

Topologie von Graphen im Kontext der Geoinformatik

Studienarbeit

am Institut für Photogrammetrie und Fernerkundung
des Karlsruher Institut für Technologie (KIT)

cand. geod.
Anna Krimmelbein

Aufgabenstellung und Betreuung:
Dr. rer. nat. Patrick Erik Bradley

Karlsruhe, den 31. März 2010

Erklärung

Ich versichere hiermit, die vorliegende Studienarbeit bis auf die dem Aufgabensteller bekannten Hilfsmittel selbstständig angefertigt, alle benutzten Hilfsmittel vollständig und genau angegeben und alles kenntlich gemacht zu haben, was aus Arbeiten Anderer unverändert oder mit Änderungen entnommen wurde.

Karlsruhe, den 31. März 2010

Unterschrift

Lizenz

Dieses Werk ist unter einem Creative Commons Namensnennung-Keine kommerzielle Nutzung-Keine Bearbeitung 3.0 Deutschland Lizenzvertrag lizenziert. Um die Lizenz anzusehen, gehen Sie bitte zu <http://creativecommons.org/licenses/by-nc-nd/3.0/de/> oder schicken Sie einen Brief an Creative Commons, 171 Second Street, Suite 300, San Francisco, California 94105, USA.

Inhaltsverzeichnis

1	Einleitung	2
2	Grundlagen	3
2.1	Begriff Topologischer Raum	3
2.2	Weitere Grundbegriffe	3
2.3	Metrische Räume	6
2.4	Eigenschaften von Topologien	8
2.5	Stetige Abbildungen	14
2.6	Graphen	16
3	Anwendungen	25
3.1	Das Königsberger Brückenproblem	26
3.2	Lösung des Brückenproblems mittels des Topologischen Datentyps	29
3.3	Bestimmung der Bettizahlen mittels des Topologischen Datentyps	32
4	Fazit und Ausblick	34
	Literatur	35

1 Einleitung

Der Begriff “Topologie” taucht im Zusammenhang mit der Geoinformatik immer wieder auf, meistens ohne dass geklärt wird, was eigentlich dahintersteckt. Als Student der “Geodäsie und Geoinformatik” weiß man über Topologie meist nur, dass es was mit den Beziehungen zwischen Objekten zu tun hat und dass Graphen verwendet werden können, um diese Beziehungen darzustellen, aber von den mathematischen Grundlagen, die dem zugrundeliegen, weiß man nichts oder sehr wenig.

Es liegt also nahe sich mal mit dem mathematischen Begriff der “Topologie” auseinanderzusetzen und sich im Zuge dessen darüber klar zu werden, wie dieser Begriff mit Graphen zusammenhängt, was also die Topologie von Graphen ist. Danach ist der nächste Schritt sich den Zusammenhang dieser, zu Beginn doch sehr theoretisch wirkenden, mathematischen Grundlagen, mit der Geoinformatik klarzumachen. Dieser Zusammenhang liegt in den Anwendungen von Graphen begründet, welche in der Geoinformatik hauptsächlich zur Analyse von geografischen Sachverhalten benutzt werden, welche man von der genauen Lage von Objekten auf die Lagebeziehung der Objekte zueinander reduzieren kann.

Die Arbeit beginnt also mit einem Überblick über die mathematischen Grundlagen, die hinter dem Begriff “Topologie” stehen, soweit diese Grundlagen zur Betrachtung der Topologie von Graphen nötig sind. Danach werden Graphen vorgestellt und der Bezug zu den vorher beschriebenen Grundlagen wird hergestellt. Dann wird auf einige Anwendungen aus dem Bereich der Geoinformatik eingegangen. Unter Anderem wird das sogenannte “Königsberger Brückenproblem” behandelt, welches Euler zu Untersuchungen zum Thema Graphen veranlasste (Originalartikel siehe Quelle [5]), was man als Geburtsstunde des mathematischen Fachgebiets “Topologie” auffassen kann.

Die Grundlage des Abschnittes über die Grundlagen bilden hauptsächlich Notizen aus einem “Topologie-Crash-Kurs” mit dem Betreuer Dr. rer. nat. Patrick Erik Bradley, welcher nötig war, weil wie erwähnt noch so gut wie gar keine Vorkenntnisse vorhanden waren. Weitere Informationen wurden aus den, im Literaturverzeichnis genannten, Quellen [1], [2], [3] und aus dem Internet (Quellen [6], [7] und [8]) entnommen. Die Idee sich in diesem Zusammenhang mit dem “Königsberger Brückenproblem” zu beschäftigen stammt aus der Quelle [4]. Die verwendeten Bilder wurden zum Großteil selbst mit “Paint” gezeichnet, oder von der Internetseite aus Quelle [10] heruntergeladen.

2 Grundlagen

2.1 Begriff Topologischer Raum

Definition 2.1 (topologischer Raum) Ein *topologischer Raum* ist ein Paar $(\mathcal{X}, \mathcal{T})$, bestehend aus einer Menge $\mathcal{X}$ und einer Menge $\mathcal{T}$ ($\mathcal{T} \subseteq \mathcal{P}(\mathcal{X}) := \{\mathcal{A} \mid \mathcal{A} \subseteq \mathcal{X}\}$) von Teilmengen von $\mathcal{X}$, für welche die folgenden Axiome gelten:

1. Axiom: Die Menge $\mathcal{X}$ und die leere Menge $\emptyset$ seien in $\mathcal{T}$ enthalten. ($\mathcal{X}, \emptyset \in \mathcal{T}$)
2. Axiom: Die Vereinigung beliebig vieler in $\mathcal{T}$ enthaltener Mengen sei Teilmenge von $\mathcal{T}$. ($\mathcal{A} \subseteq \mathcal{T} \Rightarrow \bigcup_{A \in \mathcal{A}} A \in \mathcal{T}$)
3. Axiom: Der Durchschnitt zweier Mengen aus $\mathcal{T}$ sei wieder in $\mathcal{T}$ enthalten. ($A, B \in \mathcal{T} \Rightarrow A \cap B \in \mathcal{T}$)

Sind diese drei Axiome erfüllt, so wird das Paar $(\mathcal{X}, \mathcal{T})$ als *topologischer Raum* und $\mathcal{T}$ als *Topologie* bezeichnet. Eine Menge $\mathcal{O}$ aus $\mathcal{T}$ ($\mathcal{O} \in \mathcal{T}$) heißt *offene Menge*.

Benutzt man den Begriff der offenen Menge, lassen sich die drei Axiome aus der Definition 2.1 für den topologischen Raum folgendermaßen ausdrücken:

1. Die Menge $\mathcal{X}$ und die leere Menge $\emptyset$ seien offene Mengen.
2. Die Vereinigung beliebig vieler offener Mengen sei wieder offen.
3. Der Durchschnitt zweier offener Mengen sei wieder offen.

Ein Beispiel für einen topologischen Raum ist das Paar $(\mathbb{R}, \mathcal{T})$ mit $\mathcal{T} = \{\text{Vereinigungen "offener" Intervalle } (a, b)\}$. Weitere Beispiele für topologische Räume stellen metrische Räume dar (siehe Abschnitt 2.3).

2.2 Weitere Grundbegriffe

Definition 2.2 (abgeschlossene Menge) Sei das Paar $(\mathcal{X}, \mathcal{T})$ ein topologischer Raum. Dann heißt eine Teilmenge $\mathcal{A}$ von $\mathcal{X}$ ($\mathcal{A} \subseteq \mathcal{X}$) *abgeschlossen*, falls das Komplement $(\mathcal{X} \setminus \mathcal{A})$ von $\mathcal{A}$ offen ist.

Mit Hilfe der Definition 2.2 für abgeschlossene Mengen lassen sich die drei Axiome aus Definition 2.1 folgendermaßen umformulieren:

1. Die leere Menge $\emptyset$ und die Menge $\mathcal{X}$ sind abgeschlossen.
2. Der Durchschnitt beliebig vieler abgeschlossener Mengen ist wieder abgeschlossen.
3. Die Vereinigung zweier abgeschlossener Mengen ist abgeschlossen.

Man kann sich veranschaulichen, dass diese Formulierung richtig ist, indem man die Axiome für offene Mengen betrachtet und überlegt, was dann entsprechend für die Komplemente, also für abgeschlossene Mengen gelten muss.

Das erste Axiom bleibt unverändert, da das Komplement von $\mathcal{X}$ die leere Menge $\emptyset$ und das Komplement der leeren Menge $\emptyset$ wiederum $\mathcal{X}$ ist. Die Menge $\mathcal{X}$ und die leere Menge $\emptyset$ sind also beide sowohl offen als auch abgeschlossen.

Dass beim zweiten Axiom aus der *Vereinigung* beliebig vieler *offener* Mengen, der *Durchschnitt* beliebig vieler *abgeschlossener* Mengen werden muss, kann man sich am Beispiel der Vereinigung $\mathcal{A} \cup \mathcal{B}$ zweier offener Mengen $\mathcal{A}$ und $\mathcal{B}$ verdeutlichen (siehe hierzu Abb. 1). Die Menge $\mathcal{A} \cup \mathcal{B}$ ist wieder offen, also ist ihr Komplement $\mathcal{X} \setminus (\mathcal{A} \cup \mathcal{B})$ abgeschlossen. Um dieses Komplement aus den abgeschlossenen Mengen $\mathcal{X} \setminus \mathcal{A}$ und $\mathcal{X} \setminus \mathcal{B}$ zu erhalten, muss man den Durchschnitt dieser beiden Mengen bilden, also gilt $\mathcal{X} \setminus (\mathcal{A} \cup \mathcal{B}) = (\mathcal{X} \setminus \mathcal{A}) \cap (\mathcal{X} \setminus \mathcal{B})$. Aus der Vereinigung von offenen Mengen ist also der Durchschnitt abgeschlossener Mengen geworden. Um die Gültigkeit des zweiten Axioms allgemein zu zeigen, müsste der Beweis allerdings für beliebig viele Mengen geführt werden. Er wäre also zu zeigen, dass $\mathcal{X} \setminus \bigcup_{A \in \mathcal{A}} \mathcal{A} = \bigcap_{A \in \mathcal{A}} (\mathcal{X} \setminus \mathcal{A})$.

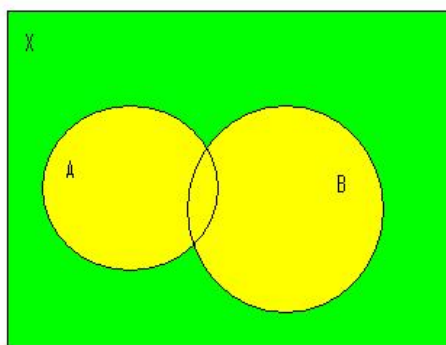


Abbildung 1: gelb: $\mathcal{A} \cup \mathcal{B}$, grün: $\mathcal{X} \setminus (\mathcal{A} \cup \mathcal{B})$

Betrachtet man nun den Durchschnitt $\mathcal{A} \cap \mathcal{B}$ zweier Mengen $\mathcal{A}$ und $\mathcal{B}$ wird schnell klar, dass beim dritten Axiom aus dem *Durchschnitt* zweier *offener*

Mengen, die *Vereinigung* zweier *abgeschlossener* Mengen werden muss (siehe Abb. 2). Bildet man die Vereinigung der beiden abgeschlossenen Mengen $\mathcal{X} \setminus \mathcal{A}$ und $\mathcal{X} \setminus \mathcal{B}$ erhält man alles außer dem Durchschnitt von $\mathcal{A}$ und $\mathcal{B}$, also genau $\mathcal{X} \setminus (\mathcal{A} \cap \mathcal{B})$. Die Menge $\mathcal{X} \setminus (\mathcal{A} \cap \mathcal{B})$ ist abgeschlossen, da ihr Komplement $\mathcal{A} \cap \mathcal{B}$ offen ist.

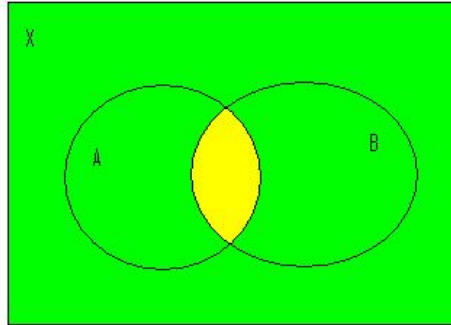


Abbildung 2: gelb: $\mathcal{A} \cap \mathcal{B}$, grün: $\mathcal{X} \setminus (\mathcal{A} \cap \mathcal{B})$

Definition 2.3 (innerer Punkt) Ein Element x einer Menge $\mathcal{X}$ ($x \in \mathcal{X}$) heißt *innerer Punkt* einer Teilmenge $\mathcal{A}$ von $\mathcal{X}$ ($\mathcal{A} \subseteq \mathcal{X}$), falls eine offene Teilmenge $\mathcal{U}$ von $\mathcal{X}$ ($\mathcal{U} \subseteq \mathcal{X}$) existiert, sodass x Element von $\mathcal{U}$ ($x \in \mathcal{U}$) und $\mathcal{U}$ Teilmenge von $\mathcal{A}$ ($\mathcal{U} \subseteq \mathcal{A}$) ist. Die Menge aller inneren Punkte von $\mathcal{A}$ heißt *Inneres* $\mathcal{A}^\circ$ von $\mathcal{A}$.

Definition 2.4 (Randpunkt) Ein Element x einer Menge $\mathcal{X}$ ($x \in \mathcal{X}$) heißt *Randpunkt* einer Teilmenge $\mathcal{A}$ von $\mathcal{X}$ ($\mathcal{A} \subseteq \mathcal{X}$), wenn jede offene Teilmenge von $\mathcal{X}$, die x enthält, sowohl Punkte aus $\mathcal{A}$, als auch Punkte aus dem Komplement von $\mathcal{A}$ enthält. Die Menge aller Randpunkte von $\mathcal{A}$ heißt *Rand* $\partial\mathcal{A}$ von $\mathcal{A}$.

Die Definitionen 2.3 und 2.4 werden in Abbildung 3 nochmal grafisch verdeutlicht.

Definition 2.5 (Abschluss) Als *Abschluss* $\bar{\mathcal{A}}$ von $\mathcal{A}$ bezeichnet man die (disjunkte) Vereinigung vom Inneren $\mathcal{A}^\circ$ von $\mathcal{A}$ und dem Rand $\partial\mathcal{A}$ von $\mathcal{A}$. ($\bar{\mathcal{A}} = \mathcal{A}^\circ \cup \partial\mathcal{A}$)

Um im Umgang mit den Begriffen aus den Definitionen 2.2 bis 2.5 ein bisschen vertrauter zu werden, kann man sich mal kurz über die folgenden Schlussfolgerungen Gedanken machen:

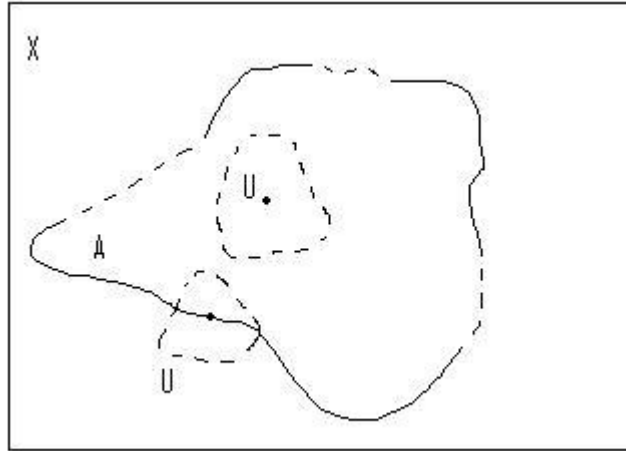


Abbildung 3: Innerer Punkt und Randpunkt

- Der Abschluss vom Abschluss einer Menge $\mathcal{A}$ ist gleich dem Abschluss von $\mathcal{A}$ ($\overline{\overline{\mathcal{A}}} = \overline{\mathcal{A}}$), weil der Rand vom Abschluss von $\mathcal{A}$ eine Teilmenge vom Abschluss von $\mathcal{A}$ ($\partial\overline{\mathcal{A}} \subseteq \overline{\mathcal{A}}$) ist und somit beim Bilden des Abschlusses vom Abschluss von $\mathcal{A}$ nichts mehr hinzukommt.
- Der Abschluss vom Inneren einer Menge $\mathcal{A}$ ist gleich dem Abschluss von $\mathcal{A}$ ($\overline{\mathcal{A}^\circ} = \overline{\mathcal{A}}$), weil der Rand vom Inneren von $\mathcal{A}$ gleich dem Rand von $\mathcal{A}$ ($\partial(\mathcal{A}^\circ) = \partial\mathcal{A}$) ist und somit beim Bilden des Abschlusses vom Inneren von $\mathcal{A}$ der "gleiche" Rand dazukommt.
- Der Rand einer Menge $\mathcal{A}$ ist gleich dem Abschluss von $\mathcal{A}$ ohne das Innere von $\mathcal{A}$ ($\partial\mathcal{A} = \overline{\mathcal{A}} \setminus \mathcal{A}^\circ$). Das ergibt sich durch Umstellen der Formel $\overline{\mathcal{A}} = \mathcal{A}^\circ \cup \partial\mathcal{A}$ für den Abschluss von $\mathcal{A}$.
- Ist das Innere einer Menge $\mathcal{A}$ Teilmenge einer Menge $\mathcal{B}$ und $\mathcal{B}$ wiederum Teilmenge vom Abschluss von $\mathcal{A}$, dann ist das Innere von $\mathcal{B}$ gleich dem Inneren von $\mathcal{A}$ und der Abschluss von $\mathcal{B}$ ist gleich dem Abschluss von $\mathcal{A}$ ($\mathcal{A}^\circ \subseteq \mathcal{B} \subseteq \overline{\mathcal{A}} \Rightarrow \mathcal{B}^\circ = \mathcal{A}^\circ$ und $\overline{\mathcal{B}} = \overline{\mathcal{A}}$). Da sich das Innere $\mathcal{A}^\circ$ und der Abschluss $\overline{\mathcal{A}}$ nur durch den Rand $\partial\mathcal{A}$ unterscheiden, kann sich $\mathcal{B}$ auch nur um Teile des Randes von $\mathcal{A}$ unterscheiden und der Rand $\partial\mathcal{B}$ muss gleich dem Rand $\partial\mathcal{A}$ sein.

2.3 Metrische Räume

Definition 2.6 (metrischer Raum) Ein *metrischer Raum* ist ein Paar (X, d) , bestehend aus einer Menge X und einer Distanzfunktion $d : X \times X \rightarrow \mathbb{R}_{\geq 0}$,

welche die folgenden Bedingungen erfüllt:

1. Ist die Distanz $d(x, y)$ gleich Null, dann folgt daraus, dass x und y gleich sind. Sind umgekehrt x und y gleich, dann muss die Distanz $d(x, y)$ gleich Null sein. ($d(x, y) = 0 \Leftrightarrow x = y$)
2. Die Distanz $d(x, y)$ sei gleich der Distanz $d(y, x)$ ($d(x, y) = d(y, x)$).
3. Die Distanz $d(x, z)$ sei kleiner oder gleich der Summe der Distanzen $d(x, y)$ und $d(y, z)$ ($d(x, z) \leq d(x, y) + d(y, z)$). Diese Bedingung wird als *Dreiecksungleichung* bezeichnet (siehe Abb. 4).

Sind diese drei Bedingungen erfüllt, dann wird (X, d) *metrischer Raum* genannt.

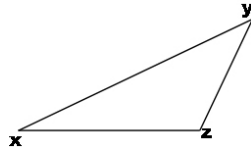


Abbildung 4: Dreiecksungleichung

Offene Mengen im Sinne der Definition topologischer Räume sind bei metrischen Räumen r -Kugeln ($B_r(a) := \{x \in X \mid d(x, a) < r\}$). Topologien auf einem metrischen Raum enthalten somit beliebig viele solcher Kugeln, deren Vereinigungen, Durchschnitte endlich vieler solcher Kugeln und diverse weitere Kombinationen von Vereinigungen und Durchschnitten, der so entstandenen offenen Mengen. Die kleinstmögliche derartige Topologie $\mathcal{T}_d$ auf X heißt die *von der Metrik d erzeugte* Topologie (mehr dazu in Abschnitt 2.4).

Beispiele für metrische Räume:

- $X = \mathbb{R}$, $d(x, y) = |x - y|$
- $X = \mathbb{R}^2$ (mit $x = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$, $y = \begin{pmatrix} y_1 \\ y_2 \end{pmatrix}$),
 $d_0(x, y) = \max\{|x_1 - y_1|, |x_2 - y_2|\}$ (siehe Abb. 5) oder
 $d_1(x, y) = |x_1 - y_1| + |x_2 - y_2| := \|x - y\|_1$ (siehe Abb. 6) oder
 $d_2(x, y) = \sqrt{|x_1 - y_1|^2 + |x_2 - y_2|^2} := \|x - y\|_2$ (siehe Abb. 7)
- $X = \mathbb{C}$, $d(z, w) = |z - w| = \sqrt{(z - w)(\bar{z} - \bar{w})}$ (siehe Abb. 7)
- $X = \mathbb{C}^2$, $d(z, w) = \sqrt{|z_1 - w_1|^2 + |z_2 - w_2|^2}$

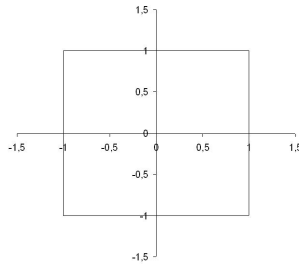


Abbildung 5:
 d_0 , 1-Kugel

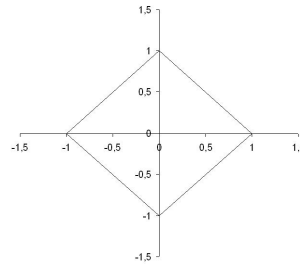


Abbildung 6:
 d_1 , 1-Kugel

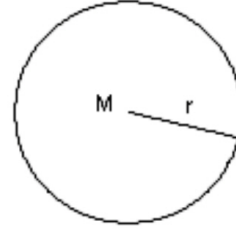


Abbildung 7:
 d_2 , r -Kugel

2.4 Eigenschaften von Topologien

Definition 2.7 (von $\mathcal{A}$ erzeugte Topologie) Seien eine Menge $\mathcal{X}$ und eine Teilmenge $\mathcal{A}$ der Potenzmenge $\mathcal{P}(\mathcal{X})$ (Menge aller Teilmengen von $\mathcal{X}$) ($\mathcal{A} \subseteq \mathcal{P}(\mathcal{X})$) gegeben. Dann sei die Topologie $\mathcal{T}(\mathcal{A})$ definiert als die kleinste Topologie $\mathcal{T}$ auf $\mathcal{X}$ mit der Eigenschaft, dass $\mathcal{A}$ Teilmenge der Topologie $\mathcal{T}$ ist ($\mathcal{A} \subseteq \mathcal{T}$). Diese Topologie $\mathcal{T}(\mathcal{A})$ heißt die *von $\mathcal{A}$ erzeugte Topologie*. Man erhält diese Topologie $\mathcal{T}(\mathcal{A})$, indem man alle Topologien $\mathcal{T}$, die $\mathcal{A}$ enthalten, schneidet ($\mathcal{T}(\mathcal{A}) = \bigcap_{\mathcal{T} \supseteq \mathcal{A}} \mathcal{T}$).

Um sich darüber klar zu werden, dass die obige Definition für die von $\mathcal{A}$ erzeugte Topologie sinnvoll ist, muss man sich als Erstes die Frage stellen, ob durch das Schneiden von Topologien auch wirklich wieder eine Topologie entsteht. Dass sich diese Frage tatsächlich mit ja beantworten lässt, kann recht gut am einfachsten Fall, nämlich dem Durchschnitt von nur zwei Topologien, verdeutlicht werden:

Gegeben seien also eine Menge $\mathcal{X}$ und zwei Topologien $\mathcal{T}_1$ und $\mathcal{T}_2$ auf $\mathcal{X}$. Es gilt zu klären, ob der Durchschnitt von $\mathcal{T}_1$ und $\mathcal{T}_2$ wieder eine Topologie ist ($\mathcal{T}_1 \cap \mathcal{T}_2 = \mathcal{T}$?). Dazu muss untersucht werden, ob die drei Axiome aus Definition 2.1 für $\mathcal{T} = \mathcal{T}_1 \cap \mathcal{T}_2$ erfüllt sind.

1. Nach dem ersten Axiom müssen die Menge $\mathcal{X}$ selbst und die leere Menge $\emptyset$ in $\mathcal{T}$ enthalten sein ($\mathcal{X}, \emptyset \in \mathcal{T}$). Dass dieses Axiom für $\mathcal{T}$ zutrifft, lässt sich recht leicht beweisen. Wenn $\mathcal{T}_1$ und $\mathcal{T}_2$ Topologien sind, dann erfüllen $\mathcal{T}_1$ und $\mathcal{T}_2$ auch das erste Axiom, d.h. sie enthalten jeweils die Menge $\mathcal{X}$ und die leere Menge $\emptyset$. Wenn die Menge $\mathcal{X}$ und die leere Menge $\emptyset$ in $\mathcal{T}_1$ und in $\mathcal{T}_2$ enthalten sind, dann sind $\mathcal{X}$ und $\emptyset$ auch im Durchschnitt von $\mathcal{T}_1$ und $\mathcal{T}_2$ enthalten, also in auch $\mathcal{T}$ ($\mathcal{X}, \emptyset \in \mathcal{T}_1; \mathcal{X}, \emptyset \in \mathcal{T}_2 \Rightarrow \mathcal{X}, \emptyset \in \mathcal{T}_1 \cap \mathcal{T}_2$). $\square$

2. Damit das zweite Axiom für $\mathcal{T}$ gilt, soll die Vereinigung beliebig vieler in $\mathcal{T}$ enthaltener Mengen Teilmenge von $\mathcal{T}$ sein ($\mathcal{A} \subseteq \mathcal{T} \Rightarrow \bigcup_{A \in \mathcal{A}} A \in \mathcal{T}$).

Dass auch dieses Axiom für $\mathcal{T}$ erfüllt ist, ist ebenfalls nicht schwierig zu beweisen. Hierfür betrachtet man eine Menge $\mathcal{A}$, welche eine Teilmenge der Potenzmenge $\mathcal{P}(\mathcal{X})$ ist. Außerdem sei die Menge $\mathcal{A}$ Teilmenge von $\mathcal{T}$, d.h. auch Teilmenge vom Durchschnitt von $\mathcal{T}_1$ und $\mathcal{T}_2$ und deshalb auch jeweils Teilmenge von $\mathcal{T}_1$ und von $\mathcal{T}_2$. Weil $\mathcal{T}_1$ und $\mathcal{T}_2$ Topologien sind, gilt das zweite Axiom für $\mathcal{T}_1$ und $\mathcal{T}_2$, also muss die Menge $\mathcal{A}$ die Vereinigung beliebig vieler jeweils in $\mathcal{T}_1$ und in $\mathcal{T}_2$ enthaltener Mengen A sein. Die Mengen A aus deren Vereinigung $\mathcal{A}$ entstanden ist, sind dann auch im Durchschnitt von $\mathcal{T}_1$ und $\mathcal{T}_2$, also in $\mathcal{T}$, enthalten. Damit ist das zweite Axiom für $\mathcal{T}$ erfüllt ($\mathcal{A} \subseteq \mathcal{T} \Rightarrow \mathcal{A} \subseteq \mathcal{T}_1$ und $\mathcal{A} \subseteq \mathcal{T}_2 \Rightarrow \bigcup_{A \in \mathcal{A}} A \in \mathcal{T}_1$ und $\bigcup_{A \in \mathcal{A}} A \in \mathcal{T}_2 \Rightarrow \bigcup_{A \in \mathcal{A}} A \in \mathcal{T}$). $\square$

3. Nach dem dritten Axiom muss der Durchschnitt zweier Mengen aus $\mathcal{T}$ wieder in $\mathcal{T}$ enthalten sein ($A, B \in \mathcal{T} \Rightarrow A \cap B \in \mathcal{T}$). Der Beweis dieses Axioms läuft ähnlich wie beim zweiten Axiom. In diesem Fall denkt man sich zwei Mengen A und B aus $\mathcal{T}$ ($A, B \in \mathcal{T}$), welche dann auch in $\mathcal{T}_1$ und in $\mathcal{T}_2$ enthalten sind. Weil das dritte Axiom für $\mathcal{T}_1$ und $\mathcal{T}_2$ als Topologien gilt, ist dann auch die Menge $A \cap B$ jeweils in $\mathcal{T}_1$ und in $\mathcal{T}_2$ enthalten. Wenn $A \cap B$ jeweils in $\mathcal{T}_1$ und in $\mathcal{T}_2$ enthalten ist, dann ist $A \cap B$ auch im Durchschnitt von $\mathcal{T}_1$ und $\mathcal{T}_2$, also in $\mathcal{T}$ enthalten. Somit ist das dritte Axiom für $\mathcal{T}$ erfüllt ($A, B \in \mathcal{T} \Rightarrow A, B \in \mathcal{T}_1$ und $A, B \in \mathcal{T}_2 \Rightarrow A \cap B \in \mathcal{T}_1$ und $A \cap B \in \mathcal{T}_2 \Rightarrow A \cap B \in \mathcal{T}$). $\square$

Der Beweis für den Durchschnitt beliebig (auch unendlich) vieler Topologien funktioniert ähnlich, nur dass man statt $\mathcal{T}_1$ und $\mathcal{T}_2$ eben “alle diese” Topologien betrachtet. Also ist z.B. das erste Axiom erfüllt, weil die Menge $\mathcal{X}$ und die leere Menge $\emptyset$ in “allen diesen” Topologien und deshalb auch im Durchschnitt “aller dieser” Topologien, enthalten sind.

Den Durchschnitt von endlich vielen Topologien kann man sich als Aneinanderreihung von Durchschnitten von jeweils zwei Topologien vorstellen, d.h. man schneidet zuerst zwei Topologien und das Ergebnis schneidet man dann mit der dritten Topologie. Das neue Ergebnis schneidet man wiederum mit der vierten Topologie usw.

Damit wäre also geklärt, dass sich durch die Formel für die von $\mathcal{A}$ erzeugte Topologie $\mathcal{T}(\mathcal{A}) = \bigcap_{\mathcal{T} \supseteq \mathcal{A}} \mathcal{T}$ tatsächlich wieder eine Topologie ergibt.

Die Frage, die man sich jetzt noch stellen könnte, ist, ob das Ergebnis dieser Formel wirklich die kleinstmögliche Topologie ist, welche die Eigenschaft $\mathcal{A} \subseteq \mathcal{T}$ hat. Diese Frage lässt sich allerdings leicht beantworten. Man kann

sich anschaulich vorstellen, dass der Durchschnitt zweier Topologien weniger Elemente enthält als beide Topologien vorher enthalten haben, nämlich nur die Elemente, welche vorher in beiden Topologien enthalten waren. Durch das Schneiden zweier Topologien wird also eine kleinere Topologie erzeugt, als die beiden Topologien vorher. Da $\mathcal{T}(\mathcal{A}) = \bigcap_{\mathcal{T} \supseteq \mathcal{A}} \mathcal{T}$ in jeder Topologie, welche die Teilmenge $\mathcal{A}$ enthält, auch enthalten sein muss, erhält man durch das Schneiden aller Topologien, welche $\mathcal{A}$ enthalten, die von $\mathcal{A}$ erzeugte Topologie.

Definition 2.8 (triviale Topologie) Als *triviale Topologie* $\mathcal{T}$ wird diejenige Topologie $\mathcal{T}$ auf der Menge $\mathcal{X}$ bezeichnet, welche nur die Menge $\mathcal{X}$ selbst und die leere Menge $\emptyset$ enthält ($\mathcal{T} = \{\mathcal{X}, \emptyset\}$).

Definition 2.9 (diskrete Topologie) Diejenige Topologie $\mathcal{T}$ auf der Menge $\mathcal{X}$, welche die gesamte Potenzmenge $\mathcal{P}(\mathcal{X})$ der Menge $\mathcal{X}$ darstellt, heißt *diskrete Topologie* ($\mathcal{T} = \mathcal{P}(\mathcal{X})$).

Auch hier kann man sich, wie bei der von $\mathcal{A}$ erzeugten Topologie, fragen, ob die triviale Topologie und die diskrete Topologie wirklich Topologien sind. Hierzu muss man wieder klären, ob die drei Axiome aus Definition 2.1 erfüllt sind.

Das erste Axiom gilt definitionsgemäß für die diskrete und die triviale Topologie. Das zweite und das dritte Axiom sind für die triviale Topologie erfüllt, da das Ergebnis von beliebig vielen Vereinigungen und Durchschnitten der Menge $\mathcal{X}$ und der leeren Menge $\emptyset$ wieder, entweder die Menge $\mathcal{X}$ selbst oder die leere Menge $\emptyset$ ist. Weil die diskrete Topologie die gesamte Potenzmenge $\mathcal{P}(\mathcal{X})$ enthält, enthält sie auch beliebige Vereinigungen und Durchschnitte von Elementen der Potenzmenge $\mathcal{P}(\mathcal{X})$, welche selbst wieder Elemente der Potenzmenge $\mathcal{P}(\mathcal{X})$ sind, deshalb gelten auch hier das zweite und das dritte Axiom.

Zwischen der (minimalen) trivialen Topologie und der (maximalen) diskreten Topologie sind noch viele weitere Topologien denkbar, welche mehr Teilmengen der Potenzmenge $\mathcal{P}(\mathcal{X})$ enthalten als die triviale Topologie, aber weniger als die diskrete Topologie. Man kann sich diesen Sachverhalt als gerichteten azyklischen Graph vorstellen, der oben mit der trivialen Topologie beginnt, sich dann zu den dazwischenliegenden Topologien verzweigt und schließlich wieder zur diskreten Topologie zusammengeführt wird. Eine Topologie, die mehr Teilmengen enthält als eine andere Topologie, wird auch als *feinere* Topologie bezeichnet. Enthält umgekehrt eine Topologie weniger Teilmengen als eine andere Topologie heißt diese *gröber*. Die triviale Topologie ist also die *größte* Topologie und die diskrete Topologie ist die *feinste* Topologie. Die von $\mathcal{A}$ erzeugte Topologie ist die *größte* mögliche Topologie, welche die Teilmenge $\mathcal{A}$ enthält.

Definition 2.10 (Spurtopologie) Sei das Paar $(\mathcal{X}, \mathcal{T}_{\mathcal{X}})$ ein topologischer Raum und die Menge $\mathcal{A}$ eine Teilmenge von $\mathcal{X}$ ($\mathcal{A} \subseteq \mathcal{X}$), dann ist die *Spurtopologie* $\mathcal{T}_{\mathcal{X}}|_{\mathcal{A}}$ (oder kurz $\mathcal{T}_{\mathcal{A}}$) definiert als die Topologie, welche gebildet wird aus dem Durchschnitt der Elemente O aus $\mathcal{T}_{\mathcal{X}}$ und $\mathcal{A}$ ($\mathcal{T}_{\mathcal{X}}|_{\mathcal{A}} := \mathcal{T}_{\mathcal{A}} := \{O \cap \mathcal{A} \mid O \in \mathcal{T}_{\mathcal{X}}\}$). Das Paar $(\mathcal{A}, \mathcal{T}_{\mathcal{A}})$ heißt dann *Teilraum* des topologischen Raumes $(\mathcal{X}, \mathcal{T}_{\mathcal{X}})$ (siehe Abb. 8).

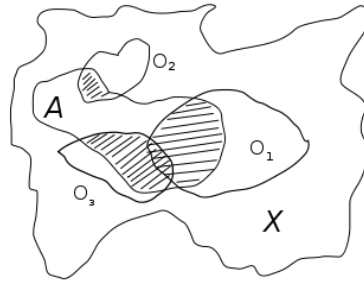


Abbildung 8: Spurtopologie

Beispiele:

- $\mathcal{X} = \emptyset$
 $\mathcal{T} = \{\emptyset\}$ ist in diesem Beispiel die einzige mögliche Topologie. Sie ist gleichzeitig die triviale und die diskrete Topologie.
- $\mathcal{X} = \{a\}$
Auch hier gibt es nur eine mögliche Topologie ($\mathcal{T} = \{\emptyset, \mathcal{X}\}$), welche zugleich die triviale und die diskrete Topologie darstellt.
- $\mathcal{X} = \{a, b\}, (a \neq b)$
Für dieses Beispiel lassen sich insgesamt vier Topologien bilden. Außer der trivialen Topologie und der diskreten Topologie gibt es noch zwei weitere mögliche Topologien, welche jeweils die Menge $\mathcal{X}$, die leere Menge $\emptyset$ und eines der beiden Elemente aus $\mathcal{X}$ enthalten.
 $\mathcal{T}_1 = \{\emptyset, \mathcal{X}\}$ (triviale Topologie)
 $\mathcal{T}_2 = \{\emptyset, \mathcal{X}, \{a\}\}$
 $\mathcal{T}_3 = \{\emptyset, \mathcal{X}, \{b\}\}$
 $\mathcal{T}_4 = \{\emptyset, \mathcal{X}, \{a\}, \{b\}\}$ (diskrete Topologie)
- $\mathcal{X} = \{a, b, c\}, (a \neq b \neq c)$
Dieses Beispiel, welches auf den ersten Blick auch noch sehr überschaubar wirkt, stellt sich bei genauerem Hinschauen doch schon als recht aufwändig heraus. Man kann nämlich insgesamt 29 Topologien finden.

Ein Weg alle diese möglichen Topologien zu bestimmen, ist die systematische Hinzunahme von Teilmengen, angefangen mit der trivialen Topologie $\mathcal{T}_1$. Die nächstfeineren Topologien beinhalten jeweils $\mathcal{X}$, $\emptyset$ und eines der drei Elemente aus $\mathcal{X}$. Durch Permutationen ergeben sich hier die Topologien $\mathcal{T}_2$ bis $\mathcal{T}_4$. Weitermachen kann man mit den Topologien, die jeweils gebildet werden aus $\mathcal{X}$, $\emptyset$ und einer zweielementigen Teilmenge von $\mathcal{X}$ ($\mathcal{T}_5$ bis $\mathcal{T}_7$). Als Nächstes kann man zu diesen jeweils ein Element hinzunehmen, welches in der zweielementigen Menge enthalten ist ($\mathcal{T}_8$ bis $\mathcal{T}_{13}$). Weitere Topologien sind solche Mengen, welche jeweils ein Element aus $\mathcal{X}$ enthalten und die zweielementige Teilmenge aus $\mathcal{X}$, welche nicht dieses Element enthält ($\mathcal{T}_{14}$ bis $\mathcal{T}_{16}$). Auch die Mengen, welche jeweils eine einelementige Teilmenge aus $\mathcal{X}$ enthalten und die beiden zweielementigen Teilmengen von $\mathcal{X}$, welche das Element aus der einelementigen Teilmenge enthalten, sind Topologien ($\mathcal{T}_{17}$ bis $\mathcal{T}_{19}$). Als Nächstes kann man Topologien bilden, welche jeweils zwei einelementige Teilmengen aus $\mathcal{X}$ und deren Vereinigung enthalten ($\mathcal{T}_{20}$ bis $\mathcal{T}_{22}$). Außerdem sind auch Mengen Topologien, welche jeweils zwei einelementige Teilmengen aus $\mathcal{X}$ und zwei zweielementige Teilmengen von $\mathcal{X}$ enthalten. Von den beiden zweielementigen Mengen muss hierbei eine diejenige sein, welche die beiden, in den einelementigen Teilmengen vorkommenden Elemente, enthält ($\mathcal{T}_{23}$ bis $\mathcal{T}_{28}$). Man könnte sich auch noch eine weitere Konstellation vorstellen, nämlich Mengen, welche zwei einelementige Teilmengen aus $\mathcal{X}$ und zwei zweielementige Teilmenge aus $\mathcal{X}$ enthalten, welche jeweils nur eines dieser einzelnen Elemente enthalten (z.B. $\mathcal{M} = \{\emptyset, \mathcal{X}, \{a\}, \{b\}, \{b, c\}, \{a, c\}\}$). Allerdings sind diese Mengen keine Topologien, weil der Durchschnitt der beiden zweielementigen Mengen, das dritte Element ergibt, welches nicht offen ist (z.B. $\{a, c\} \cap \{b, c\} = \{c\}$). Am Ende erhält man noch die diskrete Topologie ($\mathcal{T}_{29}$). Zur grafischen Veranschaulichung der hier beschriebenen möglichen Topologien dient die Abbildung 9, wobei hier die Permutationen fehlen.

$$\mathcal{T}_1 = \{\emptyset, \mathcal{X}\}$$

$$\mathcal{T}_2 = \{\emptyset, \mathcal{X}, \{a\}\}$$

$$\mathcal{T}_3 = \{\emptyset, \mathcal{X}, \{b\}\}$$

$$\mathcal{T}_4 = \{\emptyset, \mathcal{X}, \{c\}\}$$

$$\mathcal{T}_5 = \{\emptyset, \mathcal{X}, \{a, b\}\}$$

$$\mathcal{T}_6 = \{\emptyset, \mathcal{X}, \{a, c\}\}$$

$$\mathcal{T}_7 = \{\emptyset, \mathcal{X}, \{b, c\}\}$$

$$\mathcal{T}_8 = \{\emptyset, \mathcal{X}, \{a\}, \{a, b\}\}$$

$$\mathcal{T}_9 = \{\emptyset, \mathcal{X}, \{b\}, \{a, b\}\}$$

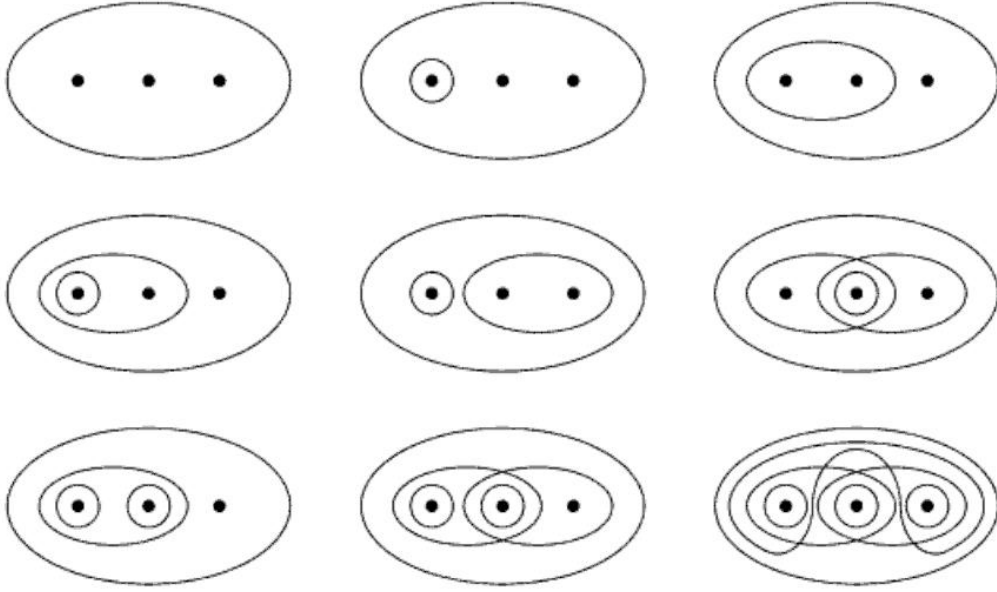


Abbildung 9: grafische Veranschaulichung der möglichen Topologien

$$\begin{aligned}
\mathcal{T}_{10} &= \{\emptyset, \mathcal{X}, \{a\}, \{a, c\}\} \\
\mathcal{T}_{11} &= \{\emptyset, \mathcal{X}, \{c\}, \{a, c\}\} \\
\mathcal{T}_{12} &= \{\emptyset, \mathcal{X}, \{b\}, \{b, c\}\} \\
\mathcal{T}_{13} &= \{\emptyset, \mathcal{X}, \{c\}, \{b, c\}\} \\
\mathcal{T}_{14} &= \{\emptyset, \mathcal{X}, \{a\}, \{b, c\}\} \\
\mathcal{T}_{15} &= \{\emptyset, \mathcal{X}, \{b\}, \{a, c\}\} \\
\mathcal{T}_{16} &= \{\emptyset, \mathcal{X}, \{c\}, \{a, b\}\} \\
\mathcal{T}_{17} &= \{\emptyset, \mathcal{X}, \{a\}, \{a, b\}, \{a, c\}\} \\
\mathcal{T}_{18} &= \{\emptyset, \mathcal{X}, \{b\}, \{a, b\}, \{b, c\}\} \\
\mathcal{T}_{19} &= \{\emptyset, \mathcal{X}, \{c\}, \{a, c\}, \{b, c\}\} \\
\mathcal{T}_{20} &= \{\emptyset, \mathcal{X}, \{a\}, \{b\}, \{a, b\}\} \\
\mathcal{T}_{21} &= \{\emptyset, \mathcal{X}, \{a\}, \{c\}, \{a, c\}\} \\
\mathcal{T}_{22} &= \{\emptyset, \mathcal{X}, \{b\}, \{c\}, \{b, c\}\} \\
\mathcal{T}_{23} &= \{\emptyset, \mathcal{X}, \{a\}, \{b\}, \{a, b\}, \{a, c\}\} \\
\mathcal{T}_{24} &= \{\emptyset, \mathcal{X}, \{a\}, \{b\}, \{a, b\}, \{b, c\}\} \\
\mathcal{T}_{25} &= \{\emptyset, \mathcal{X}, \{a\}, \{c\}, \{a, c\}, \{a, b\}\} \\
\mathcal{T}_{26} &= \{\emptyset, \mathcal{X}, \{a\}, \{c\}, \{a, c\}, \{b, c\}\} \\
\mathcal{T}_{27} &= \{\emptyset, \mathcal{X}, \{b\}, \{c\}, \{b, c\}, \{a, b\}\} \\
\mathcal{T}_{28} &= \{\emptyset, \mathcal{X}, \{b\}, \{c\}, \{b, c\}, \{a, c\}\} \\
\mathcal{T}_{29} &= \{\emptyset, \mathcal{X}, \{a\}, \{b\}, \{c\}, \{a, b\}, \{a, c\}, \{b, c\}\}
\end{aligned}$$

2.5 Stetige Abbildungen

Definition 2.11 (stetige Abbildung) Seien $(\mathcal{X}, \mathcal{T}_{\mathcal{X}})$ und $(\mathcal{Y}, \mathcal{T}_{\mathcal{Y}})$ topologische Räume. Dann ist eine *stetige Abbildung* $(\mathcal{X}, \mathcal{T}_{\mathcal{X}}) \rightarrow (\mathcal{Y}, \mathcal{T}_{\mathcal{Y}})$ definiert als eine Abbildung $f : \mathcal{X} \rightarrow \mathcal{Y}$, für die gilt, dass die Urbilder offener Mengen wieder offen sind (d.h. für jedes $\mathcal{U} \in \mathcal{T}_{\mathcal{Y}}$ existiert ein $f^{-1}(\mathcal{U}) \in \mathcal{T}_{\mathcal{X}}$).

Ein Beispiel für eine stetige Abbildung ist die Abbildung eines Teilraumes $(\mathcal{A}, \mathcal{T}_{\mathcal{X}}|_{\mathcal{A}})$ auf den gesamten topologischen Raum $(\mathcal{X}, \mathcal{T}_{\mathcal{X}})$, also $f : (\mathcal{A}, \mathcal{T}_{\mathcal{X}}|_{\mathcal{A}}) \rightarrow (\mathcal{X}, \mathcal{T}_{\mathcal{X}})$, wobei jedes Element auf sich selbst abgebildet wird ($a \rightarrow a$). Als Urbilder $f^{-1}(O)$, der in der Topologie $\mathcal{T}_{\mathcal{X}}$ offenen Elemente $O \in \mathcal{T}_{\mathcal{X}}$, ergeben sich die in der Spurtopologie $\mathcal{T}_{\mathcal{X}}|_{\mathcal{A}}$ ebenfalls offenen Mengen $f^{-1}(O) = O \cap \mathcal{A}$.

Sind die Abbildung $f : \mathcal{X} \rightarrow \mathcal{Y}$ und die Abbildung $g : \mathcal{Y} \rightarrow \mathcal{Z}$ beide stetig, dann ist auch $g \circ f : \mathcal{X} \rightarrow \mathcal{Z}$ stetig.

Man kann Definition 2.11 auch wieder für abgeschlossene Mengen umformulieren: Eine Abbildung $f : \mathcal{X} \rightarrow \mathcal{Y}$ ist genau dann stetig, wenn die Urbilder abgeschlossener Mengen wieder abgeschlossen sind. Der Beweis hierfür lässt sich in zwei Teile aufteilen, nämlich die zwei Richtungen “ $\Rightarrow$ ” und “ $\Leftarrow$ ”.

“ $\Rightarrow$ ”: Gegeben seien eine stetige Abbildung $f : \mathcal{X} \rightarrow \mathcal{Y}$ und eine abgeschlossene Teilmenge $\mathcal{A} \subseteq \mathcal{Y}$. Das Komplement $\mathcal{Y} \setminus \mathcal{A}$ muss dann offen sein. Da f stetig ist, muss somit das Urbild $f^{-1}(\mathcal{Y} \setminus \mathcal{A})$ auch offen sein. Das Urbild $f^{-1}(\mathcal{Y} \setminus \mathcal{A})$ lässt sich auch schreiben als das Urbild von $\mathcal{Y}$ ohne das Urbild von $\mathcal{A}$, also als $f^{-1}(\mathcal{Y}) \setminus f^{-1}(\mathcal{A})$. Das Urbild $f^{-1}(\mathcal{Y})$ ist wiederum gleich $\mathcal{X}$, $f^{-1}(\mathcal{Y} \setminus \mathcal{A})$ ist also gleich $\mathcal{X} \setminus f^{-1}(\mathcal{A})$. Weil $\mathcal{X} \setminus f^{-1}(\mathcal{A})$ offen ist, muss das Komplement davon $f^{-1}(\mathcal{A})$ abgeschlossen sein. $\square$

“ $\Leftarrow$ ”: Gegeben sei wieder eine stetige Abbildung $f : \mathcal{X} \rightarrow \mathcal{Y}$ und eine offene Teilmenge $\mathcal{O} \subseteq \mathcal{Y}$. Dann ist das Komplement $\mathcal{Y} \setminus \mathcal{O}$ abgeschlossen. Stimmt die Umformulierung der Definition 2.11 für abgeschlossene Mengen, dann muss das Urbild $f^{-1}(\mathcal{Y} \setminus \mathcal{O})$ wieder abgeschlossen sein. Das Urbild $f^{-1}(\mathcal{Y} \setminus \mathcal{O})$ lässt sich auch wieder als $\mathcal{X} \setminus f^{-1}(\mathcal{O})$ schreiben. Damit ist das Komplement davon, also das Urbild $f^{-1}(\mathcal{O})$, offen. $\square$

Definition 2.12 (zusammenhängend) Eine Teilmenge $\mathcal{Z}$ der Menge $\mathcal{X}$ heißt *zusammenhängend*, falls $\mathcal{Z}$ nicht die disjunkte Vereinigung zweier (nicht leerer) offener Mengen aus $\mathcal{X}$ ist.

Mit Hilfe der Definition 2.12 lässt sich eine weitere Eigenschaft stetiger Abbildungen festlegen. Ist eine Abbildung f stetig und ist die Menge $\mathcal{Z}$ zusammenhängend, dann folgt daraus, dass auch das Bild $f(\mathcal{Z})$ zusammenhängend sein muss.

Dass das wirklich zutrifft, lässt sich schön auf einfache Weise zeigen. Gegeben seien eine stetige Abbildung $f : \mathcal{X} \rightarrow \mathcal{Y}$, zwei offene, nicht leere Mengen $\mathcal{A} \subseteq \mathcal{Y}$ und $\mathcal{B} \subseteq \mathcal{Y}$ und eine zusammenhängende Menge $\mathcal{Z} \subseteq \mathcal{X}$. Die Menge $f(\mathcal{Z})$ sei die disjunkte Vereinigung von $\mathcal{A}$ und $\mathcal{B}$, also $f(\mathcal{Z}) = \mathcal{A} \dot{\cup} \mathcal{B}$. Dann muss $\mathcal{Z}$ eine Teilmenge vom Urbild $f^{-1}(\mathcal{A} \dot{\cup} \mathcal{B})$ sein. Die Urbilder von $\mathcal{A}$ und $\mathcal{B}$ müssen aufgrund der Stetigkeit offen sein und das Urbild $f^{-1}(\mathcal{A} \dot{\cup} \mathcal{B})$ der disjunkten Vereinigung von $\mathcal{A}$ und $\mathcal{B}$, muss die disjunkte Vereinigung der Urbilder $f^{-1}(\mathcal{A})$ und $f^{-1}(\mathcal{B})$ sein. Es muss also gelten $\mathcal{Z} \subseteq f^{-1}(\mathcal{A} \dot{\cup} \mathcal{B}) = f^{-1}(\mathcal{A}) \dot{\cup} f^{-1}(\mathcal{B})$. Da $\mathcal{Z}$ zusammenhängend ist, muss $\mathcal{Z}$ dann entweder Teilmenge von $f^{-1}(\mathcal{A})$ oder Teilmenge von $f^{-1}(\mathcal{B})$ sein.

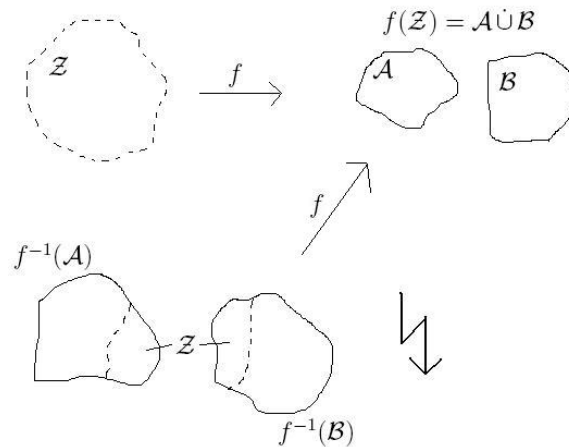


Abbildung 10: bildliche Veranschaulichung der Beweisidee

Nimmt man jetzt ohne Beschränkung der Allgemeinheit an, $\mathcal{Z}$ sei eine Teilmenge von $f^{-1}(\mathcal{A})$ und führt die Abbildung erneut aus, dann erhält man den Ausdruck $f(\mathcal{Z}) \subseteq f(f^{-1}(\mathcal{A}))$. $f(f^{-1}(\mathcal{A}))$ ist wieder $\mathcal{A}$ selbst, also müsste $f(\mathcal{Z})$ eine Teilmenge von $\mathcal{A}$ sein, was einen Widerspruch darstellt, da $f(\mathcal{Z})$ gleich $\mathcal{A} \dot{\cup} \mathcal{B}$ ist und $\mathcal{A} \dot{\cup} \mathcal{B}$ keine Teilmenge von $\mathcal{A}$ sein kann.

Nimmt man an $f(\mathcal{Z})$ wäre nicht zusammenhängend, stößt man also auf einen Widerspruch und deshalb muss für stetige Abbildungen das Bild $f(\mathcal{Z})$ einer zusammenhängenden Menge $\mathcal{Z}$ auch zusammenhängend sein. $\square$

Definition 2.13 (Homöomorphismus) Eine stetige Abbildung $f : (\mathcal{X}, \mathcal{T}_{\mathcal{X}}) \rightarrow (\mathcal{Y}, \mathcal{T}_{\mathcal{Y}})$ heißt *Homöomorphismus*, wenn eine stetige Umkehrabbildung $g = f^{-1} : (\mathcal{Y}, \mathcal{T}_{\mathcal{Y}}) \rightarrow (\mathcal{X}, \mathcal{T}_{\mathcal{X}})$ existiert. Zwei topologische Räume $(\mathcal{X}, \mathcal{T}_{\mathcal{X}})$ und $(\mathcal{Y}, \mathcal{T}_{\mathcal{Y}})$ heißen *homöomorph*, wenn ein *Homöomorphismus* zwischen ihnen existiert.

Beispiele für stetige Abbildungen:

- Die Abbildung $(\mathcal{X}, \mathcal{T}_{\mathcal{X}}) \rightarrow (\mathcal{X}, \mathcal{T}_{\mathcal{X}})$ mit $(x \rightarrow x)$ eines topologischen Raumes auf sich selbst ist stetig, da die Urbilder offener Mengen wieder offen, nämlich die Mengen selbst, sind.
- Die Abbildung $(\mathcal{X}, \{\emptyset, \mathcal{X}\}) \rightarrow (\mathcal{X}, \mathcal{T}_{\mathcal{X}})$ mit $(x \rightarrow x)$ ist nur dann stetig, wenn $\mathcal{T}_{\mathcal{X}} = \{\emptyset, \mathcal{X}\}$ ist, weil andernfalls die Urbilder aller offenen Mengen, außer der Menge $\mathcal{X}$ und der leeren Menge $\emptyset$, nicht offen sind.
- Die Abbildung $(\mathcal{X}, \mathcal{P}(\mathcal{X})) \rightarrow (\mathcal{X}, \mathcal{T}_{\mathcal{X}})$ mit $(x \rightarrow x)$ ist unabhängig von $\mathcal{T}_{\mathcal{X}}$ immer stetig, da das Urbild jeder in $\mathcal{T}_{\mathcal{X}}$ offenen Menge offen in $\mathcal{P}(\mathcal{X})$ ist.
- Die Abbildung $f : (\mathcal{X}, \mathcal{T}_{\mathcal{X}}) \rightarrow (\mathcal{Y}, \mathcal{T}_{\mathcal{Y}})$ mit $x \rightarrow c \in \mathcal{Y}$ ($c \in \mathcal{Y} = \text{konstant}$) ist stetig, weil sich für jedes $\mathcal{O} \in \mathcal{T}_{\mathcal{Y}}$ als Urbild, entweder die leere Menge $\emptyset$, falls $c \notin \mathcal{O}$, oder die Menge $\mathcal{X}$, falls $c \in \mathcal{O}$, ergeben, welche offen in $\mathcal{T}_{\mathcal{X}}$ sind.

Beispiele für Homöomorphismen:

Ein Beispiel für einen Homöomorphismus ist die Exponentialfunktion $\exp : (\mathbb{R}, \mathcal{T}_{\text{euklid}}) \rightarrow (\mathbb{R}_{>0}, \mathcal{T}_{\text{euklid}})$ mit $(x \rightarrow e^x)$. Die Exponentialfunktion ist stetig und es gibt eine stetige Umkehrabbildung, nämlich der natürliche Logarithmus $\ln : (\mathbb{R}_{>0}, \mathcal{T}_{\text{euklid}}) \rightarrow (\mathbb{R}, \mathcal{T}_{\text{euklid}})$ mit $(x \rightarrow \ln x)$. Die Exponentialfunktion ist also ein Homöomorphismus und die Räume $(\mathbb{R}, \mathcal{T}_{\text{euklid}})$ und $(\mathbb{R}_{>0}, \mathcal{T}_{\text{euklid}})$ sind homöomorph. ($\mathbb{R} \cong \mathbb{R}_{>0}$)

Weitere Beispiele für Homöomorphismen werden in Abbildung 11 dargestellt.

2.6 Graphen

Definition 2.14 (Graph) Ein *Graph* besteht anschaulich aus einer Menge von Punkten, welche von Linien verbunden werden. Die Punkte werden als *Knoten* V (von engl. “vertices”) und die Linien als *Kanten* E (von engl. “edges”) bezeichnet. Die Menge Γ , welche den gesamten Graph umfasst, ist somit die disjunkte Vereinigung der Menge aller Knoten V und der Menge aller Kanten E , also $\Gamma = V \dot{\cup} E$ (siehe Abb. 12). Ein *gerichteter* Graph ist ein Graph, bei dem die Kanten eine Richtung haben, bei dem also anschaulich aus den Linien zwischen den Knoten Pfeile werden.

Definition 2.15 (Stern) Die Menge aller Kanten $e \in E$, welche an einem Element $x \in \Gamma$ eines Graphen “hängen”, heißt zusammen mit dem Element x selbst der *Stern* von x (siehe Abb. 13). ($\text{stern}(x) = \{x\} \cup \{e \in E \mid e \text{ “hängt” an } x\}$)

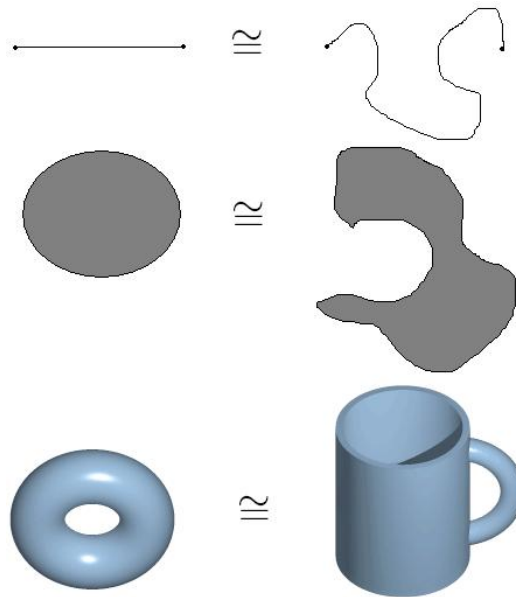


Abbildung 11: Bildbeispiele für Homöomorphismen

Die Sterne sind offene Menge im Sinne von Definition 2.1. Die Topologie eines Graphen wird also von den Sternen des Graphen erzeugt.

Definition 2.16 (Weg) Als *Weg* wird eine stetige Abbildung $\gamma : [0, 1] \rightarrow \Gamma$ bezeichnet, bei der $\gamma(0)$ und $\gamma(1)$ Knoten sind. Ein Weg wird *Rundweg* genannt, wenn Anfangs- und Endknoten gleich sind, also $\gamma(0) = \gamma(1)$ gilt.

Anschaulich kann man sich einen Weg vorstellen, als würde man in einem Graphen entlang der Kanten von einem Knoten zu einem anderen Knoten “laufen”. Kommt man am selben Knoten an, bei dem man gestartet ist, dann ergibt sich ein Rundweg. Der “kürzeste” Weg, also der Weg bei dem man wenigsten Kanten “überläuft”, zwischen zwei Knoten, wird als *Geodätische* bezeichnet. “Läuft” man auf einem Weg entlang einer Kante hin und direkt entlang dieser Kante wieder zurück, so wird diese Kante als *Stachel* bezeichnet (siehe Abb. 14).

Definition 2.17 (wegzusammenhängend) Eine Teilmenge $A \subseteq \Gamma$ eines Graphen heißt *wegzusammenhängend*, falls zwischen je zwei Knoten aus A ein Weg in A existiert.

Man kann sich auch diesen Sachverhalt anschaulich vorstellen, indem zwischen den Knoten entlang der Kanten “läuft”. Ist ein Graph (oder der Teil eines Graphen) wegzusammenhängend, dann bedeutet das anschaulich, dass

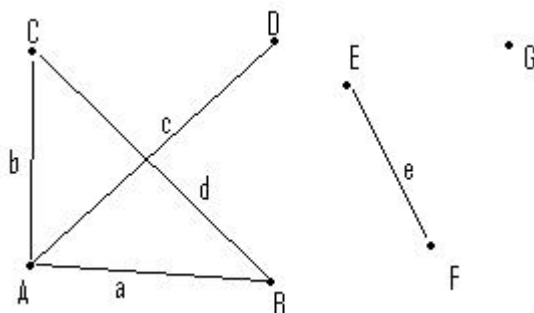


Abbildung 12: Beispiel für einen Graph

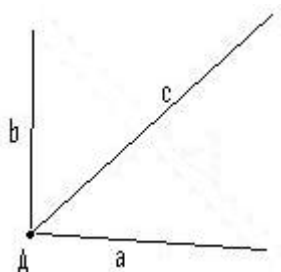


Abbildung 13: Stern von A aus Abb. 12

man von jedem Knoten im Graphen (oder im wegzusammenhängenden Teil des Graphen) zu jedem anderen Knoten im Graphen (oder im wegzusammenhängenden Teil des Graphen) “laufen” kann.

Mit Hilfe dieser Vorstellung ist leicht zu erkennen, dass der gesamte Graph aus Abbildung 12 nicht wegzusammenhängend ist. Allerdings sind mehrere Teile des Graphen wegzusammenhängend z.B. $\{E, F, e\}$. Hier lässt sich auch leicht eine zugehörige stetige Abbildung $\gamma : [0, 1] \rightarrow \{E, F, e\}$ für den Weg zwischen E und F festlegen. Man bildet z.B. 0 auf den Knoten E , 1 auf den Knoten F und alles aus dem Intervall $(1, 0)$ auf die Kante e ab.

Bei Graphen gilt: Ist ein Graph Γ zusammenhängend, dann ist Γ auch wegzusammenhängend.

Definition 2.18 (Baum) Ein *Baum* ist ein zusammenhängender Graph, in dem es keine stachelfreien Rundwege gibt (siehe Abb. 15).

In einem Baum gibt es zwischen je zwei Knoten genau eine Geodätische. Diese Geodätische ist der einzige Weg zwischen zwei Knoten, da es keine Rundwege gibt und es deshalb auch keinen alternativen Weg zwischen zwei Knoten geben kann.

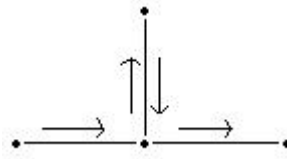


Abbildung 14: Beispiel für Stachel

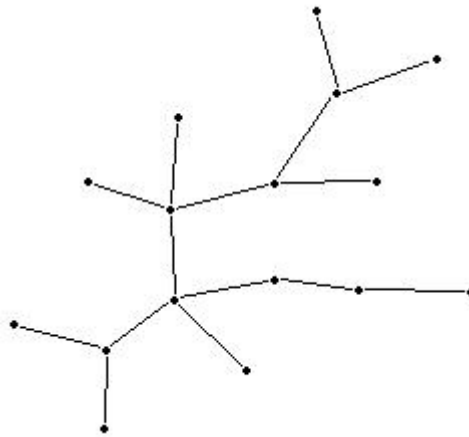


Abbildung 15: Beispiel für einen Baum

In einem zusammenhängenden Graph Γ hat jeder maximale Teilbaum diesselbe Knotenmenge wie der Graph Γ . Um aus einem Graph einen maximalen Teilbaum zu erzeugen, muss man Rundwege beseitigen, was durch Wegnahme von Kanten passiert. Die Anzahl der Knoten bleibt dabei unverändert.

Ist ein (endlicher) Graph ein Baum, dann ergibt die Differenz der Anzahl der Knoten $\#V$ und der Anzahl der Kanten $\#E$ immer eins ($\#V - \#E = 1$). Man kann sich anschaulich klarmachen, dass diese Formel tatsächlich stimmt, in dem man sich zuerst einen minimalen Baum vorstellt und dann überlegt was durch die Hinzunahme von weiteren Knoten und Kanten passiert. Ein minimaler Baum besteht aus genau einem Knoten, für den die Formel offensichtlich stimmt, nämlich $\#V = 1$ und $\#E = 0$. Damit der Graph bei Hinzunahme neuer Elemente ein Baum bleibt, muss zu jeder neuen Kante ein neuer Knoten und zu jedem neuen Knoten eine neue Kante hinzugenommen werden. Bei der Hinzunahme eines neuen Knotens ohne die Hinzunahme einer neuen Kante bleibt der Graph kein Baum, weil der Graph dann nicht mehr zusammenhängend ist und bei der Hinzunahme einer neuen Kante muss man einen neuen Knoten am Ende der Kante hinzufügen, weil der Graph sonst

kein Graph, also auch kein Baum, mehr bleibt. (siehe hierzu Abb. 16). Es kommt also immer die gleiche Anzahl Knoten und Kanten hinzu, weshalb die Differenz $\#V - \#E$ gleich bleibt.

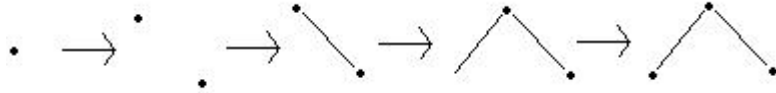


Abbildung 16: grafische Veranschaulichung der Hinzunahme von Knoten oder Kanten

Verallgemeinert man diese Formel für zusammenhängende (endliche) Graphen, so ergibt sich, dass die Anzahl der Kanten $\#E$ immer größer oder gleich der Anzahl der Knoten $\#V$ minus eins ist ($\#E \geq \#V - 1$). Im Spezialfall $\#E = \#V - 1$ ist der Graph ein Baum.

Definition 2.19 (Bettizahlen) Als *nullte Bettizahl* $b_0(\Gamma)$ wird die Anzahl der Zusammenhangskomponenten eines Graphen Γ bezeichnet. Die *erste Bettizahl* $b_1(\Gamma)$ ergibt sich als die Differenz der Anzahl der Kanten $\#E(\Gamma)$ eines Graphen und der Anzahl der Kanten $\#E(T_\Gamma)$, welche ein maximaler Teilbaum T_Γ des Graphen Γ hat. Die Bettizahlen werden auch als *topologische Invarianten* des Graphen bezeichnet.

Weil die Menge der Knoten $V(\Gamma)$ eines Graphen Γ und die Menge der Knoten $V(T_\Gamma)$ eines maximalen Teilbaums T_Γ von Γ gleich sind, kann man sich die erste Bettizahlen $b_1(\Gamma)$ als die Anzahl der "Löcher" in Γ oder auch als die Anzahl der minimalen Rundwege in Γ veranschaulichen.

Eine Möglichkeit zur Bestimmung der Anzahl der Zusammenhangskomponenten, also der nullten Bettizahl b_0 , ist die sogenannte *Tiefensuche*. Als erster Schritt wird hierfür ein beliebiger Knoten aus dem Graph ausgewählt und markiert. Vom markierten Knoten aus wird ein Nachbarknoten ausgewählt und kontrolliert, ob dieser bereits markiert ist. Wurde ein unmarkierter Knoten gefunden wird die Tiefensuche von ihm aus wieder gestartet. Wurde ein markierter Knoten gefunden, geht man zu seinem Vorgänger zurück und versucht es mit einem anderen Knoten. Das Ganze wird so lange wiederholt bis kein unmarkierter Knoten mehr gefunden wird. Sind jetzt alle Knoten des gesamten Graphen markiert, dann ist der Graph zusammenhängend. Gibt es noch unmarkierte Knoten heißt das, es gibt noch weitere Zusammenhangskomponenten. In diesem Fall wird ein noch unmarkierter Knoten ausgewählt und die Tiefensuche wird erneut gestartet. Die Anzahl dieser Suchläufe ist dann gleich der Anzahl der Zusammenhangskomponenten.

Zur Bestimmung der ersten Bettizahl b_1 muss in erster Linie ein maximaler Teilbaum des Graphen gefunden werden. Dazu gibt es zwei grundsätzliche Ansätze, die “top-down” und die “bottom-up” Methode. Bei der “top-down” Methode werden aus dem Originalgraph solange Kanten entfernt, bis es keine Rundweg mehr gibt. Bei der “bottom-up” Methode werden zuerst alle Kanten entfernt, sodass man nur noch das “Knotengerüst” hat. Dann werden wieder sukzessive Kante hinzugefügt, bis der Graph wieder zusammenhängend ist und die Hinzunahme einer weiteren Kante einen Rundweg erzeugen würde. Hat man einen maximalen Teilbaum gefunden, dann muss nur noch die Differenz der Kantenzahlen des gesamten Graphen und des maximalen Teilbaums gebildet werden.

Aus der Differenz der Bettizahlen ergibt sich die sog. Euler-Formel:

$$b_0(\Gamma) - b_1(\Gamma) = \#V(\Gamma) - \#E(\Gamma) \quad (1)$$

Um zu beweisen, dass die Euler-Formel stimmt, betrachtet man zunächst einen zusammenhängenden Graph Γ . Hier ist die nullte Bettizahl $b_0(\Gamma) = 1$, da es genau eine Zusammenhangskomponente gibt. Für die Euler-Formel ergibt sich somit

$$\#V(\Gamma) - \#E(\Gamma) = 1 - b_1(\Gamma) \quad (2)$$

Diese Formel erhält man auch, indem man einen maximalen Teilbaum T_Γ von Γ betrachtet. Für den maximalen Teilbaum T_Γ muss, wie für alle Bäume,

$$\#V(T_\Gamma) - \#E(T_\Gamma) = 1 \quad (3)$$

gelten. Mit $\#V(T_\Gamma) = \#V(\Gamma)$ erhält man

$$\#V(\Gamma) = 1 + \#E(T_\Gamma) \quad (4)$$

Zieht man hiervon auf beiden Seiten $\#E(\Gamma)$ ab, so ergibt sich

$$\#V(\Gamma) - \#E(\Gamma) = 1 - \#E(\Gamma) + \#E(T_\Gamma) = 1 - (\#E(\Gamma) - \#E(T_\Gamma)) \quad (5)$$

Dieses Ergebnis entspricht der obigen Formel 2, weil

$$\#E(\Gamma) - \#E(T_\Gamma) = b_1(\Gamma) \quad (6)$$

ist.

Der nächste Schritt ist der Übergang von einem zusammenhängenden Graph Γ zu einem allgemeinen Graph Γ , der nicht unbedingt zusammenhängend sein muss. Man betrachtet hierbei die zusammenhängenden Teilgraphen Γ_i ,

deren Anzahl man mit n bezeichnet. Der gesamte Graph Γ ist dann die disjunkte Vereinigung aller dieser zusammenhängenden Teilgraphen Γ_i , also

$$\Gamma = \bigcup_{i=1}^n \Gamma_i \quad (7)$$

Die nullte Bettizahl $b_0(\Gamma)$ ergibt sich als

$$b_0(\Gamma) = \sum_{i=1}^n b_0(\Gamma_i) = n \quad (8)$$

und die erste Bettizahl $b_1(\Gamma)$ ist

$$b_1(\Gamma) = \sum_{i=1}^n b_1(\Gamma_i) \quad (9)$$

Auch die Anzahlen der Kanten und der Knoten erhält man durch einfaches Aufsummieren der jeweiligen Anzahlen in allen Teilgraphen, also

$$\#V(\Gamma) = \sum_{i=1}^n \#V(\Gamma_i) \quad (10)$$

und

$$\#E(\Gamma) = \sum_{i=1}^n \#E(\Gamma_i) \quad (11)$$

Aus der obigen Formel 2 wird durch das Summieren

$$\sum_{i=1}^n \#V(\Gamma_i) - \sum_{i=1}^n \#E(\Gamma_i) = \sum_{i=1}^n b_0(\Gamma_i) - \sum_{i=1}^n b_1(\Gamma_i) \quad (12)$$

und somit die Euler-Formel

$$b_0(\Gamma) - b_1(\Gamma) = \#V(\Gamma) - \#E(\Gamma) \quad (13)$$

Eine weitere, vielleicht etwas anschaulichere, Möglichkeit zum Beweis der Euler-Formel ergibt sich aus der Kontraktion von Kanten auf Knoten, welche eine stetige Abbildung auf Graphen darstellt.

Hierbei gilt

$$b_0(\Gamma) = b_0(\Gamma') = b_0(\Gamma'') \quad (14)$$

$$b_1(\Gamma) = b_1(\Gamma') = b_1(\Gamma'') \quad (15)$$

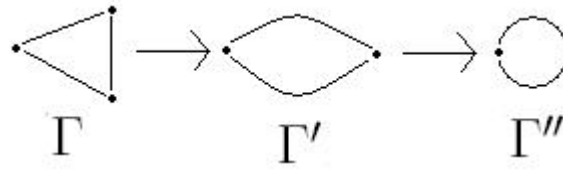


Abbildung 17: Beispiel für Kantenkontraktion

und

$$\#V(\Gamma) - \#E(\Gamma) = \#V(\Gamma') - \#E(\Gamma') = \#V(\Gamma'') - \#E(\Gamma'') \quad (16)$$

ändert sich also nichts an den Werten in der Euler-Formel. Man kann also einen Graphen immer weiter kontrahieren bis nur noch Knoten mit Schleifen dran übrigbleiben. (siehe Abb. 17).

Jetzt bleibt nur noch zu beweisen, dass jedes dieser so entstandenen Gebilde die Eulerformel erfüllt. Hierfür sei die Anzahl der verbliebenen Knoten n und die Anzahl der Kanten (Anzahl der Schleifen) sei m . Es gilt somit:

$$n - m = \#V(\Gamma) - \#E(\Gamma) \quad (17)$$

n ist aber auch die Anzahl der Zusammenhangskomponenten, also die nullte Bettizahl $b_0(\Gamma)$, weil die verbliebenen Knoten nicht mehr miteinander verbunden sind. Die erste Bettizahl $b_1(\Gamma)$ ist gleich m , da die Anzahl der Schleifen in diesem Fall auch die Anzahl der minimalen Rundwege darstellt. Mit $n = b_0(\Gamma)$ und $m = b_1(\Gamma)$ ergibt sich aus Formel 17 die Euler-Formel:

$$b_0(\Gamma) - b_1(\Gamma) = \#V(\Gamma) - \#E(\Gamma) \quad (18)$$

Definition 2.20 (Relation) Eine binäre *Relation* R ist eine Teilmenge des kartesischen Produkts zweier Mengen A und B , also $R \subseteq A \times B$ mit $A \times B := \{(a, b) \mid (a \in A) \text{ und } (b \in B)\}$. Ist $A = B$ und somit $R \subseteq A \times A$, dann nennt man die Relation *homogen*.

Man schreibt die Relation R zwischen a und b auch als $(a, b) \in R$ oder aRb . Eine Relation $R \subseteq A \times B$ und eine Relation $S \subseteq B \times C$ können verkettet, d.h. hintereinander ausgeführt, werden. Das Ergebnis schreibt man dann RS oder $R \circ S$ und es gilt $RS = R \circ S = \{(a, c) \in A \times B \mid \text{es gibt ein } b \in B \text{ mit } aRb \text{ und } bSc\}$. Durch die Verkettung einer homogenen Relation $R \circ R$ mit sich selbst ergibt sich wieder eine homogene Relation. Hierfür wird die Schreibweise $R^2 = R \circ R$ bzw. R^n (für $n \in \mathbb{N}$) verwendet, wobei als R^0 die Relation eines Elementes zu sich selbst geschrieben wird, also $R^0 = \{(a, a) \mid$

$a \in A\}$. Als *transitive und reflexive Hülle* R^* von R wird die Vereinigung aller R^n bezeichnet, also $R^* = \bigcup_{n \in \mathbb{N}} R^n$.

Homogene binäre Relationen $R \subseteq A \times A$, wie sie hier definiert wurden, können auch als gerichtete Graphen interpretiert werden, wenn die Menge A gleich der Menge V der Knoten ist. Hierbei wird durch die Relation R eine Beziehung zwischen den Knoten hergestellt, welche den gerichteten Kanten entspricht.

Einen Graph Γ bestehend aus der Menge $X = V \dot{\cup} E$ kann man z.B. durch die homogene binäre Relation $R_\Gamma \subseteq X \times X$ definieren. Hierbei soll $xR_\Gamma y$ gelten, wenn $y \in \overline{\{x\}}$ und $x \neq y$ gilt. Dadurch erhält man für jede Kante die beiden zugehörigen Knoten. Der Stern eines Elementes $x \in X$ wäre mit Hilfe dieser Relation R_Γ zu schreiben als $\text{stern}(x) = \{y \in X \mid yR_\Gamma x\}$. Dadurch erhält man alle Elemente y aus X , welche in irgendeiner Relation zu x stehen, die also an x "hängen". Die Topologie einer Relation R kann man somit verstehen als die von ihrer transitiven und reflexiven Hülle R^* erzeugte Topologie, welche anschaulich der von den Sternern erzeugten Topologie eines Graphen entspricht.

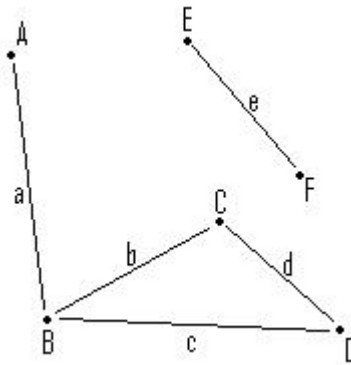


Abbildung 18: Beispielgraph

Für den Beispielgraph aus Abbildung 18 würde sich aus der Relation R_Γ z.B. die Matrix aus Abbildung 19 ergeben.

Jetzt könnten die Fragen auftauchen, was man aus solchen Matrizen über den dazugehörigen Graph ablesen kann und wozu man diese Erkenntnis eventuell zu nutzen sind. Unter Anderem mit diesen Fragen beschäftigt sich der Abschnitt 3.

R_{Γ}	A	B	C	D	E	F	a	b	c	d	e
A							x				
B							x	x	x		
C								x		x	
D									x	x	
F											x
E											x
a											
b											
c											
d											
e											

Abbildung 19: Matrizendarstellung der Relation R_{Γ} für den Graph aus Abb. 18

3 Anwendungen

In der Geoinformatik und in der Kartografie werden Graphen häufig dazu verwendet, um Sachverhalte, deren Bedeutung nicht in der genauen Lage von Objekten, sondern viel mehr in den Lagebeziehungen der Objekte zueinander, begründet liegt, zu vereinfachen, zu veranschaulichen oder zu analysieren. Ein Beispiel hierfür sind Flächenadjazenzgraphen (siehe Abb. 20), bei denen Flächenschwer- oder -mittelpunkte (oder irgendwelche anderen inneren Punkte) als Knoten mit Kanten verbunden werden. Hiermit kann z.B. die Nachbarschaft von Flurstücken vereinfacht dargestellt werden. Ein weite-

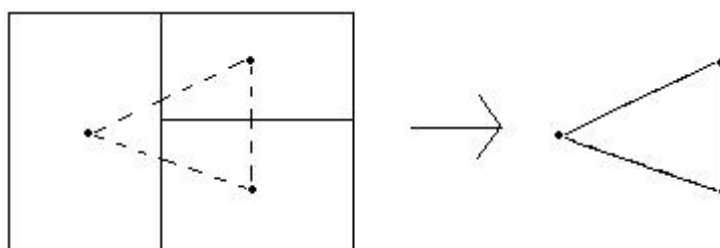


Abbildung 20: Beispiel für Flächenadjazenzgraph

res Beispiel stellen Liniennetzpläne von öffentlichen Verkehrsmitteln dar, hier sind die Haltestellen die Knoten und die Kanten werden durch die verbindenden Linien dargestellt (siehe Abb. 21).

Auch Navigationssysteme basieren häufig auf Graphen. In diesem Fall



Abbildung 21: Liniennetzplan der Wiener U-Bahn

werden *gewichtete* Graphen eingesetzt. Ein *gewichteter* Graph ist ein Graph, dessen Kanten Gewichte, also Zahlenwerte, zugewiesen worden sind. Im Fall der Navigationssysteme können diese Gewichte als die Strecke oder auch als die zeitliche Entfernung zwischen zwei Knoten aufgefasst werden. Zur Bestimmung des kürzesten bzw. schnellsten Weges zwischen zwei Knoten wird der Dijkstra-Algorithmus verwendet. Dieser funktioniert ähnlich wie die bereits erwähnte Tiefensuche, nur dass man vom Startknoten nicht irgendeinen Nachbarknoten wählt, sondern den “nähesten” unmarkierten Knoten, d.h. den Knoten der über die Kante mit dem kleinsten Gewicht zu erreichen ist. Von diesem Knoten aus sucht man den nächsten “nähesten” Knoten u.s.w. bis man am Zielknoten ankommt.

Ein sehr bekanntes Beispiel ist auch das sogenannte Königsberger Brückenproblem, welches im Abschnitt 3.1 ausführlich behandelt wird.

3.1 Das Königsberger Brückenproblem

Im 18. Jahrhundert gab es in Königsberg sieben Brücken über den Fluss Pregel, der, aufgeteilt in zwei Arme, durch Königsberg fließt (siehe Abb. 22). Die beiden Arme des Pregel umfließen eine Insel, der Kneiphof.

Einige Bürger der Stadt Königsberg stellten sich irgendwann bei einem Spaziergang die Frage, ob es wohl möglich wäre einen Spaziergang durch Königsberg zu machen, bei dem jede der sieben Brücken genau einmal über-

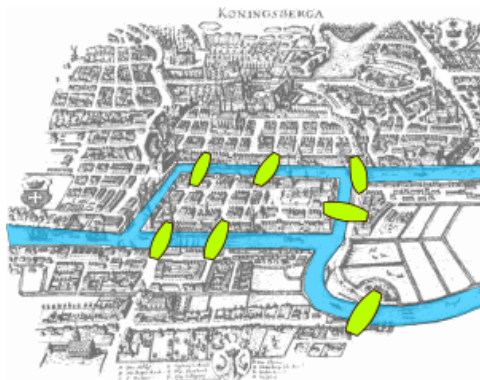


Abbildung 22: Königsberger Brücken im 18. Jh.

quert würde und ob es vielleicht sogar einen Rundweg gäbe, bei dem man beim Ausgangspunkt wieder ankäme. Durch einfaches Probieren war das Problem nicht zu lösen. Die Bürger Königsbergs hatten zwar keinen Weg gefunden, man konnte sich aber nie sicher sein, ob es nicht doch eine Lösung gäbe. Man beauftragte also einen Mathematiker, Leonhard Euler, damit das Problem zu untersuchen. Euler erkannte, dass es zur Lösung des Problems nicht auf die genaue Lage der Brücken ankommt, sondern nur darauf, welche Brücke welche Landmassen miteinander verbindet. Man kann das Königsberger Brückenproblem also als Graph darstellen, bei dem die Brücken durch die Kanten und die Landteile durch die Knoten repräsentiert werden (siehe Abb. 23).

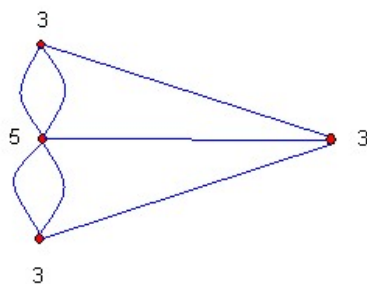


Abbildung 23: Königsberger Brücken als Graph

Mittels dieses Graphen stellte Euler fest, dass es für das Königsberger Brückenproblem keine Lösung gibt, indem er das Problem auf beliebige Graphen verallgemeinerte. Euler beschäftigte sich hierbei mit der Frage, ob es in einem beliebigen Graphen einen Weg bzw. einen Rundweg gibt, welcher genau einmal über alle Kanten führt. Euler kam zu dem Ergebnis, dass es

nur eine Lösung geben kann, wenn von allen Knoten, bis auf maximal zwei, eine gerade Anzahl von Kanten abgehen. Soll es einen Rundweg geben, dann fällt der Einschub “bis auf maximal zwei” weg. Außerdem muss der Graph zusammenhängend sein.

Das Ergebnis von Euler ist aus der anschaulichen Vorstellung gut nachzuvollziehen. Wenn es einen Weg geben soll, bei dem jede Kante genau einmal passiert wird, dann muss man jeden Knoten über eine andere Kante “verlassen” können, als über die Kante von der aus man “angekommen” ist, d.h. es muss eine gerade Anzahl von Kanten an jedem Knoten geben. Eine Ausnahme bilden der Anfangs- und der Endknoten des Weges. Diese beiden Knoten können auch über eine ungerade Anzahl von Knoten verfügen, weil man diese Knoten entweder nur “verlässt” oder nur “ankommt”. So kommt der Einschub “bis auf maximal zwei” zustande. Das gilt wiederum nicht für einen Rundweg, da für einen Rundweg Anfangs- und Endknoten gleich sind und man an diesem Knoten wieder über eine andere Kante “ankommen” muss, als über die Kante von der aus man ihn zu Beginn “verlassen” hat. Dass der Graph zusammenhängend sein muss, ist auch klar, weil sonst Teil des Graphen nicht “erreichbar” wären und so nicht alle Kanten verwendet werden könnten.

Betrachtet man mit dieser Erkenntnis den Graphen, welcher die Königsberger Brücken repräsentiert, dann kommt man zum Ergebnis, dass es keinen Weg geben kann, weil von allen Knoten eine ungerade Anzahl von Kanten abgeht.

Ein bekanntes Beispiel für einen Graph bei, dem es einem solchen Weg gibt, ist das “Haus vom Nikolaus”. Bei diesem Beispiel gehen von jedem

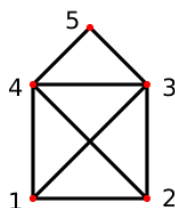


Abbildung 24: Das Haus vom Nikolaus

Knoten, außer den beiden unteren, eine gerade Anzahl Kanten ab (siehe Abb. 24: Knoten 1 und 2: jeweils 3 Kanten; Knoten 3 und 4: jeweils 4 Kanten; Knoten 5: 2 Kanten). Es gibt also einen Weg bei dem alle Kanten genau einmal benutzt werden, was man auch erwartet hat, weil man das “Haus vom Nikolaus” bekanntermaßen ohne Absetzen zeichnen kann. Rundweg gibt es allerdings keinen, weil die zwei Knoten 1 und 2 mit einer ungeraden Anzahl

von Kanten verbunden sind. Auch diesen Sachverhalt kennt man eigentlich schon vom Zeichnen des “Haus vom Nikolaus”, wobei man immer unten anfangen muss, damit es funktioniert. Das ist so, weil beiden unteren Knoten (mit der ungeraden Kantenzahl) Anfangs- und Endknoten sein müssen, damit es einen Weg gibt.

3.2 Lösung des Brückenproblems mittels des Topologischen Datentyp

Als Anwender kommt jetzt womöglich die Frage auf, wie sich ein Programm implementieren lässt, welches solche Wege bzw. Rundwege in Graphen erkennt. Ein Ansatz zur Beantwortung dieser Frage sind, wie bereits am Ende des Abschnitts 2.6 angedeutet, die durch die Anwendung der Relation $R_\Gamma \subseteq X \times X$ entstehenden Matrizen. Allerdings erscheint es bei näherer Betrachtung zweckmäßig, die Relation $D_\Gamma \subseteq E \times V$ als Reduktion von R_Γ auf die gleiche Beziehung zwischen Kanten und Knoten zu verwenden, da sich so weniger leere Felder ergeben. Die Matrizen die sich daraus ergeben, werden im Folgenden *topologische Datentypen* genannt. Für den Beispielgraph aus der Abbildung 18 ergibt sich aus der Relation D_Γ z.B. die Matrix, also der *topologische Datentyp* aus Abbildung 3.2. Die Knoten werden in der Matrix

D_Γ	a	b	c	d	e
A	x				
B	x	x	x		
C		x		x	
D			x	x	
E					x
F					x

Abbildung 25: topologischer Datentyp des Graphen aus Abb. 18

durch die Zeilen und die Kanten durch die Spalten repräsentiert. Ein Eintrag bedeutet, dass der jeweilige Knoten und die jeweilige Kante verbunden sind. Die Matrix aus Abbildung 26 ist der topologische Datentyp des Graphen aus Abbildung 12 bestätigt diese Vermutung. Hier bildet der Knoten G allerdings eine Ausnahme, weil er alleine steht und mit keiner Kante verbunden ist, d.h. es gibt in der Zeile für G keine Einträge. Für die Frage, ob es einen Weg gibt, der alle Kanten genau einmal benutzt, sind alleinstehende Knoten nicht relevant, obwohl der Graph durch sie nicht mehr zusammenhängend ist. Da solche Knoten mit keiner Kante verbunden sind und die Benutzung aller Kanten, aber nicht aller Knoten gefordert ist, kann man leere Zeilen in

D_Γ	a	b	c	d	e
A	x	x	x		
B	x			x	
C		x		x	
D			x		
E					x
F					x
G					

Abbildung 26: topologischer Datentyp des Graphen aus Abb. 12

diesem Zusammenhang aus dem topologischen Datentyp eliminieren.

Betrachtet man diese Matrizen so liegt die Vermutung nahe, dass es eine Blockzerlegung (angedeutet durch die Trennlinien in den Matrizen) gibt, so dass die Blöcke diagonal zueinander angeordnet sind und es außerhalb der Diagonalen keine weiteren Einträge gibt. Allerdings müssen diese Blöcke nicht quadratisch sein. Daraus ergibt sich weiter die Vermutung, dass die Anzahl der Blöcke gleich der Anzahl der Zusammenhangskomponenten ist.

Um diese Vermutung zu beweisen betrachtet man zunächst einen topologischen Datentyp mit n Zeilen und m Spalten. Die Knoten werden als $A_1 \dots A_n$ und die Kanten als $a_1 \dots a_m$ bezeichnet (siehe Abb. 27). Nimmt

D_Γ	$a_1 \dots a_m$
A_1	...
$\vdots$	...
A_n	...

Abbildung 27: topologischer Datentyp mit n Zeilen und m Spalten

D_Γ	$a_1 \dots a_j$	$a_{j+1} \dots a_m$
A_1	...	
$\vdots$	...	
A_i	...	
A_{i+1}		...
$\vdots$		...
A_n		...

Abbildung 28: Zerlegung in zwei Blöcke

man jetzt an es gäbe eine Zerlegung in zwei (oder mehr) Blöcke, wie in Abbildung 28 angedeutet (Punkte deuten an, wo es Einträge gibt), dann lässt

sich daraus schließen, dass diese es zwei (oder mehr) Teilgraphen gibt, welche nicht zusammenhängend sind. In den Zeilen $A_1 \dots A_i$ gibt es keine Einträge für die Spalten $a_{j+1} \dots a_m$ und umgekehrt gibt es in den Spalten $a_1 \dots a_j$ auch keine Einträge für die Zeilen $A_{i+1} \dots A_n$, d.h. es gibt keine Verbindung von den Knoten $A_1 \dots A_i$ zu den Kanten $a_{j+1} \dots a_m$ und auch keine Verbindung von den Kanten $a_1 \dots a_j$ zu den Knoten $A_{i+1} \dots A_n$. Zwischen der Knotengruppe $A_1 \dots A_i$, mit den zugehörigen Kanten $a_1 \dots a_j$, und der Knotengruppe $A_{i+1} \dots A_n$ mit den zugehörigen Kanten $a_{j+1} \dots a_m$, gibt es also keine Verbindung, denn eine Verbindung würde zusätzliche Einträge erfordern. Die beiden Gruppen von Knoten und Kanten repräsentieren also beide einen Teilgraphen, welcher jeweils keine Verbindung zum anderen Teilgraphen hat. Der gesamte Graph hat also mindestens soviele Zusammenhangskomponenten, wie die Anzahl der Blöcke.

Im nächsten Schritt gilt es noch zu klären, ob ein Block, welcher sich nicht mehr in weitere Blöcke zerlegen lässt, immer genau einen Teilgraph und somit eine Zusammenhangskomponente, repräsentiert. Hierfür stellt man sich einen Block vor, welcher sich nicht mehr zerlegen lässt, aber noch zwei (oder mehr) Zusammenhangskomponenten repräsentiert. Dies führt zu einem Widerspruch, da sich zwei (oder mehr) Zusammenhangskomponenten wieder als zwei (oder mehr) nicht zusammenhängende Teilgraphen auffassen ließen, was in der Matrix bedeuten würde, dass es keine Einträge bei den jeweils nicht verbundenen Kanten- und Knotengruppen gäbe. Das bedeutet wiederum es wäre eine weitere Blockzerlegung möglich. Lässt sich ein Block nicht weiter zerlegen, dann repräsentiert er also genau eine Zusammenhangskomponente und somit ist die Anzahl der Zusammenhangskomponenten gleich der Anzahl der Blöcke. $\square$

Ein Weg, wie er in Abschnitt 3.1 beschrieben worden ist, liegt dann vor, wenn

1. es keine Blockzerlegung gibt (der Graph also zusammenhängend ist) und
2. es in jeder Zeile, bis auf maximal zwei, der Matrix eine gerade Anzahl von Einträgen gibt.

Die zweite Bedingung spiegelt einfach die von Euler gefundene Bedingung wider, dass von jedem Knoten, bis auf maximal zwei, eine gerade Anzahl Kanten abgehen muss. Zählt man die Einträge einer Zeile, erhält man nämlich die Anzahl der Kanten, welche mit dem entsprechenden Knoten verbunden sind. Soll es einen entsprechenden Rundweg geben, dann fällt der Einschub "bis auf maximal zwei" in der zweiten Bedingung wieder weg.

Man könnte sich jetzt auch noch die Frage stellen, ob man aus dem topologischen Datentyp auch ablesen kann, ob es einen Weg zwischen zwei

beliebigen Knoten gibt. Diese Frage lässt sich auch mit Hilfe der Blockzerlegung beantworten. Möchte man wissen, ob es zwischen zwei Knoten A_i und A_j einen Weg gibt, muss man nur überprüfen, ob die zugehörigen beiden Zeilen A_i und A_j zum gleichen Block in der Matrix gehören. Ist das der Fall, dann gehören die beiden Knoten zur gleichen Zusammenhangskomponente und da zusammenhängend bei Graphen äquivalent zu wegzusammenhängend ist, muss es dann auch einen Weg zwischen den beiden Knoten geben.

3.3 Bestimmung der Bettizahlen mittels des Topologischen Datentyps

Betrachtet man den topologischen Datentyp, dann taucht vielleicht auch noch die Frage auf, wie man die, in der Euler-Formel verwendeten, Bettizahlen b_0 und b_1 daraus bestimmen kann. Die Beantwortung dieser Frage dürfte nach dem vorhergehenden Abschnitt 3.2 keine größere Schwierigkeit mehr darstellen.

Die nullte Bettizahl b_0 ist gleich der Anzahl der Zusammenhangskomponenten und im Abschnitt 3.2 wurde mit der Blockzerlegung bereits eine Möglichkeit zur Bestimmung der Anzahl der Zusammenhangskomponenten vorgestellt. Der einzige Unterschied ist hierbei, dass für die nullte Bettizahl b_0 auch alleinstehende Knoten, welche bei der Lösung des Brückenproblems unberücksichtigt geblieben sind, als Zusammenhangskomponenten mitgezählt werden müssen. Will man also die nullte Bettizahl b_0 aus dem topologischen Datentyp bestimmen, muss man zuerst die leeren Zeilen aus der Matrix eliminieren und diese dabei zählen. Dann führt man eine Blockzerlegung durch und addiert die Anzahl der Blöcke zur Anzahl der vorher eliminierten Zeilen, also zur Anzahl der alleinstehenden Knoten. So erhält man die gesamte Anzahl der Zusammenhangskomponenten.

Um die erste Bettizahl b_1 zu bestimmen, muss man die Anzahl der "Löcher", also die Anzahl der minimalen Rundwege, im Graphen bestimmen. Ein minimaler Rundweg liegt in einem Graphen dann vor, wenn es einen Rundweg gibt, für den die gleiche Anzahl von Knoten und Kanten verwendet werden und es vom gleichen Knoten aus keinen solchen Rundweg mit weniger Kanten und Knoten gibt. Anschaulich kann man sich gut vorstellen, warum ein solcher Rundweg minimal ist, da es keinen stachelfreien Rundweg gibt, in welchem weniger Kanten verwendet werden. Im Graph aus Abbildung 18 stellt z.B. das Dreieck BCD zusammen mit den verbindenden Kanten b, c und d einen minimalen Rundweg dar. Im topologischen Datentyp erkennt man solche Rundweg wieder an Teilmatrizen. Ein minimaler Rundweg liegt genau dann vor, wenn eine quadratische Teilmatrix im topologischen Daten-

typ vorliegt, welche in jeder Zeile und in jeder Spalte genau zwei Einträge hat. Für den minimalen Rundweg, welchen das Dreieck BCD aus dem Graph aus Abbildung 18 darstellt, ergibt sich z.B. die Teilmatrix aus Abbildung 29. Dass das für jeden minimalen Rundweg stimmt, ist leicht erkennbar. Die

	b	c	d
B	x	x	
C	x		x
D		x	x

Abbildung 29: Teilmatrix für minimalen Rundweg aus Abb. 18

Teilmatrix muss quadratisch sein, da die Anzahl der Knoten und der Kanten gleich ist. Dass es in jeder Zeile und in jeder Spalte genau zwei Einträge gibt, bedeutet dass jeder Knoten mit genau zwei Kanten verbunden ist und dass jede Kante wiederum genau zwei Knoten verbindet. Diese Eigenschaft ist auch anschaulich, weil bei gleicher Knoten- und Kantenanzahl in einem Rundweg jeder Knoten und jede Kante genau einmal passiert wird. Man muss also, ähnlich wie beim Brückenproblem, einen Knoten über eine andere Kante “verlassen”, welche wiederum zu einem noch nicht benutzten Knoten führt, als über die Kante von der aus man “angekommen” ist. Eine Ausnahme stellen hierbei Knoten mit einer Schleife dran dar, welche auch minimale Rundwege sind. Im topologischen Datentyp äußern sich diese als Spalten mit nur einem Eintrag. Diesen Eintrag könnte man zwar auch noch als (minimale) quadratische Teilmatrix auffassen, allerdings würde ein Algorithmus der minimale Rundwege sucht, diese nicht erkennen, da das Kriterium, dass es in jeder Zeile und in jeder Spalte genau zwei Einträge geben muss, nicht gilt. Man muss die Schleifen also gesondert behandeln, ähnlich wie bei der nullten Bettizahl b_0 die alleinstehenden Knoten.

Um die erste Bettizahl b_1 aus dem topologischen Datentyp zu bestimmen, muss man also zuerst die Spalten mit nur einem Eintrag aus der Matrix eliminieren und diese zählen. Danach bestimmt man die Anzahl der quadratischen Teilmatrizen, in denen es in jeder Zeile und in jeder Spalte genau zwei Einträge gibt und addiert diese zu der Anzahl der entfernten Spalten.

4 Fazit und Ausblick

Das Themengebiet “Topologie” ist sehr umfangreich und kann im Rahmen einer solchen Arbeit natürlich nicht vollständig abgehandelt werden. Dennoch bin ich der Meinung, dass ein guter Überblick über die Grundlagen des Themas “Topologie von Graphen” geschaffen wurde und dass ein Eindruck entsteht, wo diese Grundlagen in der Geoinformatik Anwendung finden oder finden könnten. Vor allem wurde deutlich, dass es zwischen den mathematischen Grundlagen und den Begriff “Topologie”, wie in der Geoinformatik-Lehre immer so ein wenig unpräzise “herumschwirrt”, tatsächlich ein starker Zusammenhang besteht.

Für die in den Abschnitten 3.2 und 3.3 erwähnten Lösungsansätze, werden in dieser Arbeit nur die jeweiligen Ideen beschrieben. Bei einer tatsächlichen Implementierung, welche noch austünde, eines Programmes, welches diese Ideen benutzt, müssten noch einige zusätzliche Dinge beachtet werden.

Die größte Hürde würde vermutlich der Sachverhalt darstellen, dass die Zeilen und Spalten des topologischen Datentyps nicht von vorneherein so angeordnet sein müssen, dass die Blockzerlegung für die Zusammenhangskomponenten bzw. auch die quadratischen Teilmatrizen, welche einen minimalen Rundweg repräsentieren, unmittelbar erkennbar sind, d.h. anschaulich die Teilmatrizen der Blockzerlegung bzw. die quadratischen Teilmatrizen stehen in der Matrix nicht als Blöcke zusammen. Dieser Sachverhalt führt dazu, dass man theoretisch alle denkbaren Permutationen von Zeilen und Spalten nach möglichen Blockzerlegungen bzw. quadratischen Teilmatrizen absuchen müsste. Für eine tatsächliche Implementierung wäre noch zu klären, ob sich diese “brute force” Methode vielleicht durch Ausschluss bestimmter Permutationen beschleunigen ließe und ob diese Implementierung dann wirklich schneller zu einem Ergebnis führen würde, als bereits bestehende Implementierungen.

Literatur

- [1] Jänich, Klaus: Topologie. Springer-Lehrbuch 1990
- [2] Paul, Norbert: Topologische Datenbanken für Architektonische Räume. Dissertation. Universität Karlsruhe (TH). 2008
- [3] Bradley, Patrick Erik; Paul, Norbert: Using the relational model to capture topological information of spaces. The Computer Journal, Vol. 53, No. 1 (2010), S.69-89. 2010
- [4] Singh, Simon: Fermats letzter Satz. dtv 2000
- [5] Euler, Leonhard: Solutio problematis ad geometriam situs pestinentis. Comment. acad. scient. Petrop., 8, 128-140 1741
- [6] [http://de.wikipedia.org/wiki/Relation_\(Mathematik\)](http://de.wikipedia.org/wiki/Relation_(Mathematik))
- [7] http://de.wikipedia.org/wiki/Königsberger_Brückenproblem
- [8] <http://www.xplora.org/downloads/Knoppix/MathePrisma/Start/Module/Koenigsb/index.htm>
- [9] <http://www.geoinformation.net/>
- [10] Bilder: http://commons.wikimedia.org/wiki/Main_Page