Satz folgt leicht aus Satz 8.20 aus der Analysis III. Dieser deckt die meisten Fälle ab, welche uns interessieren.

Satz 3.5. *Jede stetige Funktion $f: K \rightarrow \mathbb{R}$ auf einer kompakten Teilmenge K von $\mathbb{R}^n$ ist integrierbar.*

Der folgende Satz ist ein Spezialfall von Satz 8.17 aus der Analysis III Vorlesung.

Satz 3.6. *Es sei $f: A \rightarrow \mathbb{R}$ eine integrierbare Funktion auf einer messbaren Menge $A \subset \mathbb{R}^n$. Wenn sich $g: \mathbb{R}^n \rightarrow \mathbb{R}$ von f nur auf einer Nullmenge unterscheidet, dann ist auch g integrierbar und es gilt zudem, dass*

$$\int_A f dv = \int_A g dv.$$

Der nächste Satz erlaubt es uns den Absolutbetrag von Integralen nach oben abzuschätzen. Der Satz folgt leicht aus den Aussagen von Kapitel 8 in Analysis III.

Satz 3.7. *Es sei $A \subset \mathbb{R}^n$ eine messbare Menge und es sei $f: A \rightarrow \mathbb{R}$ eine integrierbare Funktion. Dann gilt*

$$\left| \int_A f dv \right| \leq \text{Vol}_n(A) \cdot \sup \{ |f(x)| \mid x \in A \}.$$

Der folgende Satz von Fubini erlaubt es uns in vielen Fällen das Integral im mehrdimensionalen auf mehrere eindimensionale Integral zurückzuführen.

Satz 3.8. (Fubini) Es sei $A \subset \mathbb{R}^k \times \mathbb{R}^m$ eine kompakte Teilmenge und es sei

$$\begin{aligned} f: A &\rightarrow \mathbb{R} \\ (x, y) &\mapsto f(x, y) \end{aligned}$$

eine stetige Funktion. Wir schreiben

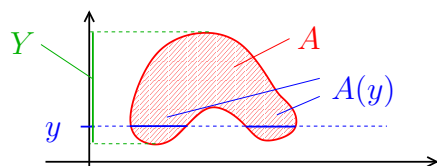
$$Y := \{y \in \mathbb{R}^m \mid \text{es gibt ein } x \in \mathbb{R}^m \text{ mit } (x, y) \in A\},$$

und für alle $y \in Y$ setzen wir

$$A(y) := \{x \in \mathbb{R}^m \mid (x, y) \in A\}.$$

Dann gilt

$$\int_A f(x, y) \, dx \, dy = \int_Y \int_{A(y)} f(x, y) \, dx \, dy.$$



nach dem Satz von Fubini gilt

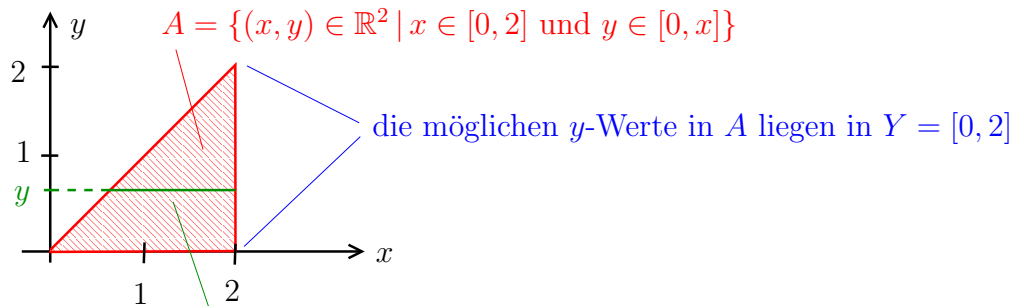
$$\int_A f(x, y) \, dx \, dy = \int_Y \int_{A(y)} f(x, y) \, dx \, dy$$

ABBILDUNG 15. Illustration vom Satz von Fubini.

Beispiel. Wir betrachten die Funktion $f(x, y) = x^2 + y^2$ auf dem Dreieck

$$A = \{(x, y) \in \mathbb{R}^2 \mid x \in [0, 2] \text{ und } y \in [0, x]\}.$$

Es folgt aus Satz 3.5, dass f integrierbar ist. Wir möchten nun das Integral von f mithilfe vom Satz von Fubini bestimmen. In diesem Fall ist



für ein gegebenes $y \in [0, 2]$ liegen die möglichen x -Werte in $A(y) = [y, 2]$

ABBILDUNG 16.

$Y =$ alle möglichen y -Werte von Punkten in $A = [0, 2]$.

Zudem gilt für jedes $y \in [0, 2]$, dass

$A(y) =$ alle x -Werte von Punkten in A mit y -Wert $y = [y, 2]$.

Es folgt also, dass

$$\int_A x^2 + y^2 dx dy = \int_{\substack{\uparrow \\ \text{Satz von Fubini}}} \int_{[0,2]} \int_{[y,2]} x^2 + y^2 dx dy = \int_{\substack{\uparrow \\ \text{Satz 3.4}}} \int_{y=0}^{y=2} \int_{x=y}^{x=2} x^2 + y^2 dx dy.$$

Der Rest der Rechnung besteht nun aus zwei einfachen Riemann-Integralen.

Der vielleicht wichtigste und auch schwierigste Satz den wir in Analysis III bewiesen hatten ist der Transformationssatz.

Satz 3.9. (Transformationssatz) *Es seien $U, V \subset \mathbb{R}^n$ offene Mengen, es sei $\Phi: U \rightarrow V$ ein Diffeomorphismus und es sei $f: V \rightarrow \mathbb{R}$ eine Funktion. Dann gilt*¹⁶

$$\int_U f(\Phi(x)) \cdot |\det D\Phi(x)| dx = \int_V f(y) dy.$$

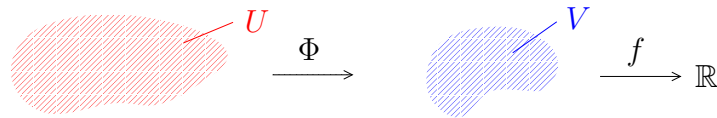


ABBILDUNG 17. Illustration der Transformationsformel.

Bevor wir den letzten Satz des Kapitels formulieren und beweisen führen wir noch folgende Sprechweisen ein, welche wir immer wieder in der Vorlesung verwenden werden.

- (1) Ein *abgeschlossener Quader* (oder oft auch nur kurz *Quader*) ist eine Teilmenge von $\mathbb{R}^n$ der Form

$$Q = [a_1, b_1] \times \cdots \times [a_n, b_n], \text{ wobei } a_i < b_i \text{ für } i = 1, \dots, n.$$

- (2) Ein *offener Quader* ist eine Menge der Form

$$Q = (a_1, b_1) \times \cdots \times (a_n, b_n), \text{ wobei } a_i < b_i \text{ für } i = 1, \dots, n.$$

Der Abschluß von einem offenen Quader ist natürlich ein abgeschlossener Quader. Zudem ist das Innere von einem abgeschlossenen Quader ein offener Quader.

Der folgende Satz besagt, dass man unter gewissen Voraussetzungen Integral und Differentiation vertauschen kann. Wir werden diesen Satz später im Beweis vom Integralsatz von Gauß benötigen.

¹⁶Die Gleichheit soll dahingehend interpretiert werden, dass das Integral auf der linken Seite existiert genau dann, wenn es auf der rechten Seite definiert ist, und wenn diese existieren, dann gilt die Gleichheit der Integrale.

Satz 3.10. Es sei $Q \subset \mathbb{R}^n$ ein abgeschlossener Quader und es sei

$$\begin{aligned} f: [a, b] \times Q &\rightarrow \mathbb{R} \\ (t, y) &\mapsto f(t, y) \end{aligned}$$

eine stetig differenzierbare¹⁷ Funktion. Dann gilt für alle $t_0 \in (a, b)$, dass

$$\underbrace{\frac{d}{dt} \int_Q f(t, y) dy}_{\text{Funktion in } t} \Big|_{t_0} = \int_Q \underbrace{\frac{\partial}{\partial t} f(t_0, y)}_{\text{partielle Ableitung von } f \text{ bei } (t_0, y)} dy.$$

Beweis. Wenn $\text{Vol}(Q) = 0$, dann folgt aus Satz 3.7, dass beide Seiten der gewünschten Gleichheit schon null sind. Es gibt also nichts zu beweisen. Wir nehmen nun an, dass $\text{Vol}(Q) > 0$.

Es sei also $t_0 \in (a, b)$. Wir wählen ein $\eta > 0$ mit $[t_0 - \eta, t_0 + \eta] \subset (a, b)$. Wir wollen also zeigen, dass

$$\frac{d}{dt} \int_Q f(t, y) dy \Big|_{t_0} - \int_Q \frac{\partial}{\partial t} f(t_0, y) dy = 0.$$

Mit anderen Worten wir wollen zeigen, dass

$$\lim_{h \rightarrow 0} \int_Q \frac{1}{h} (f(t_0 + h, y) - f(t_0, y)) - \frac{\partial}{\partial t} f(t_0, y) dy = 0.$$

Aus der Definition von partiellen Ableitungen folgt für jedes einzelne $y \in Q$, dass

$$\lim_{h \rightarrow 0} \frac{1}{h} (f(t_0 + h, y) - f(t_0, y)) - \frac{\partial}{\partial t} f(t_0, y) = 0.$$

Aber die “Geschwindigkeit” mit der dieser Grenzwert verschwindet hängt a priori von $y \in Q$ ab.

Wir wollen also zeigen, dass der Grenzwert verschwindet. Um dies zu zeigen, arbeiten wir direkt mit der Definition vom Grenzwert. Es sei also $\epsilon > 0$. Für $y \in Q$ setzen wir¹⁸

$$\delta(y) := \sup \left\{ \mu \in [0, \eta] \mid \left| \frac{1}{h} (f(t_0 + h, y) - f(t_0, y)) - \frac{\partial}{\partial t} f(t_0, y) \right| < \frac{\epsilon}{2 \text{Vol}(Q)} \text{ für alle } 0 < |h| < \mu \right\}.$$

Nachdem f stetig differenzierbar ist, folgt, dass $\delta(y) > 0$.

Am liebsten würde man jetzt das Minimum von δ auf Q betrachten. Der Definitionsbereich Q ist zwar kompakt, aber es ist nicht klar, warum δ stetig sein soll. In der Tat werden wir in Übungsblatt 2 sehen, dass δ im allgemeinen nicht stetig ist. Wir müssen deshalb mit einer schwächeren Version von Stetigkeit arbeiten.

¹⁷Wir sagen eine Funktion g auf einem abgeschlossenen Quader K ist stetig differenzierbar, wenn g stetig ist, wenn g im Inneren von K stetig differenzierbar ist, und wenn sich die partiellen Ableitungen von g vom Inneren zu einer stetigen Funktion auf dem ganzen Quader K fortsetzen lassen.

¹⁸Grob gesprochen ist also $\delta(y)$ das “maximale δ ”, welches für $y \in Q$ funktioniert.

In Übungsblatt 2 werden wir sehen, dass die Funktion $\delta: Q \rightarrow \mathbb{R}_{\geq 0}$ halbstetig nach unten ist, dies bedeutet ¹⁹

$$\forall_{y_0 \in Q} \quad \forall_{\alpha > 0} \quad \exists_{\beta > 0} \quad \forall_{y \in B_\beta(y_0) \cap Q} \quad \delta(y) > \delta(y_0) - \alpha.$$

Wir betrachten nun

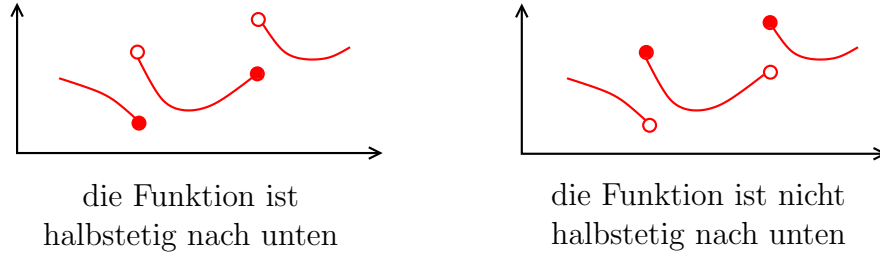


ABBILDUNG 18.

$$\mu := \inf\{\delta(y) \mid y \in Q\}.$$

In Übungsblatt 1 hatten wir schon gesehen, dass aus der Tatsache, dass $\delta(y) > 0$ für alle $y \in Q$, der Tatsache, dass δ halbstetig nach unten, und der Voraussetzung, dass Q kompakt ist, folgt, dass $\mu > 0$. Für alle $h \in (0, \mu)$ gilt nun, dass

$$\begin{aligned} & \left| \int_Q \frac{1}{h} (f(t_0 + h, y) - f(t_0, y)) - \frac{\partial}{\partial t} f(t_0, y) dt \right| \\ & \leq \underbrace{\text{Vol}(Q) \cdot \sup \left\{ \left| \frac{1}{h} (f(t_0 + h, y) - f(t_0, y)) - \frac{\partial}{\partial t} f(t_0, y) \right| \mid y \in Q \right\}}_{\leq \frac{\epsilon}{2 \text{Vol}(Q)}} < \epsilon. \\ & \quad \uparrow \\ & \text{Satz 3.7} \end{aligned}$$

Wir haben damit gezeigt, dass der Grenzwert verschwindet, und wir haben damit den Satz bewiesen. $\square$

3.2. Motivation für Integrale auf Untermannigfaltigkeiten. Für eine gegebene k -dimensionale Untermannigfaltigkeit $M \subset \mathbb{R}^n$ wollen wir nun das Integral für “vernünftige” Funktionen $f: M \rightarrow \mathbb{R}$ einführen. Dieses Integral sollte dabei sicherlich folgende Eigenschaft besitzen: wenn $M = P(v_1, \dots, v_k)$ ein k -dimensionales Parallelogramm in $\mathbb{R}^n$ ist und wenn $f: M \rightarrow \mathbb{R}$ eine konstante Funktion mit Wert C ist, dann soll gelten²⁰

$$\int_M f dv = C \cdot k\text{-dimensionales Volumen von } M.$$

¹⁹Anschaulich gesprochen kann also die Funktion δ nach oben springen, aber nicht nach unten.

²⁰Die Tatsache, dass ein Polytop streng genommen keine Untermannigfaltigkeit ist, soll uns bei dieser Diskussion nicht stören. Die inneren Punkte von einem k -dimensionalen Polytop bilden beispielsweise eine Untermannigfaltigkeit.

Die rechte Seite ist von der Bedeutung wohl intuitiv klar, aber wir wollen im Folgenden Kapitel eine präzise Definition vom k -dimensionalen Volumen eines k -dimensionalen Parallelotops geben, und wir wollen eine brauchbare Gleichung für dieses Volumen herleiten.

3.3. Einschub: Lineare Algebra. Es sei $k \leq n$. Im Folgenden identifizieren wir die k -dimensionale Ebene

$$E_k := \{(x_1, \dots, x_n) \in \mathbb{R}^n \mid x_1, \dots, x_k \in \mathbb{R}, x_{k+1} = \dots = x_n = 0\}$$

mit $\mathbb{R}^k$. Es seien $v_1, \dots, v_k$ linear unabhängige Vektoren in $\mathbb{R}^n$. Dann spannen die Vektoren das k -dimensionale Parallelotop

$$P := P(v_1, \dots, v_k) := \left\{ \sum_{i=1}^k l_i v_i \mid l_1, \dots, l_k \in [0, 1] \right\}$$

auf. Wenn $k < n$ dann folgt aus der Diskussion von Kapitel 6.2 der Analysis III, dass $\text{Vol}_n(P) = 0$. Aber zumindest anschaulich ist es klar, dass P ein “ k -dimensionales Volumen besitzt”. Wir wollen nun ein vernünftiges k -dimensionales Volumen $\text{Vol}_k(P)$ für solche k -dimensionale Polytope einführen. Hierbei soll “vernünftig” folgendes heißen:

- (1) Wenn das Parallelotop P schon in E_k liegt, dann können wir P als Teilmenge von $\mathbb{R}^k$ auffassen und $\text{Vol}_k(P)$ soll gerade das übliche k -dimensionale Volumen einer Teilmenge von $\mathbb{R}^k$ sein.
- (2) Das Volumen soll invariant unter Anwendung von orthogonalen Matrizen sein, d.h. wenn P ein k -dimensionales Parallelotop ist und $A \in O(n)$ eine orthogonale Matrix, dann soll

$$\text{Vol}_k(A \cdot P) = \text{Vol}_k(P)$$

gelten.

Wir werden jetzt diese beiden Wünsche verwenden um $\text{Vol}_k(P)$ allgemein einzuführen. Dazu benötigen wir folgenden Satz.

Satz 3.11. *Es sei $P \subset \mathbb{R}^n$ ein k -dimensionales Parallelotop. Dann gelten folgende Aussagen:*

- (1) *Es gibt eine orthogonale Matrix $A \in O(n)$, so dass $A \cdot P$ in $E_k = \mathbb{R}^k$ liegt.*
- (2) *Wenn A und B zwei orthogonale Matrizen in $O(n)$ sind, so dass $A \cdot P$ und $B \cdot P$ in $E_k = \mathbb{R}^k$ liegen, dann gibt es auch eine orthogonale Matrix $Q \in O(k)$, so dass*

$$Q \cdot (A \cdot P) = B \cdot P,$$

hierbei fassen wir $A \cdot P$ und $B \cdot P$ als Teilmengen von $\mathbb{R}^k$ auf.

- (3) *Unter der Voraussetzung von (2) gilt*

$$\text{Vol}_k(A \cdot P) = \text{Vol}_k(B \cdot P).$$

Beweis. Es sei also P ein von linear unabhängigen Vektoren $v_1, \dots, v_k \in \mathbb{R}^n$ aufgespanntes Parallelotop.

- (1) Wir müssen nun also zeigen, dass es eine orthogonale Matrix A gibt, so dass $Av_1, \dots, Av_k \in E_k = \mathbb{R}^k$. Nachdem $v_1, \dots, v_k$ linear unabhängig sind, können wir nach dem Basisergänzungssatz diese Vektoren zu einer Basis $v_1, \dots, v_n$ von $\mathbb{R}^n$ ergänzen. Wir wenden nun den Gram-Schmidt-Algorithmus auf diese Basis $v_1, \dots, v_n$ an und erhalten eine orthonormale Basis $w_1, \dots, w_n$ mit der Eigenschaft, dass für alle $m \in \{1, \dots, n\}$ gilt $\text{Span}\{v_1, \dots, v_m\} = \text{Span}\{w_1, \dots, w_m\}$. Wir bezeichnen nun mit C die orthogonale Matrix, deren Spalten gerade durch $w_1, \dots, w_n$ gegeben sind. Dann gilt $C \cdot e_i = w_i$. Insbesondere gilt

$$C \cdot E_k = C \cdot \text{Span}(e_1, \dots, e_k) = \text{Span}(w_1, \dots, w_k) = \text{Span}(v_1, \dots, v_k).$$

Daraus folgt nun auch, dass

$$C^{-1} \cdot P \subset C^{-1} \cdot \text{Span}(v_1, \dots, v_k) = \text{Span}(e_1, \dots, e_k) = E_k.$$

Also hat die orthogonale Matrix $A = C^{-1}$ die gewünschte Eigenschaft.

- (2) Es seien A und B zwei orthogonale Matrizen, so dass $A \cdot P$ und $B \cdot P$ in $E_k = \mathbb{R}^k$ liegen. Die Abbildung

$$\begin{aligned} \Phi: \mathbb{R}^n &\rightarrow \mathbb{R}^n \\ v &\mapsto B \cdot A^{-1} \cdot v \end{aligned}$$

ist dann eine Isometrie und es gilt $\Phi(A \cdot P) = B \cdot P$. Nachdem $A \cdot P$ und $B \cdot P$ in E_k liegen und E_k aufspannen, folgt es, dass $\Phi(E_k) = E_k$. Die Einschränkung von Φ auf $E_k = \mathbb{R}^k$ ist eine Isometrie und kann also durch eine orthogonale Matrix Q dargestellt werden. Diese Matrix hat die gewünschte Eigenschaft.

- (3) Diese Aussage folgt nun aus (2) und Satz 3.2, denn für jede orthogonale Matrix Q gilt $|\det(Q)| = 1$. $\square$

Es sei nun $P \subset \mathbb{R}^n$ ein k -dimensionales Parallelotop. Nach Satz 3.11 (1) existiert eine orthogonale Matrix $A \in O(n)$, so dass $A \cdot P$ in $E_k = \mathbb{R}^k$ liegt. Wir definieren nun das *k-dimensionale Volumen von P* als

$$\text{Vol}_k(P) := \text{Vol}_k(\underbrace{A \cdot P}_{\subset E_k = \mathbb{R}^k}).$$

Es folgt aus Satz 3.11 (3), dass diese Definition nicht von der Wahl der orthogonalen Matrix A abhängt.

Es stellt sich nun die Frage, ob man das Volumen von P auch “direkt” von den Vektoren $v_1, \dots, v_k$ ablesen kann, ohne zuerst eine orthogonale Matrix A wie oben finden zu müssen.

Wir machen hierzu folgende Überlegungen. Wenn die Vektoren $v_1, \dots, v_k \in E_k = \mathbb{R}^k$ liegen, dann gilt nach Satz 3.3, dass

$$\text{Vol}_k(P(v_1, \dots, v_k)) = |\det(v_1 \dots v_k)|.$$

In diesem Fall können wir also das k -dimensionale Volumen mithilfe der Koordinaten der Vektoren $v_1, \dots, v_k$ bestimmen. Nachdem sich das k -dimensionale Volumen unter Anwendung einer orthogonalen Matrix nicht ändern soll, wollen wir nun einen Ausdruck in den Koordinaten $v_1, \dots, v_k$ hinschreiben, welcher folgende zwei Eigenschaften besitzt:

- (1) wenn $v_1, \dots, v_k$ in $E_k = \mathbb{R}^k$ liegen erhalten wir den Absolutbetrag der Determinante der Matrix $(v_1 \dots v_k)$,
 (2) der Ausdruck ändert sich nicht unter Anwendung einer orthogonalen Matrix.

Die Idee ist nun die Determinante der Matrix $(v_1 \dots v_k)$ zu bestimmen durch die Skalarprodukte²¹ $v_i \cdot v_j$, $i = 1, \dots, k$, denn diese sind invariant unter Anwendung einer orthogonalen Matrix. Wir benötigen nun folgendes Lemma aus der linearen Algebra.

Lemma 3.12. *Es seien $u_1, \dots, u_k \in \mathbb{R}^k$ beliebig. Dann gilt*

$$|\det(u_1 \dots u_k)| = \sqrt{\det((u_i \cdot u_j)_{i,j=1,\dots,k})}.$$

Beweis. Wir schreiben $u_i = (u_{i1}, \dots, u_{ik})$ und wir bezeichnen mit U die $k \times k$ -Matrix $U = (u_1 \dots u_k)$. Dann gilt, dass

$$\det((u_i \cdot u_j)_{i,j=1,\dots,k}) = \det(U \cdot U^T) = \det(U) \cdot \det(U^T) = \det(U)^2.$$

↑

$$\text{denn der } ij\text{-Eintrag von } U \cdot U^T = \sum_{l=1}^k u_{il} \cdot u_{jl} = u_i \cdot u_j$$

□

Für eine $n \times k$ -Matrix V mit Spaltenvektoren $v_1, \dots, v_k \in \mathbb{R}^n$ definieren wir nun die *Gramsche Determinante* wie folgt:²²

$$\text{Gram}(V) = \text{Gram}(\underbrace{v_1 \dots v_k}_{n \times k\text{-Matrix}}) := \det(\underbrace{(v_i \cdot v_j)_{i,j=1,\dots,k}}_{k \times k\text{-Matrix}}).$$

Lemma 3.13. *Es seien $v_1, \dots, v_k$ linear unabhängige Vektoren in $\mathbb{R}^n$. Dann gilt*

$$\text{Vol}_k(P(v_1, \dots, v_k)) = \sqrt{\text{Gram}(v_1 \dots v_k)}.$$

Beweis. Es seien also $v_1, \dots, v_k$ linear unabhängige Vektoren in $\mathbb{R}^n$. Wir wählen eine orthogonale Matrix A mit $A \cdot P(v_1, \dots, v_k) \subset E_k$. Dann gilt

$$\begin{aligned} \text{Vol}_k(P(v_1, \dots, v_k))^2 &= \text{Vol}_k(\underbrace{A \cdot P(v_1, \dots, v_k)}_{\subset E_k = \mathbb{R}^k})^2 = \text{Vol}_k(P(\underbrace{Av_1, \dots, Av_k}_{Av_i \in E_k = \mathbb{R}^k}))^2 \\ &\quad \uparrow \text{per Definition} \\ &= \det(\underbrace{Av_1 \dots Av_k}_{Av_i \in E_k = \mathbb{R}^k}) = \det(\underbrace{(Av_i \cdot Av_j)_{i,j=1,\dots,k}}_{Av_i \in E_k = \mathbb{R}^k}) \\ &\quad \uparrow \text{Satz 3.3} \qquad \qquad \qquad \uparrow \text{Lemma 3.12} \\ &= \det(\underbrace{(Av_i \cdot Av_j)_{i,j=1,\dots,k}}_{Av_i \in E_k \subset \mathbb{R}^n}) = \det((v_i \cdot v_j)_{i,j=1,\dots,k}). \\ &\quad \uparrow \text{denn } A \text{ ist orthogonal} \\ &= \text{Gram}(v_1 \dots v_k)^2. \end{aligned}$$

²¹Wie üblich bezeichnen wir für Vektoren $v = (v_1, \dots, v_n)$ und $w = (w_1, \dots, w_n)$ in $\mathbb{R}^n$ mit

$$v \cdot w := \langle v, w \rangle := \sum_{i=1}^n v_i w_i \in \mathbb{R}$$

das Skalarprodukt.

²²Der Beweis von Lemma 3.12 zeigt, dass $\det((v_i \cdot v_j)_{i,j=1,\dots,k}) \geq 0$.

Wir erhalten die gewünschte Aussage durch Wurzelziehen. $\square$

Beispiel. Es seien $a = (a_1, a_2, a_3)$ und $b = (b_1, b_2, b_3)$ zwei linear unabhängige Vektoren in $\mathbb{R}^3$. Das *Kreuzprodukt* $a \times b$ ist definiert als der Vektor

$$a \times b := \sum_{i=1}^3 (-1)^{i+1} \det \left(\begin{array}{c} \text{die Matrix } (a \ b), \text{ wobei} \\ \text{die } i\text{-te Zeile gestrichen ist} \end{array} \right) = \begin{pmatrix} a_2 b_3 - a_3 b_2 \\ -a_1 b_3 + a_3 b_1 \\ a_1 b_2 - a_2 b_1 \end{pmatrix}.$$

In Übungsblatt 2 werden wir mithilfe von Lemma 3.13 zeigen, dass

$$\|a \times b\| = \text{Flächeninhalt des Parallelotops } P(a, b).$$

Im weiteren Verlauf der Vorlesung werden wir auch noch folgende Eigenschaft der Gramschen Determinante benötigen.

Lemma 3.14. *Es sei V eine $n \times k$ -Matrix und es sei A eine $k \times k$ -Matrix. Dann gilt*

$$\text{Gram}(\underbrace{V \cdot A}_{n \times k\text{-Matrix}}) = \det(A)^2 \cdot \text{Gram}(V).$$

Bemerkung. Für eine $n \times k$ -Matrix V gilt $\text{Gram}(V) = \det(V^T V)$. Wenn V nun eine quadratische Matrix ist, dann folgt die Aussage von Lemma 3.14 leicht aus dieser Bemerkung der Multiplikativität der Determinante.

Beweis. Es sei $V = (v_1 \dots v_k)$ eine $n \times k$ -Matrix und es sei A eine $k \times k$ -Matrix. Die Matrix A ist das Produkt von Elementarmatrizen, d.h. von Matrizen der Form

$$D = \begin{pmatrix} l_1 & 0 & 0 \\ 0 & \ddots & 0 \\ 0 & 0 & \lambda_k \end{pmatrix}, \quad P = \begin{pmatrix} \text{id} & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & \text{id} & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & \text{id} \end{pmatrix} \quad \text{und} \quad Q = \begin{pmatrix} \text{id} & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & \lambda & 0 \\ 0 & 0 & \text{id} & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & \text{id} \end{pmatrix}.$$

Es folgt nun leicht aus der Multiplikativität der Determinante, dass es genügt die Aussage für diese drei Typen von Matrizen zu beweisen. Dies geschieht in allen drei Fällen durch eine elementare Rechnung.

Wir führen dies nun noch etwas genauer aus. Wir betrachten dazu die drei Typen von Matrizen.

- (1) Es sei $D = (d_{ij})$ eine Diagonalmatrix mit Einträgen $\lambda_1, \dots, \lambda_k$. Dann gilt

$$\begin{aligned} \text{Gram}((v_1 \dots v_k) \cdot D) &= \text{Gram}(l_1 v_1 \dots l_k v_k) \xrightarrow{\text{Übungsblatt 2}} l_1^2 \dots l_k^2 \cdot \text{Gram}(v_1 \dots v_k) \\ &= \det(D)^2 \cdot \text{Gram}(v_1 \dots v_k). \end{aligned}$$

- (2) Es sei P wie oben, wobei die Einser in den Spalten i und j stehen. Dann ist

$$\begin{aligned} \text{Gram}((v_1 \dots v_k) \cdot P) &= \text{Gram}(v_1 \dots v_j \dots v_i \dots v_k) \xrightarrow{\text{elementare Rechnung}} \text{Gram}(v_1 \dots v_k) \\ &= \det(P)^2 \cdot \text{Gram}(v_1 \dots v_k). \end{aligned}$$

(3) Es sei Q wie oben, wobei λ in der i -ten Spalte und der j -ten Zeile steht. Dann ist

$$\begin{aligned} \text{Gram}((v_1 \dots v_k) \cdot Q) &= \text{Gram}(v_1 \dots v_i + \lambda v_j \dots v_k) \xrightarrow{\text{elementare Rechnung}} \text{Gram}(v_1 \dots v_k) \\ &= \det(Q)^2 \cdot \text{Gram}(v_1 \dots v_k). \quad \square \end{aligned}$$

Das folgende Lemma gibt uns also eine weitere Möglichkeit die Gramsche Determinante zu bestimmen.

Lemma 3.15. *Es sei A eine $n \times k$ -Matrix mit $k \leq n$. Dann gilt*

$$\text{Gram}(A) = \sum_{i_1 < \dots < i_k} \det(A_{i_1, \dots, i_k})^2.$$

Hierbei summieren wir über alle $i_1, \dots, i_k$ mit $1 \leq i_1 < \dots < i_k \leq n$ und wir bezeichnen mit $A_{i_1, \dots, i_k}$ die Matrix, welche wir aus A erhalten indem wir nur die Zeilen $i_1, \dots, i_k$ betrachten.

Beispiel. Wir betrachten die 3×2 -Matrix

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

Dann gilt

$$\begin{aligned} \text{Gram}(A) &= \sum_{i_1 < i_2} (\det A_{i_1, i_2})^2 = \det \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}^2 + \det \begin{pmatrix} 1 & 0 \\ a & b \end{pmatrix}^2 + \det \begin{pmatrix} 0 & 1 \\ a & b \end{pmatrix}^2 = 1 + a^2 + b^2. \\ &\quad \uparrow \quad \quad \quad \uparrow \\ &\quad \text{Lemma 3.15} \quad \quad \text{die einzigen Möglichkeiten sind } (i_1, i_2) = (1, 2), (1, 3), (2, 3) \end{aligned}$$

Andererseits kann man auch direkt berechnen, dass

$$\text{Gram}(A) = \det \begin{pmatrix} 1+a^2 & ab \\ ab & 1+b^2 \end{pmatrix} = (1+a^2)(1+b^2) - a^2b^2 = 1 + a^2 + b^2.$$

Nachdem dies eigentlich ein Lemma aus der Linearen Algebra ist werden wir den Beweis des Lemmas nur skizzieren.

Beweisskizze. Im Beweis von Lemma 3.12 hatten wir insbesondere gesehen, dass für eine $n \times k$ -Matrix gilt $\text{Gram}(A) = \det(A^T A)$. Es genügt nun also folgende etwas allgemeinere Aussage zu beweisen.

Behauptung. Es seien A und B zwei $n \times k$ -Matrizen mit $k \leq n$. Dann gilt

$$\det(A^T B) = \sum_{i_1 < \dots < i_k} \det(A_{i_1, \dots, i_k}) \cdot \det(B_{i_1, \dots, i_k}).$$

Es sei A eine festgewählte $n \times k$ -Matrix mit $k \leq n$. Wir wollen nun zeigen, dass die obige Gleichung für alle $n \times k$ -Matrizen gilt. Für $b_1, \dots, b_k \in \mathbb{R}^n$ bezeichnen wir im Folgenden mit $(b_1 \dots b_k)$ die $n \times k$ -Matrix deren Spalten gerade $b_1, \dots, b_k$ sind. Wir machen folgende Beobachtungen:

- (1) Eine elementare Rechnung zeigt, dass die Gleichung für alle Matrizen der Form $B = (e_{j_1} \dots e_{j_k})$ gilt, wobei $e_1, \dots, e_n$ wie üblich die Einheitsvektoren von $\mathbb{R}^n$ sind.
- (2) Wenn die Gleichung für eine Matrix der Form $B = (b_1 \dots b_k)$ gilt, dann folgt aus der Linearität der Determinante, dass die Gleichung auch für jede Matrix der Form $B' = (b_1 \dots, \lambda b_i \dots b_k)$.
- (3) Wenn die Gleichung schon für zwei Matrizen der Form $B' = (b_1, \dots, b'_i, \dots, b_k)$ und $B'' = (b_1 \dots b'_i \dots b_k)$ gilt, dann folgt aus der Linearität der Determinante, dass die Gleichung auch für die Matrix $B = (b_1 \dots b'_i + b''_i \dots b_k)$ gilt.

Jede $n \times k$ -Matrix kann man aus Matrizen der Form (1) durch mehrfache Anwendungen der Operationen (2) und (3) erhalten.²³ Wir haben damit die Behauptung bewiesen. $\square$

3.4. Integration auf Untermannigfaltigkeiten I. Es sei M eine k -dimensionale Untermannigfaltigkeit und es sei $f: M \rightarrow \mathbb{R}$ eine Funktion. In diesem Kapitel nehmen wir an, dass es eine Karte $\Phi: U \rightarrow V$ gibt, so dass f außerhalb von $U \cap M$ verschwindet.

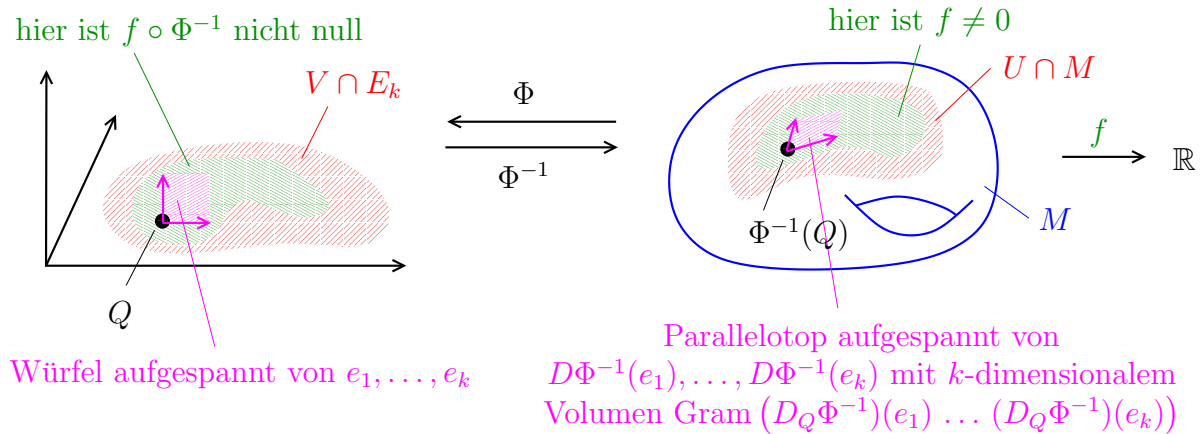


ABBILDUNG 19. Illustration der Transformationsformel.

Wir können jetzt natürlich die Funktion $f \circ \Phi^{-1}$ auf $V \cap E_k$ betrachten. Allerdings hatten wir in Lemma 3.13 gesehen, dass die Abbildung Φ^{-1} das k -dimensionale Volumen an einem Punkt $Q \subset V \cap E_k$ um den Faktor

$$\text{Gram}_{\Phi}(Q) := \text{Gram} \left(\underbrace{(D_Q \Phi^{-1})(e_1) \dots (D_Q \Phi^{-1})(e_k)}_{\text{die ersten } k \text{ Spalten von } D_Q \Phi^{-1}} \right)$$

“verzerrt”. Wir sagen jetzt $f: M \rightarrow \mathbb{R}$ ist integrierbar, wenn

$$\begin{aligned} V \cap E_k &\rightarrow \mathbb{R} \\ x &\mapsto f(\Phi^{-1}(x)) \cdot \sqrt{\text{Gram}_{\Phi}(x)} \end{aligned}$$

²³Warum ist dies der Fall?

integrierbar ist, und wir definieren^{24 25}

$$\int_M f := \int_{V \cap E_k} f(\Phi^{-1}(x)) \cdot \sqrt{\text{Gram}_\Phi(x)} dx.$$

Der folgende Satz besagt nun, dass diese Definitionen nicht von der Wahl der Karte Φ abhängen.

Satz 3.16. *Es sei $f: M \rightarrow \mathbb{R}$ eine Funktion auf einer k -dimensionalen Untermannigfaltigkeit M . Zudem seien $\Phi: U \rightarrow V$ und $\tilde{\Phi}: \tilde{U} \rightarrow \tilde{V}$ zwei Karten, so dass f außerhalb von $U \cap \tilde{U}$ verschwindet. Dann gilt²⁶*

$$\int_{V \cap E_k} f(\Phi^{-1}(x)) \cdot \sqrt{\text{Gram}_\Phi(x)} dx = \int_{\tilde{V} \cap E_k} f(\tilde{\Phi}^{-1}(y)) \cdot \sqrt{\text{Gram}_{\tilde{\Phi}}(y)} dy.$$

Beweis. O.B.d.A. können wir annehmen, dass $U = \tilde{U}$.²⁷ Wir bezeichnen mit Ψ und $\tilde{\Psi}$ die Umkehrabbildungen von Φ und $\tilde{\Phi}$. Wir betrachten die Diffeomorphismen $\Theta = \tilde{\Phi} \circ \Psi$ und $\Gamma = \Phi \circ \tilde{\Psi}$. (Zur Orientierung hilft es vielleicht Abbildung 20 zu studieren.)

Wir schreiben $V' = V \cap E_k$ und $\tilde{V}' = \tilde{V} \cap E_k$. Wir fassen V' und $\tilde{V}'$ als Teilmengen von $\mathbb{R}^k$ auf. Wir bezeichnen nun mit Ψ' und Θ' die Einschränkungen von Ψ und Θ auf $V' = V \cap E_k$ und wir bezeichnen mit $\tilde{\Psi}'$ und Γ' die Einschränkungen von $\tilde{\Psi}$ und Γ auf $\tilde{V}' = \tilde{V} \cap E_k$. Die Abbildung $\Gamma': \tilde{V}' \rightarrow V'$ ist dann ein Diffeomorphismus.²⁸

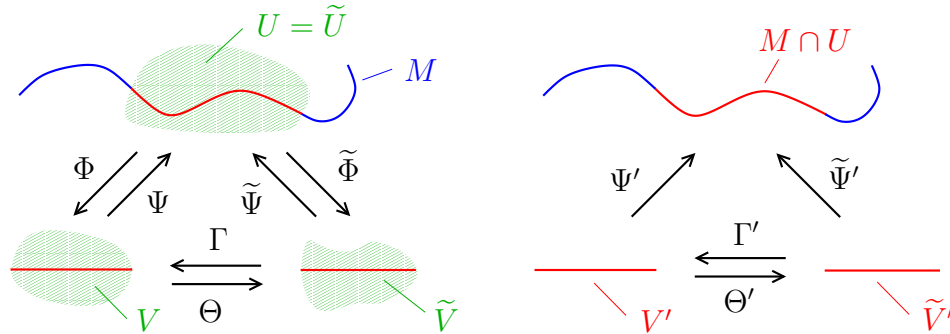


ABBILDUNG 20. Illustration zum Beweis von Satz 3.16.

²⁴Die Notation $\int_M f$ "ohne ein dx " ist kein Tippfehler, sondern beabsichtigt.

²⁵Wir fassen hierbei $V \cap E_k$ natürlich als Teilmenge von $\mathbb{R}^k$ auf, d.h. das Integral auf der rechten Seite ist das Integral einer Funktion auf einer Teilmenge von $\mathbb{R}^k$.

²⁶Genau gesprochen ist die Aussage, dass wenn das eine Integral existiert, dann existiert auch das andere, und die Integrale stimmen überein.

²⁷In der Tat, denn wir können zu $U \cap \tilde{U}$, $\Phi(U \cap \tilde{U})$ und $\tilde{\Phi}(U \cap \tilde{U})$ übergehen, ohne dass sich die Aussagen verändern.

²⁸In der Tat, denn Γ' ist offensichtlich bijektiv, und als Einschränkung einer glatten Abbildung wieder glatt. Das gleiche Argument zeigt aber auch, dass auch die Umkehrabbildung von Γ' , nämlich Θ' glatt ist.

Es ist nun

$$\begin{aligned}
\int_{V \cap E_k} f(\Phi^{-1}(x)) \cdot \sqrt{\text{Gram}_{\Phi}(x)} dx &= \int_{V'} f(\Psi'(x)) \cdot \sqrt{\text{Gram}(D_x \Psi')} dx \\
&= \int_{\tilde{V}'} f(\Psi'(\Gamma'(y))) \cdot \sqrt{\text{Gram}_{\Phi}(D_{\Gamma'(y)} \Psi')} \cdot |\det D_y \Gamma'| dy \\
&\quad \uparrow \tilde{V}' \\
&\text{Transformationssatz angewandt auf } \Gamma': \tilde{V}' \rightarrow V \\
&= \int_{\tilde{V}'} f(\tilde{\Psi}'(y)) \cdot \sqrt{\text{Gram}(D_y \tilde{\Psi}' \cdot D_{\Gamma'(y)} \Theta')} \cdot |\det D_y \Gamma'| dy \\
&\quad \uparrow \tilde{V}' \\
&\text{Kettenregel angewandt auf } \Psi' = \tilde{\Psi}' \circ \Theta' \\
&= \int_{\tilde{V}'} f(\tilde{\Psi}'(y)) \cdot |\det(D_{\Gamma'(y)} \Theta')| \cdot \sqrt{\text{Gram}(D_y \tilde{\Psi}')} \cdot |\det D_y \Gamma'| dy \\
&\quad \uparrow \tilde{V}' \\
&\text{Lemma 3.14 angewandt auf die } k \times k\text{-Matrix } D_{\Gamma'(y)} \Theta' \\
&= \int_{\tilde{V}'} f(\tilde{\Psi}'(y)) \cdot \sqrt{\text{Gram}(D_y \tilde{\Psi}')} dy \\
&\quad \uparrow \tilde{V}' \\
&\text{denn } D_{\Gamma'(y)} \Theta' = (D_y \Gamma')^{-1} \\
&= \int_{\tilde{V} \cap E_k} f(\tilde{\Phi}^{-1}(y)) \cdot \sqrt{\text{Gram}_{\tilde{\Phi}}(y)} dy. \quad \square
\end{aligned}$$

Wir wollen nun in den folgenden Kapiteln die Integration von beliebigen Funktionen auf einer Untermannigfaltigkeit auf den jetzt gerade betrachteten Fall zurückführen.

3.5. Integration auf Untermannigfaltigkeiten II. In diesem Kapitel definieren wir das Integral von Funktionen auf Untermannigfaltigkeiten, welche einen endlichen Atlas besitzen, d.h. einen Atlas, welcher aus endlich vielen Karten besteht. Beispielsweise ist es offensichtlich²⁹, dass jede kompakte Untermannigfaltigkeit einen endlichen Atlas besitzt.

Definition. Es sei M eine Untermannigfaltigkeit und es sei zudem $\mathcal{A} = \{\Phi_i: U_i \rightarrow V_i\}_{i \in I}$ ein Atlas für M . Eine *dem Atlas $\mathcal{A}$ untergeordnete messbare Zerlegung von M* ist eine Menge $\mathcal{Z}$ von Teilmengen $Z_i, i \in I$ mit folgenden Eigenschaften:

- (1) für alle i gilt $Z_i \subset U_i \cap M$,
- (2) für alle i ist $\Phi_i(Z_i)$ eine messbare Teilmenge von $E_k = \mathbb{R}^k$,
- (3) es ist³⁰

$$M = \bigsqcup_{i \in I} Z_i.$$

Satz 3.17. *Es sei M eine Untermannigfaltigkeit. Dann existiert zu jedem endlichen Atlas $\mathcal{A}$ eine dem Atlas $\mathcal{A}$ untergeordnete messbare Zerlegung von M .*

²⁹Warum ist das “offensichtlich”?

³⁰Zur Erinnerung, $X = A \sqcup B$ bedeutet, dass X die disjunkte Vereinigung von A und B ist, d.h. es ist $X = A \cup B$ und $A \cap B = \emptyset$.

Beweis. Es sei M eine Untermannigfaltigkeit und es sei $\mathcal{A} = \{\Phi_i: U_i \rightarrow V_i \mid i = 1, \dots, m\}$ ein endlicher Atlas für M . Wir setzen $Z_1 = U_1 \cap M$ und für $i = 2, \dots, m$ definieren wir iterativ

$$Z_i := (U_i \cap M) \setminus \bigcup_{j=1}^{i-1} Z_{j-1}.$$

Es ist nun offensichtlich, dass die Mengen $Z_1, \dots, Z_m$ die Eigenschaften (1) und (3) besitzen. Der Beweis, dass (2) gilt, ist eine freiwillige Übungsaufgabe. $\square$

Es sei $f: M \rightarrow \mathbb{R}$ eine Funktion auf einer Untermannigfaltigkeit, welche einen endlichen Atlas $\mathcal{A} = \{\Phi_i: U_i \rightarrow V_i \mid i = 1, \dots, m\}$ besitzt. Wir wählen eine dem Atlas $\mathcal{A}$ untergeordnete messbare Zerlegung $\mathcal{Z} = \{Z_1, \dots, Z_m\}$ von M . Wir sagen f ist integrierbar, wenn für alle $i = 1, \dots, m$ die Funktion³¹³² $f \cdot \chi_{Z_i}: M \rightarrow \mathbb{R}$ integrierbar ist. Wenn dies der Fall ist, dann definieren wir das Integral von $f: M \rightarrow \mathbb{R}$ als

$$\int_M f := \sum_{i=1}^m \underbrace{\int_M \chi_{Z_i} \cdot f}_{\text{definiert in Kapitel 3.4}}.$$

Wir müssen nun noch zeigen, dass diese Definition unabhängig von der Wahl eines endlichen Atlanten und der Wahl der untergeordneten messbaren Zerlegung ist. Dies erfolgt im folgenden Lemma.

Lemma 3.18. *Es sei $f: M \rightarrow \mathbb{R}$ eine Funktion auf einer Untermannigfaltigkeit M . Außerdem seien $\mathcal{A} = \{\Phi_i: U_i \rightarrow V_i \mid i = 1, \dots, m\}$ sowie $\mathcal{A}' = \{\Phi'_i: U'_i \rightarrow V'_i \mid i = 1, \dots, m'\}$ zwei endliche Atlanten und es es seien $\mathcal{Z} = \{Z_1, \dots, Z_m\}$ sowie $\mathcal{Z}' = \{Z'_1, \dots, Z'_{m'}\}$ den Atlanten $\mathcal{A}$ und $\mathcal{A}'$ untergeordnete messbare Zerlegungen. Dann gilt*

$$\begin{array}{ccc} \text{für alle } i = 1, \dots, m \text{ ist die Funktion} & \iff & \text{für alle } i = 1, \dots, m' \text{ ist die Funktion} \\ f \cdot \chi_{Z_i}: M \rightarrow \mathbb{R} \text{ integrierbar} & & f \cdot \chi_{Z'_i}: M \rightarrow \mathbb{R} \text{ integrierbar.} \end{array}$$

Wenn dies der Fall ist, dann gilt zudem

$$\sum_{i=1}^m \int_M \chi_{Z_i} \cdot f = \sum_{i=1}^{m'} \int_M \chi_{Z'_i} \cdot f.$$

Beweis. In dem wir zu den Schnittmengen $U_i \cap U'_j$ und $Z_i \cap Z'_j$ übergehen können wir o.B.d.A. annehmen, dass $m = m'$ und $Z_i = Z'_j$ für $i = 1, \dots, m$. Die Aussage des Lemmas folgt nun aus der schon bewiesenen Wohldefiniertheit von Integrierbarkeit und Integral von Funktionen, welche “durch eine Karte abgedeckt” werden. $\square$

³¹Zur Erinnerung, für eine Teilmenge $A \subset \mathbb{R}^n$ ist die charakteristische Funktion $\chi_A: \mathbb{R}^n \rightarrow \mathbb{R}$ gegeben durch $\chi_A(x) = 1$, falls $x \in A$ und $\chi_A(x) = 0$, falls $x \notin A$.

³²Die Funktion $f \cdot \chi_{Z_i}$ hat die Eigenschaft, dass es eine Karte, nämlich $\Phi_i: U_i \rightarrow V_i$ gibt, so dass sie außerhalb von $U_i \cap M$ verschwindet. Also ist die Integrierbarkeit dieser Funktion schon in Kapitel 3.4 definiert.

3.6. Beispiele. In diesem Kapitel wollen wir explizite Rechenformeln angeben um Integrale von stetigen Funktionen auf S^1 und S^2 zu bestimmen.

Satz 3.19. *Es sei $f: S^1 \rightarrow \mathbb{R}$ eine stetige Funktion. Dann gilt*

$$\int_{S^1} f = \int_{\varphi=0}^{\varphi=2\pi} f(\cos \varphi, \sin \varphi) d\varphi.$$

Beweis. Wir betrachten den Atlas von S^1 , welchen wir auf Seite 18 eingeführt hatten. Wir betrachten zuerst die Diffeomorphismen

$$\begin{aligned} \Psi: V := (-\pi, \pi) \times (-1, 1) &\rightarrow U := \Psi(V) \subset \mathbb{R}^2 \\ (\varphi, s) &\mapsto \begin{pmatrix} (1+s) \cos \varphi \\ (1+s) \sin \varphi \end{pmatrix}. \end{aligned}$$

und

$$\begin{aligned} \Psi': V' := (0, 2\pi) \times (-1, 1) &\rightarrow U' := \Psi'(V') \subset \mathbb{R}^2 \\ (\varphi, s) &\mapsto \begin{pmatrix} (1+s) \cos \varphi \\ (1+s) \sin \varphi \end{pmatrix}. \end{aligned}$$

Ein Atlas von S^1 ist dann gegeben durch die beiden Karten $\Phi := \Psi^{-1}: U \rightarrow V$ und $\Phi' := \Psi'^{-1}: U' \rightarrow V'$. In diesem Fall ist

$$\begin{aligned} \text{Gram}_{\Phi}(\varphi, s) &= \text{Gram}\left(1. \text{ Spalte von } D_{(\varphi, s)} \underbrace{\Phi^{-1}}_{=\Psi}\right) \\ &= \text{Gram}\left(1. \text{ Spalte von } \underbrace{\begin{pmatrix} -(1+s) \sin \varphi & \cos \varphi \\ (1+s) \cos \varphi & \sin \varphi \end{pmatrix}}_{=D_{(\varphi, s)} \Psi}\right) = \text{Gram}\left(\begin{pmatrix} -(1+s) \sin \varphi \\ (1+s) \cos \varphi \end{pmatrix}\right) = (1+s)^2. \end{aligned}$$

Natürlich gilt das gleiche Ergebnis auch für $\text{Gram}_{\Phi'}(\varphi, s)$. Nach dieser Nebenrechnung wenden wir uns jetzt der eigentlichen Berechnung zu. Als dem Atlas untergeordnete messbare Zerlegung von M wählen wir $Z = \{1\}$ und $Z' = U' \cap S^1 = S^1 \setminus \{1\}$. Dann ist

$$\begin{aligned} \int_{S^1} f &= \int_{S^1} \chi_Z \cdot f + \int_{S^1} \chi_{Z'} \cdot f \\ &\stackrel{\text{per Definition}}{=} \underbrace{\int_{\varphi \in (-\pi, \pi)} \underbrace{(\chi_Z \cdot f) \circ \Psi}_{=0 \text{ für } \varphi \neq 0} \cdot \underbrace{\sqrt{\text{Gram}_{\Phi}(\varphi, 0)}}_{=1} d\varphi}_{=0} + \int_{\varphi \in (0, 2\pi)} \underbrace{(\chi_{Z'} \cdot f) \circ \Psi'}_{=f \circ \Psi'} \cdot \underbrace{\sqrt{\text{Gram}_{\Phi'}(\varphi, 0)}}_{=1} d\varphi \\ &\stackrel{\text{per Definition, denn } V \cap E_1 = (-\pi, \pi) \text{ und } V' \cap E_1 = (0, 2\pi)}{=} \int_{\varphi \in (0, 2\pi)} (f \circ \Psi')(\varphi, 0) d\varphi = \int_{\varphi \in [0, 2\pi]} f(\cos \varphi, \sin \varphi) d\varphi = \int_{\varphi=0}^{\varphi=2\pi} f(\cos \varphi, \sin \varphi) d\varphi. \end{aligned}$$

$\uparrow$ Satz 3.6 $\uparrow$ Satz 3.4

□

Satz 3.20. *Es sei $f: S^2 \rightarrow \mathbb{R}$ eine stetige Funktion. Dann gilt*

$$\int_{S^2} f = \int_{\varphi=0}^{\varphi=2\pi} \int_{\theta=0}^{\theta=\pi} f(\cos \varphi \cdot \sin \theta, \sin \varphi \cdot \sin \theta, \cos \theta) \cdot \sin \theta \, d\theta \, d\varphi.$$

Beweis. Wir verwenden den Atlas für S^2 , welchen wir auf Seite 19 angegeben haben. Hierbei betrachten wir die Abbildung

$$\begin{aligned} \Psi: \mathbb{R}^3 &\rightarrow \mathbb{R}^3 \\ (\varphi, \theta, s) &\mapsto \begin{pmatrix} (1+s) \cos \varphi \cdot \sin \theta \\ (1+s) \sin \varphi \cdot \sin \theta \\ (1+s) \cos \theta \end{pmatrix}. \end{aligned}$$

Die Karten sind dann durch Umkehrabbildungen von Einschränkungen von Ψ gegeben. Wir berechnen zuerst

$$D_{(\varphi, \theta, s)} \Psi = \begin{pmatrix} -(1+s) \sin \varphi \cdot \sin \theta & (1+s) \cos \varphi \cdot \cos \theta & \cos \varphi \cdot \sin \theta \\ (1+s) \cos \varphi \cdot \sin \theta & (1+s) \sin \varphi \cdot \cos \theta & \sin \varphi \cdot \sin \theta \\ 0 & -(1+s) \sin \theta & \cos \theta \end{pmatrix}.$$

In diesem Fall gilt nun für jede Karte Φ , welche durch Invertieren von Ψ gegeben ist, dass

$$\begin{aligned} \text{Gram}_{\Phi}(\varphi, \theta, s) &= \text{Gram}(\text{erste zwei Spalten von } D_{(\varphi, \theta, s)} \Psi) \\ &= \text{Gram} \begin{pmatrix} -(1+s) \sin \varphi \cdot \sin \theta & (1+s) \cos \varphi \cdot \cos \theta \\ (1+s) \cos \varphi \cdot \sin \theta & (1+s) \sin \varphi \cdot \cos \theta \\ 0 & -(1+s) \sin \theta \end{pmatrix} \\ &= \det \begin{pmatrix} (1+s)^2 \sin^2 \theta & 0 \\ 0 & (1+s)^2 \end{pmatrix} = (1+s)^4 \sin^2 \theta. \\ &\quad \uparrow \\ &\text{denn } \text{Gram}(v_1 \dots v_k) = \det((v_i \cdot v_j)_{i,j=1,\dots,k}) \end{aligned}$$

Der Beweis ist nun ganz analog zum Beweis von Satz 3.20. □

3.7. Integration auf Untermannigfaltigkeiten III. Wenn M eine Untermannigfaltigkeit ist, welche einen endlichen Atlas besitzt, dann hatten wir gerade in Kapitel 3.5 für eine beliebige Funktion $f: M \rightarrow \mathbb{R}$ die Integrierbarkeit und das Integral eingeführt. Allerdings besitzt nicht jede Untermannigfaltigkeit einen endlichen Atlas. Ein Beispiel einer solchen Untermannigfaltigkeit ist gegeben durch die Fläche mit unendlich vielen “Löchern”, siehe Abbildung 21. ³³³⁴

Wir beginnen mit folgendem Lemma.

Lemma 3.21. *Jede Untermannigfaltigkeit besitzt einen abzählbaren Atlas.*

³³Es ist natürlich nicht klar, warum die abgebildete Teilmenge von $\mathbb{R}^3$ eine Untermannigfaltigkeit ist, und warum diese keinen endlichen Atlas besitzt. Beide Aussagen werden wir (eventuell) in der Vorlesung Algebraische Topologie genauer behandeln.

³⁴Ein noch einfacheres Beispiel ist gegeben durch $M = \mathbb{Z} \subset \mathbb{R}$. Dies ist eine 0-dimensionale Untermannigfaltigkeit, welche ebenfalls keinen endlichen Atlas besitzt (warum?).

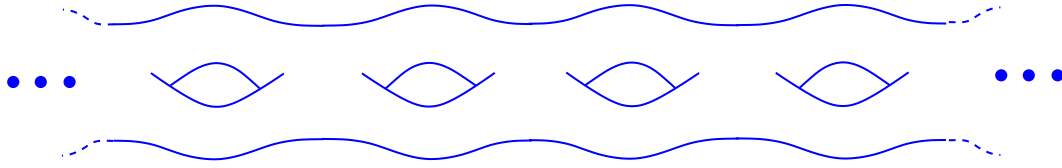


ABBILDUNG 21. Eine Untermannigfaltigkeit, welche keinen endlichen Atlas besitzt.

Beweis. Es sei also M eine Untermannigfaltigkeit von $\mathbb{R}^n$. Zu jedem $x \in M$ gibt es eine Karte $\Phi_x: U_x \rightarrow V_x$. Wir definieren nun einen *rationalen Ball* als eine Teilmenge von $\mathbb{R}^n$ der Form $B_r(y)$, wobei $y \in \mathbb{Q}^n$ und ein $r \in \mathbb{Q}_{>0}$. Zu jedem $x \in M$ wählen wir einen rationalen Ball³⁵ W_x , so dass $x \in W_x$ und $W_x \subset U_x$. Dann ist auch $\{\Phi_x: W_x \rightarrow \Phi_x(W_x)\}_{x \in M}$ ein Atlas für M . Nachdem es nur abzählbar viele rationale Bälle gibt, genügen nun schon abzählbar viele dieser Karten um M zu überdecken, d.h. M besitzt einen abzählbaren Atlas. $\square$

Der Beweis von folgendem Satz ist fast identisch zum Beweis von Satz 3.22.

Satz 3.22. *Es sei M eine k -dimensionale Untermannigfaltigkeit. Dann existiert zu jedem abzählbaren Atlas $\mathcal{A}$ eine dem Atlas $\mathcal{A}$ untergeordnete messbare Zerlegung von M .*

Es sei M eine k -dimensionale Untermannigfaltigkeit und es sei $\mathcal{A} = \{\Phi_i: U_i \rightarrow V_i\}_{i \in I}$ ein abzählbarer Atlas. Zudem sei $f: M \rightarrow \mathbb{R}$ eine beliebige Funktion. Wir wählen eine dem Atlas $\mathcal{A}$ untergeordnete messbare Zerlegung $\mathcal{Z} = \{Z_i\}_{i \in I}$ von M . Wir sagen f ist integrierbar, wenn folgende zwei Bedingungen erfüllt sind:

- (1) für alle $i \in I$ ist die Funktion $f \cdot \chi_{Z_i}: M \rightarrow \mathbb{R}$ integrierbar,
- (2) die Reihe³⁶

$$\sum_{i \in I} \left| \int_M \chi_{Z_i} \cdot f \right|$$

über die *Absolutbeträge* der einzelnen Integrale konvergiert.

³⁵Warum existiert solch ein rationaler Ball?

³⁶Es sei I eine abzählbare unendliche Menge und für jedes $i \in I$ sei $a_i \in \mathbb{R}_{\geq 0}$ gewählt. Wir wählen eine Bijektion $\varphi: \mathbb{N} \rightarrow I$ und wir setzen

$$\sum_{i \in I} a_i := \sum_{k=1}^{\infty} a_{\varphi(k)}.$$

Wir müssen zeigen, dass diese Definition nicht von der Wahl von φ abhängt. Dies folgt sofort aus dem Umordnungssatz 5.14 aus der Analysis I. Dieser besagt, dass für eine Folge $b_k, k \in \mathbb{N}$ und eine Bijektion $\psi: \mathbb{N} \rightarrow \mathbb{N}$ gilt, dass

$$\sum_{k=1}^{\infty} b_k \text{ konvergiert absolut} \implies \sum_{k=1}^{\infty} b_{\psi(k)} \text{ konvergiert absolut und } \sum_{k=1}^{\infty} b_k = \sum_{k=1}^{\infty} b_{\psi(k)}.$$

Wenn dies der Fall ist, dann definieren wir das Integral von $f: M \rightarrow \mathbb{R}$ als ³⁷

$$\int_M f := \sum_{i \in I} \int_M \chi_{Z_i} \cdot f.$$

Wir müssen nun natürlich wiederum zeigen, dass diese Definition unabhängig ist von der Wahl eines endlichen Atlanten und der untergeordneten messbaren Zerlegung. Der Beweis dazu wird beispielsweise in Kapitel 11.5. von

Königsberger: Analysis 2 (Springer-Verlag)

ausgeführt.

Zum Abschluß betrachten wir noch die beiden Extremfälle, dass $M \subset \mathbb{R}^n$ eine Untermannigfaltigkeit der Dimension 0 beziehungsweise eine Untermannigfaltigkeit der Dimension n ist. Wenn M eine 0-dimensionale Untermannigfaltigkeit ist, welche durch endlich viele Karten abgedeckt wird, dann besteht M aus endlich vielen Punkten. Für eine Funktion $f: M \rightarrow \mathbb{R}$ gilt dann

$$\int_M f = \sum_{P \in M} f(P),$$

d.h. das Integral ist die Summe der Funktionswerte.³⁸ Wenn M eine n -dimensionale Untermannigfaltigkeit von $\mathbb{R}^n$ ist, dann ist die Identitätsabbildung schon eine Karte. Offensichtlich ist $\text{Gram}_{\text{id}}(x) = 1$ für jedes $x \in M$. Also folgt sofort aus den Definitionen, dass

$$\underbrace{\int_M f}_{\text{Integral auf Untermannigfaltigkeit}} = \underbrace{\int_M f \, dv}_{\text{Integral auf Teilmenge von } \mathbb{R}^n}.$$

3.8. Das k -dimensionale Volumen.

Definition. Es sei M eine k -dimensionale Untermannigfaltigkeit. Wir sagen $A \subset M$ ist messbar mit endlichen Volumen, wenn die Funktion $\chi_A: M \rightarrow \mathbb{R}$ integrierbar ist. In diesem Fall bezeichnen wir dann

$$\text{Vol}_k(A) := \int_M \chi_A$$

als das *k -dimensionale Volumen von A* . Im Fall $k = 2$ bezeichnen wir $\text{Vol}_k(A)$ auch als den *Flächeninhalt von A* . Wenn $\text{Vol}_k(A) = 0$, dann bezeichnen wir A als *Nullmenge in M* .

³⁷Es sei I eine abzählbare unendliche Menge und für jedes $i \in I$ sei $a_i \in \mathbb{R}$ gewählt. Wir nehmen an, dass die Reihe $\sum_{i \in I} |a_i|$ konvergiert. Wir wählen eine Bijektion $\varphi: \mathbb{N} \rightarrow I$ und wir setzen

$$\sum_{i \in I} a_i := \sum_{k=1}^{\infty} a_{\varphi(k)}.$$

Es folgt wiederum aus dem Umordnungssatz 5.14 aus der Analysis I, siehe Fußnote 36, dass diese Definition nicht von der Wahl von φ abhängt.

³⁸Dies folgt aus der Tatsache, dass für eine Funktion $f: \mathbb{R}^0 = \{0\} \rightarrow \mathbb{R}$ das Lebesgue-Integral gegeben ist durch $f(0)$.

Beispiele. (1) Wir wollen nun den Flächeninhalt von S^2 bestimmen. Dieser ist gegeben durch

$$\begin{aligned} \text{Flächeninhalt}(S^2) = \int_{S^2} 1 \, d\mu &= \int_{\varphi=0}^{\varphi=2\pi} \int_{\theta=0}^{\theta=\pi} \sin \theta \, d\varphi \, d\theta = \int_{\varphi=0}^{\varphi=2\pi} \int_{\theta=0}^{\theta=\pi} \sin \theta \, d\theta \, d\varphi = \int_{\varphi=0}^{\varphi=2\pi} 2 \, d\varphi = 4\pi. \end{aligned}$$

$\uparrow$
 Satz 3.20

(2) Es seien $v_1, \dots, v_k \in \mathbb{R}^n$ linear unabhängige Vektoren. Wir bezeichnen mit

$$M(v_1, \dots, v_k) = \left\{ \sum_{i=1}^k l_i v_i \mid l_1, \dots, l_k \in (0, 1) \right\}$$

das von $v_1, \dots, v_k$ aufgespannte *offene Parallelotop*. Wir werden in Übungsblatt 3 sehen, dass dies eine k -dimensionale Untermannigfaltigkeit ist, und dass das k -dimensionale Volumen gerade dem k -dimensionalen Volumen vom entsprechenden geschlossenen Parallelotop entspricht, welches wir auf Seite 35 eingeführt hatten.

Es gelten die offensichtlichen Verallgemeinerungen der Aussagen aus der Analysis III über Nullmengen. Beispielsweise ist die abzählbare Vereinigung von Nullmengen wiederum eine Nullmenge. Zudem gilt folgendes Lemma.

Lemma 3.23. *Es sei M eine k -dimensionale Untermannigfaltigkeit. Wenn zwei Funktionen $f, g: M \rightarrow \mathbb{R}$ außerhalb einer Nullmenge übereinstimmen, dann ist f integrierbar genau dann, wenn g integrierbar ist. Im integrierbaren Fall gilt zudem*

$$\int_M f = \int_M g.$$

4. UNTERMANNIGFALTIGKEITEN MIT RAND

Wir erweitern jetzt den Begriff von Untermannigfaltigkeiten etwas:

Definition.

- (1) Für $k \leq n$ bezeichnen wir

$$H_k := \{(x_1, \dots, x_k, 0, \dots, 0) \in \mathbb{R}^n \mid x_1, \dots, x_k \in \mathbb{R} \text{ und } x_k \geq 0\},$$

als *k-dimensionale Halbebene* und wir schreiben

$$\partial H_k := \{(x_1, \dots, x_k, 0, \dots, 0) \in \mathbb{R}^n \mid x_1, \dots, x_k \in \mathbb{R} \text{ und } x_k = 0\}.$$

- (2) Eine Teilmenge $M \subset \mathbb{R}^n$ heißt *k-dimensionale Untermannigfaltigkeit*, wenn es zu jedem Punkt $a \in M$ eine *Karte um a* gibt, d.h. einen Diffeomorphismus $\Phi: U \rightarrow V$ zwischen einer offenen Umgebung $U \subset \mathbb{R}^n$ von a und einer offenen Menge $V \subset \mathbb{R}^n$, so dass entweder
- (i) $\Phi(M \cap U) = E_k \cap V$, oder
 - (ii) $\Phi(M \cap U) = H_k \cap V$ mit $\Phi(a) \in \partial H_k$.
- (3) Wir sagen $a \in M$ ist ein *innerer Punkt*, wenn es eine Karte $\Phi: U \rightarrow V$ um a vom Typ (i) gibt. Wir sagen $a \in M$ ist ein *Randpunkt*, wenn es eine Karte $\Phi: U \rightarrow V$ um a vom Typ (ii) gibt. Wir bezeichnen mit ∂M die Menge aller Randpunkte und wir nennen ∂M den *Rand von M*.
- (4) Ein *Atlas von M* ist eine Familie $\{\Phi_i: U_i \rightarrow V_i\}_{i \in I}$ von Karten mit $M \subset \bigcup_{i \in I} U_i$.

Eine Karte von Typ (i) ist also der gleiche Typ von Karte, wie wir in schon auf Seite 17 eingeführt hatten. Die Karten von Typ (ii) sind neu.

Satz 4.1. *Jeder Punkt auf einer Untermannigfaltigkeit ist entweder ein Randpunkt oder ein innerer Punkt.*

Beweis. Es sei $M \subset \mathbb{R}^n$ eine *k-dimensionale Untermannigfaltigkeit*. Nehmen wir also an, dass es einen Punkt $a \in M$ gibt, welcher sowohl ein innerer Punkt als auch ein Randpunkt ist. Dann gibt es also eine Karte $\Phi_1: U_1 \rightarrow V_1$ um a von Typ (i) und es gibt eine Karte $\Phi_2: U_2 \rightarrow V_2$ um a von Typ (ii). Indem wir zu $U_1 \cap U_2$ übergehen können wir o.B.d.A. annehmen, dass $U_1 = U_2 =: U$.

Wir schreiben $N_1 := \Phi_1(M \cap U) \subset E_k = \mathbb{R}^k$ und $N_2 := \Phi_2(M \cap U) \subset H_k \subset E_k = \mathbb{R}^k$. (Alle diese Objekte werden in Abbildung 23 skizziert.) Wir bezeichnen mit $\Psi: N_1 \rightarrow N_2$ die Einschränkung der Abbildung $\Phi_2 \circ \Phi_1^{-1}$ auf die offene Teilmenge N_1 . Dies ist eine glatte, bijektive Abbildung. Zudem ist an jedem Punkt $P \in N_1$ das Differential $D_P \Psi$ invertierbar.³⁹ Es folgt nun aus Korollar 2.2, dass $\Psi: N_1 \rightarrow N_2$ ein Diffeomorphismus ist. Insbesondere ist N_2 eine offene Teilmenge von $\mathbb{R}^k$. Aber dies ist nicht der Fall, denn $N_2 \subset H_k$ und $\Phi_2(a)$ liegt in ∂H_k , besitzt also keine offene Umgebung in N_2 . Wir haben damit einen Widerspruch erhalten. $\square$

³⁹In der Tat, denn Ψ ist die Einschränkung von dem Diffeomorphismus $\Phi_2 \circ \Phi_1^{-1}: V_1 \rightarrow V_2$ auf die Teilmenge $N_1 = V_1 \cap E_k$. Nachdem $D_P(\Phi_2 \circ \Phi_1^{-1})$ invertierbar ist, ist auch die Einschränkung auf jeden Teilraum, in diesem Fall E_k , weiterhin invertierbar.

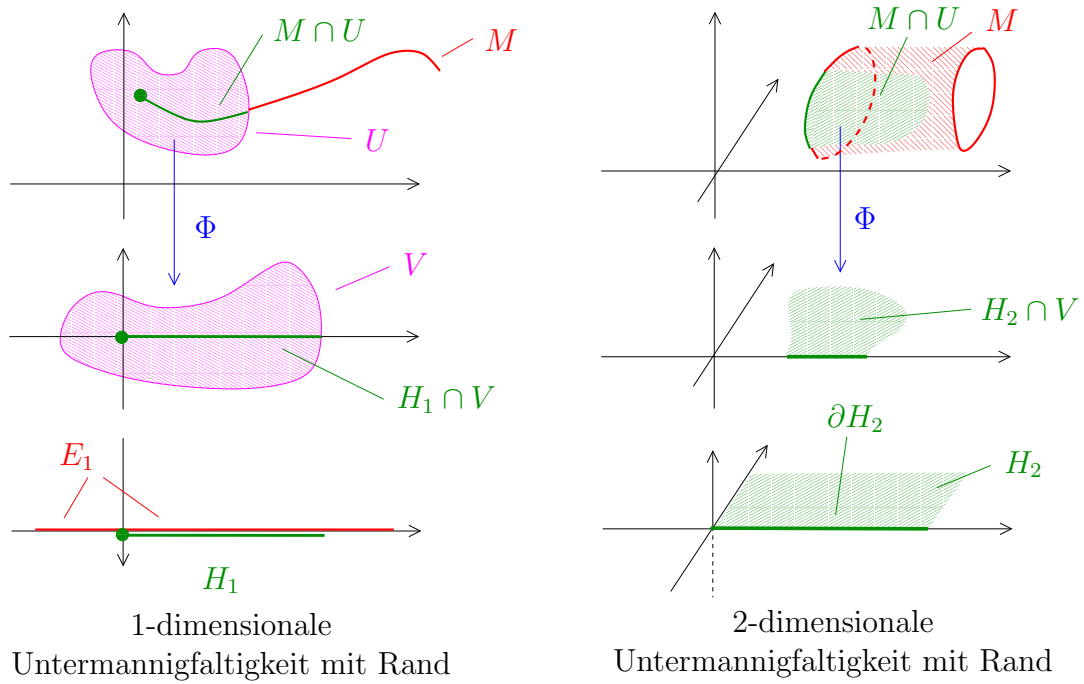


ABBILDUNG 22. Schematisches Bild von Untermannigfaltigkeiten mit Rand.

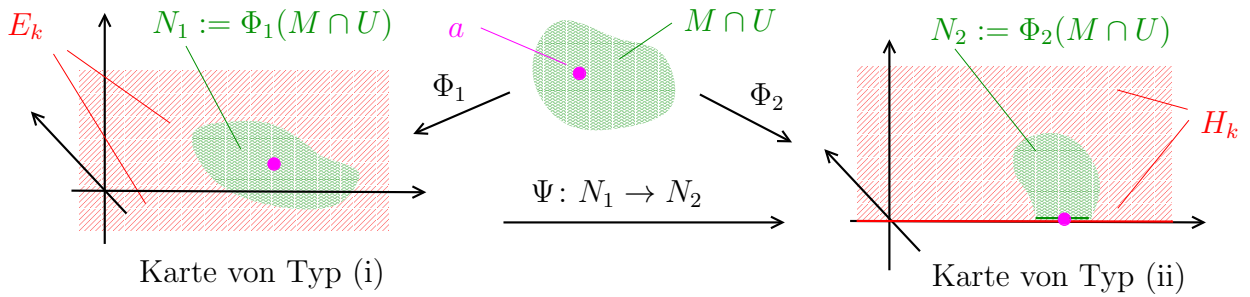


ABBILDUNG 23. Skizze zum Beweis von Satz 4.1.

Bemerkung. Jede Untermannigfaltigkeit im Sinne der Definition auf Seite 17 ist auch eine Untermannigfaltigkeit im obigen Sinne. Es folgt sofort aus Satz 4.1, dass in diesem Fall $\partial M = \emptyset$.

Beispiele.

- (1) Wir betrachten die Halbsphäre

$$M := \{(x, y, z) \in \mathbb{R}^3 \mid x^2 + y^2 + z^2 = 1 \text{ und } z \geq 0\}.$$

Wir wollen zeigen, dass dies eine Untermannigfaltigkeit ist mit Rand

$$\partial M := \{(x, y, z) \in \mathbb{R}^3 \mid x^2 + y^2 + z^2 = 1 \text{ und } z = 0\}.$$

Geeignete Einschränkungen der Karten auf Seite 19⁴⁰ geben uns eine Karte von Typ (i) für jeden Punkt $(x, y, z) \in M$ mit $z > 0$. Wir müssen nun noch Karten für die Punkte auf der xy -Ebene finden. Wir betrachten, ähnlich wie auf Seite 19, den Diffeomorphismus

$$\begin{aligned} \Psi: V := (-\pi, \pi) \times \left(-\frac{\pi}{2}, \frac{\pi}{2}\right) \times (-1, 1) &\rightarrow U := \Psi(V) \\ (\varphi, \theta, s) &\mapsto \begin{pmatrix} (1+s) \cos \varphi \cdot \sin\left(\frac{\pi}{2} - \theta\right) \\ (1+s) \sin \varphi \cdot \sin\left(\frac{\pi}{2} - \theta\right) \\ (1+s) \cos\left(\frac{\pi}{2} - \theta\right) \end{pmatrix}. \end{aligned}$$

Hierbei gilt

$$\begin{aligned} \theta \geq 0 &\iff z\text{-Koordinate von } \Psi(\varphi, \theta, s) \text{ ist } \geq 0 \\ \theta = 0 &\iff z\text{-Koordinate von } \Psi(\varphi, \theta, s) \text{ ist } = 0. \end{aligned}$$

Man sieht nun leicht, dass $\Psi(H_2 \cap V) = M \cap U$, d.h. $\Phi := \Psi^{-1}: U \rightarrow V$ ist eine Karte von Typ (ii). Diese Karte deckt alle Punkte $(x, y, 0) \neq (-1, 0, 0) \in M$ ab. Indem wir Φ geeignet spiegeln erhalten wir auch noch eine Karte für den Punkt $(-1, 0, 0)$. Anhand der Karten sieht man leicht, dass ∂M , wie oben behauptet, der “Randkreis” ist.

- (2) Man kann leicht zeigen, dass der “abgeschlossene” Zylinder

$$Z = \{(x, y, z) \mid x^2 + y^2 = 1 \text{ und } z \in [-1, 1]\}$$

eine 2-dimensionale Untermannigfaltigkeit ist mit Rand

$$\partial Z = \{(x, y, z) \mid x^2 + y^2 = 1 \text{ und } z = \pm 1\}.$$

Zudem ist der “halboffene” Zylinder

$$Z' = \{(x, y, z) \mid x^2 + y^2 = 1 \text{ und } z \in (-1, 1]\}$$

eine 2-dimensionale Untermannigfaltigkeit mit Rand

$$\partial Z' = \{(x, y, z) \mid x^2 + y^2 = 1 \text{ und } z = 1\}.$$

- (3) Die Halbebene

$$H_k = \{(x_1, \dots, x_k, 0, \dots, 0) \in \mathbb{R}^n \mid x_1, \dots, x_k \in \mathbb{R} \text{ und } x_k \geq 0\},$$

ist eine k -dimensionale Untermannigfaltigkeit mit Rand

$$\partial H_k = \{(x_1, \dots, x_k, 0, \dots, 0) \in \mathbb{R}^n \mid x_1, \dots, x_k \in \mathbb{R} \text{ und } x_k = 0\}.$$

⁴⁰Oder, alternativ eine leichte Abänderung der Karten auf Seite 23.

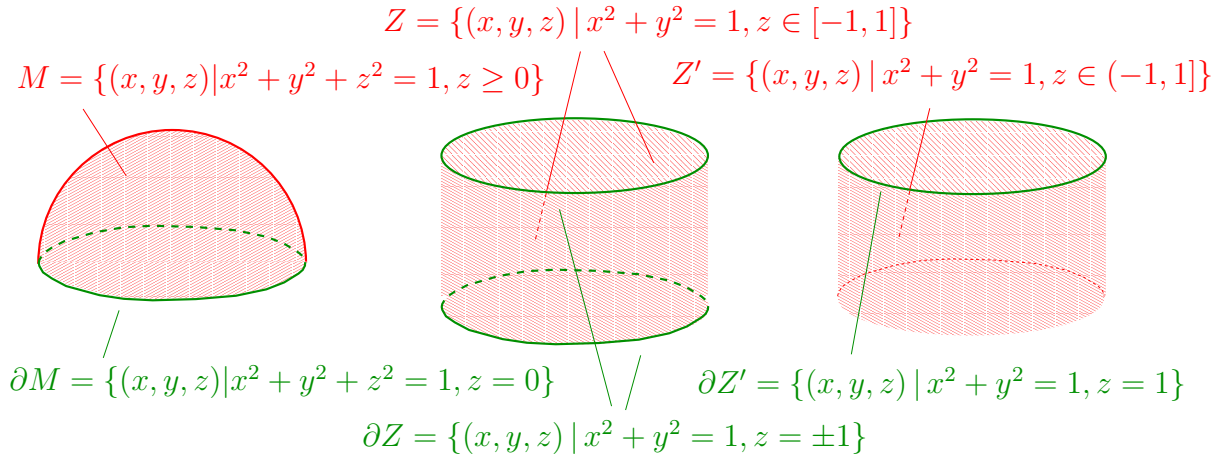


ABBILDUNG 24. Beispiele von Untermannigfaltigkeiten mit nichtleerem Rand.

Bemerkung. Unglücklicherweise gibt es nun zwei verschiedene Definitionen von “Rand”. Zur Erinnerung, in Analysis III hatten wir folgende Definition eingeführt: Es sei $A \subset \mathbb{R}^n$ eine beliebige Teilmenge, dann hatten wir den “topologischen Rand” von A in $\mathbb{R}^n$ definiert als

$$\partial A := \left\{ x \in \mathbb{R}^n \mid \begin{array}{l} \text{jede offene Umgebung von } x \text{ enthält einen Punkt in } A \\ \text{und außerdem einen Punkt außerhalb von } A \end{array} \right\}.$$

Der topologische Rand der offenen Sphäre

$$X := \{x \in \mathbb{R}^n \mid \|x\| < 1\}$$

ist also die Einheitsphäre. Der Rand von X aufgefasst als Untermannigfaltigkeit ist hingegen die leere Menge. Ein noch extremeres Beispiel ist gegeben durch

$$Y := \{(x, y) \in \mathbb{R}^2 \mid x^2 + y^2 = 1\} \subset \mathbb{R}^2.$$

In diesem Fall ist der topologische Rand von $Y \subset \mathbb{R}^2$ ganz Y , aber der Rand von Y aufgefasst als Untermannigfaltigkeit ist wiederum die leere Menge. Wenn wir ab jetzt von “Rand” reden, meinen wir immer den “Untermannigfaltigkeitsrand”.

Man kann den Satz 2.6 vom regulären Wert aus der Analysis II ohne größere Probleme verallgemeinern, und erhält dann folgenden Satz.

Satz 4.2. *Es sei $U \subset \mathbb{R}^n$ offen, es sei $f: U \rightarrow \mathbb{R}$ eine glatte Abbildung und es seien $z_1, z_2 \in \mathbb{R}$ reguläre Werte. Dann ist $M := f^{-1}([z_1, z_2]) \subset \mathbb{R}^n$ eine n -dimensionale Untermannigfaltigkeit mit*

$$\partial M = f^{-1}(z_1) \cup f^{-1}(z_2).$$

Beispiel. Wir betrachten die Funktion

$$\begin{aligned} f: \mathbb{R}^n &\rightarrow \mathbb{R} \\ (x_1, \dots, x_n) &\mapsto x_1^2 + \dots + x_n^2. \end{aligned}$$

Dann sind 1 und -1 reguläre Werte⁴¹. Es folgt aus Satz 4.2, dass

$$f^{-1}([-1, 1]) = \{(x_1, \dots, x_n) \in \mathbb{R}^n \mid x_1^2 + \dots + x_n^2 \in [-1, 1]\} = B^n$$

eine n -dimensionale Untermannigfaltigkeit ist mit Rand

$$f^{-1}(-1) \cup f^{-1}(1) = \emptyset \cup \{(x_1, \dots, x_n) \in \mathbb{R}^n \mid x_1^2 + \dots + x_n^2 = 1\} = S^{n-1}.$$

Satz 4.3. *Es sei M eine k -dimensionale Untermannigfaltigkeit mit Rand. Dann ist ∂M eine $(k-1)$ -dimensionale Untermannigfaltigkeit mit leerem Rand.*

Beweis. Es sei M eine k -dimensionale Untermannigfaltigkeit mit Rand und es sei zudem $a \in \partial M$. Dann existiert eine Karte um a vom Typ (ii), d.h. ein Diffeomorphismus $\Phi: U \rightarrow V$ zwischen einer offenen Umgebung $U \subset \mathbb{R}^n$ von a und einer offenen Menge $V \subset \mathbb{R}^n$, so dass

$$\Phi(M \cap U) = H_k \cap V,$$

wobei $\Phi(a) \in \partial H_k$. Es folgt sofort aus Satz 4.1, dass $U \cap \partial M = \Phi^{-1}(\partial H_k)$. Nachdem $\partial H_k = E_{k-1}$ sehen wir also, dass Φ auch eine Karte vom Typ (i) für ∂M um den Punkt a ist. Wir haben also bewiesen, dass ∂M eine $(k-1)$ -dimensionale Untermannigfaltigkeit ist, so dass alle Punkte eine Karte vom Typ (i) besitzen. D.h. die $(k-1)$ -dimensionale Untermannigfaltigkeit ∂M besitzt keine Randpunkte. $\square$

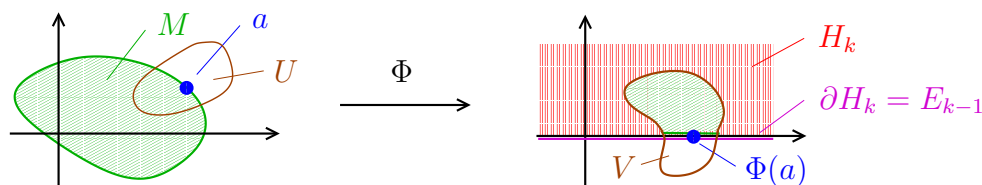


ABBILDUNG 25. Illustration zum Beweis von Satz 4.3.

Im Folgenden verwenden wir folgende Sprechweisen:

- (1) “Untermannigfaltigkeit” kann heißen, dass der Rand leer oder auch nicht leer ist.
- (2) “Untermannigfaltigkeit mit Rand” bedeutet eine Untermannigfaltigkeit M mit nicht-leerem Rand, d.h. mit $\partial M \neq \emptyset$,
- (3) “geschlossene Untermannigfaltigkeit” bedeutet eine *kompakte* Untermannigfaltigkeit M mit $\partial M = \emptyset$.

Beispiel. Beispielsweise ist die offene Kugel

$$X := \{x \in \mathbb{R}^n \mid \|x\| < 1\}$$

eine Untermannigfaltigkeit mit $\partial X = \emptyset$, aber X ist keine geschlossene Untermannigfaltigkeit, denn X ist nicht kompakt. Andererseits sind beispielsweise S^n und der auf Seite 21 eingeführte Torus geschlossene Untermannigfaltigkeiten.

⁴¹Warum ist -1 ein regulärer Wert?

Ein weiteres Beispiel von Untermannigfaltigkeiten mit Rand ist wie folgt gegeben. Es sei $U \subset \mathbb{R}^{n-1}$ eine offene Teilmenge und es sei $f: U \rightarrow \mathbb{R}$ eine glatte Abbildung. Man kann leicht zeigen, dass für jedes $a \in \mathbb{R}$

$$M = \{(x, y) \in \mathbb{R}^{n-1} \times \mathbb{R} = \mathbb{R}^n \mid a < y \leq f(x)\}$$

eine n -dimensionale Untermannigfaltigkeit ist mit

$$\partial M = \{(x, y) \in \mathbb{R}^{n-1} \times \mathbb{R} = \mathbb{R}^n \mid a < y = f(x)\}.$$

Diese Untermannigfaltigkeit wird in Abbildung 26 skizziert. Der folgende Satz besagt, dass

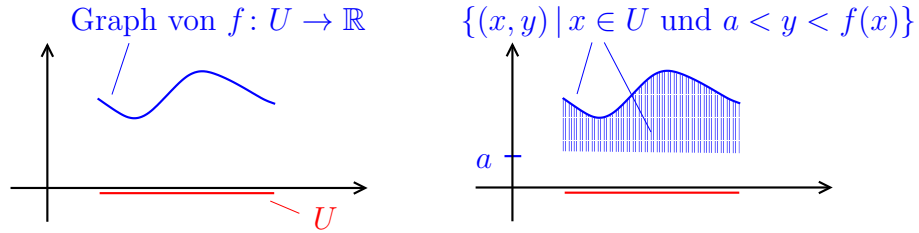


ABBILDUNG 26.

jede n -dimensionale Untermannigfaltigkeit am Rand lokal⁴² von diesem Typ ist. Die Aussage des Satzes besitzt eine nicht zu übersehende Ähnlichkeit mit Satz 2.5.

Satz 4.4. *Es sei $M \subset \mathbb{R}^n$ eine n -dimensionale Untermannigfaltigkeit und es sei $X \in \partial M$. Dann existieren eine Permutationsmatrix P , ein offener Quader $Q \subset \mathbb{R}^{n-1}$, zwei reelle Zahlen $a < b$, sowie eine glatte Funktion $f: Q \rightarrow (a, b)$, so dass $P \cdot X \in Q \times (a, b)$, so dass*

$$P \cdot M \cap (Q \times (a, b)) = \{(x, y) \mid x \in Q \text{ und } a < y \leq f(x)\}.$$

und so dass

$$P \cdot \partial M \cap (Q \times (a, b)) = \{(x, y) \mid x \in Q \text{ und } y = f(x)\}.$$

Der Satz kann in der Tat mithilfe von Satz 2.5 bewiesen werden. Wir überlassen die Durchführung des Beweises als freiwillige Übungsaufgabe.

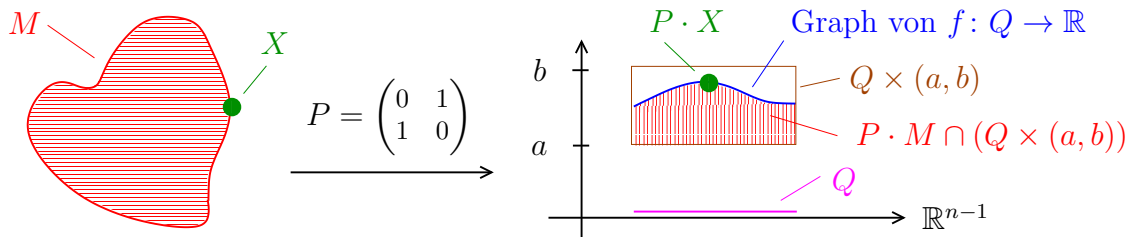


ABBILDUNG 27. Illustration von Satz 4.4.

⁴²Wir sagen ein topologischer Raum X besitzt *lokal* eine Eigenschaft P , wenn es zu jedem $x \in X$ eine offene Umgebung U , so dass U diese Eigenschaft besitzt.

Für eine Funktion $f: M \rightarrow \mathbb{R}$ auf einer Untermannigfaltigkeit mit Rand können wir die Integrierbarkeit und das Integral ganz analog zu Kapitel 3 definieren. Zudem gelten die (hoffentlich) absichtlichen Abwandlungen der Sätze aus Kapitel 3.6.

Beispiel. Wir betrachten die Halbkugel

$$M := \{(x, y, z) \in \mathbb{R}^3 \mid x^2 + y^2 + z^2 = 1 \text{ und } z \geq 0\}.$$

Ganz analog zu Satz 3.20 gilt für eine stetige Funktion $f: M \rightarrow \mathbb{R}$, dass

$$\int_M f = \int_{\varphi=0}^{\varphi=2\pi} \int_{\theta=0}^{\theta=\frac{\pi}{2}} f(\cos \varphi \cdot \sin \theta, \sin \varphi \cdot \sin \theta, \cos \theta) \cdot \sin \theta \, d\theta \, d\varphi.$$

5. ZUSAMMENHÄNGENDE TOPOLOGISCHE RÄUME

In diesem Kapitel erinnern wir an zwei Definitionen, welche wir schon in Analysis III für metrische Räume eingeführt hatten.

Definition. Es sei X ein topologischer Raum. Wir sagen X ist *zusammenhängend*, wenn gilt:

$$\begin{array}{l} X = U \sqcup V, \text{ wobei } U, V \text{ offene} \\ \text{disjunkte Teilmengen von } X \end{array} \implies U = X \text{ oder } V = X.$$

Bemerkung.

- (1) Ein topologischer Raum X ist also *nicht zusammenhängend*, wenn es disjunkte offene Teilmengen $U, V \subset X$ mit $X = U \cup V$ gibt, so dass $U \neq X$ und $V \neq X$.
- (2) Man kann leicht zeigen, dass X *zusammenhängend* ist, genau dann, wenn gilt:

$$\begin{array}{l} X = U \sqcup V, \text{ wobei } U, V \text{ abgeschlossene} \\ \text{disjunkte Teilmengen von } X \end{array} \implies U = \emptyset \text{ oder } V = \emptyset.$$

Beispiele.

- (1) Wir betrachten den topologischen Raum

$$X = [0, 1] \cup [2, 3] \subset \mathbb{R}.$$

Dieser ist nicht zusammenhängend, denn $U = [0, 1]$ und $V = [2, 3]$ sind zwei disjunkte offene Teilmengen von X .

- (2) Es sei X eine Menge. Wenn wir X mit der trivialen Topologie betrachten, dann ist dies ein zusammenhängender topologischer Raum. Wenn wir hingegen X mit der diskreten Topologie betrachten, dann ist dies nur dann ein zusammenhängender Raum, wenn X höchstens ein Element besitzt.
- (3) In Übungsblatt 3 werden wir der Frage nachgehen, ob die Gerade mit zwei Nullen zusammenhängend ist.

Definition. Es sei X ein topologischer Raum.

- (1) Es seien $x, y \in X$. Ein *Weg von x nach y* ist eine stetige Abbildung $\gamma: [a, b] \rightarrow X$ mit $\gamma(a) = x$ und $\gamma(b) = y$.
- (2) Wir sagen X ist *wegzusammenhängend*, wenn es für alle $x, y \in X$ einen Weg von x nach y gibt.

In Lemma 6.2 aus der Analysis III Vorlesung hatten wir folgendes Lemma für metrische Räume bewiesen. Der Beweis für topologische Räume ist wort-wörtlich der gleiche.

Lemma 5.1. *Jeder wegzusammenhängende topologische Raum ist auch zusammenhängend.*

Bemerkung. In Analysis III hatten wir gesehen, dass im Allgemeinen die Umkehrung von Lemma 5.1 nicht gilt, d.h. es gibt topologische Räume, welche zusammenhängend sind, aber nicht wegzusammenhängend sind. Genauer gesagt hatten wir gezeigt, dass

$$X := \{(0, y) \mid y \in [-1, 1]\} \cup \{(x, \sin(\frac{1}{x})) \mid x \in (0, \pi]\} \subset \mathbb{R}^2$$

zusammenhängend, aber nicht wegzusammenhängend ist. Siehe Abbildung 28 für eine Illustration von X .

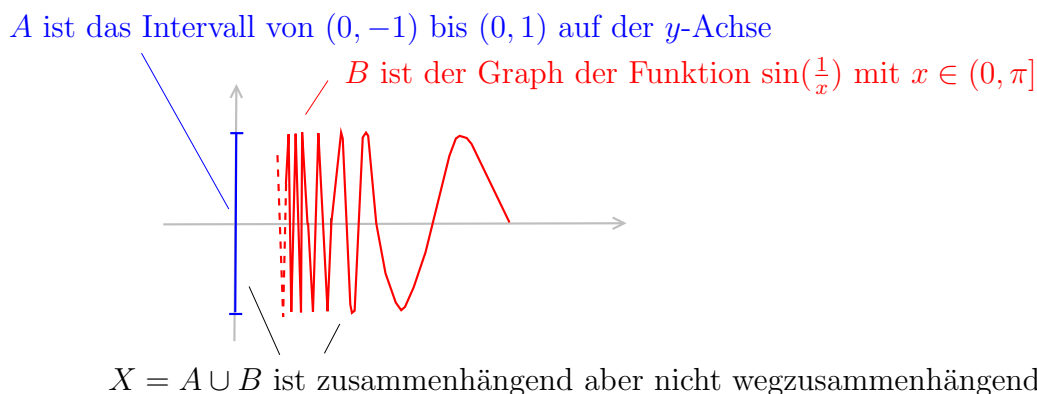


ABBILDUNG 28.

Für Untermannigfaltigkeiten gilt die Umkehrung von Lemma 5.1. Genauer gesagt gilt folgender Satz, welcher in Übungsblatt 3 bewiesen wird.

Satz 5.2. *Es sei $M \subset \mathbb{R}^n$ eine k -dimensionale Untermannigfaltigkeit. Dann gilt:*

$$M \text{ ist zusammenhängend} \iff M \text{ ist wegzusammenhängend.}$$

Folgende Definition hatten wir ganz ähnlich schon in Analysis III kennen gelernt. Es sei X ein topologischer Raum. Für $x, y \in X$ definieren wir

$$x \sim y \iff \text{es gibt einen Weg von } x \text{ nach } y.$$

Es ist leicht zu zeigen, dass dies eine Äquivalenzrelation ist. Wir bezeichnen die Äquivalenzklassen von X als die *Komponenten von X* . Wenn X aus nur einer Äquivalenzklasse besteht, dann ist X per Definition wegzusammenhängend.

Beispiel. Die Komponenten von $[0, 1] \cup [2, 3]$ sind gegeben durch $[0, 1]$ und $[2, 3]$.

Wir erinnern auch noch an den Begriff der diskreten Teilmengen.

Definition. Wir sagen eine Teilmenge D eines topologischen Raums ist *diskret*, wenn jeder Punkt $x \in D$ eine offene Umgebung U besitzt, so dass $U \cap D = \{x\}$.

Beispiel.

- (1) Jede endliche Teilmenge von $\mathbb{R}^n$ ist diskret. In der Tat, denn für eine endliche Teilmenge $A \subset \mathbb{R}^n$ und $x \in A$ setzen wir $\epsilon := \min\{\|x - y\| \mid y \in A \setminus \{x\}\}$. Dann ist $B_\epsilon(x) \cap A = \{x\}$.
- (2) Die Teilmenge $\mathbb{Z} \subset \mathbb{R}$ ist diskret.
- (3) Es sei X eine Menge, aufgefasst als topologischer Raum mit der diskreten Topologie. Dann ist jede Teilmenge diskret.

Wir beschließen das Kapitel mit folgendem elementaren Lemma, welches wir als Lemma 6.5 in Analysis II für metrische Räume bewiesen hatten. Der Beweis für topologische Räume ist auch hier wieder wort-wörtlich der gleiche.

Lemma 5.3. *Es sei $f: X \rightarrow Y$ eine stetige Abbildung zwischen zwei topologischen Räumen. Wenn X zusammenhängend ist und wenn $f(X)$ diskret ist, dann ist f eine konstante Abbildung.*

6. DER GAUSSSCHE INTEGRALSATZ

6.1. Tangentialvektoren zu Untermannigfaltigkeiten. In Analysis II hatten wir schon die Begriffe Tangentialvektoren und Tangentialraum für Untermannigfaltigkeiten ohne Rand eingeführt. Wir erweitern nun die Definition auf beliebige Untermannigfaltigkeiten.

Definition. Es sei $M \subset \mathbb{R}^n$ eine Untermannigfaltigkeit und P ein Punkt auf M .

- (1) Ein *nichttriviales* Intervall ist ein Intervall, welches nicht aus nur einem Punkt besteht.
- (2) Eine *Kurve in M durch P* ist eine stetig differenzierbare⁴³ Abbildung $\gamma: I \rightarrow M$ auf einem nichttrivialen Intervall I mit $0 \in I$ und mit $\gamma(0) = P$.
- (3) Ein Vektor $v \in \mathbb{R}^n$ heißt *Tangentialvektor* an M im Punkt P , wenn es eine Kurve γ in M durch P gibt, so dass $\gamma'(0) = v$.
- (4) Die Menge aller Tangentialvektoren am Punkt P wird mit $T_P M$ bezeichnet. Wir nennen $T_P M$ den *Tangentialraum* am Punkt P .

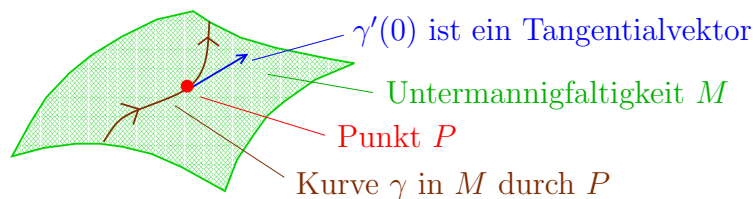


ABBILDUNG 29.

Beispiel. Wir betrachten den Punkt $P = (1, 0, 0)$ auf der Untermannigfaltigkeit

$$M := \{(x, y, z) \in \mathbb{R}^3 \mid x^2 + y^2 + z^2 = 1 \text{ und } z \geq 0\}.$$

Dann ist

$$\begin{aligned} \gamma: [-\frac{\pi}{2}, 0] &\rightarrow M \\ t &\mapsto (-\sin(t), 0, -\cos(t)) \end{aligned}$$

eine Kurve in M durch P . Insbesondere ist $\gamma'(0) = (0, 0, -1)$ ein Tangentialvektor am Punkt P .

Der folgende Satz ist eine einfache Verallgemeinerung von Satz 11.4 aus Analysis II.

Satz 6.1. *Es sei $M \subset \mathbb{R}^n$ eine k -dimensionale Untermannigfaltigkeit. Für jeden Punkt P auf M ist $T_P M$ ein k -dimensionaler Vektorraum.*

⁴³Wie in Analysis I nennen wir eine Abbildung $\gamma: I \rightarrow \mathbb{R}^n$ auf einem Intervall I stetig differenzierbar, wenn γ stetig ist, wenn γ auf den inneren Punkten von I stetig differenzierbar ist und wenn sich die Ableitung γ' zu einer stetigen Funktion auf ganz I fortsetzt. Diese Fortsetzung auf ganz I wird dann auch wieder mit γ' bezeichnet.

Bemerkung. Von der Anschauung her ist $T_P M$ der k -dimensionale Unterraum von $\mathbb{R}^n$, welcher die Untermannigfaltigkeit am Punkt P am besten approximiert.⁴⁴ Der Tangentialraum spielt also für Untermannigfaltigkeiten die Rolle der Tangenten für Graphen in $\mathbb{R}^2$.

Für $v \in \mathbb{R}^n$ schreiben wir im Folgenden

$$v^\perp = \{w \in \mathbb{R}^n \mid w \text{ ist orthogonal zum Vektor } v\}.$$

Allgemeiner schreiben wir für $A \subset \mathbb{R}^n$

$$A^\perp = \{w \in \mathbb{R}^n \mid w \text{ ist orthogonal zu allen } v \in A\}.$$

Definition. Es sei $M \subset \mathbb{R}^n$ eine Untermannigfaltigkeit und P ein Punkt auf M . Ein Vektor $v \in \mathbb{R}^n$ heißt *Normalenvektor* von M in P , wenn $v \in (T_P M)^\perp$.

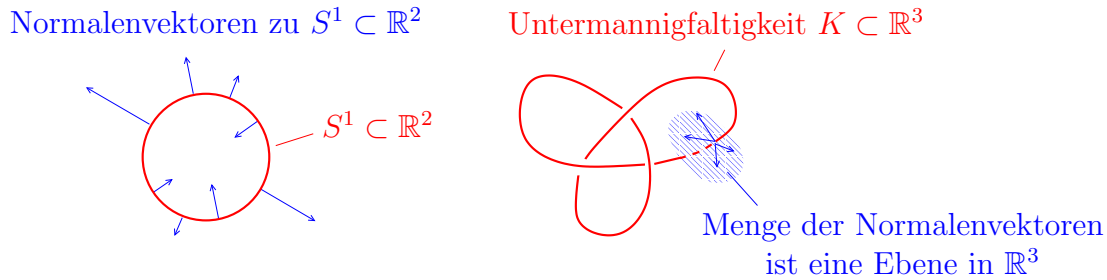


ABBILDUNG 30.

Wenn wir Untermannigfaltigkeiten betrachten, welche als Urbild eines regulären Wertes einer Funktion $f: U \rightarrow \mathbb{R}$ gegeben sind, dann besagt folgender Satz, welcher die Aussagen von Satz 11.5 und 11.6 aus Analysis II kombiniert, dass es besonders einfach ist die Tangentialräume und die Normalenvektoren zu bestimmen.

Satz 6.2. *Es sei $U \subset \mathbb{R}^n$ offen, es sei $f: U \rightarrow \mathbb{R}$ eine glatte Abbildung und es sei $z \in \mathbb{R}$ ein regulärer Wert. Dann ist $M := f^{-1}(z) \subset \mathbb{R}^n$ eine $(n-1)$ -dimensionale Untermannigfaltigkeit, und für jeden Punkt $P \in M$ gilt:*

$$T_P M = (\text{Grad } f(P))^\perp.$$

Zudem gilt

$$\text{Menge der Normalenvektoren am Punkt } P = \mathbb{R} \cdot \text{Grad } f(P).$$

Beispiel.

⁴⁴Allerdings ist der Tangentialraum ein Vektorraum, er geht also durch den Ursprung und normalerweise nicht durch den Punkt P . Um den "anschaulichen Tangentialraum" zu erhalten müssen wir $T_P M$ vom Ursprung zum Punkt P verschieben.

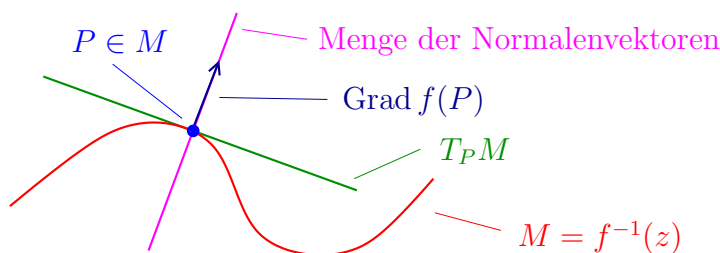


ABBILDUNG 31. Illustration von Satz 6.2.

- (1) Wir betrachten zuerst die Untermannigfaltigkeit

$$S^1 = \{(x, y) \in \mathbb{R}^2 \mid \underbrace{x^2 + y^2}_{=: f(x, y)} = 1\}.$$

Dann gilt $\text{Grad } f(x, y) = (2x, 2y)$. Es folgt aus Satz 6.2, dass der Tangentialraum von M an einem beliebigen Punkt $(x, y) \in S^1$ gegeben ist durch

$$T_{(x, y)} S^1 = (2x, 2y)^\perp = \{\lambda(y, -x) \mid \lambda \in \mathbb{R}\}.$$

- (2) Wir betrachten nun das Ellipsoid

$$M = \{(x, y, z) \in \mathbb{R}^3 \mid \underbrace{2x^2 + 3y^2 + 4z^2}_{=: f(x, y, z)} = 9\}.$$

Dann gilt $\text{Grad } f(x, y, z) = (4x, 6y, 8z)$, und es folgt aus Satz 6.2, dass der Tangentialraum von M am Punkt $(1, 1, 1)$ gegeben ist durch

$$T_{(1, 1, 1)} M = (4, 6, 8)^\perp = \{\lambda(3, -2, 0) + \mu(-2, 0, 1) \mid \lambda, \mu \in \mathbb{R}\}.$$

6.2. Vektorfelder auf Untermannigfaltigkeiten. Wir erinnern an folgende Definition aus der Analysis II Vorlesung.

Definition. Es sei $M \subset \mathbb{R}^n$ eine Untermannigfaltigkeit. Ein *Vektorfeld* auf M ist eine stetige Abbildung $v: M \rightarrow \mathbb{R}^n$. Wir sagen v ist stetig differenzierbar (beziehungsweise glatt), wenn die Koordinatenfunktionen von v stetig differenzierbar (beziehungsweise glatt) sind.⁴⁵

Etwas salopp gesprochen ordnet ein Vektorfeld jedem Punkt in M einen Vektor zu. Vektorfelder (z.B. das magnetische Feld, Gravitationsfeld, aber auch Windrichtung etc.) spielen eine wichtige Rolle in der Physik.

Definition. Es sei $M \subset \mathbb{R}^n$ eine Untermannigfaltigkeit und es sei $F: M \rightarrow \mathbb{R}^n$ ein Vektorfeld auf M .

- (1) F heißt *Tangentialvektorfeld*, wenn $F(P) \in T_P M$ für alle $P \in M$,
- (2) F heißt *Normalenfeld*, wenn $F(P) \perp T_P M$ für alle $P \in M$,
- (3) F heißt *normiert*, wenn $\|F(P)\| = 1$ für alle $P \in M$,

⁴⁵Wir hatten in Kapitel 2.5 definiert, was es heißt, dass eine reellwertige Funktion auf einer Untermannigfaltigkeit stetig, stetig differenzierbar etc. ist.

(4) ein normiertes Normalenfeld heißt *Einheitsnormalenfeld*.

In Abbildung 32 skizzieren wir zwei Einheitsnormalenfelder auf dem Kreis X von Radius zwei. Eines davon zeigt “nach außen” und das andere “nach innen”. Wir werden gleich sehen, dass X genau zwei Einheitsnormalenfelder besitzt.

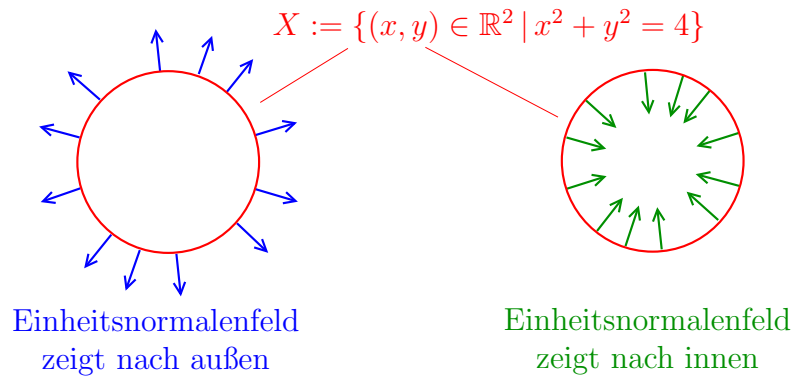


ABBILDUNG 32. Einheitsnormalenfelder für den Kreis von Radius zwei.

Definition.

- (1) Eine *Hyperfläche*⁴⁶ in $\mathbb{R}^n$ ist eine Untermannigfaltigkeit von $\mathbb{R}^n$ der Kodimension 1, d.h. eine Untermannigfaltigkeit der Dimension $n - 1$ in $\mathbb{R}^n$.
- (2) Es sei N eine Hyperfläche. Eine *Orientierung von N* ist ein Einheitsnormalenfeld auf N . Wir sagen N ist *orientierbar*, wenn N eine Orientierung besitzt.
- (3) Eine *orientierte Hyperfläche* ist ein Paar (N, ν) , wobei N eine Hyperfläche und ν ein Orientierung von N ist.

Beispiele.

- (1) Für $S^n \subset \mathbb{R}^{n+1}$ gibt es zwei Orientierungen, diese sind gegeben durch die beiden Einheitsnormalenfelder

$$\begin{array}{ccc} F: S^n & \rightarrow & \mathbb{R}^{n+1} \\ x & \mapsto & x \end{array} \quad \text{und} \quad \begin{array}{ccc} F: S^n & \rightarrow & \mathbb{R}^{n+1} \\ x & \mapsto & -x. \end{array}$$

- (2) Es sei $U \subset \mathbb{R}^{n-1}$ eine offene Teilmenge und $f: U \rightarrow \mathbb{R}$ eine glatte Funktion. Wir betrachten die Hyperfläche

$$M = \{(x, f(x)) \mid x \in U\}.$$

Für einen Punkt $P = (x, f(x)) \in \partial M$ werden wir in Übungsblatt 3 sehen, dass

$$\chi(P) = (-\text{Grad } f(x), 1)$$

⁴⁶Der Name ist vielleicht etwas unglücklich, weil es sich im Allgemeinen *nicht* um eine Fläche handelt. Aber nachdem sich der Name so eingebürgert hat, werden wir ihn auch verwenden.

ein Normalenvektor ist. (Dieser Normalenvektor wird für $n = 2$ in Abbildung 33 skizziert.) Also ist dann durch

$$\nu(P) = \frac{\chi(P)}{\|\chi(P)\|} = \frac{(-\text{Grad } f(x), 1)}{\sqrt{\|\text{Grad } f(x)\|^2 + 1}}$$

ein Einheitsnormalenfeld auf ∂M definiert.

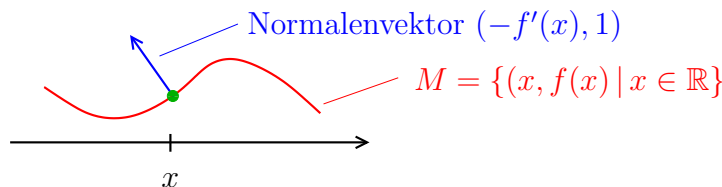


ABBILDUNG 33.

- (3) Als letztes Beispiel betrachten wir noch
- (4) das Möbiusband, welches in Abbildung 34 skizziert ist. Dies ist eine Hyperfläche,

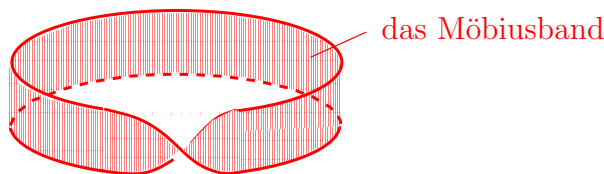


ABBILDUNG 34. Das Möbiusband besitzt kein Einheitsnormalenfeld.

aber N besitzt kein Einheitsnormalenfeld, d.h. das Möbiusband ist nicht orientierbar. Diese Aussage ist bildlich klar, ein sauberer Beweis ist allerdings etwas aufwändig und erfolgt im folgendem Lemma.

Lemma 6.3. *Das Möbiusband ist nicht orientierbar.*

Beweis. Für $\varphi, \psi \in \mathbb{R}$, $\epsilon \in \{-1, 1\}$ und $r \in (-1, 1)$ setzen wir

$$A(\varphi) := \begin{pmatrix} \cos \varphi & -\sin \varphi & 0 \\ \sin \varphi & \cos \varphi & 0 \\ 0 & 0 & 1 \end{pmatrix}, u = \begin{pmatrix} 2 \\ 0 \\ 0 \end{pmatrix}, v(r, \psi) := \begin{pmatrix} 2 + r \sin \psi \\ 0 \\ r \cos \psi \end{pmatrix} \text{ und } w_\epsilon(\psi) := \begin{pmatrix} \sin(\psi + \epsilon \frac{\pi}{2}) \\ 0 \\ \cos(\psi + \epsilon \frac{\pi}{2}) \end{pmatrix}.$$

Wir beschreiben das Möbiusband als

$$M := \{A(\varphi) \cdot v(r, \frac{1}{2}\varphi) \mid \varphi \in [0, 2\pi] \text{ und } r \in (-1, 1)\}.$$

Wir erhalten also das Möbiusband in dem wir ein offenes Intervall entlang des Kreises

$$K := \{A(\varphi) \cdot u \mid \varphi \in [0, 2\pi]\}$$

verdrillen. An einem Punkt $K(\varphi) = A(\varphi) \cdot u$ ist die Menge der normierten Normalenvektoren

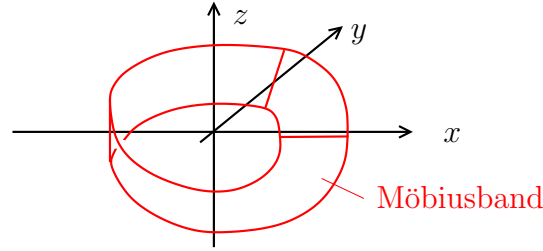


ABBILDUNG 35. Skizze zum Beweis von Lemma 6.3

gegeben durch $\tilde{K}(\varphi) := \{A(\varphi) \cdot w_\epsilon(\frac{\varphi}{2}) \mid \epsilon \in \{-1, 1\}\}$. Nehmen wir an, es gäbe doch ein Einheitsnormalenfeld ν . Indem wir ν auf K einschränken erhalten wir insbesondere eine stetige Abbildung

$$\nu: K \rightarrow \tilde{K} := \{\tilde{K}(\varphi) \mid \varphi \in [0, 2\pi]\},$$

mit $\nu(K(\varphi)) \in \tilde{K}(\varphi)$. Für $\varphi \in [0, 2\pi)$ bezeichnen wir mit $\epsilon(\varphi) \in \{-1, 1\}$ das eindeutig bestimmte Element mit $\nu(K(\varphi)) = A(\varphi) \cdot w_{\epsilon(\varphi)}(\frac{\varphi}{2})$. Die Abbildung $\epsilon: [0, 2\pi) \rightarrow \{-1, 1\}$ ist stetig, also nach Lemma 5.3 konstant. Dann gilt

$$\begin{array}{ccccccc} & & \text{da } K(0) = K(2\pi) & & \text{da } \nu \text{ stetig} & & \\ & & \downarrow & & \downarrow & & \\ A(0) \cdot w_{\epsilon(0)}(0) & = & \nu(K(0)) & = & \nu(K(2\pi)) & = & \lim_{\varphi \rightarrow 2\pi} \nu(K(\varphi)) = \lim_{\varphi \rightarrow 2\pi} A(\varphi) \cdot w_{\epsilon(\varphi)}(\frac{\varphi}{2}) \\ & = & \lim_{\varphi \rightarrow 2\pi} A(\varphi) \cdot w_{\epsilon(0)}(\frac{\varphi}{2}) & = & A(2\pi) \cdot w_{\epsilon(0)}(\pi) & = & A(0) \cdot (-w_{\epsilon(0)}(0)). \\ \uparrow & & & & & & \uparrow \\ \text{da } \epsilon \text{ konstant} & & & & & & \text{da } A(2\pi) = A(0) \text{ und } w_\epsilon(\psi) = -w_\epsilon(\psi + \pi) \end{array}$$

Wir haben also gezeigt, dass $A(0) \cdot w_{\epsilon(0)}(0) = -A(0) \cdot w_{\epsilon(0)}(0)$. Aber dieser Vektor ist nicht der Nullvektor. Wir haben also einen Widerspruch erhalten. $\square$

Satz 6.4. *Eine zusammenhängende Hyperfläche besitzt entweder keine oder genau zwei Orientierungen.*

Beweis. Es sei N eine zusammenhängende Hyperfläche. Wir nehmen an F besitzt eine Orientierung $\nu: N \rightarrow \mathbb{R}^n$. Dann ist auch $-\nu: N \rightarrow \mathbb{R}^n$ eine Orientierung. Es sei nun $\mu: N \rightarrow \mathbb{R}^n$ eine weitere Orientierung. Wir wollen zeigen, dass entweder $\mu = \nu$ oder $\mu = -\nu$. Für jeden Punkt $P \in N$ gibt es genau zwei Normalenvektoren zum Vektorraum $T_P M$ der Länge eins. Es gilt also entweder $\mu(P) = \nu(P)$ oder $\mu(P) = -\nu(P)$. Wir betrachten die Funktion

$$\begin{aligned} \varphi: N &\rightarrow \{-1, 1\} \\ P &\mapsto \begin{cases} +1, & \text{wenn } \mu(P) = \nu(P), \\ -1, & \text{wenn } \mu(P) = -\nu(P). \end{cases} \end{aligned}$$

Nachdem ν und μ stetig sind ist auch die Funktion φ stetig. Nachdem N zusammenhängend ist und nachdem $\{-1, 1\}$ eine diskrete Teilmenge von $\mathbb{R}$ ist, folgt nun aus

Lemma 5.3, dass φ konstant ist. Es ist also entweder $\mu = \nu$ oder $\mu = -\nu$. Wir haben damit bewiesen, dass N genau zwei Orientierungen besitzt. $\square$

Satz 6.5. *Es sei $M \subset \mathbb{R}^n$ eine n -dimensionale Untermannigfaltigkeit. Dann gibt es genau ein Einheitsnormalenfeld $\nu: \partial M \rightarrow \mathbb{R}^n$ mit der Eigenschaft, dass es für jedes $P \in \partial M$ ein $\epsilon > 0$ gibt, so dass $P + t \cdot \nu(P) \notin M$ für alle $t \in (0, \epsilon)$.*

Definition. Es sei $M \subset \mathbb{R}^n$ eine n -dimensionale Untermannigfaltigkeit. Wir nennen das in Satz 6.5 eindeutig bestimmte Einheitsnormalenfeld auf ∂M , das *äußere Einheitsnormalenfeld*.

Bemerkung. Anschaulich gesprochen ist das äußere Einheitsnormalenfeld an jedem Punkt $P \in \partial M$ gegeben durch den Einheitsnormalenvektor, welcher “nach außen” zeigt.

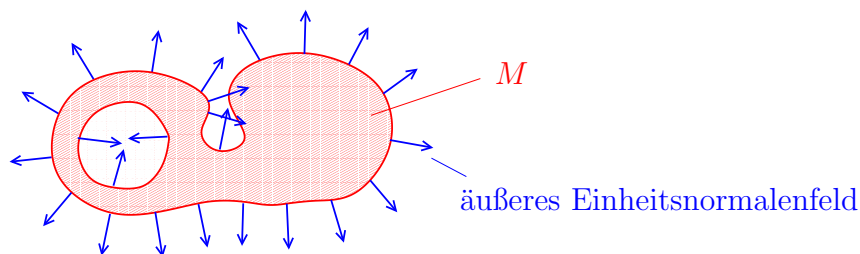


ABBILDUNG 36. Skizze für das äußere Einheitsnormalenfeld.

Beweis. Es sei also $M \subset \mathbb{R}^n$ eine n -dimensionale Untermannigfaltigkeit. Außerdem sei $P \in \partial M$. Im Hinblick auf Satz 4.4 können wir, eventuell nach Anwendung einer Permutationsmatrix, annehmen, dass es einen offenen Quader $Q \subset \mathbb{R}^{n-1}$, reelle Zahlen $a < b$ und eine glatte Funktion $f: Q \rightarrow (a, b)$ gibt, so dass

$$M \cap (Q \times (a, b)) = \{(x, y) \mid x \in Q \text{ und } y \leq f(x)\}$$

und

$$\partial M \cap (Q \times (a, b)) = \{(x, y) \mid x \in Q \text{ und } y = f(x)\}.$$

Wir wählen $x \in Q$, so dass $P = (x, f(x))$. Es folgt dann aus dem Beispiel auf Seite 61, dass

$$\nu(P) := \frac{\chi(P)}{\|\chi(P)\|} \text{ mit } \chi(P) := (-\text{Grad } f(x), 1)$$

ein normierter Normalenvektor zu ∂M ist. Es gibt zudem⁴⁷ ein $\epsilon > 0$, so dass $P + t \cdot \nu(P) \notin M$ für alle $t \in (0, \epsilon)$, während es für alle $\epsilon > 0$ ein $t \in (0, \epsilon)$ gibt, so dass $P + t \cdot (-\nu(P))$ in M liegt. Die Abbildung ν ist also in P eindeutig bestimmt. Es folgt zudem aus der Definition, dass die gerade konstruierte Abbildung $\nu: \partial M \cap (Q \times (a, b)) \rightarrow \mathbb{R}^n$ stetig ist. Also definiert ν ein Vektorfeld auf $\partial M \cap (Q \times (a, b))$. $\square$

⁴⁷Warum?

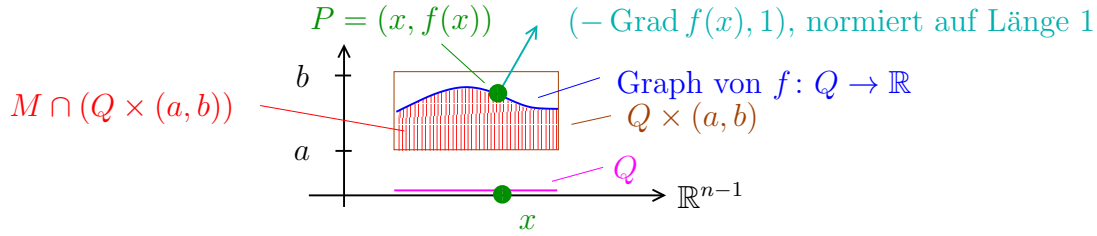


ABBILDUNG 37. Konstruktion des äußeren Einheitsnormalenfeldes.

Beispiel. Es sei $U \subset \mathbb{R}^n$ offen, es sei $f: U \rightarrow \mathbb{R}$ eine glatte Funktion und es seien $z_1, z_2 \in \mathbb{R}$ reguläre Werte. Nach Satz 4.2 ist $M := f^{-1}([z_1, z_2]) \subset \mathbb{R}^n$ eine n -dimensionale Untermannigfaltigkeit mit

$$\partial M = f^{-1}(z_1) \cup f^{-1}(z_2).$$

Mithilfe von Satz 6.2 kann man zudem leicht zeigen, dass das äußere Einheitsnormalenfeld gegeben ist durch

$$\nu(P) = \frac{1}{\|\text{Grad } f(P)\|} \text{Grad } f(P) \quad \text{für } P \in f^{-1}(z_2)$$

und ⁴⁸

$$\nu(P) = - \frac{1}{\|\text{Grad } f(P)\|} \text{Grad } f(P) \quad \text{für } P \in f^{-1}(z_1).$$

Beispielsweise ist für $(x_1, \dots, x_n) \in S^{n-1} = \partial D^n$ das äußere Einheitsnormalenfeld gegeben durch $\nu(x_1, \dots, x_n) = (x_1, \dots, x_n)$.

6.3. Formulierung des Gaußschen Integralsatzes. Bevor wir den Gaußschen Integralsatz formulieren können, benötigen wir noch zwei weitere Definitionen:

Definition. Es sei (N, ν) eine orientierte kompakte Hyperfläche in $\mathbb{R}^n$ und es sei $F: N \rightarrow \mathbb{R}^n$ ein Vektorfeld. Die Funktion

$$\begin{aligned} N &\rightarrow \mathbb{R} \\ P &\mapsto \underbrace{F(P) \cdot \nu(P)}_{\text{Skalarprodukt}} \end{aligned}$$

ist dann eine stetige Funktion auf der kompakten Untermannigfaltigkeit N , insbesondere integrierbar. Das Integral

$$\int_N F \cdot \nu$$

wird der *Fluss des Vektorfeldes F durch die orientierte Hyperfläche (N, ν)* genannt.

⁴⁸Das Minuszeichen rührt daher, dass das Normalenfeld “nach außen” zeigen soll.

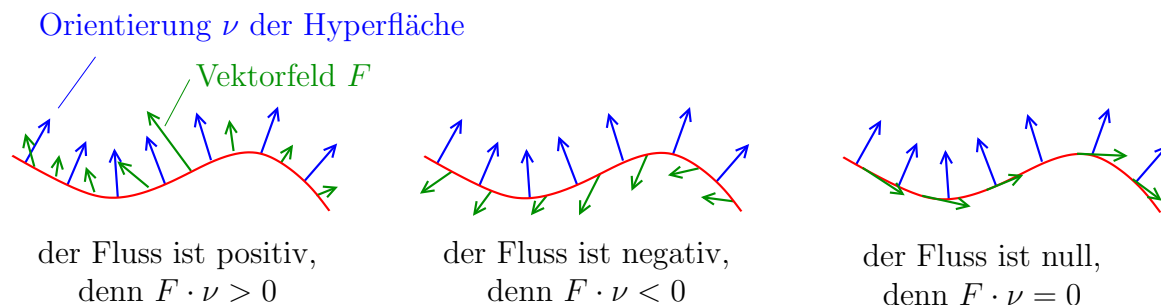


ABBILDUNG 38. Skizze zum Fluss eines Vektorfeldes durch eine orientierte Hyperfläche.

Definition. Es sei U eine offene Teilmenge von $\mathbb{R}^n$ und es sei $F: U \rightarrow \mathbb{R}^n$ ein differenzierbares Vektorfeld. Wir definieren die *Divergenz von $F = (F_1, \dots, F_n)$* als die Funktion⁴⁹

$$\begin{aligned} \operatorname{div} F: U &\rightarrow \mathbb{R} \\ P &\mapsto \operatorname{div} F(P) := \sum_{i=1}^n \frac{\partial F_i}{\partial x_i}(P). \end{aligned}$$

Wir können nun den Gaußschen Integralsatz formulieren:

Satz 6.6. (Gaußscher Integralsatz) *Es sei $M \subset \mathbb{R}^n$ eine kompakte n -dimensionale Untermannigfaltigkeit. Wir bezeichnen mit ν das äußere Einheitsnormalenfeld auf ∂M . Es sei $F: M \rightarrow \mathbb{R}^n$ ein stetig differenzierbares Vektorfeld.⁵⁰ Dann gilt*

$$\int_{\partial M} F \cdot \nu = \int_M \operatorname{div} F.$$

In Worten ausgedrückt besagt der Gaußsche Integralsatz Folgendes: für ein stetig differenzierbares Vektorfeld F auf einer kompakten Untermannigfaltigkeit M der Kodimension Null gilt

$$\begin{array}{ccc} \text{Fluss von } F \text{ durch die} & = & \text{Integral über die} \\ \text{“Oberfläche” von } M & & \text{Divergenz von } F \text{ auf } M. \end{array}$$

Beispiel. Wir wollen die Aussage des Gaußschen Integralsatz im Spezialfall $n = 1$ nachvollziehen. Dann gilt:

- (1) eine kompakte eindimensionale zusammenhängende Untermannigfaltigkeit von $\mathbb{R}$ ist ein kompaktes Intervall $M = [a, b] \subset \mathbb{R}$ mit $a < b$, wobei $\partial M = \{a\} \cup \{b\}$,
- (2) ein stetig differenzierbares Vektorfeld auf M ist eine stetig differenzierbare Abbildung $F: M \rightarrow \mathbb{R}$,

⁴⁹Diese Definition erscheint vielleicht auf den ersten Blick etwas willkürlich. Eine alternative Definition wäre, dass $\operatorname{div} F(P) = \operatorname{Spur}(DF(P))$ die Spur des Differentials der Abbildung $F: U \rightarrow \mathbb{R}^n$ ist.

⁵⁰Wir verwenden hierbei die gleichen Konvention wie in den Fußnoten 17 und 43. Genauer gesagt, ein stetig differenzierbares Vektorfeld $F: M \rightarrow \mathbb{R}^n$ ist ein stetiges Vektorfeld auf M , welches auf der offenen Menge $M \setminus \partial M$ stetig differenzierbar ist, und so dass sich die partiellen Ableitungen der Koordinatenfunktionen von F (und damit auch die Divergenz) zu stetigen Funktionen auf ganz M fortsetzen.

- (3) die Divergenz von F ist die Ableitung F' von F ,
 (4) der äußere Einheitsnormalenvektor in a ist -1 und der äußere Einheitsnormalenvektor in b ist 1 .

Dann gilt

$$\begin{aligned}
 \int_a^b F'(x) dx &= \int_{[a,b]} \operatorname{div} F = \int_{\partial[a,b]} F \cdot \nu = \int_{\{a\}} F \cdot \nu + \int_{\{b\}} F \cdot \nu \\
 &\quad \uparrow \qquad \qquad \uparrow \\
 &\text{denn } \operatorname{div} F = F' \quad \text{Gaußscher Integralsatz} \\
 &= F(a) \cdot (-1) + F(b) \cdot 1 = F(b) - F(a). \\
 &\quad \uparrow \\
 &\text{für eine nulldimensionale Untermannigfaltigkeit } N = \{P_1, \dots, P_l\} \\
 &\text{von } \mathbb{R}^n \text{ und } f: N \rightarrow \mathbb{R} \text{ eine Funktion gilt } \int_N f = f(P_1) + \dots + f(P_l)
 \end{aligned}$$

Im Falle $n = 1$ ist der Gaußsche Integralsatz also äquivalent zum Hauptsatz der Differential- und Integralrechnung aus der Analysis I.

Beispiel. Es sei $M \subset \mathbb{R}^2$ eine beliebige kompakte 2-dimensionale Untermannigfaltigkeit. Wir betrachten die Vektorfelder, welche gegeben sind durch

$$F(x, y) = (1, 0) \quad \text{und} \quad G(x, y) = (x, y) \quad \text{sowie} \quad H(x, y) = (-y, x).$$

Hierbei ist

$$\operatorname{div} F(x, y) = 0 \quad \text{und} \quad \operatorname{div} G(x, y) = 2 \quad \text{sowie} \quad \operatorname{div} H(x, y) = 0.$$

Es folgt also aus dem Gaußschen Integralsatz, dass der Fluss von F und H durch ∂M null ist, während der Fluss von G durch ∂M positiv ist.

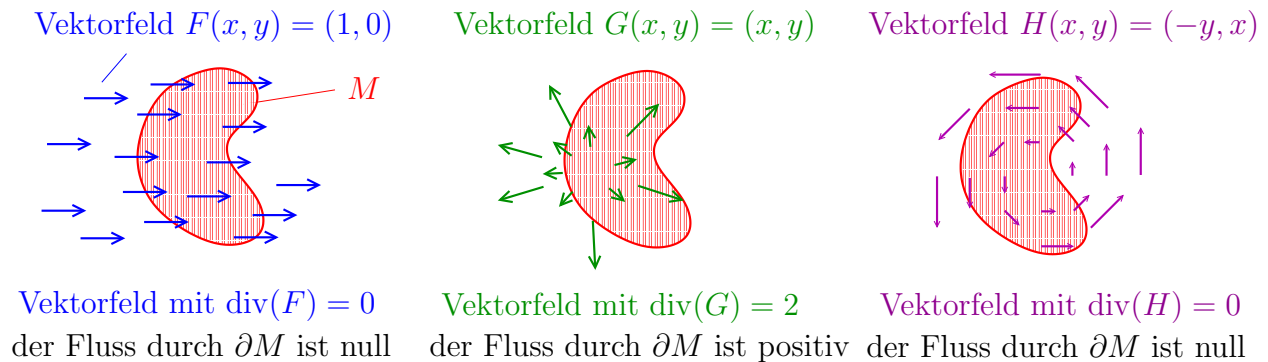


ABBILDUNG 39.

Die Beweisidee für den Beweis vom Gaußschen Integralsatz ist nun wie folgt: wir überdecken die Untermannigfaltigkeit M mit zwei Typen von offenen Quadern:

- (i) offene Quader, welche ganz in M liegen,

- (ii) offene Quader, so dass der Schnitt mit ∂M , eventuell nach Anwendung einer Permutationsmatrix, ein Graph ist.

Ganz analog zur Definition vom Integral einer Funktion auf einer Untermannigfaltigkeit führen wir nun den Beweis zurück auf die Betrachtung von Vektorfeldern auf solchen Quadern. In beiden Fällen folgt dann die Aussage mithilfe einer expliziten, wenn auch im Fall (ii) langwierigen, Rechnung. Wir führen nun diese Beweisidee in den folgenden Kapiteln aus.

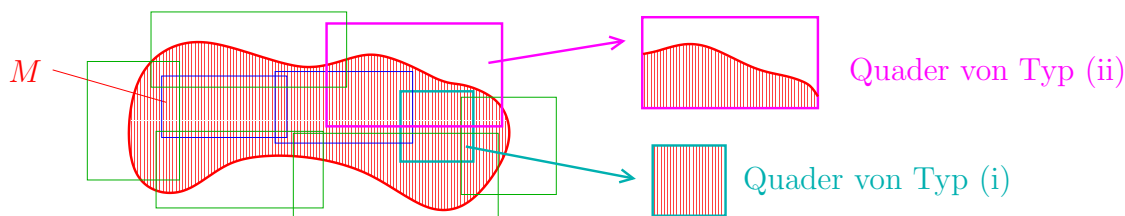


ABBILDUNG 40. Schematisches Bild der offenen Mengen $U_1, \dots, U_k$.

6.4. Beweis des Gaußschen Integralsatzes I: Quader.

Satz 6.7. *Es sei F ein stetig differenzierbares Vektorfeld auf einem n -dimensionalen Quader Q in $\mathbb{R}^n$, welches auf dem Rand ∂Q verschwindet. Dann gilt⁵¹*

$$\int_Q \operatorname{div} F = 0.$$

Satz 6.7 folgt sofort aus der Definition

$$\operatorname{div}(F_1, \dots, F_n) = \sum_{i=1}^n \frac{\partial F_i}{\partial x_i},$$

der Linearität des Integrals und folgendem Lemma.

Lemma 6.8. *Es sei $\psi: Q \rightarrow \mathbb{R}$ eine stetig differenzierbare Funktion auf einem n -dimensionalen Quader, welche auf dem Rand ∂Q verschwindet. Dann gilt für alle $i \in \{1, \dots, n\}$, dass*

$$\int_Q \frac{\partial}{\partial x_i} \psi = 0.$$

⁵¹Nachdem Q keine Untermannigfaltigkeit ist, ist dies kein Spezialfall vom Gaußschen Integralsatz. Nichtsdestotrotz ist diese Aussage ganz im Sinne vom Gaußschen Integralsatz: der Fluß durch die "Oberfläche von Q " verschwindet, denn das Vektorfeld ist dort nach Voraussetzung null, also muss auch das Integral über die Divergenz verschwinden.

Beweis. Wir betrachten zuerst den Spezialfall $n = 1$. Es sei also $Q = [a, b] \subset \mathbb{R}$ ein Quader und es sei $\varphi: Q \rightarrow \mathbb{R}$ eine stetig differenzierbare Funktion, so dass $\varphi(a) = \varphi(b) = 0$. Dann gilt

$$\int_{x=a}^{x=b} \frac{\partial}{\partial x} \varphi(x) dx = \varphi(b) - \varphi(a) = 0.$$

Wir betrachten nun den allgemeinen Fall. Es sei also $Q = [a_1, b_1] \times \cdots \times [a_n, b_n] \subset \mathbb{R}^n$ ein Quader und es sei $\psi: Q \rightarrow \mathbb{R}$ eine stetig differenzierbare Funktion, so dass ψ auf dem Rand ∂Q verschwindet. Wir beweisen nun die Aussage für $i = 1$. Es ist

$$\int_Q \frac{\partial}{\partial x_1} \psi = \int_{x_n=a_n}^{x_n=b_n} \cdots \int_{x_2=a_2}^{x_2=b_2} \underbrace{\int_{x_1=a_1}^{x_1=b_1} \frac{\partial}{\partial x_1} \psi(x_1, \dots, x_n) dx_1}_{= 0, \text{ nach dem obigen Fall angewandt auf } \varphi(x_1) = \psi(x_1, \dots, x_k)} dx_2 \dots dx_n = 0.$$

$\uparrow$
 Satz von Fubini

Die Aussage für die anderen i 's wird ganz analog bewiesen. $\square$

6.5. Beweis des Gaußschen Integralsatzes II: Stempel. Es sei $Q \subset \mathbb{R}^{n-1}$ ein $(n-1)$ -dimensionaler Quader und es sei $g: Q \rightarrow (0, \infty)$ eine glatte Funktion. Wir nennen

$$S := \{(x, y) \mid x \in Q \text{ und } y \in [0, g(x)]\}$$

einen n -dimensionalen *Stempel* und wir bezeichnen

$$\partial_0 S := \{(x, y) \mid x \in \partial Q \text{ und } y = g(x)\}$$

als die *Stempelfläche*. Wir bezeichnen mit $\partial_1 S := \partial S \setminus \partial_0 S$ den Rest des Randes von S . Wie üblich bezeichnen wir mit ν das äußere Einheitsnormalenfeld auf $\partial_0 S$

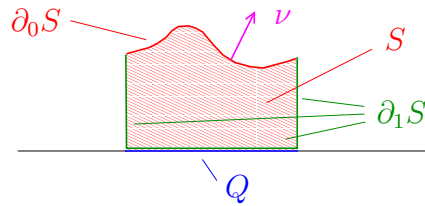


ABBILDUNG 41. Schematisches Bild eines Stempels.

Der folgende Satz behandelt nun Ausschnitte von Untermannigfaltigkeiten von Typ (ii), welche wir auf Seite 67 eingeführt hatten.

Satz 6.9. *Es sei S ein n -dimensionaler Stempel und es sei $F: S \rightarrow \mathbb{R}^n$ ein stetig differenzierbares Vektorfeld, so dass F auf $\partial_1 S$ verschwindet. Dann gilt*

$$\int_S \operatorname{div} F = \int_{\partial_0 S} F \cdot \nu.$$

Der Beweis besteht aus einer heroischen Rechnung und beruht wiederum darauf, die Aussage des Satzes auf den Hauptsatz der Differential- und Integralrechnung aus der Analysis I zurückzuführen. In dem Beweis werden wir folgendes Lemma verwenden.

Lemma 6.10. *Es seien*

$$\begin{aligned} \varphi: [a, b] \times [0, c] &\rightarrow \mathbb{R} & \text{und} & & h: [a, b] &\rightarrow [0, c] \\ (z, t) &\mapsto \varphi(z, t) & & & z &\mapsto h(z) \end{aligned}$$

*stetig differenzierbare Funktionen. Dann gilt die Gleichheit*⁵²

$$\frac{d}{dz} \int_{t=0}^{t=h(z)} \varphi(z, t) dt = \int_{t=0}^{t=h(z)} \frac{\partial}{\partial z} \varphi(z, t) dt + \varphi(z, h(z)) \cdot \frac{\partial h}{\partial z}.$$

Der Beweis von Lemma 6.10 ist eine Verallgemeinerung von Übungsaufgabe 3 (b) in Übungsblatt 2.

Beweis. Wir betrachten die beiden Funktionen

$$\begin{aligned} \Phi: [0, c] \times [a, b] &\rightarrow \mathbb{R} & \Psi: [a, b] &\rightarrow \mathbb{R}^2 \\ (x, y) &\mapsto \Phi(x, y) := \int_{t=0}^{t=x} \varphi(y, t) dt & \text{und} & & z &\mapsto \Psi(z) := \begin{pmatrix} h(z) \\ z \end{pmatrix}. \end{aligned}$$

Dann ist

$$\begin{aligned} \frac{d}{dz} \int_{t=0}^{t=h(z)} \varphi(z, t) dt &= \frac{d}{dz} (\Phi \circ \Psi)(z) \stackrel{\uparrow}{=} D\Phi(\Psi(z)) \cdot D\Psi(z) \\ &\stackrel{\text{Kettenregel}}{=} \left(\frac{\partial}{\partial x} \int_{t=0}^{t=x} \varphi(y, t) dt \quad \frac{\partial}{\partial y} \int_{t=0}^{t=x} \varphi(y, t) dt \right) \bigg|_{(x,y)=\Psi(z)} \cdot \begin{pmatrix} \frac{\partial h}{\partial z} \\ 1 \end{pmatrix} \\ &\stackrel{\uparrow}{=} \left(\varphi(y, x) \quad \int_{t=0}^{t=x} \frac{\partial}{\partial y} \varphi(y, t) dt \right) \bigg|_{(x,y)=(h(z), z)} \cdot \begin{pmatrix} \frac{\partial h}{\partial z} \\ 1 \end{pmatrix} \\ &\stackrel{\text{HDI und Satz 3.10}}{=} \varphi(z, h(z)) \cdot \frac{\partial h}{\partial z} + \int_{t=0}^{t=h(z)} \frac{\partial}{\partial z} \varphi(z, t) dt. \end{aligned}$$

Wir haben damit die gewünschte Gleichheit bewiesen. $\square$

Wir wenden uns nun dem eigentlichen Beweis von Satz 6.9 zu.

Beweis von Satz 6.9. Es sei $Q \subset \mathbb{R}^{n-1}$ ein $(n-1)$ -dimensionaler Quader und zudem sei $g: Q \rightarrow (0, \infty)$ eine glatte Funktion. Wir bezeichnen mit $S \subset \mathbb{R}^n$ den dazugehörigen Stempel. Außerdem sei $F = (F_1, \dots, F_n): M \rightarrow \mathbb{R}^n$ ein stetig differenzierbares Vektorfeld

⁵²Beide Seiten sind Funktionen in z auf $[a, b]$.

auf S , welches auf $\partial_1 S$ verschwindet. Wie wir auf Seite 61 gesehen hatten ist das äußere Einheitsnormalenfeld auf $\partial_0 S$ gegeben durch

$$\begin{aligned} \nu = (\nu_1, \dots, \nu_n): \quad \partial_0 S &\mapsto \mathbb{R}^n \\ (x, g(x)) &\mapsto \frac{1}{\sqrt{1 + \|\text{Grad } g(x)\|^2}} (-\text{Grad } g(x), 1). \end{aligned}$$

Wir wollen zuerst einen Ausdruck für das Integral einer Funktion auf der Untermannigfaltigkeit $\partial_0 S$ bestimmen. Wir betrachten dazu die Karte Φ für $\partial_0 S$, welche gegeben ist durch die Umkehrung der Abbildung⁵³

$$\begin{aligned} \Psi: Q \times \mathbb{R} &\rightarrow Q \times \mathbb{R} \\ (x_1, \dots, x_{n-1}, x_n) &\mapsto \begin{pmatrix} x_1 \\ \vdots \\ x_{n-1} \\ \text{gamma}(x_1, \dots, x_{n-1}) + x_n \end{pmatrix}. \end{aligned}$$

Für $x' = (x_1, \dots, x_{n-1}) \in \overset{\circ}{Q}$ berechnen wir

$$\begin{aligned} \text{Gram}_\Phi(x') &= \text{Gram}(\text{ersten } (n-1) \text{ Zeilen von } D_{x'} \Psi) = \text{Gram} \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & & \ddots & \vdots \\ 0 & \dots & 0 & 1 \\ \frac{\partial g}{\partial x_1} & \frac{\partial g}{\partial x_2} & \dots & \frac{\partial g}{\partial x_{n-1}} \end{pmatrix} \\ &= 1 + \|\text{Grad } g(x')\|^2. \\ &\quad \uparrow \\ &\quad \text{siehe Beispiel auf Seite 38} \end{aligned}$$

Für eine Funktion f auf $\partial_0 S$ gilt dann also nach Definition, dass

$$\int_{\partial_0 S} f = \int_{\overset{\circ}{Q}} f(\Psi(x')) \cdot \sqrt{1 + \|\text{Grad } g(x')\|^2} dx'.$$

Nachdem $\text{div}(F) = \sum_{i=1}^n \frac{\partial F_i}{\partial x_i}$ und $F \cdot \nu = \sum_{i=1}^n F_i \cdot \nu_i$ genügt es folgende Behauptung zu beweisen:

Behauptung. Für alle $i = 1, \dots, n$ gilt

$$\int_S \frac{\partial F_i}{\partial x_i} = \int_{\partial_0 S} F_i \cdot \nu_i.$$

Wir unterscheiden jetzt folgende Fälle im Beweis der Behauptung:

⁵³Im Folgenden bezeichnen wir mit Ψ auch die Abbildung

$$\begin{aligned} Q &\rightarrow Q \times \mathbb{R} \\ x' &\mapsto \Psi(x', 0). \end{aligned}$$

(A) Es sei zuerst $i = n$. Dann gilt, dass

$$\begin{aligned}
 \int_S \frac{\partial F_n}{\partial x_n} &= \int_Q \int_0^{g(x')} \frac{\partial F_n}{\partial x_n}(x', x_n) dx_n dx' = \int_Q F_n(x', g(x')) - \overbrace{F_n(x', 0)}^{=0, \text{ da } F_n|_{\partial_1 S} = 0} dx' \\
 &\quad \uparrow \text{Satz von Fubini} \quad \text{da } F_n \text{ Stammfunktion von } \frac{\partial F_n}{\partial x_n} \text{ bezüglich } x_n \\
 &= \int_Q F_n(\Psi(x')) dx' \\
 &= \int_{\overset{\circ}{Q}} F_n(\Psi(x')) \cdot \underbrace{\frac{1}{\sqrt{1+\|\text{Grad } g(x')\|^2}}}_{=\nu_n(\Psi(x'))} \cdot \sqrt{1+\|\text{Grad } g(x')\|^2} dx' = \int_{\partial_0 S} F_n \cdot \nu_n. \\
 &\quad \uparrow \text{siehe vorherige Seite}
 \end{aligned}$$

(B) Wir betrachten nun den Fall $i = 1$. Wir zerlegen den Quader Q als $Q = [a, b] \times P$, wobei $P \subset \mathbb{R}^{n-1}$ ein $(n-1)$ -dimensionaler Quader ist. Dann gilt, dass

$$\begin{aligned}
 &\text{Satz von Fubini} \\
 \int_S \frac{\partial F_1}{\partial x_1} &\downarrow \\
 &= \int_{x'' \in P} \int_{x_1=a}^{x_1=b} \int_0^{g(x_1, x'')} \frac{\partial}{\partial x_1} F_1(x_1, x'', x_n) dx_n dx_1 dx'' \\
 &\quad \text{Lemma 6.10 angewandt auf } \varphi(x_1, x_n) = F_1(x_1, x'', x_n) \text{ und } h(x_1) = g(x_1, x'') \\
 &\quad \text{mit } x'' \in P \text{ fest gew\u00e4hlt} \\
 &= \int_P \int_{x_1=a}^{x_1=b} \frac{\partial}{\partial x_1} \int_0^{g(x')} F_1(x_1, x'', x_n) dx_n - F_1(x_1, x'', g(x_1, x'')) \cdot \frac{\partial g}{\partial x_1}(x_1, x'') dx_1 dx'' \\
 &= \int_P \int_{x_1=a}^{x_1=b} \frac{\partial}{\partial x_1} \int_0^{g(x')} F_1(x_1, x'', x_n) dx_n dx_1 dx'' - \int_Q F_1(\Psi(x')) \cdot \frac{\partial g}{\partial x_1}(x') dx' \\
 &= \int_P \underbrace{\int_0^{g(x')} F_1(b, x'', x_n) dx_n}_{=0, \text{ da } F_1|_{\partial_1 S} = 0} - \underbrace{\int_0^{g(x')} F_1(a, x'', x_n) dx_n}_{=0, \text{ da } F_1|_{\partial_1 S} = 0} dx'' - \int_Q F_1(\Psi(x')) \cdot \frac{\partial g}{\partial x_1}(x') dx' \\
 &\quad \uparrow \\
 &\quad \text{da } \int_{x_1=a}^{x_1=b} h(x_1) dx_1 = h(b) - h(a) \\
 &= \int_{\overset{\circ}{Q}} F_1(\Psi(x')) \cdot \underbrace{\left(-\frac{\partial g}{\partial x_1}(x') \frac{1}{\sqrt{1+\|\text{Grad } g(x')\|^2}} \right)}_{=\nu_1(\Psi(x'))} \cdot \sqrt{1+\|\text{Grad } g(x')\|^2} dx' \\
 &= \int_{\partial_0 S} F_1 \cdot \nu_1.
 \end{aligned}$$

(C) Es sei nun $i \in \{2, \dots, n-1\}$ beliebig. Dieser Fall wird genau wie Fall (B) bewiesen, man zerlegt dabei den Quader Q in ein Intervall, welches der i -ten Koordinate entspricht und einen $(n-1)$ -dimensionalen Quader. $\square$

Es sei nun S ein Stempel und $P \in \mathbb{R}^n$ ein Punkt, wir nennen dann

$$P + S := \{P + Q \mid Q \in S\}$$

einen *verschobenen Stempel*. Es ist leicht zu sehen, dass die analoge Aussage von Satz 6.9 auch für verschobene Stempel gilt.

6.6. Beweis des Gaußschen Integralsatzes III: Der allgemeine Fall. Wir hatten schon auf Seite 68 erwähnt, dass wir folgende Strategie für den Beweis vom Gaußschen Integralsatz verfolgen: wir überdecken M durch endlich viele offene Quader und Stempel, welches es uns erlauben, den allgemeinen Fall auf die Aussagen von Satz 6.7 und Satz 6.9 zurückzuführen. Wir benötigen dazu folgendes Lemma.

Lemma 6.11. *Es sei $M \subset \mathbb{R}^n$ eine kompakte Teilmenge und es seien $U_1, \dots, U_k$ offene Quader, welche M überdecken. Dann gibt es glatte Funktionen $z_1, \dots, z_k: \mathbb{R}^n \rightarrow [0, 1]$ mit folgenden Eigenschaften:*

- (1) auf M gilt $z_1 + \dots + z_k = 1$,
- (2) für alle $i \in \{1, \dots, k\}$ gilt $\text{Träger}(z_i) \subset U_i$.⁵⁴

Die Aussage des Lemmas hat eine gewisse Ähnlichkeit zu der Aussage von Satz 3.17. Allerdings müssen wir im jetzigen Fall glatte Funktionen und nicht nur messbare Funktionen zu finden.

Beweis. Für einen offenen Quader

$$U = (a_1, b_1) \times \dots \times (a_n, b_n)$$

sowie ein $\epsilon > 0$ schreiben wir

$$U(\epsilon) = (a_1 + \epsilon, b_1 - \epsilon) \times \dots \times (a_n + \epsilon, b_n - \epsilon).$$

Der Beweis der folgenden Behauptung ist eine elementare, freiwillige Übungsaufgabe.

Behauptung. Es sei U ein offener Quader und es seien $\epsilon > \eta > 0$. Dann gibt es eine glatte Funktion $g: \mathbb{R}^n \rightarrow [0, 1]$ mit folgenden Eigenschaften:

- (1) Für alle $x \in U(\epsilon)$ ist $g(x) = 1$.
- (2) Für alle $x \notin U(\eta)$ ist $g(x) = 0$.

Die Aussage der Behauptung wird für den Fall $n = 1$ in Abbildung 42 skizziert.

Es sei nun $M \subset \mathbb{R}^n$ eine kompakte Teilmenge und es seien $U_1, \dots, U_k$ offene Quader, welche M überdecken. Nachdem die offenen Quader $U_1, \dots, U_k$ die kompakte Menge M überdecken gibt es ein $\epsilon > 0$, so dass auch noch die offenen Quader $U_1(\epsilon), \dots, U_k(\epsilon)$ die kompakte Menge M überdecken. (Die Existenz von solch einem $\epsilon > 0$ ist eine freiwillige,

⁵⁴Wir hatten den Träger einer Funktion schon in Analysis III eingeführt. Zur Erinnerung, der *Träger* einer Funktion $f: M \rightarrow \mathbb{R}$ auf einer Teilmenge $M \subset \mathbb{R}^n$ ist definiert als die Menge

$$\text{Träger}(f) := \overline{\{x \in M \mid f(x) \neq 0\}}.$$

Hierbei bilden wir den Abschluß in M . Beispielsweise ist der Träger der Sinusfunktion $\sin: \mathbb{R} \rightarrow \mathbb{R}$ ganz $\mathbb{R}$.

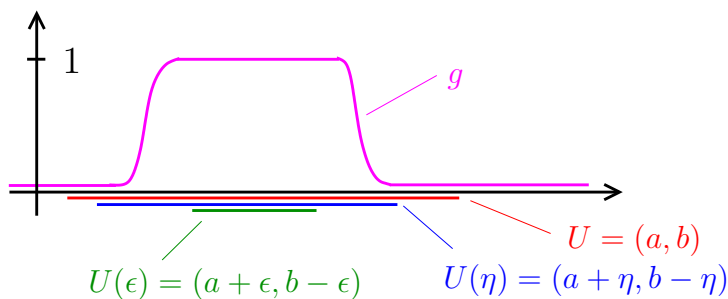


ABBILDUNG 42.

etwas weniger elementare Übungsaufgabe.) Wir wenden nun die Behauptung auf U_i , $i = 1, \dots, k$, $\eta = \frac{1}{2}\epsilon$ und ϵ selber an und erhalten glatte Funktionen $g_1, \dots, g_k$. Wir setzen jetzt

$$z_1 := g_1$$

und wir setzen iterativ für $i = 2, \dots, k$

$$z_i := g_i \cdot (1 - z_1 - \dots - z_{i-1}).$$

Es ist nun leicht zu sehen, dass $z_1, \dots, z_k$ die gewünschten Eigenschaften besitzen. $\square$

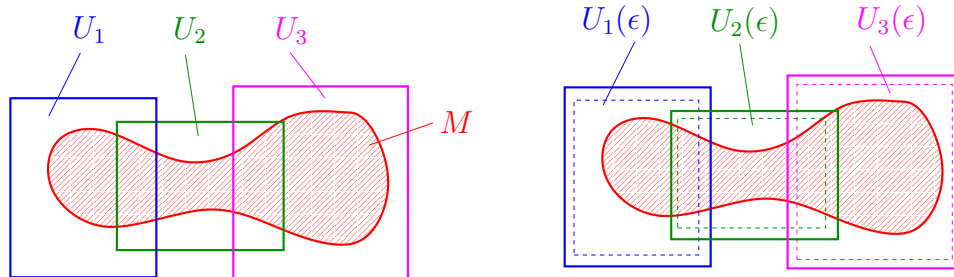


ABBILDUNG 43. Skizze zum Beweis von Lemma 6.11.

Wir sind jetzt in der Lage den Gaußschen Integralsatz zu beweisen.

Beweis vom Gaußschen Integralsatz. Es sei $M \subset \mathbb{R}^n$ eine kompakte n -dimensionale Untermannigfaltigkeit und es sei $F: M \rightarrow \mathbb{R}^n$ ein stetig differenzierbares Vektorfeld. Wir wollen zuerst folgende Behauptung beweisen:

Behauptung. Es gibt Quader $Q_1, \dots, Q_k$ und $l \in \{1, \dots, k\}$ mit folgenden Eigenschaften:

- (1) es ist $M \subset \mathring{Q}_1 \cup \dots \cup \mathring{Q}_k$,
- (2) für $i = 1, \dots, l$ ist $Q_i \subset M \setminus \partial M$,
- (3) für $i = l + 1, \dots, k$ gibt es eine Permutationsmatrix P_i , so dass $P_i \cdot (Q_i \cap M)$ ein verschobener Stempel mit Stempelfläche $P_i \cdot (\mathring{Q}_i \cap \partial M)$ ist.

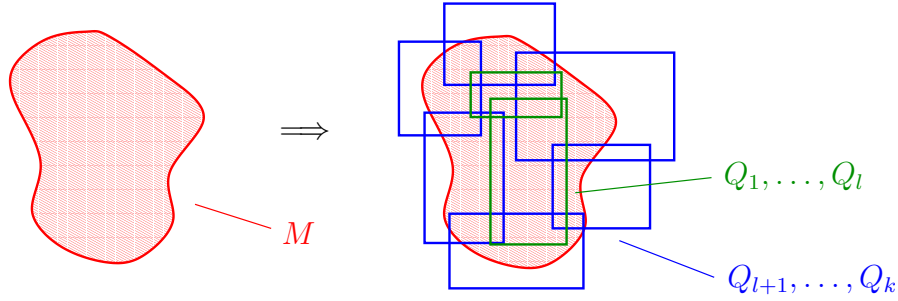


ABBILDUNG 44. Schematisches Bild der Quader $Q_1, \dots, Q_k$.

Wir beweisen nun zuerst die Behauptung. Es sei $X \in M$ beliebig.

- (A) Wenn $X \in M \setminus \partial M$ liegt, dann gibt es offensichtlich einen Quader Q_X mit $X \in \mathring{Q}_X$, so dass $Q_X \subset M \setminus \partial M$.
- (B) Es sei nun $X \in \partial M$. Es folgt aus Satz 4.4, dass es eine Permutationsmatrix P_X und einen Quader Q_X mit $X \in \mathring{Q}_X$ gibt, so dass $P_X \cdot (Q_X \cap M)$ ein verschobener Stempel ist mit Stempelfläche $P_X \cdot (\mathring{Q}_X \cap \partial M)$.

Nachdem die offenen Menge $\{\mathring{Q}_X\}_{X \in M}$ die Untermannigfaltigkeit M überdecken, und nachdem M kompakt ist, gibt es endliche viele $X_1, \dots, X_k \in M$, so dass $M \subset \mathring{Q}_{X_1} \cup \dots \cup \mathring{Q}_{X_k}$. Die dazugehörigen Quader besitzen nun (möglicherweise nach einer Umnummerierung) die gewünschten Eigenschaften. Wir haben damit die Behauptung bewiesen.

Nach Lemma 6.11 gibt es glatte Funktionen $z_1, \dots, z_k: M \rightarrow [0, 1]$ mit folgenden Eigenschaften:

- (1) $z_1 + \dots + z_k = 1$,
- (2) für alle $i \in \{1, \dots, k\}$ gilt $\text{Träger}(z_i) \subset \mathring{Q}_i$.

Dann gilt

$$\begin{array}{ll}
 \text{denn } z_1 + \dots + z_k = 1 & \text{denn } \text{Träger}(z_i) \subset \mathring{Q}_i \\
 \downarrow & \downarrow \\
 \int_M \text{div } F = \int_M \text{div} \left(\sum_{i=1}^k z_i F \right) = \sum_{i=1}^k \int_M \text{div}(z_i F) & = \sum_{i=1}^k \int_{M \cap Q_i} \text{div}(z_i F), \\
 \sum_{i=1}^l \underbrace{\int_{M \cap Q_i} \text{div}(z_i F)}_{= 0, \text{ nach Satz 6.7}} + \sum_{i=l+1}^k \int_{M \cap Q_i} \text{div}(z_i F) & = \sum_{i=l+1}^k \int_{\partial M \cap \mathring{Q}_i} (z_i F) \cdot \nu \\
 \uparrow & \uparrow \text{Satz 6.9} \\
 \text{denn } \text{Träger}(z_i) \subset \mathring{Q}_i & \text{denn } z_{l+1} + \dots + z_k = 1 \text{ auf } \partial M \\
 = \sum_{i=l+1}^k \int_{\partial M} (z_i F) \cdot \nu = \int_{\partial M} \left(\sum_{i=l+1}^k z_i F \right) \cdot \nu & = \int_{\partial M} F \cdot \nu.
 \end{array}$$

Wir haben damit den Gaußschen Integralsatz bewiesen. $\square$

6.7. Der Gaußsche Integralsatzes für Untermannigfaltigkeiten mit Kanten. In diesem Kapitel formulieren wir den Gaußschen Integralsatzes für Untermannigfaltigkeiten mit Kanten. Dieser ist sehr hilfreich für Berechnungen, aber wir werden diesen im weiteren Verlauf der Vorlesung nicht mehr verwenden. Wir erlauben uns daher gewisse Freiheiten in den Formulierungen und wir werden den Beweis nur sehr schemenhaft skizzieren.

Wir sagen $M \subset \mathbb{R}^n$ ist eine n -dimensionale *Untermannigfaltigkeiten mit Kanten*, wenn es Untermannigfaltigkeiten $M_0, \dots, M_{n-2}$ mit folgenden Eigenschaften gibt:

- (1) Für jedes $k \in \{1, \dots, n-2\}$ ist M_k eine k -dimensionale Untermannigfaltigkeit.
- (2) Für jedes $k \in \{1, \dots, n-2\}$ gilt $M_k \subset \partial M$.⁵⁵
- (3) Die Untermannigfaltigkeiten $M_0, \dots, M_{n-2}$ sind disjunkt.
- (4) Die Menge $M \setminus (M_0 \cup \dots \cup M_{n-2})$ ist eine n -dimensionale Untermannigfaltigkeit.

Wir bezeichnen dann $M_0 \cup \dots \cup M_{n-2}$ als die *Kanten von M* . Zudem bezeichnen wir den Rand von $M \setminus (M_0 \cup \dots \cup M_{n-2})$ als den Rand ∂M von M .

Beispiel.

- (1) Die Halbkugel

$$M = \{(x, y, z) \mid x^2 + y^2 + z^2 \leq 1 \text{ und } z \geq 0\}$$

ist eine Untermannigfaltigkeit mit der Kante

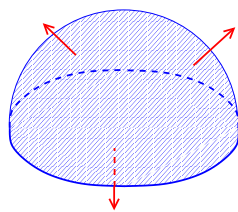
$$\{(x, y, z) \mid x^2 + y^2 + z^2 = 1 \text{ und } z = 0\}.$$

- (2) Der Würfel $M = [-1, 1]^3$ ist eine Untermannigfaltigkeit mit Kanten, hierbei ist

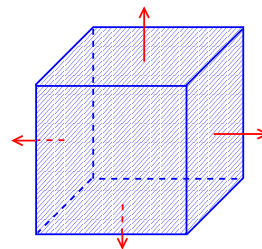
$$M_0 = \bigcup_{\epsilon, \eta \in \{0,1\}} \{\epsilon, \eta\}$$

und

$$M_1 = \bigcup_{\epsilon, \eta \in \{0,1\}} ((0, 1) \times \{\epsilon\} \times \{\eta\} \cup \{\epsilon\} \times (0, 1) \times \{\eta\} \cup \{\epsilon\} \times \{\eta\} \times (0, 1)).$$



die abgeschlossene Halbkugel
 $\{(x, y, z) \in \mathbb{R}^3 \mid x^2 + y^2 + z^2 \leq 1 \text{ und } z \geq 0\}$



der Würfel $[0, 1]^3$

ABBILDUNG 45. Beispiele von Untermannigfaltigkeiten mit Kanten.

⁵⁵In diesem Fall meinen wir mit ∂M ausnahmsweise den topologischen Rand von $M \subset \mathbb{R}^n$.

Satz 6.12. (Gaußscher Integralsatz für Untermannigfaltigkeiten mit Kanten) *Es sei $M \subset \mathbb{R}^n$ eine kompakte n -dimensionale Untermannigfaltigkeit mit Kanten und es sei $F: M \rightarrow \mathbb{R}^n$ ein stetig differenzierbares Vektorfeld. Dann gilt*

$$\text{Fluss von } F \text{ durch } M = \int_M \operatorname{div} F.$$

Wie oben schon angekündigt ist dieser Satz mit Vorsicht zu genießen. In den meisten Anwendungen kann man diesen problemlos anwenden. Es ist aber möglich, dass die Voraussetzungen an “Untermannigfaltigkeiten mit Kanten” zu schwach sind, um diesen in der angegebenen Allgemeinheit zu beweisen.

Beweisskizze. Es sei $M \subset \mathbb{R}^n$ also eine kompakte n -dimensionale Untermannigfaltigkeit mit Kanten und es sei $F: M \rightarrow \mathbb{R}^n$ ein stetig differenzierbares Vektorfeld. Man kann nun zeigen⁵⁶⁵⁷, dass es eine Folge von kompakten n -dimensionalen Untermannigfaltigkeiten W_k , $k \in \mathbb{N}$ gibt mit folgenden Eigenschaften:

- (1) für alle k ist $W_k \subset M$,
- (2) $\lim_{k \rightarrow \infty} \operatorname{Vol}(M \setminus W_k) = 0$,
- (3) $\lim_{k \rightarrow \infty} \operatorname{Vol}(\partial M \triangle \partial W_k) = 0$.⁵⁸

Wir skizzieren die Folge W_k in Abbildung 46. Wir bezeichnen mit ν das äußere Einheits-

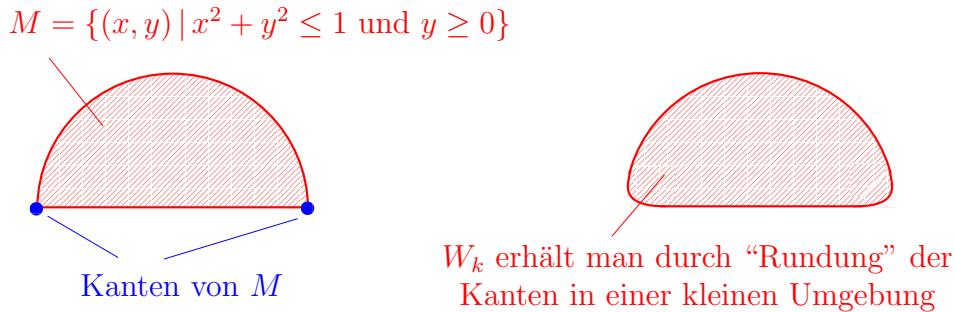


ABBILDUNG 46. Approximation einer Untermannigfaltigkeiten mit Kanten.

⁵⁶Dies ist der technisch schwierige Schritt den wir auslassen.

⁵⁷Wenn man eine Aussage nicht beweist, dann besteht immer die Gefahr, dass sie nicht richtig ist. Auch in diesem Fall kann es sein, dass man eventuell noch stärkere Voraussetzungen braucht, um pathologische Spezialfälle zu vermeiden.

⁵⁸Hier bezeichnen wir, wie in Analysis III, mit $A \triangle B = (A \setminus B) \cup (B \setminus A)$ die symmetrische Differenz von zwei Teilmengen A und B von $\mathbb{R}^n$.

normalenfeld von M und von allen W_k 's. Dann ist

$$\begin{array}{ccccccc} & & \text{da } \lim_{k \rightarrow \infty} \text{Vol}(\partial M \triangle \partial W_k) = 0 & & \text{da } \lim_{k \rightarrow \infty} \text{Vol}(M \setminus W_k) = 0 & & \\ & & \downarrow & & & & \downarrow \\ \text{Fluss von } F \text{ durch } M & = & \int_{\partial M} F \cdot \nu & = & \lim_{k \rightarrow \infty} \int_{\partial W_k} F \cdot \nu & = & \lim_{k \rightarrow \infty} \int_{W_k} \text{div } F & = & \int_M \text{div } F. \\ & & & & \uparrow & & & & \\ & & & & \text{Gaußscher Integralsatz 6.6} & & & & \square \end{array}$$

Beispiel. Es sei

$$A = \{(x, y, z) \in \mathbb{R}^3 \mid x^2 + y^2 + z^2 = 1 \text{ und } z \geq 0\}$$

die obere Hälfte der Kugel S^2 . Wir betrachten A mit dem Einheitsnormalenfeld, welches “nach außen” zeigt. Wir wollen den Fluss von dem konstanten Vektorfeld $F = (1, 0, 3)$ durch A bestimmen. Wir setzen dazu

$$B = \{(x, y, z) \in \mathbb{R}^3 \mid x^2 + y^2 \leq 1 \text{ und } z = 0\}$$

sowie

$$M = \{(x, y, z) \in \mathbb{R}^3 \mid x^2 + y^2 + z^2 \leq 1 \text{ und } z \geq 0\}.$$

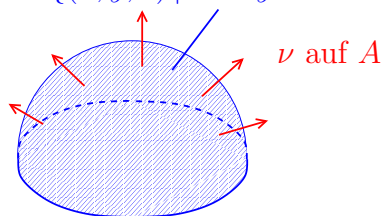
Wir betrachten B mit dem Einheitsnormalenfeld welches auch “nach oben” zeigt. Dann ist

$$\begin{aligned} \text{Fluss von } F \text{ durch } A &= \text{Fluss von } F \text{ durch } M - \text{Fluss von } F \text{ durch } B + \text{Fluss von } F \text{ durch } B \\ &= \underbrace{\text{Fluss von } F \text{ durch } \partial M}_{= 0, \text{ nach dem Gaußschen Integralsatz 6.12, da } \text{div}(F) \equiv 0} + \text{Fluss von } F \text{ durch } B \\ &= \text{Fluss von } F \text{ durch } B = (*) \end{aligned}$$

Dieser Fluss kann nun leicht explizit ausgerechnet werden, in der Tat ist

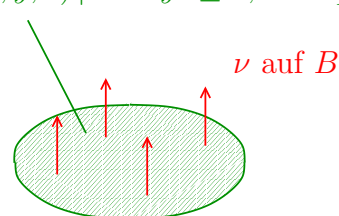
$$(*) = \int_B F \cdot \nu = \int_{\varphi=0}^{\varphi=2\pi} \int_{r=0}^{r=1} (1, 0, 3) \cdot (0, 0, 1) = 6\pi.$$

$$A = \{(x, y, z) \mid x^2 + y^2 + z^2 = 1, z \geq 0\}$$



Fluss von F durch A

$$B = \{(x, y, z) \mid x^2 + y^2 \leq 1, z = 0\}$$



Fluss von F durch B

wenn $\text{div}(F) \equiv 0$

$\downarrow$
=

ABBILDUNG 47.

6.8. **Ausblick.** Das nächste große Ziel wird sein den allgemeinen Satz von Stokes zu formulieren und zu beweisen. Der klassische Satz von Stokes besagt, dass für eine orientierte kompakte Fläche $(M, \vec{\nu})$ in $\mathbb{R}^3$ und ein stetig differenzierbares Vektorfeld $\vec{F} = (F_x, F_y, F_z)$ auf M die Gleichheit⁵⁹

$$\int_M \operatorname{rot}(\vec{F}) \cdot \vec{\nu} = \int_{\gamma=\partial M} \vec{F} \cdot d\vec{\gamma}$$

gilt. Hierbei ist die *Rotation* von $\vec{F}$ definiert als das Vektorfeld

$$\operatorname{rot}(\vec{F}) := \begin{pmatrix} \frac{\partial}{\partial x} \\ \frac{\partial}{\partial y} \\ \frac{\partial}{\partial z} \end{pmatrix} \times \begin{pmatrix} F_x \\ F_y \\ F_z \end{pmatrix} = \begin{pmatrix} \frac{\partial F_z}{\partial y} - \frac{\partial F_y}{\partial z} \\ -\frac{\partial F_z}{\partial x} + \frac{\partial F_x}{\partial z} \\ \frac{\partial F_y}{\partial x} - \frac{\partial F_x}{\partial y} \end{pmatrix}.$$

Unser Ziel ist es die “korrekte” Verallgemeinerung dieser Aussage auf den höherdimensionalen Fall zu finden, und diese dann für alle Dimensionen zu beweisen. Wir müssen dazu erst einmal relativ viele neue mathematische Objekte einführen. Der allgemeine Satz von Stokes (welcher den klassischen Fall einschließt) wird dann jedoch eine sehr kurze und elegante Formulierung besitzen, und der Beweis wird am Ende erstaunlich kurz sein.

⁵⁹Auf der rechten Seite integrieren wir über das Skalarprodukt von F mit dem Ableitungsvektor einer Parametrisierung der Kurve $\gamma = \partial M$.

7. KATEGORIEN UND FUNKTOREN

7.1. Induzierte Abbildungen auf Tangentialräumen. Es sei M eine Untermannigfaltigkeit und es sei $P \in M$. Wir sagen zwei Kurven $\alpha: I \rightarrow M$ und $\beta: J \rightarrow M$ durch P sind *äquivalent*, wenn $\alpha'(0) = \beta'(0)$. Für eine Kurve α in M durch P bezeichnen wir mit $[\alpha]$ die Äquivalenzklasse von α . Die Abbildung

$$\begin{aligned} \text{Menge der Äquivalenzklassen} & \rightarrow T_P M \\ \text{von Kurven in } M \text{ durch } P & \\ [\gamma] & \rightarrow \gamma'(0) \end{aligned}$$

ist offensichtlich eine Bijektion. Wir verwenden diese Bijektion im Folgenden um $T_P M$ mit der Menge der Äquivalenzklassen von Kurven in M durch P zu identifizieren.⁶⁰

Dieser Gesichtspunkt scheint auf den ersten Blick etwas verwirrend, aber er erleichtert oft die Notation. Es sei beispielsweise $f: M \rightarrow N$ eine stetig differenzierbare Abbildung zwischen zwei Untermannigfaltigkeiten (nicht notwendigerweise von der gleichen Dimension) und es sei $P \in M$. Wir betrachten die Abbildung⁶¹

$$\begin{aligned} Df(P): \quad T_P M & \rightarrow T_{f(P)} N \\ [\gamma: I \rightarrow M] & \rightarrow [f \circ \gamma: I \rightarrow N]. \end{aligned}$$

Wir bezeichnen die Abbildung $Df(P): T_P M \rightarrow T_{f(P)} N$ oft auch nur mit f_* und wir nennen $f_* = Df(P)$ die *induzierte Abbildung*. Bevor wir fortfahren müssen wir allerdings erst noch folgendes Lemma beweisen.

Lemma 7.1. *Die obige Abbildung $f_* = Df(P): T_P M \rightarrow T_{f(P)} N$ ist wohl-definiert.*

Beweis. Wir bezeichnen mit m die Dimension von M . Es seien $\gamma: I \rightarrow M$ und $\delta: J \rightarrow M$ zwei äquivalente Kurven durch P . O.B.d.A. können wir annehmen, dass $I = J$, und dass es eine Karte $\Phi: U \rightarrow V$ um P gibt, so dass das Bild von γ und δ in $U \cap M$ liegt. Wir schreiben $g := f \circ \Phi^{-1}$ und wir setzen $Q := \Phi(P)$. Wir betrachten nun also folgende Abbildungen

$$\begin{array}{ccccc} I & \xrightarrow[\gamma]{} & M \cap U & \xrightarrow{f} & N \\ & & \downarrow \Phi & \nearrow g = f \circ \Phi^{-1} & \\ & & E_m \cap V & & \end{array}$$

Wie immer fassen wir $E_m \cap V$ als offene Teilmenge von $\mathbb{R}^m$ auf. Dann ist

$$\begin{array}{ccccccc} (f \circ \gamma)'(0) & = & (g \circ \Phi \circ \gamma)'(0) & = & D_Q g(D_P \Phi(\gamma'(0))) & = & D_Q g(D_P \Phi(\delta'(0))) = (f \circ \delta)'(0). \\ \uparrow & & \uparrow & & \uparrow & & \uparrow \\ \text{denn } g = f \circ \Phi^{-1} & & \text{Kettenregel} & & \text{denn } \gamma'(0) = \delta'(0) & & \text{das gleiche Argument} \\ & & & & & & \text{rückwärts} \end{array}$$

Wir haben also gezeigt, dass $f \circ \gamma$ und $f \circ \delta$ ebenfalls äquivalent sind. $\square$

⁶⁰Wir verwenden insbesondere diese Bijektion um die Vektorraumstruktur auf der Menge der Äquivalenzklassen einzuführen.

⁶¹Mit der ursprünglichen Definition des Tangentialraumes wäre diese Abbildung deutlich umständlicher hinzuschreiben.

Bemerkung. Für eine offene Teilmenge $U \subset \mathbb{R}^n$ und eine stetig differenzierbare Abbildung $f: U \rightarrow \mathbb{R}^m$ erhalten wir für jedes $P \in U$ eine Abbildung

$$f_* = Df(P): T_P U = \mathbb{R}^n \rightarrow T_{f(P)} \mathbb{R}^m = \mathbb{R}^m.$$

Wie die Notation schon suggeriert ist dies gerade die Abbildung, welche dem üblichen Differential $Df(P)$ entspricht.⁶² Insbesondere ist in diesem Fall die Abbildung $f_* = Df(P)$ eine lineare Abbildung.

Satz 7.2.

- (1) Es seien $f: L \rightarrow M$ und $g: M \rightarrow N$ stetig differenzierbare Abbildungen zwischen Untermannigfaltigkeiten und es sei $P \in L$. Dann ist⁶³

$$D(g \circ f)(P) = Dg(f(P)) \circ Df(P).$$

Mit anderen Worten, es gilt

$$(g \circ f)_* = g_* \circ f_*.$$

- (2) Für jede Untermannigfaltigkeit M und jeden Punkt $P \in M$ gilt

$$(D \text{id}_M)(P) = \text{id}_{T_P M},$$

oder noch kürzer ausgedrückt, es ist

$$(\text{id}_M)_* = \text{id}_{T_P M}.$$

Beweis. Wir beweisen zuerst die erste Aussage. Es sei also $[\gamma: I \rightarrow L] \in T_P L$. Dann ist

$$\begin{aligned} D(g \circ f)(P)([\gamma]) &= [(g \circ f) \circ \gamma] = [g \circ (f \circ \gamma)] = Dg(f(P))([f \circ \gamma]) \\ &= (Dg(f(P)) \circ Df(P))([\gamma]). \end{aligned}$$

In der alternativen Schreibweise ist der Beweis noch kürzer, denn es ist

$$(g \circ f)_*([\gamma]) = [(g \circ f) \circ \gamma] = [g \circ (f \circ \gamma)] = g_*([f \circ \gamma]) = g_*(f_*([\gamma])).$$

Die zweite Aussage des Satzes ist trivial. □

Korollar 7.3. Es sei $f: M \rightarrow N$ ein Diffeomorphismus zwischen zwei Untermannigfaltigkeiten. Für alle $P \in M$ gilt

$$Df^{-1}(f(P)) = (Df(P))^{-1}.$$

Mit anderen Worten, es ist

$$(f^{-1})_* = (f_*)^{-1}.$$

Insbesondere ist $Df(P) = f_*: T_P M \rightarrow T_{f(P)} N$ eine Bijektion.

Beweis. Es ist

$$\begin{array}{ccc} f_* \circ (f^{-1})_* & = & (f \circ f^{-1})_* = \text{id}_* = \text{id} \\ \uparrow & & \uparrow \\ \text{Satz 7.2 (1)} & & \text{Satz 7.2 (2)} \end{array}$$

⁶²Warum ist dies der Fall?

⁶³Dies ist also eine Gleichheit von Abbildung $T_P L \rightarrow T_{g(f(P))} N$.

Ganz analog zeigt man auch, dass $(f^{-1})_* \circ f_* = \text{id}$, also sind $(f^{-1})_*$ und f_* zueinander inverse Abbildungen. Insbesondere ist f_* eine Bijektion. $\square$

Wir beschließen das Kapitel mit folgendem Satz.

Satz 7.4. *Es sei $f: M \rightarrow N$ eine stetig differenzierbare Abbildung zwischen zwei Untermannigfaltigkeiten. Für jedes $P \in M$ ist*

$$f_* = Df(P): T_P M \rightarrow T_{f(P)} N$$

eine lineare Abbildung.

Beweis. Es sei also $f: M \rightarrow N$ eine stetig differenzierbare Abbildung zwischen einer m -dimensionalen Untermannigfaltigkeit M und einer n -dimensionalen Untermannigfaltigkeit N . Es sei nun $P \in M$. Wir wählen eine Karte $\Phi: U \rightarrow V$ um P und wir wählen eine Karte $\Theta: X \rightarrow Y$ um $Q := f(P)$ mit $f(U) \subset X$ ⁶⁴. Wir bezeichnen mit $\Psi := \Phi^{-1}$ und $\Omega = \Theta^{-1}$ die Umkehrabbildungen. Wir betrachten die folgenden Abbildungen

$$\begin{array}{ccc} U \cap M & \xrightarrow{f} & X \cap N \\ \Phi \downarrow \nearrow \Psi & & \Omega \nearrow \downarrow \Theta \\ V \cap E_m & \xrightarrow{\Theta \circ f \circ \Psi} & Y \cap E_n. \end{array}$$

Dies ist ein kommutatives Diagramm von Abbildungen.⁶⁵ Wir erhalten das entsprechende Diagramm von induzierten Abbildungen

$$\begin{array}{ccc} T_P M & \xrightarrow{Df} & T_Q N \\ D\Phi \downarrow \nearrow D\Psi & & D\Omega \nearrow \downarrow D\Theta \\ \mathbb{R}^m & \xrightarrow{D(\Theta \circ f \circ \Psi)} & \mathbb{R}^n. \end{array}$$

(Der Übersicht halber haben wir hierbei bei den induzierten Abbildungen die Punkte weggelassen, an denen wir diese betrachten.) Es folgt aus Satz 7.2, dass auch dieses Diagramm kommutiert.⁶⁶ Es folgt aus der Bemerkung auf Seite 81, dass die Abbildungen $D\Psi$ und $D\Theta$ sowie $D(\Theta \circ f \circ \Psi)$ linear sind. Zudem folgt aus Korollar 7.3, dass $D\Psi = (D\Phi)^{-1}$ und $D\Omega = (D\Theta)^{-1}$. Insbesondere sind die vertikalen Abbildungen links beziehungsweise rechts zueinander inverse Homomorphismen von Vektorräumen. Es folgt nun, dass

$$Df = D\Omega \circ D(\Theta \circ f \circ \Psi) \circ D\Phi$$

⁶⁴Warum kann diese letzte Bedingung erfüllt werden?

⁶⁵Wir sagen ein Diagramm von Abbildungen ist kommutativ, wenn es egal ist, in welcher Reihenfolge wir die Abbildungen durchlaufen. In diesem Fall bedeutet das, dass die Verknüpfung der Abbildungen oben und rechts, die gleiche Abbildung ergibt wie die Verknüpfung der Abbildungen links und unten.

⁶⁶In der Tat, denn es ist

$$\begin{array}{ccccccc} D(\Theta \circ f \circ \Psi) \circ D\Phi & = & D(\Theta \circ f \circ \Psi \circ D\Phi) & = & D(\Theta \circ f) & = & D\Theta \circ Df. \\ & \uparrow & & & & \uparrow & \\ & \text{Satz 7.2} & & & & \text{Satz 7.2} & \end{array}$$

die Verknüpfung von drei linearen Abbildung ist. Also ist auch Df selber eine lineare Abbildung. $\square$

7.2. Duale Vektorräume. Es sei V ein reeller Vektorraum. Wir bezeichnen mit

$$V^* := \text{Hom}_{\mathbb{R}}(V, \mathbb{R})$$

den zu V *dualen Vektorraum*. Es sei nun $v_1, \dots, v_k$ eine Basis von V . Jeder Vektor $w \in V$ kann also eindeutig als Linearkombination

$$w = l_1 v_1 + \dots + l_k v_k \text{ mit } l_1, \dots, l_k \in \mathbb{R}$$

geschrieben werden. Für $i = 1, \dots, k$ betrachten wir nun die Abbildung

$$\begin{aligned} v_i^*: V &\rightarrow \mathbb{R} \\ w &\mapsto v_i^*(w) := \begin{array}{l} \text{Koeffizient von } v_i \text{ in der eindeutigen Darstellung} \\ \text{von } w \text{ als Linearkombination von } v_1, \dots, v_k. \end{array} \end{aligned}$$

Anders ausgedrückt, die Abbildungen $v_1^*, \dots, v_k^*: V \rightarrow \mathbb{R}$ sind die eindeutig bestimmten Abbildungen, welche die Eigenschaft besitzen, dass

$$w = v_1^*(w) \cdot v_1 + \dots + v_k^*(w) \cdot v_k \text{ für alle } w \in V.$$

Man kann nun leicht zeigen, dass $v_1^*, \dots, v_k^*$ lineare Abbildungen sind. Zudem gilt folgendes elementare Lemma, welches in der Linearen Algebra bewiesen wurde.

Lemma 7.5. *Es sei $v_1, \dots, v_k$ eine Basis von einem reellen Vektorraum V . Dann ist $v_1^*, \dots, v_k^*$ eine Basis von V^* .*

Es folgt aus dem Lemma, dass V^* ein reeller Vektorraum der Dimension $k = \dim(V)$ ist. Wir nennen $v_1^*, \dots, v_k^* \in V^*$ die zu $v_1, \dots, v_k$ *duale Basis*.

Bemerkung. Wenn V endlich dimensional ist, dann gibt es zu jeder *Basis* von V eine eindeutig bestimmte *duale Basis* von V^* . Allerdings kann man nicht jedem einzelnen *Vektor* in V einen eindeutig bestimmten *Vektor* in V^* zuordnen.

Es sei zum Abschluß des Kapitels $f: U \rightarrow V$ eine lineare Abbildung zwischen reellen Vektorräumen. Wir erhalten dann eine Abbildung

$$\begin{aligned} f^*: V^* &\rightarrow U^* \\ (\varphi: V \rightarrow \mathbb{R}) &\mapsto (\varphi \circ f: U \rightarrow \mathbb{R}), \end{aligned}$$

welche wir als die *induzierte Abbildung* bezeichnen. Die Abbildung f^* geht hierbei in die “umgekehrte Richtung”.

Das folgende elementare Lemma faßt einige elementare Eigenschaften von induzierten Abbildungen zusammen.

Lemma 7.6. *Es seien U, V und W reelle Vektorräume.*

(1) Es seien $f: U \rightarrow V$ und $g: V \rightarrow W$ lineare Abbildungen. Dann gilt ⁶⁷

$$(g \circ f)^* = f^* \circ g^*.$$

(2) Für die Identitätsabbildung $\text{id}: V \rightarrow V$ ist auch $\text{id}^*: V^* \rightarrow V^*$ die Identitätsabbildung, d.h. es gilt

$$\text{id}^* = \text{id}.$$

Korollar 7.7. Es sei $f: V \rightarrow W$ ein Isomorphismus zwischen reellen Vektorräumen. Dann ist auch $f^*: W^* \rightarrow V^*$ ein Isomorphismus mit $(f^*)^{-1} = (f^{-1})^*$.

Der Beweis von Korollar 7.7 ist ganz ähnlich dem Beweis von Korollar 7.3.

Beweis. Es sei $f: V \rightarrow W$ ein Isomorphismus zwischen reellen Vektorräumen. Es ist

$$\begin{array}{ccc} f^* \circ (f^{-1})^* & = & (f^{-1} \circ f)^* = \text{id}^* = \text{id} \\ \uparrow & & \uparrow \\ \text{Lemma 7.6 (1)} & & \text{Lemma 7.6 (2)} \end{array}$$

Ganz analog zeigt man auch, dass $(f^{-1})^* \circ f^* = \text{id}$, also sind $(f^{-1})^*$ und f^* zueinander inverse Abbildungen. Insbesondere ist f^* ein Isomorphismus. $\square$

7.3. Kategorien. Die Aussagen von Satz 7.2 und von Lemma 7.6 besitzen eine große formale Ähnlichkeit und wir hatten diese Ähnlichkeit benutzt, um für Korollar 7.3 und Korollar 7.7 fast identische Beweise anzugeben.

In diesem kurzen Kapitel führen wir nun “Kategorien” und “Funktoren” ein. Diese Begriffe geben uns eine Sprache um die formalen Gemeinsamkeiten der Aussagen der vorherigen beiden Kapitel herauszuarbeiten. Die Begriffe “Kategorie” und “Funktork” werden in allen Bereichen der reinen Mathematik verwendet und wir wollen uns daher in diesem Kapitel nicht nur auf Beispiele aus der Analysis beschränken.

Definition. Eine *Kategorie* $\mathcal{C}$ besteht aus folgenden Daten:

- (1) Einer Klasse⁶⁸ $\text{Ob}(\mathcal{C})$ von Objekten, welche die *Objekte der Kategorie* genannt werden,
- (2) zu jedem Paar (X, Y) von Objekten gibt es eine Menge $\text{Mor}(X, Y)$,

⁶⁷Mit anderen Worten, das folgende Diagramm von Abbildungen kommutiert

$$\begin{array}{ccc} & V^* & \\ g^* \nearrow & & \searrow f^* \\ W^* & \xrightarrow{(g \circ f)^*} & U^* \end{array}$$

⁶⁸Ich gehe hier jetzt nicht weiter auf das Wort “Klasse” ein. Anschaulich gesprochen ist eine “Klasse” eine Verallgemeinerung des Begriffs der “Menge”. Zum Beispiel macht es keinen Sinn von der “Menge aller Gruppen” zu reden, genau so wenig wie es Sinn macht von der “Menge aller Mengen” zu reden. Aber es macht Sinn von der “Klasse aller Gruppen” zu reden, siehe beispielsweise

https://de.wikipedia.org/wiki/Klasse_Mengenlehre

Wir überlassen diese heikle Geschichte lieber den Logikern.

(3) zu je drei Objekten X, Y und Z gibt es eine Abbildung

$$\begin{aligned} \text{Mor}(X, Y) \times \text{Mor}(Y, Z) &\rightarrow \text{Mor}(X, Z) \\ (f, g) &\mapsto g \circ f \end{aligned}$$

so dass folgende Axiome erfüllt sind:

(K1) (Assoziativität): Es seien $f \in \text{Mor}(W, X)$, $g \in \text{Mor}(X, Y)$ und $h \in \text{Mor}(Y, Z)$, dann gilt

$$(h \circ g) \circ f = h \circ (g \circ f).$$

(K2) (Identität): Zu jedem Objekt X gibt es einen Morphismus $\text{id}_X \in \text{Mor}(X, X)$ mit der Eigenschaft, dass gilt⁶⁹

$$\begin{aligned} \text{id}_X \circ f &= f \quad \text{für alle } Z \in \text{Ob}(\mathcal{C}) \text{ und alle } f \in \text{Mor}(Z, X), \text{ sowie} \\ f \circ \text{id}_X &= f \quad \text{für alle } Y \in \text{Ob}(\mathcal{C}) \text{ und alle } f \in \text{Mor}(X, Y). \end{aligned}$$

Die Abbildung $\text{Mor}(X, Y) \times \text{Mor}(Y, Z) \rightarrow \text{Mor}(X, Z)$ wird, wie die Notation schon suggeriert, die *Verknüpfungsabbildung* genannt.

Beispiele.

(a) Wir nennen die Kategorie $\mathcal{C}$ mit

$$\begin{aligned} \text{Ob}(\mathcal{C}) &:= \text{alle Mengen,} \\ \text{Mor}(X, Y) &:= \text{alle Abbildungen von } X \text{ nach } Y, \end{aligned}$$

mit der üblichen Verknüpfung von Abbildungen die *Kategorie aller Mengen*.

(b) Wir betrachten die Kategorie $\mathcal{C}$ mit

$$\begin{aligned} \text{Ob}(\mathcal{C}) &:= \text{alle Mengen,} \\ \text{Mor}(X, Y) &:= \begin{cases} \{\text{id}_X\}, & \text{wenn } X = Y \\ \emptyset, & \text{wenn } X \neq Y. \end{cases} \end{aligned}$$

Die Verknüpfungsabbildung muss nur definiert werden, wenn $X = Y = Z$, und in diesem Fall definieren wir $\text{id}_X \circ \text{id}_X := \text{id}_X$. Man kann sich nun leicht davon überzeugen, dass dies eine Kategorie definiert.

(c) Wir nennen die Kategorie $\mathcal{C}$ mit

$$\begin{aligned} \text{Ob}(\mathcal{V}) &:= \text{alle reellen Vektorräume,} \\ \text{Mor}(V, W) &:= \text{Hom}_{\mathbb{R}}(V, W) = \text{alle Homomorphismen von } V \text{ nach } W, \end{aligned}$$

mit der üblichen Verknüpfung von Homomorphismen die *Kategorie der reellen Vektorräume*. Ganz analog kann man natürlich auch für jeden Körper $\mathbb{K}$ die Kategorie der $\mathbb{K}$ -Vektorräume einführen.

(d) Wir nennen die Kategorie $\mathcal{G}$ mit

$$\begin{aligned} \text{Ob}(\mathcal{G}) &:= \text{alle Gruppen,} \\ \text{Mor}(G, H) &:= \text{Hom}(G, H) = \text{alle Gruppenhomomorphismen von } G \text{ nach } H, \end{aligned}$$

⁶⁹Ist id_X eindeutig bestimmt, oder nicht?

mit der üblichen Verknüpfung von Gruppenhomomorphismen die *Kategorie der Gruppen*. Ganz analog kann man auch die Kategorie der abelschen Gruppen einführen.

- (e) Wir nennen die Kategorie $\mathcal{T}$ mit

$$\begin{aligned}\mathrm{Ob}(\mathcal{T}) &:= \text{alle topologischen Räume,} \\ \mathrm{Mor}(X, Y) &:= \text{alle stetigen Abbildungen von } X \text{ nach } Y,\end{aligned}$$

mit der üblichen Verknüpfung von Abbildungen die *Kategorie der topologischen Räume*.

- (f) Wir nennen die Kategorie $\mathcal{M}$ mit

$$\begin{aligned}\mathrm{Ob}(\mathcal{M}) &:= \text{alle Untermannigfaltigkeiten,} \\ \mathrm{Mor}(M, N) &:= \text{alle glatten Abbildungen von } M \text{ nach } N,\end{aligned}$$

mit der üblichen Verknüpfung von Abbildungen die *Kategorie der Untermannigfaltigkeiten*.

- (g) Eine *punktierte Untermannigfaltigkeit* ist ein Paar (M, P) , wobei M eine Untermannigfaltigkeit ist und $P \in M$. Wir nennen die Kategorie $\mathcal{P}$ mit

$$\begin{aligned}\mathrm{Ob}(\mathcal{P}) &:= \{\text{alle punktierten Untermannigfaltigkeiten}\}, \\ \mathrm{Mor}((M, P), (N, Q)) &:= \left\{ \begin{array}{l} \text{alle glatten Abbildungen } f \\ \text{von } M \text{ nach } N \text{ mit } f(P) = Q \end{array} \right\}\end{aligned}$$

mit der üblichen Verknüpfung von Abbildungen die *Kategorie der punktierten Untermannigfaltigkeiten*.

Der Name “Morphismus” suggeriert, dass es sich bei Morphismen um Abbildungen handelt. Dies ist aber nicht notwendigerweise der Fall, wie wir im folgenden Beispiel sehen.

- (h) Es sei G eine beliebige Gruppe. Wir betrachten

$$\begin{aligned}\mathrm{Ob}(\mathcal{C}) &:= \{\text{eine Menge mit einem einzigen Element } *\}, \\ \mathrm{Mor}(*, *) &:= G.\end{aligned}$$

Wir definieren dann

$$\begin{aligned}\mathrm{Mor}(*, *) \times \mathrm{Mor}(*, *) &\rightarrow \mathrm{Mor}(*, *) \\ (f, g) &\mapsto g \circ f := gf\end{aligned}$$

mittels der Gruppenstruktur auf $\mathrm{Mor}(*, *) = G$. Die Axiome einer Kategorie folgen dann aus den Gruppenaxiomen von G .

7.4. Funktoren.

Definition. Es seien $\mathcal{C}$ und $\mathcal{D}$ Kategorien. Ein *kovarianter Funktor* $F: \mathcal{C} \rightarrow \mathcal{D}$ besteht aus einer Abbildung

$$F: \mathrm{Ob}(\mathcal{C}) \rightarrow \mathrm{Ob}(\mathcal{D})$$

und für alle $X, Y \in \mathrm{Ob}(\mathcal{C})$ gibt es zudem eine Abbildung

$$\begin{aligned}\mathrm{Mor}(X, Y) &\rightarrow \mathrm{Mor}(F(X), F(Y)) \\ \phi &\mapsto F(\phi),\end{aligned}$$

so dass folgende Axiome erfüllt sind:

(F1) Für alle $\phi \in \text{Mor}(X, Y)$ und $\psi \in \text{Mor}(Y, Z)$, gilt

$$F(\psi \circ \phi) = F(\psi) \circ F(\phi).$$

(F2) Für alle $X \in \text{Ob}(\mathcal{C})$ gilt

$$F(\text{id}_X) = \text{id}_{F(X)}.$$

Konvention. Für einen kovarianten Funktion $F: \mathcal{C} \rightarrow \mathcal{D}$ und einen Morphismus ϕ schreiben wir oft ϕ_* anstatt $F(\phi)$.

Beispiele.

- (1) Es sei $\mathcal{V}$ die Kategorie der reellen Vektorräume und es sei W ein Vektorraum, dann ist

$$\begin{aligned} F: \text{Ob}(\mathcal{V}) &\rightarrow \text{Ob}(\mathcal{V}) \\ V &\mapsto \text{Hom}_{\mathbb{R}}(W, V), \end{aligned}$$

zusammen mit den Abbildungen

$$\begin{aligned} \text{Mor}(U, V) &\mapsto \text{Mor}(F(U), F(V)) \\ (\phi: U \rightarrow V) &\mapsto \left(\begin{array}{ccc} \phi_*: \text{Hom}(W, U) &\rightarrow & \text{Hom}(W, V) \\ & f \mapsto & \phi \circ f \end{array} \right) \end{aligned}$$

ein kovarianter Funktor.⁷⁰ Folgendes Diagramm illustriert die Aussage von Axiom (F2):

$$\begin{array}{ccccc} X & \xrightarrow{\quad \phi \quad} & Y & \xrightarrow{\quad \psi \quad} & Z \\ \downarrow \wr & & \downarrow \wr & & \downarrow \wr \\ \text{Hom}(W, X) & \xrightarrow{\quad \phi_* \quad} & \text{Hom}(W, Y) & \xrightarrow{\quad \psi_* \quad} & \text{Hom}(W, Z) \end{array}$$

$\xrightarrow{\quad \psi \circ \phi \quad}$ (oben)
 $\xrightarrow{\quad (\psi \circ \phi)_* \quad}$ (unten)

- (2) Es sei $\mathcal{P}$ die Kategorie der punktierten Untermannigfaltigkeiten und es sei $\mathcal{V}$ die Kategorie der reellen Vektorräume. Es folgt aus Satz 7.2 zusammen mit Satz 7.4, dass

$$\begin{aligned} F: \text{Ob}(\mathcal{P}) &\rightarrow \text{Ob}(\mathcal{V}) \\ (M, P) &\mapsto T_P M, \end{aligned}$$

zusammen mit den Abbildungen

$$\begin{aligned} \text{Mor}((M, P), (N, Q)) &\mapsto \text{Mor}(T_P M, T_Q N) \\ f &\mapsto f_* = Df(P) \end{aligned}$$

⁷⁰In der Tat, denn sei $\phi \in \text{Mor}(X, Y)$ und $\psi \in \text{Mor}(Y, Z)$ und es sei $f \in \text{Hom}(W, X)$, dann gilt

$$F(\psi \circ \phi)(f) = (\psi \circ \phi)(f) = \psi \circ (\phi \circ f) = \psi \circ F(\phi)(f) = F(\psi)(F(\phi)(f)) = (F(\psi) \circ F(\phi))(f).$$

ein kovarianter Funktor von der Kategorie der punktierten Untermannigfaltigkeiten zu der Kategorie der Vektorräume ist.

- (3) Es sei $\mathcal{A}$ die Kategorie der abelschen Gruppen, es sei p eine Primzahl und es sei $\mathcal{F}$ die Kategorie der $\mathbb{F}_p$ -Vektorräume. Dann ist ⁷¹

$$\begin{aligned} F: \operatorname{Ob}(\mathcal{A}) &\rightarrow \operatorname{Ob}(\mathcal{F}) \\ G &\mapsto G \otimes \mathbb{F}_p, \end{aligned}$$

zusammen mit den Abbildungen

$$\begin{aligned} \operatorname{Mor}(G, H) &\mapsto \operatorname{Mor}(G \otimes \mathbb{F}_p, H \otimes \mathbb{F}_p) \\ (f: G \rightarrow H) &\mapsto (f_* := f \otimes \operatorname{id}: G \otimes \mathbb{F}_p \rightarrow H \otimes \mathbb{F}_p) \end{aligned}$$

ein kovarianter Funktor von der Kategorie der abelschen Gruppe zu der Kategorie der $\mathbb{F}_p$ -Vektorräume.

Definition. Es seien $\mathcal{C}$ und $\mathcal{D}$ Kategorien. Ein *kontravarianter Funktor* $F: \mathcal{C} \rightarrow \mathcal{D}$ besteht aus einer Abbildung

$$F: \operatorname{Ob}(\mathcal{C}) \rightarrow \operatorname{Ob}(\mathcal{D})$$

und für alle $X, Y \in \operatorname{Ob}(\mathcal{C})$ gibt es zudem eine Abbildung

$$\begin{aligned} \operatorname{Mor}(X, Y) &\rightarrow \operatorname{Mor}(F(Y), F(X)) \\ \phi &\mapsto F(\phi) \end{aligned}$$

so dass folgende Axiome erfüllt sind:

- (F1) Für alle $\phi \in \operatorname{Mor}(X, Y)$ und $\psi \in \operatorname{Mor}(Y, Z)$, gilt

$$F(\psi \circ \phi) = F(\phi) \circ F(\psi).$$

- (F2) Für alle $X \in \operatorname{Ob}(\mathcal{C})$ gilt

$$F(\operatorname{id}_X) = \operatorname{id}_{F(X)}.$$

Konvention. Für einen kovarianten Funktion $F: \mathcal{C} \rightarrow \mathcal{D}$ und einen Morphismus ϕ schreiben wir oft ϕ^* anstatt $F(\phi)$.

Bemerkung. Die Verknüpfung von zwei kovarianten Funktoren ist wiederum ein kovarianter Funktor. Eine Verknüpfung von einem kovarianten mit einem kontravarianten Funktor ist ein kontravarianter Funktion, nachdem sich “die Richtung genau einmal ändert”. Die Verknüpfung von zwei kontravarianten Funktoren ist hingegen ein kovarianter Funktor, denn “die Richtung ändert sich zwei Mal”.⁷²

Der Unterschied zur Definition eines kovarianten Funktors liegt darin, dass sich die Reihenfolge von $F(\psi)$ und $F(\phi)$ umgedreht hat. Dies erscheint auf den ersten Blick sehr unnatürlich, dies kommt aber durchaus desöfteren vor, wie wir jetzt in den Beispielen sehen werden.

⁷¹Hierbei bezeichnet $G \otimes \mathbb{F}_p$ das Tensorprodukt der abelschen Gruppen G und $\mathbb{F}_p$.

⁷²Man kann diese Beobachtung leicht verallgemeinern, die Verknüpfung von k kovarianten und n kontravarianten Funktoren ist ein kovarianter Funktor, wenn n gerade ist und es handelt sich um einen kontravarianten Funktor, wenn n ungerade ist.

Beispiele.

- (1) Es sei $\mathcal{V}$ die Kategorie der reellen Vektorräume. Wir hatten in Lemma 7.6 gezeigt, dass

$$\begin{aligned} F: \text{Ob}(\mathcal{V}) &\rightarrow \text{Ob}(\mathcal{V}) \\ V &\mapsto V^* := \text{Hom}_{\mathbb{R}}(V, \mathbb{R}) := \text{Dualraum von } V, \end{aligned}$$

zusammen mit den Abbildungen

$$\begin{aligned} \text{Mor}(U, V) &\mapsto \text{Mor}(V^*, U^*) \\ (\phi: U \rightarrow V) &\mapsto \left(\begin{array}{ccc} \phi^*: \text{Hom}(V, \mathbb{R}) &\rightarrow & \text{Hom}(U, \mathbb{R}) \\ f &\mapsto & f \circ \phi \end{array} \right) \end{aligned}$$

ein kontravarianter Funktor ist. Folgendes Diagramm illustriert die Aussage von Axiom (F2):

$$\begin{array}{ccccc} & & \psi \circ \phi & & \\ & \searrow & & \nearrow & \\ U & \xrightarrow{\phi} & V & \xrightarrow{\psi} & W \\ \downarrow \wr & & \downarrow \wr & & \downarrow \wr \\ U^* & \xleftarrow{\phi^*} & V^* & \xleftarrow{\psi^*} & W^* \\ & \nwarrow & & \swarrow & \\ & & (\psi \circ \phi)^* & & \end{array}$$

- (2) Es sei $\mathcal{T}$ die Kategorie der topologischen Räume und es sei $\mathcal{V}$ die Kategorie der reellen Vektorräume, dann ist

$$\begin{aligned} F: \text{Ob}(\mathcal{T}) &\rightarrow \text{Ob}(\mathcal{V}) \\ X &\mapsto C(X, \mathbb{R}) := \text{stetige reellwertige Funktionen auf } X, \end{aligned}$$

zusammen mit den Abbildungen

$$\begin{aligned} \text{Mor}(X, Y) &\mapsto \text{Mor}(C(Y, \mathbb{R}), C(X, \mathbb{R})) \\ (\phi: X \rightarrow Y) &\mapsto \left(\begin{array}{ccc} C(Y, \mathbb{R}) &\rightarrow & C(X, \mathbb{R}) \\ f &\mapsto & f \circ \phi \end{array} \right) \end{aligned}$$

ein kontravarianter Funktor.

8. DIFFERENTIALFORMEN ERSTER ORDNUNG

8.1. Definition von Differentialformen erster Ordnung. Es sei M eine k -dimensionale Untermannigfaltigkeit und es sei $P \in M$. Wir schreiben

$$T_P^*M := \text{Hom}_{\mathbb{R}}(T_P M, \mathbb{R})$$

für den zum Tangentialraum $T_P M$ dualen Vektorraum. Es folgt aus Satz 6.1 und aus der Diskussion vom vorherigen Kapitel, dass dies ebenfalls ein k -dimensionaler Vektorraum ist. Wir nennen Elemente in T_P^*M *Kotangentialvektoren*.

Bemerkung. Es sei wiederum $\mathcal{P}$ die Kategorie der punktierten Untermannigfaltigkeiten und und es sei $\mathcal{V}$ die Kategorie der reellen Vektorräume. Es folgt leicht aus Definitionen, dass

$$\begin{aligned} F: \text{Ob}(\mathcal{P}) &\rightarrow \text{Ob}(\mathcal{V}) \\ (M, P) &\mapsto T_P^*M, \end{aligned}$$

zusammen mit den Abbildungen

$$\begin{aligned} \text{Mor}((M, P), (N, Q)) &\mapsto \text{Mor}(T_Q^*N, T_P^*M) \\ f &\mapsto f^* = Df(P)^* \end{aligned}$$

ein kontravarianter Funktor ist.⁷³

Definition. Eine *Differentialform erster Ordnung* auf einer Untermannigfaltigkeit M ist eine Abbildung ω , welche jedem Punkt P in M einen Kotangentialvektor $\omega_P \in T_P^*M$ zuordnet.⁷⁴

Im Folgenden werden wir oft auch nur kurz “1-Form” anstatt Differentialform erster Ordnung schreiben.⁷⁵ In späteren Kapiteln werden wir dann die Differentialformen höherer Ordnung kennenlernen.

Beispiele.

- (1) Es sei F ein Vektorfeld auf einer Untermannigfaltigkeit M . Dann ist die Abbildung $\omega(F)$, welche jedem Punkt P in M den Kotangentialvektor

$$\begin{aligned} \omega(F)_P: T_P M &\rightarrow \mathbb{R} \\ v &\mapsto \underbrace{v \cdot F(P)}_{\text{Skalarprodukt}} \end{aligned}$$

zuordnet, eine 1-Form auf M .

⁷³Genauer gesagt handelt es sich bei F um die Verknüpfung des kovarianten Funktors $(M, P) \mapsto T_P M$ mit dem kontravarianten Funktor $V \mapsto V^*$, also ist F nach der Bemerkung auf Seite 88 ein kontravarianter Funktor.

⁷⁴Genauer gesagt ist eine Differentialform ω erster Ordnung eine Abbildung

$$\omega: M \rightarrow \bigcup_{P \in M} T_P^*M$$

so dass für $P \in M$ gilt, dass $\omega_P \in T_P^*M$.

⁷⁵In der Literatur werden Differentialformen erster Ordnung manchmal auch Pfaffsche Formen genannt.

- (2) Es sei $f: M \rightarrow \mathbb{R}$ eine stetig differenzierbare Funktion auf einer Untermannigfaltigkeit. Dann ist die Abbildung, welche jedem Punkt P in M den Kotangentialvektor⁷⁶

$$\begin{aligned} Df(P) = f_*: T_P M &\rightarrow \mathbb{R} = T_{f(P)} \mathbb{R} \\ v &\mapsto Df(P)(v) \end{aligned}$$

zuordnet eine 1-Form auf M . Diese Form heißt das *totale Differential von f* und wird mit df bezeichnet.⁷⁷

- (3) Es sei $M \subset \mathbb{R}^n$ eine n -dimensionale Untermannigfaltigkeit. Für $i \in \{1, \dots, n\}$ betrachten wir die 1-Form dx_i , welche an jedem Punkt $P \in M$ gegeben ist durch⁷⁸

$$\begin{aligned} dx_i: T_P M = \mathbb{R}^n &\rightarrow \mathbb{R} \\ (v_1, \dots, v_n) &\rightarrow v_i. \end{aligned}$$

Es seien nun $f_1, \dots, f_n: M \rightarrow \mathbb{R}$ Funktionen. Dann ist auch $f_1 dx_1 + \dots + f_n dx_n$ eine 1-Form.

Da $dx_1, \dots, dx_n: \mathbb{R}^n \rightarrow \mathbb{R}$ eine Basis für den Dualraum $T_P^* M = \text{Hom}_{\mathbb{R}}(\mathbb{R}^n, \mathbb{R})$ bilden, ist jeder Kotangentialvektor eine Linearkombination von $dx_1, \dots, dx_n$, und also ist auch jede 1-Form auf M von der obigen Form.

In Übungsblatt 5 werden wir folgendes elementare Lemma beweisen.

Lemma 8.1. *Es sei $f: U \rightarrow \mathbb{R}$ eine stetig differenzierbare Funktion auf einer offenen Menge $U \subset \mathbb{R}^n$. Dann gilt*

$$df = \frac{\partial f}{\partial x_1} dx_1 + \dots + \frac{\partial f}{\partial x_n} dx_n.$$

Definition. Es sei $\varphi: M \rightarrow N$ eine stetig differenzierbare Abbildung zwischen zwei Untermannigfaltigkeiten und es sei ω eine 1-Form auf N . Wir bezeichnen mit $\varphi^* \omega$ die 1-Form, welche am Punkt $P \in M$ gegeben ist durch

$$(\varphi^* \omega)_P: T_P M \xrightarrow{\varphi_*} T_{\varphi(P)} N \xrightarrow{\omega_{\varphi(P)}} \mathbb{R}.$$

Man sagt oft, dass man eine 1-Form ω von N auf M zurückgezogen hat. Man nennt $\varphi^* \omega$ den “pullback” von ω .

Das folgende Lemma ist eine Tautologie.⁷⁹ Mit anderen Worten, die Aussage des Lemmas folgt sofort aus den Definitionen, und es gibt nichts zu beweisen. Obwohl (oder gerade weil?) das Lemma so einfach ist, werden wir es erstaunlich oft verwenden.

⁷⁶Wir hatten in Satz 7.4 gezeigt, dass $Df(P) = f_*: T_P M \rightarrow \mathbb{R}$ in der Tat eine lineare Abbildung ist, d.h. $Df(P)$ ist ein Element in $T_P^* M$.

⁷⁷Eigentlich hatten wir diese Abbildung schon mit Df bezeichnet, aber in diesem Zusammenhang hat sich die Notation df eingebürgert.

⁷⁸Die Projektion $\mathbb{R}^n \rightarrow \mathbb{R}$ auf die i -te Koordinate wird oft mit p_i bezeichnet. Dann entspricht dx_i gerade der 1-Form dp_i , welche wir in (2) eingeführt hatten. Aber nachdem sich die Notation dx_i eingebürgert hat, werden wir diese verwenden, selbst wenn sie nicht ganz konsistent mit der Notation aus (2) ist.

⁷⁹

Lemma 8.2. *Es sei $\varphi: M \rightarrow N$ eine stetig differenzierbare Abbildung zwischen Untermannigfaltigkeiten, es sei ω eine 1-Form auf N , es sei $P \in M$ und es sei $v \in T_P M$. Dann gilt*

$$(\varphi^* \omega)_P(v) = \omega_{\varphi(P)}(\varphi_*(v)).$$

Definition.

- (1) Es sei $\omega = f_1 dx_1 + \dots + f_k dx_k$ eine 1-Form⁸⁰ auf einer offenen Teilmenge $W \subset \mathbb{R}^k$ beziehungsweise einer offenen Teilmenge $W \subset H_k$. Wir sagen ω ist stetig, wenn $f_1, \dots, f_k$ stetig sind.
- (2) Es sei ω eine 1-Form auf einer k -dimensionalen Untermannigfaltigkeit M .
 - (a) Es sei $\Phi: U \rightarrow V$ eine Karte für M . Wir sagen ω ist stetig bezüglich der Karte Φ , wenn $(\Phi^{-1})^*(\omega)$ eine stetige 1-Form auf $V \cap E_k$ beziehungsweise auf $V \cap H_k$ ist.⁸¹
 - (b) Wir sagen ω ist stetig auf M , wenn es einen Atlas $\{\Phi_i: U_i \rightarrow V_i\}_{i \in I}$ von M gibt, so dass ω stetig ist bezüglich jeder Karte Φ_i .

Wir verwenden im Folgenden auch die ganz analogen Definitionen bei denen wir “stetig” durch “stetig differenzierbar” beziehungsweise “glatt” ersetzen.

Der folgende Satz spielt nun eine ähnliche Rolle wie Satz 2.7.

Satz 8.3. *Es sei ω eine 1-Form auf einer Untermannigfaltigkeit M . Dann sind die folgenden Aussagen äquivalent:*

- (1) *die 1-Form ω ist stetig differenzierbar,*
- (2) *die 1-Form ω ist stetig differenzierbar bezüglich jeder Karte von M ,*
- (3) *für alle stetig differenzierbaren Tangentialvektorfelder F auf M ist die Funktion*

$$\begin{array}{ccc} M & \rightarrow & \mathbb{R} \\ P & \mapsto & \underbrace{\omega_P}_{\in T_P^* M} \underbrace{(F(P))}_{\in T_P M} \\ & & \underbrace{\hspace{1.5cm}}_{\in \mathbb{R}} \end{array}$$

stetig differenzierbar.

Die gleichen Aussagen gelten auch, wenn wir “stetig differenzierbar” durch “stetig” beziehungsweise “glatt” ersetzen.

Für den Beweis von Satz 8.3 benötigen wir folgendes Lemma.

Lemma 8.4. *Es sei $\varphi: V \rightarrow W$ eine stetig differenzierbare Abbildung von einer m -dimensionalen Untermannigfaltigkeit V von $\mathbb{R}^m$ zu einer n -dimensionalen Untermannigfaltigkeit W von $\mathbb{R}^n$. Zudem sei ω eine stetig differenzierbare 1-Form auf W . Dann ist die 1-Form $\varphi^* \omega$ auf V ebenfalls stetig differenzierbar. Die gleiche Aussage gilt auch, wenn wir “stetig differenzierbar” durch “stetig” beziehungsweise “glatt” ersetzen.*

⁸⁰Wie wir gerade in (3) angemerkt hatten ist jede 1-Form auf W von dieser Form.

⁸¹Wie immer fassen wir hierbei $V \cap E_k$ als offene Teilmenge von $E_k = \mathbb{R}^k$ auf.

Beweis. Wir beweisen die Aussage für “stetig differenzierbar”, die Aussagen für “glatt” und “stetig” werden ganz analog bewiesen. Es sei also $\omega = g_1 dx_1 + \cdots + g_n dx_n$ eine stetig differenzierbare 1-Form auf einer n -dimensionalen Untermannigfaltigkeit W von $\mathbb{R}^n$. Zudem sei V eine m -dimensionalen Untermannigfaltigkeit W von $\mathbb{R}^m$. Darüber hinaus sei $\varphi = (\varphi_1, \dots, \varphi_n): V \rightarrow W$ eine stetig differenzierbare Abbildung. Wir schreiben

$$\varphi^* \omega = f_1 dx_1 + \cdots + f_m dx_m,$$

wobei $f_1, \dots, f_m: V \rightarrow \mathbb{R}$ Funktionen sind. Wir müssen zeigen, dass $f_1, \dots, f_m$ stetig differenzierbar sind. Es sei nun $k \in \{1, \dots, m\}$. Für $P \in V$ ist

$$\begin{aligned} f_k(P) &= (f_1(P)dx_1 + \cdots + f_m(P)dx_m)(e_k) &&= (\varphi^* \omega)_P(e_k) = \omega_{\varphi(P)}(\varphi_*(e_k)) \\ &\uparrow &&\uparrow \\ &\text{denn } dx_i(e_j) = \delta_{ij}^{82} &&\text{Lemma 8.2} \\ &= \underbrace{\left(\sum_{i=1}^n (g_i \circ \varphi)(P) dx_i \right)}_{=\omega_{\varphi(P)}} \underbrace{\left(\sum_{j=1}^m \frac{\partial \varphi_j}{\partial x_k}(P) \cdot e_j \right)}_{=(D_P \varphi)(e_k) = \varphi_*(e_k)} &&= \sum_{i=1}^n (g_i \circ \varphi)(P) \cdot \frac{\partial \varphi_j}{\partial x_k}(P). \\ &\uparrow &&\uparrow \\ &\text{denn } dx_i(e_j) = \delta_{ij} &&\end{aligned}$$

Wir sehen also, dass f_k eine Verknüpfung von stetig differenzierbaren Funktionen ist. Also ist f_k auch selber stetig differenzierbar. $\square$

Das folgende Lemma beweist die Aussagen von Satz 8.3 im Spezialfall, dass M eine k -dimensionale Untermannigfaltigkeit von $\mathbb{R}^k$ ist.

Lemma 8.5. *Es sei ω eine 1-Form auf einer k -dimensionalen Untermannigfaltigkeit von $\mathbb{R}^k$. Dann sind die folgenden Aussagen äquivalent:*

- (1) *die 1-Form ω ist stetig differenzierbar,*
- (2) *für jeden Diffeomorphismus $\varphi: V \rightarrow W$ ist die 1-Form $\varphi^* \omega$ auf V stetig differenzierbar,*
- (3) *für alle stetig differenzierbaren Vektorfelder $F: W \rightarrow \mathbb{R}^k$ ist die Funktion*

$$\begin{aligned} W &\rightarrow \mathbb{R} \\ P &\mapsto \omega_P(F(P)) \end{aligned}$$

stetig differenzierbar.

Die gleichen Aussagen gelten auch, wenn wir “stetig differenzierbar” durch “stetig” beziehungsweise “glatt” ersetzen.

Beweis. Es sei $\omega = g_1 dx_1 + \cdots + g_k dx_k$ eine 1-Form auf einer k -dimensionalen Untermannigfaltigkeit von $\mathbb{R}^k$. Die Aussage “(1) $\Rightarrow$ (2)” folgt aus Lemma 8.5. Die Umkehrung “(2) $\Rightarrow$ (1)” ist natürlich trivial, wir müssen nur $\varphi = \text{id}$ setzen.

⁸²Für $i, j \in \mathbb{Z}$ ist das *Kronecker-Symbol* δ_{ij} definiert durch

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

Wir beweisen nun noch die Äquivalenz “(1) $\Leftrightarrow$ (3)”. Wir beweisen die Aussage für “stetig differenzierbar”, die Aussagen für “glatt” und “stetig” werden ganz analog bewiesen. Wir machen dazu folgende Vorbemerkung: Für ein beliebiges stetig differenzierbares Vektorfeld $F = (F_1, \dots, F_k): W \rightarrow \mathbb{R}^k$ gilt

$$\omega_P(F(P)) = \underbrace{\left(\sum_{i=1}^k g_i(P) dx_i\right)}_{=\omega_P} \underbrace{\left(\sum_{j=1}^k F_j(P) \cdot e_j\right)}_{=F(P)} \underset{\substack{\uparrow \\ \text{denn } dx_i(e_j) = \delta_{ij}}}{=} \sum_{i=1}^k g_i(P) \cdot F_i(P).$$

Aus dieser Rechnung folgt, dass wenn die Funktionen $g_1, \dots, g_k$ stetig differenzierbar sind, dann ist auch $P \mapsto \omega_P(F(P))$ stetig differenzierbar. Wir haben damit “(1) $\Rightarrow$ (3)” bewiesen. Umgedreht, wenn wir die Vektorfelder konstanten Vektorfelder $G(P) = e_i$ betrachten, dann erhalten wir $g_i(P) = \omega_P(G(P))$. Nach Voraussetzung ist die Funktion $P \mapsto \omega_P(G(P))$ stetig differenzierbar. Es folgt also, dass auch g_i stetig differenzierbar ist. Dies beweist “(3) $\Rightarrow$ (1)”. $\square$

Wir wenden uns nun dem eigentlichen Beweis von Satz 8.3 zu.

Beweis von Satz 8.3. Es sei also ω eine 1-Form auf einer k -dimensionalen Untermannigfaltigkeit M . Um die Notation etwas zu vereinfachen betrachten wir nur den Fall, dass $\partial M = \emptyset$. Wir beweisen die Aussage für “stetig differenzierbar”, die Aussagen für “glatt” und “stetig” werden ganz analog bewiesen.

Die Aussage “(2) $\Rightarrow$ (1)” ist natürlich trivial. Die Umkehraussage “(1) $\Rightarrow$ (2)” kann man, ganz ähnlich wie im Beweis von Satz 2.7, mithilfe von Karten auf den in Lemma 8.5 betrachteten Spezialfall zurückführen. Wir überlassen das Ausführen von diesem Argument als freiwillige Übungsaufgabe.

Wir beweisen nun die Aussage “(2) $\Rightarrow$ (3)”. Es sei also F ein stetig differenzierbares Tangentialvektorfeld auf M . Für $P \in M$ definieren wir $g(P) := \omega_P(F(P))$. Wir müssen nun zeigen, dass die Funktion g auf der Untermannigfaltigkeit M stetig differenzierbar ist. Wir verwenden dazu die Definition von “stetig differenzierbar” aus Kapitel 2.5. Es sei also $\Phi: U \rightarrow V$ eine Karte um P . Wir bezeichnen mit $\Psi: V \rightarrow U$ die Umkehrabbildung. Wir müssen nun zeigen, dass die Funktion $g \circ \Psi: V \cap E_k \rightarrow \mathbb{R}$ stetig differenzierbar ist. Für $Q \in V \cap E_k$ gilt

$$\begin{array}{ccc} (g \circ \Psi)(Q) & = & \omega_{\Psi(Q)}(F(\Psi(Q))) & = & ((\Psi \circ \Phi)^* \omega)_{\Psi(Q)}(F(\Psi(Q))) \\ \uparrow & & \uparrow & & \uparrow \\ \text{Definition von } g & & \text{denn } \Psi \circ \Phi = \text{id} & & \\ & = & (\Phi^*(\Psi^* \omega))_{\Psi(Q)}(F(\Psi(Q))) & = & (\Psi^* \omega)_Q(\Phi_*(F(\Psi(Q)))) \\ \uparrow & & \uparrow & & \uparrow \\ \text{denn } (\Psi \circ \Phi)^* & = & \Phi^* \circ \Psi^* & & \text{Lemma 8.2} \end{array}$$

Nach Voraussetzung ist $\Psi^* \omega$ eine stetig differenzierbare 1-Form auf $V \cap E_k$ und zudem ist $Q \mapsto \Phi_*(F(\Psi(Q)))$ ein stetig differenzierbares Vektorfeld auf $V \cap E_k$. Es folgt nun aus Lemma 8.5, dass die Funktion $g \circ \Psi$ stetig differenzierbar ist.

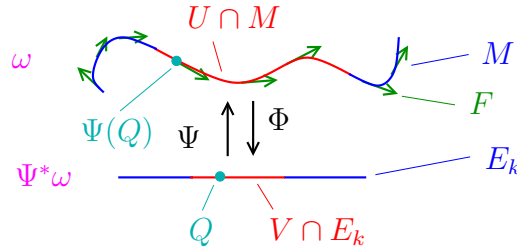


ABBILDUNG 48. Skizze zum Beweis der Aussage “(2) ⇒ (3)”.

Wir beweisen nun noch die Aussage “(3) ⇒ (1)”.⁸³ Wir nehmen also an, dass für alle stetig differenzierbaren Tangentialvektorfelder F auf M die Funktion $P \mapsto \omega_P(F(P))$ stetig differenzierbar ist. Wir müssen jetzt beweisen, dass es um jeden Punkt $P \in M$ eine Karte $\Phi: S \rightarrow T$ gibt, so dass $(\Phi^{-1})^*\omega$ eine stetig differenzierbare 1-Form auf $T \cap E_k$ ist. Es sei also $P \in M$ und es sei zuerst einmal $\Phi: U \rightarrow V$ eine beliebige Karte um P . Wir bezeichnen mit $\Psi: V \rightarrow U$ die Umkehrabbildung. Für $i \in \{1, \dots, k\}$ sei nun $G_i: E_k \cap V \rightarrow E_k = \mathbb{R}^k$ das konstante Vektorfeld, welches jedem Punkt den Standardbasisvektor e_i zuordnet. Dann ist $P \mapsto \Psi_*(G_i(P))$ ein Tangentialvektorfeld auf $M \cap U$. Dieses lässt sich jedoch im Allgemeinen nicht zu einem stetig differenzierbaren Vektorfeld auf M fortsetzen. Wir benötigen nun folgende Behauptung, welche ähnlich bewiesen wird wie Lemma 6.11. Der genaue Beweis der Behauptung ist eine freiwillige Übungsaufgabe.

Behauptung. Es existiert eine offene Umgebung S von P in M und es existiert eine glatte Funktion $z: M \rightarrow [0, 1]$ mit folgenden Eigenschaften:

- (1) $z(x) = 1$ für alle $x \in S$,
- (2) $\text{Träger}(z) \subset M \cap U$.

Wir betrachten nun auf M das Tangentialvektorfeld, welches definiert ist durch

$$F(x) = \begin{cases} z(x) \cdot \Psi_*(G_i(\Phi(x))), & \text{wenn } x \in M \cap U, \\ 0, & \text{sonst.} \end{cases}$$

Dies ist ein stetig differenzierbares Vektorfeld auf M . Wir setzen $T := \Phi(S)$. Für $Q \in T \cap E_k$ gilt nun

$$\begin{array}{ccccc} (\Psi^*\omega)_Q(G_i(Q)) & = & \omega_{\Psi(Q)}(\Psi_*(G_i(Q))) & = & \omega_{\Psi(Q)}(F(\Psi(Q))) \\ \uparrow & & \uparrow & & \\ \text{Lemma 8.2} & & \text{denn } z(Q) = 1, \text{ da } \Psi(Q) \in S & & \end{array}$$

Nach Voraussetzung ist $P \mapsto \omega_P(F(P))$ stetig differenzierbar, also ist nun auch die Funktion $Q \mapsto (\Psi^*\omega)_Q(G_i(Q))$ auf $T \cap E_k$ stetig differenzierbar. Es folgt aus dem Beweis der Aussage “(3) ⇒ (1)” in Lemma 8.5, dass $\Psi^*\omega$ auf T stetig differenzierbar ist. $\square$

⁸³Wir werden die Aussage “(3) ⇒ (1)” im weiteren Verlauf der Vorlesung nicht verwenden.

Satz 8.6. *Es sei $f: M \rightarrow N$ eine glatte Abbildung zwischen zwei Untermannigfaltigkeiten und es sei ω eine stetig differenzierbare 1-Form auf N . Dann ist $f^*\omega$ eine stetig differenzierbare 1-Form auf M . Die gleichen Aussagen gelten auch, wenn wir “stetig differenzierbar” durch “stetig” beziehungsweise “glatt” ersetzen.*

Beweis. Die Aussage des Satzes kann man mithilfe von Karten problemlos auf Lemma 8.4 zurückführen. $\square$

Bemerkung. Für eine Untermannigfaltigkeit M setzen wir nun

$$\Omega_1(M) = \text{Menge der glatten 1-Formen auf } M.$$

Zusammenfassend haben wir also insbesondere gezeigt, dass die Abbildungen

$$\text{Untermannigfaltigkeit } M \mapsto \Omega_1(M),$$

zusammen mit den Abbildungen

$$\text{glatte Abbildung } f: M \rightarrow N \mapsto (f^*: \Omega_1(N) \rightarrow \Omega_1(M))$$

einen kontravarianten Funktor von der Kategorie der Untermannigfaltigkeiten zu der Kategorie der Vektorräume bildet.

Wir beschließen dieses Kapitel mit folgendem Lemma, welches in Übungsblatt 5 bewiesen wird.

Lemma 8.7. *Es sei M eine Untermannigfaltigkeit.*

- (1) *Für jedes glatte Tangentialvektorfeld F auf M ist die zugehörige 1-Form $\omega(F)$, welche wir auf Seite 90 eingeführt hatten, glatt.*
- (2) *Zu jeder glatten 1-Form ω auf M gibt es genau ein glattes Tangentialvektorfeld F auf M mit $\omega = \omega(F)$.*

Die analoge Aussage gilt natürlich auch wenn wir “glatt” durch “stetig” beziehungsweise “stetig differenzierbar” ersetzen.

Bemerkung. Für eine Untermannigfaltigkeit M setzen wir nun

$$\tau(M) = \text{Menge der glatten Tangentialvektorfelder auf } M.$$

Dies ist ein Vektorraum und es folgt aus Lemma 8.7, dass die Abbildung

$$\begin{aligned} \tau(M) &\rightarrow \Omega_1(M) \\ F &\mapsto \omega(F) \end{aligned}$$

ein Isomorphismus von Vektorräumen ist. Wir hatten gerade angemerkt, dass $M \mapsto \Omega_1(M)$ und $f \mapsto f^*$ einen kontravarianten Funktor von der Kategorie der Untermannigfaltigkeiten zu der Kategorie der Vektorräume bildet. Es stellt sich nun die Frage, ob die Abbildungen $M \mapsto \tau(M)$ ebenfalls auf eine “natürliche” Weise einen Funktor ergeben. Aber dies ist nicht der Fall, denn eine glatte Abbildung $f: M \rightarrow N$ induziert weder eine Abbildung $\tau(M) \rightarrow \tau(N)$ noch eine Abbildung $\tau(N) \rightarrow \tau(M)$.⁸⁴

⁸⁴Warum nicht? Probieren Sie mal, solche Abbildungen hinzuschreiben. Woran scheitert das?

8.2. Integrierbarkeit von 1-Formen.

Lemma 8.8. *Es sei M eine Untermannigfaltigkeit, es sei $\gamma: [a, b] \rightarrow M$ eine Kurve auf M und es sei ω eine stetige 1-Form auf M . Dann ist die Funktion⁸⁵*

$$\begin{array}{ccc} [a, b] & \rightarrow & \mathbb{R} \\ t & \mapsto & \underbrace{\underbrace{\omega_{\gamma(t)}}_{\in T_{\gamma(t)}^* M} \left(\underbrace{\gamma'(t)}_{\in T_{\gamma(t)} M} \right)}_{\in \mathbb{R}}. \end{array}$$

stetig.

Beweis. Wir bezeichnen mit $e = 1$ den Einheitsvektor in $\mathbb{R}$. Für $t \in [a, b]$ ist nun

$$\begin{array}{ccccc} \omega_{\gamma(t)}(\gamma'(t)) & = & \omega_{\gamma(t)}(\gamma_*(e)) & = & (\gamma^*\omega)_t(e). \\ & & \uparrow & & \\ & & \text{Lemma 8.2} & & \end{array}$$

Die 1-Form $\gamma^*\omega$ auf $[a, b]$ ist nach Satz 8.6 stetig. Also folgt aus Satz 8.3 (1) $\Rightarrow$ (3), dass die Funktion $t \mapsto (\gamma^*\omega)_t(e)$ stetig ist. $\square$

Es sei nun $\gamma: [a, b] \rightarrow M$ eine Kurve auf einer Untermannigfaltigkeit M und es sei ω eine stetige 1-Form auf M . Wir definieren das *Integral der 1-Form ω über die Kurve γ* wie folgt⁸⁶

$$\int_{\gamma} \omega := \int_a^b \omega_{\gamma(t)}(\gamma'(t)) dt.$$

Beispiel. Es sei F ein stetiges Vektorfeld auf einer Untermannigfaltigkeit M und es sei $\gamma: [a, b] \rightarrow M$ eine Kurve. Wir betrachten wieder die 1-Form $\omega(F)$, welche an jedem Punkt $P \in M$ gegeben ist durch

$$\begin{array}{ccc} \omega(F)_P: T_P M & \rightarrow & \mathbb{R} \\ v & \mapsto & \omega(F)_P(v) := v \cdot F(P). \end{array}$$

Es folgt sofort aus den Definitionen, das

$$\int_{\gamma} \omega(F) = \int_a^b \underbrace{\gamma'(t) \cdot F(\gamma(t))}_{\substack{\text{Skalarprodukt der} \\ \text{Vektoren } \gamma'(t) \text{ und } F(\gamma(t))}} dt.$$

In Abbildung 49 betrachten wir anschaulich einige Spezialfälle mit $M = \mathbb{R}^2$. Wir betrachten

⁸⁵Die Funktion ist also wie folgt definiert: die Kurve γ definiert für $t \in [a, b]$ einen Vektor $\gamma'(t) \in T_{\gamma(t)} M$, eine 1-Form auf M ordnet solch einem Vektor eine reelle Zahl zu.

⁸⁶Das Integral ist definiert, denn nach Lemma 8.8 ist der Integrand eine stetige Funktion.

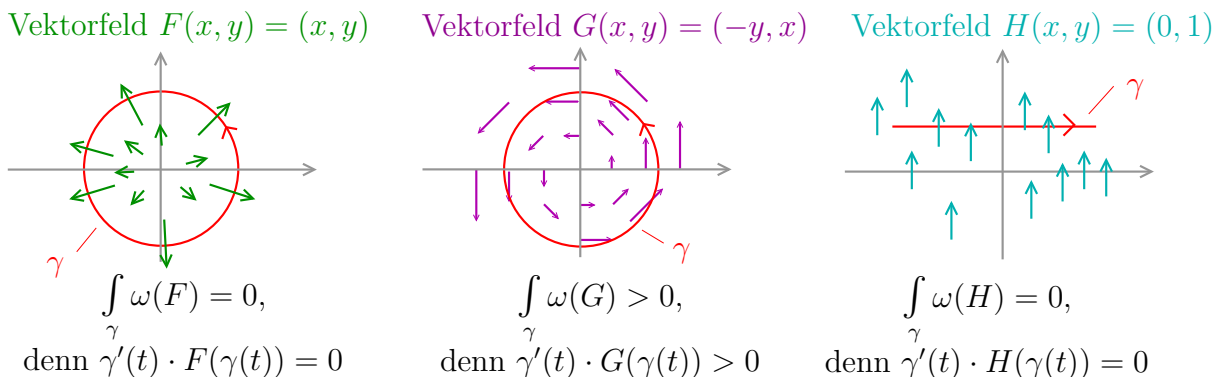


ABBILDUNG 49.

insbesondere dass Vektorfeld in der Mitte, welches gegeben ist durch $G(x, y) = (-y, x)$. Zudem betrachten wir die geschlossene Kurve

$$\begin{aligned} \gamma: [0, 2\pi] &\rightarrow \mathbb{R}^2 \\ t &\mapsto (\cos(t), \sin(t)). \end{aligned}$$

Dann gilt

$$\int_{\gamma} \omega(G) = \int_0^{2\pi} \gamma'(t) \cdot G(\gamma(t)) dt = \int_0^{2\pi} (-\sin(t), \cos(t)) \cdot (-\sin(t), \cos(t)) dt = \int_0^{2\pi} 1 dt = 2\pi.$$

Beispiel. Es sei F ein stetiges Kraftfeld auf einer offenen Teilmenge $U \subset \mathbb{R}^3$, z.B. das Gravitationsfeld oder ein Magnetfeld, welches auf einen gegebenen punktförmigen Körper K wirkt. Wir bezeichnen mit $\omega(F)$ die zugehörige 1-Form auf U . Es sei $\gamma: [a, b] \rightarrow U$ zudem eine Kurve. Dann ist

$$\int_{\gamma} -\omega(F)$$

die Energie die aufgewendet werden muss (bzw. freigesetzt wird), um den Körper K in diesem Kraftfeld auf der Bahn γ zu bewegen. Es folgt aus dem Energieerhaltungssatz, dass das Integral über einer geschlossene Kurve⁸⁷ immer verschwindet.

Der folgende Satz kann als Verallgemeinerung des Hauptsatzes der Differential- und Integralrechnung aufgefasst werden.⁸⁸

Satz 8.9. *Es sei M eine Untermannigfaltigkeit, es sei $f: M \rightarrow \mathbb{R}$ eine stetig differenzierbare Funktion und zudem sei $\gamma: [a, b] \rightarrow M$ eine Kurve auf M . Dann gilt⁸⁹*

$$\int_{\gamma} df = f(\gamma(b)) - f(\gamma(a)).$$

⁸⁷Eine Kurve heißt *geschlossen*, wenn der Anfangs- und der Endpunkt übereinstimmen.

⁸⁸In welchem Sinne ist dies eine Verallgemeinerung des Hauptsatzes der Differential- und Integralrechnung? Wie kann man diesen aus Satz 8.9 herleiten?

⁸⁹Zur Erinnerung, df bezeichnet das totale Differential, welches wir auf Seite 91 eingeführt hatten.

Insbesondere wenn γ eine geschlossene Kurve ist, dann gilt

$$\int_{\gamma} df = 0.$$

Beweis. Es sei $\gamma: [a, b] \rightarrow M$ eine Kurve auf M , dann gilt

$$\begin{aligned} \int_{\gamma} df &= \int_a^b \underbrace{(df_{\gamma(t)})(\gamma'(t))}_{=(f \circ \gamma)'(t), \text{ nach der Kettenregel}} dt = \int_a^b \frac{d}{dt} f(\gamma(t)) dt = f(\gamma(b)) - f(\gamma(a)). \end{aligned}$$

Wenn γ geschlossen ist, dann gilt $\gamma(a) = \gamma(b)$ und es folgt sofort, dass

$$\int_{\gamma} df = 0. \quad \square$$

Definition. Es sei M eine Untermannigfaltigkeit und es sei ω eine glatte 1-Form auf M . Wir sagen, ω ist *exakt*⁹⁰, wenn es eine glatte Funktion $f: M \rightarrow \mathbb{R}$ gibt mit $\omega = df$.

Beispiel. Auf Seite 98 hatten wir gezeigt, dass es auf $M = \mathbb{R}^2$ mit dem Tangentialvektorfeld $F(x, y) = (-y, x)$ eine geschlossene Kurve γ gibt, so dass

$$\int_{\gamma} \omega(F) \neq 0.$$

Es folgt also aus Satz 8.9, dass die 1-Form $\omega(F)$ nicht exakt ist.

Wir werden im Folgenden Beispiele von exakten 1-Formen geben. Hierbei verwenden wir folgendes elementare Lemma.

Lemma 8.10. *Es sei M eine n -dimensionale Untermannigfaltigkeit von $\mathbb{R}^n$ und es sei $F: M \rightarrow \mathbb{R}^n$ ein glattes Vektorfeld. Dann gilt*

$$\text{die 1-Form } \omega(F) \text{ ist exakt} \iff \begin{array}{l} \text{es gibt eine glatte Funktion,} \\ g: M \rightarrow \mathbb{R} \text{ mit } \text{Grad } g = F. \end{array}$$

Beweis. Wir beweisen zuerst die “ $\Rightarrow$ ” Richtung. Wir nehmen also an, dass es eine glatte Funktion $g: M \rightarrow \mathbb{R}$ gibt, so dass $dg = \omega(F)$. Wir wollen zeigen, dass $\text{Grad } g = F$. Für alle $v = (v_1, \dots, v_n) \in \mathbb{R}^n = T_P M$ gilt, dass

$$\begin{array}{ccccccc} F(P) \cdot v & = & (dg(P))(v) & = & \left(\sum_{i=1}^n \frac{\partial g}{\partial x_i}(P) dx_i \right) (v_1, \dots, v_n) & = & \sum_{i=1}^n \frac{\partial g}{\partial x_i}(P) \cdot v_i = \text{Grad } g(P) \cdot v. \\ \uparrow & & \uparrow & & & & \\ \text{da } \omega(F) = dg & & \text{nach Lemma 8.1} & & & & \end{array}$$

Nachdem dies für alle $v \in \mathbb{R}^n$ gilt, folgt, dass $F(P) = \text{Grad } g(P)$.

⁹⁰Exakte 1-Formen werden manchmal auch *integrierbar* genannt.

Wir beweisen nun die “ $\Leftarrow$ ” Richtung. Wir nehmen also an, dass es eine glatte Funktion $g: M \rightarrow \mathbb{R}$ gibt, so dass $F = \text{Grad } g$. Wir wollen zeigen, dass $\omega(F) = dg$. In der Tat gilt für alle $v \in \mathbb{R}^n = T_P M$, dass

$$\omega(F)_P(v) = F(P) \cdot v = \underset{\substack{\uparrow \\ \text{denn } F = \text{Grad } g}}{\text{Grad } g(P)} \cdot v = \underset{\substack{\uparrow \\ \text{siehe oben.}}}{(dg(P))}(v)$$

□

Beispiele.

- (1) Ein Vektorfeld $G: \mathbb{R}^n \setminus \{0\} \rightarrow \mathbb{R}^n$ heißt *radial*, wenn es eine Funktion $f: \mathbb{R}_{>0} \rightarrow \mathbb{R}$ gibt, so dass

$$G(x) = f(\|x\|) \cdot x \text{ für alle } x \in \mathbb{R}^n \setminus \{0\}.$$

In Übungsblatt 5 werden wir mithilfe von Lemma 8.10 zeigen, dass für jedes glatte radiale Vektorfeld $G: \mathbb{R}^n \setminus \{0\} \rightarrow \mathbb{R}^n$ die zugehörige 1-Form $\omega(G)$ exakt ist.

- (2) Das Gravitationsfeld G eines sich im Ursprung befindlichen punktförmigen Körpers in $\mathbb{R}^3$ ist ein radiales Vektorfeld. Wie wir gerade in (1) angemerkt hatten, gibt es eine Funktion $f: \mathbb{R}^3 \setminus \{0\} \rightarrow \mathbb{R}$ mit $\text{Grad } f = G$. Diese Funktion wird in der Physik das “Potential” des Gravitationsfeldes genannt. Anders ausgedrückt, an einem Punkt $P \in \mathbb{R}^3 \setminus \{0\}$ besitzt ein punktförmiger Körper der Masse M gerade die potentielle Energie $M \cdot f(P)$.
- (3) Die Diskussion in (2) gilt ganz analog auch für das elektrische Feld eines sich im Ursprung befindlichen elektrisch geladenen Körpers.

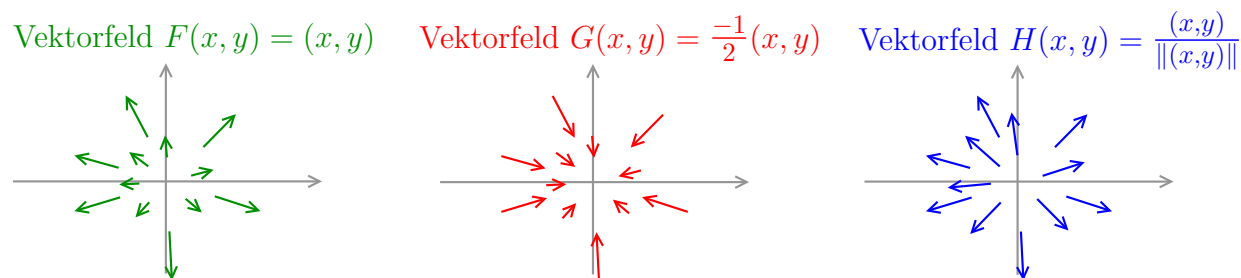


ABBILDUNG 50. Beispiele von radialen Vektorfeldern.

8.3. Integration und Parametertransformationen. Es sei $\gamma: [a, b] \rightarrow \mathbb{R}^n$ eine Kurve und es sei $\varphi: [c, d] \rightarrow [a, b]$ ein Diffeomorphismus. Dann ist

$$\begin{aligned} \gamma \circ \varphi: [c, d] &\rightarrow \mathbb{R}^n \\ t &\mapsto \gamma(\varphi(t)) \end{aligned}$$

ebenfalls eine Kurve. Diese Kurve nimmt die gleichen Punkte wie die ursprüngliche Kurve γ an. Wir sagen δ ist eine *Umparametrisierung*, welche durch die *Parametertransformation* φ aus γ hervorgeht. Wir nennen die Parametertransformation φ *orientierungserhaltend*,

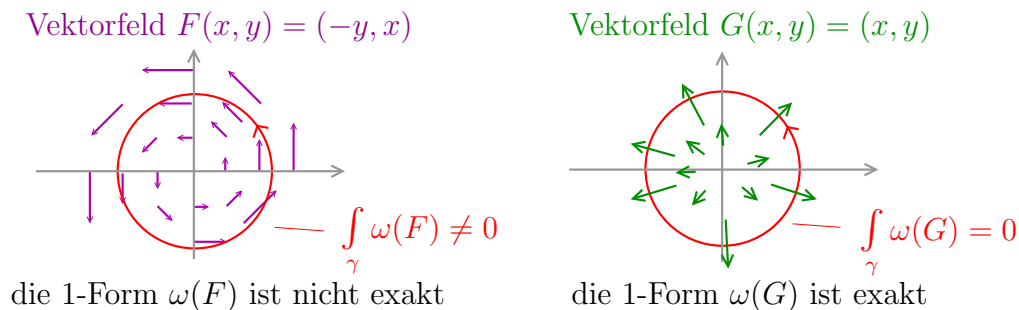


ABBILDUNG 51.

wenn $\varphi(c) = a$ und $\varphi(d) = b$ und wir nennen sie *orientierungsumkehrend*, wenn $\varphi(c) = b$ und $\varphi(d) = a$.

Beispiel. Wir betrachten

$$\begin{aligned} \gamma: [0, 2\pi] &\rightarrow \mathbb{R}^2 \\ t &\mapsto (\cos(t), \sin(t)), \end{aligned}$$

d.h. γ durchläuft den Einheitskreis “in der Zeit 2π ” “gegen den Uhrzeigersinn”. Die Kurve

$$\begin{aligned} \alpha: [0, \pi] &\rightarrow \mathbb{R}^2 \\ t &\mapsto (\cos(2t), \sin(2t)), \end{aligned}$$

ist eine Umparametrisierung von γ , wobei die Parametertransformation orientierungserhaltend ist. Die Kurve α durchläuft den Einheitskreis “in der Zeit π ” weiterhin “gegen den Uhrzeigersinn”. Die Kurve

$$\begin{aligned} \beta: [0, 4\pi] &\rightarrow \mathbb{R}^2 \\ t &\mapsto \left(\cos\left(2\pi - \frac{1}{2}t\right), \sin\left(2\pi - \frac{1}{2}t\right) \right) \end{aligned}$$

ist eine Umparametrisierung von γ , wobei die Parametertransformation orientierungsumkehrend ist. Die Kurve β durchläuft den Einheitskreis “in der Zeit 4π ” nun aber “im Uhrzeigersinn”.

Wir können nun folgenden Satz formulieren und beweisen:

Satz 8.11. *Es sei M eine Untermannigfaltigkeit und es sei ω eine stetige 1-Form auf M . Es sei γ eine Kurve auf M und es sei δ eine Kurve, welche aus γ durch eine Parametertransformation φ hervorgeht. Dann gilt*

$$\int_{\delta} \omega = \begin{cases} \int \omega, & \text{wenn } \varphi \text{ orientierungserhaltend,} \\ -\int_{\gamma} \omega, & \text{wenn } \varphi \text{ orientierungsumkehrend.} \end{cases}$$

Der Satz besagt also, dass das Integral einer 1-Form über eine Kurve nicht davon abhängt, “wie schnell wir die Punkte auf der Kurve durchlaufen”, so lange wir die Kurve “in der

gleichen Richtung", d.h. vom Endpunkt $\gamma(a) = \delta(c)$ zum Endpunkt $\gamma(b) = \delta(d)$ durchlaufen. Wenn wir die Kurve in der "entgegengesetzten Richtung" durchlaufen, dann erhalten wir bis auf ein Vorzeichen, ebenfalls das gleiche Integral.

Der Satz erinnert an Lemma 3.3 aus der Analysis III. In der Tat ist auch der Beweis fast der gleiche wie in Analysis III.

Beweis. Es sei M eine Untermannigfaltigkeit und es sei ω eine stetige 1-Form auf M . Es sei $\gamma: [a, b] \rightarrow M$ eine Kurve auf M und es sei $\delta: [c, d] \rightarrow M$ eine Kurve, welche aus γ durch eine Parametertransformation $\varphi: [c, d] \rightarrow [a, b]$ hervorgeht. Dann gilt

$$\begin{array}{ll}
 \int_{\delta} \omega &= \int_{t=c}^{t=d} \omega_{\delta(t)}(\delta'(t)) dt &= \int_{t=c}^{t=d} \omega_{(\gamma \circ \varphi)(t)}((\gamma \circ \varphi)'(t)) dt \\
 &= \int_{t=c}^{t=d} \omega_{\gamma(\varphi(t))}(\gamma'(\varphi(t)) \cdot \varphi'(t)) dt &= \int_{t=c}^{t=d} \omega_{\gamma(\varphi(t))}(\gamma'(\varphi(t))) \cdot \varphi'(t) dt \\
 \uparrow &\text{Kettenregel} &\uparrow \text{Linearität von } \omega_{\gamma(\varphi(t))}: T_{\gamma(\varphi(t))}M \rightarrow \mathbb{R} \\
 &= \int_{s=\varphi(c)}^{s=\varphi(d)} \omega_{\gamma(s)}(\gamma'(s)) ds &= \begin{cases} \int \omega, & \text{wenn } \varphi(c) = a \text{ und } \varphi(d) = b, \\ -\int_{\gamma} \omega, & \text{wenn } \varphi(c) = b \text{ und } \varphi(d) = a. \end{cases} \\
 \uparrow &\text{Substitution } s = \varphi(t) & \\
 & & \square
 \end{array}$$

8.4. Orientierte eindimensionale Untermannigfaltigkeiten. Auf Seite 61 hatten wir den Begriff einer Orientierung auf einer Hyperfläche eingeführt. Wir wollen nun den Begriff einer Orientierung einer eindimensionalen Untermannigfaltigkeit einführen.

Definition. Es sei M eine eindimensionale Untermannigfaltigkeit.

- (1) Wir sagen zwei Tangentialvektorfelder σ und τ auf M , welche nirgendwo verschwinden, sind *äquivalent*, wenn es eine positive Funktion $f: M \rightarrow \mathbb{R}_{>0}$ gibt, so dass $\sigma(P) = f(P) \cdot \tau(P)$ für alle $P \in M$.
- (2) Eine *Orientierung von M* ist eine Äquivalenzklasse von Tangentialvektorfeldern, welche nirgendwo verschwinden.
- (3) Eine *orientierte eindimensionale Untermannigfaltigkeit* ist ein Paar $(M, \mathcal{O})$, wobei M eine eindimensionale Untermannigfaltigkeit und $\mathcal{O}$ eine Orientierung von M ist. Wir unterschlagen hierbei oft die Orientierung in der Notation. Wir bezeichnen zudem mit $-M$ die orientierte eindimensionale Untermannigfaltigkeit $(M, -\mathcal{O})$.

Bemerkung. Eine eindimensionale Untermannigfaltigkeit M von $\mathbb{R}^2$ ist insbesondere auch eine Hyperfläche. In diesem Fall haben wir also zwei Begriffe von Orientierungen:

- (1) Auf Seite 61 hatten wir eine Orientierung von M definiert als ein Einheitsnormalenfeld.
- (2) Gerade eben hatten wir eine Orientierung von M definiert als eine Äquivalenzklasse von Tangentialvektorfeldern, welche nirgendwo verschwinden.

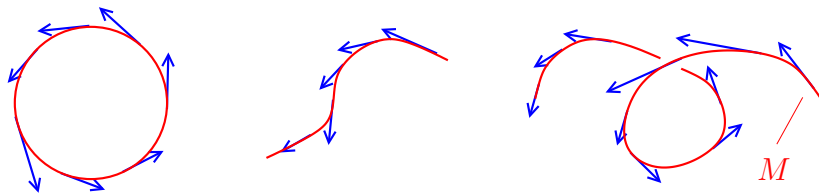


ABBILDUNG 52. Orientierungen von eindimensionalen Untermannigfaltigkeiten.

Die folgenden Abbildungen geben eine kanonische Bijektion zwischen diesen beiden Orientierungsbegriffen:

$$\begin{array}{ll}
 \text{Äquivalenzklassen von nirgendwo} & \leftrightarrow \text{Einheitsnormalenfelder auf } M \\
 \text{verschwindenden Tangentialvektorfeldern} & \\
 [\tau] & \xrightarrow{\Phi} \tau \text{ um } -\frac{\pi}{2} \text{ gedreht}^{91} \text{ und} \\
 & \text{auf Länge 1 normiert} \\
 [\nu \text{ um } \frac{\pi}{2} \text{ gedreht}] & \xleftarrow{\Psi} \nu.
 \end{array}$$

Diese beiden Abbildungen sind offensichtlich zu einander inverse. Im Folgenden werden wir kommentarlos zwischen diesen beiden Gesichtspunkten hin- und herwechseln.

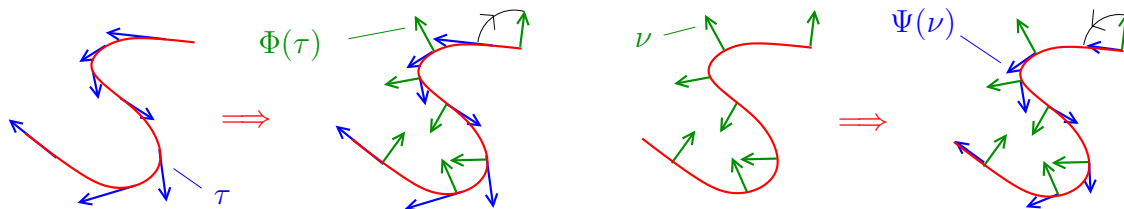


ABBILDUNG 53. Die Äquivalenz der Orientierungsbegriffe.

Es sei nun M eine zusammenhängende kompakte eindimensionale Untermannigfaltigkeit mit Rand. Es sei $\mathcal{O}$ eine Orientierung auf M und es sei ω eine 1-Form auf M . Wir definieren nun das Integral von ω auf der orientierten eindimensionalen Untermannigfaltigkeit M wie folgt. Wir wählen eine Parametrisierung $\gamma: [a, b] \rightarrow M$, welche die Eigenschaft besitzt, dass

⁹¹Wir drehen hierbei, wie üblich, gegen den Uhrzeigersinn.

$[\gamma'(t)] = \mathcal{O}$.⁹² Wir definieren dann

$$\int_M \omega := \int_{\gamma} \omega = \int_{t=a}^{t=b} \omega_{\gamma(t)}(\gamma'(t)) dt.$$

Es folgt aus Satz 8.11, dass

$$\int_{-M} \omega = - \int_M \omega.$$

Wie üblich müssen wir an dieser Stelle noch folgendes Lemma beweisen.

Lemma 8.12. *Die Definition von $\int_M \omega$ hängt nicht von der Wahl von γ ab.*

Beweis. Es sei also M eine zusammenhängende kompakte eindimensionale Untermannigfaltigkeit von $\mathbb{R}^n$ mit Rand, es sei $\mathcal{O}$ eine Orientierung auf M und es sei ω eine 1-Form auf M . Es seien $\gamma: [a, b] \rightarrow M$ und $\delta: [c, d] \rightarrow M$ zwei Parametrisierungen von M mit $[\gamma'(t)] = [\delta'(t)] = \mathcal{O}$.

Wir bezeichnen mit $\varphi: [a, b] \rightarrow [c, d]$ die Abbildung, welche gegeben ist durch die Gleichung $\varphi(t) := \gamma^{-1}(\delta(t))$. Dies ist ein Diffeomorphismus.⁹³ Aus unserer Voraussetzung $[\gamma'(t)] = [\delta'(t)] = \mathcal{O}$ folgt, dass φ eine orientierungserhaltende Umparametrisierung ist. Es folgt nun also aus Satz 8.11, dass die Integrale über γ und δ übereinstimmen. $\square$

Bemerkung. Wir haben jetzt also einen Orientierungsbegriff für Hyperflächen und für eindimensionale Untermannigfaltigkeiten. Wir werden später der Frage nachgehen, was denn eine Orientierung einer beliebigen Untermannigfaltigkeit sein soll.

Zudem haben wir das Integral einer 1-Form über einer orientierten eindimensionalen Untermannigfaltigkeit eingeführt, und wir wollen im Folgenden Kapitel das “richtige” mathematische Objekt finden, welches wir sinnvoll auf einer orientierten k -dimensionalen Untermannigfaltigkeit integrieren können.

⁹²Anschaulich gesprochen wählen wir also eine Kurve $\gamma: [a, b] \rightarrow M$, welche M einmal in der “vorgeschriebenen Richtung durchläuft”. Wir haben hier implizit die Aussage verwendet, dass es für jede zusammenhängende kompakte eindimensionale Untermannigfaltigkeit mit Rand eine solche Kurve γ gibt. Dies ist zwar anschaulich “irgendwie klar”, aber streng genommen gar nicht so leicht zu beweisen. Wir werden später noch präzise das Integral einer “ k -Form” auf einer orientierten kompakten k -dimensionalen Untermannigfaltigkeit einführen, und erlauben uns deswegen diese kleine Lücke.

⁹³Diese Aussage ist gar nicht so leicht zu beweisen. Die Abbildungen γ und δ sind Parametrisierungen, können also nach Satz 11.1 aus der Analysis II (siehe auch Satz 2.3) zu Karten Γ und Δ fortgesetzt werden. Diese sind insbesondere Diffeomorphismen zwischen offenen Teilmengen von $\mathbb{R}^n$, also ist auch $\Gamma^{-1} \circ \Delta$ ein Diffeomorphismus, also auch die Einschränkung von $\Gamma^{-1} \circ \Delta$ auf E_1 . Aber diese Einschränkung ist gerade die Abbildung $\varphi = \gamma^{-1} \circ \delta$.

9. DIFFERENTIALFORMEN HÖHERER ORDNUNG

9.1. **Motivation.** Erinnern wir uns noch einmal an den Gaußschen Integralsatz und an Satz 8.9:

Satz 9.1. (Gaußscher Integralsatz) *Es sei $M \subset \mathbb{R}^n$ eine kompakte n -dimensionale Untermannigfaltigkeit. Wir bezeichnen mit ν das äußere Einheitsnormalenfeld auf ∂M . Für jedes stetig differenzierbare Vektorfeld $F: M \rightarrow \mathbb{R}^n$ gilt*

$$\int_M \operatorname{div} F = \int_{\partial M} F \cdot \nu.$$

Wenden wir Satz 8.9 auf das Integral von 1-Formen auf eindimensionale Untermannigfaltigkeiten an (siehe Kapitel 8.3), so erhalten wir folgenden Satz:

Satz 9.2. *Es sei M eine zusammenhängende orientierte eindimensionale Untermannigfaltigkeit mit “Anfangspunkt P ” und “Endpunkt Q ”. Es sei zudem $f: M \rightarrow \mathbb{R}$ eine stetig differenzierbare Funktion. Dann gilt*

$$\int_M df = f(Q) - f(P).$$

Diese beiden Sätze sind auf den ersten Blick sehr unterschiedlich, aber sie besitzen gewisse formale Übereinstimmungen:

- (1) Beide Sätze können auf den Hauptsatz der Differential- und Integralrechnung aus der Analysis I zurückgeführt werden. Zudem können beide Sätze auch als Verallgemeinerung eben dieses Satzes aufgefasst werden, indem Sinne, dass wir für $M = [a, b]$ gerade die Aussage des Hauptsatzes der Differential- und Integralrechnung erhalten.
- (2) Wir betrachten also entweder

M eine n -Untermannigfaltigkeit von $\mathbb{R}^n$ und f ein Vektorfeld auf M , oder
 M eine 1-Untermannigfaltigkeit von $\mathbb{R}^n$ und f eine Funktion auf M .

Dann besagen die obigen Sätze, dass⁹⁴

$$\int_M \text{“Ableitung von } f\text{”} = \text{Integral von } f \text{ auf dem Rand von } M.$$

Für eine beliebige k -dimensionale Untermannigfaltigkeit wollen wir jetzt eine analoge Aussage (nämlich den Satz von Stokes) formulieren und beweisen. Dazu müssen wir uns aber zuerst überlegen, was ist das “richtige mathematische Objekt”, welches wir auf einer k -dimensionalen Untermannigfaltigkeit integrieren können.

Wir erinnern uns noch einmal an die Situation von 1-dimensionalen Untermannigfaltigkeiten. Für eine 1-Form ω auf einer 1-dimensionalen Untermannigfaltigkeit, welche eine

⁹⁴Der “klassische Satz von Stokes”, den wir schon auf Seite 79 formuliert hatten, ist übrigens von der gleichen Form.

Parametrisierung $\gamma: [a, b] \rightarrow M$ besitzt, hatten wir definiert

$$\int_M \omega := \int_{[a,b]} \underbrace{\omega_{\gamma(t)}}_{\in T_{\gamma(t)}^* M} \underbrace{(\gamma'(t))}_{\in T_{\gamma(t)} M}.$$

$\underbrace{\hspace{10em}}_{\in \mathbb{R}}$

Diese Definition macht Sinn, denn

- (1) eine Parametrisierung gibt für jeden Punkt auf M genau einen Tangentialvektor,
- (2) eine 1-Form auf M antwortet auf genau einen Tangentialvektor mit einer reellen Zahl,
- (3) wir erhalten also eine reellwertige Funktion, welche wir integrieren können,
- (4) und wir hatten gesehen, dass das Integral unabhängig von der Parametrisierung ist.

Es sei nun M eine k -dimensionale Untermannigfaltigkeit. Eine Parametrisierung $\Psi: T \rightarrow M$ liefert nun für jedes $Q \in T$ genau k Tangentialvektoren in $T_{f(Q)}M$. Wenn wir das obige Prozedere für höherdimensionale Untermannigfaltigkeiten durchziehen wollen, brauchen wir nun also auf M ein mathematisches Objekt, welches auf genau k Tangentialvektoren an einem Punkt mit einer reellen Zahl antwortet. Diese mathematischen Objekte sind die Multilinearformen, welche schon in der Linearen Algebra eingeführt wurden. Wir werden diese im nächsten Kapitel noch einmal einführen und genauer studieren.

9.2. Alternierende Multilinearformen I: Definition. Wir erinnern zuerst an folgenden Satz aus der Linearen Algebra.⁹⁵

Satz 9.3. *Es sei $n \in \mathbb{N}$. Es gibt genau eine Abbildung, genannt die Determinante,*

$$\det: M(n \times n, \mathbb{R}) \rightarrow \mathbb{R}$$

mit folgenden Eigenschaften:

- (i) *die Abbildung $\det$ ist multilinear, d.h. die Abbildung ist linear in jeder Spalte, genauer gesagt, für alle $v, v' \in \mathbb{R}^n$ und $l \in \mathbb{R}$ gilt⁹⁶*

$$\begin{aligned} \det(\dots v + v' \dots) &= \det(\dots v \dots) + \det(\dots v' \dots) \\ \det(\dots lv \dots) &= l \cdot \det(\dots v \dots), \end{aligned}$$

- (ii) *die Abbildung ist alternierend, d.h. vertauscht man zwei Spalten, so ändert sich das Vorzeichen, d.h. für alle $v, v' \in \mathbb{R}^n$ gilt*

$$\det(\dots v \dots v' \dots) = -\det(\dots v' \dots v \dots).$$

- (iii) *wenn $e_1, \dots, e_n$ die Standardbasis von $\mathbb{R}^n$ bezeichnen, dann gilt*

$$\det(e_1 \dots e_n) = 1.$$

In diesem Kapitel betrachten wir die alternierenden k -Formen auf Vektorräumen, welche als Verallgemeinerungen der Determinante aufgefasst werden können.

⁹⁵Der Satz wird in Kapitel 3.2.5 in Fischer: Lineare Algebra und analytische Geometrie bewiesen.

⁹⁶Hierbei sollen die nicht genannten Spalten auf beiden Seiten gleich sein.

Definition. Es sei V ein reeller Vektorraum. Eine *alternierende k -Form auf V* ist eine Abbildung

$$\begin{aligned}\omega: V^k &\rightarrow \mathbb{R} \\ (v_1, \dots, v_k) &\mapsto \omega(v_1, \dots, v_k)\end{aligned}$$

mit folgenden Eigenschaften:

- (i) ω ist multilinear, d.h. ω ist linear in jedem Argument. Dies bedeutet, dass für alle $v, v' \in V$ und $l \in \mathbb{R}$ gilt⁹⁷

$$\begin{aligned}\omega(\dots, v + v', \dots) &= \omega(\dots, v, \dots) + \omega(\dots, v', \dots) \\ \omega(\dots, lv, \dots) &= l \cdot \omega(\dots, v, \dots),\end{aligned}$$

- (ii) ω ist alternierend, d.h. vertauscht man zwei verschiedene Argumente, so ändert sich das Vorzeichen. Genauer gesagt, für alle $v, v' \in V$ gilt

$$\omega(\dots, v, \dots, v', \dots) = -\omega(\dots, v', \dots, v, \dots).$$

Beispiele.

- (1) Es sei $l \in \mathbb{R}$. Dann ist

$$\begin{aligned}(\mathbb{R}^n)^n &\rightarrow \mathbb{R} \\ (v_1, \dots, v_n) &\mapsto l \cdot \det(\underbrace{v_1 \dots v_n}_{n \times n\text{-Matrix}})\end{aligned}$$

eine alternierende n -Form auf $\mathbb{R}^n$. Es folgt leicht aus Satz 9.3, dass jede alternierende n -Form auf $\mathbb{R}^n$ von dieser Form ist.

- (2) Eine alternierende 1-Form ist eine lineare Abbildung $V \rightarrow \mathbb{R}$, d.h. ein Element im Dualraum $V^* = \text{Hom}(V, \mathbb{R})$. Wir nennen alternierende 1-Formen *Linearformen*.
 (3) Es sei $w \in \mathbb{R}^3$ ein Vektor. Dann ist⁹⁸

$$\begin{aligned}\mathbb{R}^3 \times \mathbb{R}^3 &\rightarrow \mathbb{R} \\ (v_1, v_2) &\mapsto w \cdot (\underbrace{v_1 \times v_2}_{\text{Kreuzprodukt}})\end{aligned}$$

eine alternierende 2-Form auf $\mathbb{R}^3$.

⁹⁷Hierbei sollen die nicht genannten Argumente auf beiden Seiten gleich sein.

⁹⁸Zur Erinnerung, dass Kreuzprodukt von zwei Vektoren

$$a = \begin{pmatrix} a_1 \\ a_2 \\ a_3 \end{pmatrix} \text{ und } b = \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix} \text{ ist definiert als } a \times b = \begin{pmatrix} a_2 b_3 - b_3 b_2 \\ -(a_1 b_3 - a_3 b_1) \\ a_1 b_2 - a_2 b_1 \end{pmatrix}$$

Das Kreuzprodukt ist bilinear und antisymmetrisch, d.h. es ist $b \times a = -a \times b$. In Übungsblatt 2 hatten wir gesehen, dass es folgende zwei geometrische Eigenschaften besitzt:

- (a) $a \times b$ ist orthogonal zu a und b ,
 (b) $\|a \times b\|$ ist das zweidimensionale Volumen des Parallelogramms, welches von a und b aufgespannt wird.

(4) Es sei $w \in \mathbb{R}^n$ ein Vektor. Dann ist

$$\begin{aligned} (\mathbb{R}^n)^{n-1} &\rightarrow \mathbb{R} \\ (v_1, \dots, v_{n-1}) &\mapsto \det \underbrace{\begin{pmatrix} w & v_1 & \dots & v_{n-1} \end{pmatrix}}_{n \times n\text{-Matrix}} \end{aligned}$$

eine alternierende $(n-1)$ -Form auf $\mathbb{R}^n$. Im Fall $n=3$ erhalten wir die gleiche alternierende 2-Form wie im vorherigen Beispiel.⁹⁹

Folgendes elementare Lemma wird später eine wichtige Rolle spielen.

Lemma 9.4. *Es sei η eine alternierende k -Form auf $\mathbb{R}^k$ und es sei A eine $k \times k$ -Matrix. Dann gilt für alle $v_1, \dots, v_k \in \mathbb{R}^k$, dass*

$$\eta(Av_1, \dots, Av_k) = \det(A) \cdot \eta(v_1, \dots, v_k).$$

Beweis. Nach der Diskussion auf Seite 107 gibt es ein $\lambda \in \mathbb{R}$, so dass

$$\eta(v_1, \dots, v_k) = \lambda \cdot \det(v_1 \dots v_k),$$

für alle $v_1, \dots, v_k \in \mathbb{R}^k$. Für beliebige $v_1, \dots, v_k \in \mathbb{R}^k$ folgt also, dass

$$\begin{aligned} \eta(Av_1, \dots, Av_k) &= \lambda \cdot \det(Av_1 \dots Av_k) \\ &= \lambda \cdot \det(A \cdot (v_1 \dots v_k)) \\ &= \lambda \cdot \det(A) \cdot \det(v_1 \dots v_k) = \det(A) \cdot \eta(v_1, \dots, v_k). \end{aligned}$$

↑

Multiplikativität der Determinante

□

Die Summe zweier alternierender k -Formen¹⁰⁰ und das Skalarprodukt einer alternierenden k -Form mit einer reellen Zahl ist wiederum eine alternierende k -Form. Die Menge der alternierenden k -Formen bildet also einen reellen Vektorraum, welchen wir mit $\wedge^k V^*$ bezeichnen.

Lemma 9.5. *Es sei V ein n -dimensionaler reeller Vektorraum. Dann gilt*

$$\begin{aligned} \wedge^1 V^* &= V^*, \\ \dim(\wedge^n V^*) &= 1, \\ \wedge^k V^* &= 0, \text{ für alle } k > n = \dim(V). \end{aligned}$$

Die ersten beiden Aussagen folgen aus der obigen Diskussion der Beispiele. Die dritte Aussage wird in Übungsblatt 5 bewiesen. Im nächsten Kapitel werden wir sehen, dass

$$\dim(\wedge^k V^*) = \binom{n}{k} \text{ für } k \in \{0, \dots, n\}.$$

⁹⁹Warum?

¹⁰⁰Wenn ω, ω' zwei alternierende k -Formen sind, dann ist die Summe $\omega + \omega'$ natürlich wie folgt definiert:

$$(\omega + \omega')(v_1, \dots, v_k) = \omega(v_1, \dots, v_k) + \omega'(v_1, \dots, v_k).$$

Im Folgenden verwenden wir auch die Konvention, dass

$$\wedge^0 V^* := \mathbb{R},$$

d.h. eine alternierende 0-Form auf V ist per Definition eine reelle Zahl.

9.3. Alternierende Multilinearformen II: Das Dachprodukt von Linearformen.

Im Folgenden sei V wiederum ein reeller Vektorraum. Es seien $\varphi_1, \dots, \varphi_k \in V^*$ Linearformen. Wir definieren

$$\begin{aligned} \varphi_1 \wedge \dots \wedge \varphi_k : V^k &\rightarrow \mathbb{R} \\ (v_1, \dots, v_k) &\mapsto \det \begin{pmatrix} \varphi_1(v_1) & \dots & \varphi_1(v_k) \\ \vdots & & \vdots \\ \varphi_k(v_1) & \dots & \varphi_k(v_k) \end{pmatrix}. \end{aligned}$$

Es ist leicht zu überprüfen, dass $\varphi_1 \wedge \dots \wedge \varphi_k$ eine alternierende k -Form ist. Wir bezeichnen diese als das *Dachprodukt*¹⁰¹¹⁰² von $\varphi_1, \dots, \varphi_k$.

Beispiel. Wir betrachten jetzt den Fall $V = \mathbb{R}^n$. Für $i = 1, \dots, n$ definieren wir¹⁰³

$$\begin{aligned} dx_i : \mathbb{R}^n &\rightarrow \mathbb{R} \\ (w_1, \dots, w_n) &\mapsto w_i. \end{aligned}$$

Dann gilt für Vektoren $v_i = (v_{i1}, \dots, v_{in}) \in \mathbb{R}^n$, $i = 1, \dots, n$, dass

$$\begin{aligned} (dx_1 \wedge \dots \wedge dx_n)(v_1, \dots, v_n) &= \det \begin{pmatrix} dx_1(v_1) & \dots & dx_1(v_n) \\ \vdots & & \vdots \\ dx_n(v_1) & \dots & dx_n(v_n) \end{pmatrix} \\ &= \det \begin{pmatrix} v_{11} & \dots & v_{n1} \\ \vdots & & \vdots \\ v_{1n} & \dots & v_{nn} \end{pmatrix} = \det \underbrace{\begin{pmatrix} v_1 & \dots & v_n \end{pmatrix}}_{n \times n\text{-Matrix}}. \end{aligned}$$

Wir haben also gezeigt, dass die alternierende n -Form

$$dx_1 \wedge \dots \wedge dx_n : (\mathbb{R}^n)^n = \mathbb{R}^{n^2} \rightarrow \mathbb{R}$$

gerade der Determinante entspricht.

Aus den Definitionen und den Eigenschaften der Determinante erhalten wir leicht folgendes Lemma.

¹⁰¹Der eigenwillige Name kommt von der Form des Symbols “ $\wedge$ ”. Im Englischen wird dieses Produkt das “wedge-product” genannt, nachdem das Symbol “ $\wedge$ ” dort als wedge = Keil interpretiert wird.

¹⁰²In der Literatur wird oft auch der Name *äußeres Produkt* verwendet.

¹⁰³Wir können also $dx_1, \dots, dx_n \in (\mathbb{R}^n)^* = \text{Hom}(\mathbb{R}^n, \mathbb{R})$ als die zur Standardbasis $e_1, \dots, e_n \in \mathbb{R}^n$ duale Basis auffassen.

Lemma 9.6. *Es V ein reeller Vektorraum. Die Abbildung*

$$\begin{aligned}(V^*)^k &\rightarrow \Lambda^k V^* \\ (\varphi_1, \dots, \varphi_k) &\rightarrow \varphi_1 \wedge \dots \wedge \varphi_k\end{aligned}$$

ist multilinear und alternierend. Insbesondere gilt für alle Linearformen φ und ψ , dass

$$\varphi \wedge \psi = -\psi \wedge \varphi.$$

Satz 9.7. *Es sei V ein n -dimensionaler reeller Vektorraum und es sei $\varphi_1, \dots, \varphi_n$ eine Basis für V^* . Es sei $k \in \{0, \dots, n\}$. Dann bilden die Dachprodukte*

$$\varphi_{i_1} \wedge \dots \wedge \varphi_{i_k} \quad \text{mit} \quad i_1 < i_2 < \dots < i_k$$

eine Basis von $\Lambda^k V^$.*

Der Satz wurde in der Linearen Algebra Vorlesung bewiesen. Wir geben einen Beweis der Vollständigkeit halber.

Beweis. Es sei $e_1, \dots, e_n$ eine Basis von V , welche zur Basis $\varphi_1, \dots, \varphi_n$ von V^* dual ist, d.h. es gilt

$$\varphi_i(e_j) = \delta_{ij} = \begin{cases} 1, & \text{wenn } i = j, \\ 0, & \text{wenn } i \neq j. \end{cases}$$

Aus der Definition des Dachproduktes folgt leicht, dass für $i_1 < \dots < i_k$ und $j_1 < \dots < j_k$, gilt

$$(*) \quad (\varphi_{i_1} \wedge \dots \wedge \varphi_{i_k})(e_{j_1}, \dots, e_{j_k}) = \begin{cases} 1, & \text{wenn } (i_1, \dots, i_k) = (j_1, \dots, j_k), \\ 0, & \text{andernfalls.} \end{cases}$$

Es folgt sofort, dass die gegebenen alternierenden k -Formen linear unabhängig sind.¹⁰⁴ Es sei nun $\omega \in \Lambda^k V^*$. Für $i_1 < \dots < i_k$ setzen wir

$$c_{i_1 \dots i_k} := \omega(e_{i_1}, \dots, e_{i_k}) \in \mathbb{R}.$$

Wir betrachten dann

$$\eta := \sum_{i_1 < \dots < i_k} c_{i_1 \dots i_k} \cdot \varphi_{i_1} \wedge \dots \wedge \varphi_{i_k}.$$

Es genügt nun zu zeigen, dass $\eta = \omega$.

- (1) Es folgt aus (*), dass ω und η die gleichen Werte auf allen $(e_{j_1}, \dots, e_{j_k}) \in V^k$ mit $j_1 < \dots < j_k$ annehmen.

¹⁰⁴In der Tat, denn für eine Linearkombination

$$0 = \sum_{i_1 < \dots < i_k} c_{i_1 \dots i_k} \cdot \varphi_{i_1} \wedge \dots \wedge \varphi_{i_k}$$

gilt für alle $j_1 < \dots < j_k$, dass

$$0 = \left(\sum_{i_1 < \dots < i_k} c_{i_1 \dots i_k} \cdot \varphi_{i_1} \wedge \dots \wedge \varphi_{i_k} \right)(e_{j_1}, \dots, e_{j_k}) = c_{j_1, \dots, j_k}.$$

- (2) Aus der definierenden Eigenschaft (ii) folgt, dass ω und η die gleichen Werte auf $(e_{j_1}, \dots, e_{j_k}) \in V^k$ für beliebige $j_1, \dots, j_k$ annehmen.
- (3) Aus der definierenden Eigenschaft (i) folgt nun, dass ω und η die gleichen Werte auf $(v_1, \dots, v_k) \in V^k$ für beliebige $v_1, \dots, v_k \in V$ annehmen.¹⁰⁵ $\square$

Wir erhalten sofort folgendes Korollar:

Korollar 9.8. *Es sei V ein n -dimensionaler reeller Vektorraum. Für jedes $k \in \{0, \dots, n\}$, gilt dass*

$$\dim(\wedge^k V^*) = \binom{n}{k} := \frac{n!}{(n-k)!k!}.$$

Beweis. Es ist

$$\dim(\wedge^k V^*) = \#\{(j_1, \dots, j_k) \mid j_1, \dots, j_k \in \{1, \dots, n\} \text{ mit } j_1 < \dots < j_k\} = \binom{n}{k}.$$

↑
die Anzahl der Möglichkeiten k verschiedene
Elemente aus einer Menge von n Elementen
zu wählen beträgt gerade $\binom{n}{k}$ $\square$

Bemerkung. Es sei V ein n -dimensionaler reeller Vektorraum und es sei $\varphi_1, \dots, \varphi_n$ eine Basis für V^* . Wir haben gesehen, dass die Dachprodukte

$$\varphi_{i_1} \wedge \dots \wedge \varphi_{i_k} \text{ mit } i_1 < i_2 < \dots < i_k$$

eine Basis von $\wedge^k V^*$ bilden. Dies bedeutet jedoch *nicht*, dass jede alternierende k -Form als Dachprodukt von k Linearformen geschrieben werden kann. Wir werden dazu ein Beispiel in Übungsblatt 6 behandeln.¹⁰⁶

Satz 9.9. *Es sei V ein endlich dimensionaler reeller Vektorraum und es seien $k, l \in \mathbb{N}_0$. Dann gibt es genau eine Abbildung*

$$\begin{aligned} \wedge^k V^* \times \wedge^l V^* &\rightarrow \wedge^{k+l} V^* \\ (\omega, \sigma) &\mapsto \omega \wedge \sigma \end{aligned}$$

mit folgenden Eigenschaften:

¹⁰⁵In der Tat, denn für $i = 1, \dots, k$ schreiben wir $v_i = (v_{i1}, \dots, v_{in})$. Dann gilt

$$\begin{aligned} \omega(v_1, \dots, v_k) &= \omega\left(\sum_{j_1=1}^n v_{1j_1} e_{j_1}, \dots, \sum_{j_k=1}^n v_{kj_k} e_{j_k}\right) \stackrel{\text{Eigenschaft (i) angewandt auf die Spalten } 1, \dots, k}{=} \sum_{j_1, \dots, j_k=1}^n v_{1j_1} \cdots v_{kj_k} \omega(e_{j_1}, \dots, e_{j_k}) \\ &= \sum_{j_1, \dots, j_k=1}^n v_{1j_1} \cdots v_{kj_k} \omega(e_{j_1}, \dots, e_{j_k}) \stackrel{\text{siehe (2) oben}}{=} \eta(v_1, \dots, v_k). \end{aligned}$$

↑
gleiche Argument rückwärts

¹⁰⁶Die ist analog zum gerne gemachten Fehler zu denken, dass in einem Tensorprodukt $V \otimes W$ jedes Element von der Form $v \otimes w$ sein muss.

- (1) die Abbildung ist bilinear, d.h. für alle $\omega, \omega_1, \omega_2 \in \wedge^k V^*$, für alle $\sigma, \sigma_1, \sigma_2 \in \wedge^l V^*$ sowie für alle $\lambda \in \mathbb{R}$ gilt

$$\begin{aligned} (\omega_1 + \omega_2) \wedge \sigma &= \omega_1 \wedge \sigma + \omega_2 \wedge \sigma, \\ \omega \wedge (\sigma_1 + \sigma_2) &= \omega \wedge \sigma_1 + \omega \wedge \sigma_2 \end{aligned}$$

und

$$(\lambda \omega) \wedge \sigma = \lambda(\omega \wedge \sigma) = \omega \wedge (\lambda \sigma).$$

- (2) für $\psi_1, \dots, \psi_k, \eta_1, \dots, \eta_l \in V^*$ gilt

$$(\psi_1 \wedge \dots \wedge \psi_k) \wedge (\eta_1 \wedge \dots \wedge \eta_l) = \psi_1 \wedge \dots \wedge \psi_k \wedge \eta_1 \wedge \dots \wedge \eta_l.$$

Auch dieser Satz wurde in der Linearen Algebra bewiesen, wir skizzieren einen Beweis der Vollständigkeit halber.

Beweis. Es sei $\varphi_1, \dots, \varphi_n$ eine Basis von V^* . Es sei $\omega \in \wedge^k V^*$ und es sei $\sigma \in \wedge^l V^*$. Nach Satz 9.7 können wir schreiben

$$\omega = \sum_{i_1 < \dots < i_k} a_{i_1, \dots, i_k} \cdot \varphi_{i_1} \wedge \dots \wedge \varphi_{i_k} \text{ und } \sigma = \sum_{j_1 < \dots < j_l} b_{j_1, \dots, j_l} \cdot \varphi_{j_1} \wedge \dots \wedge \varphi_{j_l},$$

wobei die Koeffizienten $a_{i_1, \dots, i_k}$ und $b_{j_1, \dots, j_l}$ eindeutig bestimmt sind. Wir definieren¹⁰⁷

$$\omega \wedge \sigma := \sum_{\substack{i_1 < \dots < i_k \\ j_1 < \dots < j_l}} a_{i_1, \dots, i_k} b_{j_1, \dots, j_l} \cdot \varphi_{i_1} \wedge \dots \wedge \varphi_{i_k} \wedge \varphi_{j_1} \wedge \dots \wedge \varphi_{j_l}.$$

Man kann nun leicht zeigen, dass diese Abbildung die Eigenschaften (1) und (2) besitzt. Zudem ist es ebenso leicht zu sehen, dass dies die einzige Abbildung ist, welche die Eigenschaften (1) und (2) erfüllt. $\square$

Im Folgenden bezeichnen wir $\omega \wedge \sigma$ als das *Dachprodukt* von ω und σ . Das folgende Lemma besagt, dass das Dachprodukt assoziativ ist, und dass es bis auf ein geeignetes Vorzeichen kommutativ ist.

Lemma 9.10. *Das Dachprodukt auf einem endlich dimensional reellen Vektorraum V besitzt auch folgende Eigenschaften:*

- (3) Für $\alpha \in \wedge^k V^*$, $\beta \in \wedge^l V^*$ und $\gamma \in \wedge^m V^*$ gilt

$$(\alpha \wedge \beta) \wedge \gamma = \alpha \wedge (\beta \wedge \gamma).$$

- (4) Für $\alpha \in \wedge^k V^*$ und $\beta \in \wedge^l V^*$ gilt

$$\beta \wedge \alpha = (-1)^{kl} \alpha \wedge \beta.$$

¹⁰⁷Wir definieren $\omega \wedge \sigma$ also durch den Ausdruck, in dem wir die Ausdrücke für ω und σ naiv auseinander multiplizieren. Wenn die Aussage des Satzes so gilt, dann folgt aus den Eigenschaften (1) und (2), dass diese Gleichheit gelten muss.

Beweis. Die Aussage (3) folgt leicht aus (2) unter Verwendung von Satz 9.7 und (1). Wir überlassen den genauen Beweis als freiwillige Übungsaufgabe.

Wir wenden uns nun dem Beweis der Aussage (4) zu. Wir wählen eine Basis $\varphi_1, \dots, \varphi_n$ von V^* . Aussage (4) folgt nun Satz 9.7, aus Aussage (2) und folgender Behauptung:

Behauptung. Es seien $1 \leq i_1 < \dots < i_k \leq n$ und $1 \leq j_1 < \dots < j_l \leq n$. Dann ist

$$\underbrace{\varphi_{j_1} \wedge \dots \wedge \varphi_{j_l}} \wedge \underbrace{\varphi_{i_1} \wedge \dots \wedge \varphi_{i_k}} = (-1)^{kl} \underbrace{\varphi_{i_1} \wedge \dots \wedge \varphi_{i_k}} \wedge \underbrace{\varphi_{j_1} \wedge \dots \wedge \varphi_{j_l}}.$$

Wir erinnern zuerst daran, dass $\varphi_i \wedge \varphi_j = -\varphi_j \wedge \varphi_i$. Mithilfe dieser Beobachtung erhalten wir nun

$$\begin{aligned} \varphi_{j_1} \wedge \dots \wedge \varphi_{j_l} \wedge \varphi_{i_1} \wedge \dots \wedge \varphi_{i_k} &= (-1)^l \varphi_{j_1} \wedge \dots \wedge \varphi_{j_{l-1}} \wedge \varphi_{i_1} \wedge \dots \wedge \varphi_{i_k} \wedge \varphi_{j_l} \\ &\quad \uparrow \\ &\text{denn wir benötigen } k \text{ Transpositionen um } \varphi_{j_l} \text{ "nach hinten" zu schieben} \\ &= (-1)^{kl} \varphi_{i_1} \wedge \dots \wedge \varphi_{i_k} \wedge \varphi_{j_1} \wedge \dots \wedge \varphi_{j_l}. \\ &\quad \uparrow \\ &\text{wir führen das gleiche Argument insgesamt } l \text{ mal durch} \end{aligned}$$

□

9.4. Differentialformen auf Untermannigfaltigkeiten. Es sei M eine Untermannigfaltigkeit. Eine *Differentialform k -ter Ordnung ω auf M* (oder auch kurz *k -Form ω auf M*) ordnet jedem Punkt $P \in M$ eine alternierende k -Form ω_P auf dem Tangentialraum $T_P M$ zu.¹⁰⁸

Beispiele.

- (1) Nachdem $\wedge^1 V^* = V^*$ entspricht die obige Definition einer k -Form für $k = 1$ gerade der Definition auf Seite 90.
- (2) Es folgt aus der Konvention, dass $\wedge^0 V^* = \mathbb{R}$, dass eine 0-Form auf einer Untermannigfaltigkeit M das gleiche ist wie eine Funktion $f: M \rightarrow \mathbb{R}$.
- (3) Es sei $M \subset \mathbb{R}^n$ eine Untermannigfaltigkeit beliebiger Dimension. Darüber hinaus sei $F: M \rightarrow \mathbb{R}^n$ ein Vektorfeld auf M . Für $P \in M$ betrachten wir wie in Kapitel 9.2 die alternierende $(n-1)$ -Form $\delta(F)_P$, welche gegeben ist durch¹⁰⁹

$$\begin{aligned} \delta(F)_P: (T_P M)^{n-1} &\rightarrow \mathbb{R} \\ (v_1, \dots, v_{n-1}) &\mapsto \det \underbrace{(F(P) v_1 \dots v_{n-1})}_{n \times n \text{-Matrix}}. \end{aligned}$$

¹⁰⁸Genauer gesagt ist eine Differentialform k -ter Ordnung ω eine Abbildung

$$\omega: M \rightarrow \bigcup_{P \in M} \wedge^k T_P^* M,$$

so dass für $P \in M$ gilt, dass $\omega_P \in \wedge^k T_P^* M$.

¹⁰⁹Zur Erinnerung, $T_P M$ ist ein Untervektorraum von $\mathbb{R}^n$, die Vektoren $F(P), v_1, \dots, v_{n-1}$ bilden also in der Tat eine reelle $n \times n$ -Matrix.

Insbesondere, wenn (M, ν) eine orientierte Hyperfläche in $\mathbb{R}^3$ ist, dann definieren die alternierenden 2-Formen

$$\begin{aligned} T_P M \times T_P M &\rightarrow \mathbb{R} \\ (v, w) &\mapsto \nu(P) \cdot (v \times w) \end{aligned}$$

eine 2-Form auf M .

In Kapitel 8.1 hatten wir schon diskutiert, was es heißen soll, dass eine Differentialform erster Ordnung stetig, stetig differenzierbar beziehungsweise glatt ist. Wir werden in den nächsten beiden Kapiteln diese Begriffe auf Differentialformen beliebiger Ordnung verallgemeinern.

9.5. Induzierte Abbildungen. Wir hatten schon gesehen, dass eine lineare Abbildung $f: U \rightarrow V$ zwischen reellen Vektorräumen folgende Abbildung zwischen den Dualräumen induziert:

$$\begin{aligned} f^*: V^* &\rightarrow U^* \\ (\varphi: V \rightarrow \mathbb{R}) &\mapsto \left(\begin{array}{l} U \rightarrow \mathbb{R} \\ u \mapsto \varphi(f(u)) \end{array} \right). \end{aligned}$$

Wir hatten zudem gesehen, dass $V \mapsto V^*$ und $\phi \mapsto \phi^*$ einen kontravarianten Funktor von der Kategorie der reellen Vektorräume zu der Kategorie der reellen Vektorräume bildet.

Man kann diese Konstruktion leicht verallgemeinern. Eine lineare Abbildung $f: U \rightarrow V$ zwischen reellen Vektorräumen induziert eine Abbildung

$$\begin{aligned} f^*: \wedge^k V^* &\rightarrow \wedge^k U^* \\ (\omega: V^k \rightarrow \mathbb{R}) &\mapsto \left(\begin{array}{l} f^* \omega: U^k \rightarrow \mathbb{R} \\ (u_1, \dots, u_k) \mapsto \omega(f(u_1), \dots, f(u_k)) \end{array} \right). \end{aligned}$$

Man kann mithilfe der Definitionen leicht zeigen, dass diese Abbildungen das Dachprodukt erhalten, d.h. für alle $\alpha \in \wedge^k V^*$ und $\beta \in \wedge^l V^*$ gilt¹¹⁰

$$f^*(\alpha \wedge \beta) = f^* \alpha \wedge f^* \beta.$$

Man kann nun leicht verifizieren, dass die Abbildungen $V \mapsto \wedge^k V^*$ und $\phi \mapsto \phi^*$ einen kontravarianten Funktor von der Kategorie der reellen Vektorräume zu der Kategorie der reellen Vektorräume definiert.¹¹¹

Es sei nun $f: M \rightarrow N$ eine glatte Abbildung zwischen Untermannigfaltigkeiten. Es sei $P \in M$. Wir hatten schon in Kapitel 7.1 gesehen, dass f eine lineare Abbildung

$$\begin{aligned} f_*: T_P M &\rightarrow T_{f(P)} N \\ [\gamma] &\mapsto [f \circ \gamma] \end{aligned}$$

¹¹⁰Auch diese Aussage wurde in der Linearen Algebra bewiesen.

¹¹¹Man kann diese Aussage auch noch etwas verfeinern. Für einen reellen Vektorraum V schreiben wir jetzt $\wedge^* V^* = \bigoplus_{n \in \mathbb{N}_0} \wedge^n V^*$. Dann induziert das Dachprodukt auf $\wedge^* V^*$ eine Ringstruktur. Die Abbildungen $U \mapsto \wedge^* U^*$ und $f \mapsto f^*$ definieren nun einen kontravarianten Funktor von der Kategorie der Vektorräume zu der Kategorie der Ringe.

induziert. Zudem bezeichnen wir mit f^* die Abbildung

$$\begin{aligned} T_{f(P)}^* N &\rightarrow T_P^* M \\ \varphi &\mapsto \left(\begin{array}{ccc} f^* \varphi: T_P M &\rightarrow & \mathbb{R} \\ v &\mapsto & \varphi(f_*(v)) \end{array} \right). \end{aligned}$$

Allgemeiner bezeichnen wir mit f^* auch die induzierte Abbildung

$$\begin{aligned} f^*: \wedge^k T_{f(P)}^* N &\rightarrow \wedge^k T_P^* M \\ \omega &\mapsto \left(\begin{array}{ccc} f^* \omega: (T_P M)^k &\rightarrow & \mathbb{R} \\ (v_1, \dots, v_k) &\mapsto & \omega(f_*(v_1), \dots, f_*(v_k)) \end{array} \right). \end{aligned}$$

Es ist leicht zu sehen, dass

$$\begin{array}{ccc} \text{Kategorie } \mathcal{P} \text{ der punktierten} & \rightarrow & \text{Kategorie der reellen} \\ \text{Untermannigfaltigkeiten} & & \text{Vektorräume} \end{array}$$

mit

$$(M, P) \mapsto \wedge^k T_P^* M$$

und

$$(f: (M, P) \rightarrow (N, Q)) \mapsto (f^*: \wedge^k T_Q^* N \rightarrow \wedge^k T_P^* M)$$

ein kontravarianter Funktor ist.¹¹²

9.6. Glatte Differentialformen auf Untermannigfaltigkeiten. Wir wollen nun definieren, was es heißt, dass eine Differentialform auf einer Untermannigfaltigkeit M glatt ist. Für Differentialformen erster Ordnung hatten wir dies schon auf Seite 92 ausgeführt. Die Definitionen für Differentialformen höherer Ordnung verlaufen nun ganz analog.

Definition.

- (1) Es sei ω eine l -Form auf einer offenen Teilmenge $W \subset \mathbb{R}^k$ beziehungsweise einer offenen Teilmenge $W \subset H_k$. Nach Satz 9.7 bilden die Dachprodukte $dx_{i_1} \wedge \dots \wedge dx_{i_l}$ mit $i_1 < i_2 < \dots < i_l$ eine Basis von $\wedge^l(\mathbb{R}^k)^*$. Wir können also ω wie folgt als Linearkombination betrachten:

$$\omega = \sum_{i_1 < \dots < i_l} f_{i_1 \dots i_l} dx_{i_1} \wedge \dots \wedge dx_{i_l},$$

wobei $f_{i_1 \dots i_l}$ reellwertige Funktionen auf W sind. Wir sagen nun, die l -Form ω ist glatt, wenn alle Funktionen $f_{i_1 \dots i_l}$ glatt sind.

- (2) Es sei ω eine l -Form auf einer k -dimensionalen Untermannigfaltigkeit M .
 - (a) Es sei $\Phi: U \rightarrow V$ eine Karte für M . Wir sagen ω ist glatt bezüglich der Karte Φ , falls die l -Form $(\Phi^{-1})^* \omega$ auf $V \cap E_k$ beziehungsweise auf $V \cap H_k$ glatt ist.
 - (b) Wir sagen ω ist glatt, wenn es einen Atlas von M gibt, so dass ω bezüglich allen Karten des Atlanten glatt ist.

¹¹²Beispielsweise ist dies die Verknüpfung des kovarianten Funktors $(M, P) \rightarrow T_P M$ mit dem kontravarianten Funktor $V \mapsto \wedge^k V^*$.

Wir verwenden auch die ganz analogen Definitionen mit “glatt” ersetzt durch “stetig” ersetzen.

Der folgende Satz wird ganz ähnlich wie Satz 8.3 bewiesen.

Satz 9.11. *Es sei ω eine l -Form auf einer Untermannigfaltigkeit M . Dann sind die folgenden Aussagen äquivalent:*

- (1) *die l -Form ω ist glatt auf M ,*
- (2) *die l -Form ist glatt bezüglich jeder Karte von M ,*
- (3) *für alle glatten Tangentialvektorfelder $F_1, \dots, F_l$ auf M ist die Funktion*

$$\begin{aligned} M &\rightarrow \mathbb{R} \\ P &\mapsto \omega_P(F_1(P), \dots, F_l(P)) \end{aligned}$$

glatt.

Die gleichen Aussagen gelten auch, wenn wir “glatt” durch “stetig” ersetzen.

Für eine Untermannigfaltigkeit M definieren wir nun

$$\mathcal{C}_l(M) := \{\text{stetige } l\text{-Formen auf } M\},$$

und

$$\Omega_l(M) := \{\text{glatte } l\text{-Formen auf } M\}.$$

Beispiel. Es sei $M \subset \mathbb{R}^n$ eine Untermannigfaltigkeit und es sei $F: M \rightarrow \mathbb{R}^n$ ein glattes Vektorfeld auf M . Wir betrachten wiederum die $(n-1)$ -Form auf M , welche $P \in M$ folgende alternierende $(n-1)$ -Form zuordnet:

$$\begin{aligned} (T_P M)^{n-1} &\rightarrow \mathbb{R} \\ (v_1, \dots, v_{n-1}) &\mapsto \det \underbrace{(F(P) v_1 \dots v_{n-1})}_{n \times n - \text{Matrix}}. \end{aligned}$$

In Übungsblatt 6 werden wir sehen, dass dies eine glatte $(n-1)$ -Form auf M ist.

Der Beweis von folgendem Lemma ist ganz analog zum Beweis von Satz 8.6.

Satz 9.12. *Es sei $f: M \rightarrow N$ eine glatte Abbildung zwischen zwei Untermannigfaltigkeiten und es sei ω eine l -Form auf N . Wenn ω glatt ist, dann ist auch $f^*\omega$ wiederum glatt. Die gleiche Aussage gilt auch, wenn wir “glatt” durch “stetig” ersetzen.*

Zusammenfassend haben wir also insbesondere gezeigt, dass die Abbildungen

$$\text{Untermannigfaltigkeit } M \mapsto \Omega_l(M),$$

zusammen mit den Abbildungen

$$\text{glatte Abbildung } f: M \rightarrow N \mapsto (f^*: \Omega_l(N) \rightarrow \Omega_l(M))$$

einen kontravarianten Funktor von der Kategorie der Untermannigfaltigkeiten zu der Kategorie der Vektorräume bilden. Die gleiche Aussage gilt auch, wenn wir $\Omega_l(M)$ durch $\mathcal{C}_l(M)$ ersetzen.

9.7. Integration von Differentialformen I. In Kapitel 8.4 hatten wir gesehen, dass wir eine stetige 1-Form auf einer orientierten kompakten 1-dimensionalen Untermannigfaltigkeit mit Rand integrieren können.

Es sei nun M eine kompakte k -dimensionale Untermannigfaltigkeit und es sei $\omega \in \mathcal{C}_k(M)$ eine stetige k -Form auf M . In diesem Kapitel nehmen wir erst einmal an, dass es eine Karte $\Phi: U \rightarrow V$ gibt mit $\text{Träger}(\omega) \subset U$.¹¹³ Wir bezeichnen mit $\Psi: V \rightarrow U$ die Umkehrabbildung von Φ . Es sei nun $x \in V \cap E_k$. Wir betrachten nun, ganz analog zur Diskussion in Kapitel 8.2, die Funktion

$$\begin{aligned} V \cap E_k &\rightarrow \mathbb{R} \\ x &\mapsto \underbrace{\omega_{\Psi(x)}}_{\in \wedge^k T_{\Psi(x)}^* M} \left(\underbrace{\Psi_* e_1, \dots, \Psi_* e_k}_{k \text{ Vektoren in } T_{\Psi(x)} M} \right). \end{aligned}$$

Dies ist eine stetige Funktion mit kompakten Träger, also Lebesgue-integrierbar.¹¹⁴

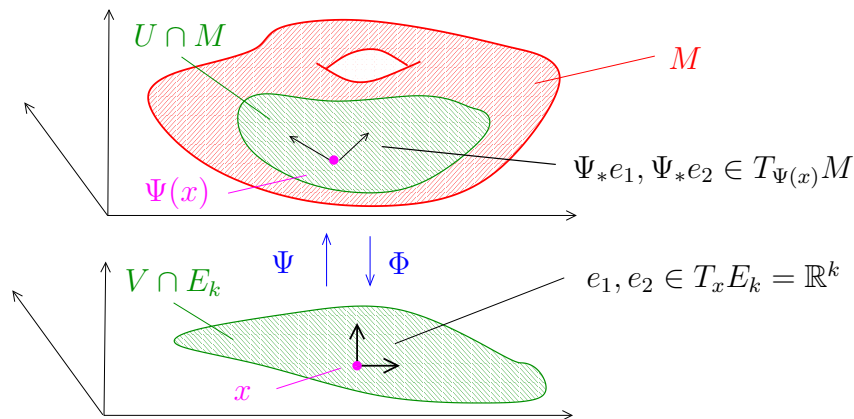


ABBILDUNG 54. Schematisches Bild für die Integration von k -Formen auf einer k -dimensionalen Untermannigfaltigkeit.

Wir können nun diese Funktion integrieren, d.h. wir betrachten

$$I(\Psi) := \int_{V \cap E_k} \omega_{\Psi(x)}(\Psi_* e_1, \dots, \Psi_* e_k) dx.$$

Es stellt sich die Frage, ob dieser Ausdruck von der Wahl der Karte Φ abhängt oder nicht.

Frage. Es sei M eine k -dimensionale Untermannigfaltigkeit und es sei $\omega \in \mathcal{C}_k(M)$. Zudem seien $\Phi: U \rightarrow V$ und $\tilde{\Phi}: \tilde{U} \rightarrow \tilde{V}$ Karten mit $\text{Träger}(\omega) \subset U \cap \tilde{U}$. Wir bezeichnen die

¹¹³Hierbei bezeichnen wir mit $\text{Träger}(\omega) \subset M$ die Menge aller Punkte $P \in M$, so dass die Differentialform auf $T_P M$ nicht die Nullform ist.

¹¹⁴Für $x \in V \cap E_k$ ist

$$\omega_{\Psi(x)}(\Psi_* e_1, \dots, \Psi_* e_k) = (\Psi^* \omega)_x(e_1, \dots, e_k),$$

die k -Form $\Psi^* \omega$ ist stetig nach Satz 9.12, also ist die Funktion stetig nach Satz 9.11. Da M kompakt ist, ist auch der Träger von ω und damit auch der Träger von $\Psi^* \omega$ kompakt.

Umkehrabbildungen von Φ und $\tilde{\Phi}$ mit Ψ und $\tilde{\Psi}$. Die Frage ist nun wie folgt: Unter welchen Voraussetzungen stimmen die Integrale

$$I(\Psi) = \int_{V \cap E_k} \omega_{\Psi(x)}(\Psi_* e_1, \dots, \Psi_* e_k) dx$$

und

$$I(\tilde{\Psi}) = \int_{\tilde{V} \cap E_k} \omega_{\tilde{\Psi}(y)}(\tilde{\Psi}_* e_1, \dots, \tilde{\Psi}_* e_k) dy$$

überein?

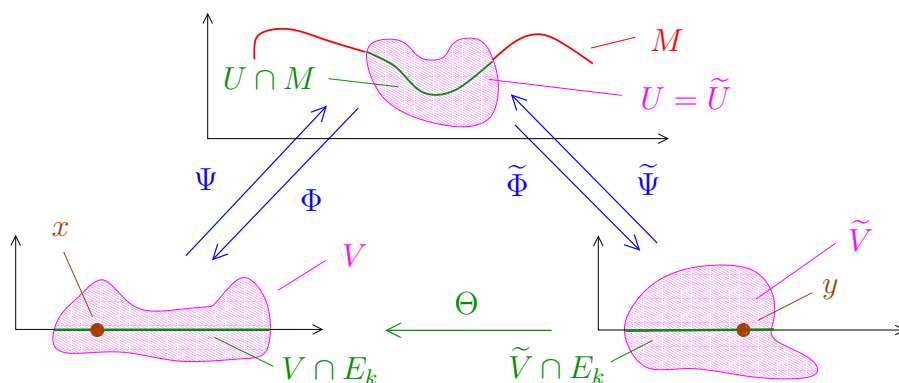


ABBILDUNG 55. Schematisches Bild für zwei Karten für M .

Antwort. Wir können o.B.d.A. annehmen, dass $U = \tilde{U}$. Wir bezeichnen mit Θ den Diffeomorphismus $\Theta := \Phi \circ \tilde{\Psi}: \tilde{V} \cap E_k \rightarrow V \cap E_k$. Zum einen gilt nun, dass

$$\begin{aligned} I(\tilde{\Psi}) &= \int_{\tilde{V} \cap E_k} \omega_{\tilde{\Psi}(y)}(D_y \tilde{\Psi}(e_1), \dots, D_y \tilde{\Psi}(e_k)) dy \\ &= \int_{\tilde{V} \cap E_k} \omega_{\tilde{\Psi}(y)}(D_{\Theta(y)} \Psi(D_y \Theta \cdot e_1), \dots, D_{\Theta(y)} \Psi(D_y \Theta \cdot e_k)) dy \\ &\quad \uparrow \tilde{V} \cap E_k \end{aligned}$$

aus der Kettenregel folgt, dass für $y \in \tilde{V} \cap E_k$ gilt:

$$D_y \tilde{\Psi} = D_y(\Psi \circ \Theta) = D_{\Theta(y)} \Psi \cdot D_y \Theta$$

$$\begin{aligned} &= \int_{\tilde{V} \cap E_k} \omega_{\Psi(\Theta(y))}(D_{\Theta(y)} \Psi(e_1), \dots, D_{\Theta(y)} \Psi(e_k)) \cdot \det(D_y \Theta) dy \\ &\quad \uparrow \tilde{V} \cap E_k \end{aligned}$$

Lemma 9.4 angewandt auf die alternierende k -Form

$$(v_1, \dots, v_k) \mapsto \eta(v_1, \dots, v_k) := \omega_{\tilde{\Psi}(y)}(D_{\Theta(y)} \Psi(v_1), \dots, D_{\Theta(y)} \Psi(v_k))$$

Andererseits gilt

$$\begin{aligned} I(\Psi) &= \int_{V \cap E_k} \omega_{\Psi(x)}((D_x \Psi)(e_1), \dots, (D_x \Psi)(e_k)) dx \\ &= \int_{\tilde{V} \cap E_k} \omega_{\Psi(\Theta(y))}((D_{\Theta(y)} \Psi)(e_1), \dots, (D_{\Theta(y)} \Psi)(e_k)) \cdot |\det D_y \Theta| dy. \\ &\quad \uparrow \end{aligned}$$

der Transformationssatz 3.9 besagt, dass für eine Funktion g auf $V \cap E_k$ gilt

$$\int_{V \cap E_k} g(x) dx = \int_{\tilde{V} \cap E_k} g(\Theta(y)) \cdot |\det D\Theta(y)| dy$$

Wir sehen also, dass die beiden Integrale $I(\Psi)$ und $I(\tilde{\Psi})$ übereinstimmen, wenn

$$\det(D_y \Theta) = |\det(D_y \Theta)| \quad \text{für alle } y \in \tilde{V} \cap E_k,$$

d.h. beide Integrale stimmen überein, wenn

$$\det(D_y \Theta) > 0 \quad \text{für alle } y \in \tilde{V} \cap E_k.$$

Wir haben damit unsere Frage beantwortet.

Wir fassen jetzt diese Rechnung als folgenden Satz zusammen.

Satz 9.13. *Es sei M eine k -dimensionale Untermannigfaltigkeit und es sei $\omega \in \mathcal{C}_k(M)$. Zudem seien $\Phi: U \rightarrow V$ und $\tilde{\Phi}: U \rightarrow \tilde{V}$ Karten mit $\text{Träger}(\omega) \subset U$. Wir nehmen an, dass $U \cap M$ zusammenhängend ist. Wir betrachten den Diffeomorphismus*

$$\Theta := \Phi \circ \tilde{\Phi}^{-1}: \tilde{V} \cap E_k \rightarrow V \cap E_k.$$

Wir setzen¹¹⁵

$$\epsilon := \begin{cases} +1, & \text{wenn } \det(D_y \Theta) > 0 \text{ für alle } y \in \tilde{V} \cap E_k, \\ -1, & \text{wenn } \det(D_y \Theta) < 0 \text{ für alle } y \in \tilde{V} \cap E_k. \end{cases}$$

Dann gilt

$$\int_{V \cap E_k} \omega_{\Psi(x)}(\Psi_* e_1, \dots, \Psi_* e_k) dx = \epsilon \cdot \int_{\tilde{V} \cap E_k} \omega_{\tilde{\Psi}(y)}(\tilde{\Psi}_* e_1, \dots, \tilde{\Psi}_* e_k) dy.$$

Nachdem also im Allgemeinen das Integral

$$I(\Psi) := \int_{V \cap E_k} \omega_{\Psi(x)}(\Psi_*(e_1), \dots, \Psi_*(e_k)) dx$$

von der Wahl der Karte abhängt, können wir den Ausdruck nicht als Definition des Integrals von einer k -Form auf einer k -dimensionalen Untermannigfaltigkeit verwenden. Andererseits haben wir gesehen, dass wir “nahe dran” sind, d.h. wir müssen noch versuchen, das “Vorzeichenproblem” unter Kontrolle zu kriegen. Um dies zu erreichen benötigen wir noch den Begriff des orientierten Vektorraumes und der orientierten Untermannigfaltigkeit.

¹¹⁵Warum tritt genau einer der beiden Fälle in der Definition von ϵ auf?

Bemerkung. Bei der Diskussion des Integrals einer Differentialform auf einer eindimensionalen Untermannigfaltigkeit M hatten wir ebenfalls gesehen, dass das Integral von der Wahl einer Orientierung von M abhängt. Es ist also nicht überraschend, dass diese Problematik auch bei der Integration von k -Formen auf k -dimensionalen Untermannigfaltigkeiten auftritt.

9.8. Orientierte Vektorräume und Untermannigfaltigkeiten. Im Folgenden werden wir bis auf Widerruf folgende Konvention verwenden: wir betrachten nur Vektorräume der Dimension größer Null, und wir betrachten nur Untermannigfaltigkeiten der Dimension größer Null.

Definition. Es sei V ein k -dimensionaler reeller Vektorraum. Wir sagen zwei Basen¹¹⁶ $\{v_1, \dots, v_k\}$ und $\{w_1, \dots, w_k\}$ sind *äquivalent*, wenn die Determinante des Basiswechsels¹¹⁷ positiv ist.

Es ist eine leichte Übungsaufgabe zu zeigen, dass dies in der Tat eine Äquivalenzrelation auf der Menge der Basen von V ist, und dass es genau zwei Äquivalenzklassen von Basen gibt.

Definition. Es sei V ein k -dimensionaler reeller Vektorraum.

- (1) Eine Äquivalenzklasse von Basen heißt *Orientierung* für V .
- (2) Wir bezeichnen V zusammen mit einer Orientierung als *orientierten Vektorraum*.
- (3) Es sei $(V, \mathcal{O})$ ein orientierter Vektorraum. Wir nennen $\{v_1, \dots, v_k\}$ eine *positive Basis*, wenn $\{v_1, \dots, v_k\}$ in $\mathcal{O}$ liegt, andernfalls sagen wir, dass $\{v_1, \dots, v_k\}$ eine *negative Basis* ist.

Beispiele.

- (1) Wenn wir nichts anderes sagen, dann betrachten wir $\mathbb{R}^n$ immer mit der Orientierung, in der die “kanonische” Basis $\{e_1, \dots, e_n\}$ eine positive Basis ist.
- (2) Wir betrachten $\mathbb{R}^2$ mit der kanonischen Orientierung, welche durch die Basis $\{e_1, e_2\}$ festgelegt ist. Es sei nun $\{v_1, v_2\}$ eine beliebige Basis für $\mathbb{R}^2$. Es sei $\varphi \in (-\pi, \pi)$ der Winkel um den man v_1 gegen den Uhrzeigersinn drehen muss, so dass v_1 in die gleiche Richtung wie v_2 zeigt.¹¹⁸ In Übungsblatt 7 werden wir sehen, dass die Basis $\{v_1, v_2\}$ genau dann eine positive Basis von $\mathbb{R}^2$ ist, wenn $\varphi \in (0, \pi)$.

¹¹⁶Als *Basis* von V betrachten wir eine “angeordnete Menge von Vektoren”, d.h. wir unterscheiden zwischen $\{v_1, \dots, v_k\}$ und einer Permutation der Vektoren. Beispielsweise sind $\{e_1, e_2\}$ und $\{e_2, e_1\}$ *nicht* äquivalente Basen von $\mathbb{R}^2$.

¹¹⁷D.h. wir betrachten die $k \times k$ -Matrix $A = (a_{ij})$, welche durch $v_i := \sum_{j=1}^k a_{ji} w_j$ definiert ist.

¹¹⁸Etwas genauer gesagt, es sei $\varphi \in (-\pi, \pi)$ der Winkel, so dass es ein $\lambda > 0$ mit

$$\begin{pmatrix} \cos \varphi & -\sin \varphi \\ \sin \varphi & \cos \varphi \end{pmatrix} v_1 = \lambda v_2$$

gibt.

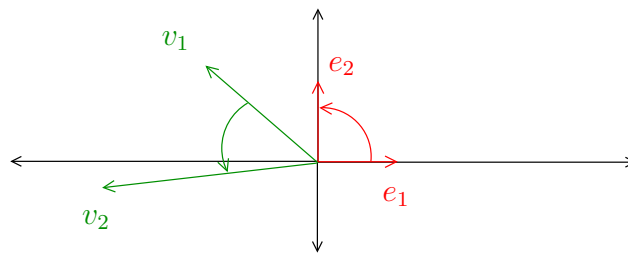


ABBILDUNG 56. Positive Basen für $\mathbb{R}^2$.

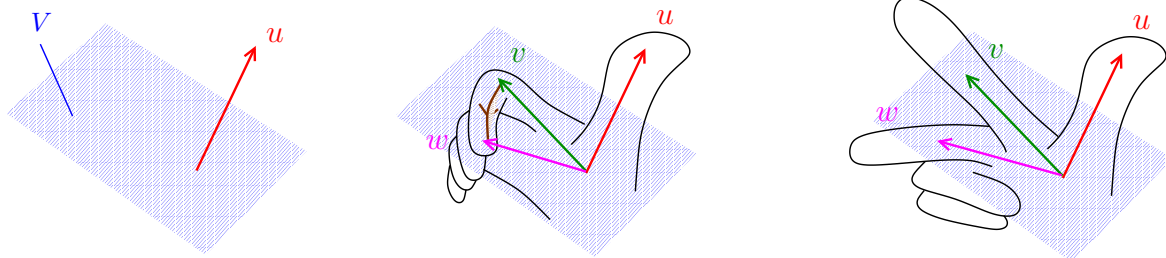
- (3) Es sei V ein 2-dimensionaler Untervektorraum von $\mathbb{R}^3$ und es sei $u \neq 0 \in \mathbb{R}^3$ ein Normalenvektor zu V . Dann legt u eine Orientierung auf V wie folgt fest:¹¹⁹

eine Basis $\{v, w\}$ ist positiv für $V \iff \det(u \ v \ w) > 0$.

Die geometrische Bedeutung der Definition ist wie folgt: eine Basis $\{v, w\}$ ist positiv, wenn man den Vektor w aus v durch Drehung um die u -Achse um einen Winkel $\varphi \in (0, \pi)$ erhält. Wir können diese Beobachtung noch anschaulicher formulieren: $\{v, w\}$ ist eine positive Basis, wenn gilt:

- (a) der Daumen zeigt in die u -Richtung,
- (b) der Zeigefinger zeigt in die v -Richtung,
- (c) der Mittelfinger zeigt in die w -Richtung.

Diese Aussage wird in Abbildung 57 illustriert.



den zweiten Basisvektor w erhalten wir durch Drehung „entlang der Finger“

der Daumen zeigt in die u -Richtung
 der Zeigefinger zeigt in die v -Richtung
 der Mittelfinger zeigt in die w -Richtung

ABBILDUNG 57.

- (4) Es sei V ein eindimensionaler Vektorraum. Dann ist jeder von 0 verschiedene Vektor eine Basis, und zwei Vektoren v und w geben die gleiche Orientierung, genau dann, wenn $v = rw$ für ein $r \in \mathbb{R}_{>0}$.

¹¹⁹Wenn man genau hinschaut, dann ist hier folgende Aussage versteckt: wenn $\{v, w\}$ und $\{v', w'\}$ positiv sind, d.h. wenn $\det(u \ v \ w) > 0$ und $\det(u \ v' \ w') > 0$, dann sind $\{v, w\}$ und $\{v', w'\}$ äquivalente Basen von V . Diese Aussage wird in Übungsblatt 7 bewiesen.

Definition. Es seien V und W orientierte Vektorräume.¹²⁰ Wir sagen ein Isomorphismus $\Phi: V \rightarrow W$ ist *orientierungserhaltend*, wenn das Bild einer positiven Basis von V eine positive Basis von W ist. Andernfalls sagen wir, dass Φ *orientierungsumkehrend* ist.¹²¹

Das folgende Lemma folgt problemlos aus den Definitionen und elementaren Eigenschaften der Determinante.

Lemma 9.14.

- (1) *Die Verknüpfung von zwei orientierungserhaltenden Isomorphismen ist wiederum orientierungserhaltend.*
- (2) *Das Inverse eines orientierungserhaltenden Isomorphismus ist wiederum orientierungserhaltend.*

Definition. Es sei M eine k -dimensionale Untermannigfaltigkeit.

- (1) Eine *Prä-Orientierung* für M ist eine Abbildung, welche jedem Tangentialraum $T_p M$ eine Orientierung zuordnet.
- (2) Wir sagen eine *Prä-Orientierung ist stetig bezüglich Karte* $\Phi: U \rightarrow V$, wenn auf jeder Komponente K von $U \cap M$ entweder gilt^{122,123}

$$\Phi_*: T_x M \rightarrow T_{\Phi(x)} E_k = \mathbb{R}^k \text{ ist orientierungserhaltend für alle } x \in K$$

oder

$$\Phi_*: T_x M \rightarrow T_{\Phi(x)} E_k = \mathbb{R}^k \text{ ist orientierungsumkehrend für alle } x \in K.$$

- (3) Wir sagen eine *Prä-Orientierung* ist eine *Orientierung*, wenn es einen Atlas für M gibt, so dass die Prä-Orientierung bezüglich allen Karten des Atlanten stetig ist.
- (4) Eine Untermannigfaltigkeit zusammen mit einer Orientierung heißt *orientierte Untermannigfaltigkeit*. Eine Untermannigfaltigkeit, für die man eine Orientierung finden kann, heißt *orientierbare Untermannigfaltigkeit*.

Bevor wir uns den Beispielen zuwenden erwähnen wir noch folgendes Lemma, welches die gleiche Rolle spielt wie Satz 2.7. Der Beweis dazu ist elementar, und verbleibt eine freiwillige Übungsaufgabe.

¹²⁰Genauer gesagt, müsste man schreiben, es seien $(V, \mathcal{O})$ und $(W, \mathcal{P})$ orientierte Vektorräume, aber wir unterdrücken “ $\mathcal{O}$ ” und “ $\mathcal{P}$ ” in der Notation, genauso wie wir sagen “Sei K ein Körper”, anstatt zu schreiben “Sei $(K, +, \cdot)$ ein Körper”.

¹²¹Es seien B und C positive Basen von V . Wenn $\Phi(B)$ eine positive Basis von W ist, warum folgt dann auch, dass $\Phi(C)$ eine positive Basis von W ist?

¹²²Hierbei betrachten wir $\mathbb{R}^k$ mit der kanonischen Orientierung, welche durch $e_1, \dots, e_k$ gegeben ist.

¹²³Anders ausgedrückt, im Hinblick auf Lemma 9.14 (2) gilt für $\Psi = \Phi^{-1}$ entweder

$$\{\Psi_*(e_1), \dots, \Psi_*(e_k)\} \text{ ist eine positive Basis von } T_{\Psi(x)} M \text{ für alle } x \in \Phi(K)$$

oder

$$\{\Psi_*(e_1), \dots, \Psi_*(e_k)\} \text{ ist eine negative Basis von } T_{\Psi(x)} M \text{ für alle } x \in \Phi(K).$$

Lemma 9.15. *Es sei M eine Untermannigfaltigkeit. Jede Orientierung ist stetig bezüglich allen Karten von M .*

Beispiele.

- (1) In Kapitel 6.2 hatten wir den Begriff der Orientierung einer Hyperfläche mithilfe von Einheitsnormalenfeldern eingeführt. Man kann leicht zeigen, dass die Definition in Kapitel 6.2 äquivalent ist zur obigen Definition. In der Tat, sei $M \subset \mathbb{R}^n$ eine Hyperfläche, dann gilt:

- (A) Es sei ν ein Einheitsnormalenfeld, d.h. ein normiertes Normalenfeld auf M . Es sei $P \in M$. Wir sagen eine Basis $\{v_1, \dots, v_{n-1}\}$ von $T_P M$ ist positiv, wenn¹²⁴

$$\det(\nu(P) v_1 \dots v_{n-1}) > 0.$$

In Übungsblatt 7 werden wir mithilfe von Satz 2.5 zeigen, dass dies eine Orientierung auf M definiert. Für M eine Hyperfläche in $\mathbb{R}^3$ erhalten wir gerade die positiven Basen der Tangentialräume, welche wir schon auf Seite 121 diskutiert hatten.

- (B) Umgekehrt, nehmen wir an, wir hätten eine Orientierung für jeden Tangentialraum $T_P M$, welche die Stetigkeitsbedingung erfüllt. Wir wollen nun ein stetiges normiertes Normalenfeld auf M konstruieren. Sei also $P \in M$. Dann haben wir schon gesehen, dass es genau zwei Vektoren w_1, w_2 gibt, mit $\|w_1\| = \|w_2\| = 1$ und $w_1, w_2 \perp T_P M$. Wir ordnen jetzt P den Vektor zu, welcher die Bedingung¹²⁵

$$\det(w_i v_1 \dots v_{n-1}) > 0$$

erfüllt für eine (und dann jede) positive Basis $v_1, \dots, v_{n-1}$ von $T_P M$. Man kann nun, mit einem ähnlich Argument wie in (A) zeigen, dass dieses normierte Normalenfeld stetig ist.

Wir werden im Folgenden zwischen diesen beiden Gesichtspunkten kommentarlos hin- und herwechseln.

- (2) Es sei $f: \mathbb{R}^n \rightarrow \mathbb{R}$ eine glatte Funktion und $z \in \mathbb{R}$ ein regulärer Wert. Wir betrachten die $(n-1)$ -dimensionale Untermannigfaltigkeit $M := f^{-1}(z)$. Es sei nun $P \in M$. Wir sagen eine Basis $\{v_1, \dots, v_{n-1}\}$ von $T_P M$ ist positiv, wenn¹²⁶

$$\det(\text{Grad } f(P) v_1 \dots v_{n-1}) > 0.$$

Dies definiert eine Orientierung für jeden Tangentialraum $T_P M$. Es folgt nun aus Satz 6.2 und dem vorherigen Beispiel, dass diese Orientierungen tatsächlich eine Orientierung für M definierten.

¹²⁴Auch hier gilt die gleiche Anmerkung wie in Fußnote 119.

¹²⁵Diese Bedingung erfüllt entweder w_1 oder w_2 .

¹²⁶Zur Erinnerung, nach Satz 6.2 gilt für $P \in M$, dass $T_P M = \{v \in \mathbb{R}^n \mid v \perp \text{Grad } f(P)\}$. Es sei nun $\{v_1, \dots, v_{n-1}\}$ eine Basis von $T_P M$, es folgt dass $\{\text{Grad } f(P), v_1, \dots, v_{n-1}\}$ eine Basis von $\mathbb{R}^n$ bildet, insbesondere gilt

$$\det(\text{Grad } f(P) v_1 \dots v_{n-1}) \neq 0.$$

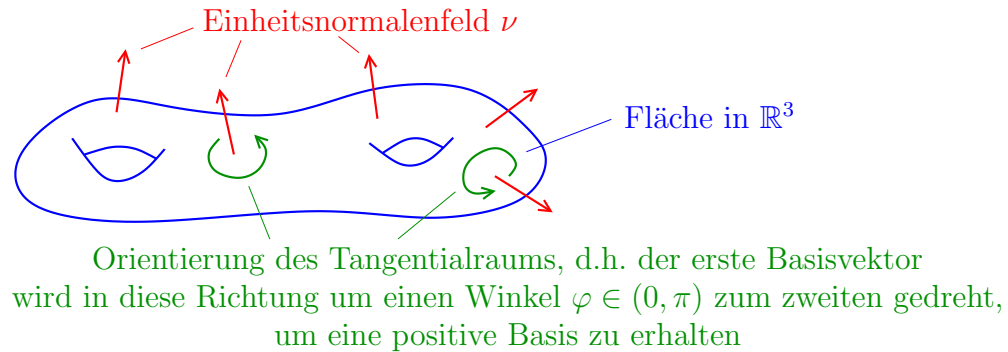


ABBILDUNG 58.

- (3) Es sei M eine zusammenhängende eindimensionale Untermannigfaltigkeit mit Rand. In Kapitel 8.4 hatten wir schon den Begriff einer Orientierung auf M eingeführt. Die beiden Gesichtspunkte sind äquivalent, denn auf einer eindimensionalen Untermannigfaltigkeit ist ein Tangentialvektor $v \neq 0$ auf $T_P M$ das gleiche wie eine Basis für $T_P M$.
- (4) Es sei M eine n -dimensionale Untermannigfaltigkeit von $\mathbb{R}^n$, z.B. eine offene Teilmenge von $\mathbb{R}^n$. Dann gilt für jeden Punkt $P \in M$, dass $T_P M = \mathbb{R}^n$. Die Standardbasis des $\mathbb{R}^n$ definiert nun eine Orientierung von M . Wenn wir nichts anderes sagen, dann betrachten wir eine n -dimensionale Untermannigfaltigkeit von $\mathbb{R}^n$ immer mit dieser Orientierung.

Definition. Es sei nun M eine k -dimensionale orientierte Untermannigfaltigkeit.¹²⁷ Wir bezeichnen dann mit $-M$ die Untermannigfaltigkeit mit der entgegengesetzten Orientierung, d.h. eine Basis $\{v_1, \dots, v_k\}$ eines Tangentialraumes ist positiv für $-M$, genau dann, wenn sie negativ für M ist.

Bemerkung.

- (1) Wenn M eine zusammenhängende orientierte Untermannigfaltigkeit ist, dann kann man ähnlich zu Satz 6.4 zeigen, dass M genau zwei verschiedene Orientierungen besitzt, nämlich M und $-M$.
- (2) Wenn M eine orientierte Untermannigfaltigkeit mit k Komponenten $M_1, \dots, M_k$ ist, dann besitzt M genau 2^k verschiedene Orientierungen, diese sind gegeben durch $\epsilon_1 M_1 \cup \dots \cup \epsilon_k M_k$ mit $\epsilon_i \in \{-1, 1\}$ für $i = 1, \dots, k$.

Definition. Es sei $\Phi: M \rightarrow N$ ein Diffeomorphismus zwischen orientierten zusammenhängenden Untermannigfaltigkeiten der Dimension $k \geq 1$. Wir sagen f ist *orientierungserhaltend*, wenn für alle $P \in M$ die Abbildung $\Phi_*: T_P M \rightarrow T_{\Phi(P)} N$ orientierungserhaltend ist. Ansonsten sagen wir, dass Φ orientierungsumkehrend ist.

¹²⁷Genauer gesagt, müsste man schreiben, es sei (M, O) eine orientierte Untermannigfaltigkeit, aber wir unterdrücken "O" in der Notation.

Bemerkung.

- (1) Man kann relativ leicht zeigen, dass ein Diffeomorphismus $\Phi: M \rightarrow N$ zwischen orientierten zusammenhängenden Untermannigfaltigkeiten genau dann orientierungserhaltend ist, wenn es ein $P \in M$ gibt, so dass $\Phi_*: T_P M \rightarrow T_{\Phi(P)} N$ orientierungserhaltend ist.
- (2) Wenn $\Phi: M \rightarrow N$ ein Diffeomorphismus zwischen zwei zusammenhängenden offenen Teilmengen von $\mathbb{R}^k$ ist, dann ist Φ orientierungserhaltend, genau dann, wenn es ein $P \in M$ gibt, so dass $\det(D_P \Phi) > 0$.

Satz 9.16. *Es sei M eine zusammenhängende orientierte Untermannigfaltigkeit. Dann ist die Abbildung*

$$\begin{aligned} \sigma: \{\text{Diffeomorphismen von } M\} &\rightarrow \{\pm 1\} \\ (\Phi: M \rightarrow M) &\mapsto \sigma(\Phi) := \begin{cases} 1, & \text{wenn } \Phi \text{ orientierungserhaltend,} \\ -1, & \text{wenn } \Phi \text{ orientierungsumkehrend} \end{cases} \end{aligned}$$

ein Gruppenhomomorphismus.^{128 129}

Beweis. Für einen Isomorphismus Φ zwischen orientierten Vektorräumen V und W definieren wir

$$\sigma(\Phi) = \begin{cases} +1, & \text{wenn } \Phi \text{ orientierungserhaltend,} \\ -1, & \text{wenn } \Phi \text{ orientierungsumkehrend.} \end{cases}$$

Ganz analog zu Lemma 9.14 kann man nun zeigen, dass für Isomorphismen $\Phi: U \rightarrow V$ und $\Psi: V \rightarrow W$ zwischen orientierten Vektorräumen gilt, dass $\sigma(\Psi \circ \Phi) = \sigma(\Psi) \cdot \sigma(\Phi)$.

Es sei nun M eine zusammenhängende orientierte Untermannigfaltigkeit. Es folgt aus der obigen Bemerkung, dass für alle $P \in M$ gilt, dass

$$\sigma(M) = \sigma(\Phi_*: T_P M \rightarrow T_{\Phi(P)} M).$$

Es seien nun $\Phi: M \rightarrow M$ und $\Psi: M \rightarrow M$ Diffeomorphismen. Es sei $P \in M$ beliebig. Dann gilt

$$\begin{aligned} \sigma(\Psi \circ \Phi) &= \sigma((\Psi \circ \Phi)_*: T_P M \rightarrow T_{(\Psi \circ \Phi)(P)} M) \\ &= \sigma((\Psi_*: T_{\Phi(P)} M \rightarrow T_{(\Psi \circ \Phi)(P)} M) \circ (\Phi_*: T_P M \rightarrow T_{\Phi(P)} M)) \\ &\quad \uparrow \\ &\text{Kettenregel} \\ &= \sigma((\Psi_*: T_{\Phi(P)} M \rightarrow T_{(\Psi \circ \Phi)(P)} M)) \cdot \sigma(\Phi_*: T_P M \rightarrow T_{\Phi(P)} M) \\ &\quad \uparrow \\ &\text{siehe obige Aussage für } \sigma \text{ von Isomorphismen von orientierten Vektorräumen} \\ &= \sigma(\Psi) \cdot \sigma(\Phi). \\ &\quad \uparrow \\ &\text{denn } \sigma(\Psi) \text{ und } \sigma(\Phi) \text{ können an jedem Punkt bestimmt werden} \end{aligned}$$

□

¹²⁸Auf gut Deutsch heißt das, dass die Verknüpfung von n Diffeomorphismen genau dann orientierungserhaltend ist, wenn eine gerade Zahl dieser Diffeomorphismen orientierungsumkehrend ist.

¹²⁹Ist die Abbildung immer ein Epimorphismus?

9.9. Integration von Differentialformen II. Es sei M eine orientierte k -dimensionale Untermannigfaltigkeit. Wir nennen eine Karte $\Phi: U \rightarrow V$ *orientierungserhaltend*, wenn die lineare Abbildung

$$\Phi_* = D_x \Phi: T_x M \rightarrow T_x(E_k) = \mathbb{R}^k$$

für alle $x \in M \cap U$ orientierungserhaltend ist. Man kann sich leicht davon überzeugen, dass eine orientierte Untermannigfaltigkeit einen Atlas besitzt, welcher nur aus orientierungserhaltenden Karten besteht.¹³⁰

Es sei nun ω eine stetige k -Form auf einer orientierten kompakten k -dimensionalen Untermannigfaltigkeit. Wir nehmen zuerst an, dass es eine orientierungserhaltende Karte $\Phi: U \rightarrow V$ von Typ (i) mit $\text{Träger}(\omega) \subset U$ gibt. Wir bezeichnen mit $\Psi: V \rightarrow U$ die Umkehrabbildung von Φ . Wir definieren nun

$$\int_M \omega := \int_{V \cap E_k} \omega_{\Psi(x)}(\Psi_*(e_1), \dots, \Psi_*(e_k)) dx.$$

Das folgende Lemma besagt nun, dass diese Definition nicht von der Wahl der orientierungserhaltenden Karte abhängt:

Lemma 9.17. *Es sei M eine orientierte kompakte k -dimensionale Untermannigfaltigkeit und es sei $\omega \in \mathcal{C}_k(M)$. Es seien zudem $\Phi: U \rightarrow V$ und $\tilde{\Phi}: \tilde{U} \rightarrow \tilde{V}$ zwei orientierungserhaltende Karten mit $\text{Träger}(\omega) \subset U \cap \tilde{U}$. Wir bezeichnen mit Ψ und $\tilde{\Psi}$ die Umkehrabbildungen von Φ und $\tilde{\Phi}$. Dann gilt*

$$\int_{V \cap E_k} \omega_{\Psi(x)}(\Psi_*(e_1), \dots, \Psi_*(e_k)) dx = \int_{\tilde{V} \cap E_k} \omega_{\tilde{\Psi}(x)}(\tilde{\Psi}_*(e_1), \dots, \tilde{\Psi}_*(e_k)) dx.$$

Beweis. Indem wir die Karten auf $U \cap \tilde{U}$ einschränken können wir o.B.d.A. annehmen, dass $U = \tilde{U}$. Wir betrachten den Diffeomorphismus

$$\Theta := \Phi \circ \tilde{\Phi}^{-1}: \tilde{V} \cap E_k \rightarrow V \cap E_k.$$

Es folgt nun aus Satz 9.13, dass es genügt zu zeigen, dass für alle $Q \in \tilde{V} \cap E_k$ gilt $\det(D_Q \Theta) > 0$, d.h. dass $D_Q \Theta: \mathbb{R}^k \rightarrow \mathbb{R}^k$ orientierungserhaltend ist. Es sei also $Q \in \tilde{V} \cap E_k$. Wir wollen zeigen, dass $D_Q \Theta: \mathbb{R}^k \rightarrow \mathbb{R}^k$ orientierungserhaltend ist. Wir setzen $P = \tilde{\Phi}^{-1}(Q)$. Dann ist

$$(D(\Phi \circ \tilde{\Phi}^{-1})_Q = D_Q \Theta: \mathbb{R}^k \rightarrow \mathbb{R}^k) = \underbrace{(D_P \Phi: T_P M \rightarrow \mathbb{R}^k)}_{\substack{\uparrow \\ \text{Kettenregel}}} \circ \underbrace{(D_Q \tilde{\Phi}^{-1}: \mathbb{R}^k \rightarrow T_P M)}_{\substack{\uparrow \\ \text{orientierungserhaltend}}}.$$

¹³⁰In der Tat, es sei $\{\Phi_i: U_i \rightarrow V_i\}_{i \in I}$ ein Atlas. Durch Einschränken auf Zusammenhangskomponenten von den V_i können wir o.B.d.A. annehmen, dass alle V_i zusammenhängend sind. Wenn $\Phi: U \rightarrow V$ eine Karte ist, so dass V zusammenhängend ist, aber so dass Φ nicht orientierungserhaltend ist, dann ist $\Theta \circ \Phi$ mit $\Theta(x_1, \dots, x_n) = (-x_1, x_2, \dots, x_n)$ eine orientierungserhaltende Karte mit dem gleichen Definitionsbereich.

Es folgt nun aus Lemma 9.14, dass die Abbildung $D_Q\Theta$ selber orientierungserhaltend ist. $\square$

9.10. Integration von Differentialformen III. In diesem Kapitel wollen wir das Integral von einer beliebigen stetigen k -Form auf einer orientierten kompakten k -dimensionalen Untermannigfaltigkeit einführen. Wir wollen die Definition, ganz analog zu Kapitel 3.5, auf den im vorherigen Kapitel diskutierten Fall zurückführen.

Um das nächste Lemma formulieren zu können benötigen wir noch folgende Definition.

Definition. Wir sagen Q ist ein *offener Quader in H_k* , wenn es einen offenen Quader in $E_k = \mathbb{R}^k$ gibt, so dass der Durchschnitt mit H_k gerade Q ist. Wir schreiben

$$\partial_0 Q := Q \cap E_{k-1},$$

und¹³¹

$$\partial_1 Q = \partial \overline{Q} \setminus \partial_0 Q.$$

Wir sagen eine stetige Abbildung $f: Q \rightarrow \mathbb{R}$ *verschwindet auf $\partial_1 Q$* , wenn sich f stetig zu einer Abbildung auf $\overline{Q} = Q \cup \partial_1 Q$ fortsetzen läßt, welche auf $\partial_1 Q$ verschwindet.

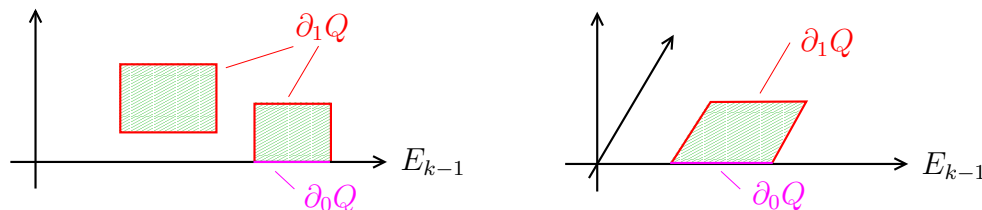


ABBILDUNG 59. Beispiele von offenen Quadern in Halbräumen.

Wir können nun folgenden Satz formulieren.

Satz 9.18. *Es sei M eine kompakte k -dimensionale Untermannigfaltigkeit. Dann gibt es einen endlichen Atlas $\{\Phi_i: U_i \rightarrow V_i\}_{i=1,\dots,r}$, sowie glatte Funktionen $z_1, \dots, z_r: M \rightarrow [0, 1]$, so dass folgende Eigenschaften erfüllt sind:*

- (1) *für alle $i \in \{1, \dots, r\}$ gilt $\text{Träger}(z_i) \subset U_i \cap M$,*
- (2) *es ist $z_1 + \dots + z_r = 1$,*

und¹³²

- (3) *für alle $i \in \{1, \dots, r\}$ gilt zudem*
 - (a) $\Phi_i(U_i \cap M) \subset H_k$,
 - (b) *die Menge $V_i \cap H_k = \Phi_i(U_i \cap M)$ ist ein offener Quader in H_k ,*
 - (c) *die Funktion $z_i \circ \Phi_i^{-1}: V_i \cap H_k \rightarrow \mathbb{R}$ verschwindet auf $\partial_1(V_i \cap H_k)$.*

¹³¹Hierbei bezeichnen wir mit $\partial \overline{Q}$ den topologischen Rand von $\overline{Q}$ betrachtet als Teilmenge von E_k .

¹³²Die dritte Eigenschaft werden wir erst etwas später benötigen.

Wenn M zudem orientiert ist, dann kann man die Karten auch orientierungserhaltend wählen.

Beweis. Wir stellen zuerst folgende Behauptung auf.

Behauptung. Es gibt einen Atlas $\{\Phi_i: U_i \rightarrow V_i\}_{i=1,\dots,r}$, so dass für alle $i \in \{1, \dots, r\}$ gilt:

- (a) $\Phi_i(U_i \cap M) \subset H_k$ und
- (b) die Menge $V_i \cap H_k$ ist ein offener Quader in H_k .

Wenn M eine orientierte Untermannigfaltigkeit ist, dann können wir zudem annehmen, dass alle Karten orientierungserhaltend sind.

Wir zeigen zuerst, dass es für jedes $X \in M$ eine Karte $\Phi_X: U_X \rightarrow V_X$ um X mit den gewünschten Eigenschaften gibt.

- (1) Wenn X ein innerer Punkt von M ist, dann wählen wir zunächst eine beliebige Karte $\Phi_X: U_X \rightarrow V_X$ um X . Nach einer eventuellen Einschränkung und Verschiebung von V_X können wir annehmen, dass $V_X \subset \{(x_1, \dots, x_n) \in \mathbb{R}^n \mid x_k > 0\}$, und dass V_X ein offener Quader ist. Dann ist auch $V_X \cap H_k$ ein offener Quader in H_k .
- (2) Wenn X ein Randpunkt von M ist, dann wählen wir auch wieder zuerst eine beliebige Karte $\Phi_X: U_X \rightarrow V_X$ um X . Nach einer eventuellen Einschränkung können wir annehmen, dass V_X ein offener Quader ist. Dann ist wiederum $V_X \cap H_k$ ein offener Quader in H_k .

Nachdem M kompakt ist, wird M schon durch endlich viele dieser Karten abgedeckt. Wenn M eine orientierte Untermannigfaltigkeit ist, dann können wir mit dem Argument von Fußnote 130 arrangieren, dass alle Karten zudem orientierungserhaltend sind. Wir haben damit also die Behauptung bewiesen.

Die Funktionen $z_1, \dots, z_r$ kann man nun mit einer leichten Abwandlung des Beweises von Lemma 6.11 konstruieren. Wir überlassen den Beweis als freiwillige Übungsaufgabe. $\square$

Bemerkung. Es sei M eine Untermannigfaltigkeit. Eine *glatte Zerlegung der Eins* ist eine Familie von glatten Funktionen $\{z_i: M \rightarrow [0, 1]\}_{i \in I}$, so dass folgende Eigenschaften erfüllt sind:

- (1) für alle $i \in I$ gibt es eine Karte $\Phi_i: U_i \rightarrow V_i$, so dass $\text{Träger}(z_i) \subset U_i \cap M$,
- (2) für alle $x \in M$ gibt es nur endlich viele Indizes $i \in I$ mit $z_i(x) \neq 0$,
- (3) für alle $x \in M$ gilt $\sum_{i \in I} z_i(x) = 1$.¹³³

Wir haben also gerade in Satz 9.18 gezeigt, dass jede kompakte Untermannigfaltigkeit eine glatte Zerlegung der Eins besitzt. Es gilt jedoch die allgemeinere Aussage, dass jede Untermannigfaltigkeit eine glatte Zerlegung der Eins besitzt. Diese Aussage wird auf Seite 362 von

Königsberger: Analysis 2 (Springer-Verlag)

¹³³Die Summe $\sum_{i \in I} z_i(x)$ macht Sinn, selbst wenn die Indexmenge I unendlich ist, denn nach (2) gibt es nur endlich viele Indizes $i \in I$ mit $z_i(x) \neq 0$.

und in in Kapitel 9.4 von

Jänich: Vektoranalysis (Springer–Verlag)

bewiesen. Man kann solche glatte Zerlegungen der Eins dazu verwenden, das Integral einer k -Form auf einer beliebigen orientierten, nicht notwendigerweise kompakten k -dimensionalen Untermannigfaltigkeit einzuführen. Wir werden diese allgemeinere Aussage allerdings nicht verwenden.

Es sei nun M eine orientierte kompakte k -dimensionale Untermannigfaltigkeit und es sei ω eine stetige k -Form auf M . Wir wählen einen Atlas $\{\Phi_i: U_i \rightarrow V_i\}_{i=1,\dots,r}$, sowie glatte Funktionen $z_1, \dots, z_r: M \rightarrow [0, 1]$ wie in Satz 9.18. Wir definieren nun

$$\int_M \omega := \sum_{i=1}^r \int_M z_i \cdot \omega,$$

wobei die Summanden auf der rechten Seite nach Kapitel 9.9 definiert sind. Das folgende Lemma besagt nun, dass diese Definition nicht von den getroffenen Wahlen abhängt. Dieses Lemma spielt die gleiche Rolle wie Lemma 3.18.

Lemma 9.19. *Das Integral einer stetigen k -Form auf einer orientierten kompakten k -dimensionalen Untermannigfaltigkeit ist wohl-definiert.*

Beweis. Es sei ω eine stetige k -Form auf einer orientierten kompakten k -dimensionalen Untermannigfaltigkeit M . Es seien $y_1, \dots, y_q$ und $z_1, \dots, z_r$ zwei Mengen von glatten Funktionen wie in Satz 9.18. Dann gilt

$$\begin{array}{ccccccc} \sum_{i=1}^p \int_M y_i \cdot \omega & = & \sum_{i=1}^p \int_M \left(\sum_{j=1}^r z_j \right) y_i \cdot \omega & = & \sum_{i=1}^p \int_M \sum_{j=1}^r z_j y_i \cdot \omega & = & \sum_{i=1}^p \sum_{j=1}^r \int_M z_j y_i \cdot \omega & = & \sum_{j=1}^r \int_M z_j \cdot \omega. \\ & \uparrow & & & \uparrow & & \uparrow & & \uparrow \\ & \text{denn } \sum_{j=1}^r z_j = 1 & & & \text{Linearität} & & \text{das gleiche Argument} & & \\ & & & & \text{des Integrals} & & \text{rückwärts} & & \square \end{array}$$

Es folgt leicht aus den Definitionen, dass für das Integral von k -Formen die “üblichen” Aussagen gelten. Beispielsweise, gilt für alle $\omega, \omega' \in \mathcal{C}_k(M)$ und $l \in \mathbb{R}$, dass

$$\int_M (\eta + \omega) = \int_M \eta + \int_M \omega \quad \text{und} \quad \int_M l\omega = l \cdot \int_M \omega.$$

Etwas interessanter ist folgender Satz:

Satz 9.20. *Es sei $M \subset \mathbb{R}^n$ eine orientierte kompakte k -dimensionale Untermannigfaltigkeit und es sei $\omega \in \mathcal{C}_k(M)$. Dann gilt*

$$\int_{-M} \omega = - \int_M \omega.$$

Beweis. Es folgt direkt aus der Definition vom Integral von k -Formen, dass es genügt den Fall zu betrachten, dass es eine Orientierungserhaltende Karte $\Phi: U \rightarrow V$ für M mit $\text{Träger}(\omega) \subset U$ gibt. Wir betrachten die durch

$$\Theta(x_1, x_2, \dots, x_n) := (-x_1, x_2, \dots, x_n)$$

definierte Abbildung. Für alle $x \in V$ gilt offensichtlich, dass $\det D_x \Theta = -1$. Die Karte $\tilde{\Phi} := \Theta \circ \Phi: U \rightarrow \tilde{V} := \Theta(V)$ ist dann eine Orientierungsumkehrende Karte für M , oder anders ausgedrückt, eine Orientierungserhaltende Karte für $-M$.

Wir bezeichnen nun mit $\Psi: V \rightarrow U$ und $\tilde{\Psi}: \tilde{V} \rightarrow U$ die Umkehrabbildungen von Φ und $\tilde{\Phi}$. Wir erhalten nun, dass

$$\begin{array}{ccccc} \text{denn } \tilde{\Psi} \text{ ist Orientierungs-} & & \text{denn } \Psi \text{ ist Orientierungs-} & & \\ \text{erhaltend für } -M & & \text{erhaltend für } M & & \\ \int_{-M} \omega & \stackrel{\downarrow}{=} & \int_{\tilde{V}} \tilde{\Psi}^* \omega & \stackrel{\uparrow}{=} & - \int_V \Psi^* \omega & \stackrel{\downarrow}{=} & - \int_M \omega. \\ & & \text{siehe Satz 9.13} & & & & \square \end{array}$$

Folgender Satz gibt uns die Transformationsformel für das Integral von k -Formen auf orientierten k -dimensionalen Untermannigfaltigkeiten. Der Beweis des Satzes ist, wenn man die Definitionen verdaut hat, völlig elementar und wird in Übungsblatt 7 ausgeführt.

Satz 9.21. *Es sei $f: M \rightarrow N$ ein Diffeomorphismus zwischen zwei orientierten kompakten k -dimensionalen Untermannigfaltigkeiten. Zudem sei ω eine stetige k -Form auf N . Dann gilt*

$$\int_M f^* \omega = \int_N \omega, \quad \text{wenn } f \text{ Orientierungserhaltend}$$

und

$$\int_M f^* \omega = - \int_N \omega, \quad \text{wenn } f \text{ Orientierungsumkehrend.}$$

Wir betrachten im Folgenden noch ein paar Spezialfälle vom Integral einer k -Form, welche wir später verwenden werden.

Lemma 9.22. *Es sei $M \subset \mathbb{R}^n$ eine n -dimensionale kompakte Untermannigfaltigkeit, d.h. eine Untermannigfaltigkeit der Kodimension 0. Zudem sei $f: M \rightarrow \mathbb{R}$ eine stetige Funktion. Dann gilt*

$$\underbrace{\int_M f \cdot dx_1 \wedge \cdots \wedge dx_n}_{\text{Integral von } n\text{-Form}} = \underbrace{\int_M f dx_1 \dots dx_n}_{\text{Integral von Funktion}}$$

Beweis. Nachdem ∂M eine Nullmenge ist, genügt es wenn wir die Gleichheit der Integrale über der offenen Teilmenge $\overset{\circ}{M} = M \setminus \partial M$ von $\mathbb{R}^n$ beweisen. Es gilt in der Tat, dass

$$\begin{array}{ccccccc} \int_{\overset{\circ}{M}} f \cdot dx_1 \wedge \cdots \wedge dx_n & = & \int_{\overset{\circ}{M}} f \cdot \underbrace{(dx_1 \wedge \cdots \wedge dx_n)(e_1, \dots, e_n)}_{=1, \text{ siehe Seite 109}} & = & \int_{\overset{\circ}{M}} f & = & \int_{\overset{\circ}{M}} f dx_1 \dots dx_n. \\ \uparrow & & \uparrow & & \uparrow & & \uparrow \\ \text{Definition angewandt auf die} & & & & & & dx_1 \dots dx_n \text{ hat keine Bedeutung.} \\ \text{Orientierungserhaltende Karte } \Phi = \text{id} & & & & & & \square \end{array}$$

Das folgende Lemma besagt, dass wir den Fluß von einem Vektorfeld durch eine orientierte Hyperfläche als Integral über eine Differentialform interpretieren können.

Lemma 9.23. *Es sei (N, ν) eine orientierte kompakte Hyperfläche in $\mathbb{R}^n$ und es sei zudem $F: N \rightarrow \mathbb{R}^n$ ein stetiges Vektorfeld. Wir bezeichnen mit $\delta(F)$ die zu F zugehörige $(n-1)$ -Form auf N , welche an jedem Punkt P gegeben ist durch*

$$\begin{aligned} \delta(F)_P: (T_P M)^{n-1} &\rightarrow \mathbb{R} \\ (v_1, \dots, v_{n-1}) &\mapsto \det(F(P) v_1 \dots v_{n-1}). \end{aligned}$$

Dann gilt¹³⁴

$$\int_N \delta(F) = \int_N F \cdot \nu.$$

Beweis. Es folgt aus den Definitionen, dass es genügt den Fall zu betrachten, dass es eine Orientierungserhaltende Karte $\Phi: U \rightarrow V$ mit $\text{Träger}(F) \subset U$ gibt.

Es seien $v_1, \dots, v_{n-1} \in \mathbb{R}^n$. Wir bezeichnen dann mit $v_1 \times \dots \times v_{n-1} \in \mathbb{R}^n$ den Vektor in $\mathbb{R}^n$ der festgelegt ist durch¹³⁵

$$k\text{-te Koordinate von } v_1 \times \dots \times v_{n-1} = (-1)^{k+1} \cdot \det \begin{pmatrix} v_1 & \dots & v_{n-1}, \\ \text{wobei } k\text{-te Zeile} \\ \text{entfernt} \end{pmatrix}.$$

In Übungsblatt 6 hatten wir sehen, dass das Kreuzprodukt $v_1 \times \dots \times v_{n-1}$ auch eindeutig bestimmt ist durch folgende Eigenschaften festgelegt:

- (a) $v_1 \times \dots \times v_{n-1}$ steht senkrecht zu $v_1, \dots, v_{n-1}$,
- (b) $\|v_1 \times \dots \times v_{n-1}\| = \sqrt{\text{Gram}(v_1, \dots, v_{n-1})}$,
- (c) $\det \underbrace{(v_1 \times \dots \times v_{n-1} \quad v_1 \dots v_{n-1})}_{n \times n\text{-Matrix}} \geq 0$.

Es sei nun $x \in V \cap E_k$. Für $i = 1, \dots, n-1$ schreiben wir $v_i := (D_x \Psi)(e_i)$ und $F := F(\Psi(x))$ sowie $\nu = \nu(\Psi(x))$. Wir erhalten, dass

denn $\text{Träger}(F) \subset U$ $\begin{aligned} \int_N \delta(F) &\stackrel{\downarrow}{=} \int_{V \cap E_{n-1}} \delta(F)(v_1, \dots, v_{n-1}) \\ &= \int_{V \cap E_{n-1}} F \cdot (v_1 \times \dots \times v_{n-1}) \\ &\stackrel{\uparrow}{=} \end{aligned}$	Definition von $\delta(F)$ $\begin{aligned} &\stackrel{\downarrow}{=} \int_{V \cap E_{n-1}} \det(F v_1 \dots v_{n-1}) \\ &= \int_{V \cap E_{n-1}} F \cdot \sqrt{\text{Gram}(v_1, \dots, v_{n-1})} \cdot \nu \\ &\stackrel{\uparrow}{=} \end{aligned}$
--	---

folgt aus der Definition von $v_1 \times \dots \times v_{n-1}$ es ist $v_1 \times \dots \times v_{n-1} = \sqrt{\text{Gram}(v_1, \dots, v_{n-1})} \cdot \nu$, denn und der Laplace-Entwicklung der Determinante beide Vektoren erfüllen (a), (b) und (c)¹³⁶

$$= \int_{V \cap E_{n-1}} F \cdot \nu \sqrt{\text{Gram}(v_1, \dots, v_{n-1})} = \int_N F \cdot \nu. \quad \square$$

¹³⁴Auf der linken Seite betrachten wir also das Integral der Differentialform $\delta(F)$, hierbei betrachten wir N als orientierte Untermannigfaltigkeit mit der durch ν festgelegten Orientierung, siehe Kapitel 9.8. Auf der rechten Seite betrachten wir das Integral der reellwertigen Funktion $x \mapsto F(x) \cdot \nu(x)$.

¹³⁵Für zwei Vektoren in $\mathbb{R}^3$ erhalten wir das übliche Kreuzprodukt.

¹³⁶In der Tat, beide Vektoren haben offensichtlich die Eigenschaften (a) und (b). Zudem ist $v_1, \dots, v_{n-1}$ eine positive Basis von $T_{\Psi(x)}M$, per Definition der Orientierung (siehe Kapitel 9.8 für die Orientierung von

9.11. Null-dimensionale Untermannigfaltigkeiten. Wir haben bisher nur Untermannigfaltigkeiten der Dimension größer Null behandelt. Es sei nun M eine kompakte null-dimensionale Untermannigfaltigkeit von $\mathbb{R}^n$, d.h. M besteht aus endlich vielen Punkten $P_1, \dots, P_k$. Wir definieren eine *Orientierung* für M als eine Funktion $M \rightarrow \{-1, 1\}$. Mit anderen Worten, eine Orientierung besteht aus der Wahl eines Vorzeichens ϵ_i für jeden Punkt P_i .

Es sei nun ω eine stetige 0-Form auf M , d.h. ω ist eine Funktion $M \rightarrow \mathbb{R}$. Wir definieren dann

$$\int_M \omega := \sum_{i=1}^k \epsilon_i f(P_i).$$

Hyperflächen) bedeutet das, dass

$$\det(\nu v_1, \dots, v_{n-1}) > 0.$$

Also haben beide Vektoren auch die Eigenschaft (c).

10. VORBEREITUNGEN FÜR DEN SATZ VON STOKES

Wir sind immer noch auf dem Weg zur Formulierung vom Satz von Stokes. Dieser ist die in Kapitel 9.1 versprochene Verallgemeinerung des Gaußschen Integralsatzes und von Satz 9.2. In Kapitel 9.1 hatten wir schon angedeutet, dass diese Verallgemeinerung in etwa von folgender Form sein soll.

Es sei M eine k -dimensionale kompakte Untermannigfaltigkeit und es sei ω ein “mathematisches Objekt” auf M , welches auf ∂M integriert werden kann. Dann soll gelten

$$\int_M \text{“Differential” oder “Ableitung” von } \omega = \int_{\partial M} \omega.$$

Die Diskussion des vorherigen Kapitels legt nahe, dass dieses mathematische Objekt und dessen “Differential” wohl Differentialformen sein sollen.

Bevor wir zu diesem Punkt gelangen, müssen wir noch zwei Fragestellungen klären:

- (1) Das Integral von Differentialformen macht Sinn über *orientierten* Untermannigfaltigkeiten. Wenn M eine orientierte Untermannigfaltigkeit ist, dann müssen wir uns davon überzeugen, dass dann auch ∂M eine orientierte Untermannigfaltigkeit ist.
- (2) Wenn ω eine $(k-1)$ -Form auf M ist, dann müssen wir das Differential $d\omega$ einführen. Dieses Differential muss ein mathematisches Objekt sein, welches wir über der k -dimensionalen Untermannigfaltigkeit M integrieren können. Es liegt nahe, dass es sich bei $d\omega$ um eine k -Form handeln soll.

In diesem Kapitel behandeln wir diese beide Themen. Im darauf folgenden Kapitel werden wir dann den Satz von Stokes formulieren und beweisen.

10.1. Orientierte Untermannigfaltigkeiten mit Rand.

Definition. Es sei M eine k -dimensionale Untermannigfaltigkeit mit Rand. Es sei $P \in \partial M$ und es sei $v \in T_P M$.

- (1) Wir sagen der Vektor v *zeigt nach außen*, wenn $v \notin T_P \partial M$ und wenn es eine Kurve $\gamma: (a, 0] \rightarrow M$ durch P gibt, so dass $\gamma'(0) = v$.
- (2) Wir sagen der Vektor v *zeigt nach innen*, wenn $v \notin T_P \partial M$ und wenn es eine Kurve $\gamma: [0, a) \rightarrow M$ durch P gibt, so dass $\gamma'(0) = v$.

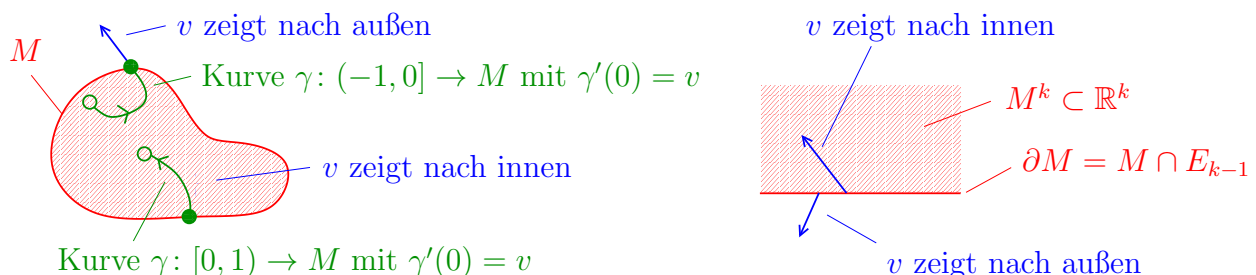


ABBILDUNG 60.

Beispiel. Es sei $M \subset H^k$ eine offene Teilmenge, dann ist M eine k -dimensionale Untermannigfaltigkeit mit Rand $\partial M = M \cap E_{k-1}$. Für $P \in \partial M$ und $(v_1, \dots, v_k) \in T_P M$ gilt dann

$$(v_1, \dots, v_k) \text{ zeigt nach außen} \iff v_k < 0,$$

und

$$(v_1, \dots, v_k) \text{ zeigt nach innen} \iff v_k > 0.$$

Es sei nun M eine orientierte k -dimensionale Untermannigfaltigkeit mit Rand und mit $k > 1$. Wir wollen nun zeigen, dass die $(k-1)$ -dimensionale Untermannigfaltigkeit ∂M ebenfalls orientierbar ist. Es sei $P \in \partial M$ und $v_1, \dots, v_{k-1}$ eine Basis von $T_P \partial M$. Wir definieren nun

$$\{v_1, \dots, v_{k-1}\} \text{ ist eine positive Basis von } T_P \partial M \iff \text{es gibt einen Vektor } w \text{ der nach außen zeigt, so dass } \{w, v_1, \dots, v_{k-1}\} \text{ eine positive Basis von } T_P M \text{ ist.}$$

Wir müssen nun noch zeigen, dass diese Orientierungen der Tangentialräume in der Tat eine Orientierung der Untermannigfaltigkeit ergeben. Dies ist die Aussage vom nächsten Lemma.

Lemma 10.1. *Die gerade eingeführten Orientierungen der Tangentialräume sind stetig bezüglich allen Karten, definieren also eine Orientierung der Untermannigfaltigkeit.*

Beweis. Wir zeigen zuerst folgende Behauptung.

Behauptung. Es sei M eine k -dimensionale Untermannigfaltigkeit. Es sei $P \in \partial M$ und es sei $v \in T_P M$. Zudem sei $\Phi: U \rightarrow V$ eine Karte um P . Wir setzen $Q = \Phi(P)$. Dann gilt

$$v \text{ zeigt nach außen} \iff \text{letzte Koordinate von } \Phi_*(v) \in T_Q E_k = \mathbb{R}^k \text{ ist negativ.}$$

Die Karte $\Phi: U \rightarrow V$ gibt insbesondere einen Diffeomorphismus zwischen der Untermannigfaltigkeit $M \cap U$ mit Rand $(\partial M) \cap U$ und der Untermannigfaltigkeit $H_k \cap V$ mit Rand $E_{k-1} \cap V$. Ein Diffeomorphismus zwischen Untermannigfaltigkeiten mit Rand schickt offensichtlich¹³⁷ Vektoren die nach außen zeigen wieder auf Vektoren die nach außen zeigen. Die Behauptung folgt nun aus dem obigen Beispiel. Wir haben damit die Behauptung bewiesen.

Für eine Karte der Untermannigfaltigkeit M kann man nun die Stetigkeit der Prä-Orientierung leicht mit der obigen Behauptung beweisen. Wir überlassen den Beweis als freiwillige Übungsaufgabe. $\square$

Konvention. Wenn M eine orientierte k -dimensionale Untermannigfaltigkeit mit Rand ist, dann betrachten wir ∂M immer mit der gerade eingeführten Orientierung.

Beispiele.

(1) Es sei

$$M := \{(x, y) \mid x^2 + y^2 \leq 1\}$$

¹³⁷Warum ist das offensichtlich?

die Einheitsscheibe im $\mathbb{R}^2$. Wir betrachten M als orientierte Untermannigfaltigkeit, wobei die Orientierung durch die übliche Orientierung von $\mathbb{R}^2$ gegeben ist. Die

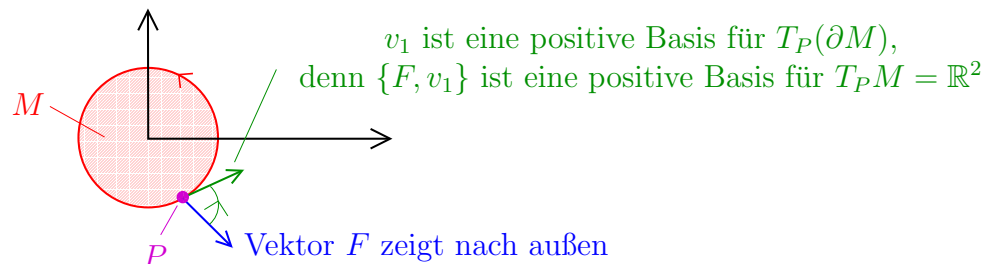


ABBILDUNG 61. Orientierung des Randes der Kreisscheibe in $\mathbb{R}^2$.

Untermannigfaltigkeit ∂M , d.h. der Einheitskreis, hat dann die Orientierung, welche “entgegen dem Uhrzeigersinn zeigt”.

- (2) Wir betrachten den Zylinder $Z = [-1, 1] \times S^1 \subset S^3$ mit der Orientierung, welche durch die Normalenvektoren gegeben sind, welche “nach außen” zeigen. Die Orientierungen auf den Randkomponente $\{-1\} \times S^1$ und $\{1\} \times S^1$ sind in Abbildung 2 skizziert.

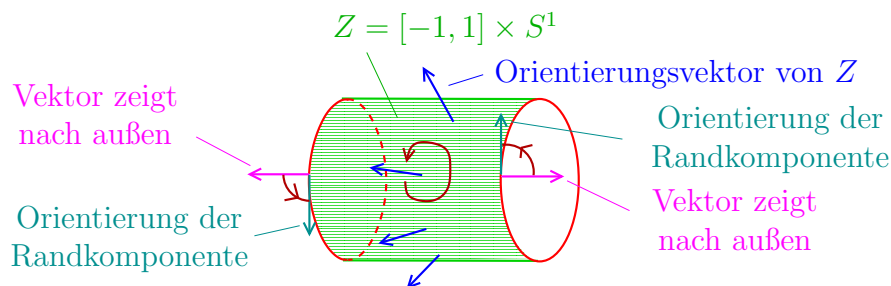


ABBILDUNG 62. Orientierung des Randes der Kreisscheibe in $\mathbb{R}^2$.

- (3) Es sei M eine kompakte k -dimensionale Untermannigfaltigkeit von $\mathbb{R}^k$. Wir hatten auf Seite 64 das *äußere Einheitsnormalenfeld* auf ∂M eingeführt. Auf Seite 123 hatten wir gesehen, dass ein solches äußere Einheitsnormalenfeld eine Orientierung auf ∂M definiert. Man kann nun leicht nachvollziehen, dass diese Orientierung gerade der Orientierung von ∂M auf Seite 134 entspricht.

Es sei zum Abschluß noch M eine zusammenhängende eindimensionale orientierte Untermannigfaltigkeit. Für $P \in \partial M$ setzen wir $\epsilon_P := 1$, wenn der Orientierungsvektor in $T_P M$ nach außen zeigt, und wir setzen $\epsilon_P := -1$, wenn der Orientierungsvektor in $T_P M$ nach innen zeigt. Wir definieren dann

$$\partial M := \bigcup_{P \text{ Randpunkt von } M} \epsilon_P \cdot P.$$

D.h. wir betrachten ∂M als 0-dimensionale Untermannigfaltigkeit mit einer Orientierung. Beispielsweise gilt für die Untermannigfaltigkeit $[a, b] \subset \mathbb{R}$ mit der üblichen Orientierung,

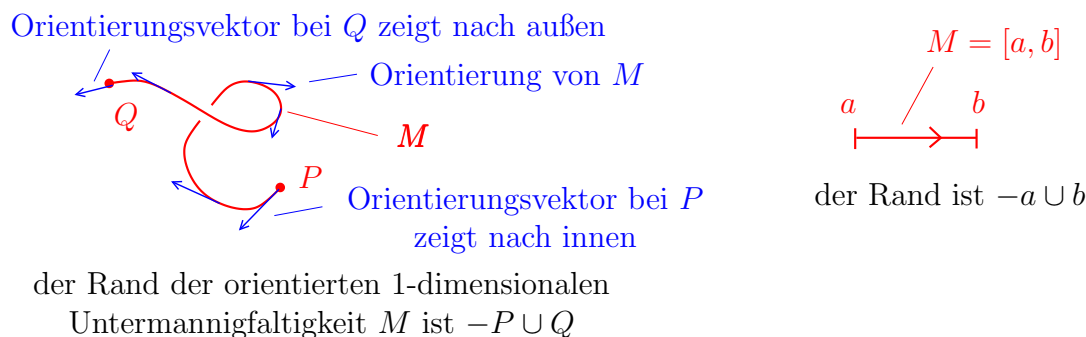


ABBILDUNG 63.

dass¹³⁸

$$\partial[a, b] = (-a) \cup b.$$

10.2. Das Differential von Differentialformen. Wir führen nun das Differential einer Differentialform ein. Wie üblich führen wir diesen Begriff erst für Differentialformen auf offenen Mengen von $\mathbb{R}^n$ ein, und definieren dann den Begriff für Differentialformen auf Untermannigfaltigkeiten mithilfe von Karten.

Es sei also im Folgenden U eine offene Teilmenge von $\mathbb{R}^k$ oder von H_k . Außerdem sei $\omega \in \Omega_l(U)$. Nach Satz 9.7 bilden die Dachprodukte $dx_{i_1} \wedge \cdots \wedge dx_{i_l}$ mit $i_1 < i_2 < \cdots < i_l$ eine Basis von $\wedge^l(\mathbb{R}^k)^*$. Wir können also ω wie folgt als Linearkombination betrachten:

$$\omega = \sum_{i_1 < \cdots < i_l} f_{i_1 \dots i_l} dx_{i_1} \wedge \cdots \wedge dx_{i_l},$$

wobei $f_{i_1 \dots i_l}$ reellwertige glatte Funktionen auf U sind. Wir definieren jetzt das *Differential von ω* ¹³⁹ als die $(l+1)$ -Form

$$d\omega := \sum_{i_1 < \cdots < i_l} \underset{\substack{\uparrow \\ \text{das totale Differential der Funktion } f_{i_1 \dots i_l}}}{df_{i_1 \dots i_l}} \wedge dx_{i_1} \wedge \cdots \wedge dx_{i_l}.$$

Beispiele.

- (1) Wenn $\omega \in \Omega_0(U)$, d.h. wenn ω eine glatte Funktion auf U ist, dann ist $d\omega$ per Definition gerade das totale Differential dieser glatten Funktion.

¹³⁸Streng genommen müsste man hier $-\{a\} \cup \{b\}$ schreiben, aber der besseren Lesbarkeit halber lassen wir die Klammern weg.

¹³⁹In der Literatur wird das Differential von ω manchmal auch die *äußere Ableitung von ω* genannt.

(2) Wir betrachten auf $\mathbb{R}^2$ die 1-Form $\omega = 3x^2y dx + x^3 dy$. Dann gilt

Definition des Differentials einer 1-Form

$$\begin{aligned}
 d\omega &= d(3x^2y dx + x^3 dy) \stackrel{\downarrow}{=} d(3x^2y) \wedge dx + d(x^3) \wedge dy \\
 &= \left(\frac{\partial 3x^2y}{\partial x} dx + \frac{\partial 3x^2y}{\partial y} dy \right) \wedge dx + \left(\frac{\partial x^3}{\partial x} dx + \frac{\partial x^3}{\partial y} dy \right) \wedge dy \\
 &\stackrel{\uparrow}{=} \text{Lemma 8.1} \\
 &= 6xy \underbrace{dx \wedge dx}_{=0} + 3x^2 \underbrace{dy \wedge dx}_{=-dx \wedge dy} + 3x^2 dx \wedge dy + 0 \cdot \underbrace{dy \wedge dy}_{=0} = 0.
 \end{aligned}$$

(3) Für $\omega = dx_i \wedge dx_j$ ist

$$d\omega = d(dx_i \wedge dx_j) = d(1 \cdot dx_i \wedge dx_j) = \underbrace{d1}_{=0} \wedge dx_i \wedge dx_j = 0.$$

Der folgende Satz faßt einige der wichtigsten Eigenschaften des Differentials von Differentialformen zusammen.

Satz 10.2. *Es sei $U \subset \mathbb{R}^k$ eine offene Teilmenge. Es seien $\omega, \omega' \in \Omega_l(U)$, $\sigma \in \Omega_m(U)$ und $l \in \mathbb{R}$. Dann gilt*

- (1) $d(\omega + \omega') = d\omega + d\omega'$,
- (2) $d(l\omega) = l \cdot d\omega$,
- (3) *die $(l+1)$ -Form $d\omega$ ist eine glatte Form,*
- (4) *es gilt die "Produktregel" für die Differentiation von Differentialformen¹⁴⁰*

$$d(\omega \wedge \sigma) = d\omega \wedge \sigma + (-1)^l \omega \wedge d\sigma,$$

- (5) *es ist $d(d\omega) = 0$.*

Genau die gleichen Aussagen gelten auch, falls U eine offene Teilmenge der Halbebene $H_k = \{(x_1, \dots, x_k) \in \mathbb{R}^k \mid x_k \geq 0\}$ ist.

Beispiel. In Übungsblatt 5 hatten wir gesehen, dass wenn

$$\omega = f(x, y) dx + g(x, y) dy$$

eine glatte 1-Form auf $\mathbb{R}^2$ ist mit

$$\frac{\partial}{\partial y} f(x, y) = \frac{\partial}{\partial x} g(x, y),$$

dann ist ω exakt, d.h. es gibt eine glatte Funktion $f: \mathbb{R}^2 \rightarrow \mathbb{R}$ mit $df = \omega$. Das Beispiel der 1-Form $\omega = 3x^2y dx + x^3 dy$ auf Seite 137 ist von dieser Form. Also folgt aus Satz 10.2 direkt, dass $d\omega = d(df) = 0$.

Beispiele. Wir betrachten die Aussagen (4) und (5) für 0-Formen, d.h. für reellwertige Funktionen.

¹⁴⁰Die Produktregel ist also ganz ähnlich der üblichen Produktregel für Ableitungen, allerdings taucht noch ein extra Vorzeichen auf.

- (1) Es seien $f, g: U \rightarrow \mathbb{R}$ zwei glatte Funktionen auf einer offenen Teilmenge $U \subset \mathbb{R}^k$. Wir können f und g als 0-Formen auffassen. Satz 10.2 (4) besagt nun also, dass¹⁴¹

$$d(fg) = df \cdot g + f \cdot dg.$$

Wir wenden nun Lemma 8.1 auf $d(fg)$, df und dg an und erhalten, dass

$$\sum_{i=1}^k \frac{\partial(fg)}{\partial x_i} dx_i = \left(\sum_{i=1}^k \frac{\partial f}{\partial x_i} dx_i \right) \cdot g + f \cdot \left(\sum_{i=1}^k \frac{\partial g}{\partial x_i} dx_i \right).$$

Durch Umsortieren erhalten wir, dass

$$\sum_{i=1}^k \frac{\partial(fg)}{\partial x_i} dx_i = \sum_{i=1}^k \left(\frac{\partial f}{\partial x_i} \cdot g + f \cdot \frac{\partial g}{\partial x_i} \right) dx_i.$$

Nachdem die 1-Formen $dx_1, \dots, dx_k$ linear unabhängig sind erhalten wir jetzt für jedes $i \in \{1, \dots, k\}$, dass

$$\frac{\partial(fg)}{\partial x_i} = \frac{\partial f}{\partial x_i} g + f \cdot \frac{\partial g}{\partial x_i}.$$

In anderen Worten, wir sehen, dass die Produktregel für 0-Formen *äquivalent* ist zur üblichen Produktformel für partielle Ableitungen.

- (2) Es sei $f: U \rightarrow \mathbb{R}$ eine glatte Funktion auf einer offenen Teilmenge $U \subset \mathbb{R}^k$. Dann ist

$$\begin{array}{ccc} d(df) & = d\left(\sum_{i=1}^k \frac{\partial f}{\partial x_i} dx_i\right) & = \sum_{i=1}^k d\left(\frac{\partial f}{\partial x_i}\right) \wedge dx_i \\ \uparrow & \text{Lemma 8.1} & \uparrow \text{Definition von } d \\ & = \sum_{i=1}^k \left(\sum_{j=1}^k \frac{\partial^2 f}{\partial x_j \partial x_i} dx_j \right) \wedge dx_i & = \sum_{i < j} \left(\frac{\partial^2 f}{\partial x_j \partial x_i} - \frac{\partial^2 f}{\partial x_i \partial x_j} \right) \cdot dx_i \wedge dx_j. \\ \uparrow & \text{Lemma 8.1} & \uparrow \\ & & \text{denn } dx_i \wedge dx_j = -dx_j \wedge dx_i \text{ und } dx_i \wedge dx_i = 0 \end{array}$$

Die 2-Formen $dx_i \wedge dx_j$ mit $i < j$ sind nach Satz 9.7 linear unabhängig. Wir sehen also, dass die Aussage $d(df) = 0$ *äquivalent* ist zu der schon in Analysis II bewiesenen Aussage, dass die partiellen Ableitungen einer glatten Funktion vertauschbar sind.

Notation. Bevor wir uns dem Beweis von Satz 10.2 zuwenden führen wir noch folgende Notation ein. Für eine aufsteigende Folge $I = 1 \leq i_1 < \dots < i_k \leq n$ schreiben wir

$$dx_I := dx_{i_1} \wedge \dots \wedge dx_{i_k}.$$

Wir bezeichnen mit $|I| = k$ die Anzahl der Elemente in der aufsteigenden Folge. Es sei nun U eine n -dimensionale Untermannigfaltigkeit von $\mathbb{R}^n$ und es sei ω eine k -Form auf U . Wir können dann ω schreiben als

$$\omega = \sum_{|I|=k} f_I dx_I,$$

¹⁴¹Das Dachprodukt von zwei 0-Formen ist einfach das übliche Produkt von Funktionen. Warum?

hierbei summieren wir über alle aufsteigenden Folgen $I = 1 \leq i_1 < \dots < i_k \leq n$ der Länge k .

Beweis von Satz 10.2. Die ersten beiden Aussagen von Satz 10.2 sind trivial. Die dritte Aussage folgt leicht aus den Definitionen und aus Lemma 8.1.

Wir wenden uns der vierten Aussage zu. Im obigen Beispiel hatten wir gesehen, dass die Produktformel für 0-Formen gerade der üblichen Produktformel für partielle Ableitungen entspricht, welche wir schon in Analysis II bewiesen hatten. Wir wollen nun den allgemeinen Fall auf den Fall von 0-Formen zurückführen.

Es sei also $\omega \in \Omega_l(U)$ und es sei $\sigma \in \Omega_m(U)$. Wir schreiben

$$\omega = \sum_{|I|=l} f_I dx_I \quad \text{und} \quad \sigma = \sum_{|J|=m} g_J dx_J$$

und erhalten

$$d\omega = \sum_{|I|=l} df_I \wedge dx_I \quad \text{und} \quad d\sigma = \sum_{|J|=m} dg_J \wedge dx_J.$$

Dann ist

$$\begin{aligned} d(\omega \wedge \sigma) &= d\left(\sum_{|I|=l} f_I dx_I \wedge \sum_{|J|=m} g_J dx_J\right) && \begin{array}{l} \text{Multilinearität des Dachprodukts} \\ \downarrow \end{array} \\ &= \sum_{I,J} d(f_I g_J) \cdot dx_I \wedge dx_J && \downarrow \\ &\uparrow && \uparrow \\ &\text{Definition des Differentials} && \text{Produktformel für Produkt von 0-Formen} \\ &= \sum_{I,J} \left((df_I \wedge dx_I) \wedge g_J dx_J + (-1)^l f_I dx_I \wedge (dg_J \wedge dx_J) \right) \\ &\uparrow \\ &\text{Lemma 9.10} \\ &= d\omega \wedge \sigma + (-1)^l \omega \wedge d\sigma. \end{aligned}$$

Wir beweisen nun noch die letzte Aussage. Im obigen Beispiel hatten wir gesehen, dass die Aussage für 0-Formen gerade der Vertauschbarkeit von partiellen Ableitungen entspricht, welche wir schon in Analysis II bewiesen hatten. Wir wollen nun auch dieses Mal den allgemeinen Fall auf den Fall von 0-Formen zurückführen.

Es sei also $\omega \in \Omega_l(U)$. Wir schreiben wiederum $\omega = \sum_{|I|=l} f_I dx_I$. Dann ist

$$\begin{aligned} d(d(\omega)) &= d\left(\sum_{|I|=l} df_I \wedge dx_I\right) && \begin{array}{l} = 0 \text{ nach obigem Beispiel} \\ \downarrow \end{array} \\ &\uparrow && \downarrow \\ &\text{Aussage (4)} && \sum_{|I|=l} \left(\overbrace{d(df_I)} \wedge dx_I + (-1) \cdot df_I \wedge \underbrace{d(dx_I)} \right) = 0. \\ &&& \uparrow \\ &&& d(dx_I) = d(1 \cdot dx_I) = d1 \wedge dx_I = 0 \quad \square \end{aligned}$$

Zudem gilt auch folgender Satz:

Satz 10.3. *Es seien $U \subset \mathbb{R}^m$ und $V \subset \mathbb{R}^n$ offenen Mengen und es sei $\Theta: V \rightarrow U$ eine glatte Abbildung. Für jede glatte k -Form ω auf U gilt dann*

$$d(\Theta^*\omega) = \Theta^*(d\omega).$$

Die analoge Aussage gilt auch für offene Teilmengen $U \subset H_m$ und $V \subset H_n$.

Der Satz besagt also, dass folgendes Diagramm von Abbildungen kommutiert:¹⁴²

$$\begin{array}{ccc} \Omega_l(U) & \xrightarrow{\Theta^*} & \Omega_l(V) \\ \downarrow d & & \downarrow d \\ \Omega_{l+1}(U) & \xrightarrow{\Theta^*} & \Omega_{l+1}(V). \end{array}$$

Beispiel. Im Beweis von Satz 10.3 werden wir sehen, dass für eine 0-Form ω die Aussage eine Umformulierung der üblichen Kettenregel ist.

Beweis. Es seien $U \subset \mathbb{R}^m$ und $V \subset \mathbb{R}^n$ offene Mengen und es sei $\Theta = (\Theta_1, \dots, \Theta_m): V \rightarrow U$ eine glatte Abbildung.

Wir betrachten zuerst den Spezialfall einer 0-Form. Es sei also $f: U \rightarrow \mathbb{R}$ eine glatte Funktion. Wir bezeichnen mit $dx_1, \dots, dx_m$ die üblichen 1-Formen auf $\mathbb{R}^m$ und wir bezeichnen mit $dy_1, \dots, dy_n$ die üblichen 1-Formen auf $\mathbb{R}^n$. Wir machen nun noch zwei Vorbemerkungen:

- (1) für eine Funktion $g: U \rightarrow \mathbb{R}$ folgt sofort aus den Definitionen, folgende Gleichheit von 1-Formen auf V :

$$\Theta^*(g \cdot dx_i) = (g \circ \Theta)d\Theta_i.$$

- (2) auf Seite 114 hatten wir schon erwähnt, dass für beliebige Differentialformen α und β auf U die Gleichheit

$$\Theta^*(\alpha \wedge \beta) = \Theta^*\alpha \wedge \Theta^*\beta$$

gilt.

Dann gilt

$$\begin{aligned} d(\Theta^*f) &= d(f \circ \Theta) = \sum_{j=1}^n \frac{\partial(f \circ \Theta)}{\partial y_j} dy_j &&= \sum_{j=1}^n \sum_{i=1}^m \left(\frac{\partial f}{\partial x_i} \circ \Theta \right) \cdot \frac{\partial \Theta_i}{\partial y_j} dy_j \\ &\quad \uparrow \text{Lemma 8.1 angewandt auf } f \circ \Theta && \uparrow \text{ Kettenregel} \\ &= \sum_{i=1}^m \left(\frac{\partial f}{\partial x_i} \circ \Theta \right) d\Theta_i &&= \Theta^* \left(\sum_{i=1}^m \frac{\partial f}{\partial x_i} dx_i \right) = \Theta^*(df). \\ &\quad \uparrow \text{Lemma 8.1 angewandt auf } \Theta_i && \uparrow \text{ siehe Vorbemerkung (1) } \quad \uparrow \text{ Lemma 8.1} \end{aligned}$$

Wir haben damit diesen Spezialfall bewiesen.

¹⁴²D.h. die Verknüpfung der Abbildung oben mit der Abbildung rechts ergibt die Verknüpfung der Abbildung links und der Abbildung unten.

Es sei nun ω eine beliebige glatte k -Form. Wir verwenden die Notation von Seite 138 und wir schreiben

$$\omega = \sum_{|I|=k} f_I dx_I.$$

Unter Verwendung der Vorbemerkungen (1) und (2) erhalten wir, dass

$$\Theta^* \omega = \sum_{|I|=k} (f_I \circ \Theta) d\Theta_I,$$

wobei $d\Theta_I = d\Theta_{i_1} \wedge \cdots \wedge d\Theta_{i_k}$. Wir berechnen nun, dass

$$\begin{aligned} d(\Theta^* \omega) &= \sum_{|I|=k} d(f_I \circ \Theta) \wedge d\Theta_I = \sum_{|I|=k} \Theta^*(df_I) \wedge \Theta^*(dx_I) \\ &\quad \uparrow \text{der obige Spezialfall einer 0-Form} \\ &= \Theta^* \left(\sum_{|I|=k} df_I \wedge dx_I \right) = \Theta^*(d\omega). \\ &\quad \uparrow \text{Vorbemerkung (2)} \end{aligned}$$

□

Wir wollen nun das Differential von Differentialformen auf Untermannigfaltigkeiten einführen. Die Idee, wie üblich, ist dabei die Betrachtung von Differentialformen auf Untermannigfaltigkeiten mithilfe von Karten auf den oben behandelten Fall von Differentialformen auf offenen Teilmengen von $\mathbb{R}^k$ und auf offenen Teilmengen von H^k zurückzuführen.

Es sei also M eine k -dimensionale Untermannigfaltigkeit und es sei $\omega \in \Omega_l(M)$. Es sei zuerst $\Phi: U \rightarrow V$ eine Karte für M von Typ (i). Wir bezeichnen mit $\Psi: V \rightarrow U$ die Umkehrabbildung von Φ . Dann ist $\Psi^* \omega$ eine l -Form auf $V \cap E_k$, also auf einer offenen Teilmenge von $E_k = \mathbb{R}^k$. Wir definieren nun

$$d\omega|_{U \cap M} := \Phi^*(d(\Psi^* \omega)).$$

Es ist eventuell hierzu hilfreich folgendes Diagramm zu betrachten

$$\begin{array}{ccc} \Omega_l(U \cap M) & \longrightarrow & \Omega_{l+1}(U \cap M) \\ \downarrow \Psi^* & & \uparrow \Phi^* \\ \Omega_l(V \cap E_k) & \xrightarrow{d} & \Omega_{l+1}(V \cap E_k). \end{array}$$

Wir definieren also die obige Abbildung mithilfe der anderen drei Abbildungen. Ganz analog definieren wir auch $d\omega|_U$ für Karten $\Phi: U \rightarrow V$ von Typ (ii). Wir haben also jetzt $d\omega$ auf den Definitionsbereichen von Karten von M definiert.

Wir müssen nun noch zeigen, dass $d\omega$ wohl-definiert ist. Genauer gesagt müssen wir folgendes Lemma beweisen.

Lemma 10.4. *Es sei M eine k -dimensionale Untermannigfaltigkeit und es sei $\omega \in \Omega_l(M)$. Es seien $\Phi: U \rightarrow V$ und $\tilde{\Phi}: \tilde{U} \rightarrow \tilde{V}$ zwei Karten für M mit Umkehrabbildungen Ψ und $\tilde{\Psi}$. Dann gilt auf $U \cap \tilde{U}$*

$$\Phi^*(d(\Psi^* \omega)) = \tilde{\Phi}^*(d(\tilde{\Psi}^* \omega)).$$

Beweis. Der Einfachheit halber betrachten wir nur den Fall, dass sowohl Φ als auch $\tilde{\Phi}$ Karten von Typ (i) sind. Wir können o.B.d.A. annehmen, dass $U = \tilde{U}$. Wir bezeichnen mit Θ den Diffeomorphismus $\Theta := \tilde{\Phi} \circ \Psi: V \cap E_k \rightarrow \tilde{V} \cap E_k$. Dann gilt

$$\begin{aligned} \Phi^*(d(\Psi^*\omega)) &= \Phi^*(d(\Psi^*(\tilde{\Phi}^*\tilde{\Psi}^*\omega))) = \Phi^*(d(\Theta^*(\tilde{\Psi}^*\omega))) = \Phi^*(\Theta^*d(\tilde{\Psi}^*\omega)) \\ &\quad \uparrow \qquad \text{denn } \tilde{\Phi}^* \circ \tilde{\Psi}^* = (\tilde{\Psi} \circ \tilde{\Phi})^* = \text{id} \qquad \qquad \qquad \uparrow \text{Satz 10.3} \\ &= \Phi^*(\Psi^*(\tilde{\Phi}^*(\tilde{\Psi}^*\omega))) = \tilde{\Phi}^*(d(\tilde{\Psi}^*\omega)). \end{aligned}$$

□

Beispiel. Es sei $\omega \in \Omega_0(M)$, dies bedeutet, dass ω eine glatte Funktion auf M ist. In diesem Fall ist $d\omega \in \Omega_1(M)$ genau die 1-Form, welche wir in Kapitel 8.1 eingeführt hatten.¹⁴³

Folgender Satz folgt leicht aus den Aussagen von Satz 10.2 und Satz 10.3.

Satz 10.5. *Es sei M eine Untermannigfaltigkeit.*

- (1) *Es gelten die gleichen Aussagen wie in Satz 10.2, insbesondere gilt für $\omega \in \Omega_l(M)$ und $\sigma \in \Omega_m(M)$ die Produktregel*

$$d(\omega \wedge \sigma) = d\omega \wedge \sigma + (-1)^l \omega \wedge d\sigma.$$

Zudem gilt für alle $\omega \in \Omega_l(M)$, dass

$$d(d\omega) = 0.$$

- (2) *Es sei $\varphi: M \rightarrow N$ eine glatte Abbildung zwischen zwei Untermannigfaltigkeiten. Für jede glatte k -Form ω auf N gilt*

$$d(\varphi^*\omega) = \varphi^*(d\omega).$$

Wir betrachten nun noch ein ausführlicheres Beispiel: Es sei M eine n -dimensionale Untermannigfaltigkeit von $\mathbb{R}^n$. Zudem sei $F = (F_1, \dots, F_n): M \rightarrow \mathbb{R}^n$ ein glattes Vektorfeld auf M . Wir hatten in Kapitel 9.4 gesehen, dass F eine $(n-1)$ -Form $\delta(F)$ auf M wie folgt definiert: Wir ordnen P in M die alternierende $(n-1)$ -Form

$$\begin{aligned} (T_P M)^{n-1} &\rightarrow \mathbb{R} \\ (v_1, \dots, v_{n-1}) &\mapsto \det(F(P) v_1 \dots v_{n-1}) \end{aligned}$$

zu. Wir wollen nun die n -Form $d\delta(F)$ auf M bestimmen.

¹⁴³Es genügt diese Aussage zu überprüfen, für 0-Formen ω deren Träger in einer Karte $\Phi: U \rightarrow V$ enthalten ist. In diesem Fall gilt für jedes $v \in T_P M$ mit $P \in U \cap M$, dass

$$\begin{array}{ccccccc} d\omega(v) = \Phi^*(d(\Psi^*\omega))(v) &= & d(\Psi^*\omega)(\Phi_*v) &= & D(\Psi^*\omega)(\Phi_*v) &= & D(\omega \circ \Psi)(\Phi_*v) = D\omega(\Psi_*\Phi_*v) = D\omega_P(v). \\ \uparrow & & \uparrow & & \uparrow & & \uparrow \\ \text{Definition von } d & & \text{Lemma 8.2} & & \text{da } \Psi^*\omega = \omega \circ \Psi & & \text{da } d = D \text{ für} \\ \text{einer } k\text{-Form auf } M & & & & \text{Funktion auf } \mathbb{R}^k & & \text{Kettenregel} \end{array}$$

Wir müssen dazu erst einmal $\delta(F)$ in die Form bringen, auf die wir die Definition des Integrals anwenden können. Wir verwenden dabei folgende hilfreich Standardnotation. Für $i = 1, \dots, n$ schreiben wir

$$dx_1 \wedge \dots \wedge \widehat{dx_i} \wedge \dots \wedge dx_n := dx_1 \wedge \dots \wedge dx_{i-1} \wedge dx_{i+1} \wedge \dots \wedge dx_n.$$

Wir können nun $\delta(F)$ umschreiben.

Behauptung. Es ist

$$\delta(F) = \sum_{i=1}^n (-1)^{i-1} F_i \cdot dx_1 \wedge \dots \wedge \widehat{dx_i} \wedge \dots \wedge dx_n.$$

Wir müssen nur zeigen, dass die linke Seite und die rechte Seite auf beliebige Vektoren $v_1, \dots, v_{n-1}$ ausgewertet das gleiche Ergebnis liefern. Dies ist in der Tat der Fall, denn für $v_i = (v_{i1}, \dots, v_{in}), i = 1, \dots, n-1$ gilt:

$$\begin{aligned} \delta(F)(v_1, \dots, v_{n-1}) &= \det \begin{pmatrix} F_1 & v_{11} & \dots & v_{n-1,1} \\ \vdots & \vdots & & \vdots \\ F_n & v_{n1} & \dots & v_{nn} \end{pmatrix} = \sum_{i=1}^n (-1)^{i-1} F_i \det \begin{pmatrix} v_1 & \dots & v_{n-1} \\ \text{mit } i\text{-ter Zeile} \\ \text{entfernt} \end{pmatrix} \\ &= \sum_{i=1}^n (-1)^{i-1} F_i \cdot (dx_1 \wedge \dots \wedge \widehat{dx_i} \wedge \dots \wedge dx_n)(v_1, \dots, v_{n-1}). \end{aligned}$$

↑
folgt dabei sofort aus der Definition von dx_j und der Definition des Dachprodukts

Wir haben damit die Behauptung bewiesen.

Es folgt nun, dass

$$\begin{aligned} d\delta(F) &= \sum_{i=1}^n d((-1)^{i-1} F_i \cdot dx_1 \wedge \dots \wedge \widehat{dx_i} \wedge \dots \wedge dx_n) \\ &= \sum_{i=1}^n (-1)^{i-1} dF_i \wedge dx_1 \wedge \dots \wedge \widehat{dx_i} \wedge \dots \wedge dx_n \\ &= \sum_{i=1}^n (-1)^{i-1} \left(\sum_{j=1}^n \frac{\partial F_i}{\partial x_j} dx_j \right) \wedge dx_1 \wedge \dots \wedge \widehat{dx_i} \wedge \dots \wedge dx_n \\ &\quad \uparrow \\ &\text{Lemma 8.1} \\ &= \sum_{i=1}^n (-1)^{i-1} \frac{\partial F}{\partial x_i} (-1)^{i-1} dx_1 \wedge \dots \wedge dx_{i-1} \wedge dx_i \wedge dx_{i+1} \wedge \dots \wedge dx_n \\ &\quad \uparrow \\ &\text{denn } dx_j \wedge dx_j = 0 \text{ und } dx_j \wedge dx_k = -dx_k \wedge dx_j \\ &= \operatorname{div}(F) \cdot dx_1 \wedge \dots \wedge dx_n. \end{aligned}$$

Zusammenfassend erhalten wir also folgendes Lemma.

Lemma 10.6. *Es sei $M \subset \mathbb{R}^n$ eine n -dimensionale Untermannigfaltigkeit und F ein glattes Vektorfeld auf M . Dann gilt*

$$d\delta(F) = \operatorname{div} F \, dx_1 \wedge \dots \wedge dx_n.$$

Wir beschließen das Kapitel mit folgendem Lemma.

Lemma 10.7. *Es sei M eine Untermannigfaltigkeit. Die glatten Funktionen $f: M \rightarrow \mathbb{R}$ mit $df = 0$ sind gerade die lokal konstanten¹⁴⁴ Funktionen auf M .*

Beweis. Es sei $f: M \rightarrow \mathbb{R}$ eine glatte Funktion auf einer k -dimensionalen Untermannigfaltigkeit M . Es sei $\Phi: U \rightarrow V$ eine Karte, wobei U und V zusammenhängend ist. Dann gilt

- (1) f ist konstant auf U , genau dann, wenn $f \circ \Phi^{-1}$ konstant auf V ist,
- (2) $df = 0$ auf U , genau dann, wenn $d(f \circ \Phi^{-1}) = 0$.¹⁴⁵

Es genügt nun die Behauptung für Funktionen auf offenen Teilmengen von $\mathbb{R}^k$ zu beweisen. Für eine konstante Funktion $g: V \rightarrow \mathbb{R}$ ist natürlich $dg = 0$. Umgekehrt, sei $g: V \rightarrow \mathbb{R}$ eine glatte Funktion mit $dg = 0$. Aus Lemma 8.1 folgt, dass die partiellen Ableitungen von g verschwinden. Lemma 8.3 aus Analysis II besagt nun, dass g auch schon konstant ist. $\square$

¹⁴⁴Eine Funktion $f: M \rightarrow \mathbb{R}$ heißt *lokal konstant*, wenn es zu jedem Punkt $P \in M$ eine offene Umgebung U gibt, so dass $f|_U$ konstant ist.

¹⁴⁵Diese Aussage folgt aus Satz 10.5 (2).

11. DER SATZ VON STOKES

11.1. Formulierung des Satz von Stokes. Wir haben nun alle Vorbereitungen für die Formulierung des Satzes von Stokes getroffen. Der Satz von Stokes lautet wie folgt:

Satz 11.1. (Satz von Stokes) *Es sei M eine orientierte k -dimensionale kompakte Untermannigfaltigkeit und es sei $\omega \in \Omega_{k-1}(M)$. Dann gilt*

$$\int_M d\omega = \int_{\partial M} \omega.$$

Insbesondere, wenn M geschlossen¹⁴⁶ ist, dann gilt¹⁴⁷

$$\int_M d\omega = 0.$$

Unser Ziel war es, eine Verallgemeinerung von Satz 9.2 und vom Gaußschen Integralsatz zu finden. Es ist offensichtlich^{148,149}, dass der Satz von Stokes im Falle $k = 1$ genau der Aussage von Satz 9.2 entspricht. Wir werden nun zeigen, dass wir den Gaußschen Integralsatz aus dem Satz von Stokes herleiten können.

Es sei also $M \subset \mathbb{R}^n$ eine kompakte n -dimensionale Untermannigfaltigkeit mit Rand. Wir bezeichnen mit ν das äußere Einheitsnormalenfeld auf ∂M . Es sei $F: M \rightarrow \mathbb{R}^n$ ein glattes Vektorfeld. Wir müssen zeigen, dass

$$\int_M \operatorname{div} F = \int_{\partial M} F \cdot \nu.$$

Wie in Kapitel 9.4 betrachten wir die zu F zugehörige $(n-1)$ -Form $\delta(F)$ auf M . Zur Erinnerung, für P ist $\delta(F)_P$ die alternierende $(n-1)$ -Form

$$\begin{aligned} (T_P M)^{n-1} &\rightarrow \mathbb{R} \\ (v_1, \dots, v_{n-1}) &\mapsto \det(F(P) v_1 \dots v_{n-1}). \end{aligned}$$

Es folgt nun, dass

$$\begin{array}{ccccccc} \int_M \operatorname{div} F & = & \int_M \operatorname{div} F \, dx_1 \wedge \dots \wedge dx_n & = & \int_M d\delta(F) & = & \int_{\partial M} \delta(F) = \int_{\partial M} F \cdot \nu. \\ & \uparrow & & \uparrow & \uparrow & \uparrow & \\ & \text{Lemma 9.22} & & \text{Lemma 10.6} & \text{Satz von Stokes} & & \text{Lemma 9.23} \end{array}$$

¹⁴⁶Zur Erinnerung, “geschlossene Untermannigfaltigkeit” bedeutet eine kompakte Untermannigfaltigkeit M mit $\partial M = \emptyset$.

¹⁴⁷In diesem Fall ist also $\partial M = \emptyset$, also verschwindet das Integral auf der rechten Seite der obigen Gleichheit.

¹⁴⁸Hierbei verwenden wir die Konvention aus Kapitel 10.1 für die Orientierung des Randes einer orientierten eindimensionalen Untermannigfaltigkeit

¹⁴⁹Diese Aussage ist zumindest offensichtlich, wenn man noch den Überblick über alle Definitionen und Konventionen besitzt.

Wir haben damit gezeigt, dass der Satz von Stokes insbesondere den Gaußschen Integralsatz 6.6 beinhaltet.¹⁵⁰

11.2. Beweis vom Satz von Stokes. Wir wollen in diesem Kapitel den Satz von Stokes beweisen. Wie für den Beweis vom Gaußschen Integralsatz wollen wir dazu erst einmal einen Spezialfall beweisen.

Für einen offenen Quader $Q \subset H_n \subset \mathbb{R}^{n151}$ schreiben wir $\partial_0 Q = Q \cap H_n$. Folgendes Lemma ist streng genommen kein Spezialfall vom Satz von Stokes, aber wir werden den Beweis vom Satz von Stokes problemlos auf dieses Lemma zurückführen.

Lemma 11.2. *Es sei Q ein offener Quader in H_n und es sei ω eine glatte $(n-1)$ -Form auf Q , welche auf $\partial_1 Q := \partial \overline{Q} \setminus \partial_0 Q$ verschwindet. Dann gilt*

$$\int_Q d\omega = \int_{\partial_0 Q} \omega.$$

Bemerkung. Wir hatten im vorherigen Kapitel gesehen, dass der Satz von Stokes für n -dimensionale Untermannigfaltigkeiten von $\mathbb{R}^n$ gerade dem Gaußschen Integralsatz 6.6 entspricht. Genau das gleiche Argument erlaubt es uns die Aussage von Lemma 11.2 auf den Gaußschen Integralsatz 6.9 für Stempel zurückzuführen. Wir werden jedoch gleich sehen, dass man Lemma 11.2 auch problemlos direkt beweisen kann, und der Beweis deutlich einfacher ist, als der etwas umständliche Beweis vom Gaußschen Integralsatz 6.9 für Stempel.

Beweis. Der Einfachheit halber beweisen wir das Lemma für $n = 2$. Der allgemeine Fall wird ganz genauso bewiesen.

Wir schreiben $\overline{Q} = [a, b] \times [c, d]$ und $\omega = f(x, y) dx + g(x, y) dy$. Nach Voraussetzung können wir $f(x, y)$ und $g(x, y)$ stetig auf den Quader $\overline{Q}$ fortsetzen, so dass $f(x, y)$ und $g(x, y)$ auf $\partial_1 Q$ verschwinden, d.h. es ist $f(x, d) = g(x, d) = 0$ für alle $x \in [a, b]$ und $f(a, y) = g(a, y) = g(b, y) = f(b, y) = 0$ für alle $y \in [c, d]$.

Wir betrachten im Folgenden den Fall, dass $c = 0$. Der Fall, dass $c > 0$ ist wird ganz ähnlich bewiesen. Es ist nun

$$\int_Q d\omega \underset{\uparrow}{=} \int_{\overline{Q}} \left(\left(\frac{\partial f}{\partial x} dx + \frac{\partial f}{\partial y} dy \right) \wedge dx + \left(\frac{\partial g}{\partial x} dx + \frac{\partial g}{\partial y} dy \right) \wedge dy \right)$$

Definition von $d\omega$ und $\overline{Q} \setminus Q$ ist eine Nullmenge

¹⁵⁰Wenn man genau hinschaut, dann stimmt das nicht ganz, der Gaußsche Integralsatz 6.6 wurde für stetig differenzierbare Vektorfelder formuliert und bewiesen, während wir den Satz von Stokes der Einfachheit halber nur für glatte Differentialformen formuliert hatten. Der Satz von Stokes gilt aber auch ganz analog, mit dem gleichen Beweis, für stetig differenzierbare Differentialformen.

¹⁵¹Wir sagen $Q \subset H_n$ ist ein offener Quader, wenn es einen offenen Quader in $\mathbb{R}^n$ gibt, dessen Durchschnitt mit H_n gerade Q ist.

$$\begin{aligned}
&= \int_{[a,b] \times [0,c]} -\frac{\partial f}{\partial y} dx \wedge dy + \int_{[a,b] \times [0,c]} \frac{\partial f}{\partial y} dx \wedge dy \\
&\uparrow \\
&\text{denn } dx \wedge dx = dy \wedge dy = 0 \text{ und } dy \wedge dx = -dx \wedge dy \\
&= \int_{x=a}^{x=b} \int_{y=0}^{y=d} -\frac{\partial f}{\partial y} dy dx + \int_{y=0}^{y=c} \int_{x=a}^{x=b} \frac{\partial g}{\partial x} dx dy \\
&\uparrow \\
&\text{Lemma 9.22 und Satz 3.8 von Fubini} \\
&= \int_{x=a}^{x=b} \underbrace{[f(x, d) - f(x, 0)]}_{=0} dx + \int_{y=0}^{y=c} \underbrace{[g(b, y) - g(a, y)]}_{=0} dy \\
&= \int_{x=a}^{x=b} f(x, 0) dx = \int_{\partial_0 Q} f dx + g dy = \int_{\partial_0 Q} \omega. \\
&\quad \uparrow \\
&\quad \text{denn } \partial_0 Q = [a, b] \times 0 \text{ und } dy(1, 0) = 0
\end{aligned}$$

Wir haben damit die gewünschte Gleichheit bewiesen. $\square$

Wir wenden uns nun dem eigentlichen Beweis vom Satz von Stokes zu.

Beweis von Satz 11.1. Es sei also M eine orientierte k -dimensionale Untermannigfaltigkeit von $\mathbb{R}^n$ und es sei $\omega \in \Omega_{k-1}(M)$. Wir müssen zeigen, dass

$$\int_M d\omega = \int_{\partial M} \omega.$$

Ganz analog zum Beweis vom Gaußschen Integralsatz wollen wir nun den allgemeinen Fall auf den oben betrachteten Spezialfall zurückführen.

Satz 9.18 besagt, dass es einen endlichen Atlas $\{\Phi_i: U_i \rightarrow V_i\}_{i=1, \dots, r}$ mit orientierungserhaltenden Karten sowie glatte Funktionen $z_1, \dots, z_r: M \rightarrow [0, 1]$ gibt, welche folgende Eigenschaften erfüllen:

- (1) für alle $i \in \{1, \dots, r\}$ gilt $\text{Träger}(z_i) \subset U_i \cap M$,
- (2) es ist $z_1 + \dots + z_r = 1$,
- (3) für alle $i \in \{1, \dots, r\}$ gilt
 - (a) $\Phi_i(U_i \cap M) \subset H_k$,
 - (b) die Menge $V_i \cap H_k = \Phi_i(U_i \cap M)$ ist ein offener Quader in H_k ,
 - (c) die Funktion $z_i \circ \Phi_i^{-1}: V_i \cap H_k \rightarrow \mathbb{R}$ verschwindet auf $\partial_1(V_i \cap H_k)$.

Es ist nun $\omega = z_1\omega + \dots + z_r\omega$. Es genügt nun den Satz von Stokes für jeden Summanden $z_i\omega$ zu beweisen. Mit anderen Worten wir können o.B.d.A. annehmen, dass es ein i gibt, so

dass $\text{Träger}(\omega) \subset U_i \cap M$. Es gilt nun, dass

$$\begin{array}{ccccc}
 \text{denn } \text{Träger}(\omega) \subset U_i \cap M & & \text{Satz 10.5} & & \\
 \int_M d\omega \stackrel{\downarrow}{=} \int_{U_i \cap M} d\omega & = & \int_{V_i \cap H_k} \Psi^*(d\omega) \stackrel{\downarrow}{=} \int_{V_i \cap H_k} d(\Psi^*\omega) & & \\
 & \uparrow & \text{Definition des Integrals einer Form} & & \\
 = \int_{\partial_0(V_i \cap H_k)} \Psi^*\omega & \stackrel{\downarrow}{=} & \int_{U_i \cap \partial M} \omega & = & \int_{\partial M} \omega. \\
 \uparrow & & & \uparrow & \\
 \text{nach Lemma 11.2, hierbei verwenden wir,} & & \text{denn } \text{Träger}(\omega) \subset U_i \cap M & & \\
 \text{dass } \Psi^*\omega \text{ auf } \partial_1(V_i \cap H_k) \text{ verschwindet} & & & &
 \end{array}$$

Wir haben damit den Satz von Stokes bewiesen. $\square$

Bemerkung. Wenn man den Beweis vom Satz von Stokes durchliest, und wenn man sich überzeugt hat, dass der Beweis von Lemma 11.2 ganz einfach ist, dann drängt sich der Verdacht auf, dass das ganze eigentlich elementar ist. Insbesondere ist der Beweis vom Satz von Stokes deutlich einfacher als der Beweis vom Gaußschen Integralsatz. Woran liegt das? Der Grund ist die harmlos ausschauende, aber eigentlich geniale Eigenschaft des Differentials, dass es mit glatten Abbildungen kommutiert, d.h. dass in unserem Fall gilt $d(\Psi^*\omega) = \Psi^*(d\omega)$. Diese Eigenschaft erlaubt es uns den Beweis auf den einfachsten Fall, nämlich den Fall, dass der Rand “horizontal” ist, zurückzuführen. Diesen Spezialfall konnten wir problemlos explizit nachweisen.

11.3. Der klassische Satz von Stokes und der Satz von Green. Es sei im Folgenden $F = (F_x, F_y, F_z)$ ein glattes Vektorfeld auf $\mathbb{R}^3$. Auf Seite 79 hatten wir die *Rotation* von F definiert als das Vektorfeld

$$\text{rot } F := \begin{pmatrix} \frac{\partial F_z}{\partial y} - \frac{\partial F_y}{\partial z} \\ -\frac{\partial F_z}{\partial x} + \frac{\partial F_x}{\partial z} \\ \frac{\partial F_y}{\partial x} - \frac{\partial F_x}{\partial y} \end{pmatrix}.$$

Es sei nun $M \subset \mathbb{R}^3$ eine kompakte orientierte Fläche. Wir definieren

$$\int_M \text{rot } F := \begin{array}{ll} \text{Integral der 2-Form } \delta(F), \text{ welche für } P \in M \text{ gegeben ist durch} & \\ T_P M \times T_P M \rightarrow \mathbb{R} & \\ (v_1, v_2) \mapsto \det(\text{rot } F(P) \ v_1 \ v_2). & \end{array}$$

Zudem definieren wir

$$\int_{\partial M} F := \begin{array}{ll} \text{Integral der 1-Form } \omega(F), \text{ welche für } P \in \partial M \text{ gegeben ist durch} & \\ T_P \partial M \rightarrow \mathbb{R} & \\ v \mapsto F(P) \cdot v. & \end{array}$$

Der klassische Satz von Stokes lautet nun wie folgt:

Satz 11.3. (Klassische Satz von Stokes) *Es sei $M \subset \mathbb{R}^3$ eine kompakte orientierte Fläche und F ein glattes Vektorfeld auf $\mathbb{R}^3$. Dann gilt*

$$\int_M \operatorname{rot} F = \int_{\partial M} F.$$

Beweis. Wir werden den klassischen Satz von Stokes mithilfe des allgemeinen Satz von Stokes beweisen. Es sei also $M \subset \mathbb{R}^3$ eine kompakte orientierte Fläche und $F = (F_x, F_y, F_z)$ ein glattes Vektorfeld auf $\mathbb{R}^3$. Wir bezeichnen mit $\omega = \omega(F)$ die 1-Form auf M , welche für $P \in M$ gegeben ist durch

$$\begin{aligned} T_P M &\rightarrow \mathbb{R} \\ v &\mapsto F(P) \cdot v. \end{aligned}$$

Der Satz von Stokes besagt, dass

$$\int_M d\omega = \int_{\partial M} \omega.$$

Der klassische Satz von Stokes folgt nun aus den Definitionen und aus der folgenden Behauptung:

Behauptung. Wir bezeichnen mit η die 2-Form auf M , welche für $P \in M$ gegeben ist durch

$$\begin{aligned} T_P M \times T_P M &\rightarrow \mathbb{R} \\ (v_1, v_2) &\mapsto \det(\operatorname{rot} F(P) v_1 v_2). \end{aligned}$$

Dann gilt $\eta = d\omega$.

Es folgt sofort aus den Definitionen, dass

$$\omega = F_x dx + F_y dy + F_z dz.$$

Wir berechnen nun, dass

$$\begin{aligned} d\omega &= d(F_x dx + F_y dy + F_z dz) \\ &= dF_x \wedge dx + dF_y \wedge dy + dF_z \wedge dz \\ &= \left(\frac{\partial F_x}{\partial x} dx + \frac{\partial F_x}{\partial y} dy + \frac{\partial F_x}{\partial z} dz \right) \wedge dx + \left(\frac{\partial F_y}{\partial x} dx + \frac{\partial F_y}{\partial y} dy + \frac{\partial F_y}{\partial z} dz \right) \wedge dy + \left(\frac{\partial F_z}{\partial x} dx + \frac{\partial F_z}{\partial y} dy + \frac{\partial F_z}{\partial z} dz \right) \wedge dz \\ &\quad \uparrow \end{aligned}$$

Lemma 8.1

$$\begin{aligned} &= \left(\frac{\partial F_y}{\partial x} - \frac{\partial F_x}{\partial y} \right) dx \wedge dy + \left(\frac{\partial F_z}{\partial x} - \frac{\partial F_x}{\partial z} \right) dx \wedge dz + \left(\frac{\partial F_z}{\partial y} - \frac{\partial F_y}{\partial z} \right) dy \wedge dz. \\ &\quad \uparrow \end{aligned}$$

denn $dx_i \wedge dx_j = -dx_j \wedge dx_i$ und $dx_i \wedge dx_i = 0$

Wir wollen nun zeigen, dass $d\omega = \eta$. Die Rechnung ist fast die gleiche wie auf Seite 143, aber der Vollständigkeit halber führen wir diese noch einmal durch.

Wir zeigen $d\omega = \eta$ indem wir zeigen, dass beide Seiten auf alle Vektoren die gleichen Ergebnisse liefern. Es seien also $v_1 = (v_{1x}, v_{1y}, v_{1z})$ und $v_2 = (v_{2x}, v_{2y}, v_{2z})$ in $\mathbb{R}^3$ beliebig.

Dann folgt aus der obigen Rechnung, dass

$$\begin{aligned}
 d\omega\left(\begin{pmatrix} v_{1x} \\ v_{1y} \\ v_{1z} \end{pmatrix}, \begin{pmatrix} v_{2x} \\ v_{2y} \\ v_{2z} \end{pmatrix}\right) &= \left(\frac{\partial F_y}{\partial x} - \frac{\partial F_x}{\partial y}\right)(v_{1x}v_{2y} - v_{1y}v_{2x}) \\
 &\quad + \left(\frac{\partial F_z}{\partial x} - \frac{\partial F_x}{\partial z}\right)(v_{1x}v_{2z} - v_{1z}v_{2x}) \\
 &\quad + \left(\frac{\partial F_z}{\partial y} - \frac{\partial F_y}{\partial z}\right)(v_{1y}v_{2z} - v_{1z}v_{2y}) \\
 &= \det \begin{pmatrix} \frac{\partial F_z}{\partial y} - \frac{\partial F_y}{\partial z} & v_{1x} & v_{2x} \\ -\frac{\partial F_z}{\partial x} + \frac{\partial F_x}{\partial z} & v_{1y} & v_{2y} \\ \frac{\partial F_y}{\partial x} - \frac{\partial F_x}{\partial y} & v_{1z} & v_{2z} \end{pmatrix} = \det(\operatorname{rot} F(P) v_1 v_2) = \eta(v_1, v_2).
 \end{aligned}$$

Wir haben damit die Behauptung bewiesen. $\square$

Der folgende Satz von Green ist ebenfalls ein klassischer Spezial vom Satz von Stokes.

Satz 11.4. (Satz von Green) *Es sei $M \subset \mathbb{R}^2$ eine kompakte 2-dimensionale Untermannigfaltigkeit und es seien $u, v: M \rightarrow \mathbb{R}$ glatte Funktionen. Dann gilt¹⁵²*

$$\int_{\partial M} \underbrace{(v dx + u dy)}_{\substack{\text{glatte 1-Form} \\ \text{auf } \partial M}} = \int_M \left(\frac{\partial u}{\partial x} - \frac{\partial v}{\partial y}\right) dx dy.$$

Bemerkung. Der Satz von Green kann dazu verwendet werden, mit einem sogenannten Planimeter den Flächeninhalt einer Teilmenge von $\mathbb{R}^2$ nur durch Abfahren der Außenlinie zu bestimmen. Details dazu kann man auf

<https://en.wikipedia.org/wiki/Planimeter>

finden.

Beweis. Es sei $M \subset \mathbb{R}^2$ eine kompakte 2-dimensionale Untermannigfaltigkeit und es seien $u, v: M \rightarrow \mathbb{R}$ glatte Funktionen. Dann gilt

$$\begin{array}{ccccccc}
 \int_{\partial M} v dx + u dy & = & \int_M d(v dx + u dy) & = & \int_M \left(\frac{\partial u}{\partial x} - \frac{\partial v}{\partial y}\right) dx \wedge dy & = & \int_M \left(\frac{\partial u}{\partial x} - \frac{\partial v}{\partial y}\right) dx dy. \\
 \uparrow & & \uparrow & & \uparrow & & \\
 \text{Satz von Stokes} & & \text{Definition des Differentials} & & \text{Lemma 9.22} & & \\
 & & \text{und Lemma 8.1} & & & &
 \end{array}$$

$\square$

11.4. Anwendung auf holomorphe Funktionen. Wir erinnern an einige Definitionen aus der Funktionentheorie. Wir identifizieren in diesem Kapitel $\mathbb{C}$ mit $\mathbb{R}^2$. Es sei $U \subset \mathbb{C}$ eine offene Teilmenge. Eine Funktion $f: U \rightarrow \mathbb{C}$ heißt *holomorph*, wenn für alle $z_0 \in U$ der Grenzwert

$$\frac{d}{dz} f(z_0) := f'(z_0) := \lim_{z \rightarrow z_0} \frac{f(z) - f(z_0)}{z - z_0} \in \mathbb{C}$$

¹⁵²Hierbei betrachten wir ∂M als orientierte Untermannigfaltigkeit mit der Konvention, welche wir auf Seite 10.1 eingeführt hatten.

existiert. Wir schreiben nun $f = u + iv: U \rightarrow \mathbb{R}$, wobei $u = \operatorname{Re}(f)$ und $v = \operatorname{Im}(f)$. In der Analysis III hatten wir gezeigt

$$f \text{ ist holomorph} \iff u \text{ und } v \text{ sind glatt mit } \frac{\partial u}{\partial x} = \frac{\partial v}{\partial y} \text{ und } \frac{\partial u}{\partial y} = -\frac{\partial v}{\partial x}.$$

Die Gleichungen auf der rechten Seite hatten wir die *Cauchy-Riemannschen Differentialgleichungen* genannt.

Für eine Kurve $\gamma: [a, b] \rightarrow U$ hatten wir das *Wegintegral von f über γ* definiert als

$$\int_{\gamma} f(z) dz := \int_a^b \underbrace{f(\gamma(t)) \cdot \gamma'(t)}_{\text{Multiplikation in } \mathbb{C}} dt \in \mathbb{C}.$$

Für $z_0 \in \mathbb{C}$ und $r > 0$ mit $\overline{D_r(z_0)} \subset U$ definieren wir zudem

$$\int_{|z-z_0|=r} f(z) dz := \int_{\substack{\gamma: [0, 2\pi] \rightarrow \mathbb{C} \\ t \mapsto z_0 + re^{it}}} f(z) dz.$$

Wir können nun folgenden Satz beweisen.

Satz 11.5. *Es sei $f: U \rightarrow \mathbb{C}$ eine holomorphe Funktion auf einer offenen Teilmenge $U \subset \mathbb{C} = \mathbb{R}^2$. Es sei $\overline{D_r(z_0)}$ eine abgeschlossene Scheibe, welche in U enthalten ist. Dann gilt*

$$\int_{|z-z_0|=r} f(z) dz = 0.$$

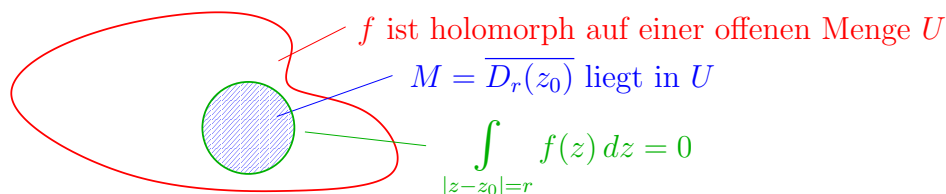


ABBILDUNG 64. Illustration von Satz 11.5.

Wir hatten diesen Satz in Analysis III als Korollar 4.4 formuliert, und dieses mithilfe des Cauchyschen Integralsatzes bewiesen. In der Tat kann man den Beweis von Satz 11.5, analog zum Beweis von Satz 6.12 abwandeln und den Cauchyschen Integralsatz 4.3 für Bilder von Rechtecken in Analysis III beweisen.

Beweis. Es sei also $f: U \rightarrow \mathbb{C}$ eine holomorphe Funktion auf einer offenen Teilmenge $U \subset \mathbb{C} = \mathbb{R}^2$ und es sei $M := \overline{D_r(z_0)}$ eine abgeschlossene Scheibe, welche in U enthalten ist. Wir setzen $u = \operatorname{Re}(f)$ und $v = \operatorname{Im}(f)$, dann ist $f = u + iv: U \rightarrow \mathbb{C}$. Wir betrachten die Kurve

$$\begin{aligned} \gamma: [0, 2\pi] &\rightarrow \mathbb{C} = \mathbb{R}^2 \\ \gamma(t) &= z_0 + re^{it} =: (\alpha(t), \beta(t)). \end{aligned}$$

Dann ist

$$\begin{aligned}
\text{Imaginärteil von } \int_{|z-z_0|=r} f(z) dz &= \text{Imaginärteil von } \int_{t=0}^{t=2\pi} f(\gamma(t)) \cdot \gamma'(t) dt \\
&= \text{Imaginärteil von } \int_{t=0}^{t=2\pi} (u(\gamma(t)) + iv(\gamma(t))) \cdot (\alpha'(t) + i\beta'(t)) dt \\
&= \int_{t=0}^{t=2\pi} (v(\gamma(t)) \cdot \alpha'(t) + u(\gamma(t)) \cdot \beta'(t)) dt \\
&= \int_{t=0}^{t=2\pi} (v(\gamma(t)) dx + u(\gamma(t)) dy)(\gamma'(t)) dt \\
&\quad \uparrow \\
&\quad \text{denn } dx(\gamma'(t)) = \alpha'(t) \text{ und } dy(\gamma'(t)) = \beta'(t) \\
&= \int_{\partial M} v dx + u dy = \int_M \left(\frac{\partial u}{\partial x} - \frac{\partial v}{\partial y} \right) dx dy = 0 \\
&\quad \uparrow \qquad \qquad \qquad \uparrow \qquad \qquad \qquad \uparrow \\
&\quad \text{Definition vom Integral einer 1-Form} \qquad \text{Satz von Green} \qquad \text{Cauchy-Riemann} \\
&\qquad \qquad \qquad \qquad \qquad \qquad \qquad \text{Differentialgleichungen}
\end{aligned}$$

Die Aussage, dass der Realteil des Integrals verschwindet wird ganz ähnlich bewiesen.¹⁵³

□

¹⁵³Eine andere Möglichkeit zu beweisen, dass der Realteil verschwindet, ist dass man die gleiche Rechnung auf die holomorphe Funktion $z \mapsto if(z)$ anwendet.

12. LÄNGE VON WEGEN UND GEODÄTEN

Wir wenden uns nun, nach der langen Diskussion des Satzes von Stokes, einem anschaulichen Thema zu, nämlich der Bestimmung von Längen von Wegen und der Bestimmung des kürzesten Weges zwischen zwei Punkten auf einer Untermannigfaltigkeit.

Definition. Ein *Weg* in einer Untermannigfaltigkeit M ist eine abschnittsweise¹⁵⁴ stetig differenzierbare Abbildung $\gamma: [a, b] \rightarrow M$.¹⁵⁵ Hierbei bedeutet “abschnittsweise stetig differenzierbar”, dass es eine Zerlegung $a = s_0 < s_1 < \dots < s_m = b$ gibt, so dass für $i = 0, \dots, m-1$ die Abbildung $\gamma|_{[s_i, s_{i+1}]}$ stetig differenzierbar ist. Wir definieren die *Länge* von γ ¹⁵⁶ als

$$\ell(\gamma) := \sum_{i=0}^{m-1} \int_{t=s_i}^{t=s_{i+1}} \|\gamma'(t)\| dt.$$

Definition. Es sei P und Q zwei Punkte auf einer Untermannigfaltigkeit M . Wenn P und Q in der gleichen Komponente von M liegen, dann definieren wir

$$d_M(P, Q) := \inf\{\ell(\gamma) \mid \gamma \text{ ist ein Weg in } M \text{ von } P \text{ nach } Q\}.$$

Ansonsten definieren wir $d_M(P, Q) := \infty$. Wir bezeichnen d_M als die *Wegmetrik* auf M .

Bemerkung. Der Name *Wegmetrik* legt nahe, dass es sich hierbei in der Tat um eine Metrik auf M handelt. Wenn M zusammenhängend¹⁵⁷ ist, dann ist dies in der Tat der Fall. Wir werden diese Aussage nicht verwenden, und daher auch nicht beweisen.¹⁵⁸

Lemma 12.1. Für $M = \mathbb{R}^n$ entspricht die Wegmetrik gerade der euklidischen Metrik.

Beweis. Um die Notation etwas zu vereinfachen beweisen wir das Lemma für den Spezialfall $n = 2$. Es seien also $P, Q \in \mathbb{R}^2$. Translation und Anwendung einer orthogonalen Matrix ändert die Länge von Vektoren und Wegen nicht. Wir können also P in den Ursprung verschieben, und Q auf die positive y -Achse drehen. Mit anderen Worten, wir können annehmen, dass $P = (0, 0)$ und $Q = (0, y)$ für ein $y \geq 0$.

Wir wollen nun zeigen, dass $d_{\mathbb{R}^2}(P, Q) = \|P - Q\| = y$. Der Weg $\gamma: [0, y] \rightarrow \mathbb{R}^2$ mit $\gamma(t) = (0, t)$ ist offensichtlich ein Weg, welcher P und Q verbindet mit $\ell(\gamma) = y$. Wir sehen also, dass $d_{\mathbb{R}^2}(P, Q) \leq \|P - Q\|$.

¹⁵⁴In der Literatur sagt man oft auch “stückweise stetig differenzierbar” anstatt “abschnittsweise stetig differenzierbar”.

¹⁵⁵Der einzige Unterschied zu einer Kurve ist, dass Kurven stetig differenzierbar sind, während Wege nur abschnittsweise stetig differenzierbar sind.

¹⁵⁶Die Definition der Länge eines Weges entspricht der Definition aus der Analysis III Vorlesung und ist äquivalent zur Definition aus der Analysis II Vorlesung.

¹⁵⁷In Satz 5.2 hatten wir bewiesen, dass eine zusammenhängende Untermannigfaltigkeit auch wegzusammenhängend ist, d.h. alle Punkte liegen in der gleichen Komponente.

¹⁵⁸Welche von den Axiomen einer Metrik bedürfen eines Beweises und welche sind offensichtlich?

Wir müssen noch zeigen, dass $d_{\mathbb{R}^2}(P, Q) \geq \|P - Q\|$. Es sei nun

$$\begin{aligned}\rho: [a, b] &\rightarrow \mathbb{R}^2 \\ t &\mapsto (\rho_x(t), \rho_y(t))\end{aligned}$$

ein Weg von $(0, 0)$ nach $(0, y)$. Wir wollen beweisen, dass $\ell(\rho) \geq y$. Wir betrachten dazu zuerst den Spezialfall, dass ρ stetig differenzierbar ist. Dann gilt

$$\begin{aligned}\ell(\rho) &= \int_a^b \|\rho'(t)\| dt = \int_a^b \sqrt{\rho'_x(t)^2 + \rho'_y(t)^2} dt \stackrel{\text{Gleichheit gilt nur wenn } \rho'_x \equiv 0}{\geq} \int_a^b \sqrt{\rho'_y(t)^2} dt \\ &= \int_a^b |\rho'_y(t)| dt \geq \left| \int_a^b \rho'_y(t) dt \right| = |\rho_y(b) - \rho_y(a)| = y.\end{aligned}$$

Der allgemeine Fall, dass ρ abschnittsweise stetig differenzierbar ist, wird ganz analog behandelt. $\square$

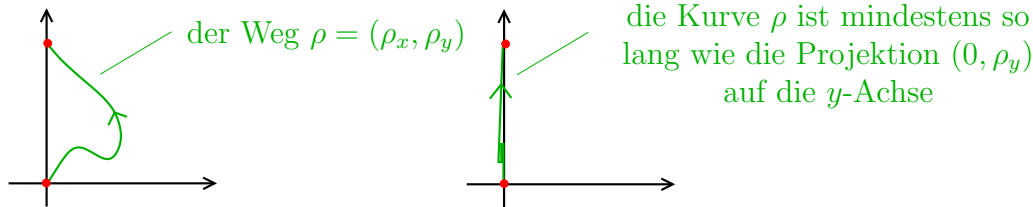


ABBILDUNG 65. Skizze zum Beweis von Lemma 12.1.

Im Beweis von Lemma 12.1 hatten wir auch schon implizit folgendes Lemma bewiesen.

Lemma 12.2. Jeder Weg $\gamma: [a, b] \rightarrow \mathbb{R}^n$ von P nach Q mit $\ell(\gamma) = \|P - Q\|$ verläuft auf der euklidischen Gerade durch P und Q .

Beweis. Wir verwenden die Notation vom vorherigen Beweis. Wenn $\delta = (\rho_x, \rho_y)$ ein Weg von $(0, 0)$ nach $(0, y)$ ist, welcher nicht auf der euklidischen Gerade durch $(0, 0)$ und $(0, y)$ verläuft, dann ist $\rho'_x(t) \neq 0$ für ein t . Dann zeigt die Rechnung aus dem vorherigen Beweis, dass $\ell(\rho) > y = \|P - Q\|$. $\square$

Definition. Es sei M eine Untermannigfaltigkeit und es sei $\gamma: [a, b] \rightarrow M$ ein Weg.

- (1) Wir sagen γ ist *minimal in M* , wenn

$$d_M(\gamma(a), \gamma(b)) = \ell(\gamma).$$

- (2) Wir nennen γ eine *Geodäte*, wenn folgende drei Bedingungen erfüllt sind:

- (a) γ ist stetig differenzierbar,
- (b) $\|\gamma'(t)\| = 1$ für alle $t \in [a, b]$ und
- (c) γ ist minimal.

- (3) Es seien P und Q zwei Punkte auf M . Eine *Geodäte von P nach Q* ist eine Geodäte $\gamma: [0, a] \rightarrow M$ mit $\gamma(0) = P$ und $\gamma(a) = Q$.
- (4) Es sei $t \in [a, b]$. Wir sagen der Weg γ ist *lokal minimal in t* , wenn es $c < d$ mit $a \leq c \leq t \leq d \leq b$ und eine offene Umgebung U von $\gamma(t)$ in M gibt, so dass $\gamma|_{[c, d]}$ in U minimal ist.
- (5) Wir nennen γ eine *lokale Geodäte*, wenn folgende drei Bedingungen erfüllt sind:
 - (a) γ ist stetig differenzierbar,
 - (b) $\|\gamma'(t)\| = 1$ für alle $t \in [a, b]$ und
 - (c) γ ist lokal minimal in allen $t \in [a, b]$.

Bemerkung. Ein minimaler Weg $\gamma: [a, b] \rightarrow M$ ist offensichtlich¹⁵⁹ auch lokal minimal in allen t . Eine Geodäte ist insbesondere auch eine lokale Geodäte. Die Umkehrung dieser Aussage gilt jedoch nicht. Beispielsweise werden wir im nächsten Kapitel, die anschaulich ziemlich offensichtliche Aussage, beweisen, dass der Weg

$$\begin{aligned} \gamma: [0, 2\pi] &\rightarrow M = S^1 \\ t &\mapsto (\cos(t), \sin(t)) \end{aligned}$$

eine lokale Geodäte ist. Genauer gesagt, wir werden beweisen, dass für alle Intervalle $[a, b]$ in $[0, 2\pi]$ mit $0 < b - a \leq \pi$ die Einschränkung $\gamma|_{[a, b]}$ minimal ist. Andererseits ist γ keine Geodäte, denn die Einschränkung von γ auf das Intervall $[0, 2\pi]$ ist nicht die kürzeste Verbindung zwischen $\gamma(0) = \gamma(2\pi) = (1, 0)$, denn der konstante Weg ist natürlich kürzer.

In Übungsblatt 8 werden wir folgendes einfache Lemma beweisen.

Lemma 12.3.

- (1) Die Einschränkung von einem minimalen Weg auf ein Teilintervall ist wiederum minimal.
- (2) Die Einschränkung von einer Geodäte auf ein Teilintervall ist wiederum eine Geodäte.

Der folgende Satz besagt, dass in $\mathbb{R}^n$ die Geodäten auf den üblichen Strecken verlaufen.

Satz 12.4. Es seien $P, Q \in \mathbb{R}^n$. Dann ist

$$\begin{aligned} \gamma: [0, \|P - Q\|] &\rightarrow \mathbb{R}^n \\ t &\mapsto \frac{t}{\|P - Q\|}P + \left(1 - \frac{t}{\|P - Q\|}\right)Q \end{aligned}$$

die einzige Geodäte von P nach Q .

Beweis. Ganz analog zum Beweis von Lemma 12.1 genügt es die Aussage zu beweisen für $n = 2$ und für $P = (0, 0)$ und $Q = (0, y)$ mit $y \geq 0$. Die Kurve

$$\begin{aligned} \gamma: [0, y] &\rightarrow \mathbb{R}^2 \\ t &\mapsto (0, t) \end{aligned}$$

ist offensichtlich glatt und offensichtlich ist $\|\gamma'(t)\| = 1$ für alle $t \in (0, y)$. Es folgt zudem aus Lemma 12.1, dass γ minimal ist.

¹⁵⁹Warum ist das offensichtlich?

Es sei nun $\delta: [0, y] \rightarrow \mathbb{R}^2$ eine weitere Geodäte von P nach Q . Aus Lemma 12.2 folgt, dass sich δ auf der y -Achse bewegt, d.h. es ist $\delta(t) = (0, \beta(t))$ für eine glatte Funktion $\beta: [0, y] \rightarrow \mathbb{R}$. Nachdem $\|\delta'(t)\| = 1$ für alle $t \in (a, b)$ ist nun $\beta'(t)$ entweder konstant 1 oder konstant -1 . Da $\beta(0) = 0$ und $\beta(y) = y > 0$ ist $\beta'(t) = 1$ für alle t . Also ist $\delta(t) = (0, \beta(t)) = (0, t) = \gamma(t)$ für alle $t \in [a, b]$. $\square$

Bemerkung. Wir hatten insbesondere gerade gesehen, dass es für zwei Punkte P und Q in $\mathbb{R}^n$ genau eine Geodäte von P nach Q gibt. Im Allgemeinen sind Geodäten jedoch nicht eindeutig bestimmt. Wir betrachten dazu wiederum die Untermannigfaltigkeit S^1 . Wenn $t \mapsto (x(t), y(t))$ eine Geodäte von $P = (-1, 0)$ nach $Q = (0, 1)$ ist, dann ist auch die Spiegelung an der x -Achse, d.h. die Abbildung $t \mapsto (x(t), -y(t))$ eine Geodäte von $P = (-1, 0)$ nach $Q = (0, 1)$.

Definition. Eine Untermannigfaltigkeit M heißt *geodätisch*, wenn es zu allen $P, Q \in M$ eine Geodäte γ gibt, welche P und Q verbindet.

Wir haben gerade gesehen, dass der euklidische Raum $\mathbb{R}^n$ geodätisch ist. Allerdings ist nicht jede Untermannigfaltigkeit geodätisch. Man kann beispielsweise leicht zeigen, dass die Untermannigfaltigkeit $\mathbb{R}^2 \setminus \{0, 0\}$ nicht geodätisch ist. Der Gedanke hierbei ist, dass man die beiden Punkte $P = (-1, 0)$ und $Q = (1, 0)$ auf der x -Achse durch Wege verbinden kann, deren Länge gegen 2 konvergiert, aber es gibt keinen Weg in $\mathbb{R}^2 \setminus \{(0, 0)\}$ der Länge 2, welcher die beiden Punkte verbindet. Andererseits kann man beweisen, dass jede

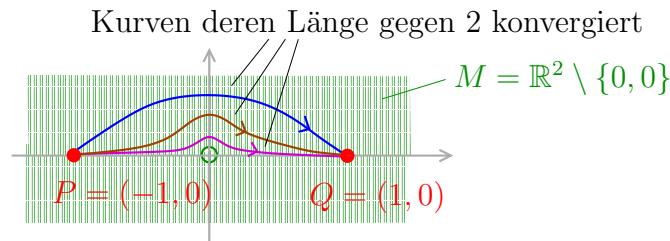


ABBILDUNG 66.

kompakte Untermannigfaltigkeit geodätisch ist. Wir werden diese Aussage aus Zeitgründen jedoch nicht beweisen. Aber wir werden zumindest im übernächsten Kapitel sehen, dass die kompakte Untermannigfaltigkeit S^2 geodätisch ist.

13. RIEMANNSCHE UNTERMANNIGFALTIGKEITEN

In Abbildung 67 sehen wir die Route von einem Flug von Houston nach München. Genauer gesagt, wir sehen die Route auf einer Karte der Erdoberfläche. Der Pilot wählt natürlich auf der Sphäre die kürzeste Route, aber wie wir in Abbildung 67 sehen, entspricht die kürzeste Route auf der Karte nicht der “direkten” Route in $\mathbb{R}^2$. Wir wollen in diesem Ka-

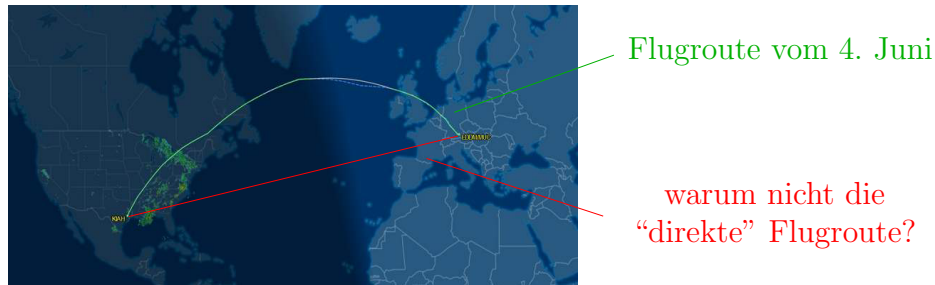


ABBILDUNG 67.

pitel unter Anderem verstehen, was die Geodäten auf einer Sphäre und allgemeiner auf Untermannigfaltigkeiten sind.

Wie viele andere Problemstellungen auf Untermannigfaltigkeiten wollen wir die Bestimmung von Geodäten auf Untermannigfaltigkeiten mithilfe von Karten auf Problemstellungen auf offenen Teilmengen von $\mathbb{R}^k$ zurückführen.

Genauer gesagt, es sei $\gamma: [a, b] \rightarrow M$ ein Weg auf einer k -dimensionale Untermannigfaltigkeit. Wir wollen überprüfen, ob γ in einem Punkt $t \in [a, b]$ lokal minimal ist. Wir wählen also eine Karte $\Phi: U \rightarrow V$ um $\gamma(t)$. Dann ist $\Phi \circ \gamma: [a, b] \rightarrow V \cap E_k$ ein Weg in einer offenen Teilmenge von $E_k = \mathbb{R}^k$.

Das Problem ist nun allerdings, wie wir schon in Abbildung 67 gesehen haben, dass die Abbildung $\Phi: U \cap M \rightarrow V \cap E_k$ im Allgemeinen nicht längenerhaltend ist. Wenn wir also wirklich die Länge von Kurven auf $U \cap M$ mithilfe von Längen von Kurven auf $V \cap E_k$ studieren wollen, dann müssen wir auf $V \cap E_k$ einen Längenbegriff verwenden, welcher sich vom üblichen Längenbegriff unterscheidet.

Diese Vorbemerkung führt uns nun direkt zu dem Begriff der Riemannschen Untermannigfaltigkeit.

13.1. Definition von Riemannschen Untermannigfaltigkeiten.

Definition.

- (1) Eine *Bilinearform* auf einem reellen Vektorraum V ist eine Abbildung

$$\begin{aligned} g: V \times V &\rightarrow \mathbb{R} \\ (v_1, v_2) &\mapsto g(v_1, v_2), \end{aligned}$$

welche in beiden Einträgen linear ist. Wir sagen die Bilinearform ist *symmetrisch*, wenn $g(v_2, v_1) = g(v_1, v_2)$ für alle $v_1, v_2 \in V$, und wir sagen die Bilinearform ist *positiv definit*, wenn für alle $v \neq 0 \in V$ gilt, dass $g(v, v) > 0$.

- (2) Es sei h eine Bilinearform auf einem reellen Vektorraum W und es sei $\varphi: V \rightarrow W$ ein Homomorphismus. Wir bezeichnen mit φ^*h die Bilinearform auf V , welche gegeben ist durch¹⁶⁰¹⁶¹

$$(\varphi^*h)(v_1, v_2) := h(\varphi(v_1), \varphi(v_2)).$$

- (3) Es seien g, h Bilinearformen auf reellen Vektorräumen V und W . Eine *Isometrie* von (V, g) nach (W, h) ist ein Isomorphismus $\varphi: V \rightarrow W$, so dass $\varphi^*h = g$, oder mit anderen Worten, so dass

$$h(\varphi(v_1), \varphi(v_2)) = g(v_1, v_2)$$

für alle $v_1, v_2 \in V$.

Beispiel. Auf dem Vektorraum $\mathbb{R}^n$ ist das übliche Skalarprodukt

$$\begin{aligned} \mathbb{R}^n \times \mathbb{R}^n &\rightarrow \mathbb{R} \\ ((v_1, \dots, v_n), (w_1, \dots, w_n)) &\mapsto v \cdot w = \langle v, w \rangle = \sum_{i=1}^n v_i w_i \end{aligned}$$

eine symmetrische, positiv definite Bilinearform. In diesem Kapitel werden wir nur die Notation $\langle v, w \rangle$ verwenden um Verwechslungen in der Notation zu vermeiden.

Bemerkung. Auf dem Vektorraum $\mathbb{R}^n$ entsprechen die (positiv definiten) Bilinearformen gerade den (positiv definiten) symmetrischen reellen $n \times n$ -Matrizen. Genauer gesagt haben wir zueinander inverse Bijektionen

$$\left\{ \begin{array}{c} \text{(positiv definite)} \\ \text{Bilinearformen auf } \mathbb{R}^n \end{array} \right\} \leftrightarrow \left\{ \begin{array}{c} \text{(positiv definite) symmetrische} \\ \text{reelle } n \times n\text{-Matrizen} \end{array} \right\}$$

welche gegeben sind durch

$$g \mapsto (g(e_i, e_j))_{i,j=1,\dots,n}$$

und

$$\left(\begin{array}{c} \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R} \\ (v, w) \mapsto v^T A w \end{array} \right) \leftarrow A.$$

Das übliche Skalarprodukt auf $\mathbb{R}^n$ entspricht hierbei der Identitätsmatrix.

Definition.

- (1) Eine *prä-Riemannsche Struktur* auf einer Untermannigfaltigkeit M ist eine Abbildung g , welche jedem Punkt $P \in M$ eine positiv definite symmetrische Bilinearform g_P auf dem Vektorraum $T_P M$ zuordnet.

¹⁶⁰Die Notation legt natürlich schon die Tatsache nahe, dass die “üblichen Regeln” für die induzierten Abbildungen gelten. Insbesondere gilt für lineare Abbildungen $\psi: U \rightarrow V$ und $\varphi: V \rightarrow W$ und eine Bilinearform h auf W , dass $(\varphi \circ \psi)^*h = \psi^*(\varphi^*h)$.

¹⁶¹Wenn h positiv definit ist, dann ist φ^*h genau dann positiv definit, wenn φ injektiv ist.

- (2) Eine Abbildung $\varphi: M \rightarrow N$ zwischen zwei Untermannigfaltigkeiten ist ein *lokaler Diffeomorphismus*, wenn es zu jedem $P \in M$ offene Umgebungen U von P und V von $\varphi(P)$ gibt, so dass $\varphi: U \rightarrow V$ ein Diffeomorphismus ist.
- (3) Es sei $\varphi: M \rightarrow N$ ein lokaler Diffeomorphismus und es sei g eine *prä-Riemannsche Struktur* auf N . Für $P \in M$ betrachten wir dann die Bilinearform

$$\begin{aligned} T_P M \times T_P M &\rightarrow \mathbb{R} \\ (v, w) &\mapsto g_{\varphi(P)}(\varphi_* v, \varphi_* w). \end{aligned}$$

Wir erhalten also eine prä-Riemannsche Struktur¹⁶² auf M , welche wir mit φ^*g bezeichnen. Wir nennen φ^*g manchmal den “Pullback von g ”.

Wir wollen nun definieren, was es heißen soll, dass eine prä-Riemannsche Struktur glatt ist. Die Definition ist ganz ähnlich der Definition von einer glatten Differentialform, welche wir auf Seite 115 gegeben hatten.

Definition.

- (1) Es sei W eine offene Teilmenge von $\mathbb{R}^k$ oder von H_k . Wir sagen eine Abbildung

$$\begin{aligned} W &\rightarrow M(k \times k, \mathbb{R}) \\ P &\mapsto A(P) = (a_{ij}(P))_{i,j=1,\dots,k} \end{aligned}$$

ist *glatt*, wenn alle Einträge a_{ij} glatt in P sind. Wir sagen eine prä-Riemannsche Struktur g auf W ist glatt, wenn die entsprechende Abbildung

$$W \rightarrow M(k \times k, \mathbb{R})$$

glatt ist.

- (2) Es sei g eine prä-Riemannsche Struktur auf einer k -dimensionalen Untermannigfaltigkeit M .
- (a) Es sei $\Phi: U \rightarrow V$ eine Karte für M . Wir sagen g ist *glatt bezüglich der Karte Φ* , falls die prä-Riemannsche Struktur $(\Phi^{-1})^*g$ auf $V \cap E_k$ beziehungsweise $V \cap H_k$ glatt ist.
- (b) Wir sagen g ist *glatt*, wenn es einen Atlas von M gibt, so dass g bezüglich allen Karten des Atlanten glatt ist.
- (3) Eine glatte prä-Riemannsche Struktur wird *Riemannsche Struktur* genannt.
- (4) Eine *Riemannsche Untermannigfaltigkeit* ist ein Paar (M, g) bestehend aus einer Untermannigfaltigkeit M und einer Riemannschen Struktur g auf M .

Der folgende Satz wird ganz ähnlich wie Satz 8.3 bewiesen.

Satz 13.1. *Es sei g eine prä-Riemannsche Struktur auf einer Untermannigfaltigkeit M . Dann sind die folgenden Aussagen äquivalent:*

- (1) *die prä-Riemannsche Struktur ist eine Riemannsche Struktur,*

¹⁶²Nachdem φ ein lokaler Diffeomorphismus ist, ist insbesondere für jedes $P \in M$ die induzierte Abbildung $\varphi_*: T_P M \rightarrow T_{\varphi(P)} N$ ein Isomorphismus. Also folgt aus Fußnote 161, dass φ^*g in der Tat an jedem Punkt wiederum positiv definit ist.

- (2) die prä-Riemannsche Struktur ist glatt bezüglich jeder Karte von M ,
 (3) für alle glatten Tangentialvektorfelder F, G auf M ist die Funktion

$$\begin{aligned} M &\rightarrow \mathbb{R} \\ P &\mapsto g_P(F(P), G(P)) \end{aligned}$$

glatt.

Der folgende Satz wird zudem ganz ähnlich wie Satz 8.6 bewiesen.

Satz 13.2. *Es sei $\varphi: M \rightarrow N$ ein lokaler Diffeomorphismus zwischen Untermannigfaltigkeiten und es sei g eine Riemannsche Struktur auf N . Dann ist φ^*g eine Riemannsche Struktur auf M .*

Das folgende Lemma gibt eines der wichtigsten Beispiele einer Riemannschen Struktur.

Lemma 13.3. *Es sei M eine Untermannigfaltigkeit von $\mathbb{R}^n$. Für jedes $P \in M$ betrachten wir die positiv definite symmetrische Bilinearform*

$$\begin{aligned} g_P: T_P M \times T_P M &\rightarrow \mathbb{R} \\ (v_1, v_2) &\mapsto g_P(v_1, v_2) := \underbrace{\langle v_1, v_2 \rangle}_{\substack{\uparrow \\ \text{Skalarprodukt in } \mathbb{R}^n}}. \end{aligned}$$

Dies ist eine Riemannsche Struktur auf M .

Wir verwenden im weiteren Verlauf der Vorlesung die Konvention, dass wir eine Untermannigfaltigkeit als Riemannsche Untermannigfaltigkeit bezüglich der Riemannschen Struktur aus Lemma 13.3 auffassen, außer wir sagen explizit etwas anderes.

Beweis. Es sei M eine k -dimensionale Untermannigfaltigkeit von $\mathbb{R}^n$. Wir bezeichnen mit g die prä-Riemannsche Struktur, welche gegeben ist dadurch, dass wir das Skalarprodukt von $\mathbb{R}^n$ auf die Tangentialräume $T_P M$ einschränken. Es sei $\Phi: U \rightarrow V$ eine beliebige Karte für M . Wir bezeichnen mit $\Psi: V \rightarrow U$ die Umkehrabbildung. Wir müssen zeigen, dass Ψ^*g eine glatte prä-Riemannsche Struktur auf $V \cap E_k$ ist. Wir identifizieren dazu wiederum die Bilinearformen auf $\mathbb{R}^k$ mit $k \times k$ -Matrizen. Eine Abbildung zu den $k \times k$ -Matrizen ist per Definition glatt, wenn alle Einträge glatt sind.

In unserem Fall gilt für alle $i, j \in \{1, \dots, k\}$ und alle $Q \in V \cap E_k$, dass

$$\begin{array}{ccccc} & & \downarrow & & \\ ij\text{-Eintrag von } (\Psi^*g)(Q) & = & (\Psi^*g)_Q(e_i, e_j) & & \\ & = & g_{\Psi(Q)}(D_Q \Psi(e_i), D_Q \Psi(e_j)) & = & \langle D_Q \Psi(e_i), D_Q \Psi(e_j) \rangle. \\ & \uparrow & & \uparrow & \\ & \text{Definition von } \Psi^*g & & \text{Definition von } g & \end{array}$$

Diese Funktion ist offensichtlich¹⁶³ eine glatte Abbildung in Q . □

¹⁶³Warum ist das offensichtlich?

Beispiel. Viele interessante Beispiele von Riemannschen Untermannigfaltigkeiten erhalten wir schon dadurch, dass wir als Untermannigfaltigkeit M eine offene Teilmenge von $\mathbb{R}^n$ wählen, aber wir diese mit einer anderen als der üblichen Riemannschen Struktur versehen.

- (1) Es sei $U \subset \mathbb{R}^2$ eine offene Teilmenge und es sei

$$\begin{aligned} A: U &\rightarrow \text{symmetrische positiv definite } 2 \times 2\text{-Matrizen} \\ P &\mapsto A(P) \end{aligned}$$

eine glatte Abbildung. Für $P \in U$ betrachten wir dann auf $T_P U = \mathbb{R}^2$ die positiv definite symmetrische Bilinearform

$$\begin{aligned} \mathbb{R}^2 \times \mathbb{R}^2 &\rightarrow \mathbb{R} \\ (v, w) &\mapsto v^T A(P) w. \end{aligned}$$

Dies ist dann eine Riemannsche Struktur auf U , welche wir mit $g(A)$ bezeichnen.

- (2) Es sei U wiederum eine offene Teilmenge von $\mathbb{R}^2$ und es sei zudem $\rho: U \rightarrow \mathbb{R}_{>0}$ eine glatte Funktion. Für $P \in U$ betrachten wir auf $T_P U = \mathbb{R}^2$ die positiv definite symmetrische Bilinearform

$$\begin{aligned} \mathbb{R}^2 \times \mathbb{R}^2 &\rightarrow \mathbb{R} \\ (v, w) &\mapsto \rho(P) \cdot \langle v, w \rangle. \end{aligned}$$

Dies ist eine Riemannsche Struktur auf U , welche wir mit $g(\rho)$ bezeichnen.¹⁶⁴

Definition. Es sei (M, g) eine Riemannsche Untermannigfaltigkeit.

- (1) Für $P \in M$ und $v \in T_P M$ bezeichnen wir

$$\|v\|_g := \sqrt{g_P(v, v)}$$

als die *Länge von v bezüglich der Riemannschen Struktur g* .

- (2) Es sei nun $\gamma: [a, b] \rightarrow M$ ein Weg, d.h. eine abschnittsweise stetig differenzierbare Abbildung. Wir wählen eine Zerlegung $a = s_0 < s_1 < \dots < s_m = b$ gibt, so dass für $i = 0, \dots, m-1$ die Abbildung $\gamma|_{[s_i, s_{i+1}]}$ stetig differenzierbar ist. Wir definieren die *Länge von γ bezüglich g* als¹⁶⁵

$$\ell_g(\gamma) := \sum_{i=0}^{m-1} \int_{t=s_i}^{t=s_{i+1}} \|\gamma'(t)\|_g dt.$$

- (3) Für $P, Q \in M$ definieren wir zudem¹⁶⁶

$$d_M^g(P, Q) := \inf\{\ell_g(\gamma) \mid \gamma \text{ ist ein Weg von } P \text{ nach } Q\}.$$

¹⁶⁴Dieses Beispiel ist nur ein Spezialfall vom vorherigen Beispiel, mit $A(P) = \begin{pmatrix} \rho(P) & 0 \\ 0 & \rho(P) \end{pmatrix}$.

¹⁶⁵Für jedes i ist die Abbildung $[s_i, s_{i+1}] \rightarrow \mathbb{R}, t \mapsto \|\gamma'(t)\|_g$ stetig, also existiert das Integral. Der Beweis der Tatsache, dass die Abbildung stetig ist, verbleibt eine freiwillige (nicht ganz einfache) Übungsaufgabe.

¹⁶⁶Hierbei verwenden wir natürlich wieder die Konvention, dass $d_M^g(P, Q) = \infty$, wenn P und Q nicht in der gleichen Komponente von M liegen.

- (4) Eine *Geodäte* in (M, g) ist ein Weg $\gamma: [a, b] \rightarrow M$, welcher folgende drei Bedingungen erfüllt:
- (a) γ ist stetig differenzierbar,
 - (b) $\|\gamma'(t)\|_g = 1$ für alle $t \in [a, b]$ und
 - (c) $d_M^g(\gamma(a), \gamma(b)) = \ell_g(\gamma)$.

Beispiele.

- (1) Wir betrachten $U = \mathbb{R} \times (0, \pi)$ und für $(x, y) \in Q$ betrachten wir

$$A(x, y) := \begin{pmatrix} \sin^2 y & 0 \\ 0 & 1 \end{pmatrix}.$$

Dann gilt für jeden horizontalen Weg

$$\begin{aligned} \gamma_y: [a, b] &\rightarrow Q \\ t &\mapsto (t, y), \end{aligned}$$

dass die Länge von γ_y in der Riemannschen Untermannigfaltigkeit $(U, g(A))$ gegeben ist durch

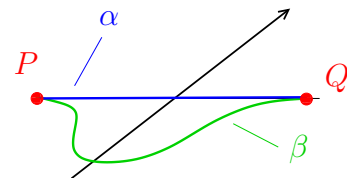
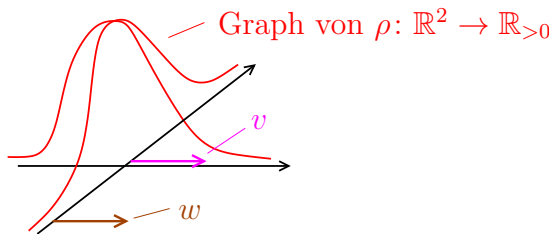
$$\ell_{g(A)}(\gamma_y) = \int_{t=a}^{t=b} \sqrt{g(A)_{\gamma_y(t)}(\gamma'_y(t), \gamma'_y(t))} dt = \int_{t=a}^{t=b} \sqrt{(1, 0) \begin{pmatrix} \sin^2 y & 0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 \\ 0 \end{pmatrix}} dt = \sin(y)(b-a).$$

Wir sehen also, dass in der Riemannschen Untermannigfaltigkeit $(U, g(A))$ die Länge des horizontalen Weges von der y -Koordinate abhängt.

- (2) In Abbildung 68 betrachten wir eine glatte Funktion $\rho: \mathbb{R}^2 \rightarrow \mathbb{R}_{>0}$, welche am Ursprung ein relatives Maximum besitzt. Für einen Tangentialvektor $v \in T_P \mathbb{R}^2$ gilt dann, dass

Länge von v bezüglich $g(\rho) = \sqrt{\rho(P)} \langle v, v \rangle = \sqrt{\rho(P)} \cdot \text{euklidische Länge von } v$.

Um von einem Punkt auf der x -Achse P zum gegenüberliegenden Punkt Q zu kommen, ist es bezüglich der Riemannschen Metrik $g(\rho)$ besser, den Ursprung zu umgehen.



euklidische Länge von v = euklidische Länge von w
 Länge von v in $(\mathbb{R}^2, g(\rho)) >$ Länge von w in $(\mathbb{R}^2, g(\rho))$

der Weg β ist bezüglich $g(\rho)$
 kürzer als der Weg α

ABBILDUNG 68. Länge von Wegen in $(\mathbb{R}^2, g(\rho))$.

- (3) Das vorherige Beispiel ist durchaus realitätsnah. Es sei $B \subset S^2$ die Teilmenge, welche dem Freistaat Bayern entspricht und es sei g die übliche Riemannsche Struktur auf U , welche durch das Skalarprodukt in $\mathbb{R}^3$ gegeben ist. Dann gilt für $P, Q \in B$, dass

$$\text{Distanz zwischen } P \text{ und } Q = d_B^g(P, Q).$$

Es sei nun $\rho: B \rightarrow \mathbb{R}_{>0}$ die Funktion, welche an jedem Punkt gerade die Verkehrsdichte mißt. Dann gilt, zugegebenermaßen grob vereinfacht, dass

$$\text{Fahrzeit von } P \text{ nach } Q = d_{g(\rho)}(P, Q).$$

Die Funktion ρ im vorherigen Beispiel modelliert die Verkehrsdichte im Großraum München. Um schnellst möglich vom Westen Münchens in den Osten zu kommen ist es am besten die Innenstadt weiträumig zu umfahren.

Das nächste Beispiel formulieren wir als Lemma, weil wir später bei zwei Gelegenheiten darauf zurückgreifen wollen. Der Beweis der ersten Aussage ist eine Abwandlung des Beweises von Lemma 12.2 und erfolgt in Übungsblatt 9.

Lemma 13.4. *Es seien $-\infty \leq a < b \leq \infty$ und $-\infty \leq c < d \leq \infty$. Es sei $Q = (a, b) \times (c, d)$ der dazugehörige offene Quader in $\mathbb{R}^2$. Darüber hinaus seien $p, q: (c, d) \rightarrow \mathbb{R}_{>0}$ positive glatte Funktionen. Wir betrachten die glatte Abbildung*

$$\begin{aligned} A: Q &\rightarrow \text{symmetrische positiv definite } 2 \times 2\text{-Matrizen} \\ (x, y) &\mapsto A(x, y) := \begin{pmatrix} p(y) & 0 \\ 0 & q(y) \end{pmatrix} \end{aligned}$$

und die zugehörige Riemannsche Struktur $g(A)$.

- (1) Für zwei Punkte $S = (x, s)$ und $T = (x, t)$ in Q mit der gleichen x -Koordinate verläuft jede Geodäte in der Riemannschen Untermannigfaltigkeit $(Q, g(A))$ auf der vertikalen Gerade $\{(x, y) \mid y \in (c, d)\}$.
- (2) Im Allgemeinen verlaufen die Geodäten in der Riemannschen Untermannigfaltigkeit $(Q, g(A))$ nicht mehr auf euklidischen Geraden.

Definition. Es seien (M, g) und (N, h) Riemannsche Untermannigfaltigkeiten. Eine Abbildung $\varphi: M \rightarrow N$ heißt (lokale) *Isometrie*, wenn φ ein (lokaler) Diffeomorphismus ist und wenn $\varphi^*h = g$.¹⁶⁷

Beispiele.

- (1) Für jede orthogonale Matrix $A \in O(3)$ ist die Abbildung

$$\begin{aligned} S^2 &\rightarrow S^2 \\ P &\mapsto A \cdot P \end{aligned}$$

eine Isometrie. In der Tat, denn die Riemannsche Struktur ist durch das Skalarprodukt definiert, und orthogonale Matrizen erhalten das Skalarprodukt.

¹⁶⁷Mit anderen Worten, für alle $P \in M$ und $v, w \in T_P M$ soll gelten, dass $h(\varphi_* v, \varphi_* w) = g(v, w)$.

(2) Wir betrachten die Abbildung

$$\begin{aligned}\varphi: (-2, 2) \times \mathbb{R} &\rightarrow (-2, 2) \times S^1 \\ (s, t) &\mapsto \begin{pmatrix} s \\ \cos t \\ \sin t \end{pmatrix}.\end{aligned}$$

Anschaulich gesprochen “rollt” die Abbildung den Streifen $(-2, 2) \times \mathbb{R}$ wie einen Teppich zu einem Zylinder zusammen. Die Abbildung φ ist ein lokaler Diffeomor-

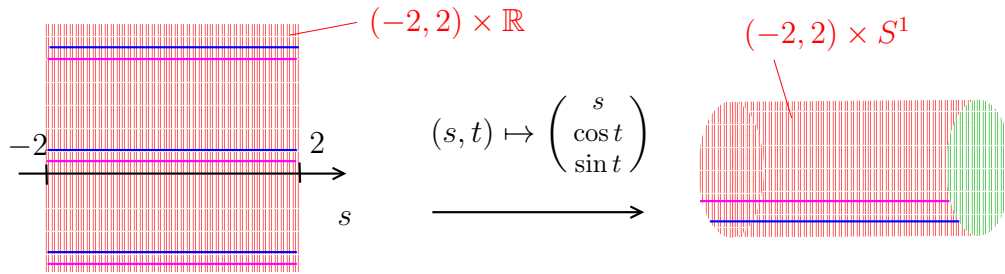


ABBILDUNG 69.

phismus, denn für alle $(s, t) \in \mathbb{R}$ ist die Einschränkung von f auf $(-2, 2) \times (t - \pi, t + \pi)$ ein Diffeomorphismus auf ihr Bild. Wir betrachten nun $(-2, 2) \times \mathbb{R}$ und $(-2, 2) \times S^1$ mit den üblichen Riemannschen Strukturen g und h , welche durch das Skalarprodukt auf $\mathbb{R}^2$ beziehungsweise $\mathbb{R}^3$ gegeben sind. Dann ist die obige Abbildung sogar eine lokale Isometrie. In der Tat, denn für alle $(s, t) \in (-2, 2) \times \mathbb{R}$ und $v, w \in T_{(s,t)}((-2, 2) \times \mathbb{R}) = \mathbb{R}^2$ gilt

$$\begin{aligned}\begin{array}{ccc} \text{Definition von } h \text{ und } \varphi_* & & \text{denn } \langle a, b \rangle = a^T b \\ \downarrow & & \downarrow \\ h(\varphi_* v, \varphi_* w) & = \langle D\varphi_{(s,t)} v, D\varphi_{(s,t)} w \rangle & = v^T D\varphi_{(s,t)}^T D\varphi_{(s,t)} w \\ & = v^T \underbrace{\begin{pmatrix} 1 & 0 & 0 \\ 0 & -\sin t & \cos t \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & -\sin t \\ 0 & \cos t \end{pmatrix}}_{= \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}} w & = v^T w = g(v, w). \end{array}\end{aligned}$$

(3) Genau die gleiche Rechnung wie in (2) zeigt, dass die Abbildung

$$\begin{aligned}\varphi: \mathbb{R} &\rightarrow S^1 \\ t &\mapsto (\cos t, \sin t)\end{aligned}$$

eine lokale Isometrie ist.

(4) Es sei $\varphi: M \rightarrow N$ ein (lokaler) Diffeomorphismus zwischen zwei Untermannigfaltigkeiten. Es sei h eine Riemannsche Struktur auf N . Dann ist $\varphi^* h$ nach Satz 13.2 eine Riemannsche Struktur auf M und $\varphi: (M, \varphi^* h) \rightarrow (N, h)$ ist eine (lokale) Isometrie.

Lemma 13.5. *Es sei $\varphi: M \rightarrow N$ eine glatte Abbildung zwischen zwei Riemannschen Untermannigfaltigkeiten (M, g) und (N, h) .*

- (1) *Wenn φ eine lokale Isometrie ist, dann gilt für jeden Weg $\gamma: [a, b] \rightarrow M$, dass*

$$\ell_g(\gamma) = \ell_h(\varphi \circ \gamma).$$

- (2) *Wenn φ eine Isometrie ist, dann gilt für alle $P, Q \in M$, dass*

$$d_M^g(P, Q) = d_N^h(\varphi(P), \varphi(Q)).$$

- (3) *Die Abbildung φ gibt eine Bijektion zwischen den Geodäten von (M, g) und den Geodäten von (N, h) .*

Der Beweis der ersten Aussage ist eine Übungsaufgabe in Übungsblatt 9. Die zweite und die dritte Aussage folgen leicht aus der ersten Aussage und den Definitionen.

Beispiel. Wir betrachten die “rechte offene Halbsphäre”

$$H := \{(x, y, z) \in S^2 \mid x > 0\}.$$

Wir bezeichnen wiederum mit g die übliche Riemannsche Struktur auf H , welche durch das Skalarprodukt auf $\mathbb{R}^3$ gegeben ist. Wir betrachten zudem den Diffeomorphismus¹⁶⁸

$$\begin{aligned} \Psi: V := \left(-\frac{\pi}{2}, \frac{\pi}{2}\right) \times (0, \pi) &\rightarrow H \\ (\varphi, \theta) &\mapsto \begin{pmatrix} \cos \varphi \cdot \sin \theta \\ \sin \varphi \cdot \sin \theta \\ \cos \theta \end{pmatrix}. \end{aligned}$$

Es sei nun $(\varphi, \theta) \in V$ und es seien $v, w \in T_{(\varphi, \theta)}V = \mathbb{R}^2$. Dann gilt per Definition, dass

$$\begin{aligned} &\text{Definition von } \Psi^*g \\ &\downarrow \\ (\Psi^*g)_{(\varphi, \theta)}(v, w) &= g_{\Psi(\varphi, \theta)}(\Psi_*v, \Psi_*w) = g_{\Psi(\varphi, \theta)}(D_{(\varphi, \theta)}\Psi(v), D_{(\varphi, \theta)}\Psi(w)) \\ &= \langle D_{(\varphi, \theta)}\Psi(v), D_{(\varphi, \theta)}\Psi(w) \rangle = v^T (D_{(\varphi, \theta)}\Psi)^T \cdot D_{(\varphi, \theta)}\Psi \cdot w \\ &\uparrow \qquad \qquad \qquad \uparrow \\ &\text{Definition von } g \qquad \qquad \text{für beliebige Vektoren } a \text{ und } b \text{ gilt } \langle a, b \rangle = a^T b \\ &= v^T \begin{pmatrix} -\sin \varphi \cdot \sin \theta & \cos \varphi \cdot \sin \theta & 0 \\ \cos \varphi \cdot \cos \theta & \sin \varphi \cdot \cos \theta & -\sin \theta \end{pmatrix} \begin{pmatrix} -\sin \varphi \cdot \sin \theta & \cos \varphi \cdot \cos \theta \\ \cos \varphi \cdot \sin \theta & \sin \varphi \cdot \cos \theta \\ 0 & -\sin \theta \end{pmatrix} w \\ &= v^T \begin{pmatrix} \sin^2 \theta & 0 \\ 0 & 1 \end{pmatrix} w. \end{aligned}$$

Insbesondere ist die Riemannsche Struktur Ψ^*g auf V von der Form des Beispiels, welches wir in Lemma 13.4 betrachtet hatten. Wir wollen diese etwas abstrakte Rechnung noch an dem Beispiel von Längen von Wegen nachvollziehen. Für $-\frac{\pi}{2} < a < b < \frac{\pi}{2}$ und fest gewähltes $\theta \in (0, \pi)$ betrachten wir den horizontalen Weg

$$\begin{aligned} \gamma_\theta: [a, b] &\rightarrow V \\ t &\mapsto (t, \theta), \end{aligned}$$

¹⁶⁸Dies ist ein Diffeomorphismus zwischen den Untermannigfaltigkeiten $(-\frac{\pi}{2}, \frac{\pi}{2}) \times (0, \pi)$ und H .

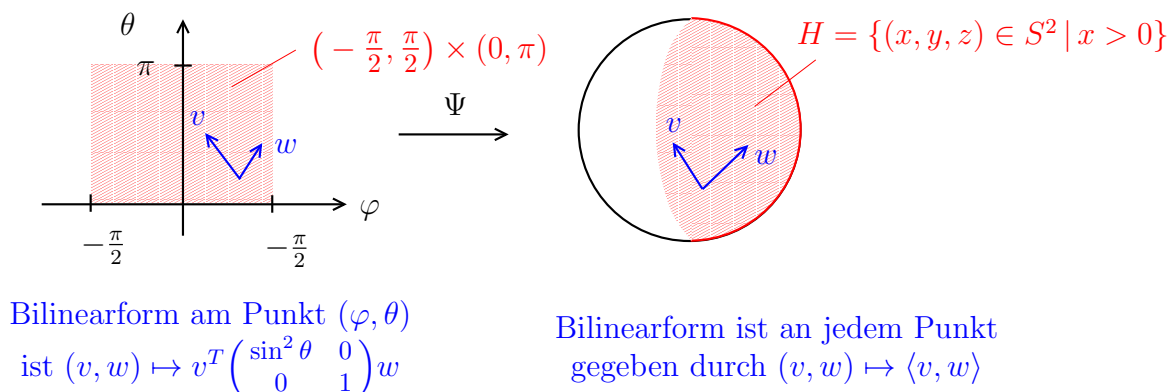


ABBILDUNG 70.

Auf Seite 162 hatten wir schon gesehen, dass

$$\ell_{\Psi^*g}(\gamma_\theta) = \sin(\theta) \cdot (b - a).$$

Dies entspricht gerade der Tatsache, dass auf der Sphäre auf einem festgewählten Breitengrad¹⁶⁹ θ der Weg vom Längengrad a zum Längengrad b die Länge $(b - a) \sin \theta$ besitzt.

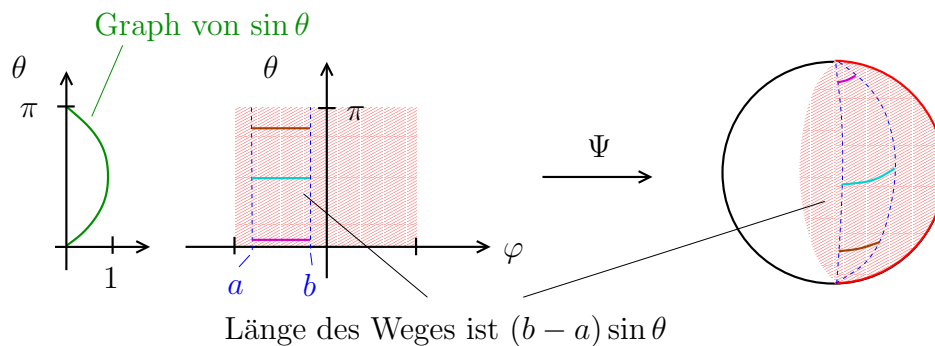


ABBILDUNG 71.

Wenn wir die Geodäten von H verstehen wollen haben wir nun also im Hinblick auf Lemma 13.5 die Wahlfreiheit:

- (1) entweder wir arbeiten direkt mit der Halbsphäre H und der einfachen Riemannschen Struktur, welche durch das übliche Skalarprodukt gegeben ist, oder
- (2) wir arbeiten mit dem offenen Rechteck $V = (-\frac{\pi}{2}, \frac{\pi}{2}) \times (0, \pi)$, aber dann müssen wir in Kauf nehmen, dass die Riemannsche Struktur etwas komplizierter ist.

¹⁶⁹Der Längengrad misst die “West-Ost”-Koordinate, und der Breitengrad entspricht der “Nord-Süd”-Koordinate auf der Sphäre.

13.2. Geodäten auf der Sphäre. Zur Erinnerung, ein *Großkreis* auf S^2 ist der Schnitt von S^2 mit einer Ebene, welche den Ursprung enthält. Für zwei Punkte $P \neq \pm Q$ auf S^2 gibt es genau eine Ebene, welche den Ursprung sowie die Punkte P und Q enthält. Mit anderen Worten, für $P \neq \pm Q$ gibt es genau einen Großkreis durch P und Q .

Unser Ziel ist es nun folgenden Satz zu beweisen.

Satz 13.6. *Es seien P und Q zwei Punkte auf S^2 . Jede Geodäte von P nach Q liegt auf einem Großkreis.*

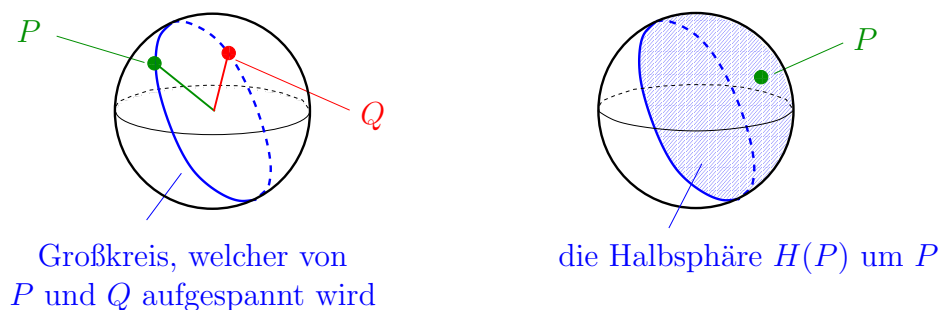


ABBILDUNG 72.

Bevor wir uns dem eigentlichen Beweis von Satz 13.6 zuwenden betrachten wir erst einmal Geodäten in offenen Halbsphären. Es sei nun $P \in S^2$. Wir bezeichnen

$$H(P) := \{Q \in S^2 \mid \langle P, Q \rangle > 0\}$$

als die *Halbsphäre um den Punkt P* .¹⁷⁰

Lemma 13.7. *Es sei $P \in S^2$ und es sei $Q \neq P \in H(P)$. Jede Geodäte von P nach Q innerhalb von $H(P)$ liegt auf dem Großkreis durch P und Q .*¹⁷¹

Beweis. Es sei also $P \in S^2$ und es sei $Q \neq P \in H(P)$. Nach Anwendung einer orthogonalen Matrix können wir o.B.d.A. annehmen, dass $P = (1, 0, 0)$.¹⁷² In diesem Fall ist

$$H(P) = \{(x, y, z) \in S^2 \mid x > 0\}$$

die “rechte offene Halbsphäre”. Nach Anwendung einer Rotation um die x -Achse können wir zudem annehmen, dass $Q = (\sin \alpha, 0, \cos \alpha)$ für ein $\alpha \in (0, \pi)$. Wie auf Seite 165

¹⁷⁰Mit anderen Worten, $H(P)$ ist die Menge aller Punkte auf S^2 , welche näher an P als an $-P$ liegen.

¹⁷¹Das Lemma schließt zumindest a priori nicht aus, dass es eine Geodäte in S^2 gibt, welche $H(P)$ verläßt, und nicht auf einem Großkreis liegt.

¹⁷²Hierbei verwenden wir, dass wir schon auf Seite 163 gesehen hatten, dass die Multiplikation mit einer orthogonalen Matrizen eine Isometrie von S^2 ist. Zudem verwenden wir, dass nach Lemma 13.5 Isometrien Geodäten in Geodäten überführen.

betrachten wir den Diffeomorphismus

$$\begin{aligned} \Psi: V := \left(-\frac{\pi}{2}, \frac{\pi}{2}\right) \times (0, \pi) &\rightarrow H \\ (\varphi, \theta) &\mapsto \begin{pmatrix} \cos \varphi \cdot \sin \theta \\ \sin \varphi \cdot \sin \theta \\ \cos \theta \end{pmatrix}. \end{aligned}$$

In diesem Fall ist $\Psi(0, \frac{\pi}{2}) = P$ und $\Psi(0, \alpha) = Q$.

Wir bezeichnen mit g die Riemannsche Metrik auf $H \subset S^2$, welche wie immer durch das Skalarprodukt auf $\mathbb{R}^3$ gegeben ist. Nachdem Ψ ein Diffeomorphismus ist, ist Ψ per Definition auch eine Isometrie zwischen (V, Ψ^*g) und (H, g) . Lemma 13.5 besagt insbesondere, dass Ψ eine Bijektion zwischen den Geodäten auf (V, Ψ^*g) und den Geodäten auf (H, g) induziert. Auf Seite 165 hatten wir gesehen, dass die Riemannsche Metrik Ψ^*g auf V an jedem Punkt

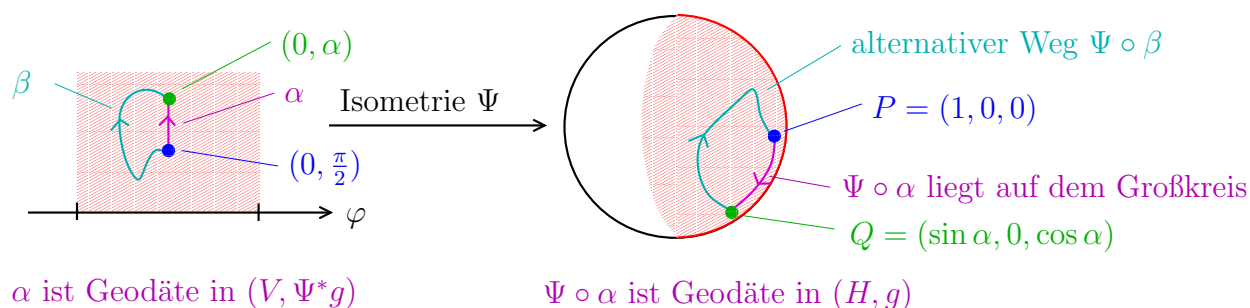


ABBILDUNG 73. Skizze zum Beweis von Lemma 13.7.

(φ, θ) gegeben ist durch

$$\begin{aligned} \mathbb{R}^2 \times \mathbb{R}^2 &\mapsto \mathbb{R} \\ (v, w) &\mapsto v^T \begin{pmatrix} \sin^2 \theta & 0 \\ 0 & 1 \end{pmatrix} w. \end{aligned}$$

Es folgt nun aus Lemma 13.4, dass in (V, Ψ^*g) , jede Geodäte von $(0, \frac{\pi}{2})$ nach $(0, \alpha)$ vertikal verläuft. Aber das Bild einer solchen vertikalen Gerade in V unter der Abbildung Ψ liegt gerade auf dem Großkreis durch P und Q . $\square$

Wir wenden uns nun dem eigentlichen Beweis von Satz 13.6 zu.

Beweis von Satz 13.6. Es sei $\gamma: [a, b] \rightarrow S^2$ eine Geodäte zwischen zwei Punkten P und Q auf S^2 . Wenn $P = Q$, dann gibt es nichts zu beweisen. Wir nehmen jetzt also an, dass $P \neq Q$.

Es sei $t \in (a, b)$ beliebig. Wir wählen ein $\epsilon > 0$, so dass $[t - \epsilon, t + \epsilon] \subset [a, b]$, und so dass¹⁷³ $\gamma([t - \epsilon, t + \epsilon]) \subset H(\gamma(t))$. Nach Lemma 13.7 liegt $\gamma|_{[t - \epsilon, t + \epsilon]}$ auf einem eindeutig bestimmten Großkreis $G(t)$.¹⁷⁴

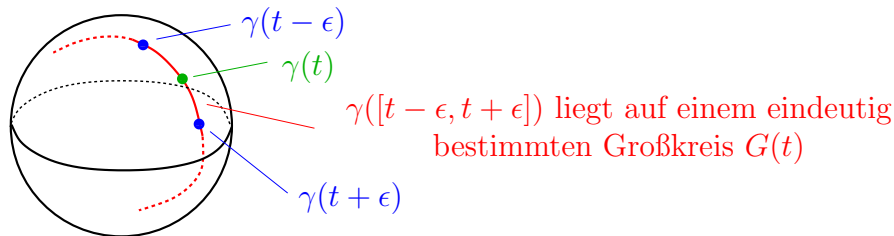


ABBILDUNG 74. Skizze zur Definition von $G(t)$.

Behauptung.

- (1) Die Definition von $G(t)$ hängt nicht von der Wahl von $\epsilon > 0$ ab.
- (2) Die Abbildung

$$\begin{aligned} G: (a, b) &\rightarrow \text{Menge der Großkreise} \\ t &\mapsto G(t) \end{aligned}$$

ist lokal konstant, d.h. zu jedem $t \in (a, b)$ gibt es ein $\eta > 0$, so dass G auf $(t - \eta, t + \eta)$ konstant ist.

Wir beweisen zuerst die Aussage (1). Es sei also $t \in (a, b)$. Es seien $\epsilon > 0$ und $\eta > 0$ mit den vorgeschriebenen Eigenschaften. Wir bezeichnen die entsprechenden Großkreise kurzzeitig mit G_ϵ und G_η . Wir müssen zeigen, dass $G_\epsilon = G_\eta$. O.B.d.A. sei $\eta < \epsilon$. Dann gilt $\gamma|_{[t - \eta, t + \eta]} \subset G_\eta$ und auch $\gamma|_{[t - \eta, t + \eta]} \subset \gamma|_{[t - \epsilon, t + \epsilon]} \subset G_\epsilon$. Aus der Eindeutigkeitsaussage von Lemma 12.3 folgt, dass $G_\eta = G_\epsilon$.

Wir beweisen nun Aussage (2). Es sei also $t \in (a, b)$ und es sei $\epsilon > 0$ mit den vorgeschriebenen Eigenschaften. Wir setzen $\eta = \frac{1}{2}\epsilon$. Nun sei $s \in (t - \eta, t + \eta)$. Dann gilt

$$\gamma([s - \eta, s + \eta]) \subset \gamma([t - \epsilon, t + \epsilon]) \subset G(t).$$

$\uparrow \qquad \qquad \qquad \uparrow$
 denn $\eta = \frac{1}{2}\epsilon$ und $s \in (t - \eta, t + \eta)$ Definition von $G(t)$

Also gilt per Definition, dass $G(s) = G(t)$. Wir haben damit die Behauptung bewiesen.

Behauptung. Es gibt genau einen Großkreis G , so dass $\gamma((a, b)) \subset G$.

¹⁷³Es ist $\gamma(t) \in H(\gamma(t))$. Nachdem $H(\gamma(t))$ eine offene Teilmenge von S^2 ist und nachdem γ stetig ist, existiert insbesondere solch ein $\epsilon > 0$.

¹⁷⁴Nach Voraussetzung ist γ eine Geodäte in S^2 . Also ist nach Lemma 12.3 auch $\gamma|_{[t - \epsilon, t + \epsilon]}$ eine Geodäte in S^2 . Nachdem $\gamma|_{[t - \epsilon, t + \epsilon]}$ ganz in $H(\gamma(t))$ liegt, ist dann $\gamma|_{[t - \epsilon, t + \epsilon]}$ auch eine Geodäte in der offenen Halbsphäre $H(\gamma(t))$. Wir können also in der Tat Lemma 13.7 anwenden. Der Großkreis ist eindeutig, weil aus Lemma 12.3 sofort folgt, dass $\gamma(t - \epsilon) \neq \gamma(t + \epsilon)$.

Es sei X die Menge der Großkreise. Wir betrachten X mit der diskreten Topologie. Es folgt aus der obigen Behauptung, dass die Abbildung

$$\begin{aligned} (a, b) &\rightarrow X \\ t &\mapsto G(t) \end{aligned}$$

stetig¹⁷⁵ ist. Also ist die Abbildung nach Lemma 5.3 auch konstant. Wir haben damit die Behauptung bewiesen.

Aus Stetigkeitsgründen¹⁷⁶ folgt nun aber auch, dass $\gamma([a, b]) \subset G$. □

Für P und Q auf S^2 sei nun $\alpha(P, Q) \in [0, \pi]$ der Winkel, welcher von dem Ursprung und den Punkten P und Q aufgespannt wird. Etwas genauer, $\alpha = \alpha(P, Q)$ ist die eindeutig bestimmte reelle Zahl α in $[0, \pi]$ mit

$$\langle P, Q \rangle = \cos(\alpha).$$

Mit ein klein bißchen Mehraufwand kann man, ganz analog zum Beweis von Satz 12.4, nun noch folgenden Satz beweisen. Wir überlassen die genaue Ausführung des Beweises als Übungsaufgabe.

Satz 13.8. *Es seien P und Q zwei Punkte auf S^2 . Dann gilt*

- (1) $d_{S^2}(P, Q) = \alpha(P, Q)$.
- (2) Wenn $P \neq \pm Q$, dann gibt es genau eine Geodäte von P nach Q , diese bewegt sich auf dem eindeutig bestimmten Großkreis durch P und Q .
- (3) Wenn $P = -Q$, dann gibt es unendlich viele Großkreise durch P und Q , und jeder Halbgroßkreis durch P und Q entspricht gerade einer Geodäte von P nach Q .

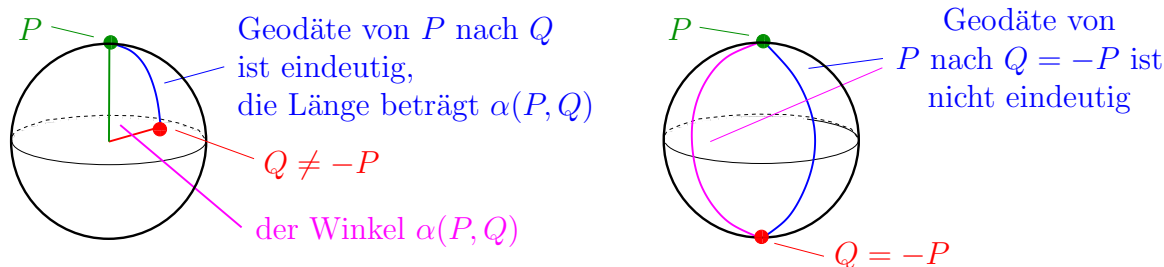


ABBILDUNG 75.

¹⁷⁵Warum folgt das aus der obigen Behauptung?

¹⁷⁶Was soll das jetzt heißen? Warum folgt die Aussage aus “Stetigkeitsgründen”?

14. HYPERBOLISCHE GEOMETRIE

14.1. Der hyperbolische Raum. In diesem Kapitel identifizieren wir $\mathbb{R}^2$ mit $\mathbb{C}$, und wechseln kommentarlos von einem Gesichtspunkt zum anderen.

Definition. Der 2-dimensionale hyperbolische Raum ist die Riemannsche Untermannigfaltigkeit, welche gegeben ist durch

$$\mathbb{H} = \{z = x + iy \in \mathbb{C} = \mathbb{R}^2 \mid \operatorname{Im}(z) > 0\}$$

zusammen mit der Riemannschen Struktur g , welche an jedem Punkt $z = x + iy \in \mathbb{H}$ gegeben ist durch¹⁷⁷¹⁷⁸

$$\begin{aligned} g: \mathbb{R}^2 \times \mathbb{R}^2 &\rightarrow \mathbb{R} \\ (v, w) &\mapsto \frac{1}{\operatorname{Im}(z)^2} \langle v, w \rangle = \frac{1}{y^2} \langle v, w \rangle, \end{aligned}$$

wobei $\langle v, w \rangle$ das übliche Skalarprodukt auf $\mathbb{R}^2$ bezeichnet. Für einen Weg $\gamma: [a, b] \rightarrow \mathbb{H}$ bezeichnen wir $\ell_g(\gamma)$ als die *hyperbolische Länge* von γ . Zudem bezeichnen wir die Geodäten in $(\mathbb{H}, g)$ als die *hyperbolischen Geodäten*.

Beispiel. Es sei $\gamma: [a, b] \rightarrow \mathbb{H}$ eine stetig differenzierbare Abbildung. Dann ist natürlich

$$\text{euklidische Länge von } \gamma = \int_a^b |\gamma'(t)| dt,$$

während direkt aus den Definitionen folgt, dass

$$\text{hyperbolische Länge von } \gamma = \int_a^b \frac{1}{\operatorname{Im}(\gamma(t))} |\gamma'(t)| dt.$$

Im Folgenden wollen wir insbesondere die Geodäten im hyperbolischen Raum bestimmen. Wir verfahren dabei ganz ähnlich zum euklidischen und zum sphärischen Fall: wenn zwei Punkte P und Q in $\mathbb{H}$ gegeben sind, dann wollen wir P und Q mithilfe einer Isometrie von $\mathbb{H}$ in eine besonders günstige Lage überführen, in der es dann einfach ist die Geodäten zu bestimmen. Um dieses Programm durchzuführen müssen wir nun “genügend” Isometrien von $\mathbb{H}$ bestimmen. Wir bewerkstelligen dies, indem wir uns an die Möbiustransformationen aus der Analysis III entsinnen.

Lemma 14.1.

¹⁷⁷Die Riemannsche Struktur ist also vom Typ $g(\rho)$ mit $\rho(z) = \frac{1}{\operatorname{Im}(z)^2}$, siehe Seite 161.

¹⁷⁸Diese Definition verallgemeinert sich leicht auf höhere Dimensionen. Der *n-dimensionale hyperbolische Raum* ist die Riemannsche Untermannigfaltigkeit, welche gegeben ist durch

$$\mathbb{H}^n = \{(x_1, \dots, x_n) \in \mathbb{R}^n \mid x_n > 0\}$$

zusammen mit der Riemannschen Struktur g , welche an jedem Punkt $(x_1, \dots, x_n) \in \mathbb{H}^n$ gegeben ist durch

$$\begin{aligned} g: \mathbb{R}^n \times \mathbb{R}^n &\rightarrow \mathbb{R} \\ (v, w) &\mapsto \frac{1}{x_n^2} \langle v, w \rangle. \end{aligned}$$

In dieser Vorlesung betrachten wir der Einfachheit halber nur den Fall $n = 2$.

(1) Es seien $a, b, c, d \in \mathbb{R}$ mit $ad - bc = 1$. Die Abbildung

$$\begin{aligned} \mathbb{H} &\rightarrow \mathbb{H} \\ z &\mapsto \frac{az + b}{cz + d}, \end{aligned}$$

genannt Möbiustransformation, ist ein Diffeomorphismus.

(2) Die Verknüpfung von zwei Möbiustransformationen und die Umkehrabbildung einer Möbiustransformation sind wiederum Möbiustransformationen.

Beweis. Es seien $a, b, c, d \in \mathbb{R}$ mit $ad - bc = 1$. Wir zeigen zuerst, dass das Bild der Abbildung

$$\begin{aligned} \Phi: \mathbb{H} &\rightarrow \mathbb{H} \\ z &\mapsto \frac{az + b}{cz + d}, \end{aligned}$$

wiederum in $\mathbb{H}$ liegt. Dies ist in der Tat der Fall, denn für $z = x + iy$ mit $y > 0$ gilt

$$\begin{aligned} \operatorname{Im} \left(\frac{az + b}{cz + d} \right) &= \operatorname{Im} \left(\frac{(c(x - iy) + d)(a(x + iy) + b)}{(c\bar{z} + d)(cz + d)} \right) = \frac{yad - ybc}{\|cz + d\|^2} = \frac{y}{\|cz + d\|^2} > 0. \\ &\quad \uparrow \qquad \qquad \qquad \uparrow \\ &\quad \text{denn } a, b, c, d \in \mathbb{R} \qquad \text{denn } ad - bc = 1 \end{aligned}$$

Wir bezeichnen nun mit $M(\mathbb{H})$ die Menge der glatten Abbildungen von $\mathbb{H} \rightarrow \mathbb{H}$. Dies ist eine Halbgruppe bezüglich der Verknüpfung, d.h. das Assoziativgesetz gilt und es gibt ein neutrales Element, nämlich die Identität.

Wir beweisen nun folgende Behauptung.

Behauptung. Die Abbildung

$$\begin{aligned} (2, \mathbb{R}) &\rightarrow M(\mathbb{H}) \\ A := \begin{pmatrix} a & b \\ c & d \end{pmatrix} &\rightarrow \left(\begin{array}{ccc} \Phi(A): \mathbb{H} &\rightarrow & \mathbb{H} \\ & z &\mapsto & \frac{az+b}{cz+d} \end{array} \right) \end{aligned}$$

ist ein Homomorphismus von Halbgruppen.

Es ist offensichtlich, dass $\Phi(\operatorname{id}_2)$ die Identitätsabbildung auf $\mathbb{H}$ ist. Wir müssen also nur noch zeigen, dass die Abbildung multiplikativ ist. Es seien also

$$A := \begin{pmatrix} a & b \\ c & d \end{pmatrix} \quad \text{und} \quad A' := \begin{pmatrix} a' & b' \\ c' & d' \end{pmatrix}$$

zwei Matrizen in $(2, \mathbb{R})$. Dann ist

$$\begin{aligned} (\Phi(A) \circ \Phi(A'))(z) &= \Phi(A) \left(\frac{a'z + b'}{c'z + d'} \right) = \frac{\frac{a'z + b'}{c'z + d'} + b}{\frac{a'z + b'}{c'z + d'} + d} = \frac{a(a'z + b') + b(c'z + d')}{c(a'z + b') + d(c'z + d')} \\ &= \frac{(aa' + bc')z + ab' + bd'}{(ca' + dc')z + cb' + dd'} = \Phi \left(\begin{pmatrix} aa' + bc' & ab' + bd' \\ ca' + dc' & cb' + dd' \end{pmatrix} \right)(z) = \Phi(A \cdot A')(z). \end{aligned}$$

Wir haben damit die Behauptung bewiesen.

Jede Möbiustransformation ist offensichtlich glatt. Es folgt aus der obigen Behauptung, dass die Umkehrabbildung der Möbiustransformation $\Phi(A)$ gegeben ist durch die Möbiustransformation $\Phi(A^{-1})$. Also ist jede Möbiustransformation ein Diffeomorphismus.

Nachdem das Produkt von zwei Matrizen in $(2, \mathbb{R})$ wiederum eine Matrix in $(2, \mathbb{R})$ ist, folgt nun aus der obigen Behauptung, dass die Verknüpfung von zwei Möbiustransformationen wiederum eine Möbiustransformation ist. $\square$

Wir sind im Folgenden besonders an den folgenden drei Typen von Möbiustransformationen interessiert:¹⁷⁹

- (1) für $r > 0$ ist die Streckung $z \mapsto rz = \frac{\sqrt{r} \cdot z + 0}{0 \cdot z + \frac{1}{\sqrt{r}}}$ eine Möbiustransformation,
- (2) für $d \in \mathbb{R}$ ist die horizontale Verschiebung $z \mapsto z + d = \frac{1 \cdot z + d}{0 \cdot z + 1}$ eine Möbiustransformation,
- (3) die Abbildung $z \mapsto -\frac{1}{z} = \frac{0 \cdot z + 1}{(-1) \cdot z + 0}$ ist eine Möbiustransformation.

Wir betrachten die letzte Möbiustransformation noch etwas genauer. Diese ist die Verknüpfung der Inversion am Einheitskreis $z \mapsto \frac{1}{z}$ mit der Spiegelung $z \mapsto -z$ am Ursprung. Die Abbildung wird in Abbildung 76 skizziert. Der Einfachheit halber nennen wir die Abbildung $z \mapsto -\frac{1}{z}$ ebenfalls eine Inversion.

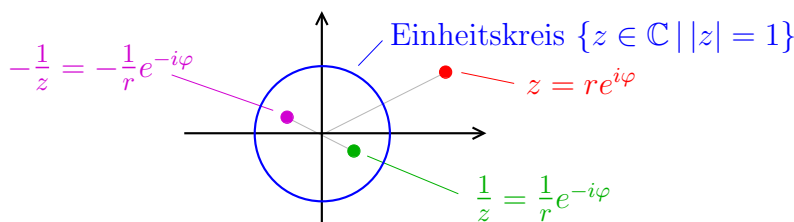


ABBILDUNG 76.

Satz 14.2. Jede Möbiustransformation kann geschrieben werden als eine Verknüpfung von Streckungen, Verschiebungen und Inversionen am Einheitskreis.

Beweis. Es seien $a, b, c, d \in \mathbb{R}$ mit $ad - bc = 1$ und es sei

$$\Phi(z) = \frac{az + b}{cz + d}$$

die zugehörige Möbiustransformation. Wenn $c = 0$, dann ist

$$\Phi(z) = \frac{a}{d}z + \frac{b}{d}$$

¹⁷⁹Warum sind dies eigentlich Möbiustransformationen?

die Verknüpfung der Streckung $z \mapsto \frac{a}{d}z$ und der Verschiebung $z \mapsto \frac{b}{d}z$. Wenn $c \neq 0$, dann ist

$$\Phi(z) = \frac{az+b}{cz+d} = \frac{acz+bc}{c^2z+dc} = \frac{acz+ad-1}{c^2z+cd} = \frac{-1+a(cz+d)}{c^2z+cd} = \frac{-1}{c^2z+cd} + \frac{a}{c}.$$

$\uparrow$
 denn $ad-bc=1$

Wir sehen also, dass Φ die Verknüpfung der folgenden Abbildungen ist:

$$z \mapsto c^2z, \quad z \mapsto z+dc, \quad z \mapsto \frac{-1}{z} \quad \text{und} \quad z \mapsto z + \frac{a}{c}. \quad \square$$

Satz 14.3. *Jede Möbiustransformation ist eine orientierungserhaltende Isometrie von $\mathbb{H}$.*

Bemerkung. Es gilt auch die interessantere Aussage, dass jede orientierungserhaltende Isometrie von $\mathbb{H}$ eine Möbiustransformation ist. Dies wird beispielsweise in Theorem 3.19 von Anderson: Hyperbolic Geometry bewiesen.

Beweis. Nach Satz 9.16 und Satz 14.2 genügt es zu zeigen, dass Streckungen, Verschiebungen und Inversionen orientierungserhaltende Isometrien sind.

Wir wenden uns zuerst den Streckungen zu. Es sei also $r > 0$. Wir bezeichnen dann mit $\Phi(z) = rz$ die zugehörige Streckung. Es sei $z \in \mathbb{H}$. Dann ist

$$D_z\Phi = \begin{pmatrix} r & 0 \\ 0 & r \end{pmatrix},$$

diese Matrix hat eine positive Determinante, also ist Φ orientierungserhaltend. Wir müssen nun noch zeigen, dass Φ eine Isometrie ist. Es sei also $z \in H$ und es seien $v, w \in T_z\mathbb{H} = \mathbb{R}^2$. Dann ist

$$\begin{aligned} (\Phi^*g)_z(v, w) &= g_{\Phi(z)}(D_z\Phi(v), D_z\Phi(w)) = g_{rz}(rv, rw) = \frac{1}{\text{Im}(rz)^2} \langle rv, rw \rangle \\ &\quad \uparrow \\ &\quad \text{Definition der Riemannschen Struktur auf } \mathbb{H} \\ &= \frac{1}{r^2 \text{Im}(z)^2} \cdot r^2 \cdot \langle v, w \rangle = \frac{1}{\text{Im}(z)^2} \langle v, w \rangle = g_z(v, w). \end{aligned}$$

Wir haben also gezeigt, dass jede Streckung eine Isometrie ist.

Ganz analog, und noch etwas leichter, kann man verifizieren, dass Verschiebungen orientierungserhaltende Isometrien von $\mathbb{H}$ sind.

Die Aussage, dass die Inversion $z \mapsto -\frac{1}{z}$ eine orientierungserhaltende Isometrie von $\mathbb{H}$ ist wird in Übungsblatt 10 bewiesen. $\square$

Lemma 14.4. *Zu je zwei Punkten $P, Q \in \mathbb{H}$ gibt es immer eine Möbiustransformation φ mit $\varphi(P) = Q$.*

Man kann das Lemma einfach dadurch beweisen, dass man zwei Punkte P und Q mithilfe von Verschiebungen und Streckungen ineinander überführt. Der Beweis wird in Übungsblatt 9 ausgeführt.

Wir führen jetzt folgende, auf den ersten Blick sehr ungewöhnliche Definition ein. Wir werden später begründen, warum der Name “hyperbolische Gerade” durchaus angebracht ist.

Definition.

- (1) Für jedes $a \in \mathbb{R}$ bezeichnen wir

$$\{z \in \mathbb{H} \mid \operatorname{Re}(z) = a\}$$

als *hyperbolische Gerade mit Endpunkten a und ∞* .

- (2) Für $s \in \mathbb{R}$ und $r > 0$ bezeichnen wir

$$\{z \in \mathbb{H} \mid |z - s| = r\}$$

als *hyperbolische Gerade mit Endpunkten $s - r$ und $s + r$* .

Eine hyperbolische Gerade vom ersten Typ ist also eine euklidische Halbgerade, welche senkrecht auf der x -Achse steht. Eine hyperbolische Gerade vom zweiten Typ entspricht einem euklidischen Halbkreis, dessen Mittelpunkt auf der x -Achse liegt.

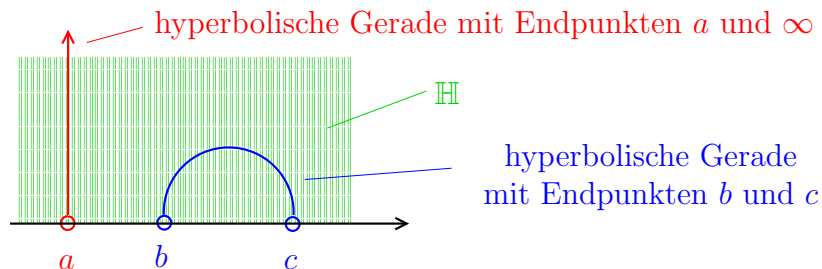


ABBILDUNG 77. Beispiele von hyperbolischen Geraden.

Lemma 14.5. Zu je zwei verschiedenen Punkten P und Q in $\mathbb{H}$ gibt es genau eine hyperbolische Gerade $g(P, Q)$ durch P und Q .

Beweis. Es seien also $P \neq Q$ zwei Punkte in $\mathbb{H}$. Wenn $\operatorname{Re}(P) = \operatorname{Re}(Q)$, dann liegt g auf der hyperbolischen Gerade, welche durch $\operatorname{Re}(z) = \operatorname{Re}(P) = \operatorname{Re}(Q)$ beschrieben ist. Wenn $\operatorname{Re}(P) \neq \operatorname{Re}(Q)$, dann gibt es genau einen euklidischen Halbkreis durch P und Q , dessen Mittelpunkt auf der x -Achse liegt, siehe Abbildung 78. Die Eindeutigkeit folgt aus den klassischen Sätzen der euklidischen Geometrie oder aus einer leichten Rechnung. $\square$

Satz 14.6. Es sei Φ eine Möbiustransformation. Dann gilt¹⁸⁰

$$\Phi \left(\begin{array}{c} \text{hyperbolische Gerade mit} \\ \text{Endpunkten } P \text{ und } Q \end{array} \right) = \begin{array}{c} \text{hyperbolische Gerade mit} \\ \text{Endpunkten } \Phi(P) \text{ und } \Phi(Q). \end{array}$$

Beispiel. In Abbildung 79 sehen wir, dass die Inversion angewandt auf eine Gerade von Typ (1) eine Gerade von Typ (2) ergibt.¹⁸¹

¹⁸⁰Hierbei verwenden wir die üblichen “Rechenregeln” für ∞ , z.B. ist $\frac{1}{0} = \infty$ und $\frac{1}{\infty} = 0$.

¹⁸¹Diese Aussage stimmt nicht ganz, für welche Gerade von Typ (1) erhalten wir keine Gerade von Typ (2)?

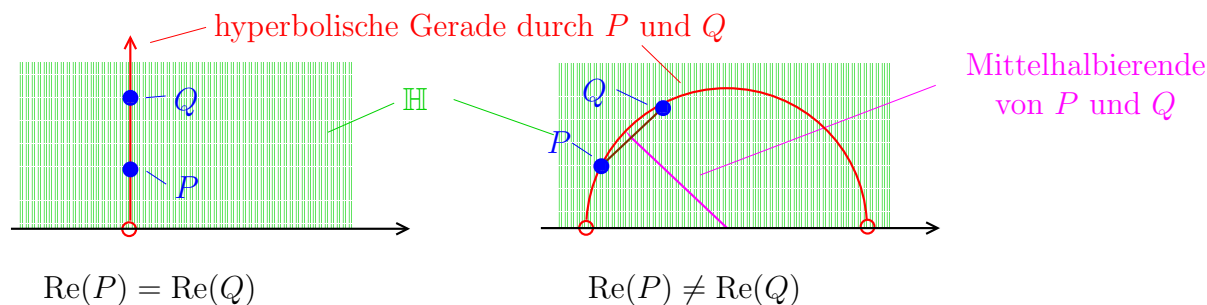
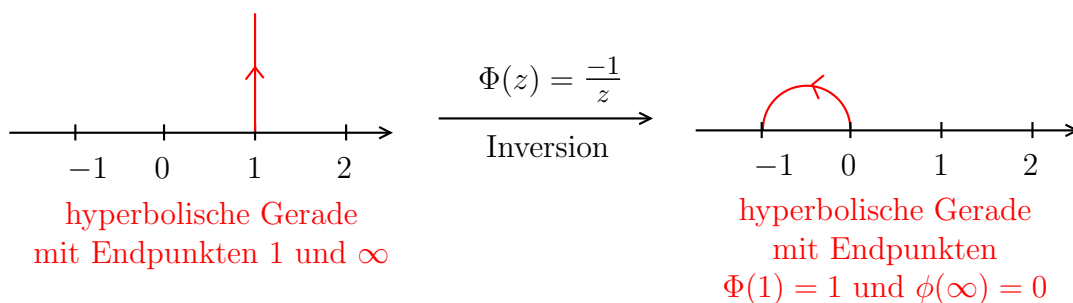
ABBILDUNG 78. Hyperbolischen Gerade durch zwei Punkte P und Q .

ABBILDUNG 79.

Beweis. Nach Satz 14.2 genügt es den Satz für Streckungen, Verschiebungen und Inversionen zu beweisen. Die Aussage ist offensichtlich für Streckungen und Verschiebungen.

Es sei nun $\Phi(z) = -\frac{1}{z}$ die Inversion. Bevor wir uns dem eigentlichen Beweis zuwenden machen wir noch folgende Vorbemerkung: Für jedes $a \in \mathbb{R}$ ist

$$\{z \in \mathbb{H} \mid \text{Re}(z) = a\} = \{z \in \mathbb{H} \mid z + \bar{z} = 2a\}$$

und für $s \in \mathbb{R}$ und $r > 0$ ist

$$\{z \in \mathbb{H} \mid |z - s| = r\} = \{z \in \mathbb{H} \mid z\bar{z} - sz - s\bar{z} = r^2 - s^2\}.$$

$$\uparrow$$

$$\text{denn } |z - s|^2 = (z - s)(\bar{z} - s)$$

Es sei zuerst g eine hyperbolische Gerade vom zweiten Typ, d.h. es gibt $s \in \mathbb{R}$ und $r > 0$ mit

$$g = \{z \in \mathbb{H} \mid z\bar{z} - sz - s\bar{z} = r^2 - s^2\}.$$

Dann ist

$$\Phi(g) = \left\{z \in \mathbb{H} \mid \frac{1}{z}\frac{1}{\bar{z}} + s\frac{1}{z} + s\frac{1}{\bar{z}} = r^2 - s^2\right\} = \left\{z \in \mathbb{H} \mid (r^2 - s^2)z\bar{z} - s\bar{z} - sz = 1\right\}.$$

$\uparrow$
Durchmultiplizieren mit $z \cdot \bar{z}$

Wenn $r^2 - s^2 = 0$, dann beschreibt dies eine hyperbolische Gerade vom ersten Typ und wenn $r^2 - s^2 \neq 0$, dann beschreibt dies eine hyperbolische Gerade vom zweiten Typ.¹⁸² Der Fall, dass g eine hyperbolische Gerade vom ersten Typ ist, wird ganz analog bewiesen. Die Aussage über die Endpunkte kann man ebenso leicht nachvollziehen. $\square$

Folgendes Lemma wird in Übungsblatt 10 bewiesen.

Lemma 14.7. *Zu je zwei hyperbolischen Geraden g und h gibt es immer eine Möbiustransformation, welche g in h überführt.*

Satz 14.8. *Es seien P und Q zwei verschiedene Punkte auf $\mathbb{H}$. Jede Geodäte von P nach Q verläuft auf der eindeutigen hyperbolischen Gerade durch P und Q .*

Wie wir gleich sehen werden, ist der Beweis von Satz 14.8 formal fast der gleiche wie der Beweis von Lemma 12.1 und von Lemma 13.7.

Beweis. Es seien also P und Q zwei verschiedene Punkte auf $\mathbb{H}$. Nach Satz 14.5 gibt es eine hyperbolische Gerade g , welche P und Q verbindet. Nach Lemma 14.7 gibt es eine Möbiustransformation Φ , welche g in die hyperbolische Gerade $h = \{z \in \mathbb{H} \mid \operatorname{Re}(z) = 0\}$ überführt. Wir schreiben $\Phi(P) = (0, p)$ und $\Phi(Q) = (0, q)$.

Nach Satz 14.3 ist die Möbiustransformation Φ insbesondere eine Isometrie und führt damit nach Lemma 13.5 insbesondere Geodäten in Geodäten über. Es genügt nun zu zeigen, dass jede Geodäte in $\mathbb{H}$ von $(0, p)$ zu $(0, q)$ auf h liegt. Aber dies folgt sofort aus Lemma 13.4, denn die Riemannsche Struktur g auf $\mathbb{H}$ ist von der Form

$$g_{x+iy}(v, w) = \frac{1}{y^2} \langle v, w \rangle = \frac{1}{y^2} v^T w = v^T \underbrace{\begin{pmatrix} \frac{1}{y^2} & 0 \\ 0 & \frac{1}{y^2} \end{pmatrix}}_{=: A(x, y)} w. \quad \square$$

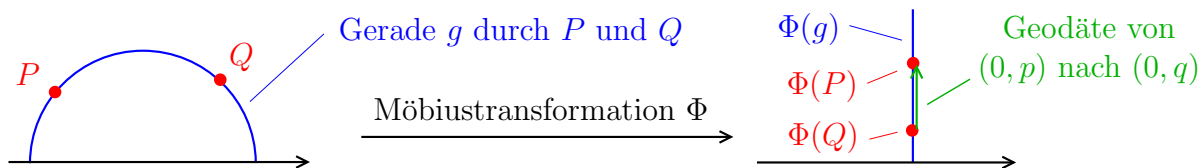


ABBILDUNG 80. Skizze zum Beweis von Satz 14.8.

¹⁸²Wenn man genau hinschaut, dann ist das Argument im zweiten Fall etwas flott. Eine Teilmenge der Form $\{z \in \mathbb{H} \mid z\bar{z} + az + a\bar{z} = b\}$ mit $a, b \in \mathbb{R}$ ist eine hyperbolische Gerade genau dann, wenn $b + a^2 > 0$. In unserem Fall ist

$$\{z \in \mathbb{H} \mid (r^2 - s^2)z\bar{z} - s\bar{z} - sz = 1\} = \{z \in \mathbb{H} \mid z\bar{z} - \frac{s}{r^2 - s^2}\bar{z} - \frac{s}{r^2 - s^2}z = \frac{1}{r^2 - s^2}\}$$

glücklicherweise von dieser Form.

Wenn man genau hinschaut, dann sieht man, dass wir eigentlich noch gar nicht gezeigt haben, dass es überhaupt immer eine Geodäte zwischen zwei Punkten in $\mathbb{H}$ gibt. Wir holen dies nun nach.

Satz 14.9.

- (1) Es seien $(0, a)$ und $(0, b)$ Punkte in $\mathbb{H}$ mit $b > a$. Dann ist

$$\begin{aligned} \gamma: [0, \ln(b) - \ln(a)] &\rightarrow \mathbb{H} \\ t &\mapsto \begin{pmatrix} 0 \\ \exp(t + \ln(a)) \end{pmatrix} \end{aligned}$$

die einzige Geodäte von $(0, a)$ nach $(0, b)$.

- (2) Es ist

$$\text{hyperbolischer Abstand zwischen } (0, a) \text{ und } (0, b) = \ln(b) - \ln(a).$$

- (3) Zu je zwei Punkten P und Q in $\mathbb{H}$ gibt es genau eine Geodäte von P nach Q .

Beweis.

- (1) Die Kurve γ ist offensichtlich glatt. Außerdem sieht man sofort durch Einsetzen, dass $\gamma(0) = (0, a)$ und $\gamma(\ln(b) - \ln(a)) = (0, b)$. Zudem gilt für alle t , dass

$$\begin{aligned} \|\gamma'(t)\|_{\mathbb{H}} &= \frac{1}{y\text{-Koordinate von } \gamma(t)^2} \langle \gamma'(t), \gamma'(t) \rangle \\ &= \frac{1}{\exp(t + \ln(a))^2} \cdot \left\langle \begin{pmatrix} 0 \\ \exp(t + \ln(a)) \end{pmatrix}, \begin{pmatrix} 0 \\ \exp(t + \ln(a)) \end{pmatrix} \right\rangle = 1. \end{aligned}$$

Der Rest vom Argument ist nun ganz analog zum Beweis von Satz 12.4.

- (2) Diese Aussage folgt sofort aus (1), denn für jede Geodäte $\gamma: [a, b] \rightarrow \mathbb{H}$ gilt per Definition einer Geodäte, dass $d_{\mathbb{H}}(\gamma(a), \gamma(b)) = b - a$.
 (3) Diese Aussage kann man, wie üblich, mithilfe einer Möbiustransformation auf den ersten Fall zurückführen. $\square$

Viele Definitionen und Aussagen der klassischen euklidischen Geometrie übertragen sich auf die hyperbolische Geometrie. Beispielsweise werden wir im Folgenden Spiegelungen in Geraden im hyperbolischen Raum einführen.

Wir erinnern uns dazu erst noch einmal an die Definition im euklidischen Fall. Es sei also g eine euklidische Gerade und P ein Punkt, welcher nicht auf g liegt. Wir definieren die Spiegelung $s_g(P)$ wie folgt:

- (1) Wir bezeichnen mit h die eindeutig bestimmte euklidische Gerade h durch P , welche g orthogonal in einem Punkt S schneidet.
 (2) Wir definieren dann $Q = s_g(P)$ als den eindeutig bestimmten Punkt auf h , mit $d(S, Q) = d(S, P)$ und mit $P \neq Q$.

Für $P \in g$ definieren wir zudem $s_g(P) = P$.

Wir führen nun die gleiche Konstruktion hyperbolisch durch. Es sei also g eine hyperbolische Gerade und P ein Punkt, welcher nicht auf g liegt. Wir definieren die Spiegelung $s_g(P)$ wie folgt:

- (1) Man kann leicht zeigen, dass es eine eindeutig bestimmte hyperbolische Gerade h durch P gibt, welche g orthogonal¹⁸³ in einem Punkt S schneidet.¹⁸⁴
- (2) Wir definieren dann $Q = s_g(P)$ als den eindeutig bestimmten Punkt auf h , mit $d_{\mathbb{H}}(S, Q) = d_{\mathbb{H}}(S, P)$ und mit $P \neq Q$.

Für $P \in g$ definieren wir wiederum $s_g(P) = P$. Diese beide Konstruktionen von Spiegelungen sind in Abbildung 81 skizziert.

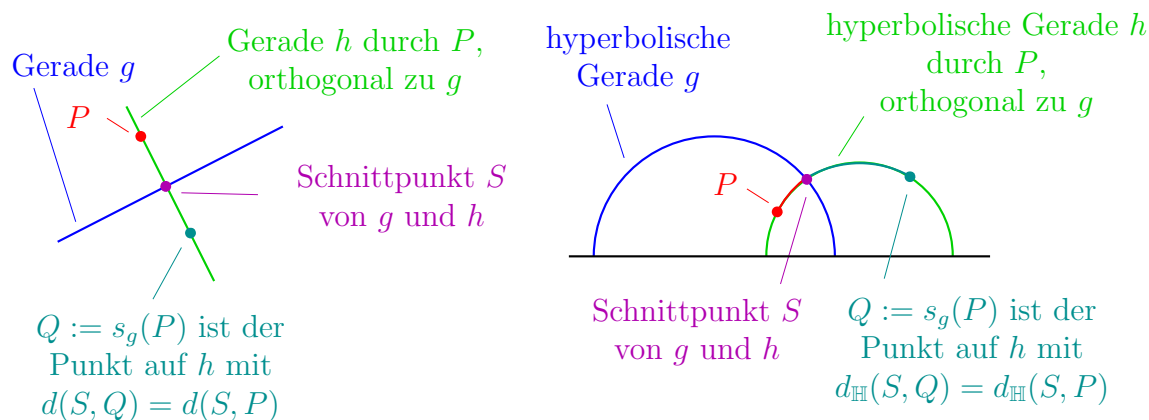


ABBILDUNG 81. Spiegelung in der hyperbolischen Gerade g .

Lemma 14.10. *Es seien g und h zwei hyperbolische Geraden und es sei Φ eine Möbiustransformation mit $\Phi(g) = h$. Dann gilt*

$$s_g = \Phi^{-1} \circ s_h \circ \Phi.$$

Beweis. Das Lemma folgt aus der in Satz 14.3 bewiesenen Tatsache, dass Möbiustransformationen Isometrien sind, insbesondere winkelerhaltend sind.

Genauer gesagt, es sei P ein Punkt, welcher nicht auf g liegt. Wir müssen zeigen, dass

$$\Phi(s_g(P)) = s_h(\Phi(P)).$$

Es sei k die hyperbolische Gerade durch P , welche g orthogonal schneidet. Es sei S der Schnittpunkt von k und g . Dann ist $s_g(P) := Q \neq P$ der Punkt, welcher auf k liegt mit $d_{\mathbb{H}}(P, S) = d_{\mathbb{H}}(Q, S)$. Also ist $\Phi(Q)$ der Punkt, welcher auf $\Phi(k) =: l$ liegt mit $\Phi(Q) \neq \Phi(P)$ und mit $d_{\mathbb{H}}(\Phi(P), \Phi(S)) = d_{\mathbb{H}}(\Phi(Q), \Phi(S))$. Aber $\Phi(k)$ ist orthogonal zu $\Phi(g) = h$ und geht durch $\Phi(P)$. Also erhalten wir per Definition von der Spiegelung $s_h(\Phi(P))$ die Gleichheit $\Phi(Q) = s_h(\Phi(P))$. $\square$

¹⁸³Hierbei meinen wir mit “orthogonal”, dass sich die hyperbolischen Geraden, welche ja euklidische Halbkreise oder Halbgeraden sind, euklidisch orthogonal schneiden.

¹⁸⁴Die Beweisidee ist immer die gleiche. Man beweist die Aussage zuerst für die hyperbolische Gerade $g = \{iy \mid y \in \mathbb{R}_{>0}\}$ und führt dann den allgemeinen Fall mithilfe einer Möbiustransformation auf den Spezialfall zurück.

Wir können die Spiegelungen an hyperbolischen Geraden auch explizit angeben.

Lemma 14.11.

(1) Für die Gerade $h = \{z \in \mathbb{H} \mid \operatorname{Re}(z) = a\}$ ist die Spiegelung in h gerade die Abbildung

$$\begin{aligned} s_h: \mathbb{H} &\rightarrow \mathbb{H} \\ x + iy &\mapsto 2a - x + iy, \end{aligned}$$

dies ist also die übliche Spiegelung an der euklidischen Gerade.

(2) Für die Gerade $h = \{z \in \mathbb{H} \mid |z-s| = r\}$ ist die Spiegelung in h gerade die Abbildung

$$\begin{aligned} s_h: \mathbb{H} &\rightarrow \mathbb{H} \\ z &\mapsto \frac{r^2}{\bar{z}-s} + s, \end{aligned}$$

dies ist also die übliche Spiegelung am entsprechenden euklidischen Kreis.

Beweis. Für die Gerade $h = \{z \in \mathbb{H} \mid \operatorname{Re}(z) = 0\}$ ist der Beweis in Abbildung 82 skizziert. Der allgemeine Beweis von (1) verläuft ganz analog.

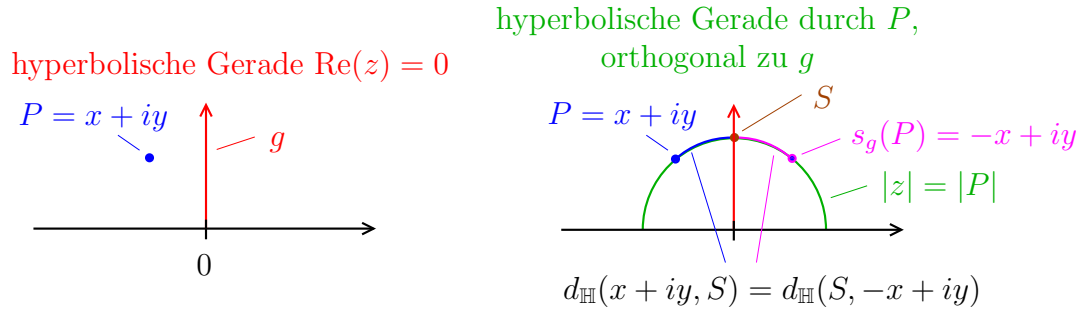


ABBILDUNG 82.

Wir betrachten nun die Gerade $g = \{z \in \mathbb{H} \mid |z| = 1\}$ mit Endpunkten ± 1 . Es sei

$$\Phi(z) = \frac{z+1}{z-1}$$

die Möbiustransformation, welche nach Satz 14.6 die Eigenschaft besitzt, dass $\Phi(g) = h$. Dann gilt

$$\begin{aligned} s_g(z) &= (\Phi^{-1} \circ s_h \circ \Phi)(z) &&= \Phi\left(s_h\left(\frac{z+1}{z-1}\right)\right) &&= \Phi\left(-\frac{\bar{z}+1}{\bar{z}-1}\right) &&= \frac{-\frac{\bar{z}+1}{\bar{z}-1} + 1}{-\frac{\bar{z}+1}{\bar{z}-1} - 1} = \frac{1}{\bar{z}}. \\ \uparrow &&&\uparrow &&\uparrow && \\ \text{Lemma 14.10} &&&\text{denn } \Phi^{-1} = \Phi &&\text{wie oben berechnet ist } s_h(w) = -\bar{w} \end{aligned}$$

Der allgemeine Fall von (2) wird ganz analog behandelt. $\square$

Satz 14.12.

(1) Jede Spiegelung ist eine orientierungsumkehrende Isometrie von $\mathbb{H}$.

- (2) Jede Möbiustransformation ist die Verknüpfung von einer geraden Anzahl an Spiegelungen.¹⁸⁵

Bemerkung. Wir hatten auf Seite 174 schon erwähnt, dass jede orientierungserhaltende Isometrie von $\mathbb{H}$ eine Möbiustransformation ist. Satz 14.12 besagt also insbesondere, dass jede orientierungserhaltende Isometrie von $\mathbb{H}$ die Verknüpfung von einer geraden Anzahl an Spiegelungen ist. Genau die gleiche Aussage gilt auch für den euklidischen Raum: jede orientierungserhaltende Isometrie von $\mathbb{R}^2$ ist die Verknüpfung von einer geraden Anzahl an euklidischen Spiegelungen.¹⁸⁶

Beweis.

- (1) Man kann problemlos mithilfe der Definitionen beweisen, dass die Spiegelung in der positiven y -Achse, d.h. die Abbildung $s_h(x + iy) := -x + iy$ eine orientierungsumkehrende Isometrie von $\mathbb{H}$ ist. Es sei nun g eine beliebige hyperbolische Gerade. Nach Lemma 14.7 existiert eine Möbiustransformation Φ mit $\Phi(g) = h$. Dann gilt nach Lemma 14.10, dass $s_g = \Phi^{-1} \circ s_h \circ \Phi$. Also ist s_g die Verknüpfung von s_h mit zwei Möbiustransformationen ist. Es folgt nun aus Satz 9.16, dass jede Spiegelung eine orientierungsumkehrende Isometrie von $\mathbb{H}$ ist.
- (2) Nach Satz 14.2 ist jede Möbiustransformation eine Verknüpfung von Streckungen, Verschiebungen und Inversionen. Es genügt nun also zu zeigen, dass jede dieser Abbildungen die Verknüpfung von zwei Spiegelungen ist.

Für eine euklidische Gerade g , welche orthogonal zur x -Achse ist und für einen euklidischen Kreis g , dessen Mittelpunkt auf der x -Achse liegt bezeichnen wir nun kurzzeitig mit $s(g)$ die Spiegelung an der entsprechenden hyperbolischen Gerade.

Man kann nun mithilfe von Lemma 14.11 leicht folgende Aussagen zeigen:

- (a) für $d \in \mathbb{R}$ ist

$$(z \mapsto z + d) = s(\operatorname{Re}(z) = d) \circ s(\operatorname{Re}(z) = \tfrac{1}{2}d),$$

- (b) für $r \in \mathbb{R}_{>0}$ ist

$$(z \mapsto rz) = s(|z| = r) \circ s(|z| = \sqrt{r}),$$

- (c) es ist

$$(z \mapsto -\tfrac{1}{z}) = s(\operatorname{Re}(z) = 0) \circ s(|z| = 1). \quad \square$$

14.2. Das Poincaré-Model für den hyperbolischen Raum. In diesem Kapitel führen wir das Poincaré-Model $\mathbb{D}$ für den hyperbolischen Raum ein. Dies ist eine Riemannsche Untermannigfaltigkeit, welche isometrisch zu $\mathbb{H}$ ist. Je nach Problemstellung ist es manchmal einfacher mit $\mathbb{H}$ oder mit $\mathbb{D}$ zu arbeiten.

Definition.

¹⁸⁵Ist es möglich, dass eine Möbiustransformation als Verknüpfung von einer ungeraden Anzahl an Spiegelungen geschrieben werden kann? Was besagt Satz 9.16 in diesem Fall?

¹⁸⁶Dann stellt sich natürlich die Frage, warum diese Aussage im euklidischen Fall wahr ist.

- (1) Das *Poincaré-Model für den hyperbolischen Raum* ist die Riemannsche Untermannigfaltigkeit, welche gegeben ist durch die offene Scheibe

$$\mathbb{D} = \{z \in \mathbb{C} \mid |z|^2 < 1\}$$

zusammen mit der Riemannschen Struktur g , welche an jedem Punkt $z \in \mathbb{D}$ auf dem Tangentialraum $T_z\mathbb{D} = \mathbb{R}^2$ gegeben ist durch

$$g_z: \mathbb{R}^2 \times \mathbb{R}^2 \rightarrow \mathbb{R} \\ (v, w) \mapsto \frac{2}{(1 - |z|^2)^2} \langle v, w \rangle.$$

- (2) Eine *Gerade in $\mathbb{D}$* ist eine Teilmenge von $\mathbb{D}$ der folgenden Form:
- (a) der Schnitt $g \cap \mathbb{D}$, wobei g eine euklidische Gerade durch 0 ist,
 - (b) der Schnitt $K \cap \mathbb{D}$, wobei K ein euklidischer Kreis ist, welcher den Kreis $S^1 = \partial\mathbb{D}$ orthogonal schneidet.

Die Schnittpunkte von g beziehungsweise K mit S^1 sind die *Endpunkte der Gerade*.

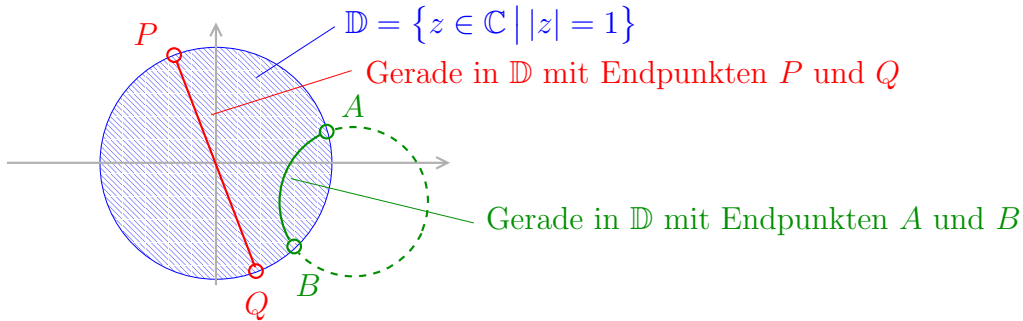


ABBILDUNG 83. Beispiele von Geraden im Poincaré-Model.

Satz 14.13. *Die Abbildungen*

$$\begin{aligned} \Phi: \mathbb{H} &\rightarrow \mathbb{D} & \text{und} & & \Psi: \mathbb{D} &\rightarrow \mathbb{H} \\ z &\mapsto \frac{z-i}{z+i} & & & w &\mapsto \frac{i+iw}{1-w} \end{aligned}$$

sind zueinander inverse orientierungserhaltende Isometrien. Diese induzieren zudem Bijektionen zwischen den hyperbolischen Geraden in $\mathbb{H}$ und den Geraden in $\mathbb{D}$.¹⁸⁷

Beweis. In Analysis III Lemma 13.2 hatten wir schon bewiesen, dass $\Psi: \mathbb{D} \rightarrow \mathbb{H}$ ein Diffeomorphismus ist. Wir zeigen nun, dass Ψ eine Isometrie ist.

Wir beginnen mit ein paar Vorbemerkungen. Für $v, w \in \mathbb{C} = \mathbb{R}^2$ ist $\langle v, w \rangle = \operatorname{Re}(v \cdot \bar{w})$. Zudem entspricht das Differential $D_z\Psi$ gerade der Ableitung $\Psi'(z) \in \mathbb{C}$.

¹⁸⁷Eine hyperbolische Gerade g mit Endpunkten P und Q wird hierbei auf eine Gerade in $\mathbb{D}$ mit Endpunkten $\Phi(P)$ und $\Phi(Q)$ geschickt. Hierbei bezeichnen wir mit Φ die offensichtliche Fortsetzung von der bisherigen Abbildung Φ zu einer Abbildung $\mathbb{C} \cup \{\infty\} \rightarrow \mathbb{C} \cup \{\infty\}$.

Wir beginnen nun mit dem eigentlichen Beweis, dass Ψ eine Isometrie ist. Es sei also $z = x + iy \in \mathbb{D}$ und es seien $v, w \in T_z \mathbb{D} = \mathbb{R}^2$. Wir bezeichnen mit g die Riemannsche Struktur auf $\mathbb{D}$ und wir bezeichnen mit h die Riemannsche Struktur auf $\mathbb{H}$. Zudem ist

$$\Psi'(z) = \frac{1+i}{(1-z)^2}.$$

Für $z \in \mathbb{D}$ ist nun

$$\begin{aligned} (\Psi^*h)_z(v, w) &= h_{\Psi(z)}(\Psi'(z)v, \Psi'(z)w) = \frac{1}{\operatorname{Im}(\Psi(z))^2} \operatorname{Re}(\Psi'(z)v \cdot \overline{\Psi'(z)w}) \\ &= \frac{1}{\operatorname{Im}\left(\frac{1+i}{1-z}\right)^2} \Psi'(z) \overline{\Psi'(z)} \cdot \operatorname{Re}(v \cdot \bar{w}) \\ &= \frac{1}{-\frac{1}{4}\left(\frac{-i-i\bar{z}}{1-\bar{z}} - \frac{i+iz}{1-z}\right)^2} \frac{2}{(1-z)^2(1-\bar{z})^2} \cdot \operatorname{Re}(v \cdot \bar{w}) \\ &\quad \uparrow \\ &\text{denn } \operatorname{Im}(w) = \frac{1}{2i}(w + \bar{w}) \\ &= \frac{-8}{\left((-i-i\bar{z})(1-z) - (i+iz)(1-\bar{z})\right)^2} \cdot \operatorname{Re}(v \cdot \bar{w}) \\ &= \frac{-8}{(-i+iz-i\bar{z}+iz\bar{z}-i-iz+i\bar{z}+iz\bar{z})^2} \cdot \operatorname{Re}(v \cdot \bar{w}) \\ &= \frac{-8}{(-2i+iz\bar{z})^2} \cdot \operatorname{Re}(v \cdot \bar{w}) = \frac{2}{(1-|z|^2)^2} \cdot \operatorname{Re}(v \cdot \bar{w}) = g_z(v, w). \end{aligned}$$

Den Beweis der Aussage über die Geraden ist ganz ähnlich zum Beweis von Satz 14.6. Wir überlassen die Ausführung des Beweises als freiwillige Übungsaufgabe. $\square$

Nachdem Isometrien unter anderem auch Geodäten in Geodäten überführen, erhalten wir aus Satz 14.8 und Satz 14.13 folgendes Korollar.

Korollar 14.14. *Es seien P und Q zwei verschiedene Punkte auf $\mathbb{D}$. Jede Geodäte von P nach Q verläuft auf der eindeutigen Gerade durch P und Q .*

Die Definition von Spiegelung in einer Gerade überträgt sich auf das Poincaré-Modell für den hyperbolischen Raum. Ganz ähnlich zu Lemma 14.11 wir folgendes Lemma beweisen.

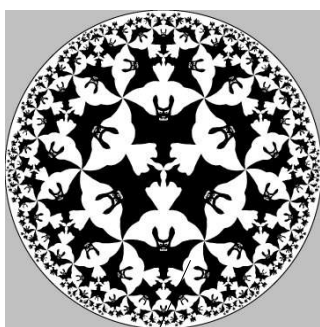
Lemma 14.15. *Es sei $z \in S^1$ und es sei $s_z: \mathbb{D} \rightarrow \mathbb{D}$ die Spiegelung an der hyperbolischen Gerade $\{tz \mid t \in (-1, 1)\}$. Dann ist s_z gegeben durch die übliche euklidische Spiegelung.*

Zum Abschluß des Kapitels betrachten wir Abbildung 84. Diese zeigt eine vereinfachte Version von Escher's Circle Limit IV. Das Bild besteht letztendlich aus Sechsecken in $\mathbb{D}$, welche durch Spiegelungen in Geraden ineinander übergeführt werden können. In Abbildung 85 sehen wir noch weitere Graphiken von Escher, welche mit Methoden der hyperbolischen Geometrie entworfen wurden. Mehr über die Mathematik hinter Escher's Bildern kann man beispielsweise auf

<http://www.math.cornell.edu/~mec/Winter2009/Mihai/index.html>

und

<http://www.math-art.eu/Documents/pdfs/Dunham.pdf>



Escher: Circle Limit IV

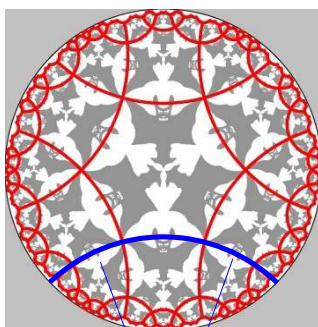
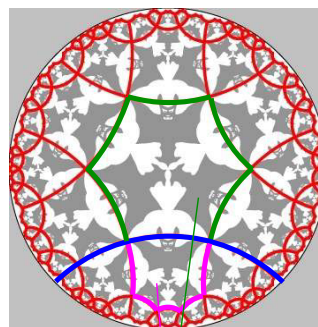
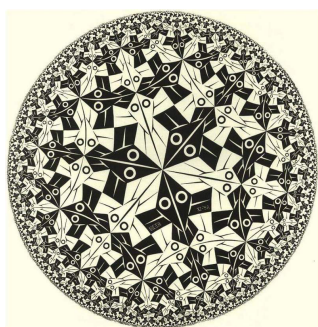
Gerade g in $\mathbb{D}$ die beiden Sechsecke
sind Spiegelbilder in g

ABBILDUNG 84.



Circle Limit I



Circle Limit III



Circle Limit IV

ABBILDUNG 85.

nachlesen. Es gibt auch eine lange Liste von Videos zu hyperbolischer Geometrie, beispielsweise

https://www.youtube.com/watch?v=eGEQ_UuQtYs

15. VOLUMEN UND KRÜMMUNG AUF RIEMANNSCHEN UNTERMANNIGFALTIGKEITEN

Wir führen nun folgende Bezeichnungen für Riemannsche Untermannigfaltigkeiten ein:

- (1) $\mathbb{E}^n$ bezeichnet $\mathbb{R}^n$ zusammen mit der euklidischen Metrik.
- (2) $\mathbb{S}^2$ bezeichnet die Sphäre S^2 zusammen mit der sphärischen Metrik, d.h. mit der üblichen Riemannschen Struktur, welche durch das Skalarprodukt in $\mathbb{R}^3$ gegeben ist.
- (3) Wie zuvor bezeichnet $\mathbb{H}$ die obere Halbebene zusammen mit der hyperbolischen Metrik und $\mathbb{D}$ bezeichnet die offene Scheibe $D_1(0) \subset \mathbb{C} = \mathbb{R}^2$ zusammen mit der Riemannschen Struktur, welche wir auf Seite 182 eingeführt hatten.

In diesem Kapitel werden wir hauptsächlich den Flächeninhalt und die Krümmung auf den obigen Riemannschen Untermannigfaltigkeiten studieren. In diesem Kapitel werden wir an manchen Stellen der Verständlichkeit halber etwas anschaulich argumentieren.

15.1. Volumen auf Riemannschen Untermannigfaltigkeiten. Es sei $f: M \rightarrow \mathbb{R}$ eine Funktion auf einer k -dimensionalen Untermannigfaltigkeit M . Wir nehmen zuerst an, dass es eine Karte $\Phi: U \rightarrow V$ gibt, so dass $\text{Träger}(f) \subset U$. Wir bezeichnen mit $\Psi: V \rightarrow U$ die Umkehrabbildung. Wir hatten, siehe Seiten 36 und 40, das Integral von f auf M definiert als

$$\int_M f := \int_{V \cap E_k} f(\Psi(x)) \cdot \sqrt{\det(\langle \Psi_* e_i, \Psi_* e_j \rangle_{i,j=1,\dots,k})}.$$

Mithilfe von messbaren Zerlegung hatten wir den allgemeinen Fall, d.h. das Integral einer beliebigen Funktionen, auf diesen Spezialfall zurückgeführt. Für eine Riemannsche Untermannigfaltigkeit können wir nun das Skalarprodukt durch die Riemannsche Struktur ersetzen. Mit anderen Worten wir definieren

$$\int_{(M,g)} f := \int_{V \cap E_k} f(\Psi(x)) \cdot \sqrt{\det(g_{\Psi(x)}(\Psi_* e_i, \Psi_* e_j)_{i,j=1,\dots,k})}.$$

Das Argument von Satz 3.16 zeigt, dass diese Definition nicht von der Wahl der Karte abhängt. Wir verallgemeinern diese Definition, ganz analog zur Diskussion in den Kapiteln 3.5 und 3.7, mithilfe von messbaren Zerlegungen auf Funktionen, deren Träger nicht in einer Karte enthalten ist. Für eine Teilmenge $A \subset M$ definieren wir insbesondere

$$\text{Vol}_{(M,g)}(A) := \int_{(M,g)} \chi_A,$$

vorausgesetzt das Integral auf der rechten Seite ist definiert. Der folgende Satz folgt leicht aus den Definitionen.¹⁸⁸

Satz 15.1. *Es sei $\varphi: (M, g) \rightarrow (N, h)$ eine Isometrie zwischen Riemannschen Untermannigfaltigkeiten.*

¹⁸⁸Der Beweis der ersten Aussage ist ganz ähnlich dem Beweis der Übungsaufgabe 4 in Übungsblatt 9.

(1) Für jede Funktion $f: N \rightarrow \mathbb{R}$ gilt

$$\int_{(M,g)} \varphi \circ f = \int_{(N,h)} f.$$

(2) Für eine Teilmenge A von M gilt

$$\text{Vol}_{(M,g)}(A) = \text{Vol}_{(N,h)}(\varphi(A)).$$

Wir betrachten nun noch etwas ausführlicher den Flächeninhalt von Teilmengen von $\mathbb{H}$. Für eine Teilmenge A von $\mathbb{H}$ bezeichnen wir $\text{Vol}_{\mathbb{H}}(A)$ als den *hyperbolischen Flächeninhalt*. Für $A \subset \mathbb{H}$ gilt

$$\text{hyperbolischer Flächeninhalt} = \int_{\mathbb{H}} \chi_A(y) \cdot \frac{1}{y^2} dx dy = \int_A \frac{1}{y^2} dx dy.$$

folgt aus der Karte $\Phi = \text{id}$ und $g_{(x,y)}(e_i, e_j) = \frac{1}{y^2} \langle e_i, e_j \rangle$

Beispiel. Wir betrachten die Teilmenge Δ von $\mathbb{H}$, welche in Abbildung 86 skizziert ist, wobei $0 < \theta < \varphi < \pi$. Es ist

$$\begin{aligned} \text{hyperbolischer} \\ \text{Flächeninhalt von } \Delta &= \int_{\Delta} \frac{1}{y^2} dx dy = \int_{x=\cos(\varphi)}^{x=\cos(\theta)} \int_{y=\sqrt{1-x^2}}^{\infty} \frac{1}{y^2} dy dx = \int_{x=\cos(\varphi)}^{x=\cos(\theta)} \left[\frac{-1}{y} \right]_{y=\sqrt{1-x^2}}^{y=\infty} dx \\ &\quad \uparrow \text{Satz von Fubini} \\ &= \int_{x=\cos(\varphi)}^{x=\cos(\theta)} \frac{1}{\sqrt{1-x^2}} dx = \int_{u=\varphi}^{u=\theta} -du = \varphi - \theta. \\ &\quad \uparrow \text{Substitution } u = \arccos(x), \\ &\quad \text{wobei } \frac{du}{dx} = -\frac{1}{\sqrt{1-x^2}} \end{aligned}$$

Wir sehen also, dass der hyperbolische Flächeninhalt von Δ endlich ist, obwohl der euklidische Flächeninhalt von Δ unendlich ist.

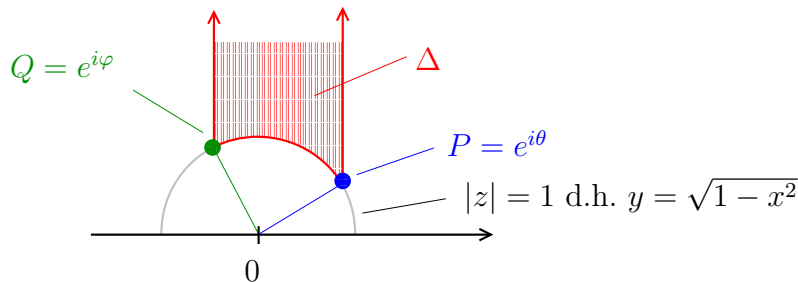


ABBILDUNG 86.

15.2. Winkel und Dreiecke auf Riemannschen Untermannigfaltigkeiten.

Definition. Es sei (M, g) eine Riemannsche Untermannigfaltigkeit, es sei $P \in M$ und es seien $v, w \in T_P M \setminus \{0\}$. Dann gibt es genau ein $\varphi \in [0, \pi]$ mit

$$\cos(\varphi) = \frac{g(v, w)}{\sqrt{g(v, v) \cdot g(w, w)}},$$

$\in [-1, 1]$ nach der Cauchy-Schwartz
Ungleichung

welches wir als den *Winkel zwischen v und w bezeichnen*.

Beispiele.

- (1) Es sei $M = \mathbb{E}$, $P \in \mathbb{E}^2 = \mathbb{R}^2$ und $v, w \in T_P \mathbb{E}^2 = \mathbb{R}^2$. Es sei $\varphi \in [0, \pi]$ der Winkel zwischen v und w . Dies bedeutet gerade, dass

$$\langle v, w \rangle = \cos(\varphi) \cdot \|v\| \cdot \|w\|.$$

Wir erhalten also die übliche Definition für den Winkel zwischen zwei Vektoren in $\mathbb{R}^2$.

- (2) Für $M = \mathbb{S}^2$ ist die Riemannsche Struktur gegeben, durch die Einschränkung des Skalarprodukts von $\mathbb{R}^3$ auf den jeweiligen Tangentialraum. Es folgt also aus dem gleichen Argument wie in (1), dass der Winkel zwischen zwei Tangentialvektoren gerade der übliche euklidische Winkel ist.
- (3) Es sei $U \subset \mathbb{R}^2$ eine offene Teilmenge und es sei $\rho: U \rightarrow \mathbb{R}_{>0}$ eine nichtnegative glatte Funktion. Wir betrachten dann auf U die Riemannsche Struktur $g = g(\rho)$, welche wir schon auf Seite 161 eingeführt hatten, und welche an einem Punkt $P \in U$ gegeben ist durch

$$\begin{aligned} g(\rho)_P: T_P U \times T_P U &\rightarrow \mathbb{R} \\ (v, w) &\mapsto \rho(P) \cdot \langle v, w \rangle. \end{aligned}$$

Für $P \in U$ und $v, w \in T_P U = \mathbb{R}^2$ gilt dann

$$\frac{g(v, w)}{\sqrt{g(v, v) \cdot g(w, w)}} = \frac{\rho(P) \langle v, w \rangle}{\sqrt{\rho(P) \langle v, v \rangle \cdot \rho(P) \langle w, w \rangle}} = \frac{\langle v, w \rangle}{\sqrt{\langle v, v \rangle \cdot \langle w, w \rangle}} = \frac{\langle v, w \rangle}{\|v\| \cdot \|w\|}.$$

$\uparrow$
denn $\rho(P)$ kürzt sich weg

Also ist

Winkel zwischen v und w in $(U, g(\rho))$ = euklidischer Winkel zwischen v und w .

Wenn wir diese Beobachtung auf den hyperbolischen Raum $\mathbb{H}$ anwenden, erhalten wir, dass für $P \in \mathbb{H}$ und $v, w \in T_P \mathbb{H} = \mathbb{R}^2$ gilt, dass

hyperbolischer Winkel zwischen v und w = euklidischer Winkel zwischen v und w .

Genau das gleiche Argument zeigt auch, dass in $\mathbb{D}$ der hyperbolische Winkel ebenfalls gerade dem euklidischen Winkel entspricht.

Das nächste Lemma folgt sofort aus den Definitionen.

Lemma 15.2. Jede lokale Isometrie $\varphi: (M, g) \rightarrow (N, h)$ zwischen zwei Riemannschen Untermannigfaltigkeiten ist winkelerhaltend¹⁸⁹, d.h. für alle $P \in M$ und alle $v, w \in T_P M$ gilt

Winkel zwischen v und w in $T_P M$ = Winkel zwischen $\varphi_* v$ und $\varphi_* w$ in $T_{\varphi(P)} N$.

Definition. Es sei (M, g) eine Riemannsche Untermannigfaltigkeit.

- (1) Es seien $A, B \in M$. Eine *Strecke* von A nach B ist das Bild einer Geodäte von A nach B .
- (2) Es sei b eine Strecke von A nach B und es sei c eine Strecke von A nach C . Wir bezeichnen mit β und γ die entsprechenden Geodäten. Wir definieren

Winkel zwischen b und c := Winkel zwischen $\beta'(0)$ und $\gamma'(0)$.

Beispiel. Für $(M, g) = \mathbb{E}^2$ erhalten wir nach Satz 12.4 den üblichen Begriff von einer Strecke und es gibt genau eine Strecke zwischen zwei Punkten. In $\mathbb{S}^2$ gibt es nach Satz 13.8 für $A \neq -B$ ebenfalls genau eine Strecke von A nach B . Wenn A und B Antipoden sind, d.h. wenn $A = -B$, dann gibt es jedoch unendlich viele Strecken von A nach B . Auf $\mathbb{H}$ gibt es nach Satz 14.9 immer nur genau eine Strecke zwischen zwei Punkten.

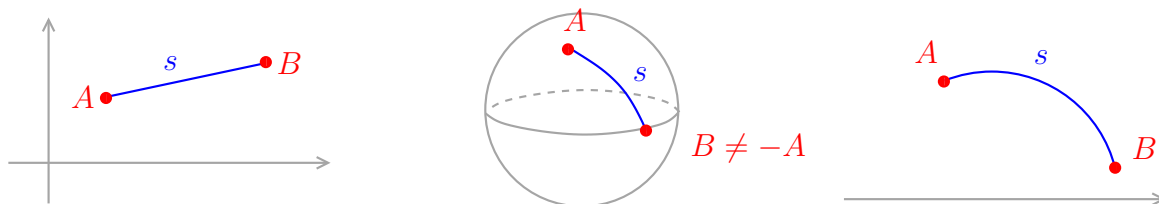


ABBILDUNG 87. Beispiele von Strecken in $\mathbb{E}^2$, $\mathbb{S}^2$ und $\mathbb{H}$.

Definition. Es sei (M, g) eine Riemannsche Untermannigfaltigkeit. Ein *Dreieck* Δ_{ABC} ist eine offene¹⁹⁰ Teilmenge von M mit folgender Eigenschaft: Der Rand¹⁹¹ kann geschrieben werden als

$$\partial\Delta_{ABC} = a \cup b \cup c,$$

wobei gilt

- (1) a ist eine Strecke von B nach C ,
- (2) b ist eine Strecke von A nach C ,
- (3) c ist eine Strecke von A nach B ,
- (4) es ist $a \cap b = \{C\}$, $a \cap c = \{B\}$ und $b \cap c = \{A\}$.

¹⁸⁹In der Literatur werden winkelerhaltende Abbildungen gerne auch *konform* genannt.

¹⁹⁰Warum habe ich das Dreieck nicht als kompakte Teilmenge mit den entsprechenden Eigenschaften definiert?

¹⁹¹Hierbei bezeichnen wir mit $\partial\Delta_{ABC}$ den topologischen Rand von der Teilmenge Δ_{ABC} von M

Wir bezeichnen A, B und C als die *Eckpunkte von Δ_{ABC}* , wir bezeichnen a, b, c als die *Seiten von Δ_{ABC}* und wir bezeichnen die Winkel zwischen den Strecken als die *Innenwinkel von Δ_{ABC}* .

Beispiel. Für $(M, g) = \mathbb{E}^2$ erhalten wir den üblichen Begriff von einem Dreieck. In $\mathbb{S}^2$ wird ein Dreieck durch Großkreise beschrieben und in $\mathbb{H}$ wird ein Dreieck durch hyperbolische Geraden beschrieben. Beispiele von solchen Dreiecken werden in Abbildung 88 skizziert.

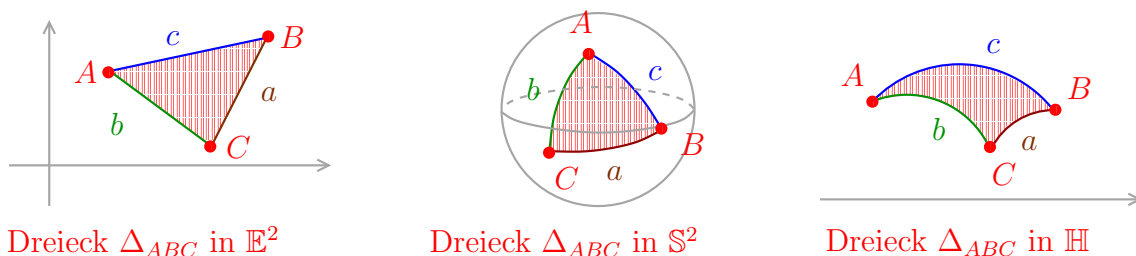


ABBILDUNG 88.

Satz 15.3. *Es sei Δ_{ABC} ein Dreieck auf $\mathbb{S}^2$, welches sich in einer offenen Halbsphäre befindet. Dann gilt*

$$\text{Flächeninhalt von } \Delta_{ABC} = \text{Summe der Innenwinkel} - \pi.$$

Bemerkung.

- (1) Satz 15.3 besagt also, dass der Flächeninhalt von einem Dreieck auf $\mathbb{S}^2$ durch seine Innenwinkel festgelegt ist, oder mit anderen Worten, aus seinen Innenwinkeln bestimmt werden kann.
- (2) Die vorherige Bemerkung auch auf der Erde und gilt insbesondere für alle Dreiecke, welche man in der Schule gezeichnet und abgemessen hat. Diese Dreiecke sind jedoch im Vergleich zur Erdoberfläche so klein, dass man den Unterschied zwischen der Summe der Innenwinkel und π nicht messen kann.

Beispiel. Wir betrachten ein Dreieck, welches aufgespannt wird durch den “Nordpol” $(0, 0, 1)$ und zwei Punkte auf dem “Äquator” $\{(x, y, z) \in S^2 \mid z = 0\}$, siehe Abbildung 89. Es sei α der Innenwinkel am Nordpol. Die Innenwinkel an den Punkten auf dem Äquator sind gerade $\frac{\pi}{2}$. Wir sehen also, dass der Flächeninhalt des Dreiecks gegeben ist durch

$$\text{Summe der Innenwinkel} - \pi = \alpha + \frac{\pi}{2} + \frac{\pi}{2} - \pi = \alpha.$$

Das gleiche Ergebnis hätten wir natürlich auch mit der Rechenmethode auf Seite 47 bestimmen können.

Beweis. Wir beginnen den Beweis mit einer einfachen Vorbemerkung. Ein *Zweieck mit Winkel $\alpha \in (0, \pi)$* ist eine Teilmenge von S^2 , welche durch zwei Großkreise begrenzt wird, welche den Winkel α einschlagen. Diese Definition wird in Abbildung 90 illustriert. Mit fast

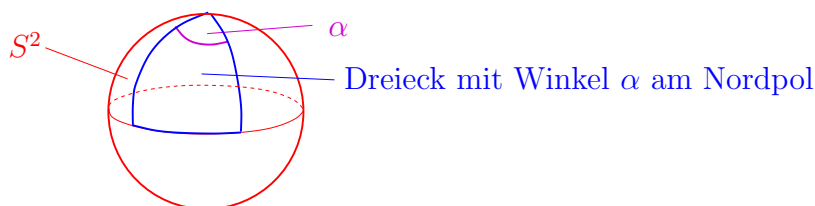
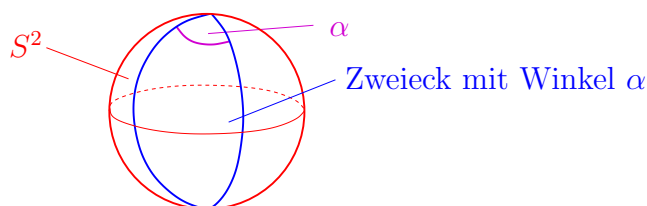


ABBILDUNG 89.

ABBILDUNG 90. Ein Zweieck mit Innenwinkel α auf S^2 .

dem gleichen Argument wie auf Seite 47 sieht man, dass

$$\text{Flächeninhalt von Zweieck mit Winkel } \alpha = 2\alpha.$$

Wir beginnen nun mit dem eigentlichen Beweis von Satz 15.3. Es sei also $\Delta = \triangle_{ABC}$ ein Dreieck auf S^2 . Die drei Großkreise, welche durch die Seiten des Dreiecks gebildet werden, zerteilen S^2 in folgende acht Dreiecke:¹⁹²

- (1) das ursprüngliche Dreieck Δ ,
- (2) das Dreieck X , welches auf der anderen Seite der Strecke BC liegt,¹⁹³
- (3) das Dreieck Y , welches auf der anderen Seite der Strecke AC liegt,
- (4) das Dreieck Z , welches auf der anderen Seite der Strecke AB liegt,
- (5) die Spiegelungen Δ', X', Y', Z' der Dreiecke Δ, X, Y, Z im Ursprung.

Diese acht Dreiecke sind in Abbildung 91 skizziert. Wir bezeichnen mit $x, y, z, \delta, x', y', z', \delta'$ die Flächeninhalte von $X, Y, Z, \Delta, X', Y', Z', \Delta'$. Offensichtlich ist $\Delta \cup X$ ein Zweieck mit Winkel α .¹⁹⁴ Damit folgt aus der Vorbemerkung, dass $x + \delta = 2\alpha$. Analog gilt $x + \delta = 2\beta$

¹⁹²Zwei Großkreise zerlegen S^2 in vier Zweiecke, der dritte Großkreis zerteilt nun jedes Zweieck in zwei Dreiecke, wir erhalten also insgesamt acht Dreiecke.

¹⁹³Etwas genauer, X ist das Dreieck, welches von B, C und vom Spiegelbild $-A$ aufgespannt wird.

¹⁹⁴Wenn man ganz genau hinschaut, dann sieht man, dass die Aussage so nicht ganz stimmt. Wir hatten Dreiecke und Zweiecke als offene Teilmengen definiert. Also ist die Vereinigung von Δ mit X und mit einer Kante ein Zweieck. Nachdem wir uns nur für Flächeninhalte interessieren ignorieren wir durchgehend Nullmengen in diesem Beweis.

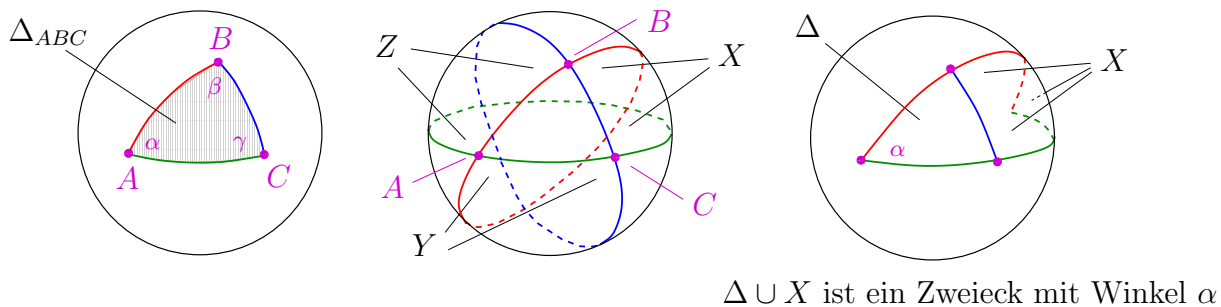


ABBILDUNG 91.

und $z + \delta = 2\gamma$. Es folgt nun, dass

$$\begin{aligned}
 4\pi = \text{Flächeninhalt von } S^2 &= \delta + \delta' + x + x' + y + y' + z + z' &&= 2\delta + 2x + 2y + 2z \\
 &\quad \uparrow && \uparrow \\
 &\text{Zerlegung von } S^2 \text{ in die 8 Dreiecke} && \Delta \text{ und } \Delta' \text{ sind Spiegelbilder, also ist } \delta = \delta', \\
 &&& \text{ganz analog gilt } x = x', y = y', z = z' \\
 &= 2(\underbrace{\delta + x}_{=2\alpha}) + 2(\underbrace{\delta + y}_{=2\beta}) + 2(\underbrace{\delta + z}_{=2\gamma}) - 4\delta &&= 4(\alpha + \beta + \gamma) - 4\delta.
 \end{aligned}$$

Durch Auflösen nach δ erhalten wir jetzt die gewünschte Aussage.¹⁹⁵ □

Wir wenden uns nun den hyperbolischen Dreiecken zu. Auch in diesem Fall erhalten wir eine denkbar einfache Formel.

Satz 15.4. *Es sei Δ_{ABC} ein Dreieck in $\mathbb{H}$. Dann gilt*

$$\text{hyperbolischer Flächeninhalt von } \Delta_{ABC} = \pi - \text{Summe der Innenwinkel.}$$

Bemerkung. Satz 15.4 besagt also insbesondere, dass der Flächeninhalt von einem hyperbolischen Dreieck immer kleiner als π ist.

Wir skizzieren im Folgenden den Beweis von Satz 15.4. Wir erweitern dazu den Begriff von einem Dreieck.

- (1) Es sei $P \in \mathbb{H}$ und $Q \in \mathbb{R} \cup \{\infty\}$. Dann gibt es eine eindeutig bestimmte hyperbolische Gerade g durch P mit Endpunkt Q .¹⁹⁶ Der *Halbstrahl von P nach Q* ist die Menge aller Punkte auf g , welche zwischen P und Q liegen.
- (2) Es seien P und Q zwei verschiedene Punkte in $\mathbb{H}$ und es sei $R \in \mathbb{R} \cup \{\infty\}$. Das *ideale Dreieck Δ_{PQR}* ist definiert als die offene Teilmenge von $\mathbb{H}$, welche begrenzt

¹⁹⁵In der Voraussetzung von Satz 15.3 steht: es sei Δ_{ABC} ein Dreieck auf S^2 , welches sich in einer offenen Halbsphäre befindet. Haben wir diese Voraussetzung irgendwo verwendet? Gilt die Aussage des Satzes auch ohne die Voraussetzung?

¹⁹⁶Warum ist dies der Fall?

ist durch

Strecke von P nach $Q \cup$ Halbstrahl von P nach $R \cup$ Halbstrahl von Q nach R .

Der Innenwinkel am Punkt R ist hierbei per Definition der Nullwinkel.

Die Definitionen werden in Abbildung 92 illustriert.

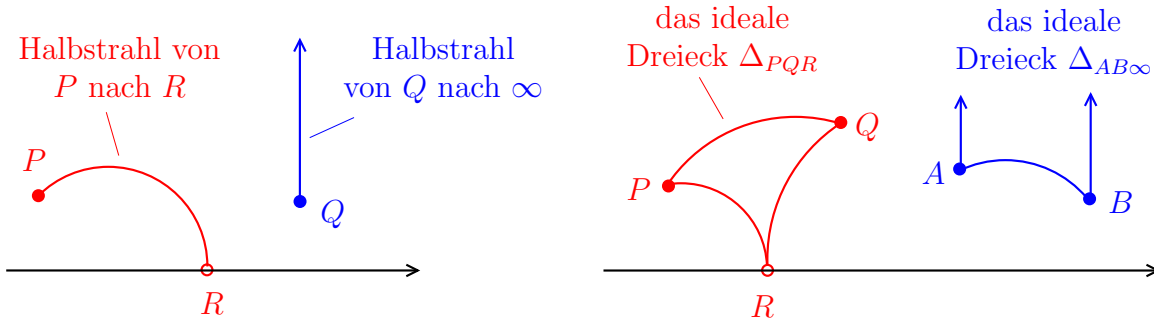


ABBILDUNG 92.

Lemma 15.5. *Es sei Δ_{ABC} ein ideales Dreieck. Dann gibt es eine Möbiustransformation Φ , so dass*

$$\Phi(\Delta_{ABC}) = \Delta_{PQ\infty},$$

wobei P und Q auf der hyperbolischen Gerade $\{z \in \mathbb{H} \mid |z| = 1\}$ liegen.

Beweis. Es sei Δ_{ABC} ein ideales Dreieck. Wir betrachten zuerst den Fall, dass $C = \infty$. Durch Verschieben und Strecken können wir den die hyperbolische Gerade durch A und B in den Einheitskreis $\{z \in \mathbb{C} \mid |z| = 1\}$ überführen. Aber dann haben wir auch Δ_{ABC} in die gewünschte Form übergeführt.

Wir betrachten nun noch den Fall, dass $C \in \mathbb{R}$. Nach einer Verschiebung können wir annehmen, dass $C = 0$. Wir wenden die Inversion $z \mapsto \frac{-1}{z}$ an. Unter dieser Abbildung wird $C = 0$ in den Punkt ∞ übergeführt. Wir fahren jetzt wie im ersten Fall fort. $\square$

Beweis von Satz 15.4. Wir beweisen zuerst folgende etwas einfachere Behauptung.

Behauptung. Es seien $A, B \in \mathbb{H}$ und es sei $D \in \mathbb{R} \cup \{\infty\}$. Dann gilt

$$\text{Flächeninhalt von } \Delta_{ABD} = \pi - \text{Summe der Innenwinkel}.$$

Nach Lemma 15.5 gibt es eine Möbiustransformation Φ , so dass $\Phi(\Delta_{ABD}) = \Delta_{PQ\infty}$, wobei $P = e^{i\theta}$ und $Q = e^{i\varphi}$ mit $0 < \theta < \varphi < \pi$. Möbiustransformationen sind nach Satz 14.3 Isometrien von $\mathbb{H}$ und erhalten daher nach Lemma 13.5 und Satz 15.1 Winkel und Flächeninhalte. Es genügt daher die Aussage für $\Delta_{PQ\infty}$ zu beweisen. Hierbei ist

$$\text{Flächeninhalt von } \Delta_{PQ\infty} = \varphi - \theta = \pi - \text{Summe der Innenwinkel}.$$

$\uparrow$ $\uparrow$
 siehe Berechnung auf Seite 186 siehe Abbildung 93

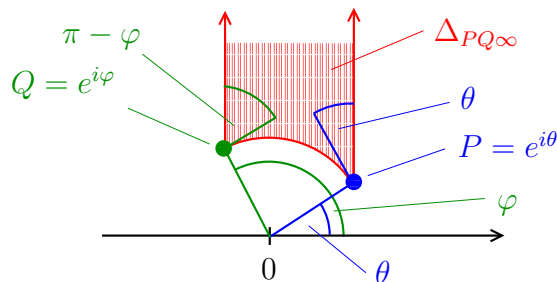


ABBILDUNG 93.

Wir haben damit die Behauptung bewiesen.

Es sei nun Δ_{ABC} ein Dreieck in $\mathbb{H}$. Wir bezeichnen mit R den Punkt auf $\mathbb{R} \cup \{\infty\}$, den wir dadurch erhalten, dass wir die Strecke von A nach C zu einem Halbstrahl verlängern. Wir betrachten die Winkel, welche in Abbildung 94 skizziert sind. Wir erhalten nun, dass

$$\begin{aligned} \text{Flächeninhalt}(\Delta_{ABC}) &= \text{Flächeninhalt}(\Delta_{ABR}) - \text{Flächeninhalt}(\Delta_{CBR}) \\ &= (\pi - \alpha - (\beta + \beta')) - (\pi - \gamma' - \beta') = -\alpha - \beta + \gamma' = \pi - \alpha - \beta - \gamma. \end{aligned}$$

$\uparrow$
 obige Behauptung angewandt auf die
 idealen Dreiecke Δ_{ABR} und Δ_{CBR}

$\uparrow$
 denn $\gamma' + \gamma = \pi$

□

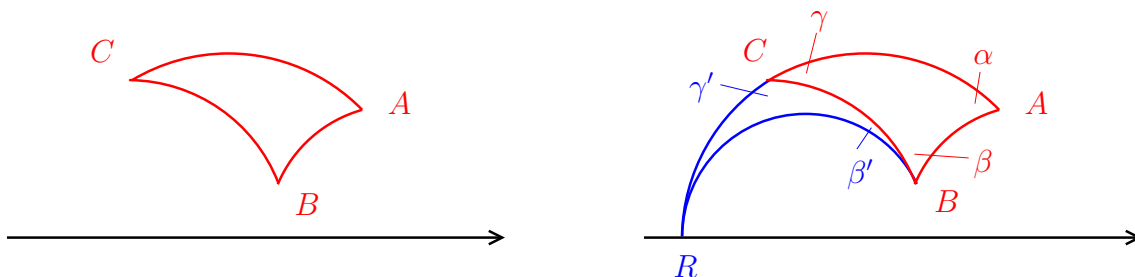


ABBILDUNG 94.

Zum Abschluß der Diskussion von Dreiecken formulieren wir folgenden vertrauten Satz.

Satz 15.6. *Es seien Δ_{ABC} und Δ_{DEF} zwei Dreiecke in $\mathbb{H}$ mit $\angle ABC = \angle DEF > 0$ und mit $d(B, A) = d(E, D)$ und $d_{\mathbb{H}}(B, C) = d_{\mathbb{H}}(E, F)$. Dann gibt es genau eine Möbiustransformation Φ mit $\Phi(B) = E$ und $\Phi(\Delta_{ABC}) = \Delta_{DEF}$.*

Beweisskizze. Wir sagen ein Dreieck Δ_{PQR} ist *positiv orientiert*, wenn die Tangentialvektoren der Strecken PQ und QR am Punkt Q eine positive Basis von $T_Q\mathbb{H} = \mathbb{R}^2$ bilden.

Es seien nun Δ_{ABC} und Δ_{DEF} zwei Dreiecke in $\mathbb{H}$ mit $\angle ABC = \angle DEF > 0$ und mit $d(B, A) = d(E, D)$ und $d(B, C) = d(E, F)$. Indem wir notfalls A und C beziehungsweise D

und F vertauschen können wir o.B.d.A. annehmen, dass Δ_{ABC} und Δ_{DEF} positiv orientiert sind.

Es genügt nun den Fall zu betrachten¹⁹⁷, dass $E = i$, dass $D = iy$ für $y > 1$, und dass zudem $F \in \mathbb{H}$ mit $\operatorname{Re}(C) < 0$. Wir zeigen zuerst, dass es eine Möbiustransformation Φ mit $\Phi(B) = E$ und $\Phi(\Delta_{ABC}) = \Delta_{DEF}$ gibt. Nach Lemma 14.7 existiert eine Möbiustransformation Φ mit $\Phi(g(B, A)) = g(D, E)$. Nach Anwendung einer Streckung und eventuell einer Inversion können wir zudem annehmen, dass $\Phi(B) = E$ und $\Phi(A) = iy'$ für ein $y' > 1$. Aus $d_{\mathbb{H}}(A, B) = d_{\mathbb{H}}(D, E)$ folgt dann $y = y'$, d.h. $\Phi(B) = E$. Der Halbstrahl $\Phi(\overrightarrow{BC})$ schlägt den Winkel $\angle ABC = \angle DEF$ mit dem Halbstrahl $\overrightarrow{BA} = \overrightarrow{ED}$ ein. Aus der Tatsache, dass beide Dreiecke positiv orientiert sind folgt nun, dass $\Phi(\overrightarrow{BC}) = \overrightarrow{EF}$. Nachdem $d_{\mathbb{H}}(C, B) = d_{\mathbb{H}}(F, E)$ folgt dann $\Phi(C) = F$. Zusammengefasst gilt also $\Phi(\Delta_{ABC}) = \Delta_{DEF}$.

Die Eindeutigkeit von Φ wird in Übungsblatt 10 bewiesen. $\square$

15.3. Krümmung auf Riemannschen Untermannigfaltigkeiten. Es gibt eine sehr gut entwickelte Theorie von “Krümmung” auf Riemannschen Untermannigfaltigkeiten. Wir haben leider keine Zeit, diese in dieser Vorlesung zu entwickeln. Wir begnügen uns deshalb mit einer sehr anschaulichen Definition auf 2-dimensionalen Riemannschen Untermannigfaltigkeiten.

Definition. Es sei (M, g) eine 2-dimensionale Riemannsche Untermannigfaltigkeit und es sei $P \in M$. Wir sagen die *Krümmung in P ist positiv*, wenn es eine Umgebung U von P gibt, so dass für jedes Dreieck in U gilt, dass

$$\text{Summe der Innenwinkel} > \pi.$$

Ganz analog definieren wir auch, dass die Krümmung in P negativ ist.

Bemerkung. Die obige Definition ist weit entfernt von der Standarddefinition der Krümmung einer 2-dimensionalen Riemannschen Untermannigfaltigkeit. Mit mehr Aufwand kann man jedem Punkt P einer 2-dimensionalen Riemannschen Untermannigfaltigkeit eine Krümmung $k(P) \in \mathbb{R}$ zuordnen. Das Vorzeichen von $k(P)$ entspricht dann der obigen Konvention.

Beispiele.

- (1) Es folgt aus Satz 15.3, dass die Krümmung in jedem Punkt auf $\mathbb{S}^2$ positiv ist.
- (2) Es folgt aus Satz 15.4, dass die Krümmung in jedem Punkt auf $\mathbb{H}$ negativ ist.
- (3) In den Abbildungen 95 und 96 sind mehrere Punkte von positiven und negativer Krümmung skizziert. Für einen Punkt P auf einer Fläche M in $\mathbb{R}^3$ gilt anschaulich gesprochen folgende Regel:
 - (a) die Krümmung ist positiv, wenn die Fläche in der Nähe von P wie eine Ausbuchtung aussieht,
 - (b) die Krümmung ist negativ, wenn die Fläche in der Nähe von P wie ein Sattel aussieht.

¹⁹⁷Warum?

In Riemannschen Untermannigfaltigkeiten hängt im Allgemeinen die Krümmung vom Punkt ab.

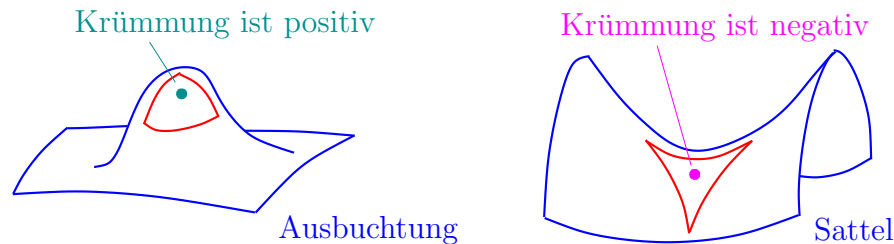


ABBILDUNG 95.

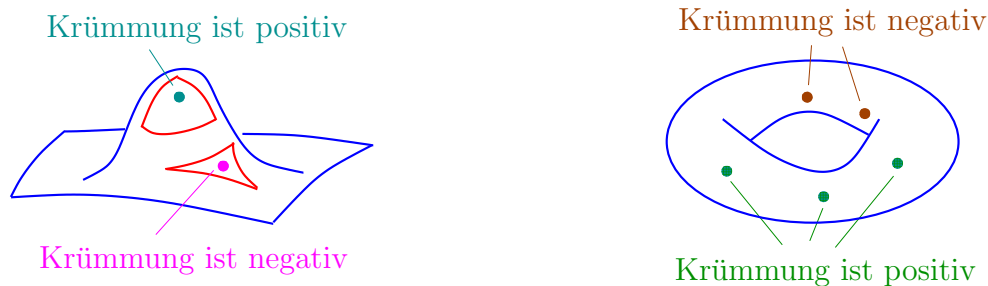


ABBILDUNG 96.

Definition. Es sei (M, g) eine 2-dimensionale Riemannsche Untermannigfaltigkeit.

- (1) Wir sagen (M, g) ist *lokal isometrisch* zu einer 2-dimensionalen Riemannschen Untermannigfaltigkeit (N, h) , wenn es für jedes $P \in M$ eine offene Umgebung U und eine offene Teilmenge $V \subset N$ sowie eine Isometrie $\varphi: (U, g) \rightarrow (V, h)$ gibt.
- (2) Wir sagen (M, g) besitzt *konstante Krümmung 0*, wenn (M, g) lokal isometrisch zu $\mathbb{E}^2$ ist.
- (3) Wir sagen (M, g) besitzt *konstante Krümmung +1*, wenn (M, g) lokal isometrisch zu $\mathbb{S}^2$ ist.
- (4) Wir sagen (M, g) besitzt *konstante Krümmung -1*, wenn (M, g) lokal isometrisch zu $\mathbb{H}$ ist.

Bemerkung.

- (1) Wenn eine Riemannsche Untermannigfaltigkeit (M, g) konstante Krümmung +1 besitzt, dann folgt leicht aus Satz 15.3, dass die Krümmung in jedem Punkt auf M positiv ist. Die analoge Aussage gilt nach Satz 15.4 auch, wenn wir “+1” und “positiv” durch “-1” und “negativ” ersetzen.

- (2) Es sei $\chi \in \{-1, 0, +1\}$ und es sei (M, g) eine Riemannsche Untermannigfaltigkeit, welche konstante Krümmung χ besitzt. Dann folgt aus Satz 15.3, Satz 15.4 und aus klassischer Schulmathematik, dass für jedes genügend kleine Dreieck in (M, g) die Gleichung

$$\chi \cdot \text{Flächeninhalt des Dreiecks} = \text{Summe der Innenwinkel} - \pi$$

gilt.

- (3) Der Zylinder

$$Z := \{(x, y, z) \in \mathbb{R}^3 \mid x^2 + y^2 = 1\}$$

besitzt konstante Krümmung 0. In der Tat zeigt die Rechnung von Seite 164, dass die Abbildung

$$\begin{aligned} \mathbb{R}^2 &\rightarrow Z \\ (s, t) &\mapsto \begin{pmatrix} \cos s \\ \sin s \\ t \end{pmatrix}. \end{aligned}$$

eine lokale Isometrie ist. Diese Abbildung ist anschaulich dadurch gegeben, dass wir ein Blatt Papier zu einem Zylinder “Zusammenrollen”. Auf dem Blatt Papier ist die Winkelsumme immer π , nachdem das Zusammenrollen eine lokale Isometrie ist, gilt dies auch für Dreiecke auf dem Zylinder.

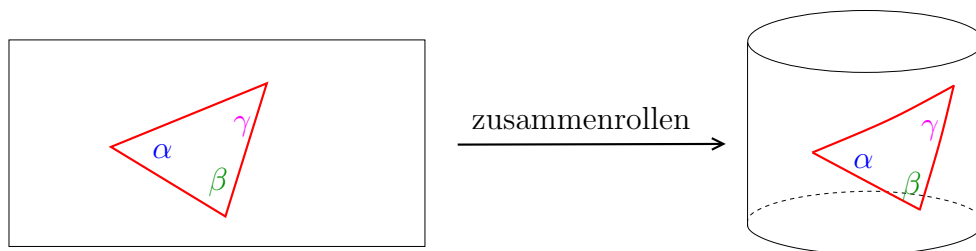


ABBILDUNG 97.

Wir haben gerade gesehen, dass der Zylinder konstante Krümmung 0 besitzt. Andererseits hatten wir in Abbildung 95 anschaulich gesehen, dass der Standardtorus in $\mathbb{R}^3$ keine konstante Krümmung besitzt. Es stellt sich jetzt folgende Frage:

Frage. *Gibt es geschlossene 2-dimensionale Riemannsche Untermannigfaltigkeiten, außer S^2 , welche konstante Krümmung besitzen?*

16. TOPOLOGISCHE RÄUME II

Wir haben uns jetzt also ausführlich mit Untermannigfaltigkeiten beschäftigt. Die Theorie der Untermannigfaltigkeiten ist aus mehreren Gründen etwas unbefriedigend:

- (1) Unser Universum schaut zwar in unserer Umgebung aus wie eine Teilmenge von $\mathbb{R}^3$, und vermutlich ist dies an den meisten Orten des Universums der Fall, aber es ist nicht klar, dass unser Universum als ganzes als Teilmenge des $\mathbb{R}^3$ aufgefasst werden kann.¹⁹⁸ Andererseits ist unser Universum erst Recht nicht eine 3-dimensionale Untermannigfaltigkeit eines größeren $\mathbb{R}^n$'s. Wir würden deshalb gerne eine Theorie von "3-dimensionalen Objekten" aufbauen, welche nicht Teilmenge von einem $\mathbb{R}^n$ sind.
- (2) Jede Fläche von Geschlecht g , wie in Abbildung 98, sollte eigentlich eine 2-dimensionale Untermannigfaltigkeit sein. Aber mit der Sprache der Untermannigfaltigkeiten ist es sehr schwierig eine mathematisch saubere Definition und Beschreibung von Flächen von Geschlecht g zu geben.

Wir werden im Folgenden den Begriff einer "Mannigfaltigkeit" einführen, welcher den Begriff von "Untermannigfaltigkeit" erweitert. Um diesen Begriff einführen zu können, müssen wir die topologischen Räume noch einmal genauer studieren. Insbesondere werden wir in diesem Kapitel verschiedene Konstruktionen von topologischen Räumen kennenlernen.

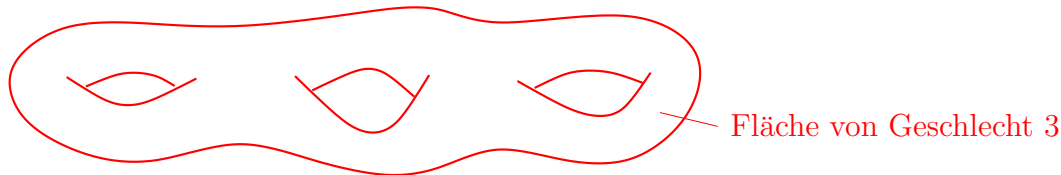


ABBILDUNG 98.

16.1. Definition eines topologischen Raums und Beispiele. Wir erinnern noch einmal an die Definition von einem topologischen Raum.

Definition. Ein *topologischer Raum* ist ein Paar $(X, \mathcal{T})$, wobei X eine Menge ist und $\mathcal{T}$ eine *Topologie auf X* , d.h. $\mathcal{T}$ ist eine Menge von Teilmengen von X mit folgenden Eigenschaften:

- (T1) die leere Menge und die ganze Menge X sind in $\mathcal{T}$ enthalten,
- (T2) der Durchschnitt von *endlich vielen* Mengen in $\mathcal{T}$ ist wiederum eine Menge in $\mathcal{T}$,
- (T3) die Vereinigung von *beliebig vielen* Mengen in $\mathcal{T}$ ist wiederum eine Menge in $\mathcal{T}$.

Die Mengen in $\mathcal{T}$ werden als *offen bezüglich $\mathcal{T}$* , oder, wenn keine Verwechslungsgefahr besteht, einfach nur als *offen* bezeichnet.

¹⁹⁸Wenn unser Universum eine Teilmenge von $\mathbb{R}^3$ wäre und es anscheinend wohl endlich ist, was soll dann das Komplement vom Universum sein? Hmm.

16.2. Basis einer Topologie. Es sei X eine Menge und es sei $\mathcal{B} = \{B_i\}_{i \in I}$ eine Familie von Teilmengen von X . Wir sagen $\mathcal{B}$ besitzt die *Basiseigenschaft*, wenn gilt:

- (B1) Zu jedem $x \in X$ existiert ein $i \in I$, so dass $x \in B_i$.¹⁹⁹
- (B2) Es seien $i, j \in I$ und es sei $x \in B_i \cap B_j$. Dann existiert ein $k \in I$, so dass $x \in B_k$ und $B_k \subset B_i \cap B_j$.

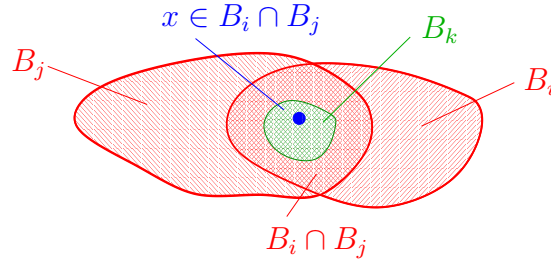


ABBILDUNG 99. Schematisches Bild für die Basiseigenschaft (B2).

Lemma 16.1. Es sei X eine Menge und es sei $\mathcal{B} = \{B_i\}_{i \in I}$ eine Familie von Teilmengen von X , welche die Basiseigenschaft besitzt. Wir definieren

$$\mathcal{T} := \{V \subset X \mid \text{zu jedem } x \in V \text{ gibt es ein } i \in I, \text{ so dass } x \in B_i \subset V\}.$$

Dann ist $\mathcal{T}$ eine Topologie auf X .

Wir nennen $\mathcal{T}$ die *von $\mathcal{B}$ erzeugte Topologie auf X* . Umgekehrt, sagen wir, dass $\mathcal{B}$ eine *Basis der Topologie $\mathcal{T}$* ist.

Beispiel. Es sei X ein metrischer Raum und es sei $\mathcal{B}$ die Menge aller offenen Kugeln in X , d.h.

$$\mathcal{B} := \{B_\epsilon(x) \mid x \in X \text{ und } \epsilon > 0\},$$

dann besitzt $\mathcal{B}$ die Basiseigenschaft.²⁰⁰ Die von $\mathcal{B}$ erzeugte Topologie ist dann die übliche Topologie auf einem metrischen Raum.

Beweis von Lemma 16.1. Wir zeigen nun, dass die Axiome (T1), (T2) und (T3) erfüllt sind.

- (T1) Nach (B1) gibt es zu jedem $x \in X$ ein $i \in I$ mit $x \in B_i$. Es folgt also per Definition, dass $X \in \mathcal{T}$. Zudem ist $\emptyset \in \mathcal{T}$, nachdem die Vereinigung von null Mengen, die Nullmenge ist.
- (T2) Wir müssen zeigen, dass der Durchschnitt von endlich vielen Mengen in $\mathcal{T}$ wieder in $\mathcal{T}$ liegt.

Wir betrachten zuerst den Fall von zwei Mengen in $\mathcal{T}$. Es seien also $U, V \in \mathcal{T}$. Wir wollen zeigen, dass $U \cap V \in \mathcal{T}$ ist. Es sei also $x \in U \cap V$. Nachdem U und

¹⁹⁹Mit anderen Worten, es ist $X = \bigcup_{i \in I} B_i$.

²⁰⁰Warum besitzt $\mathcal{B}$ die Basiseigenschaft (B2)?

V in $\mathcal{T}$ liegen, gibt es also $i, j \in I$, so dass $x \in B_i \subset U$ und $x \in B_j \subset V$. Nach Eigenschaft (B2) von $\mathcal{B}$ existiert ein $k \in I$, so dass $x \in B_k$ und $B_k \subset B_i \cap B_j$. Insbesondere gilt dann auch, dass $x \in B_k \subset U \cap V$. Also ist $U \cap V \in \mathcal{T}$.

Es seien nun $U_1 \cap \dots \cap U_k$ offene Mengen. Es folgt aus der gerade bewiesenen Aussage und einem einfachen Induktionsargument, dass $U_1 \cap \dots \cap U_k$ ebenfalls in $\mathcal{T}$ liegt.

- (T3) Es sei also $U_j, j \in J$ eine Familie von Mengen in $\mathcal{T}$. Wir müssen zeigen, dass die Vereinigungsmenge $\bigcup_{j \in J} U_j$ in $\mathcal{T}$ liegt. Es sei also $x \in \bigcup_{j \in J} U_j$. Dann gibt es insbesondere ein $j \in J$, so dass $x \in U_j$. Nachdem $U_j \in \mathcal{T}$ gibt es ein $i \in I$, so dass $x \in B_i \subset U_j$. Dann gilt aber auch, dass

$$x \in B_i \subset U_k \subset \bigcup_{j \in J} U_j. \quad \square$$

Lemma 16.2. *Es sei X eine Menge und es sei $\mathcal{B} = \{B_i\}_{i \in I}$ eine Familie von Teilmengen von X , welche die Basiseigenschaft besitzt. Es sei $\mathcal{T}$ die von $\mathcal{B}$ erzeugte Topologie auf X . Dann gilt*

$$V \subset X \text{ ist offen bezüglich } \mathcal{T} \iff V \text{ ist die Vereinigung von Mengen in } \mathcal{B}.$$

Beweis. Wir nehmen zuerst an, dass $V = \bigcup_{j \in J} B_j$ die Vereinigung von Mengen in $\mathcal{B}$ ist. Nachdem die Mengen in $\mathcal{B}$ offen bezüglich $\mathcal{T}$ sind, und nachdem $\mathcal{T}$ eine Topologie ist, folgt auch, dass $V \in \mathcal{T}$.

Nehmen wir nun an, dass $V \subset X$ offen ist bezüglich $\mathcal{T}$. Für jedes $x \in V$ existiert dann, per Definition von $\mathcal{T}$, ein $i_x \in I$, so dass $x \in B_{i_x} \subset V$. Dann folgt aber, dass

$$V = \bigcup_{x \in V} \{x\} \subset \bigcup_{x \in V} B_{i_x} \subset V, \quad \text{es folgt also, dass } V = \bigcup_{x \in V} B_{i_x}. \quad \square$$

Basen von Topologien können erstaunlich praktisch sein, beispielsweise besagt folgendes Lemma, dass es genügt die Stetigkeitseigenschaft einer Abbildung zwischen topologischen Räumen für offene Mengen in einer Basis zu überprüfen:

Lemma 16.3. *Es sei $f: X \rightarrow Y$ eine Abbildung zwischen topologischen Räumen. Es sei $\mathcal{B} = \{B_i\}_{i \in I}$ eine Basis für die Topologie von Y . Dann gilt*

$$f \text{ ist stetig} \iff \text{für jedes } i \in I \text{ ist } f^{-1}(B_i) \text{ offen in } X.$$

Beweis. Die Richtung " $\Rightarrow$ " ist offensichtlich. Nehmen wir nun an, dass für jedes $i \in I$ das Urbild $f^{-1}(B_i)$ offen in X ist. Wir wollen zeigen, dass f stetig ist. Es sei also $U \subset Y$ eine

offene Teilmenge. Nach Lemma 16.2 gibt es eine Teilmenge J von I , so dass

$$U = \bigcup_{j \in J} B_j.$$

Dann ist

$$f^{-1}(U) = f^{-1}\left(\bigcup_{j \in J} B_j\right) = \bigcup_{j \in J} \underbrace{f^{-1}(B_j)}_{\substack{\text{Vereinigung von offenen} \\ \text{Mengen, also offen in } X}}.$$

□

Lemma 16.4. *Es sei $(X, \mathcal{T})$ ein topologischer Raum und es sei $\mathcal{C}$ eine Familie von offenen Mengen in $(X, \mathcal{T})$ mit folgender Eigenschaft:*

Zu jeder offenen Menge U und zu jedem $x \in U$ gibt es ein $C \in \mathcal{C}$, so dass $x \in C \subset U$.

Dann ist $\mathcal{C}$ eine Basis für die Topologie $\mathcal{T}$.

Beispiel. Es sei $X = \mathbb{R}^2$ und es sei

$$\mathcal{C} := \{(a, b) \times (c, d) \mid a, b, c, d \in \mathbb{R}\}$$

die Menge aller offenen Rechtecke in $\mathbb{R}^2$. Dann kann man sich leicht davon überzeugen, siehe Abbildung 100, dass $\mathcal{C}$ die gewünschte Eigenschaft in Lemma 16.4 besitzt, und daher eine Basis für die Standardtopologie auf $\mathbb{R}^2$ bildet.

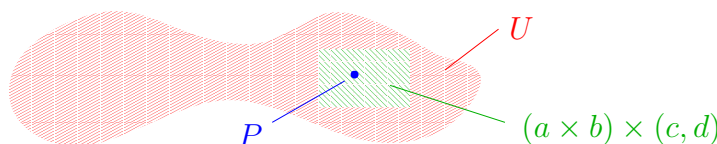


ABBILDUNG 100.

Beweis. Wir zeigen zuerst, dass $\mathcal{C}$ die Basiseigenschaft erfüllt:

- (B1) Wenden wir die Eigenschaft von $\mathcal{C}$ auf $U = X$ an, dann erhalten wir sofort, dass es zu jedem $x \in X$ ein $C \in \mathcal{C}$ gibt, so dass $x \in C$.
- (B2) Es sei nun $x \in C_1 \cap C_2$, wobei $C_1, C_2 \in \mathcal{C}$. Nachdem C_1 und C_2 offen sind, ist auch $U = C_1 \cap C_2$ offen. Wir wenden die definierende Eigenschaft von $\mathcal{C}$ auf $U = C_1 \cap C_2$ an, und erhalten ein $C \in \mathcal{C}$, so dass $x \in C \subset C_1 \cap C_2$.

Wir bezeichnen nun mit SS die Topologie auf X , welche von $\mathcal{C}$ erzeugt wird. Wir müssen zeigen, dass $SS = \mathcal{T}$.

Wir zeigen zuerst die Inklusion $SS \subset \mathcal{T}$. Nachdem alle Mengen in $\mathcal{C}$ offen bezüglich $\mathcal{T}$ sind, folgt per Definition und aus Lemma 16.2, dass in der Tat $SS \subset \mathcal{T}$.

Wir zeigen nun die umgekehrte Inklusion $\mathcal{T} \subset SS$. Es sei also $U \in \mathcal{T}$. Wir wollen zeigen, dass $U \in SS$. Sei also $x \in U$. Dann existiert nach Voraussetzung ein $C \in \mathcal{C}$, so

dass $x \in C \subset U$. Also folgt per Definition, dass $U \in SS$. Wir haben also gezeigt, dass $\mathcal{T} \subset SS$. $\square$

Wir erinnern nun an den Begriff der Teilraumtopologie, welchen wir schon auf Seite 10 eingeführt hatten. Es sei also $(X, \mathcal{T})$ ein topologischer Raum und es sei $A \subset X$ eine Teilmenge. Die Teilraumtopologie auf A ist dann gegeben durch die Topologie

$$SS := \{A \cap U \mid U \in \mathcal{T}\}.$$

Die Definition besagt also, dass eine Teilmenge U offen in A ist, genau dann, wenn es eine offene Menge $V \subset X$ gibt, so dass $U = A \cap V$.

Lemma 16.5. *Es sei X ein topologischer Raum und $A \subset X$ eine Teilmenge. Es sei $\mathcal{B}$ eine Basis für die Topologie von X . Dann ist*

$$\{B \cap A \mid B \in \mathcal{B}\}$$

eine Basis für die Topologie von A .

Beweis. Wir beweisen das Lemma mithilfe von Lemma 16.4. Es sei also $U \subset A$ eine offene Menge in A und $x \in U$. Nach Voraussetzung existiert eine offene Menge $V \subset X$ in X mit $U = V \cap A$. Nachdem $\mathcal{B}$ eine Basis für die Topologie von X ist, existiert ein $B \in \mathcal{B}$, so dass $x \in B \subset V$. Dann folgt aber auch, dass $x \in B \cap A \subset V \cap A = U$. Es folgt aus Lemma 16.4, dass

$$\{B \cap A \mid B \in \mathcal{B}\}$$

eine Basis für die Topologie von A ist. $\square$

Beispiele.

(1) Wir betrachten

$$S^1 = \{(x, y) \in \mathbb{R}^2 \mid x^2 + y^2 = 1\} \subset \mathbb{R}^2.$$

Eine Basis für die Topologie auf $\mathbb{R}^2$ ist gegeben durch

$$\mathcal{B} = \{B_r(x, y) \mid (x, y) \in \mathbb{R}^2 \text{ und } r > 0\}.$$

Die Schnittmenge von einem offenen Ball in $\mathbb{R}^2$ mit S^1 ist ein “offenes Intervall auf S^1 ”, d.h. eine Teilmenge der Form $\{e^{i\varphi} \mid \varphi \in (a, b)\}$. Nach Lemma 16.5 bildet also die Menge $\mathcal{C}$ der offenen Intervalle auf S^1 eine Basis der Topologie auf S^1 .

(2) Wir betrachten jetzt wieder die “Gerade mit einem Punkt im Unendlichen”, d.h. $X = \mathbb{R} \cup \{\infty\}$ mit der Topologie, welche wir auf Seite 8 eingeführt hatten. Man kann sich leicht davon überzeugen, dass eine Basis der Topologie von X gegeben ist durch

$$\mathcal{D} := \{(a, b) \mid a < b \in \mathbb{R}\} \cup \{(-\infty, C) \cup \{\infty\} \cup (D, \infty) \mid C, D \in \mathbb{R}\}.$$

Wir betrachten nun die Abbildung

$$\begin{aligned} f: X &\rightarrow S^1 \\ x &\mapsto \begin{cases} e^{i2 \arctan(x)}, & \text{wenn } x \in \mathbb{R} \\ (-1, 0), & \text{wenn } x = \infty. \end{cases} \end{aligned}$$

Diese Abbildung ist offensichtlich eine Bijektion. Man sieht auch leicht, dass

$$U \in \mathcal{D} \iff f(U) \in \mathcal{C}.$$

Es folgt nun aus Lemma 16.3, dass die Abbildung f ein Homöomorphismus ist. Insbesondere ist X homöomorph zu S^1 .

16.3. Das Produkt von topologischen Räumen.

Lemma 16.6. *Es seien X und Y topologische Räume. Dann besitzt*

$$\mathcal{B} := \{U \times V \mid U \text{ offen in } X \text{ und } V \text{ offen in } Y\}$$

die Basiseigenschaft.

Beweis. Wir müssen nun also nachweisen, dass $\mathcal{B}$ in der Tat die beiden Basiseigenschaften besitzt:

- (B1) Es ist $X \times Y \in \mathcal{B}$, also gibt es zu jedem $(x, y) \in X \times Y$ eine Teilmenge in $\mathcal{B}$, nämlich $X \times Y$, welche (x, y) enthält. $\mathcal{B}$ erfüllt also die Basiseigenschaft (B1).
- (B2) Es seien nun $U_1 \times V_1$ und $U_2 \times V_2$ aus $\mathcal{B}$, dann gilt

$$(U_1 \times V_1) \cap (U_2 \times V_2) = (U_1 \cap U_2) \times (V_1 \cap V_2) \in \mathcal{B}.$$

Insbesondere erfüllt also $\mathcal{B}$ auch die Basiseigenschaft (B2). □

Definition. Es seien X und Y topologische Räume. Die *Produkttopologie auf $X \times Y$* ist die von

$$\mathcal{B} := \{U \times V \mid U \text{ offen in } X \text{ und } V \text{ offen in } Y\}$$

erzeugte Topologie auf $X \times Y$.

Konvention. Im Folgenden, wenn X und Y topologische Räume sind, dann verstehen wir $X \times Y$ immer mit der Produkttopologie.

Lemma 16.7. *Es sei $\mathcal{B}$ eine Basis für den topologischen Raum X und es sei $\mathcal{C}$ eine Basis für den topologischen Raum Y . Dann ist*

$$\mathcal{D} = \{B \times C \mid B \in \mathcal{B} \text{ und } C \in \mathcal{C}\}$$

eine Basis für $X \times Y$.

Beispiel. Es sei $X = Y = \mathbb{R}$ mit der Standardtopologie. Die offenen endlichen Intervalle bilden eine Basis für den topologischen Raum $X = Y = \mathbb{R}$. Es folgt also aus Lemma 16.7, dass die Produkttopologie auf $\mathbb{R} \times \mathbb{R}$ von “offenen Quadern” der Form $(a, b) \times (c, d)$ erzeugt wird. Wie wir auf Seite 200 gesehen hatten, erzeugen die offenen Quader gerade die Standardtopologie auf $\mathbb{R}^2 = \mathbb{R} \times \mathbb{R}$.

Beweis. Wir beweisen das Lemma mithilfe von Lemma 16.4. Es sei also $W \subset X \times Y$ eine offene Menge und $(x, y) \in W$. Nach Definition der Produkttopologie existieren offene Mengen $U \subset X$ und $V \subset Y$, so dass $(x, y) \in U \times V$ und $U \times V \subset W$.

Nachdem $\mathcal{B}$ eine Basis für X ist, folgt, dass es $B \in \mathcal{B}$ mit $x \in B \subset U$ gibt. Ganz analog zeigt man, dass es $C \in \mathcal{C}$ mit $y \in C \subset V$ gibt. Es folgt, dass

$$(x, y) \in B \times C \subset U \times V \subset W.$$

Es folgt also aus Lemma 16.4, dass

$$\mathcal{D} = \{B \times C \mid B \in \mathcal{B} \text{ und } C \in \mathcal{C}\}$$

eine Basis für den topologischen Raum $X \times Y$ ist. $\square$

Lemma 16.8. *Es seien X und Y zwei nichtleere topologische Räume. Dann gilt*

$$(a) \quad X \text{ und } Y \text{ sind Hausdorff} \iff X \times Y \text{ ist Hausdorff}$$

und

$$(b) \quad X \text{ und } Y \text{ sind kompakt} \iff X \times Y \text{ ist kompakt.}$$

Beweis. Wir beweisen Teil (a) des Lemmas, Teil (b) ist eine Übungsaufgabe in Übungsblatt 11.

Es seien also X und Y zwei topologische Räume. Wir nehmen zuerst an, dass X und Y Hausdorff sind. Es seien nun (x, y) und (x', y') zwei Punkte in $X \times Y$. Nach Voraussetzung existieren offene Umgebungen U von x und U' von x' sowie V von y und V' von y' , so dass $U \cap U' = \emptyset$ und $V \cap V' = \emptyset$. Also ist $U \times V \cap U' \times V' = \emptyset$. Per Definition der Topologie auf $X \times Y$ sind $U \times V$ und $U' \times V'$ offen in $X \times Y$. Wir haben also disjunkte offene Umgebungen um (x, y) und (x', y') gefunden. Also ist $X \times Y$ Hausdorff. (Der Beweis ist in Abbildung 101 links skizziert.)

Wir nehmen nun an, dass $X \times Y$ Hausdorff ist. Wir wollen zeigen, dass X Hausdorff ist. Es seien also x und x' zwei verschiedene Punkte in X . Nachdem Y nichtleer ist existiert ein $y \in Y$. Nachdem $X \times Y$ Hausdorff ist existieren disjunkte offene Umgebungen W und W' von (x, y) und (x', y) . Per Definition der Produkttopologie gibt es offene Umgebungen U von x und U' von x' sowie V von y und V' von y , so dass $(U \times V) \cap (U' \times V') = \emptyset$. Insbesondere ist dann auch $(U \times \{y\}) \cap (U' \times \{y\}) = \emptyset$. Dann gilt aber auch $U \cap U' = \emptyset$. Insbesondere sind U und U' disjunkte offene Umgebungen von x und x' . Also ist X Hausdorff. (Der Beweis ist in Abbildung 101 rechts skizziert.) Der Beweis, dass auch Y Hausdorff, ist natürlich genau der gleiche. $\square$

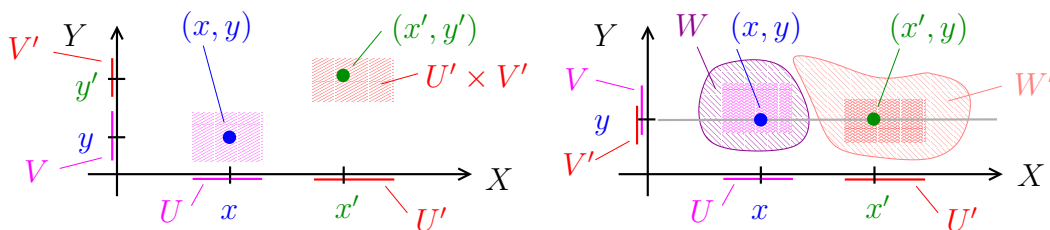


ABBILDUNG 101. Skizzen zum Beweis von Lemma 16.8.

Wir beschließen die Diskussion von Produkträumen mit dem Beweis, dass $S^1 \times S^1$ homöomorph zum Torus

$$T := \{((2 + \sin \theta) \cos \varphi, (2 + \sin \theta) \sin \varphi, \cos \theta) \mid \theta, \varphi \in \mathbb{R}\}$$

ist. Wir betrachten hierzu die Abbildung

$$\begin{aligned} f: T &\rightarrow S^1 \times S^1 \\ ((2 + \sin \theta) \cos \varphi, (2 + \sin \theta) \sin \varphi, \cos \theta) &\mapsto (e^{i\theta}, e^{i\varphi}). \end{aligned}$$

Man kann leicht zeigen, dass die Abbildung f bijektiv. Mithilfe von Lemma 16.7 und der Basis $\mathcal{C}$ der Topologie von S^1 , welche wir auf Seite 201 eingeführt hatten, kann man zudem ohne größere Probleme zeigen, dass die Abbildung f stetig ist. Nachdem T kompakt ist und nachdem $S^1 \times S^1$ nach Lemma 16.8 Hausdorff ist, folgt nun aus Satz 1.13, dass f ein Homöomorphismus ist.

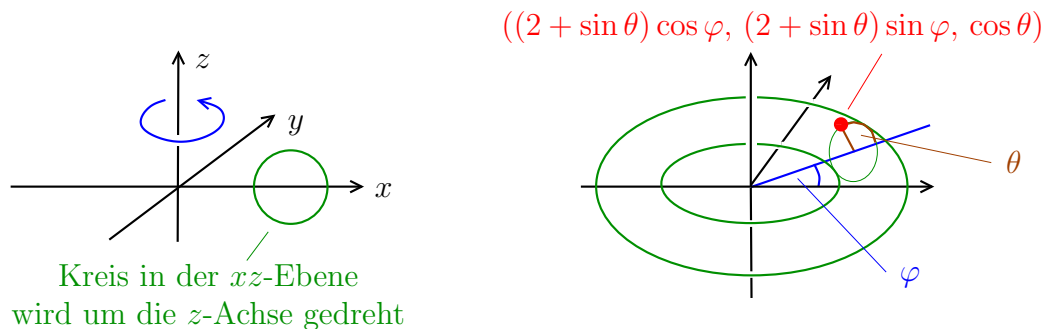


ABBILDUNG 102. Der Torus T in $\mathbb{R}^3$.

17. QUOTIENTENRÄUME

17.1. Äquivalenzrelationen. Es sei X eine Menge. Wir werden in diesem Kapitel ausführlich mit Äquivalenzrelationen und Äquivalenzklassen arbeiten. Wir erinnern deshalb zuerst noch einmal an die wichtigsten Definitionen. Eine *Äquivalenzrelation* auf X ist eine Teilmenge $M \subset X \times X$, so dass für alle $x, y, z \in X$ gilt

$$\begin{aligned} (x, x) &\in M \\ (x, y) \in M &\implies (y, x) \in M \quad (\text{Symmetrie}) \\ (x, y) \in M \text{ und } (y, z) \in M &\implies (x, z) \in M \quad (\text{Transitivität}). \end{aligned}$$

Wir schreiben im Folgenden

$$x \sim y \quad :\Longleftrightarrow \quad (x, y) \in M,$$

und wir sagen x und y sind *äquivalent*. Mithilfe dieser Notation können wir die obigen definierenden Eigenschaften wie folgt umformulieren: $\sim$ ist eine Äquivalenzrelation auf einer Menge A , wenn für alle $x, y, z \in A$ gilt

$$\begin{aligned} x &\sim x \\ x \sim y &\implies y \sim x \\ x \sim y \text{ und } y \sim z &\implies x \sim z. \end{aligned}$$

Beispiele.

(a) Für $x, y \in \mathbb{R}$ definieren wir

$$x \sim y \quad :\Longleftrightarrow \quad (x - y) \in \mathbb{Z}.$$

Dies ist eine Äquivalenzrelation auf $\mathbb{R}$.

(b) Für $P, Q \in S^2 := \{(x, y, z) \in \mathbb{R}^3 \mid x^2 + y^2 + z^2 = 1\}$ schreiben wir

$$P \sim Q \quad :\Longleftrightarrow \quad P = Q \text{ oder } P = -Q$$

Dies ist eine Äquivalenzrelation auf S^2 .

(c) Für $P, Q \in D^2 := \{(x, y) \in \mathbb{R}^2 \mid x^2 + y^2 \leq 1\}$ definieren wir

$$P \sim Q \quad :\Longleftrightarrow \quad \begin{aligned} &P = Q \text{ oder} \\ &P \text{ und } Q \text{ liegen auf } S^1 := \{(x, y) \mid x^2 + y^2 = 1\}. \end{aligned}$$

Dies ist eine Äquivalenzrelation auf D^2 .

(d) Wir verallgemeinern jetzt Beispiel (c). Genauer gesagt, es sei X eine Menge und $A \subset X$ eine Teilmenge. Wir definieren dann

$$P \sim Q \quad :\Longleftrightarrow \quad P = Q \text{ oder } P \text{ und } Q \text{ liegen in } A.$$

Dies ist eine Äquivalenzrelation auf X .

Es sei $\sim$ eine Äquivalenzrelation auf einer Menge X . Eine *Äquivalenzklasse* ist eine Teilmenge $Y \subset X$, so dass gilt

- (1) für alle $y, y' \in Y$ gilt: $y \sim y'$,
- (2) wenn $z \sim y$ für ein $y \in Y$, dann gilt auch $z \in Y$,

Anders ausgedrückt, eine Äquivalenzklasse von X ist eine Teilmenge Y , so dass je zwei Elemente in Y äquivalent sind, und welches jedes Element von X enthält, welches zu einem Element in Y äquivalent ist.

In Analysis III hatten wir folgendes (elementare) Lemma bewiesen:

Lemma 17.1. *Es sei $\sim$ eine Äquivalenzrelation auf einer Menge X . Dann ist X die disjunkte Vereinigung aller Äquivalenzklassen ist, d.h.*

- (1) *jedes $x \in X$ liegt in einer Äquivalenzklasse,*
- (2) *wenn A und B Äquivalenzklassen sind, dann gilt entweder $A = B$ oder $A \cap B = \emptyset$.*

Es sei $\sim$ eine Äquivalenzrelation auf einer Menge X . Wir bezeichnen dann mit $X/\sim$ die Menge der Äquivalenzklassen. Wenn $x \in X$, dann bezeichnen wir die Äquivalenzklasse von x ²⁰¹ normalerweise mit $[x] \in X/\sim$ oder $\bar{x} \in X/\sim$.

Beispiele.

- (a) In Beispiel (a) ist jede Äquivalenzklasse von der Form

$$t + \mathbb{Z} := \{t + n \mid n \in \mathbb{Z}\}$$

wobei $t \in \mathbb{R}$. Zudem gilt $s + \mathbb{Z} = t + \mathbb{Z}$ genau dann, wenn $s - t \in \mathbb{Z}$. Die Menge der Äquivalenzklassen wird mit $\mathbb{R}/\mathbb{Z}$ bezeichnet. Dies ist gerade die Quotientengruppe der abelschen Gruppe $\mathbb{R}$ bezüglich der Untergruppe $\mathbb{Z}$.

- (b) In Beispiel (b) ist jede Äquivalenzklasse von der Form $\{P, -P\}$, wobei $P \in S^2$.
- (c) In Beispiel (c) ist jeder Punkt $(x, y) \in \mathbb{R}^2$ mit $x^2 + y^2 < 1$ eine Äquivalenzklasse, und alle Punkte in $S^1 = \{(x, y) \mid x^2 + y^2 = 1\}$ bilden eine Äquivalenzklasse.
- (d) In Beispiel (d) ist jeder Punkt $x \in X \setminus A$ eine Äquivalenzklasse, und A bildet eine Äquivalenzklasse. Wir bezeichnen die Menge der Äquivalenzklassen oft auch mit X/A .²⁰²

17.2. Die Quotiententopologie. Es sei nun $(X, \mathcal{T})$ ein topologischer Raum und es sei $\sim$ eine Äquivalenzrelation auf X . Wir werden nun eine Topologie auf der Menge $X/\sim$ der Äquivalenzklassen einführen.

Wir bezeichnen im Folgenden

$$\begin{aligned} p: X &\rightarrow X/\sim \\ x &\mapsto [x] \end{aligned}$$

²⁰¹Die Äquivalenzklasse von x ist die nach Lemma 17.1 eindeutig bestimmte Äquivalenzklasse von X , welche x enthält. Man sieht leicht, dass

$$[x] = \{y \in X \mid y \sim x\}.$$

²⁰²Die Notation ist manchmal etwas gefährlich. Beispielsweise hat $\mathbb{R}/\mathbb{Z}$ nun zwei verschiedene Bedeutungen, je nachdem ob wir der Notation aus Beispiel (a) oder der Notation aus Beispiel (d) folgen. Vom Kontext her sollte es aber normalerweise klar sein, welche Konvention wir verwenden.

als die *Projektionsabbildung*. Wir betrachten

$$SS := \{U \subset X/\sim \mid p^{-1}(U) \text{ ist offen in } X\}.$$

Dann kann man leicht überprüfen, dass SS eine Topologie auf $X/\sim$ definiert. Diese Topologie wird die *Quotiententopologie* genannt. Wir betrachten im Folgenden $X/\sim$ durchgehend als topologischen Raum bezüglich der Quotiententopologie.

Aus der Definition der Quotiententopologie folgt sofort folgendes Lemma:

Lemma 17.2. *Es sei X ein topologischer Raum und es sei $\sim$ eine Äquivalenzrelation auf X . Dann ist die Projektionsabbildung $p: X \rightarrow X/\sim$ stetig.*

Die Tatsache, dass die Projektionsabbildung $X \rightarrow X/\sim$ stetig ist, ist natürlich nicht überraschend, wir haben die Topologie auf $X/\sim$ genauso gewählt, dass die Projektion $X \rightarrow X/\sim$ stetig ist.

Bemerkung. Wenn X ein kompakter topologischer Raum ist, dann folgt aus Satz 1.10 und Lemma 17.2, dass auch jeder Quotientenraum $X/\sim$ wiederum kompakt ist. In Übungsblatt 10 werden wir jedoch ein Beispiel von einem Hausdorff Raum X mit einer Äquivalenzrelation $\sim$ sehen, so dass der Quotientenraum $X/\sim$ nicht Hausdorff ist.

Definition. Eine Abbildung $f: X \rightarrow Y$ zwischen topologischen Räumen heißt *offen*, wenn das Bild jeder offenen Menge in X offen in Y ist.

Beispiele.

- (1) Die Projektionsabbildung $p: \mathbb{R}^2 \rightarrow \mathbb{R}$, $p(x, y) := x$ ist offen, während die Inklusionsabbildung $i: \mathbb{R} \rightarrow \mathbb{R}^2$, $i(x) = (x, 0)$ offensichtlich nicht offen ist.
- (2) Es sei $X = \mathbb{R}$ und es sei $\sim$ die Äquivalenzrelation, welche gegeben ist durch

$$P \sim Q \iff P = Q \text{ oder } P \text{ und } Q \text{ liegen in } [-2, 2].$$

Dann ist die Projektionsabbildung $p: X \rightarrow X/\sim$ *nicht offen*. In der Tat, wir betrachten $U = (-1, 1) \subset \mathbb{R}$. Dann ist U natürlich offen in $\mathbb{R}$ aber $p(U)$ ist nicht offen in $\mathbb{R}/\sim$, denn $p^{-1}(p(U)) = [-2, 2]$ ist nicht offen in $\mathbb{R}$.

Lemma 17.3. *Es sei X ein topologischer Raum, es sei $\sim$ eine Äquivalenzrelation auf X und es sei $\mathcal{B}$ eine Basis der Topologie auf X . Wenn die Projektionsabbildung $X \rightarrow X/\sim$ offen ist, dann ist²⁰³*

$$p(\mathcal{B}) = \{p(B) \mid B \in \mathcal{B}\}$$

eine Basis der Topologie auf $X/\sim$.

Beweis. Wir verwenden das Kriterium von Lemma 16.4 um zu zeigen, dass $p(\mathcal{B})$ eine Basis der Topologie auf $X/\sim$ ist. Es sei also $U \subset X/\sim$ eine offene Menge und es sei $y \in U$. Wir müssen zeigen, dass es ein $B \in \mathcal{B}$ mit $y \in p(B) \subset U$ gibt.

²⁰³Die Mengen $p(B)$ sind offen in $X/\sim$, weil nach Voraussetzung die Projektionsabbildung $X \rightarrow X/\sim$ offen ist.

Wir wählen ein $x \in X$ mit $p(x) = y$. Nachdem $\mathcal{B}$ eine Basis der Topologie auf X und nachdem $p^{-1}(U)$ offen in X ist, gibt es nach Lemma 16.4 ein $B \in \mathcal{B}$ mit $x \in B \subset p^{-1}(U)$. Dann gilt aber auch, dass $y = p(x) \in p(B) \subset p(p^{-1}(U)) = U$. $\square$

Lemma 17.4. *Es sei $\sim$ eine Äquivalenzrelation auf einer Menge X und es sei $f: X \rightarrow Y$ eine Abbildung mit der Eigenschaft, dass $f(x) = f(y)$, wenn $x \sim y$. Dann existiert genau eine Abbildung $g: X/\sim \rightarrow Y$, so dass $f = g \circ p$, d.h. so dass das folgende Diagramm von Abbildungen kommutiert:*

$$\begin{array}{ccc} X & \xrightarrow{p} & X/\sim \\ & \searrow f & \downarrow g \\ & & Y. \end{array}$$

Wenn f stetig ist, dann ist zudem auch g stetig.

Bemerkung.

- (1) Es sei $f: X \rightarrow Y$ eine Abbildung wie in Lemma 17.4. Wir nennen die Abbildung $g: X/\sim \rightarrow Y$ die *induzierte Abbildung*. Wir bezeichnen die induzierte Abbildung oft mit $\bar{f}$, oder auch, wenn keine Verwechslungsgefahr besteht, einfach nur mit f .
- (2) Die Aussage von Lemma 17.4 ist ganz ähnlich zur folgenden, hoffentlich aus der Algebra bekannten, Aussage: es sei $\varphi: G \rightarrow H$ ein Gruppenhomomorphismus und $K \subset G$ eine normale Untergruppe mit $K \subset \text{Kern}(\varphi)$. Dann induziert φ einen eindeutig bestimmten Homomorphismus $\psi: G/K \rightarrow H$, so dass $\varphi = \psi \circ \rho$, wobei $\rho: G \rightarrow G/K$ die Projektionsabbildung ist.

Beweis von Lemma 17.4. Es sei $\sim$ eine Äquivalenzrelation auf einer Menge X und es sei $f: X \rightarrow Y$ eine Abbildung mit der Eigenschaft, dass $f(x) = f(y)$, wenn $x \sim y$. Es sei nun a ein Element in $X/\sim$. Dann existiert ein $x \in X$ mit $a = [x]$. Wir setzen

$$g(a) := f(x).$$

Diese Definition hängt nicht von der Wahl von x ab, nachdem $f(x) = f(y)$, wenn immer $x \sim y$. Diese Abbildung $g: X/\sim \rightarrow Y$ hat dann offensichtlich die Eigenschaft, dass $f = g \circ p$. Nachdem die Projektion $X \rightarrow X/\sim$ surjektiv ist, ist die Abbildung g auch eindeutig bestimmt.

Nehmen wir nun an, dass f stetig ist. Wir wollen zeigen, dass g stetig ist. Es sei also $U \subset Y$ offen. Wir müssen zeigen, dass $g^{-1}(U) \subset X/\sim$ offen ist. Per Definition der Quotiententopologie müssen wir also zeigen, dass $p^{-1}(g^{-1}(U))$ offen in X ist. Es ist aber

$$p^{-1}(g^{-1}(U)) = (g \circ p)^{-1}(U) = f^{-1}(U)$$

offen, nachdem f stetig ist. $\square$

Beispiel. Wir betrachten noch einmal Beispiel (a). Zur Erinnerung, für $x, y \in \mathbb{R}$ gilt

$$x \sim y \quad :\Leftrightarrow \quad (x - y) \in \mathbb{Z}.$$

Auf Übungsblatt 10 werden wir sehen, dass $\mathbb{R}/\sim$ kompakt ist. Wir betrachten

$$\begin{aligned} \varphi: \mathbb{R} &\rightarrow S^1 \\ x &\mapsto (\cos(2\pi x), \sin(2\pi x)). \end{aligned}$$

Es gilt offensichtlich, dass $\varphi(x) = \varphi(y)$, wenn immer $x \sim y$. Nach Lemma 17.4 existiert daher genau eine Abbildung

$$\bar{\varphi}: \mathbb{R}/\mathbb{Z} := \mathbb{R}/\sim \rightarrow S^1,$$

mit der Eigenschaft, dass $\bar{\varphi}(\bar{x}) = \varphi(x)$ für alle $x \in \mathbb{R}$. Man kann sich leicht davon überzeugen, dass $\bar{\varphi}$ bijektiv ist.

Nachdem $\mathbb{R}/\sim$ kompakt ist, und nachdem S^1 Hausdorff ist, folgt nun aus Satz 1.13, dass $\bar{\varphi}: \mathbb{R}/\mathbb{Z} \rightarrow S^1$ ein Homöomorphismus ist.

Beispiel. Zum Abschluß betrachten wir noch Beispiel (d). Zur Erinnerung, für $P, Q \in D^2$ hatten wir folgende Äquivalenzrelation eingeführt:

$$P \sim Q \quad :\Leftrightarrow \quad \begin{aligned} &P = Q \text{ oder} \\ &P \text{ und } Q \text{ liegen auf } S^1 := \{(x, y) \mid x^2 + y^2 = 1\}. \end{aligned}$$

Wir betrachten zuerst die stetige Abbildung

$$\begin{aligned} \varphi: D^2 &\rightarrow S^2 \\ (x, y) &\mapsto (\cos(\varphi) \cdot \sin(\pi r), \sin(\varphi) \cdot \sin(\pi r), \cos(\pi r)). \\ &\quad \uparrow \end{aligned}$$

wir schreiben $(x, y) = (r \cos \varphi, r \sin \varphi)$ mit $r \in [0, 1]$ und $\varphi \in [0, 2\pi]$

Man kann sich leicht vergewissern, dass diese Abbildung in der Tat wohl definiert ist, d.h. dass der Ausdruck auf der rechten Seite nicht von der Wahl der Polarkoordinaten (r, φ) für einen Punkt $(x, y) \in D^2$ abhängt, und dass der Punkt auf der rechten Seite in der Tat in S^2 liegt.

Es gilt zudem, dass $\varphi(x, y) = (0, 0, -1)$ für alle Punkte mit $r = 1$, d.h. für alle Punkte (x, y) in $S^1 = \{(x, y) \mid x^2 + y^2 = 1\}$. Es folgt also insbesondere, dass $\varphi(P) = \varphi(Q)$ für alle Punkte $P, Q \in D^2$ mit $P \sim Q$. Nach Lemma 17.4 existiert daher genau eine Abbildung $g: D^2/\sim \rightarrow S^2$, mit der Eigenschaft, dass $g([P]) = \varphi(P)$ für alle $P \in D^2$. Man kann sich problemlos davon überzeugen, dass g bijektiv ist.

Der topologische Raum D^2 ist kompakt. Aus Satz 1.10 folgt dann auch, dass $D^2/\sim$ kompakt. Nachdem S^2 zudem Hausdorff ist, folgt nun aus Satz 1.13, dass g ein Homöomorphismus ist.

Die anschauliche Darstellung des Arguments ist in Abbildung 103 skizziert. In $D^2/\sim$ fassen wir alle Punkte auf dem Rand zu einem Punkt zusammen. Das “Endergebnis” ist eine 2-dimensionale Sphäre.

17.3. Beispiele: Zylinder, Torus, das Möbius Band und die Kleinsche Flasche. Wir betrachten im Folgenden noch viele weitere Beispiele von Äquivalenzrelationen auf

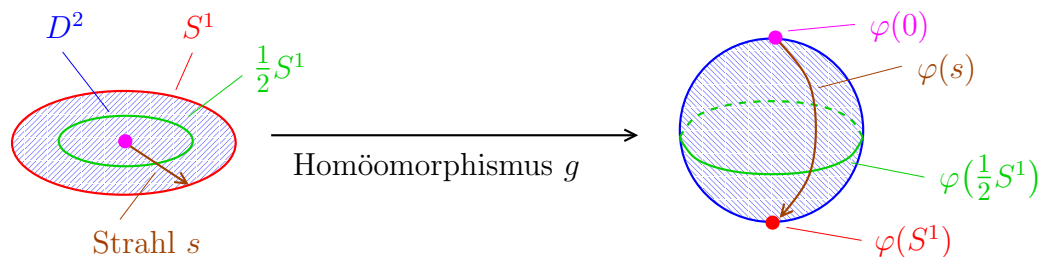


ABBILDUNG 103.

topologischen Räumen. Bei der Beschreibung der Beispiele ist es hilfreich folgende Sprechweise einzuführen. Es sei $\sim$ eine Relation auf X , d.h. es sei eine Teilmenge von $X \times X$ gegeben. Wir sagen $x, y \in X$ sind äquivalent, wenn es $x = x_1, \dots, x_k = y$ in X gibt, so dass für alle $i = 1, \dots, k-1$ gilt: entweder ist $x_i \sim x_{i+1}$ oder es ist $x_{i+1} \sim x_i$. Wir schreiben dann wiederum $x \sim y$, wenn x und y äquivalent sind.

Beispiel. Die Relationen $x \sim (x+1)$ mit $x \in \mathbb{R}$ auf $\mathbb{R}$ erzeugen die schon verwendete Äquivalenzrelation $x \sim y \Leftrightarrow x - y \in \mathbb{Z}$.

(A) Wir betrachten $X = [0, 1] \times [0, 1] \subset \mathbb{R}^2$ und die Äquivalenzrelation, welche erzeugt wird von

$$(x, 0) \sim (x, 1) \quad \text{für alle } x \in [0, 1].^{204}$$

Der Quotientenraum $X/\sim$ entsteht also aus dem Quadrat X , indem wir jeden Punkt auf der oberen Kante mit dem entsprechenden Punkt auf der unteren Kante identifizieren. Etwas anschaulicher, wir erhalten den Quotientenraum indem wir die "obere Kante mit der unteren Kante verkleben". Die Abbildung 104 veranschaulicht diese Operation und suggeriert, dass der Quotientenraum $X/\sim$ ein Zylinder ist. Wir bezeichnen diesen topologischen Raum manchmal auch als Annulus.

Wir können diese Aussage auch leicht beweisen. Wir betrachten zuerst die stetige surjektive Abbildung

$$\begin{aligned} \varphi: X = [0, 1] \times [0, 1] &\rightarrow [0, 1] \times S^1 \\ (x, y) &\mapsto (x, \cos(2\pi y), \sin(2\pi y)). \end{aligned}$$

Nachdem $\varphi(a) = \varphi(b)$ für alle $a \sim b$ induziert φ eine Abbildung $\psi: X/\sim \rightarrow [0, 1] \times S^1$. Diese Abbildung ist stetig und (wie man leicht sieht), bijektiv. Es folgt aus Satz 1.10, dass $X/\sim$ kompakt ist. Nachdem $[0, 1] \times S^1 \subset \mathbb{R}^3$ zudem Hausdorff ist, folgt wiederum aus Satz 1.13, dass ψ ein Homöomorphismus ist, d.h. $X/\sim$ ist in der Tat homöomorph zu einem Zylinder.

²⁰⁴Die Äquivalenzrelation ist also gegeben durch

$$P \sim Q \quad :\Leftrightarrow \quad P = Q \text{ oder } P = (x, 1) \text{ und } Q = (x, 0) \text{ für ein } x \in [0, 1].$$

die Punkte $(x, 1)$ und $(x, 0)$ sind äquivalent

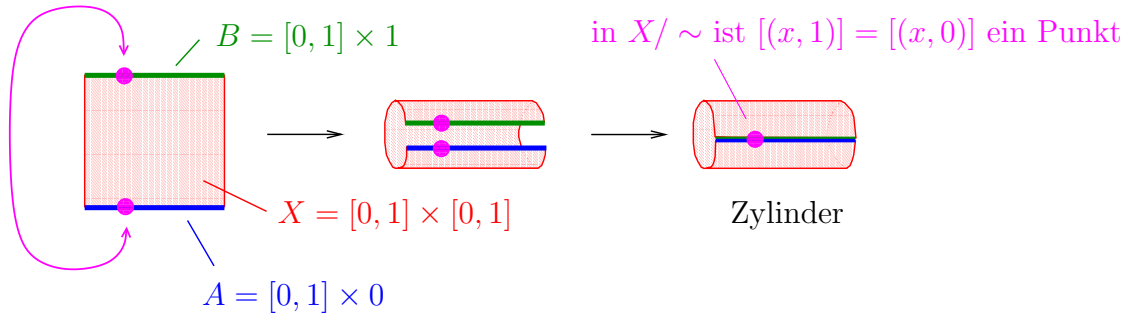


ABBILDUNG 104. Verkleben von zwei Seiten eines Quadrats ergibt einen Zylinder.



der Ausgang links entspricht dem Ausgang rechts, Pacman bewegt sich also auf einem Zylinder

ABBILDUNG 105.

(B) Wir betrachten jetzt $X = [0, 1] \times (0, 1) \subset \mathbb{R}^2$, dieses Mal mit der Äquivalenzrelation, welche erzeugt wird von

$$(0, y) \sim (1, 1 - y) \quad \text{für alle } y \in (0, 1).$$

Der Quotientenraum $X/\sim$ entsteht also aus dem Quadrat X , indem wir die “linke Kante nach einer Verdrillung mit der rechten Kante verkleben”. Den topologischen Raum $X/\sim$ nennen wir das *Möbiusband*. Die Abbildung 106 veranschaulicht diese Operation. Wir überlassen es als Übungsaufgabe nachzuweisen, dass $X/\sim$ homöomorph ist zum Teilraum von $\mathbb{R}^3$, welchen wir im Beweis von Lemma 6.3 eingeführt hatten.

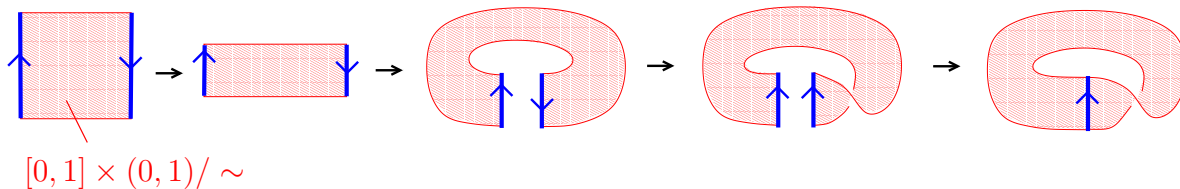


ABBILDUNG 106. Das Möbiusband.

(C) Wir betrachten jetzt wiederum $X = [0, 1] \times [0, 1] \subset \mathbb{R}^2$, aber dieses Mal mit der Äquivalenzrelation, welche erzeugt wird von

$$(x, 0) \sim (x, 1) \quad \text{für alle } x \in [0, 1]$$

und von

$$(0, y) \sim (1, y) \quad \text{für alle } y \in [0, 1].$$

Der Quotientenraum $X/\sim$ entsteht also aus dem Quadrat X , indem wir die “obere Kante mit der unteren Kante verkleben” und die “linke Kante mit der rechten Kante verkleben”. Die Abbildung 107 veranschaulicht diese Operation und suggeriert, dass der Quotientenraum $X/\sim$ ein Torus ist.

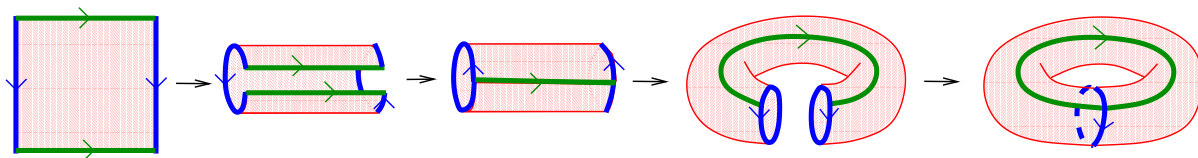


ABBILDUNG 107. Verkleben von den gegenüberliegenden Kanten eines Quadrats ergibt einen Torus.

Man kann nun ohne größere Probleme, ganz analog zum Argument auf Seite 210 zeigen, dass X in der Tat homöomorph ist zum topologischen Raum $S^1 \times S^1$.

(D) Wir betrachten jetzt noch einmal $X = [0, 1] \times [0, 1] \subset \mathbb{R}^2$, jedoch dieses Mal mit der Äquivalenzrelation, welche erzeugt wird von

$$(x, 0) \sim (x, 1) \quad \text{für alle } x \in [0, 1]$$

und von

$$(0, y) \sim (1, 1 - y) \quad \text{für alle } y \in [0, 1].$$

Der Quotientenraum $X/\sim$ entsteht also aus dem Quadrat X , indem wir die “obere Kante mit der unteren Kante verkleben” und die “linke Kante nach einer Verdrillung mit der rechten Kante verkleben”. Den topologischen Raum $X/\sim$ nennen wir die *Kleinsche Flasche*.²⁰⁵ Die Abbildung 108 veranschaulicht diese Operation. In Abbildung 109 sehen wir die Skizze einer stetigen Abbildung von der Kleinschen Flasche in $\mathbb{R}^3$. Diese Abbildung ist nicht injektiv, zwei Kreise auf der Kleinschen Flasche bilden auf den gleichen Kreis in $\mathbb{R}^3$ ab. In der Tat kann man zeigen, aber dies geht weit über diese Vorlesung hinaus, dass es keine injektive stetige Abbildung von der Kleinschen Flasche nach $\mathbb{R}^3$ geben kann.

Es gibt allerdings eine injektive stetige Abbildung von der Kleinschen Flasche nach $\mathbb{R}^4$. Die Konstruktion ist in Abbildung 110 skizziert. Genauer gesagt, es sei $\Phi: X/\sim \rightarrow \mathbb{R}^3$ die Abbildung, welche in Abbildung 109 skizziert ist. Wir wählen eine weitere stetige Abbildung

²⁰⁵Die Kleinsche Flasche ist nach dem Mathematiker Felix Klein (1849-1925) benannt.

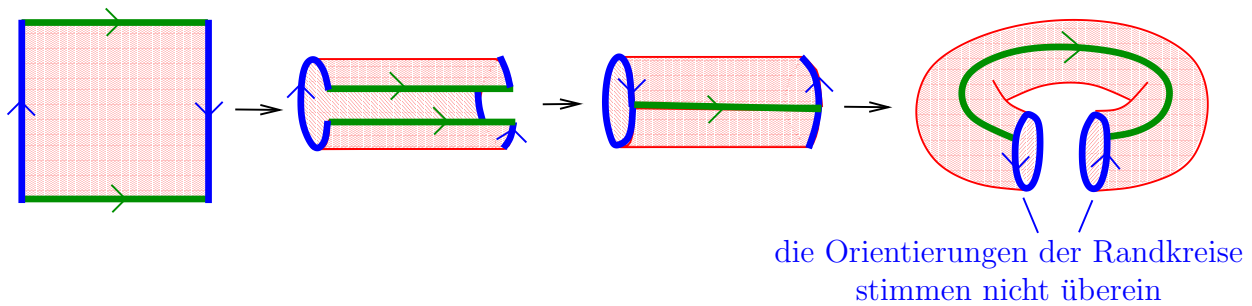


ABBILDUNG 108. Die Definition der Kleinschen Flasche.

jede Abbildung der Kleinschen Flasche nach $\mathbb{R}^3$ hat einen Selbstschnitt

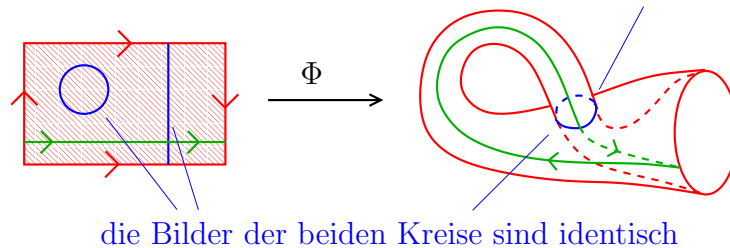


ABBILDUNG 109. Die Kleinsche Flasche kann nicht als Teilmenge von $\mathbb{R}^3$ dargestellt werden.

$\Psi(x, y) = f(x): X/\sim \rightarrow \mathbb{R}$, welche die Eigenschaft hat, dass die Werte für alle Punkte auf A echt kleiner sind als die Werte auf B . Dann ist die Abbildung

$$\begin{aligned} X/\sim &\rightarrow \mathbb{R}^4 = \mathbb{R}^3 \times \mathbb{R} \\ [(x, y)] &\mapsto (\Phi(x, y), f(x)) \end{aligned}$$

injektiv. In der Tat, denn wenn $\Phi(x, y') = \Phi(x, y)$ mit $(x, y) \neq (x', y')$, dann liegen (x, y) und (x', y') auf zwei verschiedenen Kreisen A und B , also unterscheiden sich die f -Werte.

17.4. Beispiele: Flächen von höherem Geschlecht. Wir wollen im folgenden eine mathematische saubere Definition von Flächen von Geschlecht 2 geben. Um diese zu motivieren beginnen wir zuerst mit einer etwas informellen Diskussion.

Wir betrachten erst einmal den topologischen Raum, welcher auf der linken Seite von Abbildung 111 skizziert ist. Wir betrachten hierbei ein Pentagon, bei dem jeweils zwei Seiten mit entgegengesetzter Orientierung identifiziert werden.

In dem Beispiel sind die fünf Eckpunkte äquivalent.²⁰⁶ Daher ist dieser topologische Raum homöomorph zum topologischen Raum, welchen wir aus dem Torus $X = [0, 1] \times [0, 1]/\sim$

²⁰⁶In der Tat, denn $1 \sim 4$ (wegen rot) und $4 \sim 3$ (wegen blau) und $3 \sim 2$ (wegen rot) und $2 \sim 5$ (wegen blau).

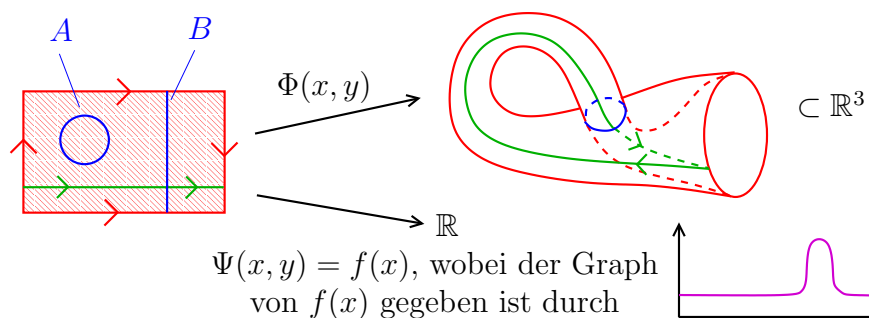


ABBILDUNG 110. Die Kleinsche Flasche kann jedoch als Teilmenge von $\mathbb{R}^4$ aufgefasst werden.

erhalten, indem wir eine offene Scheibe entfernen. Wir erhalten also eine Fläche mit einer Randkomponente.

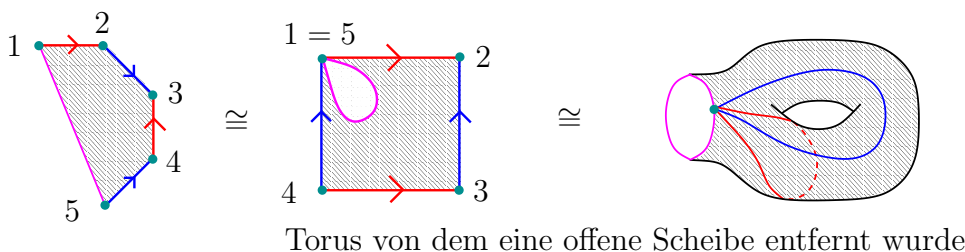


ABBILDUNG 111.

Wir betrachten jetzt das reguläre Oktagon²⁰⁷ E_8 , welches in Abbildung 112 skizziert ist. Wie für den Torus wählen wir eine Äquivalenzrelation, so dass jeweils zwei Kanten, welche durch genau eine Kante getrennt sind, mit "entgegengesetzter Orientierung" äquivalent werden. Das Oktagon mit den Verklebungen wird in Abbildung 112 skizziert. Der topolo-

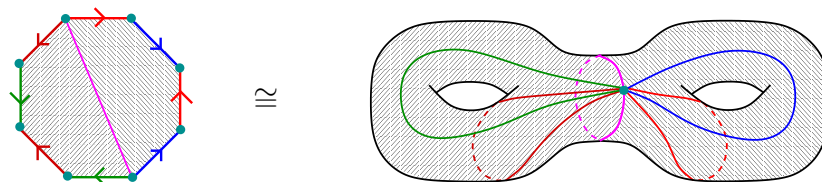


ABBILDUNG 112. Die Fläche von Geschlecht 2.

gische Raum $E_8 / \sim$, den wir dadurch erhalten, ist anschaulich die Fläche von Geschlecht

²⁰⁷Ein n -Eck E im $\mathbb{R}^2$ heißt regulär, wenn alle Kanten die gleiche Länge besitzen und wenn alle Innenwinkel gleich sind.

2. Dies kann man wie folgt sehen: Wir betrachten in E_8 zuerst die rechte obere Hälfte, wir identifizieren dann jeweils zwei Kanten, und erhalten, wie wir gerade gesehen hatten einen “Torus minus eine Scheibe”. Genau das gleiche gilt auch für die linke untere Hälfte. Wir erhalten jetzt $E_8/\sim$, indem wir zwei solche “Tori minus eine Scheibe” am Rand verkleben. Das Ergebnis ist, wie in Abbildung 112 illustriert, eine Fläche von Geschlecht 2.

Wir kehren jetzt zurück zu präziser Mathematik. Genauer gesagt, wir beschreiben jetzt E_8 und die obige Äquivalenzrelation präzise. Dabei bezeichnen wir mit E_8 das reguläre Oktagon mit den Eckpunkten $Q_k = e^{2\pi i k/16}$, wobei $k = 1, 3, \dots, 15$. Für Punkte $A, B \in \mathbb{R}^2 = \mathbb{C}$ bezeichnen wir wie üblich mit $\overline{AB}$ die euklidische Strecke von A nach B . Zudem bezeichnen wir für $z \in S^1$ mit $s_z: \mathbb{C} \rightarrow \mathbb{C}$ die Spiegelung an der euklidischen Geraden $\{tz \mid t \in \mathbb{R}\}$. Wir bezeichnen mit $\sim$ die Äquivalenzrelation auf E_8 , welche erzeugt ist durch

$$P \in \overline{Q_{2k-1}Q_{2k+1}} \sim s_{e^{2\pi i(2k+2)/16}}(P) \in \overline{Q_{2k+3}Q_{2k+5}}$$

für $k = 0, 1, 4, 5$. Diese vier Relationen werden in Abbildung 113 skizziert.²⁰⁸ Die obige Diskussion motiviert jetzt die Konvention, dass wir den topologischen Raum $E_8/\sim$ als die *Fläche von Geschlecht 2* bezeichnen.

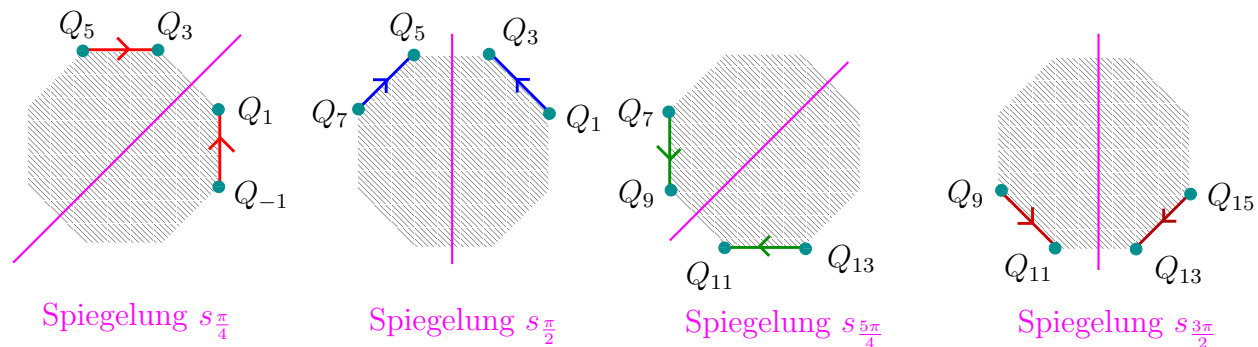


ABBILDUNG 113. Jeweils zwei Kanten werden mithilfe einer Spiegelung identifiziert.

Ganz analog kann man für jedes $g \geq 3$ auch die Fläche von Geschlecht g definieren. Genauer gesagt, wir starten in diesem Fall mit einem regulären $4g$ -Eck und identifizieren für $j = 1, \dots, g$ die Kante $4j + 1$ mit der Kante $4j + 3$ und die Kante $4j + 2$ mit der Kante $4j + 4$, wobei die Identifizierung jedes Mal gegeben ist durch eine Spiegelung. Für $g = 3$ ist diese Konstruktion in Abbildung 114 skizziert.

²⁰⁸Es ist eine amüsante Aufgabe sich davon zu überzeugen, dass bei dieser Äquivalenzrelation alle Eckpunkte des Oktagon äquivalent sind.

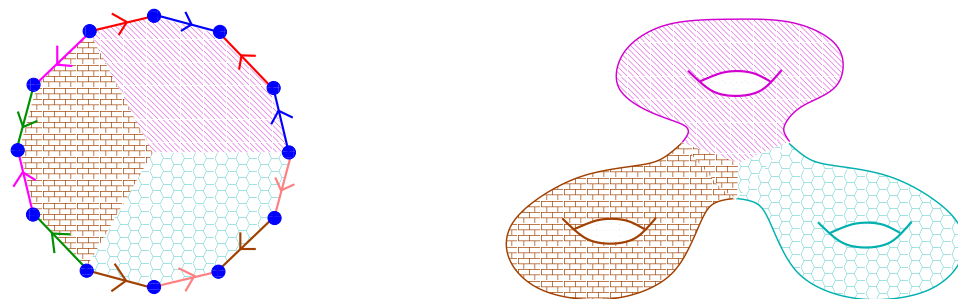


ABBILDUNG 114. Die Fläche von Geschlecht drei.

18. TOPOLOGISCHE MANNIGFALTIGKEITEN

18.1. Zweitabzählbare topologische Räume. Bevor wir die topologischen Mannigfaltigkeiten einführen können, benötigen wir noch folgende etwas technische Definition.

Definition. Ein topologischer Raum heißt *zweitabzählbar*, wenn es eine abzählbare Basis für die Topologie gibt.

Lemma 18.1.

- (1) $\mathbb{R}^n$, mit der üblichen Topologie, ist zweitabzählbar,
- (2) jede Teilmenge von $\mathbb{R}^n$, mit der üblichen Teilraumtopologie, ist zweitabzählbar.

Beweis. Wir betrachten zuerst

$$\mathcal{B} := \{\text{alle Teilmengen von } \mathbb{R}^n \text{ der Form } B_\epsilon(p) \text{ mit } \epsilon \in \mathbb{Q}_{>0} \text{ und } p \in \mathbb{Q}^n\}.$$

Es folgt leicht aus Lemma 16.4, dass $\mathcal{B}$ eine Basis der Topologie von $\mathbb{R}^n$ ist.²⁰⁹ Nachdem $\mathcal{B}$ aus nur abzählbar vielen Mengen besteht ist also $\mathbb{R}^n$ zweitabzählbar. Es sei nun A eine Teilmenge von $\mathbb{R}^n$. Es folgt dann aus Lemma 16.5, dass

$$\mathcal{C} := \{\text{alle Teilmengen von } A \text{ der Form } A \cap X \text{ mit } X \in \mathcal{B}\}$$

eine Basis der Topologie von A ist. Nachdem $\mathcal{C}$ offensichtlich abzählbar ist, folgt, dass A zweitabzählbar ist. $\square$

Lemma 18.2. *Das Produkt von zwei zweitabzählbaren topologischen Räumen ist wiederum zweitabzählbar.*

Beweis. Es seien also X und Y zwei zweitabzählbare topologische Räume. Zudem sei $\{B_i\}_{i \in I}$ beziehungsweise $\{C_j\}_{j \in J}$ eine abzählbare Basis der Topologie von X beziehungsweise Y . Es folgt aus Lemma 16.7, dass $\{B_i \times C_j\}_{(i,j) \in I \times J}$ eine Basis für $X \times Y$ ist. Das Produkt der beiden abzählbaren Mengen I und J ist wiederum abzählbar²¹⁰, also ist $\{B_i \times C_j\}_{(i,j) \in I \times J}$ eine abzählbare Basis für $X \times Y$. $\square$

Folgendes Lemma folgt sofort aus Lemma 17.3.

Lemma 18.3. *Es sei X ein zweitabzählbarer topologischer Raum und es sei $\sim$ eine Äquivalenzrelation auf X . Wenn die Projektionsabbildung $p: X \rightarrow X/\sim$ offen ist, dann ist auch $X/\sim$ zweitabzählbar.*

Beispiel. Die bisherigen beiden Lemmas besagen, dass die meisten Beispiele von topologischen Räumen, welche wir kennen, in der Tat zweitabzählbar sind. Allerdings sind nicht alle topologischen Räume zweitabzählbar. Beispielsweise ist $\mathbb{R}$ mit der *diskreten* Topologie nicht zweitabzählbar.²¹¹

²⁰⁹In der Tat, denn sei $U \subset \mathbb{R}^n$ offen und es sei $x \in U$. Dann gibt es per Definition ein $\epsilon > 0$, so dass $B_\epsilon(x) \subset U$. Wir wählen nun eine rationale Zahl $\eta < \frac{\epsilon}{2}$ und wir wählen einen Punkt $y \in \mathbb{Q}^n$ mit $\|x - y\| < \eta$. Dann gilt $x \in B_\eta(y) \subset B_{\eta+\|x-y\|}(x) \subset B_\epsilon(x) \subset U$. Es folgt nun aus Lemma 16.4, dass $\mathcal{B}$ eine Basis der Topologie von $\mathbb{R}^n$ ist.

²¹⁰Warum ist das Produkt von zwei abzählbaren Mengen wiederum abzählbar?

²¹¹Warum nicht?

18.2. Definition von topologischen Mannigfaltigkeiten. Wir können jetzt den Begriff der “topologischen Mannigfaltigkeit” einführen.

Definition. Es sei X ein topologischer Raum.

- (1) Eine n -dimensionale Karte für X ist ein Homöomorphismus $\Phi: U \rightarrow V$ zwischen einer offenen Teilmenge $U \subset X$ und einer offenen Teilmenge von $\mathbb{R}^n$.
- (2) Ein n -dimensionaler Atlas für X ist eine Familie von n -dimensionalen Karten $\{\Phi_i: U_i \rightarrow V_i\}_{i \in I}$, so dass $\bigcup_{i \in I} U_i = X$.
- (3) Wir sagen X ist eine n -dimensionale topologische Mannigfaltigkeit, wenn X ein zweitabzählbarer Hausdorff-Raum ist und wenn es zu jedem $x \in X$ eine n -dimensionale Karte $\Phi: U \rightarrow V$ mit $x \in U$ gibt.²¹²

Beispiele.

- (1) Jede k -dimensionale Untermannigfaltigkeit M mit $\partial M = \emptyset$ ist eine k -dimensionale topologische Mannigfaltigkeit. In der Tat, denn es gilt:
 - (a) M ist eine Teilmenge von $\mathbb{R}^n$, also insbesondere nach Seite 10 Hausdorff und nach Lemma 18.1 zudem zweitabzählbar,
 - (b) zudem schränkt sich jede Karte $\Phi: U \rightarrow V$, im Sinne der Definition auf Seite 17, auf einen Homöomorphismus $\Phi: U \cap M \rightarrow V \cap E_k \subset E_k = \mathbb{R}^k$ ein. Diese Aussage wird in Abbildung 115 illustriert.
- (2) Es sei X die Gerade mit zwei Nullen, welche wir auf Seite 8 kennengelernt hatten. In Übungsblatt 12 werden wir sehen, dass es zu jedem $x \in X$ eine Karte $\Phi: U \rightarrow V$ mit $x \in U$ gibt. Andererseits hatten wir schon gesehen, dass X nicht Hausdorff ist. Dies zeigt, dass die Hausdorff-Eigenschaft nicht aus der Existenz von Karten folgt.

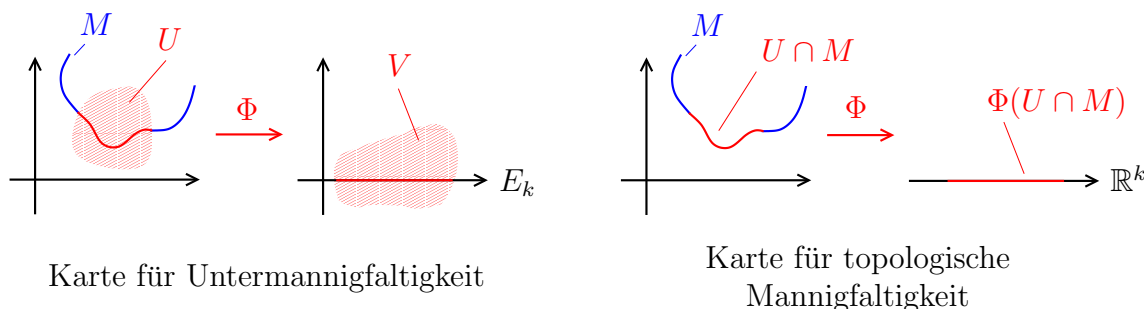


ABBILDUNG 115.

Bemerkung. Die Bedingung, dass eine topologische Mannigfaltigkeit X zweitabzählbar sein soll ist sicher unerwartet. Wir geben hier eine Begründung, und etwas später geben wir noch eine weitere Begründung, warum man diese Einschränkung vornimmt.

²¹²Mit anderen Worten, ein zweitabzählbarer Hausdorff-Raum X ist eine topologische n -dimensionale Mannigfaltigkeit, wenn X einen n -dimensionalen Atlas besitzt.

Man kann sich fragen, was für eindimensionale zusammenhängende topologische Mannigfaltigkeiten es gibt. Wir kennen natürlich S^1 und $\mathbb{R}$, und es erscheint auf den ersten Blick vernünftig, dass jede eindimensionale zusammenhängende topologische Mannigfaltigkeit zu einer der beiden Beispiele homöomorph ist.²¹³ Erstaunlicherweise gibt es jedoch noch einen weiteren topologischen Raum, welcher Hausdorff ist, und welcher einen Atlas besitzt, welcher jedoch nicht zu S^1 oder $\mathbb{R}$ homöomorph ist. Dies ist die sogenannte “lange Gerade”, siehe

[http://en.wikipedia.org/wiki/Long_line_\(topology\)](http://en.wikipedia.org/wiki/Long_line_(topology))

oder

http://de.wikipedia.org/wiki/Lange_Gerade

oder

<http://pages.uoregon.edu/koch/math431/LongLine.pdf>

Dieser topologische Raum ist jedoch nicht zweitabzählbar, und damit keine topologische Mannigfaltigkeit in unserem Sinne. Um dieses exotische Beispiel auszuschließen haben wir in der Definition gefordert, dass eine topologische Mannigfaltigkeit zweitabzählbar sein soll.

Das folgende Lemma gibt ein weiteres Beispiel einer topologischen Mannigfaltigkeit.

Lemma 18.4. *Das Möbiusband ist eine 2-dimensionale topologische Mannigfaltigkeit.*

Bemerkung. Mit den Methoden des Beweises von Lemma 18.4 kann man auch problemlos zeigen, dass auch die anderen topologischen Räume, welche wir in Kapitel 17.3 eingeführt hatten, nämlich der Zylinder, der Torus und die Kleinsche Flasche ebenfalls 2-dimensionale topologische Mannigfaltigkeiten sind. Wir werden den Fall einer Fläche von Geschlecht 2 etwas später auch noch explizit betrachten und wir werden dann zeigen, dass es sich hierbei auch um eine topologische Mannigfaltigkeit handelt.

Beweis. Wir betrachten wie auf Seite 211 den topologischen Raum $X = [0, 1] \times (0, 1) \subset \mathbb{R}^2$ mit der Äquivalenzrelation, welche erzeugt wird von

$$(0, y) \sim (1, 1 - y) \quad \text{für alle } y \in (0, 1).$$

Wir bezeichnen mit

$$\begin{aligned} p: X &\rightarrow X/\sim \\ Q &\mapsto p(Q) = [Q] \end{aligned}$$

die Projektionsabbildung. In Übungsblatt 12 wird gezeigt, dass $X/\sim$ zweitabzählbar ist. Zudem wird in Übungsblatt 11 gezeigt, dass $X/\sim$ Hausdorff ist.²¹⁴

Wir müssen nun also nur noch zeigen, dass es zu jedem Punkt $Q \in X/\sim$ eine 2-dimensionale Karte um Q gibt. Es sei also $Q \in X/\sim$.

²¹³Beispielsweise ist $(-1, 1)$ homöomorph zu $\mathbb{R}$, via $t \mapsto \tan(\frac{\pi}{2}t)$. Andererseits entspricht die obige Definition von topologischer Mannigfaltigkeit der ursprünglichen Definition von Untermannigfaltigkeit auf Seite 17, d.h. wir erlauben keinen Rand.

²¹⁴Nicht alle Eigenschaften einer topologischen Mannigfaltigkeit sind gleich wichtig. Die Hausdorff-Eigenschaft und die Zweitabzählbarkeit kann man in den meisten vernünftigen Fällen problemlos nachweisen. Der Knackpunkt ist normalerweise die Existenz der Karten.

1. Fall Wir nehmen zuerst an, dass $Q = p(x, y)$ mit $x \neq 0, 1$. Wir setzen $V_1 = (0, 1) \times (0, 1)$. Die Einschränkung von p auf V_1 ist ein Homöomorphismus²¹⁵ und $U_1 := p(V_1)$ ist eine offene Umgebung von Q . Also ist $\Phi_1 := p^{-1}: U_1 \rightarrow V_1$ eine Karte um Q .

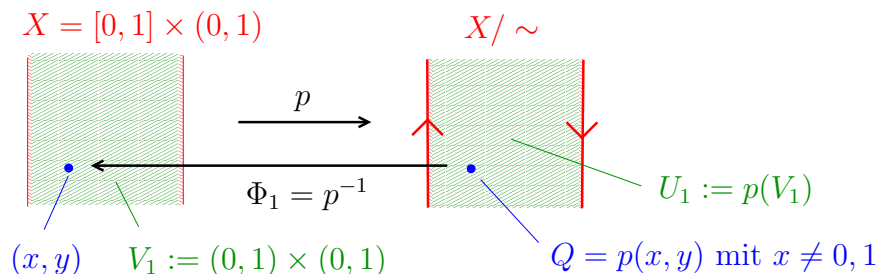


ABBILDUNG 116. Skizze zum Beweis von Lemma 18.4: 1. Fall

2. Fall Wir nehmen nun an, dass $Q = p(0, y)$ für ein $y \in (0, 1)$. Wir setzen

$$W := [0, \frac{1}{4}] \times (0, 1) \cup (\frac{3}{4}, 1] \times (0, 1).$$

Dies ist offensichtlich eine offene Menge von X . Wir setzen nun $U_2 := p(W)$. Nachdem $p^{-1}(U_2) = p^{-1}(p(W)) = W$ offen ist, ist U_2 eine offene Umgebung von $Q = p(0, y)$. Wir betrachten nun die Abbildung

$$\begin{aligned} \Phi_2: U_2 = p(W) &\rightarrow V_2 := (-\frac{1}{4}, \frac{1}{4}) \times (0, 1) \\ p(a, b) &\mapsto \begin{cases} (a, b), & \text{wenn } (a, b) \in [0, \frac{1}{4}] \times (0, 1), \\ (a - 1, 1 - b), & \text{wenn } (a, b) \in (\frac{3}{4}, 1] \times (0, 1). \end{cases} \end{aligned}$$

Diese Abbildung ist, wie man leicht zeigen kann, wohl-definiert²¹⁶, bijektiv, stetig und die Umkehrabbildung ist ebenfalls stetig. Also ist $\Phi_2: U_2 \rightarrow V_2$ eine Karte um $Q = p(0, y)$.

Wir haben nun also einen Atlas $\{\Phi_i: U_i \rightarrow V_i\}_{i=1,2}$ für das Möbiusband gefunden.²¹⁷ $\square$

18.3. Gruppenoperationen. Um weitere Beispiele von topologischen Mannigfaltigkeiten zu konstruieren, führen wir den Begriff von Gruppenoperationen ein.

Definition. Es sei X eine Menge und G eine Gruppe mit trivialem Element e .

- (1) Eine *Gruppenoperation* (oder kurz *Operation*) von G auf X ist eine Abbildung

$$\begin{aligned} G \times X &\rightarrow X \\ (g, x) &\mapsto g \cdot x \end{aligned}$$

mit folgenden Eigenschaften

$$\begin{aligned} e \cdot x &= x, & \text{für alle } x \in X, \\ g \cdot (h \cdot x) &= (gh) \cdot x, & \text{für alle } x \in X \text{ und } g, h \in G. \end{aligned} \quad ^{218}$$

²¹⁵Warum?

²¹⁶D.h. wenn $p(a, b) = p(a', b')$, dann gilt auch nach der obigen Definition $\Phi_2(p(a, b)) = \Phi_2(p(a', b'))$.

²¹⁷Warum müssen wir nicht auch noch den Fall betrachten, dass $Q = p(1, y)$ für ein $y \in (0, 1)$?

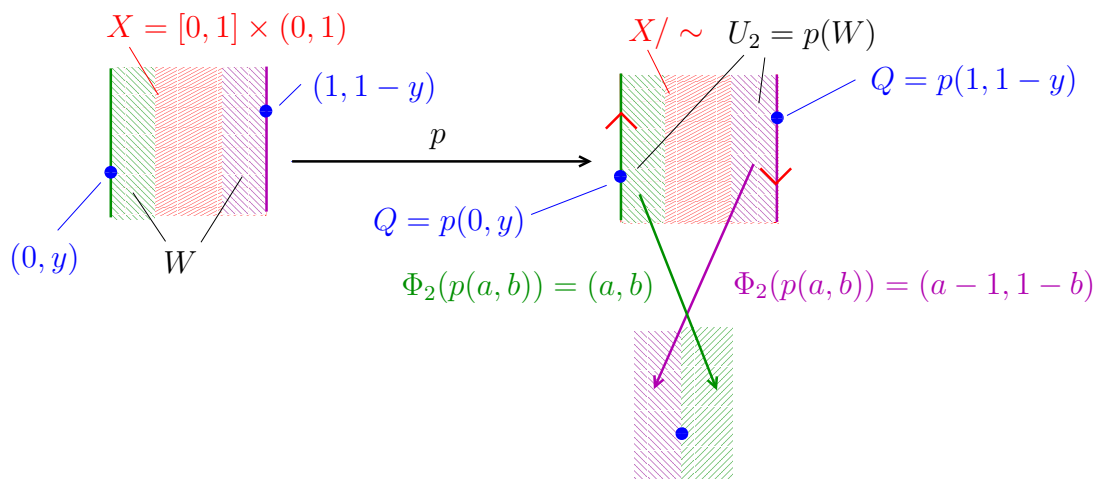


ABBILDUNG 117. Skizze zum Beweis von Lemma 18.4: 2. Fall.

- (2) Die Operation heißt *transitiv*, wenn es zu allen $x, y \in X$ ein $g \in G$ mit $g \cdot x = y$ gibt.
- (3) Die Operation heißt *frei*, wenn $g \cdot x = x$ für ein $x \in X$ impliziert, dass $g = e$.

Beispiele.

- (1) Die orthogonale Gruppe $O(n)$ operiert auf $\mathbb{R}^n$ durch die übliche Multiplikation. Diese Operation ist weder transitiv noch frei.²¹⁹
- (2) Die Möbiustransformationen bilden nach Lemma 14.1 eine Gruppe, welche nach Satz 14.6 auf der Menge der hyperbolischen Geraden operiert. Lemma 14.7 besagt, dass diese Operation transitiv ist. Allerdings ist die Operation nicht frei.²²⁰
- (3) Es sei X die Menge aller unangeordneten Basen von $\mathbb{R}^n$. Dann ist

$$\begin{aligned} \text{GL}(n, \mathbb{R}) \times X &\mapsto X \\ (A, \{v_1, \dots, v_n\}) &\mapsto \{Av_1, \dots, Av_n\} \end{aligned}$$

eine transitive Operation, welche jedoch nicht frei ist.

- (4) Es sei S_n die Permutationsgruppe der Menge $\{1, \dots, n\}$. Dann ist

$$\begin{aligned} S_n \times \mathbb{R}^n &\rightarrow \mathbb{R}^n \\ (\sigma, (v_1, \dots, v_n)) &\mapsto (v_{\sigma(1)}, \dots, v_{\sigma(n)}) \end{aligned}$$

eine Operation, welche weder frei noch transitiv ist.

²¹⁸Links wenden wir also zwei Mal die Operation an, während wir rechts erst die beiden Elemente in der Gruppe multiplizieren, und dann auf x anwenden.

²¹⁹Warum nicht?

²²⁰Auch hier, warum nicht?

- (5) Jede Gruppe G operiert auf sich selbst durch Linksmultiplikation. Diese Operation ist transitiv und frei.

Folgendes Lemma folgt leicht aus den Definitionen und verbleibt eine freiwillige Übungsaufgabe.

Lemma 18.5. *Es sei X eine Menge und G eine Gruppe, welche auf X operiert. Dann ist*

$$x \sim y \iff \text{es existiert ein } g \in G, \text{ sodass } g \cdot x = y$$

eine Äquivalenzrelation auf X .

Es sei X eine Menge und G eine Gruppe, welche auf X operiert. Wir bezeichnen mit $\sim$ die Äquivalenzrelation aus Lemma 18.5. Wir definieren dann

$$X/G := X/\sim.$$

Im weiteren Verlauf der Vorlesung werden wir folgendes elementare Lemma mehrmals verwenden.

Lemma 18.6. *Es sei G eine Gruppe, welche auf einer Menge X operiert. Wir bezeichnen mit $p: X \rightarrow X/G$ die Projektionsabbildung. Es seien A und B Teilmengen von X . Dann gilt*

$$p(A) \cap p(B) = \emptyset \iff \text{für alle } g \in G \text{ ist } gA \cap B = \emptyset$$

und die äquivalente Aussage

$$p(A) \cap p(B) \neq \emptyset \iff \text{es gibt ein } g \in G \text{ mit } gA \cap B \neq \emptyset.$$

Beweis. Wir beweisen die zweite Aussage, die erste Aussage ist zu dieser äquivalent. Es gilt

$$\begin{aligned} p(A) \cap p(B) \neq \emptyset &\iff \text{es gibt } a \in A \text{ und } b \in B \text{ mit } p(a) = p(b) \in X/G \\ &\iff \text{es gibt } a \in A \text{ und } b \in B \text{ und ein } g \in G \text{ mit } ga = b \\ &\iff \text{es gibt ein } g \in G \text{ mit } gA \cap B \neq \emptyset. \end{aligned} \quad \square$$

18.4. Stetige Operationen.

Definition. Es sei X ein topologischer Raum und G eine Gruppe, welche auf X operiert. Wir sagen die Operation ist *stetig*, wenn für jedes $g \in G$ die Abbildung

$$\begin{aligned} X &\rightarrow X \\ x &\mapsto g \cdot x \end{aligned}$$

stetig ist.

Beispiele.

- (A) Es sei $X = \mathbb{R}^n$ und $G = \mathbb{Z}^n$. Dann ist

$$\begin{aligned} \mathbb{Z}^n \times \mathbb{R}^n &\rightarrow \mathbb{R}^n \\ (z, v) &\mapsto z + v \end{aligned}$$

eine stetige und freie Operation.

(B) Es sei $X = \mathbb{R} \times (-1, 1)$ und $G = \mathbb{Z}$. Dann ist

$$\begin{aligned} \mathbb{Z} \times (\mathbb{R} \times (-1, 1)) &\rightarrow \mathbb{R} \times (-1, 1) \\ (n, (x, y)) &\mapsto (x + n, (-1)^n y) \end{aligned}$$

eine Operation²²¹, welche ebenfalls stetig und frei ist.

(C) Es sei $X = \mathbb{R}$ und $G = \{\pm 1\}$. Dann ist

$$\begin{aligned} \{\pm 1\} \times \mathbb{R} &\rightarrow \mathbb{R} \\ (\epsilon, x) &\mapsto \epsilon \cdot x \end{aligned}$$

eine stetige Operation. Diese ist jedoch nicht frei, denn $(-1) \cdot 0 = 0$, aber -1 ist nicht das triviale Element der Gruppe $G = \{\pm 1\}$.

(D) Es sei $X = S^n$ und $G = \{\pm 1\}$. Dann ist

$$\begin{aligned} \{\pm 1\} \times S^n &\rightarrow S^n \\ (\epsilon, P) &\mapsto \epsilon \cdot P \end{aligned}$$

eine stetige und freie Operation.

(E) Es sei $X = \mathbb{R}$ und $G = \mathbb{Q}$. Dann ist

$$\begin{aligned} \mathbb{Q} \times \mathbb{R} &\rightarrow \mathbb{R} \\ (r, x) &\mapsto r + x, \end{aligned}$$

ganz analog zu Beispiel (A) eine stetige und freie Operation.

(F) Die Abbildung

$$\begin{aligned} (2, \mathbb{R}) \times \mathbb{H} &\rightarrow \mathbb{H} \\ \left(\begin{pmatrix} a & b \\ c & d \end{pmatrix}, z \right) &\mapsto \frac{az + b}{cz + d} \end{aligned}$$

ist eine stetige und transitive Operation.²²²

²²¹Warum ist dies eine Operation?

²²²Die Aussage, dass dies in der Tat eine Operation ist folgt aus folgender nicht besonders erhellenden Rechnung: für $z \in \mathbb{H}$ ist

$$\text{id} \cdot z = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \cdot z = \frac{1 \cdot z + 0}{0 \cdot z + 1} = z.$$

Zudem gilt für

$$G = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \quad \text{und} \quad G' = \begin{pmatrix} a' & b' \\ c' & d' \end{pmatrix}$$

in $(2, \mathbb{R})$ und $z \in \mathbb{H}$, dass

$$\begin{aligned} G' \cdot (G \cdot z) &= \begin{pmatrix} a' & b' \\ c' & d' \end{pmatrix} \cdot \left(\begin{pmatrix} a & b \\ c & d \end{pmatrix} z \right) = \begin{pmatrix} a' & b' \\ c' & d' \end{pmatrix} \cdot \frac{az + b}{cz + d} = \frac{a' \frac{az+b}{cz+d} + b'}{c' \frac{az+b}{cz+d} + d'} \\ &= \frac{(aa' + b'c)z + (a'b + b'd)}{(c'a + d'c)z + (c'b + d'd)} = \begin{pmatrix} aa' + b'c & a'b + b'd \\ c'a + d'c & c'b + d'd \end{pmatrix} \cdot z = (G' \cdot G) \cdot z. \end{aligned}$$

Die Aussage, dass diese eine transitive Operation ist, ist eine Umformulierung von Lemma 14.4.

Lemma 18.7. *Es sei G eine Gruppe, welche stetig auf einem topologischen Raum X operiert. Für jedes $g \in G$ ist die Abbildung*

$$\begin{aligned} X &\rightarrow X \\ x &\mapsto g \cdot x \end{aligned}$$

ein Homöomorphismus.

Beweis. Es sei G eine Gruppe, welche stetig auf einem topologischen Raum X operiert und es sei $g \in G$. Die Abbildung

$$\begin{aligned} X &\rightarrow X \\ x &\mapsto g^{-1} \cdot x \end{aligned}$$

ist nach Voraussetzung ebenfalls stetig. Zudem ist diese Abbildung die Umkehrabbildung der gegebenen Abbildung, denn für alle $x \in X$ ist

$$\begin{array}{ccccc} g^{-1} \cdot (g \cdot x) & = & (g^{-1}g) \cdot x & = & e \cdot x = x. \\ \uparrow & & & & \uparrow \\ \text{zweites Axiom einer Operation} & & & & \text{erstes Axiom einer Operation.} \end{array}$$

Genau das gleiche Argument zeigt auch, dass $g \cdot (g^{-1} \cdot x) = x$ für alle x . $\square$

Es sei X ein topologischer Raum und G eine Gruppe, welche auf X stetig operiert. Wir fassen dann X/G immer als topologischen Raum bezüglich der Quotiententopologie auf, welche wir in Kapitel 17.2 eingeführt hatten. Zur Erinnerung, dies bedeutet, dass $U \subset X/G$ offen ist, wenn unter der Projektionsabbildung $p: X \rightarrow X/G$ das Urbild $p^{-1}(U) \subset X$ offen ist.

Wir betrachten nun noch die Quotientenräume von einigen der vorherigen Beispiele.

Beispiele.

- (A) Wir betrachten wiederum die stetige Operation von $\mathbb{Z}^n$ auf $\mathbb{R}^n$. Auf Seite 209 hatten wir gesehen, dass $\mathbb{R}/\mathbb{Z}$ homöomorph zu S^1 ist. Genau der gleiche Beweis zeigt nun, dass

$$\begin{aligned} \mathbb{R}^n/\mathbb{Z}^n &\rightarrow (S^1)^n = \overbrace{S^1 \times \dots \times S^1}^{n\text{-Mal}} \\ [(x_1, \dots, x_n)] &\mapsto (e^{2\pi i x_1}, \dots, e^{2\pi i x_n}) \end{aligned}$$

ein Homöomorphismus ist. Wir bezeichnen $\mathbb{R}^n/\mathbb{Z}^n \cong (S^1)^n$ Raum als den *n-dimensionalen Torus*.

- (B) Es sei $X = [0, 1] \times (0, 1)$ und es sei $\sim$ die Äquivalenzrelation von Seite 211, welche das Möbiusband definiert. Man kann nun leicht zeigen, dass die Abbildung

$$\begin{aligned} X/\sim &\rightarrow (\mathbb{R} \times (-1, 1))/\mathbb{Z} \\ (x, y) &\mapsto [(x, 1 - 2y)] \end{aligned}$$

wohldefiniert ist, und dass dies ein Homöomorphismus ist. Mit anderen Worten, $(\mathbb{R} \times (-1, 1))/\mathbb{Z}$ ist homöomorph zum Möbiusband. Im weiteren Verlauf bezeichnen wir oft auch $(\mathbb{R} \times (-1, 1))/\mathbb{Z}$ als das Möbiusband.

- (C) Wir betrachten noch einmal die Operation von $G = \{\pm 1\}$ auf $X = \mathbb{R}$. Man kann nun leicht zeigen, dass

$$\begin{aligned} \mathbb{R}/\{\pm 1\} &\rightarrow [0, \infty) \\ [x] &\mapsto |x| \end{aligned}$$

ein Homöomorphismus ist.

- (D) Wir bezeichnen den Quotientenraum $S^n/\{\pm 1\}$ als den *n-dimensionalen projektiven Raum* $\mathbb{R}P^n$. Der reelle *n*-dimensionale projektive Raum wird normalerweise definiert als

$$\mathbb{R}P^n = \text{Menge aller Geraden in } \mathbb{R}^{n+1} = (\mathbb{R}^{n+1} \setminus \{0\})/\sim,$$

wobei $v \sim w$, wenn $v = \lambda w$ für ein $\lambda \in \mathbb{R}$. Die Abbildung $S^n/\sim \rightarrow (\mathbb{R}^{n+1} \setminus \{0\})/\sim$, welche gegeben ist durch $[v] \mapsto [v]$, ist jedoch offensichtlich eine Bijektion. Der projektive Raum spielt eine wichtige Rolle in der “algebraischen Geometrie” und der “Topologie”.

- (E) Wir betrachten wiederum die Operation

$$\begin{aligned} \mathbb{Q} \times \mathbb{R} &\rightarrow \mathbb{R} \\ (r, x) &\mapsto r + x. \end{aligned}$$

In diesem Fall ist der Quotientenraum $\mathbb{R}/\mathbb{Q}$ kein “alter Bekannter”, sondern ein eher eigenwilliger Raum. Beispielsweise ist $\mathbb{R}/\mathbb{Q}$ kein Hausdorff-Raum. In der Tat, denn es seien P und Q in $\mathbb{R}/\mathbb{Q}$ und es seien U und V offene Umgebungen von P und Q . Wir wollen zeigen, dass $U \cap V \neq \emptyset$. Wir setzen $A = p^{-1}(U)$ und $B = p^{-1}(V)$. Dann ist $U = p(p^{-1}(U)) = p(A)$ und $V = p(p^{-1}(V)) = p(B)$. Wir müssen also zeigen, dass $p(A) \cap p(B) \neq \emptyset$. Nachdem A eine offene Teilmenge von $\mathbb{R}$ ist gibt es ein $r \in \mathbb{Q}$ mit $(r + A) \cap B \neq \emptyset$. Also folgt aus Lemma 18.6, dass $p(A) \cap p(B) \neq \emptyset$.

Wir wollen jetzt ein Kriterium dafür finden, dass für einen Hausdorff-Raum X auch der Quotientenraum X/G wiederum Hausdorff ist.

Definition. Es sei X ein topologischer Raum und G eine Gruppe, welche auf X operiert. Wir sagen, G *operiert eigentlich*, wenn es für alle x und y in X offene Umgebungen U von x und V von y gibt, so dass die Menge $\{g \in G \mid gU \cap V \neq \emptyset\}$ endlich ist.²²³

Wir betrachten wiederum einige der vorherigen Beispiele.

Beispiele.

²²³Diese Definition ist nicht die Standarddefinition von einer eigentlichen Operation. Normalerweise sagt man, dass eine Gruppe G eigentlich operiert, wenn für alle kompakten Teilmengen K und L die Menge $\{g \in G \mid gK \cap L \neq \emptyset\}$ endlich ist. Wir interessieren uns im Folgenden nur für Operationen auf topologischen Mannigfaltigkeiten, und mit etwas Aufwand kann man zeigen, dass in diesem Spezialfall die beiden Definitionen äquivalent sind.

- (A) Die Operation von $G = \mathbb{Z}^n$ auf $X = \mathbb{R}^n$ ist eigentlich. In der Tat, denn es seien P und Q zwei Punkte in $\mathbb{R}^n$. Dann gilt für jede Wahl von beschränkten Umgebungen U und V , dass $\{z \in \mathbb{Z}^n \mid z + U \cap V \neq \emptyset\}$ endlich ist.²²⁴
- (B) Ein ähnliches Argument wie in (A) zeigt, dass die Operation der Gruppe $G = \mathbb{Z}$ auf $X = \mathbb{R} \times (-1, 1)$ eigentlich ist.
- (C&D) Jede stetige Operation einer endlichen Gruppe ist offensichtlich eigentlich. Insbesondere sind die Operationen der Gruppe $G = \{\pm 1\}$ auf $X = \mathbb{R}$ und $X = S^n$ eigentlich.
- (E) Die Operation von $G = \mathbb{Q}$ auf $X = \mathbb{R}$ ist nicht eigentlich. Diese Aussage kann man leicht ganz explizit nur mithilfe der Definitionen zeigen. Die Aussage folgt auch aus der oben bewiesenen Tatsache, dass $\mathbb{R}/\mathbb{Q}$ nicht Hausdorff ist, zusammen mit dem nächsten Satz.

Satz 18.8. *Es sei G eine Gruppe, welche auf einem Hausdorff-Raum X stetig operiert. Wenn die Operation zudem eigentlich ist, dann ist der Quotientenraum X/G ebenfalls Hausdorff.*

Im Beweis von Satz 18.8 werden wir folgende zwei Lemmas verwenden.

Lemma 18.9. *Es sei G eine Gruppe, welche stetig auf einem topologischen Raum X operiert. Dann ist die Projektionsabbildung $X \rightarrow X/G$ offen.*²²⁵

Beweis von Lemma 18.9. Wir müssen also zeigen, dass p offen ist. Es sei also $U \subset X$ offen. Wir müssen zeigen, dass $p(U)$ offen in X/G ist, d.h. wir müssen zeigen, dass $p^{-1}(p(U))$ offen in X ist. Es folgt aus den Definitionen, dass²²⁶

$$p^{-1}(p(U)) = \bigcup_{g \in G} gU.$$

Nachdem G stetig operiert, ist die Multiplikationsabbildung $x \mapsto gx$ nach Lemma 18.7 ein Homöomorphismus. Also ist gU offen in X . Es folgt, dass $p^{-1}(p(U))$ als Vereinigung von offenen Mengen, offen in X ist. $\square$

Wir benötigen auch noch folgendes Lemma.

Lemma 18.10. *Es sei G eine Gruppe, welche auf einem Hausdorff-Raum X stetig und eigentlich operiert. Es seien a und b zwei Punkte in X . Dann gibt es offene Umgebungen A von a und B von b mit folgender Eigenschaft: für alle $g \in G$ mit $ga \neq b$ gilt auch $gA \cap B = \emptyset$.*

²²⁴Da U und V beschränkt sind gibt es ein $C \geq 0$, so dass $\|u\| \leq d$ und $\|v\| \leq d$ für alle $u \in U$ und $v \in V$. Für $z \in \mathbb{Z}^n$ mit $z + U \cap V \neq \emptyset$ gibt es also $u \in U$ und $v \in V$ mit $z = v - u$, also folgt aus der Dreiecksungleichung, dass $\|z\| \leq 2d$. Aber es gibt nur endlich viele $z \in \mathbb{Z}^n$ mit $\|z\| \leq 2d$.

²²⁵Zur Erinnerung, wir hatten auf Seite 207 eine Abbildung $f: X \rightarrow Y$ zwischen topologischen Räumen *offen* genannt, wenn das Bild jeder offenen Menge in X offen in Y ist.

²²⁶Warum folgt das "aus den Definitionen"?

Beweis von Lemma 18.10. Es sei G eine Gruppe, welche auf einem Hausdorff-Raum X stetig und eigentlich operiert. Es seien a und b zwei Punkte in X .

Nachdem G eigentlich auf X operiert gibt es offene Umgebungen U von a und V von b , so dass $\{g \in G \mid gU \cap V \neq \emptyset\}$ eine endliche Menge ist. Wir bezeichnen die Elemente in dieser Menge mit $g_1, \dots, g_r$. Wir setzen nun $I := \{i = 1, \dots, r \mid g_i a \neq b\}$. Nachdem X ein

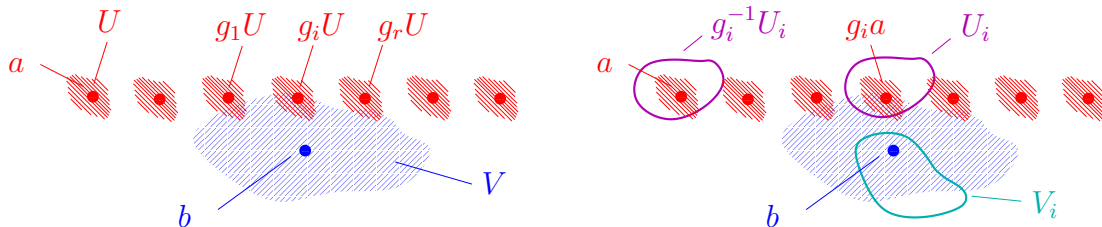


ABBILDUNG 118. Skizze zum Beweis von Lemma 18.10 mit $I = \{1, \dots, r\}$.

Hausdorff-Raum ist, existieren für jedes $i \in I$ eine offene Umgebung U_i von $g_i a$ und eine offene Umgebung V_i von b , so dass $U_i \cap V_i = \emptyset$. Wir setzen nun

$$A := U \cap \bigcap_{i \in I} g_i^{-1} U_i \quad \text{und} \quad B := V \cap \bigcap_{i \in I} V_i.$$

Die Mengen A und B sind als Durchschnitt von endlich vielen offenen Mengen²²⁷ wiederum offen. Insbesondere ist also A eine offene Umgebung von a und B ist eine offene Umgebung von b .

Man kann nun leicht verifizieren, dass A und B die gewünschten Eigenschaften besitzen. In der Tat, es sei $g \in G$ mit $ga \neq b$. Wenn $g \notin \{g_1, \dots, g_r\}$, dann ist schon $gU \cap V = \emptyset$, also ist $gA \cap B = \emptyset$. Wenn $g \in \{g_1, \dots, g_r\}$, dann ist $g = g_i$ für ein $i \in I$. Es folgt, dass $g_i(g_i^{-1}U_i) \cap V_i = U_i \cap V_i = \emptyset$, also ist auch $g_i A \cap B = \emptyset$. $\square$

Wir wenden uns nun dem eigentlichen Beweis von Satz 18.8 zu.

Beweis von Satz 18.8. Es sei G eine Gruppe, welche auf einem Hausdorff-Raum X stetig und eigentlich operiert. Wir wollen zeigen, dass der Quotientenraum X/G ebenfalls Hausdorff ist. Wir bezeichnen mit $p: X \rightarrow X/G$ die Projektionsabbildung. Es seien x und y zwei verschiedene Punkte in X/G . Wir wählen a und b in X mit $p(a) = x$ und $p(b) = y$. Nachdem $p(a) = x \neq y = p(b)$ gilt $ga \neq b$ für alle $g \in G$. Nach Lemma 18.10 gibt es nun offene Umgebungen A von a und B von b , so dass $gA \cap B = \emptyset$ für alle $g \in G$. Es folgt aus Lemma 18.6, dass $p(A)$ und $p(B)$ disjunkt sind.

Nachdem A und B offen sind folgt zudem aus Lemma 18.9, dass $p(A)$ und $p(B)$ offene Teilmengen von X/G und insbesondere offene Umgebungen von $x = p(a)$ und $y = p(b)$ sind. Wir haben also die gewünschten disjunkten Umgebungen von x und y in X/G gefunden. $\square$

²²⁷Hierbei verwenden wir, dass G stetig operiert, denn dies impliziert, dass die Mengen $h^{-1}U_i$ wiederum offen sind.

18.5. Gruppenoperationen auf topologischen Mannigfaltigkeiten. Es sei nun M eine topologische Mannigfaltigkeit und G eine Gruppe, welche auf M operiert. Der folgende Satz besagt nun, dass unter vernünftigen Voraussetzungen der Quotient M/G wiederum eine topologische Mannigfaltigkeit ist.

Satz 18.11. *Es sei M eine n -dimensionale topologische Mannigfaltigkeit und es sei G eine Gruppe, welche auf M stetig, frei und eigentlich operiert. Dann ist auch M/G eine n -dimensionale topologische Mannigfaltigkeit.*

Für den Beweis von Satz 18.11 benötigen wir noch folgende zwei, etwas technische Lemmas.

Lemma 18.12. *Es sei X ein Hausdorff-Raum und es sei G eine Gruppe, welche auf X stetig, frei und eigentlich operiert. Dann gibt es zu jedem $x \in X$ eine offene Umgebung U , so dass $gU \cap U = \emptyset$ für alle $g \neq e$.²²⁸*

Beweis von Lemma 18.12. Es sei X also ein Hausdorff-Raum und es sei G eine Gruppe, welche auf X stetig, frei und eigentlich operiert. Zudem sei x ein Punkt in X . Nachdem G frei operiert gilt $gx \neq x$ für alle $g \neq e$.

Wir wenden Lemma 18.10 auf $a = b = x$ an und erhalten offene Umgebungen A und B von x mit $gA \cap B = \emptyset$ für alle $g \neq e$. Aber dann besitzt $U := A \cap B$ die gewünschte Eigenschaft. $\square$

Lemma 18.13. *Es sei X ein Hausdorff-Raum und es sei G eine Gruppe, welche auf X stetig operiert. Wir bezeichnen mit $p: X \rightarrow X/G$ die Projektionsabbildung. Es sei nun $U \subset X$ eine offene Teilmenge, so dass die Abbildung $p: U \rightarrow X/G$ injektiv ist. Dann ist die Abbildung $p: U \rightarrow p(U)$ ein Homöomorphismus.*

Beweis. Die Projektionsabbildung $p: X \rightarrow X/G$ ist nach Lemma 17.2 stetig, und damit ist auch die Einschränkung von p auf U stetig. Die Abbildung $p: U \rightarrow p(U)$ ist nach Voraussetzung injektiv und offensichtlich surjektiv, also eine Bijektion.

Es verbleibt zu zeigen, dass $q := p^{-1}: p(U) \rightarrow U$ stetig ist. Es sei also $W \subset U$ offen. Dann ist $q^{-1}(W) = (p^{-1})^{-1}(W) = p(W)$ offen nach Lemma 18.9. Also ist $p^{-1}: p(U) \rightarrow U$ stetig. $\square$

Wir sind nun in der Lage Satz 18.11 zu beweisen.

Beweis von Satz 18.11. Es sei M eine n -dimensionale topologische Mannigfaltigkeit und es sei G eine Gruppe, welche auf M stetig, frei und eigentlich operiert. Wir wollen zeigen, dass M/G eine n -dimensionale topologische Mannigfaltigkeit ist. Wir müssen also folgende drei Aussagen beweisen:

²²⁸Ganz allgemein sagt man, dass eine Gruppe G auf einem topologischen Raum X *diskret* operiert, wenn es für jedes $x \in X$ eine offene Umgebung U gibt, so dass $gU \cap U = \emptyset$ für alle $g \neq e$. Das Lemma besagt also, dass eine Gruppe, welche frei, stetig und eigentlich auf einem Hausdorff-Raum operiert auch diskret operiert. Gilt diese Aussage auch ohne die Voraussetzung "Hausdorff-Raum" oder ohne die Voraussetzung "frei"?

- (1) M/G ist zweitabzählbar,
- (2) M/G ist Hausdorff, und
- (3) um jeden Punkt $y \in M/G$ gibt es eine n -dimensionale Karte.

Die erste Aussage folgt aus Lemmas 18.3 und 18.9. Die zweite Aussage folgt sofort aus Satz 18.8. Wir wenden uns nun der dritten Aussage zu. Wir bezeichnen mit $p: M \rightarrow M/G$ die Projektionsabbildung. Es sei $y \in M/G$. Wir wählen ein $x \in M$ mit $p(x) = y$.

- (1) Nachdem M eine n -dimensionale topologische Mannigfaltigkeit ist, existiert eine n -dimensionale Karte $\Phi: V \rightarrow W$ um x .
- (2) Lemma 18.12 besagt, dass es eine offene Umgebung U von x gibt, so dass $gU \cap U \neq \emptyset$ für alle $g \neq e$. Dies ist äquivalent zu der Aussage, dass die Einschränkung von $p: M \rightarrow M/G$ auf U injektiv ist.²²⁹

Indem wir U und V durch $U \cap V$ ersetzen können wir o.B.d.A. annehmen, dass $U = V$. Wir betrachten nun die Abbildungen

$$p(U) \xrightarrow{p^{-1}} U \xrightarrow{\Phi} W.$$

Die Abbildung Φ ist ein Homöomorphismus, und es folgt aus Lemma 18.13, dass auch die Abbildung $p: U \rightarrow p(U) \subset M/G$ ein Homöomorphismus ist. Also ist die obige Abbildung $\Phi \circ p^{-1}: p(U) \rightarrow W$ ein Homöomorphismus, und insbesondere eine n -dimensionale Karte für M/G um y . Wir haben damit bewiesen, dass M/G eine n -dimensionale topologische Mannigfaltigkeit ist. $\square$

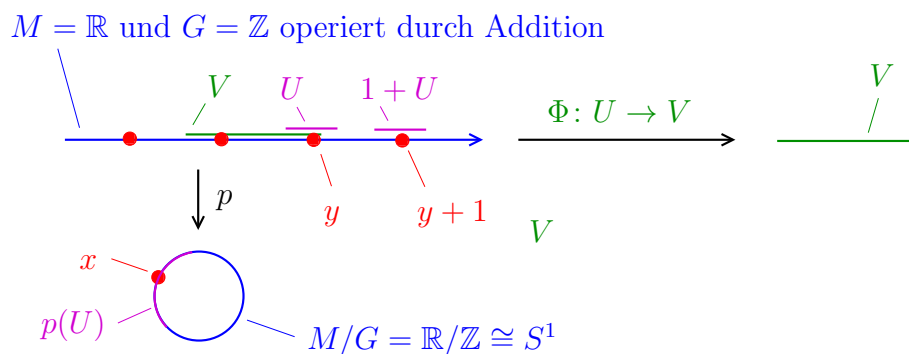


ABBILDUNG 119. Schematische Skizze für den Beweis von Satz 18.11.

Beispiele.

- (A) Wir hatten schon gesehen, dass die Operation von $G = \mathbb{Z}^n$ auf $X = \mathbb{R}^n$ stetig, frei und eigentlich ist. Es folgt also aus Satz 18.11, dass $M = \mathbb{R}^n/\mathbb{Z}^n$ eine n -dimensionale topologische Mannigfaltigkeit ist. Insbesondere erhalten wir einen

²²⁹Es ist eine gute Übung sich zu überlegen, dass dies in der Tat der Fall ist.

weiteren²³⁰ Beweis für die Aussage, dass der 2-dimensionale Torus $\mathbb{R}^2/\mathbb{Z}^2 = S^1 \times S^1$ eine 2-dimensionale topologische Mannigfaltigkeit ist.

- (B) Ganz analog zu (A) sehen wir nun, dass $(\mathbb{R} \times (-1, 1))/\mathbb{Z}$ eine topologische Mannigfaltigkeit ist. Nachdem dieser topologische Raum homöomorph zum Möbiusband ist erhalten wir nun nach Lemma 18.4, einen weiteren Beweis für die Aussage, dass das Möbiusband eine topologische Mannigfaltigkeit ist. Der Beweis in Lemma 18.4 war “ad hoc”, während sich der Beweis mithilfe von Satz 18.11, wie wir gerade auf dieser Seite sehen, auf viele weitere Beispiele verallgemeinert.
- (C) Wir hatten schon gesehen, dass

$$\begin{aligned} \{\pm 1\} \times S^n &\rightarrow S^n \\ (\epsilon, P) &\mapsto \epsilon \cdot P \end{aligned}$$

eine stetige, freie und eigentliche Operation ist. Also ist der n -dimensionalen projektiven Raum $\mathbb{R}P^n = S^n/\{\pm 1\}$ eine n -dimensionale topologische Mannigfaltigkeit.

²³⁰Wir hatten auf Seite 21 schon mal angedeutet, dass der 2-dimensionale Torus eine Untermannigfaltigkeit ist. Allerdings war der Beweis so umständlich, dass wir gleich verzichtet hatten, diesen auszuführen.

19. MANNIGFALTIGKEITEN

19.1. Definition und Beispiele von Mannigfaltigkeiten. Wir erinnern zuerst an die Definition einer n -dimensionalen topologischen Mannigfaltigkeit M .

Definition. Es sei M ein topologischer Raum.

- (1) Eine n -dimensionale Karte für X ist ein Homöomorphismus $\Phi: U \rightarrow V$ zwischen einer offenen Teilmenge $U \subset X$ und einer offenen Teilmenge von $\mathbb{R}^n$.
- (2) Ein n -dimensionaler Atlas für X ist eine Familie von n -dimensionalen Karten $\{\Phi_i: U_i \rightarrow V_i\}_{i \in I}$, so dass $\bigcup_{i \in I} U_i = X$.
- (3) Eine n -dimensionale topologische Mannigfaltigkeit ist ein zweitabzählbarer Hausdorff-Raum, welcher einen n -dimensionalen Atlas besitzt.

Es sei nun $f: M \rightarrow \mathbb{R}$ eine Funktion auf einer n -dimensionalen topologischen Mannigfaltigkeit. Nachdem M ein topologischer Raum ist, macht es Sinn von Stetigkeit von f zu sprechen. Können wir auch von Differenzierbarkeit von f sprechen?

Man könnte ganz naiv folgende "Definition" einführen: wir sagen $f: M \rightarrow \mathbb{R}$ ist im Punkt $x \in M$ differenzierbar, wenn für eine Karte $\Phi: U \rightarrow V$ mit $x \in U$ die Abbildung $f \circ \Phi^{-1}: V \rightarrow \mathbb{R}$ im Punkt $\Phi(x)$ differenzierbar ist. Auf den ersten Blick macht diese Definition Sinn, denn $f \circ \Phi^{-1}$ ist eine Abbildung von einer offenen Teilmenge in $\mathbb{R}^n$ nach $\mathbb{R}$, d.h. es macht Sinn von der Differenzierbarkeit von $f \circ \Phi^{-1}$ zu reden.

Aber was passiert, wenn wir eine andere Karte $\tilde{\Phi}: U \rightarrow W$ um x wählen? Wir verfahren wie im Beweis von Satz 2.7. Genauer gesagt wir schreiben

$$f \circ \tilde{\Phi}^{-1} = (f \circ \Phi^{-1}) \circ (\Phi \circ \tilde{\Phi}^{-1}).$$

Wir wissen, dass $\Phi \circ \tilde{\Phi}^{-1}$ stetig ist, aber a priori wissen wir nicht, dass $\Phi \circ \tilde{\Phi}^{-1}$ auch differenzierbar ist. Im Allgemeinen folgt aus der Differenzierbarkeit von $f \circ \Phi^{-1}$ also nicht notwendigerweise die Differenzierbarkeit von $f \circ \tilde{\Phi}^{-1}$.

Wenn wir hingegen nur Karten $\{\Phi_i\}_{i \in I}$ betrachten würden, so dass jeder Kartenwechsel $\Phi_j \circ \Phi_i^{-1}$ glatt ist, dann hätten wir kein Problem. Diese Beobachtung gibt uns die Idee für folgende Definition.

Definition.

- (1) Ein Atlas $\{\Phi: U_i \rightarrow V_i\}_{i \in I}$ für eine topologische Mannigfaltigkeit M heißt *glatt*, wenn für alle $i, j \in I$ der Kartenwechsel

$$\Phi_j \circ \Phi_i^{-1}: \Phi_i(U_i \cap U_j) \rightarrow \Phi_j(U_i \cap U_j)$$

glatt ist.

- (2) Eine n -dimensionale Mannigfaltigkeit²³¹²³² ist ein Paar $(M, \mathcal{A})$, wobei M eine topologische n -dimensionale Mannigfaltigkeit und $\mathcal{A}$ ein glatter Atlas für M ist.

Beispiel. Jede k -dimensionale Untermannigfaltigkeit M mit $\partial M = \emptyset$ ist eine k -dimensionale Mannigfaltigkeit. In der Tat, denn es sei $\{\Phi_i: U_i \rightarrow V_i\}_{i \in I}$ ein Atlas für M im Sinne der Definition auf Seite 17. Dann ist $\{\Phi_i: U_i \cap M \rightarrow V_i \cap E_k\}_{i \in I}$ ein glatter Atlas im obigen Sinne.²³³

Lemma 19.1. *Das Möbiusband ist eine 2-dimensionale Mannigfaltigkeit.*²³⁴

Bemerkung. Man kann mit ähnlichen Methoden wie im Beweis von Lemma 19.1 nachweisen, dass auch die Kleinsche Flasche eine 2-dimensionale Mannigfaltigkeit ist.

Beweis. Wir hatten schon in Lemma 18.4 gezeigt, dass das Möbiusband eine 2-dimensionale topologische Mannigfaltigkeit ist.

Es genügt nun zu zeigen, dass der Atlas aus dem Beweis von Lemma 18.4 glatt ist. Wir betrachten also jetzt $X = [0, 1] \times (0, 1) \subset \mathbb{R}^2$ mit der Äquivalenzrelation, welche erzeugt ist durch $(0, y) \sim (0, 1 - y)$ für alle $y \in (0, 1)$.

Wir bezeichnen mit $p: X \rightarrow X/\sim$ die Projektionsabbildung. Im Beweis von Lemma 18.4 hatten wir gesehen, dass²³⁵

$$\begin{array}{ccc} \Phi_1: \overbrace{p\left(\left(\frac{1}{8}, \frac{7}{8}\right) \times (0, 1)\right)}^{=:U_1} & \rightarrow & \overbrace{\left(\frac{1}{8}, \frac{7}{8}\right) \times (0, 1)}^{=:V_1} \\ p(a, b) & \mapsto & (a, b) \end{array}$$

und

$$\begin{array}{ccc} \Phi_2: \overbrace{p\left(\left([0, \frac{3}{8}] \cup \left(\frac{5}{8}, 1\right]\right) \times (0, 1)\right)}^{=:U_2} & \rightarrow & \overbrace{\left(-\frac{3}{8}, \frac{3}{8}\right) \times (0, 1)}^{=:V_2} \\ p(a, b) & \mapsto & \begin{cases} (a, b), & \text{wenn } (a, b) \in [0, \frac{3}{8}] \times (0, 1), \\ (a - 1, 1 - b), & \text{wenn } (a, b) \in (\frac{5}{8}, 1] \times (0, 1). \end{cases} \end{array}$$

ein Atlas für $X/\sim$ ist. Wir wollen nun zeigen, dass dieser Atlas sogar glatt ist.

²³¹Manchmal wird eine Mannigfaltigkeit auch als “differenzierbare Mannigfaltigkeit” oder “glatte Mannigfaltigkeit” bezeichnet, um diese stärker von den topologischen Mannigfaltigkeiten abzugrenzen.

²³²Jetzt, nach gerade Mal 12 Wochen Vorlesung haben wir nun also endlich Mannigfaltigkeiten eingeführt, nach denen eigentlich die ganze Vorlesung benannt ist.

²³³Dies sieht man wie folgt. Es sei $\Phi_i: U_i \rightarrow V_i$ eine der k -dimensionalen Untermannigfaltigkeit M . Dies ist also ein Diffeomorphismus zwischen offenen Teilmengen von $\mathbb{R}^n$. Insbesondere ist auch jeder Kartenwechsel $\Phi_j \circ \Phi_i^{-1}: \Phi_i(U_i \cap U_j) \rightarrow \Phi_j(U_i \cap U_j)$ ein Diffeomorphismus.

Für eine Karte $\Phi_i: U_i \rightarrow V_i$ der k -dimensionalen Untermannigfaltigkeit M bezeichnen wir im Folgenden mit $\tilde{\Phi}_i: \tilde{U}_i := U_i \cap M \rightarrow \tilde{V}_i := V_i \cap E_k$ die Einschränkung von Φ_i auf $U_i \cap M$. Wir müssen zeigen, dass die entsprechenden Kartenwechsel alle glatt sind. Der Kartenwechsel $\tilde{\Phi}_j \circ \tilde{\Phi}_i^{-1}: \tilde{\Phi}_i(\tilde{U}_i \cap \tilde{U}_j) \rightarrow \tilde{\Phi}_j(\tilde{U}_i \cap \tilde{U}_j)$ ist die Einschränkung von dem Diffeomorphismus $\Phi_j \circ \Phi_i^{-1}: \Phi_i(U_i \cap U_j) \rightarrow \Phi_j(U_i \cap U_j)$ auf die Teilmenge $\Phi_i(U_i \cap U_j) \cap E_k = \tilde{\Phi}_i(\tilde{U}_i \cap \tilde{U}_j) \subset E_k = \mathbb{R}^k$. Dies ist also wiederum ein Diffeomorphismus.

²³⁴Etwas genauer müsste man sagen, “das Möbiusband besitzt einen glatten Atlas”.

²³⁵Wir haben hierbei U_1 und U_2 etwas abgeändert, damit es leichter ist $U_1 \cap U_2$ zu skizzieren, mathematisch macht das keinen Unterschied.

Wir zeigen zuerst, dass der Kartenwechsel

$$\Phi_2 \circ \Phi_1^{-1}: \Phi_1(U_1 \cap U_2) \rightarrow \Phi_2(U_1 \cap U_2)$$

glatt ist. In unserem Fall ist dies gerade die Abbildung

$$\begin{aligned} \left(\frac{1}{8}, \frac{3}{8}\right) \times (0, 1) \cup \left(\frac{5}{8}, \frac{7}{8}\right) \times (0, 1) &\rightarrow \left(-\frac{3}{8}, -\frac{1}{8}\right) \times (0, 1) \cup \left(\frac{1}{8}, \frac{3}{8}\right) \times (0, 1) \\ (a, b) &\mapsto \begin{cases} (a, b), & \text{wenn } (a, b) \in \left(\frac{1}{8}, \frac{3}{8}\right) \times (0, 1), \\ (a - 1, 1 - b), & \text{wenn } (a, b) \in \left(\frac{5}{8}, \frac{7}{8}\right) \times (0, 1). \end{cases} \end{aligned}$$

Aber diese Abbildung ist in der Tat glatt, denn die Abbildung ist offensichtlich auf den beiden Komponenten $\left(\frac{1}{8}, \frac{3}{8}\right) \times (0, 1)$ und $\left(\frac{5}{8}, \frac{7}{8}\right) \times (0, 1)$ des Definitionsbereichs glatt.

Genau das gleiche Argument zeigt, dass auch der Kartenwechsel

$$\Phi_1 \circ \Phi_2^{-1}: \Phi_2(U_1 \cap U_2) \rightarrow \Phi_1(U_1 \cap U_2)$$

glatt ist. Wir haben also gezeigt, dass $\{\Phi_i: U_i \rightarrow V_i\}_{i=1,2}$ ein glatter Atlas für das Möbiusband ist. $\square$

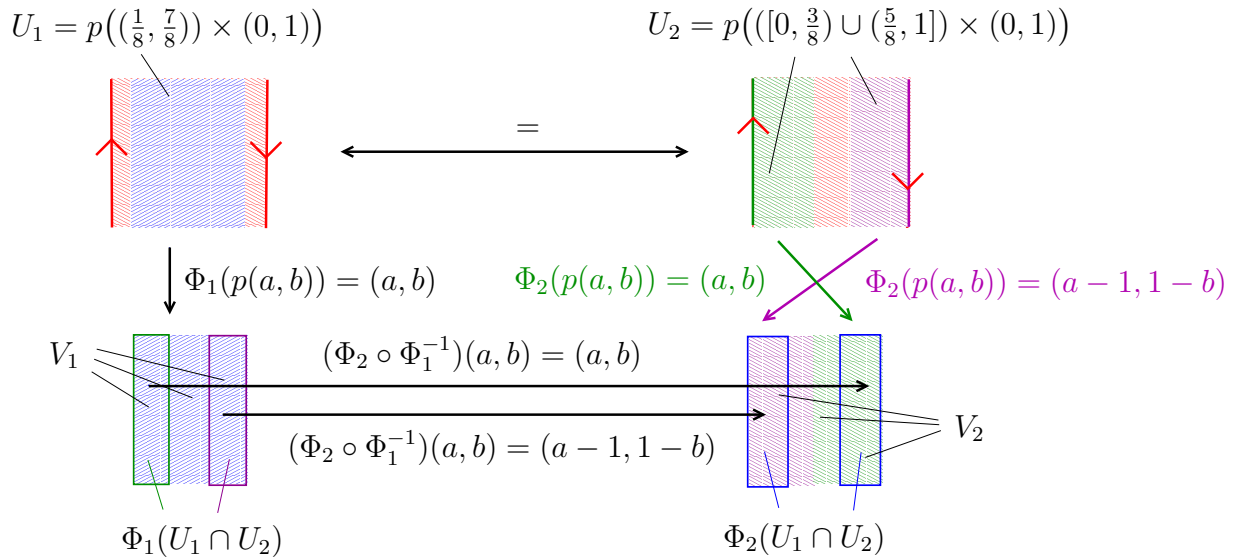


ABBILDUNG 120. Schematische Skizze für den Beweis von Lemma 19.1.

Definition. Es sei $(M, \mathcal{A})$ eine m -dimensionale Mannigfaltigkeit und es sei $(N, \mathcal{B})$ eine n -dimensionale Mannigfaltigkeit. Es sei $f: M \rightarrow N$ eine stetige Abbildung.

- (1) Wir sagen f ist glatt, wenn für alle Karten $\Phi: U \rightarrow V$ aus $\mathcal{A}$ und $\Psi: W \rightarrow X$ aus $\mathcal{B}$ die Abbildung

$$\begin{array}{ccccccc} \underbrace{\Phi(f^{-1}(W) \cap U)}_{\substack{\uparrow \\ \text{offene Teilmenge von } \mathbb{R}^m}} & \xrightarrow{\Phi^{-1}} & f^{-1}(W) \cap U & \xrightarrow{f} & W & \xrightarrow{\Psi} & X \\ & & & & \uparrow & & \text{Teilmenge von } \mathbb{R}^n \end{array}$$

glatt²³⁶ ist.

- (2) Wir sagen $f: M \rightarrow N$ ist ein *Diffeomorphismus*, wenn f ein Homöomorphismus ist und wenn zudem sowohl f als auch f^{-1} glatte Abbildungen sind.

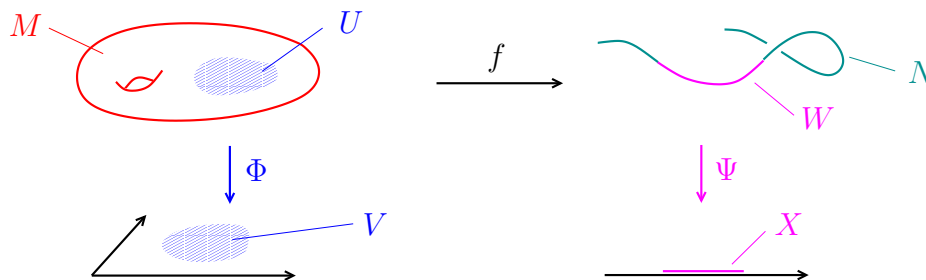


ABBILDUNG 121. Schematische Skizze für die Definition von glatten Abbildungen zwischen Mannigfaltigkeiten.

Beispiel.

- (1) Es sei $(M, \mathcal{A})$ eine n -dimensionale Mannigfaltigkeit und es sei $f: M \rightarrow \mathbb{R}^n$ eine stetige Abbildung. Wir betrachten $\mathbb{R}^n$ als n -dimensionale Mannigfaltigkeit mit dem Atlas, welcher nur aus der Identitätsabbildung besteht. Die Abbildung $f: M \rightarrow \mathbb{R}^n$ ist dann per Definition genau dann glatt, wenn für alle Karten $\Phi: U \rightarrow V$ im Atlas $\mathcal{A}$ die Abbildung $f \circ \Phi^{-1}: V \rightarrow \mathbb{R}^n$ glatt ist.
- (2) Es sei $X = [0, 1] \times (0, 1)$ und es sei $\sim$ wiederum die Äquivalenzrelation auf X , welche durch $(0, y) \sim (1, 1 - y)$ erzeugt wird. Mit anderen Worten, $X/\sim$ ist das Möbiusband. In Übungsblatt 12 werden wir sehen, dass die Funktion

$$\begin{aligned} X/\sim &\rightarrow \mathbb{R} \\ [(x, y)] &\mapsto (1 - y)y \cdot \cos(2\pi x) \end{aligned}$$

glatt ist.

Definition. Es sei M eine Mannigfaltigkeit²³⁷ und G eine Gruppe. Eine *glatte Operation* von G auf M ist eine Operation

$$\begin{aligned} G \times M &\rightarrow M \\ (g, x) &\mapsto g \cdot x, \end{aligned}$$

so dass für jedes $g \in G$ die Abbildung

$$\begin{aligned} M &\rightarrow M \\ x &\mapsto g \cdot x \end{aligned}$$

eine glatte Abbildung ist.

²³⁶Ganz analog könnte man nun auch “differenzierbar” und “glatt” einführen, aber da wir dafür keine Verwendung haben, unterlassen wir es diese Begriffe einzuführen.

²³⁷Wir unterschlagen hierbei den glatten Atlas in der Notation.

Beispiel. Die Gruppe $G = \mathbb{Z}^n$ operiert durch Addition glatt auf der n -dimensionalen Mannigfaltigkeit $\mathbb{R}^n$. Ganz analog sind auch alle weiteren Beispiele, welche wir auf Seite 222 eingeführt hatten, Beispiele von glatten Operationen auf Mannigfaltigkeiten.

Satz 19.2. *Es sei $(M, \mathcal{A})$ eine Mannigfaltigkeit und es sei G eine Gruppe, welche auf M frei, eigentlich und glatt operiert.*

- (1) *Es gibt einen glatten Atlas auf M/G , so dass die Projektionsabbildung $p: M \rightarrow M/G$ eine glatte Abbildung ist.*
- (2) *Die Projektionsabbildung $p: M \rightarrow M/G$ ist ein lokaler Diffeomorphismus.*

Beweis.

- (1) Es sei $\{\Phi_i: U_i \rightarrow V_i\}_{i \in I}$ der Atlas für M/G , welchen wir wie im Beweis von Satz 18.11, ausgehend von den Karten in $\mathcal{A}$, konstruieren. Man sieht nun leicht, dass für alle $i, j \in I$ der entsprechende Kartenwechsel $\Phi_j \circ \Phi_i^{-1}: \Phi_i(U_i \cap U_j) \rightarrow \Phi_j(U_i \cap U_j)$ gegeben ist durch Verknüpfungen von folgenden Abbildungen:
 - (a) Karten und deren Umkehrabbildungen,
 - (b) die Operation von einem $g \in G$ auf offenen Teilmengen von M .
 Alle diese Abbildungen sind glatte Abbildungen zwischen Mannigfaltigkeiten.²³⁸ Die Verknüpfung von glatten Abbildungen zwischen Mannigfaltigkeiten ist aber wiederum eine glatte Abbildung.
- (2) Es sei also $P \in M$. Nachdem G stetig, frei und eigentlich auf M operiert gibt es nach Lemma 18.12 eine offene Umgebung U , so dass $gU \cap U \neq \emptyset$ für alle $g \neq e$. Es folgt aus Lemma 18.13, dass die Einschränkung $p: U \rightarrow p(U)$ ein Homöomorphismus ist. Es folgt leicht aus den Definitionen, dass $p: U \rightarrow p(U)$ sogar ein Diffeomorphismus ist. $\square$

Beispiele. Wir kehren ein letztes Mal zu den Beispielen von Seite 222 zurück.

- (A) Die Gruppe $G = \mathbb{Z}^n$ operiert frei, eigentlich und glatt auf der n -dimensionalen Mannigfaltigkeit $M = \mathbb{R}^n$. Also besagt Satz 19.2, dass wir den n -dimensionalen Torus $\mathbb{R}^n/\mathbb{Z}^n$ als n -dimensionale Mannigfaltigkeit auffassen können.
- (B) Die Gruppe $G = \mathbb{Z}$ operiert frei, eigentlich und glatt auf der 2-dimensionalen Mannigfaltigkeit $M = \mathbb{R} \times (-1, 1)$, also ist $M/G = (\mathbb{R} \times (-1, 1))/\mathbb{Z}$ eine 2-dimensionale Mannigfaltigkeit. Dies gibt also nach Lemma 19.1 einen weiteren Beweis dafür, dass das Möbiusband eine 2-dimensionale Mannigfaltigkeit ist.
- (D) Die Gruppe $G = \{\pm 1\}$ operiert frei, eigentlich und glatt auf der n -dimensionalen Mannigfaltigkeit $M = S^n$, also ist der projektive Raum $M/G = S^n/\{\pm 1\} = \mathbb{R}P^n$ eine n -dimensionale Mannigfaltigkeit.

19.2. Untermannigfaltigkeiten und Produkte von Mannigfaltigkeiten.

Definition. Es sei $(M, \mathcal{A})$ eine n -dimensionale Mannigfaltigkeit. Wir sagen $N \subset M$ ist eine k -dimensionale Untermannigfaltigkeit von M , wenn es zu jedem $P \in N$ eine Karte

²³⁸Die Mannigfaltigkeiten sind hierbei entweder offene Teilmengen von M oder offene Teilmengen von $\mathbb{R}^n$.

$\Phi: U \rightarrow V$ in $\mathcal{A}$ um P gibt, so dass

$$\Phi(U \cap N) = V \cap E_k.$$

Bemerkung. Es sei N eine k -dimensionale Untermannigfaltigkeit von einer n -dimensionalen Mannigfaltigkeit $(M, \mathcal{A})$. Die Einschränkungen der Karten in $\mathcal{A}$ auf N bilden, wie wir schon in Fußnote 233 gesehen hatten, einen glatten Atlas für N . Wir können deshalb N wiederum als Mannigfaltigkeit auffassen.

Beispiele.

- (1) Es sei $M = \mathbb{R}^n$ und es sei $\mathcal{A}$ der glatte Atlas, welcher gegeben ist durch alle Diffeomorphismen $\Phi: U \rightarrow V$ zwischen offenen Teilmengen von $\mathbb{R}^n$. Dann ist eine Teilmenge von $\mathbb{R}^n$ eine Untermannigfaltigkeit der Mannigfaltigkeit $(\mathbb{R}^n, \mathcal{A})$ genau dann, wenn sie eine Untermannigfaltigkeit im Sinne der Definition von Kapitel 2.1 ist.
- (2) In Übungsblatt 13 werden wir sehen, dass jeder Großkreis auf S^2 eine 1-dimensionale Untermannigfaltigkeit von S^2 ist.

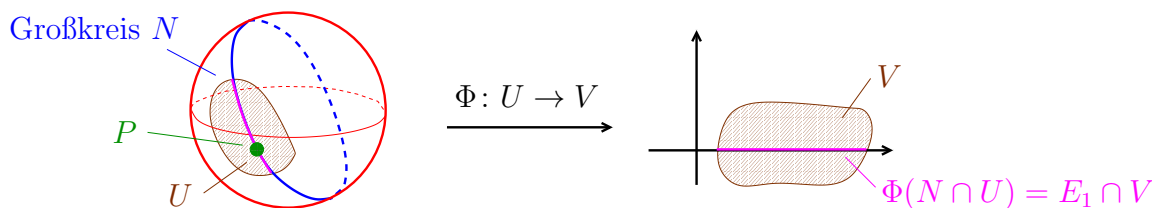


ABBILDUNG 122. Jeder Großkreis ist eine Untermannigfaltigkeit von S^2 .

Lemma 19.3. *Es sei $(M, \mathcal{A})$ eine m -dimensionale Mannigfaltigkeit und es sei $(N, \mathcal{B})$ eine n -dimensionale Mannigfaltigkeit. Dann gelten folgende Aussagen:*

- (1) *Es sei $\Phi: U \rightarrow V$ eine Karte aus $\mathcal{A}$ und es sei $\Psi: W \rightarrow X$ eine Karte aus $\mathcal{B}$. Dann ist*

$$\begin{aligned} \Phi \times \Psi: U \times W &\rightarrow V \times X \\ (p, q) &\mapsto (\Phi(p), \Psi(q)) \end{aligned}$$

eine $(m+n)$ -dimensionale Karte für $M \times N$ und alle solche Karten $\Phi \times \Psi$ bilden einen glatten Atlas für $M \times N$.

- (2) *Das Produkt $M \times N$ ist eine $(m+n)$ -dimensionale Mannigfaltigkeit.*

Beispiel. Es folgt beispielsweise aus Lemma 19.3, dass das Produkt $S^k \times S^l$ von zwei Sphären S^k und S^l eine $(k+l)$ -dimensionale Mannigfaltigkeit ist.

Beweis. Die erste Aussage kann man problemlos überprüfen. Die zweite Aussage folgt aus der ersten Aussage zusammen mit Lemmas 16.8 und 18.2. $\square$

Wir haben jetzt viele Beispiele von Mannigfaltigkeiten kennengelernt. Insbesondere können wir den Torus auf verschiedene Weisen als Mannigfaltigkeit auffassen:

- (1) Auf Seite 21 hatten wir den Torus als “explizite” Teilmenge von $\mathbb{R}^3$ hingeschrieben:

$$T^2 := \{((2 + \sin \theta) \cos \varphi, (2 + \sin \theta) \sin \varphi, \cos \theta) \mid \theta, \varphi \in \mathbb{R}\}.$$

Wir hatten zumindest angedeutet, dass man mit etwas Mühe beweisen kann, dass T^2 eine Untermannigfaltigkeit ist, also nach der obigen Diskussion auch eine Mannigfaltigkeit.

- (2) Wir können den Torus als Quotient $\mathbb{R}^2/\mathbb{Z}^2$ auffassen. Dies ist nach Satz 19.2 ebenfalls eine Mannigfaltigkeit.
- (3) Wir können den Torus als Produkt $S^1 \times S^1$ von zwei Mannigfaltigkeiten auffassen. Nach Lemma 19.3 ist dies eine Mannigfaltigkeit.

Mit einer elementaren, aber etwas langwierigen, Rechnung kann man nun zeigen, dass die drei Mannigfaltigkeiten T , $\mathbb{R}^2/\mathbb{Z}^2$ und $S^1 \times S^1$ diffeomorph sind.

Wenn wir Mannigfaltigkeiten studieren, dann ist es in den allermeisten Fällen irrelevant mit welchem glatten Atlas wir arbeiten, wir werden diesen deswegen ab sofort fast immer in der Notation weglassen.²³⁹

19.3. Die Fläche von Geschlecht 2 als Mannigfaltigkeit.

Satz 19.4. *Die Fläche von Geschlecht g ist eine 2-dimensionale Mannigfaltigkeit.*

Der Beweis von Satz 19.4 ist eine etwas aufwändigere Variante des Beweises von Lemma 18.4 und Lemma 19.1.

Beweis. Wir beweisen den Satz für den Fall $g = 2$. Der allgemeine Fall wird ganz analog bewiesen. Wir betrachten also wieder das Oktagon E_8 mit den Eckpunkten $Q_k = e^{2\pi i k/16}$, $k = 1, 3, \dots, 15$ und der Äquivalenzrelation $\sim$, welche wir in Kapitel 17.4 eingeführt hatten. Das Oktagon mit der Äquivalenzrelation sind in Abbildung 123 skizziert.

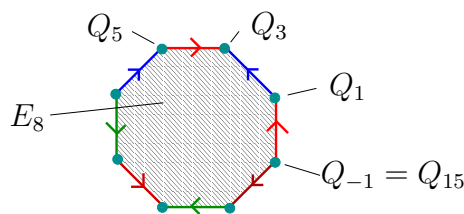


ABBILDUNG 123.

Man kann leicht nachweisen, dass $E_8/\sim$ zweitzählbar und Hausdorff ist. Wir überlassen die Ausführung des Arguments als freiwillige Übungsaufgabe.

²³⁹Allerdings ist es nicht immer so, dass wir uns keinen Gedanken darüber machen müssen, mit welchem glatten Atlas wir arbeiten. Selbst auf höherdimensionalen Sphären ist Vorsicht geboten:

https://en.wikipedia.org/wiki/Exotic_sphere

Wir zeigen als nächstes, dass $E_8/\sim$ eine topologische 2-dimensionale Mannigfaltigkeit ist. Genauer gesagt, wir bestimmen explizit einen Atlas für $E_8/\sim$. Wir werden danach zeigen, dass es sich bei diesem Atlas um einen glatten Atlas handelt.

Es sei zuerst P ein Punkt im Inneren $\mathring{E}_8$ von E_8 . Die Abbildung

$$\begin{aligned} \Phi: U := p(\mathring{E}_8) &\rightarrow V = \mathring{E}_8 \subset \mathbb{R}^2 \\ p(X) &\mapsto X \end{aligned}$$

ist dann eine Karte um $p(P)$. Wir bezeichnen diese als Karte vom 1. Typ.

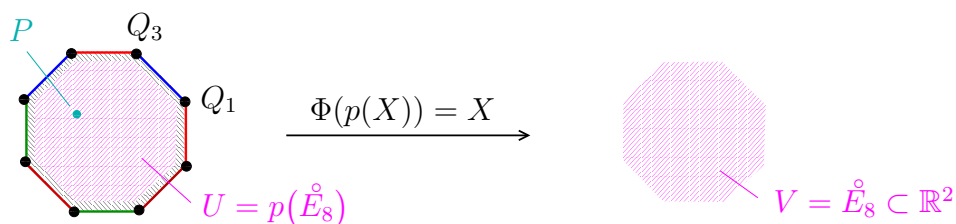


ABBILDUNG 124. Karte vom 1. Typ für einen Punkt im Inneren von E_8 .

Es sei nun P ein Punkt, welcher im Inneren einer Kante von E_8 liegt. Wir wollen eine Karte um $p(P)$ finden. O.B.d.A. sei $P \in \overline{Q_{-1}Q_1}$. In dem Beweis verwenden wir folgende Notation: wir bezeichnen wir mit $s_{\frac{\pi}{4}}: \mathbb{C} \rightarrow \mathbb{C}$ die Spiegelung an der Gerade $\{re^{i\frac{\pi}{4}} \mid r \in \mathbb{R}\}$ und wir bezeichnen mit $P' = s_{\frac{\pi}{4}}(P)$ den Punkt, welcher mit P identifiziert wird. Wir wählen ein $\epsilon > 0$, so dass die offene Scheibe $B_\epsilon(P)$ keinen Eckpunkt des Oktagons berührt. Wir schreiben $U := B_\epsilon(P) \cap E_8$ und $U' := B_\epsilon(P') \cap E_8$. Dann ist $p(U \cup U')$ eine offene Umgebung um P . Es sei $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ die Spiegelung an der Geraden, welche durch die Kante $\overline{Q_{-1}Q_1}$ definiert wird. Es folgt nun leicht aus den Definitionen, dass die Abbildung

$$\begin{aligned} p(U \cup U') &\rightarrow B_\epsilon(P) \subset \mathbb{R}^2 \\ p(X) &\mapsto \begin{cases} X, & \text{wenn } X \in U \\ f(s_{\frac{\pi}{4}}(X)), & \text{wenn } X \in U' \end{cases} \end{aligned}$$

wohl-definiert ist, und dass diese Abbildung ein Homöomorphismus ist. Insbesondere ist dies also eine Karte für E_8 um $p(P)$. Wir bezeichnen diese als Karte vom 2. Typ.

Wir wollen nun noch eine Karte um $p(Q_1) = p(Q_3) = \dots = p(Q_{15})$ finden. Wir wählen ein $\eta > 0$, so dass sich die η -Scheiben um $Q_1, \dots, Q_{15}$ nicht schneiden. Für $k = 1, 3, \dots, 15$ betrachten wir jetzt $U_k := B_\eta(Q_k) \cap E_8$. Die Mengen $U_1, U_3, \dots, U_{15}$ sind disjunkt und man kann sich leicht davon überzeugen, dass

$$p(U_1) \cup p(U_3) \cup \dots \cup p(U_{15})$$

eine offene Umgebung von $p(Q_1) = p(Q_3) = \dots = p(Q_{15})$ ist.

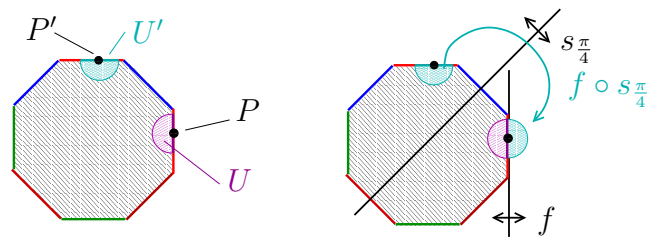


ABBILDUNG 125. Karte vom 2. Typ für einen Punkt im Inneren einer Kante von E_8 .

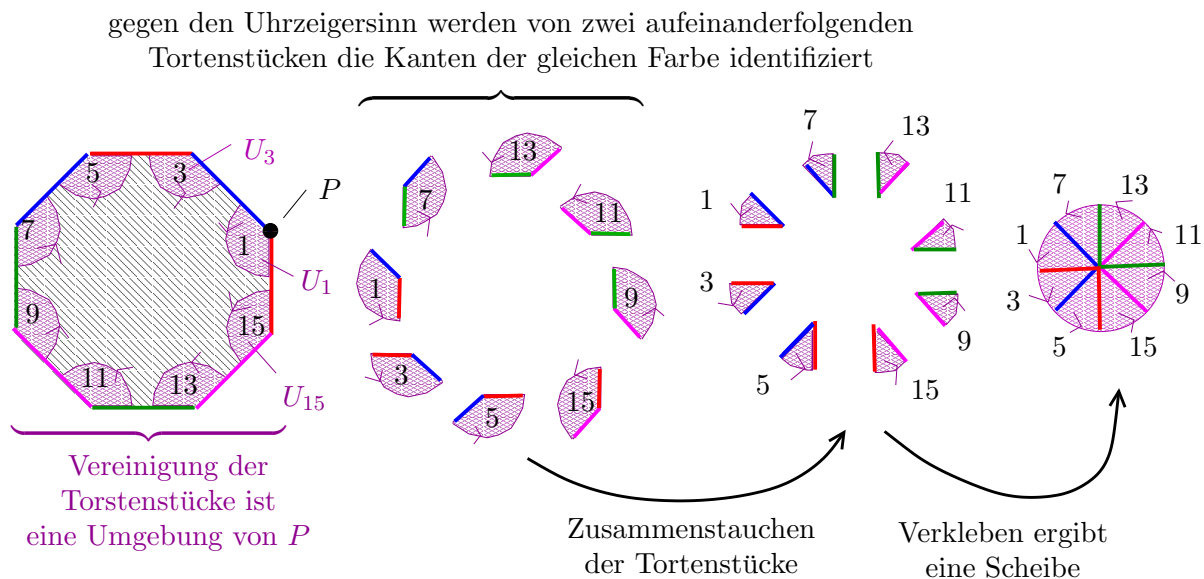


ABBILDUNG 126. Karte vom 3. Typ für einen Eckpunkt von E_8 .

Bevor wir die Karte mathematisch sauber aufschreiben, wollen wir uns die Lage veranschaulichen. In Abbildung 126 sehen wir links die Mengen $U_1, \dots, U_{15}$. Wenn wir diese wie im zweiten Bild anordnen, dann sehen wir die “Tortenstücke” mit den Öffnungswinkeln $\frac{3\pi}{4}$ mit den jeweiligen Identifikationen am Rand. Die Abbildung vom zweiten Bild zum dritten Bild ist dadurch gegeben, dass wir jedes “Tortenstück” auf ein Drittel des Öffnungswinkels zusammenstauchen. Diese acht Tortenstücke, welche jeweils den Öffnungswinkel $\frac{\pi}{4}$ besitzen, fügen sich dann zusammen zu einer Scheibe in $\mathbb{R}^2$.

Wir führen jetzt die gerade beschriebene Idee sauber aus. Wir betrachten die Abbildung²⁴⁰

$$\begin{aligned} \Phi: p(U_1) \cup \dots \cup p(U_{15}) &\rightarrow U_\eta(0) \\ \underbrace{p(Q_k + re^{i\alpha})}_{\in p(U_k)} &\mapsto re^{i(\frac{1}{3}(\alpha - \beta(k)) + \gamma(k))}, \text{ für } k = 1, 3, \dots, 15, \end{aligned}$$

wobei $\beta(k)$ durch folgende Tabelle gegeben ist:

k	1	3	5	7	9	11	13	15
$\beta(k)$	$\frac{3}{4}\pi$	π	$\frac{5}{4}\pi$	$\frac{3}{2}\pi$	$\frac{7}{4}\pi$	0	$\frac{1}{4}\pi$	$\frac{1}{2}\pi$
$\gamma(k)$	$\frac{3}{4}\pi$	π	$\frac{5}{4}\pi$	$\frac{1}{2}\pi$	$\frac{7}{4}\pi$	0	$\frac{1}{4}\pi$	$\frac{3}{2}\pi$

Man kann nun leicht überprüfen, dass die Abbildung Φ wohl-definiert ist und das dies in der Tat ein Homöomorphismus ist.

Wir bezeichnen diese Abbildung als Karte vom 3. Typ.

Wir haben bisher nur bewiesen, dass die Fläche von Geschlecht 2 eine topologische 2-dimensionale Mannigfaltigkeit ist. Wir müssen nun noch zeigen, dass die obigen Karten einen glatten Atlas bilden. Dies ist in der Tat der Fall. Der Kartenwechsel zwischen den Karten von den ersten beiden Typen vergleichen ist entweder die Identität oder eine Verknüpfung von zwei Spiegelungen. In beiden Fällen ist der Kartenwechsel glatt. Wenn wir die Karte um den “Eckpunkt” mit den anderen Karten vergleichen, dann sieht man, dass die Kartenwechsel im Überlappungsgebiet Verknüpfungen von

- (1) Spiegelungen,
- (2) Translationen,
- (3) und der Abbildung

$$\begin{aligned} \{re^{i\alpha} \mid r \in (0, \eta), \alpha \in (0, \frac{\pi}{4})\} &\rightarrow \{re^{i\alpha} \mid r \in (0, \eta), \alpha \in (0, \frac{3\pi}{4})\} \\ re^{i\varphi} &\mapsto re^{3i\varphi}, \end{aligned}$$

- (4) und der Umkehrfunktion aus (3)

sind. Diese Abbildungen sind jedoch glatt, also ist der Kartenwechsel wie gewünscht glatt. $\square$

²⁴⁰Für jedes Tortenstück besitzt die Randkurve eine Orientierung. Es macht daher Sinn von der Ausgangskante zu reden. Die Abbildung verfährt nun also wie folgt mit Tortenstück k :

- (1) das Tortenstück wird in den Ursprung verschoben,
- (2) es wird gegen den Uhrzeigersinn um $\beta(k)$ gedreht, so dass die Ausgangskante auf der x -Achse liegt,
- (3) dann wird das Tortenstück auf ein Drittel des Öffnungswinkels gestaucht,
- (4) dann wird das Tortenstück um den Winkel $\beta(k)$ im Uhrzeigersinn zum Winkel γ gedreht.

20. RIEMANNSCHE MANNIGFALTIGKEITEN

Wir wollen nun natürlich die Begriffe “Integral einer Funktion”, “Länge einer Kurve”, “Geodäten” und “ k -Formen” von Untermannigfaltigkeiten auf Mannigfaltigkeiten verallgemeinern.

Es sei M eine k -dimensionale Untermannigfaltigkeit von $\mathbb{R}^n$ mit einer Karte $\Phi: U \rightarrow V$ mit Umkehrabbildung $\Psi: V \rightarrow U$.

- (1) In der Definition vom “Integral einer Funktion” hatten wir insbesondere die Skalarprodukte

$$\langle D\Psi(e_i), D\Psi(e_j) \rangle = \langle \text{Ableitung von } \Psi(Q + te_i), \text{Ableitung von } \Psi(Q + te_j) \rangle$$

verwendet.

- (2) In der Definition der Länge einer Kurve γ hatten wir das Skalarprodukt $\langle \gamma'(t), \gamma'(t) \rangle$ verwendet.
- (3) Eine m -Form war definiert als eine glatte alternierende m -Form auf den Tangentialräumen $T_P M$. Insbesondere hatten wir in der Definition des Integrals einer k -Form ω auf M diese ausgewertet auf den Vektoren

$$D\Psi(e_i) = \text{Ableitung von } \Psi(Q + te_i)$$

für $i = 1, \dots, k$.

Wir wollen nun diese Begriffe auf Mannigfaltigkeiten verallgemeinern. Das Problem hierbei ist, dass Mannigfaltigkeiten nicht notwendigerweise in einem $\mathbb{R}^n$ liegen, dass heißt, es macht keinen Sinn von Ableitungen von Kurven zu reden. Zudem hatten wir das Skalarprodukt auf dem “umgebenden” $\mathbb{R}^n$ verwendet. Dieses steht uns für Mannigfaltigkeiten nicht zur Verfügung.

Für beide Probleme hatten wir schon mehr oder minder gesehen, was die Lösung ist.

- (1) In Kapitel 7.1 hatten wir schon gesehen, dass wir den Tangentialraum einer Untermannigfaltigkeit von $\mathbb{R}^n$ als die Menge von Äquivalenzklassen von Kurven definieren können. Diese Sichtweise werden wir ohne größere Probleme auf Mannigfaltigkeiten übertragen.
- (2) In Kapitel 13 hatten wir schon gesehen, dass wir anstatt mit dem Skalarprodukt in $\mathbb{R}^n$ auch mit einer beliebigen Riemannschen Struktur, d.h. mit einer Wahl von positiv definiten symmetrischen Bilinearformen auf den Tangentialräumen arbeiten können. Diesen Gesichtspunkt können wir problemlos auf Mannigfaltigkeiten übertragen.

20.1. Der Tangentialraum einer Mannigfaltigkeit.

Definition. Es sei M eine Mannigfaltigkeit und P ein Punkt auf M .

- (1) Eine *Kurve in M durch P* ist eine glatte Abbildung $\gamma: I \rightarrow M$ auf einem offenen Intervall I mit $0 \in I$ und mit $\gamma(0) = P$.

- (2) Wir sagen zwei Kurven γ und δ durch P sind *äquivalent*, wenn es eine Karte²⁴¹ $\Phi: U \rightarrow V$ um P gibt, so dass²⁴² $(\Phi \circ \gamma)'(0) = (\Phi \circ \delta)'(0)$.
- (3) Die Menge aller Äquivalenzklassen von Kurven durch P wird mit $T_P M$ bezeichnet. Wir nennen $T_P M$ den *Tangentialraum* am Punkt P .

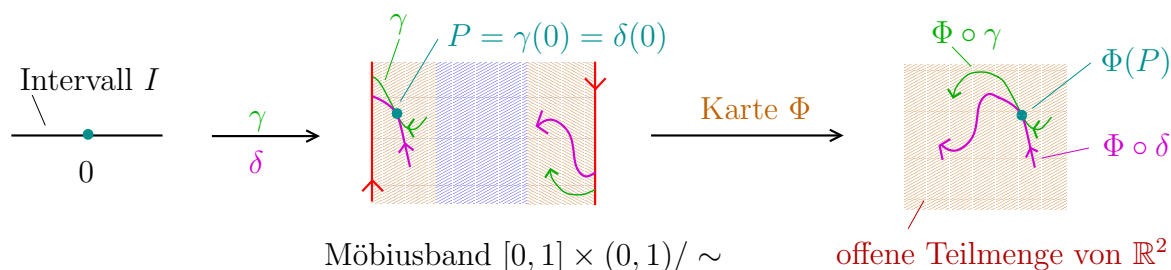


ABBILDUNG 127. Äquivalenz von Kurven auf einer Mannigfaltigkeit.

Lemma 20.1. *Es sei M eine Mannigfaltigkeit. Es seien γ und δ zwei Kurven durch einen Punkt P auf M . Dann gilt*

$$\gamma \text{ und } \delta \text{ sind äquivalent} \iff \text{für jede Karte } \Psi: W \rightarrow X \text{ um } P \text{ gilt} \\ (\Psi \circ \gamma)'(0) = (\Psi \circ \delta)'(0).$$

Beweis. Wir müssen natürlich nur die “ $\Rightarrow$ ”-Richtung beweisen. Es seien also γ und δ zwei äquivalente Kurven durch P auf M . Es gibt also eine Karte $\Phi: U \rightarrow V$ um P , so dass $(\Phi \circ \gamma)'(0) = (\Phi \circ \delta)'(0)$. Es sei nun $\Psi: W \rightarrow X$ eine weitere Karte um P . Dann ist

$$\begin{aligned} (\Psi \circ \gamma)'(0) &= ((\Psi \circ \Phi^{-1}) \circ (\Phi \circ \gamma))'(0) = D_{\Phi(P)}(\Psi \circ \Phi^{-1})((\Phi \circ \gamma)'(0)) \\ &\quad \uparrow \\ &\quad \text{Kettenregel} \\ &= D_{\Phi(P)}(\Psi \circ \Phi^{-1})((\Phi \circ \delta)'(0)) = (\Psi \circ \delta)'(0). \\ &\quad \uparrow \qquad \qquad \qquad \uparrow \\ &\quad \text{nach Voraussetzung} \qquad \qquad \text{gleiche Argument rückwärts} \end{aligned}$$

Hierbei haben wir verwendet, dass $\Psi \circ \Phi^{-1}$ eine glatte Abbildung ist. Dies ist der Fall, denn Ψ und Φ liegen in dem festgewählten glatten Atlas der Mannigfaltigkeit M . $\square$

Bemerkung. Es sei nun M eine k -dimensionale Untermannigfaltigkeit von $\mathbb{R}^n$ und es sei $P \in M$. Wir haben nun zwei verschiedene Definitionen vom Tangentialraum $T_P M$. Zum einen haben wir den klassischen Begriff, welchen wir auf Seite 3 eingeführt hatten, zum anderen haben wir den abstrakteren Begriff, welchen wir gerade erst eingeführt hatten.

²⁴¹Eine “Karte einer Mannigfaltigkeit” ist hierbei immer eine Karte aus dem differenzierbaren Atlas.

²⁴²Die Abbildungen $\Phi \circ \gamma$ und $\Phi \circ \delta$ sind glatte Abbildungen von einem offenen Intervall, welches 0 enthält, zu einer offenen Teilmenge von $\mathbb{R}^n$. Es macht also Sinn, die Ableitungen zu bestimmen.

Man kann nun aber ohne Probleme zeigen, dass ganz analog zur Diskussion in Kapitel 7.1, die Abbildung

$$\begin{aligned} \text{Menge der Äquivalenzklassen} &\rightarrow \text{Ableitungen von Kurven durch } P \\ \text{von Kurven durch } P & \\ [\gamma] &\rightarrow \gamma'(0) \end{aligned}$$

eine Bijektion ist. Wir verwenden im Folgenden diese Bijektion um die beiden Definitionen von $T_P M$ zu identifizieren.

Es sei $f: M \rightarrow N$ eine glatte Abbildung zwischen zwei Mannigfaltigkeiten und es sei $P \in M$. Ganz analog zur Diskussion auf Seite 80 betrachten wir die Abbildung

$$\begin{aligned} f_*: T_P M &\rightarrow T_{f(P)} N \\ [\gamma: I \rightarrow M] &\rightarrow [f \circ \gamma: I \rightarrow N]. \end{aligned}$$

Fast das gleiche Argument wie im Beweis von Lemma 20.1 zeigt, dass diese Abbildung nicht von der Wahl der Kurve γ abhängt. Für eine glatte Abbildung zwischen zwei Untermannigfaltigkeiten von $\mathbb{R}^n$ entspricht diese Definition gerade der Definition aus Kapitel 7.1. Wort-wörtlich der gleiche Beweis wie von Satz 7.2 liefert uns nun auch folgenden Satz.

Satz 20.2.

- (1) Für eine Mannigfaltigkeit M und $P \in M$ gilt

$$\text{id}_* = \text{id}: T_P M \rightarrow T_P M.$$

- (2) Es seien $f: L \rightarrow M$ und $g: M \rightarrow N$ glatte Abbildungen zwischen Mannigfaltigkeiten und es sei $P \in L$. Dann ist

$$(g \circ f)_* = g_* \circ f_*: T_P L \rightarrow T_{g(f(P))} N.$$

Es sei nun M eine n -dimensionale Mannigfaltigkeit und es sei P ein Punkt auf M . Wir hatten gerade den Tangentialraum $T_P M$ eingeführt. Aber es ist von der Definition her nicht klar, in wie weit $T_P M$ eigentlich ein Vektorraum ist. Wir holen dies nun nach und definieren eine Vektorraumstruktur auf $T_P M$. Wir wählen dazu eine Karte²⁴³ $\Phi: U \rightarrow V$ um P . Für $v, w \in T_P M$ definieren wir nun

$$v + w := \Phi_*^{-1}(\underbrace{\Phi_*(v) + \Phi_*(w)}_{\in T_{\Phi(P)} V = \mathbb{R}^n}),$$

und für $v \in T_P M$ und $\lambda \in \mathbb{R}$ definieren wir

$$\lambda \cdot v := \Phi_*^{-1}(\underbrace{\lambda \cdot \Phi_*(v)}_{\in T_{\Phi(P)} V = \mathbb{R}^n}).$$

Wie üblich müssen wir nun zeigen, dass diese Definitionen nicht von der Wahl der Karte abhängt. Das folgende Lemma besagt, dass die Addition und die Skalarmultiplikation auf einem Tangentialraum wohldefiniert sind.

²⁴³Wie immer muss die Karte in unserem festgewählten glatten Atlas liegen.

Lemma 20.3. *Es sei M eine n -dimensionale Mannigfaltigkeit, es sei P ein Punkt auf M und es seien $\Phi: U \rightarrow V$ und $\Psi: W \rightarrow X$ zwei Karten um P . Dann gilt für alle $v, w \in T_P M$, dass*

$$\Phi_*^{-1}(\Phi_*(v) + \Phi_*(w)) = \Psi_*^{-1}(\Psi_*(v) + \Psi_*(w)).$$

Zudem gilt für alle $v \in T_P M$ und $\lambda \in \mathbb{R}$, dass

$$\Phi_*^{-1}(\lambda \cdot \Phi_*(v)) = \Psi_*^{-1}(\lambda \cdot \Psi_*(v)).$$

Beweis. Wir beweisen nur die erste Aussage, die zweite Aussage wird ganz analog bewiesen. Es seien also $v, w \in T_P M$. Dann ist

$$\begin{aligned} & \text{denn } \Psi_*^{-1} \circ \Phi_* = \text{id} \\ & \downarrow \\ \Phi_*^{-1}(\Phi_* v + \Phi_* w) &= (\Psi_*^{-1} \circ \Phi_*)(\Phi_*^{-1}(\Phi_* v + \Phi_* w)) \\ &= \Psi_*^{-1}((\Psi_* \circ \Phi_*^{-1})(\Phi_* v + \Phi_* w)) \\ &= \Psi_*^{-1}((\Psi_* \circ \Phi_*^{-1})(\Phi_* v)) + \Psi_*^{-1}((\Psi_* \circ \Phi_*^{-1})(\Phi_* w)) = \Phi_*^{-1}(\Phi_* v + \Phi_* w). \\ & \uparrow \\ & \Psi_* \circ \Phi_*^{-1} = (\Psi \circ \Phi^{-1})_* \text{ ist eine lineare Abbildung,} \\ & \text{denn } \Psi \circ \Phi^{-1} \text{ ist eine glatte Abbildung zwischen offenen Teilmengen von } \mathbb{R}^n \end{aligned} \quad \square$$

Das folgende Lemma wird ähnlich wie Satz 7.4 bewiesen.

Lemma 20.4. *Es sei $f: M \rightarrow N$ eine glatte Abbildung zwischen zwei Mannigfaltigkeiten und es sei $P \in M$. Dann ist die Abbildung*

$$f_*: T_P M \rightarrow T_{f(P)} N$$

eine lineare Abbildung.

Korollar 20.5. *Es sei M eine n -dimensionale Mannigfaltigkeit und es sei $P \in M$ ein Punkt. Dann ist $T_P M$ ein n -dimensionaler Vektorraum.*

Beweis. Wir wählen eine Karte $\Phi: U \rightarrow V$ um P . Es folgt aus Lemma 20.4, dass die induzierte Abbildung $\Phi_*: T_P M = T_P U \rightarrow T_{\Phi(P)} V$ eine lineare Abbildung ist und es folgt leicht aus Satz 20.2, dass $(\Phi^{-1})_*: T_{\Phi(P)} V \rightarrow T_P M = T_P U$ eine Umkehrabbildung zu Φ_* ist. Also ist $\Phi_*: T_P M = T_P U \rightarrow T_{\Phi(P)} V$ ein Isomorphismus von Vektorräumen.

Nachdem $V \subset U$ eine offene Teilmenge von $\mathbb{R}^n$ ist, folgt aus der Bemerkung auf Seite 242, dass $T_{\Phi(P)} V = T_{\Phi(P)} \mathbb{R}^n = \mathbb{R}^n$. Zusammengefasst ist also $T_P M$ isomorph zum n -dimensionalen Vektorraum $\mathbb{R}^n$. $\square$

Wir betrachten die *Kategorie der Mannigfaltigkeiten* $\mathcal{M}$, welche gegeben ist durch

$$\begin{aligned} \text{Ob}(\mathcal{M}) &:= \text{alle Mannigfaltigkeiten,} \\ \text{Mor}(M, N) &:= \text{alle glatten Abbildungen von } M \text{ nach } N. \end{aligned}$$

Eine *punktierte Mannigfaltigkeit* ist eine Paar (M, P) , wobei M eine Mannigfaltigkeit ist und P ein Punkt auf M . Wir betrachten die *Kategorie der punktierten Mannigfaltigkeiten*

$\mathcal{P}$, welche gegeben ist durch

$$\begin{aligned}\mathrm{Ob}(\mathcal{P}) &:= \{\text{alle punktierten Mannigfaltigkeiten}\}, \\ \mathrm{Mor}((M, P), (N, Q)) &:= \left\{ \begin{array}{l} \text{alle glatten Abbildungen } f \\ \text{von } M \text{ nach } N \text{ mit } f(P) = Q \end{array} \right\}.\end{aligned}$$

Satz 20.2 und Lemma 20.4 besagen, dass

$$(M, P) \mapsto T_P M$$

und

$$(f: (M, P) \rightarrow (N, Q)) \mapsto (f_*: T_P M \rightarrow T_Q N)$$

einen kovarianten Funktor von der Kategorie der punktierten Mannigfaltigkeiten zu der Kategorie der Vektorräume definiert.

20.2. Riemannsche Mannigfaltigkeiten und Integration von Funktionen. Wir wollen nun Funktionen auf Mannigfaltigkeiten integrieren. Nachdem uns das Skalarprodukt nicht mehr zur Verfügung steht müssen wir, wie schon auf Seite 241 angedeutet, diesen Wegfall durch die Wahl einer positiv definiten symmetrischen Bilinearform für jeden Tangentialraum kompensieren. Dies führt uns zum Begriff der Riemannschen Mannigfaltigkeit. Deren Definition ist fast wort-wörtlich die gleiche wie die Definition von Riemannschen Untermannigfaltigkeiten von $\mathbb{R}^n$, welche wir auf Seite 158 gegeben hatten.

Etwas genauer gesagt, eine *Riemannsche Struktur* auf einer Mannigfaltigkeit M ist eine glatte Abbildung g , welche jedem Punkt $P \in M$ eine positiv definite symmetrische Bilinearform g_P auf $T_P M$ zuordnet. Die Definition von “glatt” erfolgt hierbei über Karten, ganz analog zur Definition auf Seite 158. Eine *Riemannsche Mannigfaltigkeit* ist ein Paar (M, g) , wobei M eine Mannigfaltigkeit und g eine Riemannsche Struktur auf M ist.

Es sei nun $f: M \rightarrow \mathbb{R}$ eine Funktion auf einer k -dimensionalen Riemannschen Mannigfaltigkeit (M, g) . Wir nehmen zuerst an, dass es eine Karte $\Phi: U \rightarrow V$ aus dem gegebenen Atlas gibt, so dass $\mathrm{Träger}(f) \subset U$. Wir bezeichnen mit $\Psi: V \rightarrow U$ die Umkehrabbildung. Dann definieren wir, ganz analog zur Definition auf Seite 185

$$\int_{(M, g)} f := \int_V f(\Psi(x)) \cdot \sqrt{\det (g_{\Psi(x)}(\underbrace{\Psi_* e_i, \Psi_* e_j}_{e_i, e_j \text{ liegen in } \mathbb{R}^k = T_x V \text{ und } \Psi_* e_i, \Psi_* e_j \text{ liegen daher in } T_{\Psi(x)} M}))_{i,j=1,\dots,k}}.$$

Das Argument von Satz 3.16 zeigt, dass diese Definition nicht von der Wahl der Karte abhängt.

Wir verallgemeinern nun diese Definition auf die Integration von Funktionen auf beliebigen Riemannschen Mannigfaltigkeiten:

- (1) Ganz analog zur Diskussion in Kapitel 3.5 erweitern wir die Definition auf Mannigfaltigkeiten, welche durch endliche viele Karten aus dem vorgegebenen Atlas abgedeckt werden.

- (2) In Kapitel 3.7 hatten wir das Integral einer Funktion auf einer beliebigen Untermannigfaltigkeit von $\mathbb{R}^n$ eingeführt. Hierbei hatten wir Lemma 3.21 verwendet, welches besagt, dass jede Untermannigfaltigkeit von $\mathbb{R}^n$ einen abzählbaren Atlas besitzt.

An dieser Stelle verwenden wir jetzt die Tatsache, dass per Definition jede Mannigfaltigkeit M zweitabzählbar ist. Der Beweis von Lemma 3.21 zeigt dann, dass M auch einen abzählbaren glatten Atlas besitzt.²⁴⁴ Der Rest von Kapitel 3.7 kann nun problemlos übernommen werden.

20.3. Mannigfaltigkeiten mit Rand. In Kapitel 2.1 hatten wir zuerst Untermannigfaltigkeiten von $\mathbb{R}^n$ “ohne Rand” eingeführt. In Kapitel 4 hatten wir diesen ursprünglichen Begriff von Untermannigfaltigkeiten erweitert, und wir hatten Untermannigfaltigkeiten von $\mathbb{R}^n$ zugelassen, welche einen nichtleeren Rand besitzen. Mit fast wort-wörtlich der gleichen Vorgehensweise kann man nun auch topologische Mannigfaltigkeiten mit Rand sowie Mannigfaltigkeiten mit Rand einführen. Im folgenden kann eine (topologische) Mannigfaltigkeit auch immer einen Rand besitzen.

Es gelten die offensichtlichen Verallgemeinerungen der Aussagen über Untermannigfaltigkeiten mit Rand und den Rand einer Untermannigfaltigkeit. Beispielsweise gilt, ganz analog zu Satz 4.3, dass der Rand ∂M einer k -dimensionalen Mannigfaltigkeit M mit Rand eine $(k - 1)$ -dimensionale Mannigfaltigkeit mit leerem Rand ist.

Beispiel. Wir betrachten jetzt $Y = [0, 1] \times [0, 1] \subset \mathbb{R}^2$, wobei die Äquivalenzrelation erzeugt ist durch

$$(0, y) \sim (1, 1 - y) \quad \text{für alle } y \in [0, 1].$$

Dies ist also fast die gleiche Definition wie die Definition vom Möbiusband auf Seite 211. Aber dieses Mal betrachten wir das abgeschlossene Quadrat $[0, 1] \times [0, 1]$ anstatt dem halboffenen Quadrat $[0, 1] \times (0, 1)$. Der Quotientenraum $Y / \sim$ entsteht wie zuvor aus dem Quadrat Y , indem wir die “linke Kante nach einer Verdrillung mit der rechten Kante verkleben”. Den topologischen Raum $Y / \sim$ nennen wir das *abgeschlossene Möbiusband*.

Fast das gleiche Argument wie in Lemmas 18.4 und 19.1 zeigt, dass das abgeschlossene Möbiusband eine 2-dimensionale Mannigfaltigkeit ist mit Rand

$$(\{0, 1\} \times [0, 1]) / \sim = ((\{0\} \times [0, 1]) \cup (\{1\} \times [0, 1])) / \sim,$$

wobei $(0, 0) \sim (1, 1)$ und $(0, 1) \sim (1, 0)$. Man kann nun leicht zeigen, dass der Rand von $Y / \sim$ diffeomorph zum Kreis S^1 ist. Die Abbildung 128 veranschaulicht das abgeschlossene Möbiusband und seinen Rand.

²⁴⁴Der Vollständigkeit halber führen wir das Argument aus. Es sei also M eine Mannigfaltigkeit und es sei $\mathcal{A}$ ein glatter Atlas für M . Zudem sei $\mathcal{B}$ eine abzählbare Basis der Topologie. Zu jedem $x \in M$ gibt es eine Karte $\Phi_x: U_x \rightarrow V_x$ aus dem Atlas $\mathcal{A}$. Per Definition einer Basis der Topologie, siehe Seite 198, gibt es zu jedem $x \in M$ ein $W_x \in \mathcal{B}$ mit $x \in W_x$ und $W_x \subset U_x$. Dann ist auch $\{\Phi_x: W_x \rightarrow \Phi_x(W_x)\}_{x \in M}$ ein Atlas für M . Nachdem es nur abzählbar viele offene Mengen in $\mathcal{B}$ gibt, genügen nun schon abzählbar viele dieser Karten um M zu überdecken, d.h. M besitzt einen abzählbaren Atlas.

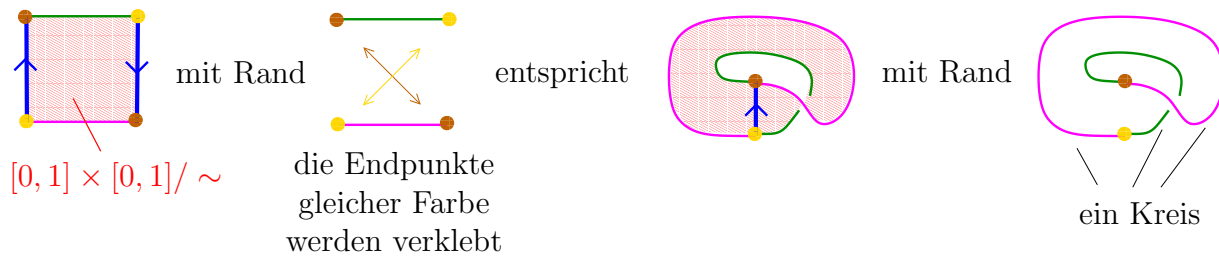


ABBILDUNG 128. Das abgeschlossene Möbiusband.

20.4. Differentialformen auf Mannigfaltigkeiten. In den Kapiteln 9 und 10 hatten wir für eine Untermannigfaltigkeit M von $\mathbb{R}^n$ folgende Begriffe eingeführt:

- (1) stetige und glatte Differentialformen,
- (2) die Menge $\Omega_k(M)$ der glatten k -Formen auf M ,
- (3) Orientierungen auf M ,
- (4) das Integral einer stetigen k -Form auf einer kompakten orientierten k -dimensionalen Untermannigfaltigkeit von $\mathbb{R}^n$,
- (5) das Differential $d\omega$ einer glatten k -Form.

In allen diesen Definitionen hatten wir immer mit den Tangentialräumen gearbeitet, aber wir hatten ansonsten nie verwendet, dass M einen “umgebenden Raum” besitzt.²⁴⁵ Insbesondere hatten wir bei diesen Definitionen und Sätzen nie das Skalarprodukt des $\mathbb{R}^n$ in irgendeiner Weise verwendet.

Nachdem wir jetzt für Mannigfaltigkeiten ebenfalls einen Begriff von Tangentialräumen haben, können wir nun diese Begriffe und alle Aussagen aus Kapitel 9 und 10 fast wortwörtlich zu Mannigfaltigkeiten übertragen.²⁴⁶ Insbesondere erhalten wir folgendes Analogon zu Satz 10.2, Satz 9.12 und Satz 10.3.

Satz 20.6.

- (1) Für jede glatte k -Form ω auf einer Mannigfaltigkeit ist $d\omega$ eine glatte $(k+1)$ -Form und es gilt $d d\omega = 0$.
- (2) Es sei $f: M \rightarrow N$ eine glatte Abbildung zwischen zwei Mannigfaltigkeiten und es sei $\omega \in \Omega_k(N)$. Dann ist $f^*\omega$ eine glatte k -Form auf M und es gilt

$$d(f^*\omega) = f^*(d\omega).$$

²⁴⁵Die Definitionen und Sätze wurden auch immer formuliert in der Form “Es sei M eine Untermannigfaltigkeit”, und nie von der Form “Es sei M eine Untermannigfaltigkeit von $\mathbb{R}^n$ ”, weil eben der $\mathbb{R}^n$ nie eine Rolle gespielt hat.

²⁴⁶Allerdings müssen wir uns bei Karten dann immer auf Karten des Atlanten einschränken. Beispielsweise heißt eine k -Form auf einer Mannigfaltigkeit $(M, \mathcal{A})$ glatt, wenn für alle Karten $\Phi: U \rightarrow V$ im Atlas $\mathcal{A}$ die k -Form $(\Phi^{-1})^*\omega$ auf V glatt ist.

Bemerkung. Der Satz besagt also insbesondere, dass für eine glatte Abbildung $f: M \rightarrow N$ zwischen zwei Mannigfaltigkeiten folgendes Diagramm von Abbildungen kommutiert:

$$\begin{array}{ccc} \Omega_k(N) & \xrightarrow{f^*} & \Omega_k(M) \\ \downarrow d & & \downarrow d \\ \Omega_{k+1}(N) & \xrightarrow{f^*} & \Omega_{k+1}(M). \end{array}$$

Wir können auch problemlos das Integral von stetigen k -Formen auf orientierten kompakten k -dimensionalen Mannigfaltigkeiten einführen. Das Integral wird durch den folgenden Satz vollständig charakterisiert.

Satz 20.7. *Es gibt genau eine Abbildung, genannt das Integral, welche jeder stetigen k -Form ω auf einer orientierten kompakten k -dimensionalen Mannigfaltigkeit M eine reelle Zahl*

$$\int_M \omega$$

zuordnet, welche folgende Eigenschaften besitzt:

- (1) *Das Integral ist linear, d.h. für jede orientierte kompakte k -dimensionale Mannigfaltigkeit M und alle stetigen k -Formen ω, σ auf M sowie alle $\lambda \in \mathbb{R}$ gilt*

$$\int_M (\omega + \sigma) = \int_M \omega + \int_M \sigma \quad \text{und} \quad \int_M \lambda \omega = \lambda \int_M \omega.$$

- (2) *Wenn M eine k -dimensionale kompakte Untermannigfaltigkeit von $\mathbb{R}^k$ ist, und wenn $f: M \rightarrow \mathbb{R}$ eine stetige Funktion ist, dann gilt*

$$\underbrace{\int_M f \cdot dx_1 \wedge \cdots \wedge dx_k}_{\text{Integral einer } k\text{-Form}} = \underbrace{\int_M f dx_1 \dots dx_k}_{\text{Lebesgue-Integral}}.$$

- (3) *Für jeden Orientierungserhaltenden Diffeomorphismus $f: M \rightarrow N$ zwischen zwei kompakten orientierten k -dimensionalen Mannigfaltigkeiten und jede stetige k -Form auf N gilt*

$$\int_M f^* \omega = \int_N \omega.$$

Beweis. Die Tatsache, dass das Integral von k -Formen diese Eigenschaften besitzt folgt leicht aus den Definitionen, sowie aus Lemma 9.22 und Satz 9.21, umformuliert für Mannigfaltigkeiten. Die Tatsache, dass das Integral durch diese Eigenschaften charakterisiert ist, folgt aus der Beobachtung, dass man, ganz analog zur Diskussion auf Seite 129, jede glatte k -Form auf einer kompakten Mannigfaltigkeit als Summe von glatten k -Formen schreiben kann, deren Träger jeweils im Definitionsbereich einer Karte enthalten sind. Für solche k -Formen ist das Integral durch die Eigenschaften (2) und (3) charakterisiert. $\square$

Wir erhalten auch folgende Verallgemeinerung vom Satz 11.1 von Stokes für Untermannigfaltigkeiten von $\mathbb{R}^n$ auf Mannigfaltigkeiten.

Satz 20.8. (Satz von Stokes) *Es sei M eine orientierte k -dimensionale kompakte Mannigfaltigkeit und es sei $\omega \in \Omega_{k-1}(M)$. Dann gilt*

$$\int_M d\omega = \int_{\partial M} \omega.$$

Insbesondere, wenn $\partial M = \emptyset$, dann gilt

$$\int_M d\omega = 0.$$

Im Folgenden führen wir noch eine Möglichkeit ein, Differentialformen auf Quotientenmannigfaltigkeiten zu konstruieren. Wir benötigen für die Formulierung des nächsten Lemmas folgende zwei Definitionen.

Definition.

- (1) Es sei G eine Gruppe, welche glatt auf einer Mannigfaltigkeit M operiert. Für $g \in G$ bezeichnen wir mit $\phi_g: M \rightarrow M$ den Diffeomorphismus, welcher gegeben ist durch $x \mapsto g \cdot x$. Wir sagen eine k -Form ω auf M ist G -invariant, wenn für alle $g \in G$ gilt, dass $\phi_g^*(\omega) = \omega$.
- (2) Eine Differentialform ω auf einer Mannigfaltigkeit M wird *geschlossen* genannt, wenn $d\omega = 0$.

Lemma 20.9. *Es sei M eine Mannigfaltigkeit und G eine Gruppe, welche auf M frei, glatt und eigentlich operiert. Wir bezeichnen mit $p: M \rightarrow M/G$ die Projektionsabbildung. Es sei ω eine k -Form auf M , welche G -invariant ist. Dann gibt es genau eine glatte k -Form $\bar{\omega}$ auf der Mannigfaltigkeit M/G mit*

$$\omega = p^*(\bar{\omega}).$$

Wenn ω geschlossen ist, dann ist auch die induzierte k -Form $\bar{\omega}$ auf der Quotientenmannigfaltigkeit M/G geschlossen.

Beispiel. Wir bezeichnen mit dx und dy die 1-Formen auf $\mathbb{R}^2$, welche an jedem $P \in \mathbb{R}^2$ gegeben sind durch

$$\begin{aligned} dx: T_P \mathbb{R}^2 = \mathbb{R}^2 &\rightarrow \mathbb{R} & \text{und} & & dy: T_P \mathbb{R}^2 = \mathbb{R}^2 &\rightarrow \mathbb{R} \\ (v_x, v_y) &\mapsto v_x & & & (v_x, v_y) &\mapsto v_y. \end{aligned}$$

Diese 1-Formen sind invariant unter der Operation von $\mathbb{Z}^2$ auf $\mathbb{R}^2$. Diese beiden 1-Formen induzieren also 1-Formen $\bar{dx}$ und $\bar{dy}$ auf $\mathbb{R}^2/\mathbb{Z}^2$. Nachdem dx und dy geschlossen sind, sind auch die 1-Formen $\bar{dx}$ und $\bar{dy}$ geschlossen.

Die 2-Form $dx \wedge dy$ auf $\mathbb{R}^2$ ist ebenfalls invariant unter der Operation von $\mathbb{Z}^2$ und induziert damit eine 2-Form $\bar{dx} \wedge \bar{dy}$ auf $\mathbb{R}^2/\mathbb{Z}^2$. Man sieht leicht, dass $\bar{dx} \wedge \bar{dy} = \overline{dx \wedge dy}$. Dann gilt

beispielsweise²⁴⁷

$$\int_{\mathbb{R}^2/\mathbb{Z}^2} \overline{dx} \wedge \overline{dy} = 1.$$

Beweis. Es sei M eine Mannigfaltigkeit und G eine Gruppe, welche auf M frei, glatt und eigentlich operiert. Wir bezeichnen mit $p: M \rightarrow M/G$ die Projektionsabbildung. Es sei ω eine k -Form auf M , welche G -invariant ist.

Es sei nun x ein Punkt in M/G . Wir wählen ein $y \in M$ mit $p(y) = x$. Nach Satz 19.2 gibt es eine offene Umgebung U von y , so dass $p: U \rightarrow p(U)$ ein Diffeomorphismus ist. Wir bezeichnen mit $q: p(U) \rightarrow U$ die Umkehrabbildung und wir setzen $\overline{\omega}_x = q^*\omega_y$. Wir zeigen nun, dass $\omega = p^*(\overline{\omega})$ für alle Punkte in $p^{-1}(y) = G \cdot y = \{gy \mid g \in G\}$. Es sei also $g \in G$ beliebig. Wir bezeichnen mit $\phi_{g^{-1}}: M \rightarrow M$ die Abbildung, welche durch $\phi_{g^{-1}}(x) := g^{-1} \cdot x$ definiert ist. Dann gilt in der Tat

$$\begin{array}{ccccccccccc} & & \text{Definition von } \overline{\omega} & & & \text{da } p \circ q = \phi_{g^{-1}} \text{ auf } gU & & & & & \\ & & \downarrow & & & \downarrow & & & & & \\ p^*(\overline{\omega})_{gy} & = & p^*(\overline{\omega}_{p(gy)}) & = & p^*(\overline{\omega}_x) & = & p^*(q^*(\omega_y)) & = & (p \circ q)^*(\omega_y) & = & \phi_{g^{-1}}^*(\omega_y) & = & \omega_{gy}. \\ & & \uparrow & & & & \uparrow & & & & \uparrow & & \\ & & \text{denn } p(gy) = p(y) = x & & & \text{nach Seite 115 ist dies ein} & & & & & \text{denn } \omega \text{ ist} & & \\ & & & & & \text{kontravarianter Funktor} & & & & & G\text{-invariant} & & \end{array}$$

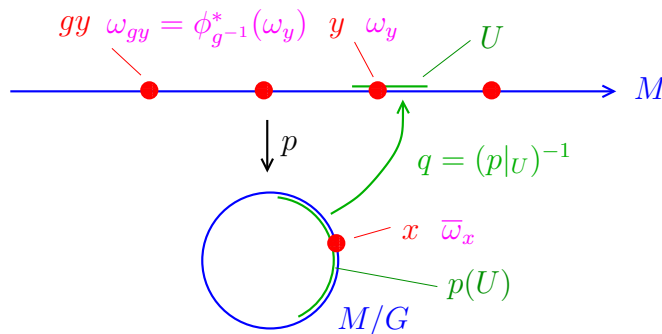


ABBILDUNG 129. Skizze zum Beweis von Lemma 20.9.

Wir nehmen nun noch an, dass ω geschlossen ist. Wir wollen zeigen, dass auch $\overline{\omega}$ geschlossen ist. Dies sieht man wie folgt. Es genügt zu zeigen, dass es für jedes $b \in M/G$ eine offene Umgebung gibt, so dass $\overline{\omega}$ auf dieser offenen Umgebung geschlossen ist. Es

²⁴⁷Dies sieht man wie folgt: wir schreiben $U = (0, 1) \times (0, 1)$, dann ist $p: U \rightarrow p(U)$ ein Diffeomorphismus und es ist

$$\begin{array}{ccccccc} \int_{\mathbb{R}^2/\mathbb{Z}^2} \overline{dx} \wedge \overline{dy} & = & \int_{p(U)} \overline{dx} \wedge \overline{dy} & = & \int_U p^*(\overline{dx} \wedge \overline{dy}) & = & \int_U dx \wedge dy = \text{Flächeninhalt von } U = (0, 1)^2 = 1. \\ & & \uparrow & & \uparrow & & \uparrow \\ & & \text{denn } \mathbb{R}^2/\mathbb{Z}^2 \setminus p(U) & & \text{Satz 20.7(3)} & & \text{Satz 20.7(2)} \\ & & \text{ist eine "Nullmenge"} & & \text{obige Diskussion} & & \end{array}$$

sei also $b \in M/G$. Wir bezeichnen mit $p: M \rightarrow M/G$ die Projektionsabbildung und wir wählen ein $a \in M$ mit $p(a) = b$. Nach Satz 19.2 gibt es eine offene Umgebung U von a , so dass $p: U \rightarrow p(U)$ ein Diffeomorphismus ist. Wir bezeichnen mit $q = p^{-1}: p(U) \rightarrow U$ die Umkehrabbildung. Es folgt nun, dass

$$\begin{array}{ccccccc}
 d\bar{\omega} & = & d(q^*p^*\bar{\omega}) & = & dq^*(\omega) & = & q^*(d\omega) & = & 0. \\
 \uparrow & & \uparrow & & \uparrow & & \uparrow & & \\
 \text{da } q^* \circ p^* = (p \circ q)^* = \text{id} & & \text{da } \omega = p^*(\bar{\omega}) & & \text{Satz 20.6 (2)} & & \text{da } d\omega = 0
 \end{array}$$

□

21. KRÜMMUNG AUF MANNIGFALTIGKEITEN

Der Begriff der Länge eines Weges und der Begriff der Geodäte, welche wir in Kapitel 13 für Riemannsche Untermannigfaltigkeiten von $\mathbb{R}^n$ eingeführt hatten überträgt sich problemlos auf Riemannsche Mannigfaltigkeiten. Etwas genauer gesagt, es sei (M, g) eine Riemannsche Mannigfaltigkeit.

- (1) Für $P \in M$ und $v \in T_P M$ bezeichnen wir $\|v\|_g := \sqrt{g_P(v, v)}$ als die *Länge von v bezüglich der Riemannschen Struktur g* .
- (2) Es sei nun $\gamma: [a, b] \rightarrow M$ ein Weg, d.h. eine abschnittsweise glatte Abbildung. Wir wählen eine Zerlegung $a = s_0 < s_1 < \dots < s_m = b$ gibt, so dass für $i = 0, \dots, m-1$ die Abbildung $\gamma|_{[s_i, s_{i+1}]}$ glatt ist. Wir definieren die *Länge von γ bezüglich g* als

$$\ell_g(\gamma) := \sum_{i=0}^{m-1} \int_{t=s_i}^{t=s_{i+1}} \underbrace{\|\gamma_*(e)\|_g}_{\substack{\uparrow \\ \text{es ist } e = (1) \in \mathbb{R} = T_t \mathbb{R} \\ \text{und daher ist } \gamma_*(e) \in T_{\gamma(t)} M}} dt.$$

- (3) Eine *Geodäte in (M, g)* ist ein Weg $\gamma: [a, b] \rightarrow M$, welcher folgende drei Bedingungen erfüllt:
 - (a) γ ist glatt,
 - (b) $\|\gamma'(t)\|_g = 1$ für alle $t \in [a, b]$ und
 - (c) jeder andere Weg von $\gamma(a)$ nach $\gamma(b)$ ist mindestens so lang wie γ .

Wir können zudem die Krümmungsbegriffe auf Riemannschen Untermannigfaltigkeiten von $\mathbb{R}^n$, welche wir in Kapitel 15.3 für 2-dimensionale Riemannsche Untermannigfaltigkeiten von $\mathbb{R}^n$ eingeführt hatten, auf 2-dimensionale Riemannsche Mannigfaltigkeiten verallgemeinern. Beispielsweise sagen wir nun, dass eine 2-dimensionale Riemannsche Mannigfaltigkeit (M, g) gekrümmt ist, wenn es ein Dreieck gibt, dessen Winkelsumme $\neq \pi$ ist.

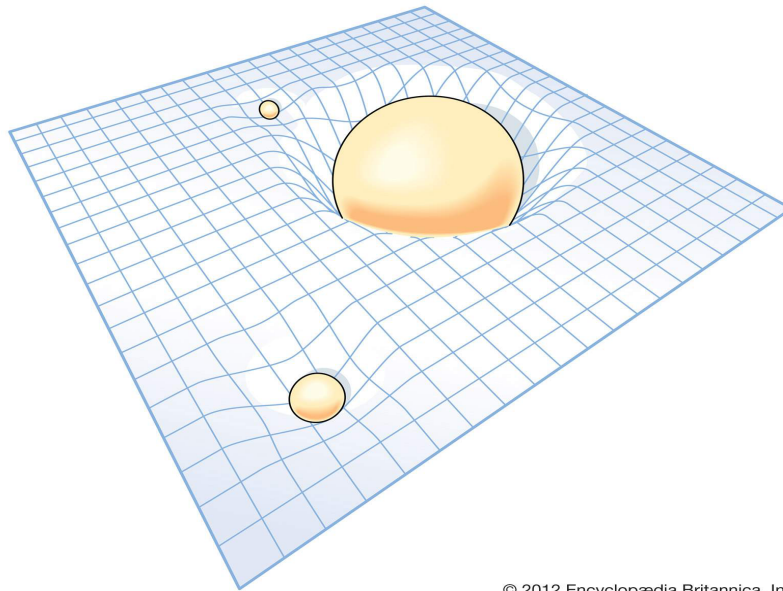
21.1. Krümmung des Raums. Dieses kurze Kapitel behandelt die Krümmung des physikalischen Raums. Nachdem dies keine Physikvorlesung ist, sollte man dieses Kapitel nicht zu ernst nehmen.

Seit Einstein wissen wir, dass “der physikalische Raum” gekrümmt ist, genauer gesagt, dass Masse den Raum krümmt. In der Literatur wird das oft wie in Abbildung 130 illustriert. Das Bild zeigt das Universum als Teilmenge von “einem größeren $\mathbb{R}^n$ ”, in dem es gekrümmt ist. Das ist natürlich nicht der Fall, das Universum ist besser beschrieben als eine 3-dimensionale Riemannsche²⁴⁸ Mannigfaltigkeit. Lichtstrahlen bewegen sich auf

²⁴⁸Es ist wohl jetzt ziemlich klar, warum wir uns unser Universum als 3-dimensionale Mannigfaltigkeit vorstellen können. Allerdings ist es vielleicht weniger offensichtlich, warum wir es uns als Riemannsche Mannigfaltigkeit vorstellen können. Eine Bilinearform an einem Punkt ist äquivalent zu folgenden zwei Begriffen:

- (1) Längenmessung von Vektoren,
- (2) den Begriff des Winkels zwischen zwei Vektoren.

Beide Begriffe sind in jedem Punkt des Universums gegeben.



© 2012 Encyclopædia Britannica, Inc.

ABBILDUNG 130.

Geodäten, und diese sind gekrümmt. Es sei beispielsweise ein Dreieck gegeben, welches durch Lichtstrahlen gebildet ist, und in dessen Mitte sich ein schwerer Körper befindet. Die Summe der Innenwinkel wird dann nicht π betragen. Beispielsweise zeigt die Abbildung 131 zwei Lichtstrahlen, ausgehend vom gleichen Objekt, welche nicht parallel sind, aber sich dennoch treffen.

Es gibt Vermutungen, dass schon Gauß experimentell versucht hat zu überprüfen, ob das Universum gekrümmt ist. Allerdings waren die Meßfehler viel zu groß um eine Aussage treffen zu können. Details zum Experiment von Gauß kann man auf folgender Webseite finden:

<https://thatsmaths.com/2014/07/10/gausss-great-triangle-and-the-shape-of-space/>

Die Theorie von Einstein wurde 1919 experimentell während einer Sonnenfinsternis bestätigt. Die Geschichte dieses Experiments kann man auf folgender Webseite nachlesen:

<https://www.mpg.de/9236014/eddington-sonnenfinsternis-1919>

21.2. Krümmung von 2-dimensionalen Mannigfaltigkeiten. Folgende Definition für 2-dimensionale Riemannsche Mannigfaltigkeiten ist wort-wörtlich die gleiche Definition, welche wir schon auf Seite 195 für Riemannsche Untermannigfaltigkeiten von $\mathbb{R}^n$ eingeführt hatten.

Definition. Es sei (M, g) eine 2-dimensionale Riemannsche Mannigfaltigkeit.

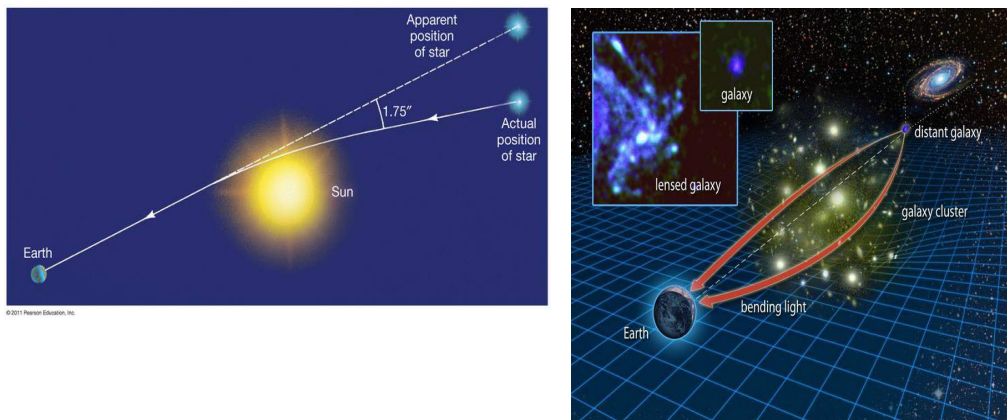
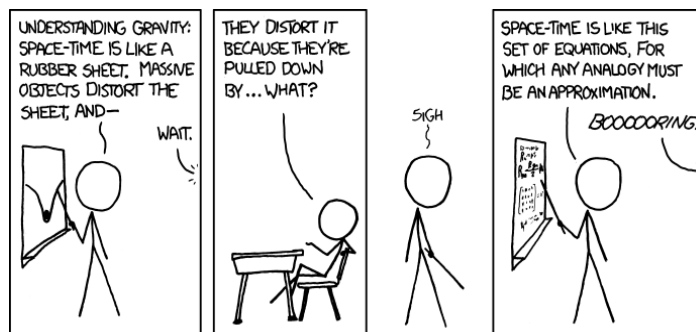


ABBILDUNG 131.

ABBILDUNG 132. Quelle: <http://xkcd.com/>

- (1) Wir sagen (M, g) ist *lokal isometrisch* zu einer 2-dimensionalen Riemannschen Mannigfaltigkeit (N, h) , wenn es für jedes $P \in M$ eine offene Umgebung U um P und eine offene Teilmenge $V \subset N$ sowie eine Isometrie $f: (U, g) \rightarrow (V, h)$ gibt.
- (2) Wir sagen (M, g) besitzt *konstante Krümmung* 0, wenn (M, g) lokal isometrisch zu $\mathbb{E}^2$ ist.
- (3) Wir sagen (M, g) besitzt *konstante Krümmung* +1, wenn M lokal isometrisch zu $\mathbb{S}^2$ ist.
- (4) Wir sagen (M, g) besitzt *konstante Krümmung* -1, wenn M lokal isometrisch zu $\mathbb{H}$ (oder äquivalent, zu $\mathbb{D}$) ist.

Wir erinnern an folgende Frage, welche wir schon auf Seite 196 gestellt hatten.

Frage. Welche (geschlossenen) 2-dimensionale Riemannsche (Unter-) Mannigfaltigkeiten besitzen konstante Krümmung?

Die Sphäre S^2 besitzt natürlich, per Definition, konstante Krümmung $+1$. Wir hatten zudem in Kapitel 15.3 gesehen, dass der Zylinder

$$Z = \{(x, y, z) \in \mathbb{R}^3 \mid x^2 + y^2 = 1\}$$

konstante Krümmung 0 besitzt. Andererseits hatten wir in Abbildung 95 anschaulich gesehen, dass der Standardtorus in $\mathbb{R}^3$ keine konstante Krümmung besitzt.

Folgender Satz besagt nun, dass der Torus mit der “richtigen” Metrik konstante Krümmung 0 besitzt.

Satz 21.1. *Auf den folgenden Mannigfaltigkeiten gibt es eine Riemannsche Metrik mit konstanter Krümmung 0:*

- (1) *der Torus,*
- (2) *der Zylinder $S^1 \times \mathbb{R}$,*
- (3) *das Möbiusband,*
- (4) *die Kleinsche Flasche.*

In dem Beweis von Satz 21.1 werden wir folgendes Lemma verwenden.

Lemma 21.2. *Es sei M eine Mannigfaltigkeit und G eine Gruppe, welche auf M frei, glatt und eigentlich operiert. Es sei g eine Riemannsche Struktur auf M . Wenn g invariant ist unter der Operation durch G ²⁴⁹, dann gibt es genau eine Riemannsche Struktur $\bar{g}$ auf M/G mit²⁵⁰*

$$g = p^*(\bar{g}).$$

Insbesondere ist die Projektionsabbildung $p: M \rightarrow M/G$ eine lokale Isometrie.

Beweis von Lemma 21.2. Die eindeutige Existenz von $\bar{g}$ auf M/G wird ganz ähnlich wie Lemma 20.9 bewiesen. Satz 19.2 besagt, dass die Projektionsabbildung $p: M \rightarrow M/G$ ein lokaler Diffeomorphismus ist. Mit der Wahl der Riemannschen Strukturen ist dies nun eine lokale Isometrie. $\square$

Wir wenden uns nun dem Beweis von Satz 21.1 zu.

Beweis von Satz 21.1.

- (1) Wir betrachten nun den Torus geschrieben als $T = \mathbb{R}^2/\mathbb{Z}^2$. Wir bezeichnen jetzt mit $p: \mathbb{R}^2 \rightarrow \mathbb{R}^2/\mathbb{Z}^2$ die Projektionsabbildung und wir bezeichnen mit g die übliche Riemannsche Struktur auf $\mathbb{R}^2$. Diese ist offensichtlich invariant unter der Operation von $\mathbb{Z}^2$ auf $\mathbb{R}^2$. Es folgt also aus Lemma 21.2, dass es genau eine Riemannsche Struktur $\bar{g}$ auf $\mathbb{R}^2/\mathbb{Z}^2$ mit $g = p^*(\bar{g})$ gibt. Zudem besagt Lemma 21.2, dass die Abbildung $p: \mathbb{R}^2 \rightarrow \mathbb{R}^2/\mathbb{Z}^2$ eine lokale Isometrie ist. Insbesondere ist der Torus $\mathbb{R}^2/\mathbb{Z}^2$ lokal isometrisch zu $\mathbb{E}^2$.

²⁴⁹Hierbei definieren wir G -invariante Riemannsche Strukturen ganz analog zu den G -invarianten Differentialformen, welche wir auf Seite 249 eingeführt hatten. Eine Riemannsche Struktur ist per Definition genau dann G -invariant, wenn Multiplikation mit einem $g \in G$ eine Isometrie der gegebenen Riemannschen Mannigfaltigkeit ist.

²⁵⁰Hierbei ist $\bar{g}$ wie auf Seite 159 definiert.

- (2) Wir betrachten folgende Operation von $G = \mathbb{Z}$ auf $M = \mathbb{R}^2$:

$$\begin{aligned}\mathbb{Z} \times \mathbb{R}^2 &\rightarrow \mathbb{R}^2 \\ (n, (x, y)) &\mapsto (x + n, y).\end{aligned}$$

Die Gruppe $G = \mathbb{Z}$ operiert durch Isometrien auf $\mathbb{E}^2$. Es folgt wie in (2), dass der Quotient $\mathbb{R}^2/\mathbb{Z}$ lokal isometrisch zu $\mathbb{E}^2$ ist. Andererseits ist der Quotient $\mathbb{R}^2/\mathbb{Z}$ offensichtlich²⁵¹ homöomorph zu $S^1 \times \mathbb{R}$.²⁵²

- (3) Wir betrachten noch einmal, wie auf Seite 222, folgende Operation von $G = \mathbb{Z}$ auf $X = \mathbb{R} \times (-1, 1)$:

$$\begin{aligned}\mathbb{Z} \times (\mathbb{R} \times (-1, 1)) &\rightarrow \mathbb{R} \times (-1, 1) \\ (n, (x, y)) &\mapsto (x + n, (-1)^n y).\end{aligned}$$

Die Gruppe G operiert also durch Isometrien. Wie in (1) sehen wir also, dass das Möbiusband $(\mathbb{R} \times (-1, 1))/\mathbb{Z}$ lokal isometrisch zu $\mathbb{R} \times (-1, 1)$ ist. Insbesondere ist es lokal isometrisch zu $\mathbb{E}^2$.

- (4) Wir betrachten die folgenden beiden Isometrien von $\mathbb{E}^2$:

$$\begin{aligned}A: \mathbb{R}^2 &\rightarrow \mathbb{R}^2 & \text{und} & & B: \mathbb{R}^2 &\rightarrow \mathbb{R}^2 \\ (x, y) &\mapsto (x + 1, 1 - y) & & & (x, y) &\mapsto (x, y + 1).\end{aligned}$$

Wir bezeichnen mit G die Untergruppe aller Diffeomorphismen von $\mathbb{R}^2$, welche von A und B erzeugt wird.²⁵³ Man kann nun, mit etwas Aufwand, zeigen, dass G auf $\mathbb{R}^2$ glatt, frei und eigentlich operiert. Es folgt nun, ganz analog zu (1), dass $\mathbb{R}^2/G$ lokal isometrisch zu $\mathbb{E}^2$ ist.

Wir betrachten nun $X = [0, 1] \times [0, 1]$ mit der Äquivalenzrelation, welche erzeugt wird durch $(0, y) \sim (1, 1 - y)$ und $(x, 0) \sim (x, 1)$. Man kann nun zeigen, dass die Abbildung

$$\begin{aligned}X/\sim &\rightarrow \mathbb{R}^2/G \\ [(x, y)] &\mapsto [(x, y)]\end{aligned}$$

ein Homöomorphismus ist. Mit anderen Worten, $\mathbb{R}^2/G$ ist homöomorph zur Kleinschen Flasche. $\square$

Bemerkung. Wenn wir den Torus $S^1 \times S^1$ als Teilmenge von $\mathbb{R}^2 \times \mathbb{R}^2 = \mathbb{R}^4$ auffassen, dann kann man, ganz analog zur Diskussion auf den Seiten 164 und 196 zeigen, dass das übliche Skalarprodukt auf $\mathbb{R}^4$ eine Riemannsche Metrik auf $S^1 \times S^1$ mit konstanter Krümmung 0 definiert.

Nachdem die Spiegelung $P \mapsto -P$ eine Isometrie von $\mathbb{S}^2$ ist, kann man ganz analog zu Satz 21.1 auch leicht folgenden Satz beweisen.

²⁵¹Was ist der Homöomorphismus?

²⁵²Der Zylinder $S^1 \times \mathbb{R}$ ist natürlich auch homöomorph zu $Z = \{(x, y, z) \in \mathbb{R}^3 \mid x^2 + y^2 = 1\}$. Wie gerade erst erwähnt hatten wir schon in Kapitel 15.3 gesehen, dass der Zylinder $Z \cong S^1 \times \mathbb{R}$ eine Riemannsche Metrik mit konstanter Krümmung 0 besitzt.

²⁵³D.h. die Elemente in G sind Diffeomorphismen die man durch Verknüpfungen von A, A^{-1} sowie B und B^{-1} erhält.

Satz 21.3. *Es gibt eine Riemannsche Metrik auf dem projektiven Raum $\mathbb{R}P^2 = S^2/\{\pm 1\}$ mit konstanter Krümmung $+1$.*

21.3. Flächen von höherem Geschlecht. Wir wollen nun noch der Frage nachgehen, ob die Flächen von höherem Geschlecht ebenfalls eine Riemannsche Metrik konstanter Krümmung besitzen. Nachdem wir die Flächen von höherem Geschlecht über Verklebungen eingeführt hatten, wollen wir nun nach Lemma 21.2 eine weitere Möglichkeit kennenlernen, um eine Riemannsche Metrik konstanter Krümmung zu konstruieren. Wir führen dazu folgende Definitionen ein. Es sei N eine Mannigfaltigkeit. In unseren Anwendungen ist $N = \mathbb{E}^2$, $N = \mathbb{H} = \mathbb{D}$ oder $N = \mathbb{S}^2$. Ein N -Atlas für eine Mannigfaltigkeit M ist eine Familie $\{\Phi_i: U_i \rightarrow V_i\}_{i \in I}$ von Diffeomorphismen von offenen Teilmengen in M zu offenen Teilmengen in N , so dass $\bigcup_{i \in I} U_i = M$. Für $i, j \in I$ bezeichnen wir dann die Abbildung

$$\Phi_i(U_i \cap U_j) \xrightarrow{(\Phi_i|_{U_i \cap U_j})^{-1}} U_i \cap U_j \xrightarrow{\Phi_j|_{U_i \cap U_j}} \Phi_j(U_i \cap U_j)$$

als den *Kartenwechsel* von Φ_i zu Φ_j .

Lemma 21.4. *Es sei M eine 2-dimensionale Mannigfaltigkeit. Folgende beiden Aussagen sind äquivalent:*

- (1) *M besitzt eine Riemannsche Metrik von konstanter Krümmung 0.*
- (2) *M besitzt einen $\mathbb{E}^2$ -Atlas, so dass alle Kartenwechsel Isometrien von Teilmengen von $\mathbb{E}^2$ sind.*

Die analogen Aussagen gelten auch für $\mathbb{D}$ und konstante Krümmung -1 sowie für $\mathbb{S}^2$ und konstante Krümmung $+1$.

Beispiel. Im Beweis von Lemma 19.1 hatten wir gezeigt, dass das Möbiusband einen Atlas bestehend aus zwei Karten besitzt. Der einzige Kartenwechsel ist gegeben durch die Identität (auf einer der beiden Komponenten) und die Abbildung $(a, b) \mapsto (a - 1, 1 - b)$ (auf der anderen Komponente). Beide Abbildungen sind Isometrien von $\mathbb{E}^2$. Also erhalten wir nach Satz 21.1 einen weiteren Beweis dafür, dass das Möbiusband eine Riemannsche Metrik von konstanter Krümmung 0 besitzt. Ganz analog sieht man auch, dass die Kleinsche Flasche eine Riemannsche Metrik mit konstanter Krümmung 0 besitzt.

Beweis. Es sei M also eine 2-dimensionale Mannigfaltigkeit. Wir bezeichnen mit h die übliche Riemannsche Metrik auf $\mathbb{R}^2$.

- (1) Wir nehmen zuerst an, dass M eine Riemannsche Metrik g von konstanter Krümmung 0 ist. Per Definition bedeutet dies, dass es einen Atlas $\{\Phi_i: U_i \rightarrow V_i\}_{i \in I}$ gibt, so dass jedes Φ_i eine Isometrie von (U_i, g) zu (V_i, h) ist, wobei (V_i, h) eine Teilmenge von $\mathbb{E}^2$ ist. Dann sind aber auch alle Kartenwechsel Verknüpfungen von Isometrien, also selber schon Isometrien. Also ist (2) erfüllt.
- (2) Wir nehmen nun an, dass M einen $\mathbb{E}^2$ -Atlas besitzt, so dass alle Kartenwechsel Isometrien von $\mathbb{E}^2$ sind. Wir definieren jetzt eine Riemannsche Struktur auf M wie folgt. Für einen Punkt $P \in M$ wählen wir nun eine Karte $\Phi: U \rightarrow V$ um P

in unserem $\mathbb{E}^2$ -Atlas und wir definieren $g_P := \Phi^*(h_{\Phi(P)})$. Hätten wir eine andere Karte in dem $\mathbb{E}^2$ -Atlas gewählt, hätten wir die gleiche Metrik erhalten, nachdem die Kartenwechsel Isometrien sind. Man kann nun leicht zeigen, dass g eine Metrik von konstanter Krümmung 0 ist.²⁵⁴

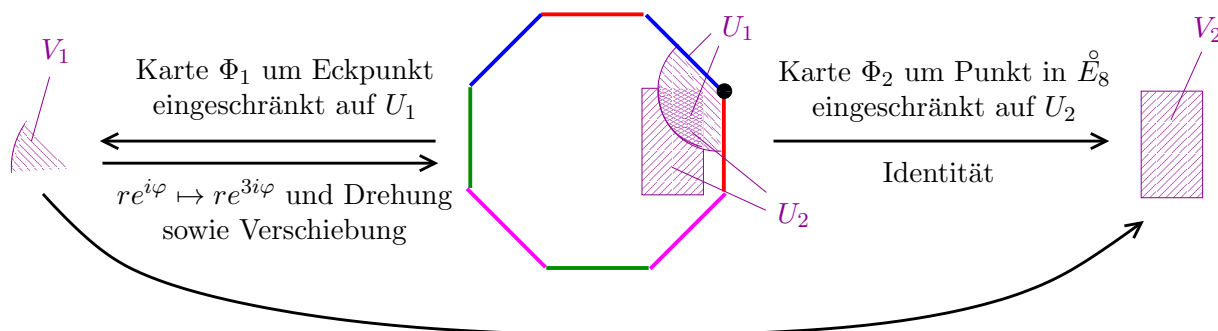
Die Aussagen für $\mathbb{D}$ und $\mathbb{S}^2$ werden natürlich ganz analog bewiesen. $\square$

21.4. Reguläre n -Ecke und hyperbolische Geometrie. Im Beweis von Satz 19.4 hatten wir einen Atlas angegeben, welcher beweist, dass die Fläche von Geschlecht 2 eine 2-dimensionale Mannigfaltigkeit ist. Der Kartenwechsel von der Karte um einen Eckpunkt zu der Karte für das Innere des Oktagons ist hierbei, wie in Abbildung 133 skizziert, gegeben durch die Verknüpfung der folgenden drei Typen von Abbildungen:

- (1) Verschiebung,
- (2) Drehung,
- (3) die Einschränkung der Abbildung $re^{i\varphi} \mapsto re^{3i\varphi}$ auf eine offene Teilmenge von $\mathbb{C}$.

Die ersten beiden Abbildungen sind Isometrien von $\mathbb{E}^2$, aber die dritte Abbildung ist keine Isometrie von $\mathbb{E}^2$. Der Atlas im Beweis von Satz 19.4 erfüllt also *nicht* die Bedingung von Lemma 21.4 (2).

Die Abbildung $re^{i\varphi} \mapsto re^{3i\varphi}$ wiederum taucht deswegen auf, weil der Innenwinkel in einem euklidischen Oktagon $\frac{3}{4}\pi$ beträgt, und wir den Innenwinkel auf $\frac{1}{4}\pi$ reduzieren müssen, damit wir die acht “Tortenstücke” zu einem einzigen Kreis zusammenfassen können.



Kartenwechsel ist Verknüpfung von $z \mapsto z^3$ und Drehungen und Verschiebungen

ABBILDUNG 133. Kartenwechsel für die Fläche von Geschlecht 2.

Wir könnten dieses Problem umschiffen, wenn wir ein Oktagon mit Innenwinkel $\frac{1}{4}\pi$ hätten. Dies ist im euklidischen Fall natürlich nicht möglich. Aber wir werden jetzt sehen, dass dies im hyperbolischen Fall kein Problem ist.

²⁵⁴In der Tat, denn sei $P \in M$. Wir wählen eine Karte $\Phi: U \rightarrow V$ aus dem $\mathbb{E}^2$ -Atlas um P . Dann gilt nicht nur $g_P := \Phi^*(h_{\Phi(P)})$ sondern auch $g_Q = \Phi^*(h_{\Phi(Q)})$ für alle $Q \in U$. Also ist Φ eine Isometrie von einer offenen Teilmenge von (M, g) zu einer offenen Teilmenge von $\mathbb{E}^2 = (\mathbb{R}^2, h)$.

Lemma 21.5. *Es gibt ein $r > 0$, so dass das von $Q_k = e^{2\pi i k/16}$, $k = 1, 3, \dots, 15$ in $\mathbb{D}$ aufgespannte hyperbolische Oktagon an allen Eckpunkten den Innenwinkel $\frac{\pi}{4}$ besitzt.*

Beweisskizze. Wir geben im Folgenden nur eine Skizze für den Beweis von Lemma 21.5. Die Beweisidee ist leicht nachvollziehbar, aber es ist etwas umständlich, diesen Beweis dann wirklich sauber aufzuschreiben.

Es sei $r \in (0, 1)$. Wir bezeichnen mit O_r das hyperbolische Oktagon, welches von den Punkten $re^{2\pi i k/16}$, $k = 1, 3, \dots, 15$ in $\mathbb{D}$ aufgespannt wird. Dieses Oktagon wird für r nahe an 0 und für r nahe an 1 in Abbildung 134 skizziert. Für $r \rightarrow 0$ konvergiert der Flächeninhalt gegen null. Mithilfe von Satz 15.4 kann man nun zeigen, dass der Innenwinkel von O_r gegen den euklidischen Innenwinkel, d.h. gegen $\frac{3\pi}{4}$ konvergiert. Für $r \rightarrow 1$ konvergieren die Eckpunkte gegen Punkte auf dem Randkreis S^1 , aber nachdem hyperbolische Geraden senkrecht auf dem Randkreis S^1 stehen, konvergiert der Innenwinkel mit $r \rightarrow 1$ gegen 0. Aus Stetigkeitsgründen²⁵⁵ muss es nun auch ein $r \in (0, 1)$ geben, so dass der Innenwinkel genau $\frac{\pi}{4}$ beträgt. $\square$

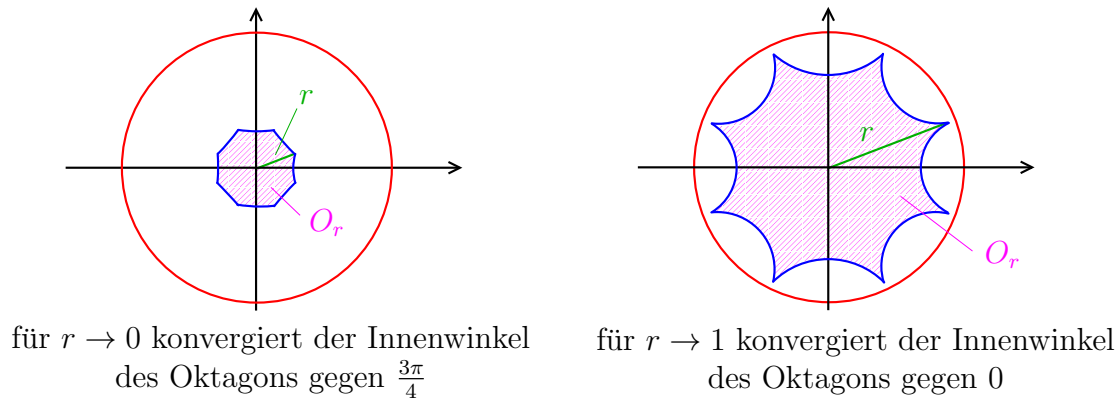


ABBILDUNG 134. Hyperbolische Oktagon von verschiedener Größe.

21.5. Hyperbolische Metriken auf Flächen von Geschlecht 2. Unser Ziel in diesem Kapitel ist es folgenden Satz zu beweisen:

Satz 21.6. *Die Fläche von Geschlecht 2 besitzt eine Riemannsche Metrik mit konstanter Krümmung -1 .*

Wie sich im letzten Kapitel schon angedeutet hat, werden wir Satz 21.6 dadurch beweisen, dass wir uns die Fläche von Geschlecht 2 nicht mehr als Quotient von einem euklidischen Oktagon, sondern als Quotient von einem hyperbolischen Oktagon vorstellen.

Bevor wir mit dem eigentlichen Beweis von Satz 21.6 anfangen, wollen wir noch eine Definition einführen.

²⁵⁵Hier ist der Haken bei dieser Beweisskizze: warum ist der Innenwinkel, beziehungsweise der Flächeninhalt, stetig in r ?

Definition.

- (1) Die Kreisscheibe in $\mathbb{D}$ mit Mittelpunkt $P \in \mathbb{D}$ und Radius $r \in \mathbb{R}_{>0}$ ist die Teilmenge

$$U_r(P) = \{Q \in \mathbb{D} \mid d_{\mathbb{D}}(P, Q) \leq r\}.$$

Wir bezeichnen jede hyperbolische Gerade durch P als *Durchmesser* der Kreisscheibe.²⁵⁶

- (2) Eine Teilmenge einer Kreisscheibe, welche durch zwei Durchmesser g und h begrenzt ist, wird *Tortenstück* genannt. Wir bezeichnen den Winkel zwischen den beiden Durchmessern als den *Öffnungswinkel* und wir bezeichnen den Radius des Kreises als die *Seitenlänge* des Tortenstücks.

Ganz analog definieren wir auch Tortenstücke in $\mathbb{E}^2$ und Tortenstücke in $\mathbb{H}$.

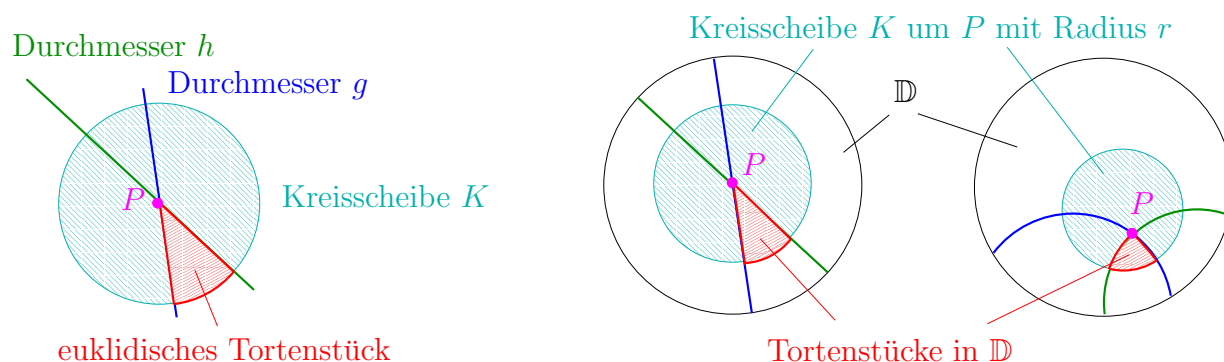


ABBILDUNG 135. Euklidische und hyperbolische Tortenstücke.

Wir erinnern zuerst an folgenden Satz, welchen man schon mithilfe der Schulmathematik beweisen kann.

Lemma 21.7. *Es seien K, L zwei Tortenstücke in $\mathbb{E}^2$ mit gleicher Seitenlänge und gleichem Öffnungswinkel. Dann gibt es genau eine orientierungserhaltende Isometrie f von $\mathbb{E}^2$ mit $f(K) = L$.*

Beweisskizze. Wir werden die Aussage nicht verwenden, wir werden diese deswegen auch nicht präzise beweisen. Aber es sollte klar sein, dass man mithilfe einer Translation und einer Drehung K in L überführen kann. Es verbleibt die Eindeutigkeit zu beweisen. Es

²⁵⁶Jede hyperbolische Kreisscheibe $U_r(P) = \{Q \in \mathbb{D} \mid d_{\mathbb{D}}(P, Q) \leq r\}$ in $\mathbb{D}$ ist auch eine euklidische Kreisscheibe $D_s(R) = \{Q \in \mathbb{D} \mid \|R - Q\| \leq s\}$, allerdings ist $r \neq s$ und für $P \neq 0$ ist auch $R \neq P$. Die Tatsache, dass $U_r(P)$ eine euklidische Kreisscheibe ist, ist klar für $P = 0$. Für $P \neq 0$ folgt die Aussage aus dem Spezialfall $P = 0$ und folgenden beiden Tatsachen:

- (a) Möbiustransformation operieren transitiv auf hyperbolischen Kreisscheiben von gleichem Radius, und
- (b) Möbiustransformation führen, ganz analog zu Satz 14.6, euklidische Kreise in $\mathbb{D}$ in euklidische Kreise in $\mathbb{D}$ über.

genügt zu zeigen, dass für jedes Tortenstück T die Identität die einzige orientierungserhaltende Isometrie ist, welche T auf T schickt. Diese Aussage wiederum folgt leicht aus der Tatsache, dass jede orientierungserhaltende Isometrie von $\mathbb{E}^2$ durch die Verknüpfung von einer Translation und einer Drehung gegeben ist. Wir überlassen die Durchführung des Arguments als Übungsaufgabe. $\square$

Die wichtigsten Symmetrien vom hyperbolischen Raum sind die Möbiustransformationen. Für $\mathbb{H}$ hatten wir diese schon in Lemma 14.1 kennengelernt. Wir führen diese nun auch noch auf $\mathbb{D}$ ein. Wir bezeichnen dazu mit $\Phi: \mathbb{H} \rightarrow \mathbb{D}$, $\Phi(z) = \frac{z-i}{z+i}$ die Isometrie aus Satz 14.13. Wir sagen eine Abbildung $\alpha: \mathbb{D} \rightarrow \mathbb{D}$ ist eine *Möbiustransformation auf $\mathbb{D}$* , wenn es eine Möbiustransformation $\beta: \mathbb{H} \rightarrow \mathbb{H}$ gibt, so dass $\alpha = \Phi \circ \beta \circ \Phi^{-1}$.

Wir können jetzt das hyperbolische Analogon zu Lemma 21.7 formulieren.

Lemma 21.8.

- (1) *Es seien K, L Tortenstücke in $\mathbb{H}$ mit gleicher Seitenlänge und gleichem Öffnungswinkel. Dann gibt es genau eine Möbiustransformation f von $\mathbb{H}^2$ mit $f(K) = L$.*
- (2) *Es seien K, L Tortenstücke in $\mathbb{D}$ mit gleicher Seitenlänge und gleichem Öffnungswinkel. Dann gibt es genau eine Möbiustransformation f von $\mathbb{D}^2$ mit $f(K) = L$.*

Beweis.

- (1) Es seien K, L zwei Tortenstücke in $\mathbb{H}$ mit gleicher Seitenlänge und gleichem Öffnungswinkel. In Lemma ?? wurde schon folgende Aussage bewiesen: Es seien Δ_{ABC} und Δ_{DEF} zwei Dreiecke in $\mathbb{H}$ mit $\sphericalangle ABC = \sphericalangle DEF$ und mit $d(B, A) = d(E, D)$ und $d(B, C) = d(E, F)$. Dann gibt es genau eine Möbiustransformation f mit $f(B) = E$ und $f(\Delta_{ABC}) = \Delta_{DEF}$. Wir wenden nun die Aussage an auf die Dreiecke, welche von den drei Eckpunkten der Tortenstücke aufgespannt werden und erhalten die gewünschte eindeutig bestimmte Möbiustransformation.
- (2) Diese Aussage folgt sofort aus (1). $\square$

Wir wenden uns nun dem eigentlichen Beweis von Satz 21.6 zu.

Beweis von Satz 21.6. Wir betrachten jetzt das reguläre hyperbolische Oktagon H_8 in $\mathbb{D}$, welches wir im Beweis von Lemma 21.5 konstruiert hatten. Die Eckpunkte des Oktagon sind die Punkte $Q_k = re^{2\pi ik/16}$ mit $k = 1, 3, \dots, 15$, wobei r so gewählt ist, dass die Innenwinkel alle gleich $\frac{\pi}{4}$ sind. Für Punkte $A, B \in \mathbb{D}^2$ bezeichnen wir mit $\overline{AB}$ die hyperbolische Strecke von A nach B . Zudem bezeichnen wir für $\varphi \in [0, 2\pi]$ mit $s_\varphi: \mathbb{D} \rightarrow \mathbb{D}$ die Spiegelung an der hyperbolischen Geraden $\{re^{i\varphi} \mid r \in (-1, 1)\} \subset \mathbb{D}$. (Wir hatten schon in Lemma 14.15 gesehen, dass diese Spiegelung durch die übliche euklidische Spiegelung gegeben ist.) Wir bezeichnen mit $\sim$ die Äquivalenzrelation auf H_8 , welche, ganz analog zur Definition auf Seite 215, erzeugt wird durch die Relationen

$$P \in \overline{Q_{2k-1}Q_{2k+1}} \sim s_{2\pi(2k+2)/16}(P) \in \overline{Q_{2k+3}Q_{2k+5}}$$

für $k = 0, 1, 4, 5$. Es ist leicht zu sehen, dass der topologische Raum $H_8/\sim$ homöomorph zu

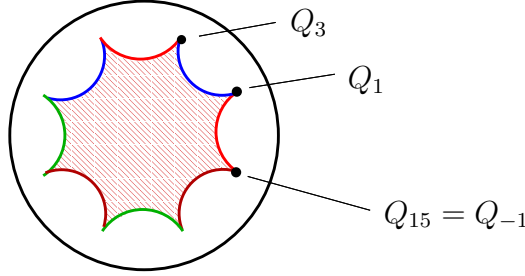


ABBILDUNG 136. Hyperbolisches Oktagon mit Innenwinkel $\frac{\pi}{4}$ mit identifizierten Seiten.

$E_8/\sim$ ist.²⁵⁷ Der topologische Raum $H_8/\sim$ ist also homöomorph zur Fläche von Geschlecht 2. Es genügt daher für $H_8/\sim$ einen $\mathbb{D}$ -Atlas zu finden, so dass alle Kartenwechsel durch Isometrien von $\mathbb{D}$ gegeben sind.

Wir bezeichnen mit p die Projektionsabbildung $H_8 \rightarrow H_8/\sim$. Wir verfahren im Folgenden soweit wie möglich wie auf Seite 238.

Es sei also zuerst P ein Punkt im Inneren $\mathring{H}_8$ von H_8 . Die Abbildung

$$\begin{aligned} \Phi: U := p(\mathring{H}_8) &\rightarrow V = \mathring{H}_8 \subset \mathbb{C} = \mathbb{R}^2 \\ p(X) &\mapsto X \end{aligned}$$

ist dann eine Karte um $p(P)$.

Es sei nun P ein Punkt, welcher im Inneren einer Kante von H_8 liegt. Wir wollen eine Karte um $p(P)$ finden. O.B.d.A. sei $P \in \overline{Q_{-1}Q_1}$. Wir bezeichnen mit $P' = s_{\frac{\pi}{4}}(P)$ den Punkt, welcher mit P identifiziert wird. Wir wählen ein $\epsilon > 0$, so dass $U_\epsilon(P)$ keinen Eckpunkt des Oktagons berührt und so dass $U_\epsilon(P) \cap U_\epsilon(P') = \emptyset$. Wir schreiben jetzt $U := U_\epsilon(P) \cap H_8$ und $U' := U_\epsilon(P') \cap H_8$. Dann ist $p(U \cup U')$ eine offene Umgebung um $p(P)$. Es sei $f: \mathbb{D}^2 \rightarrow \mathbb{D}^2$ die Spiegelung an der hyperbolischen Gerade, welche durch die Kante $\overline{Q_{-1}Q_1}$ definiert wird. Es folgt nun leicht aus den Definitionen, dass die Abbildung

$$\begin{aligned} p(U \cup U') &\rightarrow U_\epsilon(P) \\ p(X) &\mapsto \begin{cases} X, & \text{wenn } X \in U, \\ f(s_{\frac{\pi}{4}}(X)), & \text{wenn } X \in U' \end{cases} \end{aligned}$$

wohl-definiert ist, und dass diese Abbildung ein Homöomorphismus ist. Insbesondere ist dies also eine Karte für H_8 um $p(P)$. Diese Karte wird in Abbildung 137 skizziert.

Wir betrachten nun noch $p(Q_1) = p(Q_3) = \dots = p(Q_{15})$, das Bild der Eckpunkte des Oktagons. Wir wählen ein $\epsilon > 0$, so dass die hyperbolischen ϵ -Umgebungen um die Eckpunkte $Q_1, Q_3, \dots, Q_{15}$ disjunkt sind. Für $k = 1, 3, \dots, 15$ betrachten wir dann jeweils

²⁵⁷Für einen Punkt $P \neq 0$ in E_8 sei $e(P)$ die eindeutig bestimmte positive reelle Zahl mit $t \cdot P \in \partial E_8$ und ganz analog definieren wir $h(P)$ auf $H_8 \setminus \{0\}$. Dann induziert die Abbildung $0 \mapsto 0$ und $P \mapsto \frac{h(P)}{e(P)} \cdot P$ für $P \neq 0$ einen Homöomorphismus $E_8/\sim \rightarrow H_8/\sim$.

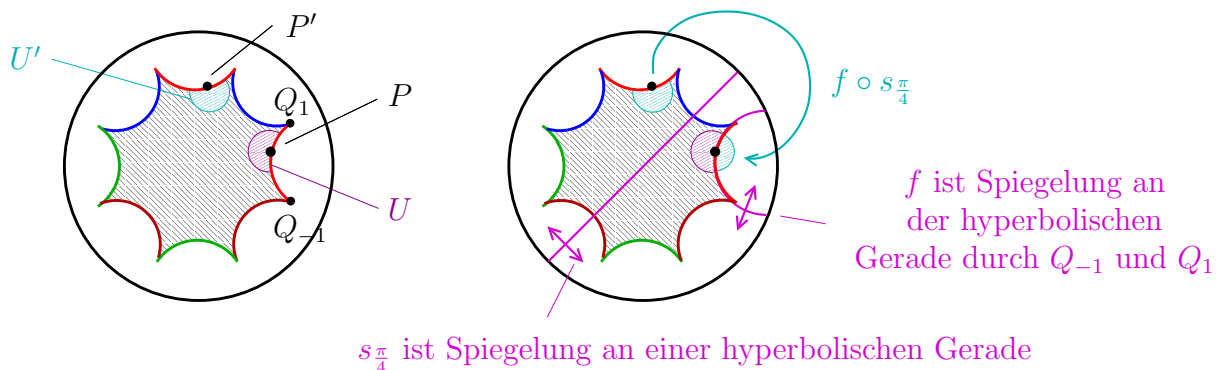


ABBILDUNG 137. Karte für einen Punkt im Inneren einer Kante von H_8 .

das Tortenstück²⁵⁸ $U_k := U_\epsilon(Q_k) \cap H_8$. Die Mengen $U_1, U_3, \dots, U_{15}$ sind disjunkt und es folgt aus den Definitionen, dass

$$p(U_1) \cup p(U_3) \cup \dots \cup p(U_{15})$$

eine offene Umgebung von $p(Q_1) = p(Q_3) = \dots = p(Q_{15})$ ist.

Für $\varphi \in [0, 2\pi]$ bezeichnen wir dann mit

$$T(\varphi) := U_\epsilon(0) \cap \{re^{i\alpha} \mid \alpha \in [\varphi, \varphi + \frac{\pi}{4}]\}$$

das hyperbolische Tortenstück²⁵⁹ um den Ursprung mit Öffnungswinkel $\frac{\pi}{4}$ und hyperbolischer Seitenlänge ϵ , welches um den Winkel φ zur x -Achse verdreht ist. Das Tortenstück $T(\varphi)$ ist in Abbildung 138 skizziert.

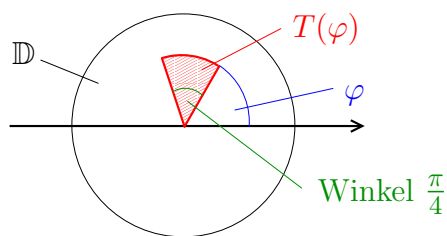


ABBILDUNG 138.

²⁵⁸Dies ist in der Tat ein Tortenstück im Sinne der Definition auf Seite 260, denn es ist eine Teilmenge der Kreisscheibe $U_\epsilon(Q_k)$, welche durch die beiden hyperbolischen Geraden begrenzt wird, welche durch die beiden Kanten am Punkt Q_k definiert sind.

²⁵⁹Auch dies ist ein Tortenstück im Sinne der Definition auf Seite 260.

Für $k \in \{1, 3, \dots, 15\}$ definieren wir nun $\gamma(k)$ durch die folgende Tabelle:

k	1	3	5	7	9	11	13	15
$\gamma(k)$	$\frac{3}{4}\pi$	π	$\frac{5}{4}\pi$	$\frac{1}{2}\pi$	$\frac{7}{4}\pi$	0	$\frac{1}{4}\pi$	$\frac{3}{2}\pi$

Für $k = 1, 3, \dots, 15$ besitzen die Tortenstücke U_k und $T(\gamma(k))$ die gleiche hyperbolische Seitenlänge ϵ und den gleichen Öffnungswinkel $\frac{\pi}{4}$. Es gibt also nach Lemma 21.8 jeweils genau eine Möbiustransformation f_k mit $f_k(U_k) = T(\gamma(k))$.

Wir betrachten jetzt die Abbildung

$$\begin{aligned} \Phi: p(U_1) \cup \dots \cup p(U_{15}) &\rightarrow U_\epsilon(0) \\ p(X) &\mapsto f_k(X), \quad \text{wenn } X \in U_k. \end{aligned}$$

Man kann nun leicht überprüfen, dass die Abbildung Φ wohl-definiert ist, und das dies in der Tat ein Homöomorphismus ist.

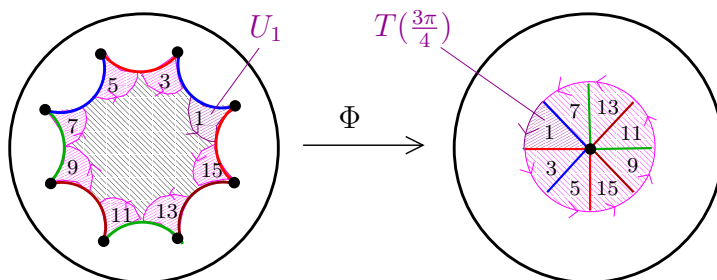


ABBILDUNG 139. Karte für die Projektion der Eckpunkte von H_8 .

Wir haben jetzt also wiederum bewiesen, dass $H_8/\sim$ eine 2-dimensionale Mannigfaltigkeit ist. Genauer gesagt, wir haben einen $\mathbb{D}$ -Atlas für $H_8/\sim$ gefunden. Alle Kartenwechsel sind nun gegeben durch Verknüpfungen von Spiegelungen und Möbiustransformation. Aber nach Satz 14.3 und Satz 14.12 sind beides Isometrien. Mit anderen Worten, alle Kartenwechsel sind gegeben durch Isometrien von $\mathbb{D}^2$.

Es folgt daher aus Lemma 21.4, dass $H_8/\sim$ eine Metrik besitzt, welche konstante Krümmung -1 besitzt. Wir haben damit Satz 21.6 bewiesen. $\square$

21.6. Die Klassifizierung von geschlossenen Flächen. Wir führen nun noch folgende Bezeichnungen ein:

- (1) eine *Fläche* ist eine geschlossene 2-dimensionale Mannigfaltigkeit,
- (2) wir bezeichnen die Sphäre als *Fläche von Geschlecht 0*,
- (3) wir bezeichnen den Torus als *Fläche von Geschlecht 1*.

Folgender Satz klassifiziert orientierbare Flächen bis auf einen Diffeomorphismus.

Satz 21.9. *Zu jeder orientierbaren Fläche M gibt es genau ein $g \in \mathbb{Z}_{\geq 0}$, so dass M diffeomorph zur Fläche von Geschlecht g ist.*

Der Satz beinhaltet genau betrachtet zwei Aussagen:

- (1) Jede orientierbare Fläche ist diffeomorph zu einer Fläche von Geschlecht g .
- (2) Für $g \neq h$ sind die Flächen von Geschlecht g und h nicht diffeomorph.

Wir werden im nächsten Kapitel des Skripts zeigen, dass die Sphäre nicht diffeomorph zum Torus ist. Mit anderen Worten, die Fläche von Geschlecht 0 ist nicht diffeomorph zur Fläche von Geschlecht 1. Die allgemeine Aussage (2) werden wir in der Vorlesung algebraische Topologie beweisen. Aus Zeitgründen können wir leider die Aussage (1) nicht beweisen.

Bemerkung. Mehr zur Klassifikation von Flächen kann man auf

[https://en.wikipedia.org/wiki/Surface_\(topology\)#Classification_of_closed_surfaces](https://en.wikipedia.org/wiki/Surface_(topology)#Classification_of_closed_surfaces) finden. Insbesondere wird dort auch die Klassifikation von *nicht* orientierbaren Flächen angegeben.²⁶⁰

Wir führen auch noch folgende Bezeichnungen ein:

- (1) eine Riemannsche Metrik mit konstanter Krümmung $+1$ heißt *sphärisch*,
- (2) eine Riemannsche Metrik mit konstanter Krümmung 0 heißt *euklidisch*,
- (3) eine Riemannsche Metrik mit konstanter Krümmung -1 heißt *hyperbolisch*.

Den folgenden Satz von Gauß-Bonnet können wir leider auch nicht beweisen. Mehr Informationen dazu kann man hier finden:

https://en.wikipedia.org/wiki/Gauss-Bonnet_theorem

Wir werden den Satz von Gauß-Bonnet in Algebraischer Topologie I beweisen.

Satz 21.10. (Gauß-Bonnet) *Es sei M eine Fläche von Geschlecht g mit einer Riemannschen Metrik h von konstanter Krümmung $\epsilon \in \{-1, 0, 1\}$. Dann gilt*

$$\epsilon \cdot \text{Vol}(M, h) = 2 - 2g.$$

Das Volumen einer Riemannschen Mannigfaltigkeit ist immer positiv. Der Satz von Gauß-Bonnet besagt also, dass wenn eine Fläche von Geschlecht g eine sphärische Metrik besitzt, dann ist $2 - 2g > 0$, d.h. es ist $g = 0$. Mit anderen Worten, S^2 ist bis auf Diffeomorphismus die einzige orientierbare Fläche, welche eine sphärische Metrik besitzt. Ganz ähnlich sieht man, dass der Torus bis auf Diffeomorphismus die einzige orientierbare Fläche ist, welche eine euklidische Metrik besitzt.²⁶¹

Wir haben also gerade gesehen, dass wenn M eine Fläche von Geschlecht $g \geq 2$ ist, dann kann M weder eine sphärische noch eine euklidische Metrik besitzen. Wir hatten gerade in Satz 21.6 gesehen, dass die Fläche von Geschlecht 2 eine hyperbolische Metrik besitzt.

²⁶⁰Wir kennen bisher die Kleinsche Flasche und den projektiven Raum $\mathbb{R}P^2$ als Beispiele von nichtorientierbaren Flächen. Gibt es noch weitere Beispiele?

²⁶¹Man kann zudem auch zeigen, dass (bis auf Diffeomorphie) der projektive Raum $\mathbb{R}P^2$ die einzige nicht orientierbare sphärische Fläche ist, und dass die Kleinsche Flasche die einzige nicht orientierbare euklidische Fläche ist.

Eine kleine Verallgemeinerung dieses Beweises zeigt, dass die gleiche Aussage auch für jede orientierbare Fläche von Geschlecht ≥ 2 gilt.

Zusammengefasst gilt also folgender Satz:

Satz 21.11. *Es gelten folgende Aussagen:*

- (1) *Die Sphäre besitzt eine sphärische Metrik,*
- (2) *der Torus besitzt eine euklidische Metrik,*
- (3) *jede Fläche von Geschlecht $g \geq 2$ besitzt eine hyperbolische Metrik.*

Darüber hinaus kann keine dieser Flächen eine Riemannsche Metrik der anderen beiden Arten besitzen.

21.7. Die Geometrie von 3-dimensionalen Mannigfaltigkeiten. Für $n \in \mathbb{N}$ ist der n -dimensionale hyperbolische Raum definiert als die Riemannsche Mannigfaltigkeit, welche gegeben ist durch

$$\mathbb{H}^n = \{(x_1, \dots, x_n) \in \mathbb{R}^n \mid x_n > 0\},$$

zusammen mit der Riemannschen Struktur g , welche an jedem Punkt $(x_1, \dots, x_n) \in \mathbb{H}^n$ gegeben ist durch

$$\begin{aligned} g: \mathbb{R}^n \times \mathbb{R}^n &\rightarrow \mathbb{R} \\ (v, w) &\mapsto \frac{1}{x_n^2} \langle v, w \rangle. \end{aligned}$$

Wir sagen eine Metrik g auf einer n -dimensionalen Mannigfaltigkeit M ist *hyperbolisch*, wenn (M, g) lokal isometrisch zu $\mathbb{H}^n$ ist.

Wir haben gerade gesehen, dass, bis auf zwei Ausnahmen, alle orientierbaren 2-dimensionalen orientierbaren geschlossenen Mannigfaltigkeiten eine hyperbolische Metrik besitzen. Jetzt kann man sich natürlich fragen, wie es den in anderen Dimensionen ausschaut. Für $n \geq 4$ sind die “meisten”²⁶² geschlossenen n -dimensionalen Mannigfaltigkeiten nicht hyperbolisch.

Für geschlossene 3-dimensionale Mannigfaltigkeiten ist es sehr schwierig explizit eine hyperbolische Metrik anzugeben. In der Tat gab es bis etwa 1970 nur ein Beispiel von einer geschlossenen 3-dimensionalen Mannigfaltigkeit, welche eine hyperbolische Metrik besitzt, nämlich die sogenannte Seifert–Weber Mannigfaltigkeit:

https://en.wikipedia.org/wiki/Seifert-Weber_space

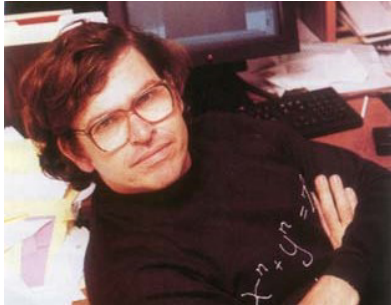
In der Mitte der 70er Jahre hat William Thurston die damals sehr gewagte Vermutung aufgestellt, dass die “meisten” geschlossenen 3-dimensionalen Mannigfaltigkeiten eine hyperbolische Metrik besitzen. Er hat diese Vermutung Ende der 70er Jahre für eine sehr große Klasse von 3-Mannigfaltigkeiten bewiesen. Die Vermutung von Thurston wurde dann endgültig in 2002 von Grigori Perelman bewiesen. Es gilt also folgender Satz:

Satz 21.12. (Thurston-Perelman) *Die “meisten” geschlossenen 3-dimensionalen Mannigfaltigkeiten besitzen eine hyperbolische Metrik.*

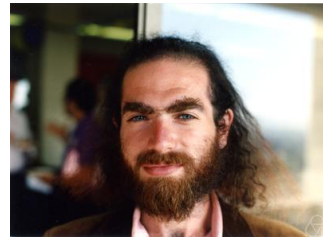
²⁶²Wir wollen und können jetzt die genaue Bedeutung von “meisten” nicht präzise machen.

Die genaue Formulierung von die “meisten” ist nicht so schwierig, aber für die jetzige Vorlesung doch zu umfangreich. Wir werden die Formulierung in der Vorlesung “Introduction to knot theory” nachreichen. Allerdings gehen die Beweise von Thurston und Perelman weit jenseits dessen, was man in einer (oder sogar mehreren) Vorlesungen behandeln kann.

Thurston erhielt 1982 für seine Arbeiten die Fields-Medaille und Perelman hätte diese 2006 bekommen, wenn er diese nicht abgelehnt hätte.



William Thurston (1946-2012)



Grigori Perelman *1966

ABBILDUNG 140.

22. DIE DE RHAM-KOHOMOLOGIE VON MANNIGFALTIGKEITEN

In diesem Kapitel für wir die de Rham-Kohomologiegruppen von Mannigfaltigkeiten ein. Diese werden es uns unter anderem ermöglichen zu zeigen, dass die Sphäre S^2 nicht diffeomorph zum Torus ist.

22.1. Quotientenvektorräume. In diesem kurzen Kapitel erinnern wir an die Definition der Quotientenvektorräume und an einige Eigenschaften derselben. Diese wurden schon in der Linearen Algebra II eingeführt.

Es sei im Folgenden V ein reeller Vektorraum und $U \subset V$ ein Untervektorraum. Für $v, w \in V$ schreiben wir

$$v \sim w \quad :\Longleftrightarrow \quad v - w \in U.$$

Es ist leicht nachzuprüfen, dass dies eine Äquivalenzrelation auf V definiert. Jede Äquivalenzklasse²⁶³ ist von der Form

$$v + U := \{v + u \mid u \in U\}$$

für ein $v \in V$. Wir bezeichnen nun mit V/U die Menge aller Äquivalenzklassen.

Wir wollen jetzt eine Addition auf V/U definieren. Es seien also $x, y \in V/U$. Dann gibt es $a, b \in V$, so dass $x = a + U$ und $y = b + U$. Wir definieren

$$x + y := a + b + U.$$

Wir müssen natürlich zeigen, dass diese Definition nicht von der Wahl von a und b abhängt. Wir wählen $a', b' \in V$ mit $x = a' + U$ und $y = b' + U$. Dann gilt

$$a' + b' + U = a + b + \underbrace{a' - a + b' - b}_{\in U} + U = a + b + U.$$

Ganz analog können wir eine Skalarmultiplikation auf V/U einführen. In der Tat, es sei $x \in V/U$ und $l \in \mathbb{R}$. Wir wählen $a \in V$ mit $x = a + U$. Wir definieren

$$l \cdot x := l \cdot a + U.$$

Wie oben kann man zeigen, dass diese Definition nicht von der Wahl von a abhängt.

Man kann nun leicht nachweisen, dass mit dieser Addition und dieser Skalarmultiplikation V/U wiederum ein reeller Vektorraum ist. Die “Null” des Vektorraums V/U ist dabei die Äquivalenzklasse $0 + U = U$. Dieser Vektorraum wird auch *Quotientenvektorraum* genannt.

Die Abbildung

$$\begin{aligned} V &\rightarrow V/U \\ v &\mapsto v + U \end{aligned}$$

ist surjektiv und ein Homomorphismus. Diese Abbildung wird die *Projektionsabbildung* von V auf V/U genannt.

²⁶³Äquivalenzklassen werden in diesem Zusammenhang oft auch *Restklassen* genannt.

22.2. Die Definition der de Rham-Kohomologie. Es sei M eine Mannigfaltigkeit. Für $k \in \mathbb{N}_0$ betrachten wir

$$\Omega_k(M) := \{\text{glatte } k\text{-Formen auf } M\}.$$

Wir setzen zudem $\Omega_{-1}(M) = \{0\}$. Nach Satz 20.6 (1) definiert das Differential für jedes k eine lineare Abbildung

$$d: \Omega_{k-1}(M) \rightarrow \Omega_k(M).$$

Wir sagen eine glatte Differentialform ω ist *geschlossen*, wenn $d\omega = 0$ und wir sagen ω ist *exakt*, wenn es eine Differentialform σ mit $d\sigma = \omega$ gibt. Für jedes k betrachten wir nun

$$Z_k(M) := \text{Kern}\{d: \Omega_k(M) \rightarrow \Omega_{k+1}(M)\} = \{\text{alle geschlossenen } k\text{-Formen}\}$$

und

$$B_k(M) := \text{Im}\{d: \Omega_{k-1}(M) \rightarrow \Omega_k(M)\} = \{\text{alle exakten } k\text{-Formen}\}.$$

Nach Satz 20.6 (1) gilt $dd\omega = 0$ für jede Differentialform $\omega \in \Omega_{k-1}(M)$. Mit anderen Worten, $B_k(M)$ ist ein Untervektorraum von $Z_k(M)$. Wir können also den Quotientenvektorraum

$$H^k(M) := Z_k(M)/B_k(M)$$

bilden. Der Vektorraum $H^k(M)$ wird die *k-te de Rham-Kohomologiegruppe von M* genannt.²⁶⁴

22.3. Berechnungen der de Rham-Kohomologie. Die de Rham-Kohomologie ordnet also jeder Mannigfaltigkeit M die de Rham-Kohomologiegruppen $H^k(M)$, $k \in \mathbb{Z}_{\geq 0}$ zu. Wir wollen im Folgenden einige de Rham-Kohomologiegruppen von Mannigfaltigkeiten bestimmen.

Lemma 22.1. *Für eine Mannigfaltigkeit M mit m Komponenten gilt*

$$H^0(M) \cong \mathbb{R}^m.$$

Insbesondere gilt für jede zusammenhängende Mannigfaltigkeit M , dass

$$H^0(M) \cong \mathbb{R}.$$

Beweis. Es sei also M eine Mannigfaltigkeit. Dann ist

$$\begin{aligned} H^0(M) &= \text{Kern}\{\Omega_0(M) \rightarrow \Omega_1(M)\} = \{f: M \rightarrow \mathbb{R} \mid df = 0\}. \\ &\quad \uparrow \qquad \qquad \qquad \uparrow \\ &\text{da } \Omega_{-1}(M) = 0 \qquad \text{da } \Omega_0(M) = \{\text{glatte Funktionen } M \rightarrow \mathbb{R}\} \\ &= \{\text{alle lokal konstanten Funktionen auf } M\} \\ &\quad \uparrow \\ &\text{Lemma 10.7 angewandt auf Mannigfaltigkeiten} \end{aligned}$$

Das Lemma folgt nun aus der Beobachtung, dass die Dimension vom Vektorraum der lokal konstanten Funktion auf M gerade die Anzahl der Komponenten von M ist.²⁶⁵ $\square$

²⁶⁴Georges de Rham (1903-1990) war ein schweizer Mathematiker.

²⁶⁵Warum ist dies der Fall?

Bemerkung. Es sei M eine Mannigfaltigkeit von Dimension $n \geq 1$. Der Vektorraum $H^k(M)$ ist für $k \in \{1, \dots, n\}$ der Quotient von zwei unendlich (sogar überabzählbar) dimensionalen Vektorräumen. Wenn M kompakt ist, dann gilt allerdings überraschenderweise, dass die de Rham-Kohomologiegruppen $H^k(M)$ endlich dimensionale Vektorräume sind. Wir können diese Aussage in dieser Vorlesung aus Zeitgründen leider nicht beweisen.

Lemma 22.2. *Für jede k -dimensionale Mannigfaltigkeit gilt $H^i(M) = 0$ für $i > k$.*

Beweis. Jeder Tangentialraum einer k -dimensionalen Mannigfaltigkeit ist nach Korollar 20.5 ein k -dimensionaler Vektorraum. Für $i > k$ folgt also aus Lemma 9.5, dass $\wedge^i T_P^* M = 0$ für alle $P \in M$. Insbesondere ist $\Omega^i(M) = 0$, und damit auch $H^i(M) = 0$. $\square$

Andererseits gilt folgender Satz.

Satz 22.3. *Es sei M eine orientierbare k -dimensionale geschlossene Mannigfaltigkeit. Dann gilt $\dim H^k(M) \geq 1$.*

Bemerkung. Es gilt die stärkere Aussage, dass die Dimension von $H^k(M)$ gerade gegeben ist durch die Anzahl der Komponenten von M . Wir können diese Aussage in dieser Vorlesung nicht beweisen.

Beweis. Es ist

$$H^k(M) = Z_k(M)/B_k(M) = \text{Kern}\{d: \Omega_k(M) \rightarrow \Omega_{k+1}(M)\}/B_k(M) = \Omega_k(M)/B_k(M).$$

$\uparrow$
nach dem Beweis von Lemma 22.2 ist $\Omega_{k+1}(M) = 0$

Wir betrachten die lineare Abbildung

$$\begin{aligned} \varphi: \Omega_k(M) &\rightarrow \mathbb{R} \\ \omega &\mapsto \int_M \omega. \end{aligned}$$

Da M orientierbar und geschlossen ist, folgt aus dem Satz 20.8 von Stokes, dass für alle $\omega \in \Omega_{k-1}(M)$ gilt $\varphi(\omega) = 0$. Mit anderen Worten, es ist $B_k(M) \subset \text{Kern}(\varphi)$. Die lineare Abbildung φ induziert also eine wohl-definierte lineare Abbildung

$$\begin{aligned} \varphi: H^k(M) = \Omega_k(M)/B_k(M) &\rightarrow \mathbb{R} \\ [\omega] &\mapsto \int_M \omega. \end{aligned}$$

Es genügt nun also zu zeigen, dass es eine k -Form ω auf M gibt mit $\varphi(\omega) \neq 0$. Es sei $P \in M$ ein beliebiger Punkt. Wir wählen eine Karte $\Phi: U \rightarrow V$ und eine nichtnegative glatte Funktion $z: \mathbb{R}^n \rightarrow \mathbb{R}_{\geq 0}$ mit $z(\Phi(P)) > 0$ und $\text{Träger}(z) \subset V$.²⁶⁶ Die Karte Φ und die Funktion z sind in Abbildung 141 skizziert. Wir betrachten nun die k -Form ω auf M ,

²⁶⁶Warum gibt es solch eine Funktion?

dass es eine $(k-1)$ -Form δ auf M gibt mit

$$\omega - \sigma = d\delta.$$

Dann gilt auch

$$f^*\omega - f^*\sigma = f^*(\omega - \sigma) = f^*(d\delta) = d(f^*\delta)$$

$\uparrow$
Satz 20.6

Dies impliziert, dass $[f^*\omega] = [f^*\sigma] \in H^k(M)$. Wir haben damit die zweite Behauptung bewiesen.

Es ist leicht zu sehen, dass die Abbildung wohl-definiert und linear ist. $\square$

Satz 22.5. Für jedes k definieren die Abbildungen

$$M \mapsto H^k(M)$$

und

$$(f: M \rightarrow N) \mapsto (f^*: H^k(N) \rightarrow H^k(M))$$

einen kontravarianten Funktor von der Kategorie der Mannigfaltigkeiten zu der Kategorie der Vektorräume.

Beweis. Im Hinblick auf Satz 22.4 genügt es folgende zwei Aussagen zu bewiesen:

- (1) für eine Mannigfaltigkeit M gilt $\text{id}_M^* = \text{id}_{H^k(M)}$,
- (2) für glatte Abbildungen $f: L \rightarrow M$ und $g: M \rightarrow N$ gilt $(f \circ g)^* = g^* \circ f^*$.

Beide Aussagen folgen leicht aus den Definitionen. $\square$

Das folgende Korollar besagt insbesondere, dass diffeomorphe Mannigfaltigkeiten isomorphe de Rham-Kohomologiegruppen besitzen.

Korollar 22.6. Es sei $f: M \rightarrow N$ ein Diffeomorphismus zwischen Mannigfaltigkeiten, dann ist für alle k die induzierte Abbildung $f^*: H^k(N) \rightarrow H^k(M)$ ein Isomorphismus.

Der Beweis von Korollar 22.6 ist fast identisch mit dem Beweis von Korollar 7.3.

Definition. Es sei $f: M \rightarrow N$ eine glatte Abbildung zwischen Mannigfaltigkeiten. Wir sagen f besitzt ein *Linksinverse*, wenn es eine glatte Abbildung $g: N \rightarrow M$ gibt, so dass $g \circ f = \text{id}_M$.

Folgendes Lemma ist ganz ähnlich der Aussage von Übungsaufgabe 4 in Übungsblatt 4.

Lemma 22.7. Es sei $f: M \rightarrow N$ eine glatte Abbildung zwischen Mannigfaltigkeiten, welche ein Linksinverse besitzt. Dann ist für alle k die induzierte Abbildung $f^*: H^k(N) \rightarrow H^k(M)$ ein Epimorphismus.

Beispiel. Die Abbildung

$$\begin{array}{lll} f: S^1 \rightarrow S^1 \times S^1 & & g: S^1 \times S^1 \rightarrow S^1 \\ x \mapsto (x, 1) & \text{besitzt das Linksinverse} & (x, y) \mapsto x. \end{array}$$

Es folgt aus Lemma 22.7, dass $f^*: H^1(S^1 \times S^1) \rightarrow H^1(S^1)$ ein Epimorphismus ist. Satz 22.3 besagt, dass $H^1(S^1) \neq 0$, also ist auch $H^1(S^1 \times S^1) \neq 0$.

Beweis. Es sei $f: M \rightarrow N$ eine glatte Abbildung zwischen Mannigfaltigkeiten und es sei $g: N \rightarrow M$ eine glatte Abbildung mit $g \circ f = \text{id}_M$. Dann ist

$$\begin{array}{ccccccc} \text{id}_{H^1(M)} & = & \text{id}_M^* & = & (g \circ f)^* & = & f^* \circ g^* \\ & & \uparrow & & \uparrow & & \\ & & \text{Satz 22.5} & & \text{Satz 22.5} & & \end{array}$$

Nachdem $f^* \circ g^* = \text{id}$ insbesondere ein Epimorphismus ist, muss auch f^* ein Epimorphismus sein. $\square$

22.5. Das Dachprodukt auf de Rham-Kohomologie. Mithilfe von Satz 10.2 (4) kann man problemlos folgenden Satz beweisen.

Satz 22.8. *Es sei M eine Mannigfaltigkeit und es seien $\omega \in \Omega_k(M)$, $\sigma \in \Omega_l(M)$. Es ist*

$$d(\omega \wedge \sigma) = d\omega \wedge \sigma + (-1)^k \omega \wedge d\sigma.$$

Mithilfe von Satz 9.9 und Satz 22.8 kann man nun leicht folgendes Korollar beweisen.

Korollar 22.9. *Es sei M eine Mannigfaltigkeit. Dann ist*

$$\begin{aligned} \wedge: H^k(M) \times H^l(M) &\rightarrow H^{k+l}(M) \\ ([\omega], [\sigma]) &\mapsto [\omega] \wedge [\sigma] := [\omega \wedge \sigma] \end{aligned}$$

eine wohldefinierte Abbildung. Diese hat zudem folgende Eigenschaften:

(1) *die folgende Abbildung ist bilinear:*

$$\begin{aligned} \wedge: H^k(M) \times H^l(M) &\rightarrow H^{k+l}(M) \\ (a, b) &\mapsto a \wedge b \end{aligned}$$

(2) *für alle $a \in H^k(M)$, $b \in H^l(M)$ und $c \in H^m(M)$ ist*

$$a \wedge (b \wedge c) = (a \wedge b) \wedge c.$$

(3) *für alle $a \in H^k(M)$ und $b \in H^l(M)$ ist*

$$b \wedge a = (-1)^{kl} \cdot b \wedge a.$$

Beispiel. Wir betrachten die beiden 1-Formen $\overline{dx}$ und $\overline{dy}$ auf dem Torus $T = \mathbb{R}^2/\mathbb{Z}^2$, welche wir auf Seite 249 eingeführt hatten. Wir hatten auf Seite 249 bewiesen, dass $\overline{dx}$ und $\overline{dy}$ geschlossen sind. Hierbei ist $[\overline{dx}] \wedge [\overline{dy}] = [\overline{dx} \wedge \overline{dy}]$. Wir hatten auf Seite 249 gezeigt, dass

$$\int_{\mathbb{R}^2/\mathbb{Z}^2} \overline{dx} \wedge \overline{dy} = 1.$$

Es folgt nun aus dem Beweis von Seite 270, dass $[\overline{dx} \wedge \overline{dy}] \neq 0 \in H^2(T)$. Zudem folgt aus Korollar 22.9 (1), dass $[\overline{dx}] \neq 0 \in H^1(T)$ und $[\overline{dy}] \neq 0 \in H^1(T)$. Außerdem folgt aus

Korollar 22.9 (3), dass $[\overline{dx}]$ und $[\overline{dy}]$ linear unabhängig in $H^1(T)$ sind.²⁶⁷ Wir haben also gezeigt, dass $\dim H^1(T) \geq 2$.²⁶⁸

Für eine n -dimensionale Mannigfaltigkeit M setzen wir nun²⁶⁹

$$H^*(M) = \bigoplus_{i=0}^n H^i(M).$$

Für $a_i, b_i \in H^i(M)$, $i = 0, \dots, n$ definieren wir

$$\left(\sum_{i=0}^n a_i \right) \wedge \left(\sum_{i=0}^n b_i \right) := \sum_{i=0}^n \sum_{j=0}^n \underbrace{a_i \wedge b_j}_{\in H^{i+j}(M)}.$$

Es folgt aus Korollar 22.9, dass mit dieser Multiplikation $H^*(M)$ ein (nicht notwendigerweise kommutativer) Ring ist.²⁷⁰

Auf Seite 9.5 hatten wir schon erwähnt, dass für eine lineare Abbildung $f: U \rightarrow V$ zwischen Vektorräumen und für alle $\alpha \in \wedge^k V^*$ und $\beta \in \wedge^l V^*$ gilt, dass

$$f^*(\alpha \wedge \beta) = f^*\alpha \wedge f^*\beta.$$

Mithilfe dieser Bemerkung kann man nun, ganz analog zu Satz 22.4, folgenden Satz beweisen.

²⁶⁷In der Tat, denn wenn $[\overline{dx}]$ und $[\overline{dy}]$ linear abhängig wären, dann gäbe es insbesondere ein $\lambda \in \mathbb{R} \setminus \{0\}$ mit $[\overline{dx}] = \lambda[\overline{dy}]$. Daraus wiederum folgt, dass

$$\begin{aligned} [\overline{dx}] \wedge [\overline{dy}] &= \lambda[\overline{dy}] \wedge [\overline{dy}] = -\lambda[\overline{dy}] \wedge [\overline{dy}] = -[\overline{dx}] \wedge [\overline{dy}], \\ &\quad \uparrow \\ &\quad \text{Korollar 22.9 (3)} \end{aligned}$$

d.h. es ist $[\overline{dx}] \wedge [\overline{dy}] = 0$.

²⁶⁸In der Tat gilt $H^0(T) \cong \mathbb{R}$, $H^1(T) \cong \mathbb{R}^2$ und $H^2(T) \cong \mathbb{R}$. Aber wir können das mit den bisher erarbeiteten Methoden nicht beweisen.

²⁶⁹Für Vektorräume $V_0, \dots, V_n$ bezeichnen wir hierbei mit

$$\bigoplus_{i=0}^n V_i := V_0 \oplus \dots \oplus V_n := \{(v_1, \dots, v_n) \mid v_i \in V_i\}$$

die *direkte Summe* von $V_0, \dots, V_n$. Die direkte Summe besitzt eine offensichtliche Vektorraumstruktur. Für jedes $i \in \{0, \dots, n\}$ ist

$$V_i \rightarrow \bigoplus_{i=0}^n V_i = V_0 \oplus \dots \oplus V_n \quad w \mapsto (0, \dots, 0, \underbrace{w}_{\substack{\uparrow \\ i\text{-te Koordinate}}, 0, \dots)$$

ein Monomorphismus von Vektorräumen und wir identifizieren V_i mit dem Bild dieser Abbildung, d.h. wir fassen V_i als Untervektorraum von $\bigoplus_{i=0}^n V_i = V_0 \oplus \dots \oplus V_n$ auf. Für $v_i \in V_i$, $i = 0, \dots, n$ schreiben wir also insbesondere

$$v_0 + \dots + v_n = (v_0, \dots, v_n).$$

²⁷⁰Besitzt der Ring $H^*(M)$ ein multiplikativ neutrales Element?

Satz 22.10. *Die Abbildungen*

$$M \mapsto H^*(M)$$

und

$$(f: M \rightarrow N) \mapsto (f^*: H^*(N) \rightarrow H^*(M))$$

definieren einen kontravarianten Funktor von der Kategorie der Mannigfaltigkeiten zu der Kategorie der Ringe.

22.6. Die Sphäre und der Torus sind nicht diffeomorph. Wir wollen nun zum Abschluß der Vorlesung beweisen, dass die Sphäre und der Torus nicht diffeomorph sind. Wir werden dabei folgenden Satz verwenden.

Satz 22.11. *Es ist $H^1(\mathbb{R}^2) = 0$.*

Bemerkung.

- (1) Mit etwas mehr Aufwand kann man auch beweisen, dass für alle $n \geq 0$ und für alle $k \geq 1$ gilt, dass $H^k(\mathbb{R}^n) = 0$.
- (2) Eine Teilmenge $U \subset \mathbb{R}^n$ heißt *sternförmig*, wenn es ein $x \in U$ gibt, so dass für alle $y \in U$ die Verbindungsstrecke $(1-s)x + sy$, $s \in [0, 1]$ von x nach y ebenfalls ganz in U liegt. Beispielsweise sind konvexe Teilmengen von $\mathbb{R}^n$ sternförmig. Insbesondere ist der ganze $\mathbb{R}^n$ sternförmig. In Forster: Analysis III wird gezeigt, dass für jede sternförmige Teilmenge U von $\mathbb{R}^n$ die Kohomologiegruppen $H^k(U)$ für $k \geq 1$ verschwinden. Die Aussage wird oft als das Poincaré-Lemma bezeichnet.

Beweis. Wir müssen also zeigen, dass jede geschlossene 1-Form auf $\mathbb{R}^2$ auch schon exakt ist. Es sei also

$$\omega = f(x, y) dx + g(x, y) dy$$

eine geschlossene 1-Form. Dies bedeutet, dass

$$\begin{aligned} 0 = d\omega &= d(f \cdot dx + g \cdot dy) = df \wedge dx + dg \wedge dy \\ &= \left(\frac{\partial f}{\partial x} dx + \frac{\partial f}{\partial y} dy \right) \wedge dx + \left(\frac{\partial g}{\partial x} dx + \frac{\partial g}{\partial y} dy \right) \wedge dy = \left(-\frac{\partial f}{\partial y} + \frac{\partial g}{\partial x} \right) dx \wedge dy. \end{aligned}$$

$\uparrow$ Lemma 8.1 $\uparrow$ denn $dy \wedge dx = -dx \wedge dy$
und $dx \wedge dx = dy \wedge dy = 0$

Es folgt also, dass

$$\frac{\partial f}{\partial y} = \frac{\partial g}{\partial x}.$$

Wir hatten schon in Übungsblatt 5 gesehen, dass dies bedeutet, dass ω exakt ist.

Der Vollständigkeit halber führen wir das Argument noch einmal aus. Es genügt nun eine Funktion c zu finden mit

$$\frac{\partial c}{\partial x} = f(x, y) \quad \text{und} \quad \frac{\partial c}{\partial y} = g(x, y),$$

denn nach Lemma 8.1 ist dann $dc = \omega$. Es liegt nahe, als erstes eine Stammfunktion von f bezüglich x zu betrachten. Für $(x, y) \in U$ definieren wir nun

$$a(x, y) := \int_{s=0}^{s=x} f(s, y) ds.$$

Dann gilt

$$\frac{\partial a}{\partial x} = \frac{\partial}{\partial x} \int_{s=0}^{s=x} f(s, y) ds = f(x, y).$$

$\uparrow$
 HDI

Zudem gilt

$$\frac{\partial a}{\partial y} = \frac{\partial}{\partial y} \int_{s=0}^{s=x} f(s, y) ds = \int_{s=0}^{s=x} \frac{\partial}{\partial y} f(s, y) ds = \int_{s=0}^{s=x} \frac{\partial}{\partial s} g(s, y) ds = g(x, y) - g(0, y).$$

$\uparrow$ $\uparrow$
 Satz 3.10 angewandt denn $\frac{\partial}{\partial y} f(s, y) = \frac{\partial}{\partial s} g(s, y)$ auf f
 auf das Rechteck $[0, x] \times [0, y]$

Wir erhalten also die gewünschte partielle Ableitung bezüglich x . Wir erhalten fast die gewünschte partielle Ableitung bezüglich y , aber im Moment erhalten wir noch den Extraterm $-g(0, y)$. Wir können nun jedoch noch eine geeignet gewählte Funktion in y hinzuzufügen. In der Tat kann man nun leicht nachprüfen, dass

$$c(x, y) := a(x, y) + \int_{s=0}^{s=y} g(0, y) ds$$

die gewünschte Eigenschaft besitzt. □

Wir beschließen das Kapitel mit folgendem Satz.

Satz 22.12. Die Sphäre S^2 ist nicht diffeomorph zum Torus $T = S^1 \times S^1 = \mathbb{R}^2/\mathbb{Z}^2$.

Beweis. Nehmen wir also an, es gibt einen Diffeomorphismus $\Phi: S^2 \rightarrow T := S^1 \times S^1$. Wir bezeichnen mit $P = (0, 0, 1)$ den “Nordpol” von S^2 und wir schreiben $\Phi(P) = (z, w)$. Die Einschränkung von Φ auf $S^2 \setminus P \rightarrow T \setminus \Phi(P)$ ist ebenfalls ein Diffeomorphismus von Mannigfaltigkeiten.

Korollar 22.6 besagt, dass diffeomorphe Mannigfaltigkeiten isomorphe Kohomologiegruppen besitzen. Es genügt nun also folgende Behauptung zu beweisen.

Behauptung. Es ist $H^1(S^2 \setminus P) = 0$ und es ist $H^1(T \setminus \Phi(P)) \neq 0$.

Wie auf Seite 20 betrachten wir nun die stereographische Projektion

$$\begin{aligned} \Psi: S^2 \setminus P &\rightarrow \mathbb{R}^2 \\ (x, y, z) &\mapsto \left(\frac{x}{1-z}, \frac{y}{1-z} \right). \end{aligned}$$

Dies gibt uns einen Diffeomorphismus von der Mannigfaltigkeit $S^2 \setminus P$ zu $\mathbb{R}^2$. Es folgt nun aus Satz 22.11 und aus Korollar 22.6, dass $H^1(S^2 \setminus P) = 0$.

Der Beweis, dass $H^1(T \setminus \Phi(P)) \neq 0$ ist fast identisch zu dem Beweis auf Seite 273, dass $H^1(T) \neq 0$. Wir wählen ein $v \neq w$. Dann ist $S^1 \times \{v\} \subset S^1 \times S^1 \setminus (z, w) = T \setminus \Phi(P)$. Die Abbildung

$$\begin{array}{ccc} f: S^1 & \rightarrow & S^1 \times S^1 \setminus \Phi(P) \\ x & \mapsto & (x, v) \end{array} \quad \text{besitzt das Linksinverse} \quad \begin{array}{ccc} g: S^1 \times S^1 \setminus \Phi(P) & \rightarrow & S^1 \\ (x, y) & \mapsto & x. \end{array}$$

Es folgt aus Lemma 22.7, dass $f^*: H^1(S^1 \times S^1 \setminus \Phi(P)) \rightarrow H^1(S^1)$ ein Epimorphismus ist. Satz 22.3 besagt, dass $H^1(S^1) \neq 0$. Also ist auch $H^1(S^1 \times S^1 \setminus \Phi(P)) \neq 0$. $\square$

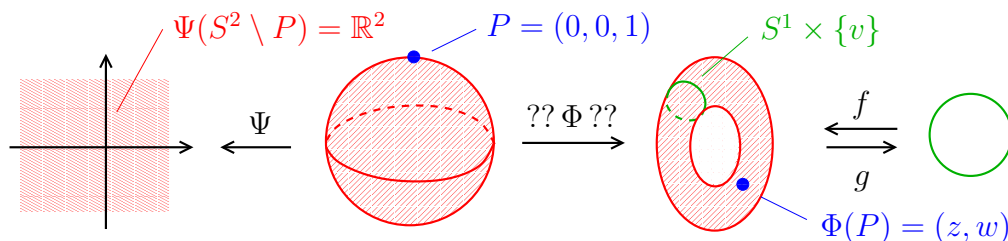


ABBILDUNG 142. Skizze zum Beweis von Satz 22.12.

Bemerkung. Der vorherige Satz besagt also, dass die Sphäre S^2 *nicht diffeomorph* zum Torus $T = S^1 \times S^1 = \mathbb{R}^2/\mathbb{Z}^2$ ist. A priori könnte es allerdings sein, dass die Sphäre S^2 *homöomorph* zum Torus T ist. In der Tat gibt es Mannigfaltigkeiten, welche homöomorph aber nicht diffeomorph sind. Wir werden in der Vorlesung Algebraische Topologie I sehen, dass dies in diesem Fall jedoch nicht der Fall ist.

23. AUSBLICK

Im letzten Kapitel haben wir gezeigt, dass die 2-dimensionale Sphäre und der Torus nicht diffeomorph sind. Wir haben dies dadurch bewerkstelligt, dass wir einen Funktor von der Kategorie der Mannigfaltigkeiten zu der Kategorie der Vektorräume eingeführt hatten. Dieser Funktor hat es uns dann ermöglicht Mannigfaltigkeiten mithilfe von der (deutlich einfacheren Theorie) der Vektorräume zu studieren.

In den Vorlesungen Algebraische Topologie I und II werden wir ganz ähnlich verfahren. Wir werden Funktoren von der Kategorie der topologischen Räume in die Kategorie der Gruppen und in die Kategorie der Vektorräume einführen. Diese ermöglichen es uns topologische Probleme mithilfe von algebraischen Methoden zu studieren.

Insbesondere werden wir in der Algebraische Topologie II die “singuläre Kohomologiegruppen” für beliebige topologische Räume einführen. Für Mannigfaltigkeiten sind diese isomorph zu der de Rham-Kohomologie. Wir werden zudem viele Eigenschaften der Kohomologiegruppen kennenlernen, welche es ermöglichen diese für viele Mannigfaltigkeiten zu bestimmen.