

Überlagerungen simplizialer Flächenkomplexe durch Faltungen

Coverings of simplicial surfaces using foldings

MASTERARBEIT

vorgelegt der

FAKULTÄT FÜR MATHEMATIK, INFORMATIK
UND NATURWISSENSCHAFTEN

angefertigt am

LEHRSTUHL B FÜR MATHEMATIK
RWTH AACHEN UNIVERSITY

von

MARKUS BAUMEISTER

betreut von

PROF. DR. WILHELM PLESKEN
PROF. DR. ALICE NIEMEYER

30.09.2016

Inhaltsverzeichnis

1	Einleitung	5
1.1	Herangehensweise	6
1.2	Aufbau dieser Arbeit	7
1.3	Vorkenntnisse und Notation	7
2	Grundlegende Modellierung	9
2.1	Homogene Komplexe	12
2.1.1	Lokale Simplität	16
2.2	Überlagerungskomplexe	17
2.3	Zusammenfalten	19
2.3.1	Primitive Erweiterungen	23
2.3.2	Faltung benachbarter Flächen	28
2.3.3	Kommutativität	30
2.4	Entfalten	33
2.5	Einbettungen	36
2.5.1	Färbbarkeit	39
2.5.2	Verträglichkeit mit Faltung	42
3	Faltkomplexe	45
3.1	Definition des Faltzustandes	45
3.1.1	Definition von Fächern	46
3.1.2	Definition von Oberflächen	49
3.1.3	Abhängigkeiten zwischen Fächern	53
3.1.4	Definition von Randstücken	57
3.1.5	Vollständige Definition	62
3.2	Definition der Faltung	63
3.2.1	Definition der Fächersumme	65
3.2.2	Primitive Faltung vom Grad $n - 1$	74
3.3	Mehrfachfaltung	78
3.4	Entfaltungen	85
3.5	Faltpläne	91
3.5.1	Kommutativität	93
3.5.2	Entfaltung	96
4	Die Einbettungsfrage	103
4.1	Überlagerung eines Dreiecks	107
4.1.1	Kreisdarstellung	108
4.1.2	Einfaches Beispiel	117
4.1.3	Gruppenoperation	120
4.1.4	Tatsächliche Faltung	124
4.1.5	Maximale Einbettung	127
4.1.6	Orientierbarkeit	128

4.1.7	Algorithmisches Vorgehen	133
4.2	Überlagerung einiger ebener Gitter	136
4.2.1	Beschreibung des hexagonalen Gitters	137
4.2.2	Unendliche Faltungen	140
4.2.3	Automorphismen des hexagonalen Gitters	146
4.2.4	Beschreibung des quadratischen Gitters	147
5	Modellerweiterungen	151
5.1	Allgemeine m -Ecke als Flächen	151
5.2	Unregelmäßige Flächen	156
	Literatur	159
	Index	160
	Danksagung	162
	Eidesstattliche Erklärung	163

1 Einleitung

Faltungen haben eine reiche Geschichte in der Kunst, werden häufig praktisch angewendet und nehmen eine zunehmend größere Rolle in der Mathematik ein. In der Kunst ist die Faltung von Papier als Origami bekannt. Dort startet man häufig mit einem quadratischen Blatt Papier, das man beliebig falten darf. Dabei ist es in der Regel nicht erlaubt, zu schneiden oder zu kleben. Mit diesen Einschränkungen haben Künstler wie Robert Lang sehr eindrucksvolle Werke erstellt, wie z.B. das Orchester aus Abbildung 1 (jede Figur ist dabei aus einem einzelnen Blatt Papier gefaltet).



Abb. 1: Orchester von Robert Lang ([Lan04])

Auch außerhalb der Kunst finden Faltungen vielfache Anwendungen. Aus dem Alltag kennt man das Falten von Landkarten und Kleidungsstücken. In der Raumfahrt werden Faltungen standardmäßig verwendet, da sie eine sehr platzsparende Aufbewahrung erlauben. Auch Verpackungen und Airbags werden häufig gefaltet, um nur einige der verbreitetsten Anwendungen zu nennen.

Vergleicht man den künstlerischen mit dem industriellen Zugang, dann sehen die industriellen Faltungen immer etwas blasser aus. Während die Kunst filigran und detailliert ist, sehen sie gröber und robuster aus. Beispielsweise sind bei der Faltung einer Landkarte alle möglichen Faltkanten schon vorher vorgegeben, während dies bei der Faltung der Rose aus Abbildung 2 nicht der Fall ist. Das hat einen sehr einfachen Grund: Die Kunst beschäftigt sich damit, Unikate zu schaffen und damit Grenzen zu erweitern, während praktische Anwendungen eher von ihrer vielseitigen Anwendbarkeit und daher Einfachheit profitieren. Daher wirken solche praktischen Anwendungen im direkten Vergleich oft sehr blass, obwohl sie einen deutlich größeren Einfluss auf unseren Alltag haben.



Abb. 2: Miura-ken Rose von Robert Lang ([Lan06])

In allen diesen Situationen fällt dem Betrachter auf, dass es Strukturen beim Falten gibt. Es gibt einige Faltungen, die man ausführen kann und andere, deren Ausführung unmöglich

ist. Um nur zwei der möglichen Fragestellungen aus der Literatur zu nennen: Man kann sich fragen, ob man eine Landkarte falten kann, wenn man vorgibt, in welche Richtung die Kanten geknickt werden (vgl. [ABD⁺01]) oder untersuchen, wann man ein Muster an Faltkanten flach falten kann (vgl. [Hul94]). Wenn man sich mit solchen Fragen beschäftigt, gelangt man zur mathematischen Beschreibung von Faltungen. Diese ist ebenso divers wie die Anwendungsgebiete von Faltungen.

In dieser Arbeit beschäftigen wir uns mit der Situation, dass die möglichen Faltkanten bereits feststehen, wie bei einer Landkarte. Unser Leitstern wird die Frage nach allen möglichen Faltungen sein, die wir aus einem solchen Muster erhalten können. Um diese Frage anzugehen, ist es nötig, ein Modell von Faltungen anzugeben, das sowohl aussagestark als auch anwendbar ist.

1.1 Herangehensweise

Zur Modellierung und Untersuchung von Papierfaltungen starten wir mit einem Muster von Polygonen auf Papier, wobei die Kanten dieses Musters die möglichen Faltkanten darstellen (d. h. alle möglichen Faltkanten sind von Beginn an vorgegeben). Während der Faltung sollen die Polygone starr bleiben, also isometrisch transformiert werden. In meiner Bachelorarbeit [Bau14] wurden solche Muster als Teilmenge des $\mathbb{R}^2$ definiert, die sich dann im $\mathbb{R}^3$ (eingeschränkt) bewegen können. Allerdings enthält dieser Ansatz eine fundamentale Limitierung:

Eine zentrale Modelleigenschaft ist die Durchdringungsfreiheit von Papier. Sobald man aber zwei Polygone zusammenfaltet, belegen diese im Modell dieselbe Punktmenge des $\mathbb{R}^3$. Insbesondere geht die Information, in welcher Reihenfolge diese Polygone aufeinander liegen, verloren. Solange niemals entfaltet wird, kann diese Einschränkung unter Umständen akzeptiert werden. Da wir aber Faltungen allgemein verstehen wollen, unter anderem Flexagone (vgl. [Whe58], [PH97] und [COO57]), ist eine Beschreibung von Entfaltungen zwingend erforderlich.

Bereits in [DO10] wird dieses Problem erkannt und eine Lösung vorgeschlagen. Diese Lösung kann aber weder theoretisch noch praktisch sinnvoll verwendet werden, was auch daran erkennbar ist, dass sie ihren Ansatz im restlichen Buch fast nicht verwenden. Ihr Modell wäre zwar aussagekräftig genug, um das Problem zu lösen, aber da es nicht angewandt werden kann, ist es zur Beantwortung unserer Fragen nicht geeignet. Wir müssen also nach einer anderen Lösung suchen. Im Punkt der Anwendbarkeit kann man die expliziten Darstellungen im $\mathbb{R}^3$ anklagen, die in eigentlich allen bestehenden Arbeiten als grundlegende Modelle angenommen werden. Mit diesen ist es in der Regel sehr aufwendig, Faltbewegungen explizit auszurechnen, selbst wenn man das Muster gut kennt (das liegt im Wesentlichen daran, dass das Problem der globalen Überschneidungsfreiheit bislang noch unverstanden ist).

Um diese Nachteile zu umgehen, ist ein grundlegender Wechsel in der Herangehensweise angebracht: Wir beobachten, dass eine Faltbewegung zwei Faltzustände ineinander überführt (wobei wir verschiedene Faltzustände nur dadurch unterscheiden, dass verschiedene Polygone zusammengefaltet sind). Auch wenn auf die Beschreibung der Bewegung selbst verzichtet wird, bleibt ein abstrakter Zusammenhang zwischen den beiden Faltzuständen bestehen, der von der Faltbewegung induziert wird. Auf Basis dieses Zusammenhangs werden wir unser

Modell aufbauen.

Indem wir darauf verzichten, die Faltbewegung zu beschreiben, sind wir nicht mehr strikt an den $\mathbb{R}^3$ gebunden. Wir können nun feststellen, dass die Frage „Wie kann man dieses Muster falten?“ zu einem großen Teil nur von diesem Muster selbst abhängt und nicht von der konkreten Realisierung im $\mathbb{R}^3$.

1.2 Aufbau dieser Arbeit

In Kapitel 2 werden wir ein grundlegendes mengentheoretisches Modell dieser abstrakten Faltungen entwickeln, indem wir das Konzept simplizialer Flächen geeignet verallgemeinern. Wir werden diese falten, entfalten und auch auf ihre Einbettungen in den dreidimensionalen Raum eingehen. Während dieses Modell relativ einfach und schön strukturiert ist, leidet es grundlegend daran, dass es keine Möglichkeit hat, über die Reihenfolge von Flächen zu sprechen.

In Kapitel 3 geben wir ein wenig Anwendbarkeit auf (die vernünftige Beschreibung von Reihenfolgen stellt sich als nicht trivial heraus) und gewinnen dafür die Ausdrucksschärfe, die dem rein mengentheoretischen Modell fehlte. Wir können dort zum ersten Mal sowohl Faltung als auch Entfaltung präzise beschreiben. Kapitel 3.5 ist einer kompakteren Darstellung von Faltungen und Entfaltungen geschuldet. Dort werden wir einige nützliche Erkenntnisse für die Struktur aller möglichen Faltungen sammeln.

Schließlich kommen wir darauf zurück, dass wir unsere Faltungen auch im $\mathbb{R}^3$ darstellen möchten. Die Frage, ob dies allgemein möglich ist, stellt sich als zu umfangreich heraus. Wir werden sie in Kapitel 4 für einige Spezialfälle untersuchen. Schließlich dient Kapitel 5 der Verfolgung der Frage, inwieweit man unser abstraktes Modell auf allgemeinere Situationen anwenden kann.

Interessant ist, dass man mit diesem Modell sowohl Faltungen im $\mathbb{R}^2$ als auch im $\mathbb{R}^3$ beschreiben kann. Der Grund dafür ist, dass es den Formalismus nicht grundsätzlich ändert, ob man Strecken isometrisch in der Ebene bewegt (um Punkte herum) oder ob man Polygone isometrisch im Raum bewegt (um Faltkanten herum). Es stellt sich heraus, dass sich die meisten Erkenntnisse sogar allgemein im $\mathbb{R}^n$ beschreiben lassen. Davon inspiriert (und weil wir nicht jede Aussage in zwei Versionen präsentieren möchten) konstruieren wir unser Modell in beliebigen Dimensionen, wobei der dreidimensionale Fall stets von primärem Interesse ist. Dadurch erzwingen wir, dass sich alle Beschreibungen und Beweise auf strukturelle Eigenschaften der Faltung stützen müssen, anstatt spezielle Umstände im dreidimensionalen Raum auszunutzen.

1.3 Vorkenntnisse und Notation

Diese Arbeit entwickelt ein Modell zum Falten, das weitestgehend unabhängig von der bestehenden Literatur ist. Daher sind kaum Vorkenntnisse des Lesers erforderlich. Wir setzen voraus, dass der Leser weiß, was Partitionen und Äquivalenzrelationen sind und mit diesen umgehen kann. Kenntnisse simplizialer Komplexe erleichtern die Lektüre, sind aber nicht notwendig für das Verständnis der Inhalte.

Wir bezeichnen mit $\mathbb{N}$ die natürlichen Zahlen, startend bei 1.

In Abschnitt 3.1 benötigen wir außerdem elementare Kenntnisse aus der linearen Algebra (z.B. Orthogonalraum und Determinante) und der Gruppentheorie (z.B. Linksoperationen, Bahnen und Morphismen). In Abschnitt 5.1 benötigen wir Kenntniss der Diedergruppen, sowie (in einem Beweis) von freien Gruppen, Normalteilern und Präsentationen.

Die **symmetrische Gruppe** einer Menge M bezeichnen wir mit

$$\text{Sym}(M) := \{\alpha : M \rightarrow M : \alpha \text{ bijektiv}\}.$$

Im Spezialfall $M = \{1, \dots, n\}$ setzen wir $S_n := \text{Sym}(\{1, \dots, n\})$. Eine **Transposition** ist eine Permutation $\tau \in S_n$, die genau $n - 2$ Fixpunkte hat. Jede Permutation $\pi \in S_n$ kann in ein Produkt von Transpositionen zerlegt werden¹:

$$\pi = \tau_1 \circ \dots \circ \tau_t$$

Das **Signum** (bzw. Vorzeichen) der Permutation π ist dann als $(-1)^t$ definiert. Damit ist ein wohldefinierter Gruppenhomomorphismus festgelegt. Die **alternierende Gruppe** ist dann

$$A_n := \{\pi \in S_n \mid \text{sign}(\pi) = +1\}.$$

¹Diese Aussage ist auch als „Alptraum des Bibliothekars“ bekannt.

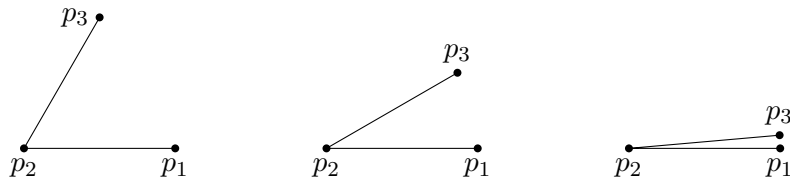
2 Grundlegende Modellierung

In diesem Kapitel beschreiben wir das grundlegende Modell unserer abstrakten (d. h. von einer Einbettung unabhängigen) Faltungen. Dazu brauchen wir ein elementares Verständnis der Faltungen, die wir abstrahieren wollen. Der Einfachheit halber führen wir dies nur für zweidimensionale Faltungen aus, der dreidimensionale Fall geht analog.

Beispiel 2.1. Wir betrachten die einfachste nicht-triviale zweidimensionale Faltung im $\mathbb{R}^2$, die aus drei Punkten $\{p_1, p_2, p_3\}$ und zwei Kanten $\{\overline{p_1 p_2}, \overline{p_2 p_3}\}$ besteht, wobei jede Kante die Länge 1 hat. Legen wir ohne Einschränkung $p_2 = (0, 0)$ fest, dann liegen p_1 und p_3 auf dem Kreis $S^1 := \{(x_1, x_2) \in \mathbb{R}^2 \mid x_1^2 + x_2^2 = 1\}$. Setzen wir $p_1 = (1, 0)$, dann können wir eine glatte Bewegung von p_3 angeben, die alle Punkte des Kreises durchläuft:

$$\alpha : [0, 2\pi] \rightarrow S^1 : \varphi \mapsto (\cos(\varphi), \sin(\varphi))$$

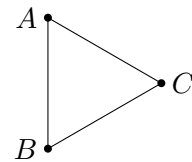
Betrachten wir einige der Bilder dieser Abbildung:



Aus diesen Bildern wird es einsichtig, diese Bewegung als zweidimensionale Faltung aufzufassen, welche sich allerdings selbst durchdringen kann. Insbesondere ist der Abstand zwischen p_1 und p_3 während dieser Faltung nicht konstant.

Die Veränderung des Abstandes zwischen p_1 und p_3 durch eine glatte Abbildung in Beispiel 2.1 ist notwendig für unser Konzept von Faltungen. Als nächstes betrachten wir ein Beispiel, in dem eine solche Veränderung niemals auftreten kann.

Beispiel 2.2. Betrachten wir die Kantenkonfiguration aus der unten stehenden Abbildung. Wenn wir die Längen der Strecken $\overline{AB}$, $\overline{AC}$ und $\overline{BC}$ festlegen (in der Graphik sind sie alle auf denselben Wert gesetzt), ist es nicht möglich, eine Bewegung wie in Beispiel 2.1 auszuführen – einfach dadurch, dass wir die Längen aller vorkommenden Strecken festgelegt haben. Das Muster ist also in dem Sinne **starr**, dass keine Faltungen ausgeführt werden können.

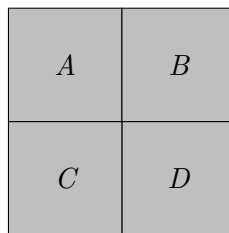


Interessant ist, dass wir die Starrheit dieses Musters (die wir wie in Beispiel 2.2 als das Nicht-Vorhandensein einer Faltbewegung wie in Beispiel 2.1 verstehen) allein daraus ablesen können, wie die Punkte und Kanten zusammenhängen. Das Zusammenfallen der Kanten $\overline{AC}$ und $\overline{BC}$ würde implizieren, dass die Punkte A und B zusammenfallen. Damit hätte die Kante $\overline{AB}$ nach der Faltung die Länge Null, während diese vor der Faltung positiv war. Es handelt sich dabei also um eine sehr eklatante Verletzung der Längenerhaltung. Unser abstrakter Formalismus soll diese extremen Verletzungen erkennen können, aber er wird nicht in der Lage sein, subtilere Verletzungen zu sehen (da er die Längen nicht modellieren wird).

Unser Modell muss also definitiv Punkte und Kanten (sowie im dreidimensionalen Flächen) speichern, ebenso wie deren Inzidenzbeziehung (also welcher Punkt in welcher Kante liegt usw.). Das erinnert an die Strukturen schlichter Graphen (in zwei Dimensionen) und simplizialer Flächen (in drei Dimensionen), also allgemeiner an simpliziale Komplexe. Wir werden aber gleich ein Phänomen untersuchen, dass uns dazu bringt, einen etwas allgemeineren Formalismus zu verwenden.

Bei diesem Phänomen handelt es sich um **erzwungene Faltungen**. Diese treten dann auf, wenn die Faltung von zwei Flächen die Faltung von zwei anderen Flächen nach sich zieht. Anders gesagt bedeutet dies, dass zwei Paare von Flächen nur simultan gefaltet werden können. Wir demonstrieren dieses Phänomen an einem Beispiel.

Beispiel 2.3. *Betrachten wir die folgende Flächenkonfiguration aus vier Quadraten:*

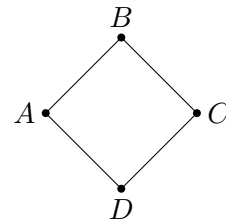


Es ist anschaulich klar: Wenn wir Quadrat A auf Quadrat C falten, müssen wir auch Quadrat B auf Quadrat D falten.

Dieses Phänomen tritt auch schon bei Faltungen in zwei Dimensionen auf.

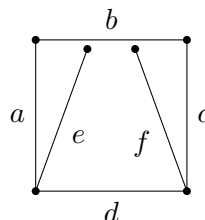
Beispiel 2.4. *Betrachten wir die folgende Kantenkonfiguration:*

Wenn wir die Strecken $\overline{AB}$ und $\overline{AD}$ zusammenfallen, müssen die Punkte B und D zusammenfallen, womit auch die Strecken $\overline{BC}$ und $\overline{CD}$ zusammengefallen werden müssen. Damit erzwingt die Faltung der ersten beiden Strecken die Faltung der letzten beiden.



Eine abstrakte Beschreibung des Faltvorgangs muss mit erzwungenen Faltungen umgehen können. Obwohl die Beispiele 2.3 und 2.4 relativ harmlos wirken, verkomplizieren erzwungene Faltungen die Untersuchung von Faltungen deutlich:

Beispiel 2.5. *Betrachten wir die folgende Kantenkonfiguration:*



Wenn wir die Kanten a und b zusammenfallen wollen (sodass diese nach der Faltung benachbart sind), müssen alle anderen Kanten zusammenfallen (die inneren Kanten müssen zwischen den Kanten c und d bleiben, die ihrerseits wie in Beispiel 2.4 zusammengefaltete werden müssen).

Es ist aber nicht eindeutig, in welcher Reihenfolge die Kanten $\{c, d, e, f\}$ liegen. Die Reihenfolgen $c-e-f-d$ und $c-f-e-d$ sind beide möglich. Wir können diesen Unterschied nicht ignorieren, da Reihenfolgen eine zentrale Rolle in der Modellierung von Faltungen spielen.

Wir können diese Muster aber auch durch eine andere Abfolge von Faltungen erhalten: Wir falten zuerst e und d , sowie f an c , bevor wir a und b falten.

Beispiel 2.5 demonstriert die Probleme, die von erzwungenen Faltungen ausgelöst werden:

- Erzwungene Faltungen führen nicht zu einem eindeutigen Faltzustand (wenn man damit startet, zwei Flächen zusammenzufalten). Dieses Problem verkompliziert auch auf die Untersuchung von Entfaltungen, da es sich in diesen Kontext überträgt.
- In Beispiel 2.5 haben wir gesehen, dass wir einen Übergang zwischen zwei Faltzuständen sowohl durch eine als auch durch drei aufeinanderfolgende Faltungen erreichen können. Da wir nicht eindeutig sagen können, wie viele Faltungen wir benötigen, um von einem Faltzustand zu einem andern überzugehen, werden strukturelle Betrachtungen der Gesamtheit aller Faltungen aufwendiger.

Beide Probleme lassen sich beheben, wenn man annimmt, dass jede Faltung genau zwei Flächen zusammenführt. Unser neuer Ansatz ist es, diesen Spezialfall zum neuen Standard zu machen. Damit das funktioniert, dürfen erzwungene Faltungen nicht ausgeführt werden. Da die Faltung aus Beispiel 2.4 weiterhin legitim sein soll, müssen wir Zwischenzustände erlauben, in denen es zwei verschiedene Kanten mit denselben Eckpunkten gibt. Das ist eine signifikante Einsicht, da sie mit der Intuition bricht, dass die Kanten und Flächen einer Faltung *gerade* (im Unterschied zu *krummlinig*) sein müssen. Da dies auch die Intuition ist, die uns zu Beginn die Modellierung als simpliziale Komplexe nahegelegt hat, schließen wir nun, dass simpliziale Komplexe zur Modellierung von Faltzuständen ungeeignet sind.

Da sich die Zwischenzustände von den ursprünglichen Zuständen unterscheiden lassen, lässt sich die Frage nach den erzwungenen Faltungen wie folgt umformulieren: In welchen „echten“ Faltzuständen ist die Faltung ausgeführt, die die erzwungenen Faltungen auslöst? Da wir also die gleichen Fragen beantworten können, verlieren wir durch die Einführung dieser Zwischenzustände keine Aussagekraft.

Auf Basis dieser Überlegungen ist dieses Kapitel wie folgt aufgebaut: In Abschnitt 2.1 werden wir die angesprochenen Zwischenzustände definieren. Mit diesen können wir zwar den Faltzustand beschreiben, aber sie sind zur Beschreibung des Faltvorgangs nur schlecht geeignet. Um den Faltvorgang rudimentär einzufangen, werden wir in Abschnitt 2.2 eine mengentheoretische Erweiterung dieser Zustände betrachten. Die Reihenfolge der Flächen in der Faltung können wir damit noch nicht beschreiben, aber wir erhalten eine schöne Theorie für das Zusammenfallen (Abschnitt 2.3) und das Entfallen (Abschnitt 2.4). Vor allem während der Klassifikation der Entfaltungen wird uns aber vor Augen geführt, dass die Einbeziehung der Reihenfolgen von großer Wichtigkeit ist.

Im Anschluss an diese abstrakten Konstruktionen beginnen wir mit der Betrachtung des Einbettungsproblems. In Abschnitt 2.5 werden wir zuerst klären, was es bedeutet, eine gefaltete Fläche in einen $\mathbb{R}^d$ einzubetten. Wir möchten aber eigentlich eine Einbettung finden, in der sich die Faltungen des Komplexes in Faltbewegungen übersetzen. Dabei handelt es sich um ein sehr schweres Problem, weshalb wir uns hier mit einigen prinzipiellen Aussagen zufrieden geben müssen. In diesem Kontext ist es natürlich relevant zu wissen, welche Muster man nicht weiter zusammenfalten kann, da diese in der Regel die minimale Zahl von Freiheitsgraden für unsere Einbettung vorgeben. Wir werden diese minimalen Komplexe in Abschnitt 2.5.1 durch einen Färbungsbegriff klassifizieren.

2.1 Homogene Komplexe

Nachdem wir im letzten Abschnitt die Probleme mit simplizialen Komplexen aufgelistet haben, führen wir in diesem Abschnitt die in der Einleitung angesprochene Verallgemeinerung aus. Da wir unser Modell nicht nur auf Dreiecke beschränken möchten (wir wollen später z. B. auch mit quadratischen Flächen arbeiten können), legen wir die Form der Flächen noch nicht direkt fest. Wir werden in Abschnitt 2.1.1 erläutern, wie sich unser Modell auf den Fall von Dreiecken spezialisiert. In Abschnitt 5.1 werden wir über die Modifikationen für allgemeine m -Ecke sprechen.

Wir müssen uns also überlegen, welche Eigenschaften wir benötigen, um eine Faltstruktur darzustellen, die aus Punkten, Kanten und Flächen besteht. Grundlegend ist natürlich, dass es eine Punktmenge $\mathcal{K}_0$, eine Kantenmenge $\mathcal{K}_1$ und eine Flächenmenge $\mathcal{K}_2$ gibt (der Index bezeichnet dabei die Dimension der entsprechenden Elemente). Wir brauchen auch eine Relation, die aussagt, welcher Punkt auf welcher Kante und welche Kante auf welcher Fläche liegt. Diese Relation muss transitiv sein, weswegen sie eine Teilmenge von

$$(\mathcal{K}_0 \times \mathcal{K}_1) \uplus (\mathcal{K}_0 \times \mathcal{K}_2) \uplus (\mathcal{K}_1 \times \mathcal{K}_2)$$

ist. Damit haben wir aber noch kaum Einschränkungen, denn es könnte z. B. eine Kante geben, die keinen Punkt enthält und in keiner Fläche liegt. Während wir nicht im Detail fordern wollen, wie viele Punkte eine Fläche enthält, ist es sinnvoll, eine gewisse Regularität vorauszusetzen.

Definition 2.6 (homogener Komplex). *Eine Menge $\mathcal{K}$ sei eine disjunkte Vereinigung $\mathcal{K} = \mathcal{K}_0 \uplus \dots \uplus \mathcal{K}_n$ mit $\mathcal{K}_n \neq \emptyset$ und $\prec$ eine transitive Relation auf $\mathcal{K}$, die Teilmenge von $\biguplus_{i < j} \mathcal{K}_i \times \mathcal{K}_j$ ist.*

*Dann ist $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec)$ ein **homogener n -Komplex**, falls gilt:*

- *Zu jedem $x \in \mathcal{K}_i$ (mit $i < n$) gibt es ein $y \in \mathcal{K}_{i+1}$ mit $x \prec y$.*
- *Zu jedem $y \in \mathcal{K}_j$ (mit $j > 0$) gibt es ein $x \in \mathcal{K}_{j-1}$ mit $x \prec y$.*

*Die Elemente in $\mathcal{K}_i$ werden **von Dimension i** genannt, im Speziellen heißen die Elemente von $\mathcal{K}_0$ Punkte, die von $\mathcal{K}_1$ Kanten und die von $\mathcal{K}_2$ Flächen. Wir nennen n die **Dimension** des Komplexes.*

Da wir mit dieser Definition formalisieren wollen, wie Punkte, Kanten und Flächen geometrisch zusammenhängen, interpretieren wir $x \prec y$ so, dass x in y liegt. Man kann auch sagen, dass x in y enthalten ist oder dass y das Element x enthält.

Die Regularitätsforderung der Definition erzwingt bei einem homogenen 1-Komplex, dass jeder Punkt in einer Kante liegen und jede Kante einen Punkt enthalten muss. Für jede Kante eines homogenen 2-Komplexes muss es einen Punkt geben, der auf der Kante liegt, und eine Fläche, in der die Kante enthalten ist. Allgemeiner können wir für jedes Element eines homogenen 2-Komplexes eine Kette von Punkt–Kante–Fläche konstruieren, wobei der Punkt in der Kante und die Kante in der Fläche liegt:

Folgerung 2.7. *Sei $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec)$ ein homogener n -Komplex, dann gilt:*

1. *Zu jedem $x_j \in \mathcal{K}_j$ gibt es $x_i \in \mathcal{K}_i$ (für $i \neq j$), sodass es eine Kette*

$$x_0 \prec x_1 \prec \cdots \prec x_j \prec \cdots \prec x_n$$

gibt.

2. *Die Länge jeder solchen Kette ist um eins größer als die Dimension des Komplexes.*

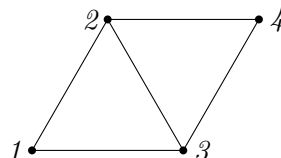
Gemäß Folgerung 2.7 können wir die Dimension des homogenen n -Komplexes bereits aus der Relation $\prec$ ablesen. Die Ketten aus der Folgerung sind nämlich in dem Sinn maximal, dass es zu $x_i \prec x_{i+1}$ kein z mit $x_i \prec z \prec x_{i+1}$ gibt. Es genügt also, eine maximale solche Kette zu bestimmen, um auf die Dimension des homogenen n -Komplexes schließen zu können.

Auch wenn unsere Definitionen für beliebige Dimensionen gelten werden, sind wir primär an den Faltungen von Papier interessiert. Dabei werden Flächen gefaltet, sodass wir uns im Fall $n = 2$ befinden. Der Fall $n = 1$ (die Faltung von Kanten in der Ebene) ist hauptsächlich zur Illustration von Interesse, da man an diesem schon viele Phänomene erkennen kann, die auch bei $n = 2$ auftreten.

Die Komplexe, an denen wir interessiert sind, bestehen in der Regel aus endlichen Mengen. In Abschnitt 4.2.2 werden wir uns aber mit der Faltung unendlicher Komplexe befassen, weswegen wir die Endlichkeit nicht in die Definition aufgenommen haben.

Da die Definition der homogenen n -Komplexe sehr allgemein ist, umfasst sie auch einige entartete Fälle, von denen einige in Beispiel 2.9 dargestellt sind. In Abschnitt 2.1.1 werden wir die Definition so einschränken, dass diese Fälle nicht mehr vorkommen können. Bevor wir uns aber mit den entarteten Fällen beschäftigen, listen wir in Beispiel 2.8 einige Standardfälle von homogenen n -Komplexen auf.

Beispiel 2.8. 1. *Dieses Bild stellt einen homogenen 1-Komplex dar, wobei wir lediglich die Punkte und Kanten, nicht jedoch die Flächen betrachten (dass wir die Flächen nicht betrachten, verdeutlichen wir dadurch, dass wir diesen Bereich weiß lassen):*



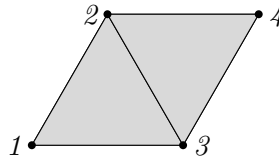
Formal ist der 1-Komplex durch die Mengen

$$\mathcal{K}_0 := \{\{1\}, \{2\}, \{3\}, \{4\}\}$$

$$\mathcal{K}_1 := \{\{1, 2\}, \{1, 3\}, \{2, 3\}, \{2, 4\}, \{3, 4\}\}$$

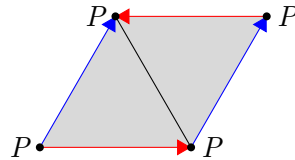
und die Teilmengenrelation gegeben.

2. Fügt man die Flächen hinzu, so erhält man einen homogenen 2-Komplex. Dazu ergänzt man $\mathcal{K}_2 := \{\{1, 2, 3\}, \{2, 3, 4\}\}$:



Es ist nicht möglich, $\mathcal{K}_2 := \{\{1, 2, 3\}\}$ zu definieren, da dann $\{2, 4\} \in \mathcal{K}_1$ in keinem Element aus $\mathcal{K}_2$ enthalten wäre.

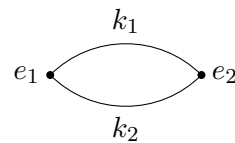
3. Wir können eine *Kleinsche Flasche*² als 2-Komplex mit zwei Flächen darstellen, indem wir die Kanten geeignet identifizieren (parallele Pfeile werden miteinander identifiziert):



4. In den Fällen (1) und (2) bestand $\mathcal{K}_i$ aus i -elementigen Teilmengen einer festen Menge. Während die *Kleinsche Flasche* aus Fall (3) das nicht zulässt (dann wären alle ihre Kanten gleich), ist diese ein eher degeneriertes Beispiel, mit dem wir nicht viel falten werden. Das folgende Beispiel eines homogenen 1-Komplexes leidet nicht unter diesem Manko, da es tatsächlich häufig in der Berechnung von Faltungen auftaucht (z. B. in Beispiel 2.4):

$$\mathcal{K}_0 := \{e_1, e_2\}$$

$$\mathcal{K}_1 := \{k_1, k_2\}$$



mit $\prec := \mathcal{K}_0 \times \mathcal{K}_1$.

Nach diesen Standardbeispielen ist hier eine mögliche Entartung der Definition aufgezeigt:

Beispiel 2.9. Sei eine Menge $\mathcal{K}$ wie in Definition 2.6 gegeben. Dann ist

$$\left((\mathcal{K}_i)_{0 \leq i \leq n}, \biguplus_{i < j} \mathcal{K}_i \times \mathcal{K}_j \right)$$

²Eine der nicht-orientierbaren geschlossenen Flächen.

ein homogener n -Komplex.

Wenn man hier z. B. $\mathcal{K}_0 := \{1, 2, 3\}$ und $\mathcal{K}_1 := \{a\}$ wählt, sieht man, dass die Kante a mehr als zwei Punkte enthält, nämlich jeden Punkt in $\mathcal{K}_0$.

Da es sich bei homogenen n -Komplexen um den Grundstein des gesamten Formalismus handelt, wird es häufig nötig sein, verschiedene Komplexe miteinander zu vergleichen. Diese Vergleiche lassen sich durch Morphismen beschreiben, insbesondere durch Isomorphismen.

Definition 2.10 (Komplex-Morphismus). Für homogene n -Komplexe³ $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec^{\mathcal{S}})$ und $\mathcal{T} = ((\mathcal{L}_i)_{0 \leq i \leq n}, \prec^{\mathcal{T}})$ ist ein **Komplex-Morphismus** $\varphi : \mathcal{S} \rightarrow \mathcal{T}$ eine Familie von Abbildungen

$$(\varphi_i : \mathcal{K}_i \rightarrow \mathcal{L}_i)_{0 \leq i \leq n},$$

sodass $x \prec^{\mathcal{S}} y$ impliziert, dass $\varphi(x) \prec^{\mathcal{T}} \varphi(y)$ gilt.

Falls alle φ_i injektiv sind, heißt φ **Komplex-Monomorphismus**.

Falls alle φ_i surjektiv sind, heißt φ **Komplex-Epimorphismus**.

Falls es einen Komplex-Morphismus $\psi : \mathcal{T} \rightarrow \mathcal{S}$ mit $\varphi \circ \psi = \text{id}_{\mathcal{T}}$ und $\psi \circ \varphi = \text{id}_{\mathcal{S}}$ gibt, heißt φ **Komplex-Isomorphismus**.

Die Definition eines Komplex-Isomorphismus unterscheidet sich aus gutem Grund von denen von Komplex-Monomorphismus und Komplex-Epimorphismus. Nur zu fordern, dass jedes φ_i bijektiv ist, wäre nämlich nicht genug, wie Beispiel 2.11 demonstriert.

Beispiel 2.11. Betrachten wir die beiden homogenen 1-Komplexe



aus den Mengen

$$\begin{aligned}\mathcal{K}_0 &:= \{A, B\} \\ \mathcal{K}_1 &:= \{a, b\}\end{aligned}$$

$$\begin{aligned}\mathcal{L}_0 &:= \{A, B\} \\ \mathcal{L}_1 &:= \{a, b\}\end{aligned}$$

und den aus den Abbildungen ersichtlichen $\prec$ -Relationen bestehend. Dann bilden die Abbildungen $\varphi_i : \mathcal{K}_i \rightarrow \mathcal{L}_i, x \mapsto x$ für $i \in \{1, 2\}$ einen Komplex-Morphismus. Obwohl beide φ_i bijektiv sind, würde man diese beiden Komplexe niemals als isomorph ansehen.

Dieses Beispiel beruht auf den möglichen Entartungen unserer Definition von homogenen n -Komplexen. Wir können die Definition eines Komplex-Isomorphismus dennoch direkt in Bezug auf die Abbildungen φ_i formulieren. Das ist hilfreich, da wir häufig nachweisen müssen, dass eine Abbildung ein Komplex-Isomorphismus ist.

Bemerkung 2.12. Seien $\mathcal{S}, \mathcal{T}$ zwei homogene n -Komplexe. Ein Komplex-Morphismus $\varphi : \mathcal{S} \rightarrow \mathcal{T}$ ist genau dann ein Komplex-Isomorphismus, wenn gilt:

1. Jedes φ_i ist bijektiv.
2. $x \prec^{\mathcal{S}} y$ ist äquivalent zu $\varphi(x) \prec^{\mathcal{T}} \varphi(y)$.

³Es würde an dieser Stelle genügen, dass $\mathcal{T}$ ein m -Komplex mit $m \geq n$ ist. Dann wären die φ_i mit $i > n$ leere Abbildungen.

2.1.1 Lokale Simplizität

Bislang haben wir homogene n -Komplexe nur in völliger Allgemeinheit studiert und festgestellt, dass dadurch einige Entartungen auftreten. In diesem Abschnitt wird beschrieben, wie sich der Formalismus einschränkt, wenn jede Fläche ein Dreieck sein muss. Dadurch wird die Entartung aus Beispiel 2.9 vermieden.

Wir müssen dazu festhalten, was es bedeutet, dass eine Fläche ein Dreieck ist. Man kann dies dadurch zu beschreiben versuchen, dass jede Fläche aus drei Kanten besteht und jede Kante aus zwei Punkten, aber dann muss man noch viele Interaktionsbedingungen formulieren, um sicherzustellen, dass tatsächlich ein Dreieck beschrieben wird. Dieses Problem wird in höheren Dimensionen noch umfangreicher, weswegen dieser Ansatz in einem dimensionsunabhängigen Rahmen zu vermeiden ist.

Stattdessen sehen wir Kanten und Dreiecke als spezielle Simplizes an. Im dreidimensionalen Raum ist eine Kante (oder 1-Simplex) die konvexe Hülle von zwei verschiedenen Punkten, während ein Dreieck (oder 2-Simplex) die konvexe Hülle von drei Punkten ist, die nicht auf einer Geraden liegen. Ein 3-Simplex ist dann die konvexe Hülle von vier Punkten, die nicht alle in der selben Ebene liegen. Solche 3-Simplizes würden bei Faltung im vierdimensionalen Raum eine zentrale Rolle spielen. Wenn wir von den konkreten Positionen der Punkte im Raum zu einer abstrakten Betrachtung von diesen übergehen, überträgt sich das Konzept der konvexen Hülle dahingehend, dass zwischen je zwei Punkten eine Kante und je drei Punkten ein Dreieck existiert (usw.).

Definition 2.13 (Simplex). *Sei M eine endliche Menge mit $|M| = n + 1$. Dann nennen wir den homogenen n -Komplex $((\text{Pot}_{i+1}(M))_{0 \leq i \leq n}, \subsetneq)$ einen n -**Simplex**. Dabei bezeichnet $\text{Pot}_{i+1}(M)$ die Menge der $(i + 1)$ -elementigen Teilmengen von M .*

Dass jede Fläche eines homogenen 2-Komplexes ein Dreieck ist, ist eine lokale Bedingung, d. h. sie ist nur von dieser Fläche und allen Punkten und Kanten *in* dieser Fläche abhängig. Wir müssen also zu $x \in \mathcal{K}_2$ nur die Elemente y betrachten, die $y \preccurlyeq x$ erfüllen. Die Betrachtung eines solchen Teilkomplexes können wir allgemein definieren:

Definition 2.14 (Einschränkung). *Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec)$ ein homogener n -Komplex und $x \in \mathcal{K}_i$. Dann ist die **Einschränkung von $\mathcal{S}$ auf x** , genannt $\mathcal{S}_{\preccurlyeq x}$, definiert als der homogene i -Komplex $(\{y \in \mathcal{K}_j \mid y \preccurlyeq x\})_{0 \leq j \leq i}$, zusammen mit der Einschränkung von $\prec$ auf diese Teilmenge⁴.*

Damit können wir beschreiben, dass ein homogener n -Komplex lokal die Gestalt eines Simplex hat. Wenn jede Fläche $x \in \mathcal{K}_2$ ein Dreieck sein soll, muss die Einschränkung des Komplexes auf x ein 2-Simplex sein.

Definition 2.15 (lokal simplizial). *Ein homogener n -Komplex $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec)$ heißt **lokal simplizial**, wenn die Einschränkung $\mathcal{S}_{\preccurlyeq y}$ für alle $y \in \mathcal{K}_n$ isomorph zu einem n -Simplex ist.*

Mit dieser Modifikation verschwindet die Entartung aus Beispiel 2.9.

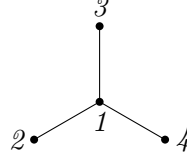
⁴Da $x \in \mathcal{K}_i$, ist $\{y \in \mathcal{K}_i \mid y \preccurlyeq x\} = \{x\}$ nicht leer.

Beispiel 2.16. 1. Die homogenen Komplexe aus den Fällen (1), (2) und (4) in Beispiel 2.8 sind lokal simplizial. Die Kleinsche Flasche aus Fall (3) ist nicht lokal simplizial, da jede Kante nur einen Punkt enthält – nicht die zwei, die wir erwarten würden, wenn es sich um einen 1-Simplex handeln würde.

2. An einem Punkt können drei Kanten zusammenstoßen, wie der homogene 1-Komplex

$$\begin{aligned}\mathcal{K}_0 &:= \{\{1\}, \{2\}, \{3\}, \{4\}\} \\ \mathcal{K}_1 &:= \{\{1, 2\}, \{1, 3\}, \{1, 4\}\}\end{aligned}$$

(mit Teilmengenrelation) demonstriert.



Dieses Beispiel zeigt einen 1-Komplex, der keine Kurve ist. Ebenso gibt es 2-Komplexe, die keine Flächen sind. Daher sind wir in einer Situation, die nicht mehr als Mannigfaltigkeit dargestellt werden kann. Die Ideen aus der diskreten Differentialgeometrie können also nicht direkt auf Faltungen übertragen werden.

2.2 Überlagerungskomplexe

In Abschnitt 2.1 haben wir ein Modell für unsere Faltzustände entwickelt, die lokal simplizialen n -Komplexe. Diese erlauben aber noch keine Faltungen, da es nicht möglich ist, zusammengefaltete Flächen zu erkennen. Wenn man nur die Mittel des vorherigen Abschnitts 2.1 verwendet, kann man versuchen, die Faltung durch einen Komplex-Epimorphismus $\mathcal{S} \rightarrow \mathcal{T}$ zu beschreiben. Allerdings ist es dann keine Entfaltung möglich, wenn nur $\mathcal{T}$ bekannt ist. Das widerspricht unserer Intuition, dass wir allein aus einem Faltzustand schließen können, wie sich dieser entfalten lässt.

Betrachten wir also wieder reale Faltungen aus Papier. Bei diesen wird der Anfangszustand der Faltung anders angeordnet, aber nicht grundlegend verändert (wir schneiden keine Dreiecke aus oder kleben Kanten zusammen). Daher werden wir im Folgenden den homogenen n -Komplex während der Faltung als konstant ansehen. Aber dann benötigen wir eine zusätzliche Struktur, um die Faltung dieses Komplexes zu modellieren.

Eine zentrale Eigenschaft zweier zusammengefalteter Flächen ist, dass diese Flächen an derselben Stelle im Raum liegen. Für einen gegebenen homogenen 2-Komplex $((\mathcal{K}_0, \mathcal{K}_1, \mathcal{K}_2), \prec)$ können wir dies als eine Abbildung

$$\iota : \mathcal{K}_2 \rightarrow \text{Pot}(\mathbb{R}^3)$$

modellieren, sodass zwei Flächen $f, g \in \mathcal{K}_2$ genau dann zusammengefalteter sind, wenn $\iota(f) = \iota(g)$ gilt. Folglich haben wir eine Äquivalenzrelation auf den Flächen $\mathcal{K}_2$ vorliegen, die wir informell als „zusammengefalteter sein“ beschreiben können.

Wenn zwei Flächen zusammengefalteter sind, wissen wir aber noch nicht, welche Kanten und Punkte dieser Flächen zusammengefalteter sind. Daher definieren wir auch auf den Kanten und Punkten die Äquivalenzrelation „zusammengefalteter sein“. Damit können wir eindeutig spezifizieren, wie zwei Flächen zusammengefalteter werden.

Allerdings stellen wir fest, dass es für die Äquivalenzrelation keinen Unterschied macht, ob sie auf Punkten, Kanten oder Flächen definiert ist. Sie verhält sich immer gleich und kann daher keine dimensionsspezifischen Eigenschaften beschreiben. Da sie daher die Reihenfolge der Flächen nicht abbilden kann, ist diese Beschreibung notwendig unvollständig. Dennoch beschäftigen wir uns ausführlich mit diesem Modell, da es eine schöne Theorie ermöglicht und sich das spezifischere Modell (das wir in Kapitel 3 entwickeln werden) als Modifikation dieser Situation verstehen lässt.

Definition 2.17 (Überlagerungskomplex). *Sei ein homogener n -Komplex $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec)$ zusammen mit einer Äquivalenzrelation $\sim$ auf $\bigcup_{i=0}^n \mathcal{K}_i$, die die folgenden Eigenschaften für alle $x, y \in \bigcup_{i=0}^n \mathcal{K}_i$ erfüllt, gegeben:*

1. (**Dimensionalität**) *Wenn $x \sim y$ gilt, gibt es ein $0 \leq i \leq n$ mit $x \in \mathcal{K}_i$ und $y \in \mathcal{K}_i$.*
2. (**Abwärtskompatibilität**) *Falls $x \sim y$ gilt, dann existiert für jedes $a \prec x$ ein $b \prec y$ mit $a \sim b$.*
3. (**Irreduzibilität**) *Falls $a \prec x$ und $b \prec x$ gelten, so ist entweder $a = b$ oder $a \not\sim b$.*

Dann nennen wir $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ einen n -**Überlagerungskomplex**.

Die Äquivalenzklasse von x bezeichnen wir mit $[x]_\sim$ oder $[x]_S$, sowie $[x]$, wenn der Kontext klar ist. Wir schreiben $[x] \prec [y]$, falls es ein $a \sim x$ und ein $b \sim y$ mit $a \prec b$ gibt.

Wir bemerken, dass wir die Äquivalenzrelation $\sim$, die wir als „zusammengefasst sein“ interpretieren, auf $\bigcup_{i=0}^n \mathcal{K}_i$ anstatt auf jedem $\mathcal{K}_i$ einzeln definiert haben. Das hat es uns z. B. ermöglicht, sowohl Abwärtskompatibilität als auch Irreduzibilität ohne die Verwendung vieler Indizes zu formulieren.

Wir bemerken auch, dass wir in der Definition nicht gefordert haben, dass der homogene n -Komplex lokal simplizial ist. Auch wenn das eine Annahme ist, die wir an die meisten unserer Beispiele stellen, ist sie nicht universell. Insbesondere in Abschnitt 5.1 verallgemeinern wir diese Bedingung, um z. B. über quadratische Flächen sprechen zu können. Um klarer herauszustellen, an welchen Stellen die lokale Simplizität formal benötigt wird, haben wir sie an dieser Stelle nicht in die Definition aufgenommen. In Kapitel 5 wird dieses Vorgehen es uns einfacher machen, zu erkennen, welche Aussagen wir verallgemeinern müssen.

Um nachzuweisen, dass ein n -Überlagerungskomplex vorliegt, ist es gelegentlich hilfreicher, die Abwärtskompatibilität und die Irreduzibilität simultan nachzuweisen.

Bemerkung 2.18 (starke Abwärtskompatibilität). *Die Bedingungen (2) und (3) in Definition 2.17 sind äquivalent zur **starken Abwärtskompatibilität**:*

- *Für alle $x \sim y$ (mit $x, y \in \bigcup_{i=0}^n \mathcal{K}_i$) mit $a \prec x$ gibt es genau ein $b \prec y$ mit $a \sim b$.*

Beweis. Die Irreduzibilität folgt aus der starken Abwärtskompatibilität, wenn man $x = y$ wählt. $\square$

Als nächstes stellt sich die Frage, wie man die Faltung von n -Überlagerungskomplexen beschreibt. Nach unserer Modellierung ändert sich der zugrundeliegende homogene n -Komplex

nicht. Der einzige Unterschied liegt also in den Äquivalenzrelationen. Betrachten wir also zwei n -Überlagerungskomplexe $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec^{\mathcal{S}}, \sim^{\mathcal{S}})$ und $\mathcal{T} = ((\mathcal{L}_i)_{0 \leq i \leq n}, \prec^{\mathcal{T}}, \sim^{\mathcal{T}})$. Wir wollen beschreiben, dass $\mathcal{T}$ eine Faltung von $\mathcal{S}$ ist. Da bei einer Faltung keine bereits zusammengefalteten Elemente wieder entfaltet werden, muss aus $x \sim^{\mathcal{S}} y$ bereits $x \sim^{\mathcal{T}} y$ folgen. Wir sagen dann, dass $\mathcal{S}$ eine Teilüberlagerung von $\mathcal{T}$ ist.

Definition 2.19 (Teilüberlagerung). Für n -Überlagerungskomplexe $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim^{\mathcal{S}})$ und $\mathcal{T} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim^{\mathcal{T}})$ über demselben homogenen n -Komplex $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec)$ heißt $\mathcal{S}$ **Teilüberlagerung** von $\mathcal{T}$, falls für alle $x, y \in \bigsqcup_{i=0}^n \mathcal{K}_i$ aus $x \sim^{\mathcal{S}} y$ bereits $x \sim^{\mathcal{T}} y$ folgt.

Es kann natürlich sein, dass es $x, y \in \bigsqcup_{i=0}^n \mathcal{K}_i$ gibt, für die $x \not\sim^{\mathcal{S}} y$ gilt (d. h. x und y sind in $\mathcal{S}$ nicht zusammengefasst), aber $x \sim^{\mathcal{T}} y$ (d. h. x und y sind in $\mathcal{T}$ zusammengefasst). Das sind genau die Situationen, in denen $\mathcal{T}$ aus Faltung durch $\mathcal{S}$ hervorgeht. Diese Definition erlaubt es uns aber zunächst nur zu überprüfen, ob eine notwendige Bedingung für eine Faltung vorliegt. Wir werden in Abschnitt 2.3 sehen, dass diese Bedingung auch hinreichend für die Existenz einer Faltung ist, indem wir explizit konstruieren durch welche Faltungen wir von $\mathcal{S}$ zu $\mathcal{T}$ gelangen. Die Ergebnisse dieses Abschnitts können dann auch algorithmisch verwendet werden.

2.3 Zusammenfalten

Nachdem wir in Abschnitt 2.2 die homogenen n -Komplexe um eine Äquivalenzrelation erweitert haben und ein Konzept von Faltung in Definition 2.19 (Teilüberlagerung) vorgelegt haben, beschäftigen wir uns jetzt intensiver mit dem Zusammenfalten. Unser Hauptziel ist die direkte Konstruktion von Faltungen. Dazu betrachten wir wieder Papierfaltungen. Um dort zwei Flächen zusammenzufalten, führen wir die folgenden Schritte aus:

1. Wir wählen zwei Flächen des Modells, die noch nicht zusammengefasst sind.
2. Wir entscheiden uns, wie wir diese Flächen zusammenfügen wollen.
3. Wir fügen die Flächen zusammen.

Auf der formalen Seite stellen wir das Papiermodell durch einen 2-Überlagerungskomplex $\mathcal{S} := ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ dar. Die Schritte aus dem Papiermodell übersetzen sich wie folgt:

1. Wir wählen zwei $x, y \in \mathcal{K}_2$ mit $x \not\sim y$.
2. Wir wählen einen Komplex-Isomorphismus $\varphi : \mathcal{S}_{\prec x} \rightarrow \mathcal{S}_{\prec y}$.
3. Wir fügen $a \sim \varphi(a)$ für jedes $a \prec x$ zur Äquivalenzrelation $\sim$ hinzu.

Falls dieses Verfahren wieder einen 2-Überlagerungskomplex liefert, wird dadurch eine Teilüberlagerung induziert. Wenn wir den abstrakten Ablauf des Verfahrens studieren, fällt auf, dass wir im ersten Schritt auch $x, y \in \mathcal{K}_1$ oder sogar $x, y \in \mathcal{K}_0$ hätten wählen können, ohne die nachfolgenden Schritte zu gefährden. Dem würde im Papiermodell das Zusammenfalten von zwei Kanten bzw. Punkten entsprechen. Da wir in Abschnitt 2.3.1 ausführlich von

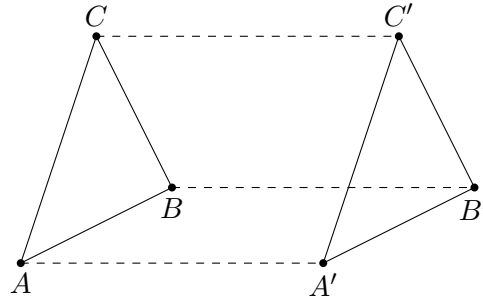
diesen Faltungen Gebrauch machen werden, formulieren wir unsere Definitionen für diesen allgemeineren Fall.

Da diese Konstruktion eine zentrale Rolle spielt, geben wir dem verwendeten Isomorphismus einen eigenen Namen:

Definition 2.20 (Identifikation). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein n -Überlagerungskomplex. Weiter seien $f, g \in \mathcal{K}_j$ (mit $0 \leq j \leq n$) gegeben. Ein Komplex-Isomorphismus $\varphi : \mathcal{S}_{\prec f} \rightarrow \mathcal{S}_{\prec g}$ heißt **Identifikation (vom Grad j) auf $\mathcal{S}$** .

Das Hinzufügen von $a \sim \varphi(a)$ zur Äquivalenzrelation $\sim$ lässt sich theoretisch sehr einfach formulieren: Man fügt das Paar $(a, \varphi(a))$ zur Relation hinzu und nimmt den transitiven und symmetrischen Abschluss.

Da wir aber auch an einer Implementierung dieses Prozesses interessiert sind, suchen wir nach einer konkreteren Beschreibung. Betrachten wir das Zusammenfalten von zwei Dreiecken aus Papier, stellen wir fest, dass genau zwei Flächen zusammengefaltet werden. Wir müssen noch herausfinden, wie viele Kanten und Punkte zusammengefaltet werden. Dazu betrachten wir die folgende Graphik, die das Zusammenfalten der Dreiecke $\{A, B, C\}$ und $\{A', B', C'\}$ illustriert.



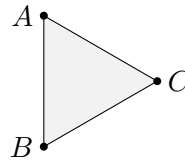
Dabei werden die Paare $\{A, A'\}$, $\{B, B'\}$ und $\{C, C'\}$ jeweils zusammengefaltet. Es werden also jeweils höchstens zwei Punktklassen vereinigt (falls $A \sim A'$ gilt, gehören A und A' schon vor der Faltung zur selben Punktklasse). Man überzeugt sich leicht, dass auch jeweils höchstens zwei Kantenklassen vereinigt werden. Insgesamt werden bei einer Faltung aus Papier also jeweils höchstens zwei Äquivalenzklassen vereinigt. Das ist in der bisherigen Formulierung aber nicht gewährleistet, wie Beispiel 2.21 zeigt.

Beispiel 2.21. Wir betrachten den lokal simplizialen 2-Überlagerungskomplex $\mathcal{S}$:

$$\mathcal{K}_0 := \{\{A\}, \{B\}, \{C\}\}$$

$$\mathcal{K}_1 := \{\{A, B\}, \{A, C\}, \{B, C\}\}$$

$$\mathcal{K}_2 := \{\{A, B, C\}\}$$



mit $\prec := \subseteq$ und $\sim$ als Gleichheit.

Die Abbildung $\varphi : \mathcal{S} \rightarrow \mathcal{S}$ mit $\varphi(A) = B$, $\varphi(B) = C$ und $\varphi(C) = A$ ist eine Identifikation vom Grad 2, da $\mathcal{S} = \mathcal{S}_{\prec \{A, B, C\}}$ gilt. Das Hinzufügen dieser Relationen zu $\sim$ würde dazu führen, dass die drei einelementigen Äquivalenzklassen auf $\mathcal{K}_0$ zusammenfallen würden.

Wir sehen aber auch, dass wir in diesem Beispiel keinen n -Überlagerungskomplex erhalten, da wir die Irreduzibilität der Äquivalenzrelation brechen. Da wir sowieso nur an solchen Identifikationen interessiert sind, die wieder zu n -Überlagerungskomplexen führen, verbieten wir die einfachsten solcher Regelverletzungen, indem wir fordern, dass die Identifikation bereits zusammengefaltete Elemente aufeinander abbilden muss. Falls wir also eine Identifikation $\varphi : \mathcal{S}_{\preccurlyeq f} \rightarrow \mathcal{S}_{\preccurlyeq g}$ vorliegen haben, muss für alle $a \prec f$ und $b \prec g$, die $a \sim b$ erfüllen, schon $\varphi(a) = b$ gelten.

Definition 2.22 (konstant auf dem Schnitt). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein n -Überlagerungskomplex. Eine Identifikation $\varphi : \mathcal{S}_{\preccurlyeq f} \rightarrow \mathcal{S}_{\preccurlyeq g}$ heißt **konstant auf dem Schnitt**, falls für $a \prec f, b \prec g$ mit $a \sim b$ schon $\varphi(a) = b$ folgt (dabei sind $a, b \in \bigsqcup_{i=0}^n \mathcal{K}_i$).

Mit dieser Einschränkung können wir die neuen Äquivalenzklassen explizit angeben. Die Überlegung, dass jeweils höchstens zwei Klassen vereinigt werden, stellt sich als zutreffend heraus.

Definition 2.23 (Erweiterung). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein n -Überlagerungskomplex. Weiter seien $f, g \in \mathcal{K}_j$ (mit $0 \leq j \leq n$) gegeben, zusammen mit einer Identifikation $\varphi : \mathcal{S}_{\preccurlyeq f} \rightarrow \mathcal{S}_{\preccurlyeq g}$, die konstant auf dem Schnitt ist.

Die (**formale**) **Erweiterung von $\sim$ um φ** ist die Äquivalenzrelation $\sim_\varphi \subseteq \bigcup_{0 \leq i \leq n} \mathcal{K}_i \times \mathcal{K}_i$, bei der $x \sim_\varphi y$ genau dann gilt, wenn mindestens eine dieser Bedingungen erfüllt ist:

- $x \sim y$
- Es gibt ein $z \preccurlyeq f$, sodass $x \sim z$ und $\varphi(z) \sim y$ gilt.
- Es gibt ein $z \preccurlyeq g$, sodass $x \sim z$ und $\varphi^{-1}(z) \sim y$ gilt.

Falls $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim_\varphi)$ ein n -Überlagerungskomplex ist, heißt er **Erweiterung von $\mathcal{S}$ um φ** .

Wohldefiniertheit. Wir müssen zeigen, dass $\sim_\varphi$ tatsächlich eine Äquivalenzrelation auf den $\mathcal{K}_i$ ist. Reflexivität und Symmetrie sind klar.

Für die Transitivität betrachten wir $x \sim_\varphi y$ und $y \sim_\varphi z$. Falls $x \sim y$ oder $y \sim z$ gilt, sind wir fertig. Damit verbleiben im Wesentlichen zwei Fälle:

1. Es gelten:

$$\begin{array}{ccccccc} x & \sim & v & & & & \\ & & \varphi(v) & \sim & y & \sim & w \\ & & & & & & \varphi(w) \sim z \end{array}$$

Wegen $w \preccurlyeq f$ und $\varphi(v) \preccurlyeq g$ (zweite Zeile) folgt mit der Bedingung an φ , dass $\varphi(v) = \varphi(w)$ gilt, also $v = w$. Folglich erhalten wir $x \sim_\varphi z$.

2. Es gelten:

$$\begin{array}{ccc} \textcolor{red}{x} & \sim & v \\ \varphi(v) & \sim & \textcolor{red}{y} \sim \varphi(w) \end{array} \quad \begin{array}{ccc} w & \sim & \textcolor{red}{z} \end{array}$$

Da $\varphi(v) \preceq g$ und $\varphi(w) \preceq g$, folgt $\varphi(v) = \varphi(w)$ aus der Irreduzibilität. Also haben wir auch hier $\textcolor{red}{x} \sim_{\varphi} \textcolor{red}{z}$, sogar $\textcolor{red}{x} \sim \textcolor{red}{z}$. $\square$

Eine Äquivalenzrelation kann stets formal erweitert werden, ein n -Überlagerungskomplex jedoch nur, falls dadurch wieder ein n -Überlagerungskomplex entsteht. Daraus ergeben sich zwei Fragen:

- Unter welchen Voraussetzungen ist $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim_{\varphi})$ wieder ein n -Überlagerungskomplex?
- Können wir jeden n -Überlagerungskomplex durch Iteration dieser Konstruktion erhalten?

Da sowohl Erweiterung als auch Teilüberlagerung ein Konzept von Faltung beschreiben, sollten wir die Interaktion dieser Begriffe klären.

Bemerkung 2.24. *Sei $\mathcal{S}$ ein n -Überlagerungskomplex und $\mathcal{T}$ eine Erweiterung von $\mathcal{S}$. Dann ist $\mathcal{S}$ eine Teilüberlagerung von $\mathcal{T}$.*

Die Umkehrung dieser Bemerkung ist aber nicht offensichtlich. Wir formulieren sie daher als weitere Frage.

- Lässt sich jede Teilüberlagerung durch wiederholte Erweiterungen darstellen?

Bevor wir uns diesen Fragen in Abschnitt 2.3.1 zuwenden, sollten wir klarstellen, wieso die lokale Simplität in der Definition der Erweiterung keine Rolle gespielt hat. Da die Konstruktion einer Erweiterung nur auf der Existenz einer Identifikation beruht, ist sie auch für die entarteten Fälle aus Beispiel 2.9 definiert. Lokal simpliziale Komplexe zeichnen sich dadurch aus, dass es immer Identifikationen von zwei Flächen gibt, die konstant auf dem Schnitt sind, wie Lemma 2.25 zeigt.

Lemma 2.25. *Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex und $f, g \in \mathcal{K}_j$ (mit $0 \leq j \leq n$). Dann gibt es mindestens eine Identifikation $\varphi : \mathcal{S}_{\preceq f} \rightarrow \mathcal{S}_{\preceq g}$, die konstant auf dem Schnitt ist.*

Beweis. Wir beschreiben zunächst alle möglichen Identifikationen $\varphi : \mathcal{S}_{\preceq f} \rightarrow \mathcal{S}_{\preceq g}$. Dazu betrachten wir die (aufgrund der lokalen Simplität) gleichmächtigen Mengen

$$\begin{aligned} \mathcal{K}_0^{\preceq f} &:= \{p \in \mathcal{K}_0 \mid p \preceq f\} \\ \mathcal{K}_0^{\preceq g} &:= \{p \in \mathcal{K}_0 \mid p \preceq g\}. \end{aligned}$$

Jede Bijektion $\beta : \mathcal{K}_0^{\preceq f} \rightarrow \mathcal{K}_0^{\preceq g}$ induziert eine Identifikation $\varphi : \mathcal{S}_{\preceq f} \rightarrow \mathcal{S}_{\preceq g}$, da $\mathcal{S}_{\preceq f}$ und $\mathcal{S}_{\preceq g}$ isomorph zu j -Simplizes sind.

Um die Konstanz auf dem Schnitt einzuarbeiten, betrachten wir

$$G := \{(a, b) \in \mathcal{K}_0^{\preceq f} \times \mathcal{K}_0^{\preceq g} \mid a \sim b\}$$

Aufgrund der Irreduzibilität von $\mathcal{S}$ folgt für zwei verschiedene $(a_1, b_1), (a_2, b_2) \in G$, dass $\{a_1, b_1\} \cap \{a_2, b_2\} = \emptyset$ gilt. Folglich lässt sich G zu einer Bijektion $\beta : \mathcal{K}_0^{\preceq f} \rightarrow \mathcal{K}_0^{\preceq g}$ fortsetzen.

Wir zeigen nun, dass jede solche Fortsetzung eine Identifikation induziert, die konstant auf dem Schnitt ist. Seien dazu $x, y \in \mathcal{K}_i$ mit $x \preceq f$, $y \preceq g$ und $x \sim y$ gegeben. Da nach Konstruktion (und Abwärtskompatibilität) die Menge $\{p \in \mathcal{K}_0^{\preceq f} \mid p \preceq x\}$ auf $\{p \in \mathcal{K}_0^{\preceq g} \mid p \preceq y\}$ abgebildet wird, bildet die entsprechende Identifikation auch x auf y ab. $\square$

2.3.1 Primitive Erweiterungen

In diesem Abschnitt beantworten wir die Konstruktionsfragen des vorherigen Abschnitts:

- Wann liefert die formale Erweiterung einen n -Überlagerungskomplex?
- Können wir jeden n -Überlagerungskomplex durch wiederholte Erweiterungen erhalten?
- Lässt sich jede Teilüberlagerung durch wiederholte Erweiterungen darstellen?

Wir werden primär untersuchen, unter welchen Voraussetzungen die formale Erweiterung aus Definition 2.23 einen n -Überlagerungskomplex liefert. Diese Untersuchung wird es uns ermöglichen, die beiden anderen Fragen zustimmend zu beantworten.

Wir werden den Fall allgemeiner n -Überlagerungskomplexe betrachten, da die Annahme der lokalen Simplität nur einzelne Details verändert. Diese Änderungen sind ab Seite 27 aufgeführt. Eine Abschwächung der lokalen Simplität (sodass beliebige m -Ecke als Flächen zugelassen sind) untersuchen wir in Kapitel 5.1. Keine dieser Modifikationen wird stark von den Ergebnissen dieses Abschnitts abweichen.

Um zu verstehen, wann die Erweiterung einer Äquivalenzrelation einen n -Überlagerungskomplex liefert, müssen wir uns näher mit den Details der Identifikation befassen. Bei jeder Identifikation fallen jeweils maximal zwei Äquivalenzklassen zusammen. Sobald es mehrere solcher Vereinigungen gibt, müssen wir uns mit den Zusammenhängen zwischen diesen beschäftigen. Um solche Interaktionen gering zu halten, wäre es einfacher, sich nur solche Identifikationen anzusehen, in denen *genau* zwei Äquivalenzklassen vereinigt werden. Erweiterungen, die von solchen Identifikationen stammen, nennen wir *primitiv*.

Definition 2.26 (primitive Erweiterung). *Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein n -Überlagerungskomplex. Die Erweiterung von $\sim$ um die Identifikation φ heißt **primitiv**, falls es genau ein $0 \leq i \leq n$ mit*

- $|\mathcal{K}_i / \sim| = 1 + |\mathcal{K}_i / \sim_\varphi|$ (genau zwei Klassen fallen zusammen) und
- $\mathcal{K}_j / \sim = \mathcal{K}_j / \sim_\varphi$ für $j \neq i$ (alle anderen Klassen bleiben gleich)

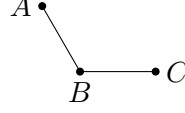
*gibt. Sie wird auch als **primitive Erweiterung vom Grad i** bezeichnet.*

*Falls $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim_\varphi)$ ein n -Überlagerungskomplex ist, dann nennen wir diese Erweiterung von $\mathcal{S}$ **primitiv** und bezeichnen sie mit $\mathcal{S}_\varphi$.*

Wenn wir uns ausschließlich mit primitiven Erweiterungen befassen wollen, müssen wir jede Identifikation so zerlegen, dass in jedem Schritt nur ein Paar von Äquivalenzklassen vereinigt wird. In Beispiel 2.27 wird eine solche Zerlegung exemplarisch ausgeführt.

Beispiel 2.27. Wir betrachten den folgenden lokal simplizialen 1-Überlagerungskomplex $\mathcal{S}$:

$$\begin{aligned}\mathcal{K}_0 &:= \{A, B, C\} \\ \mathcal{K}_1 &:= \{\{A, B\}, \{B, C\}\}\end{aligned}$$



mit $\prec := \in$ und $\sim$ als Gleichheit.

Dann ist $\varphi : \mathcal{S}_{\prec\{A,B\}} \rightarrow \mathcal{S}_{\prec\{B,C\}}$ mit $\varphi(A) = C$ eine Identifikation, die konstant auf dem Schnitt ist. Die Erweiterung um diese Identifikation würde die folgenden Äquivalenzklassen vereinigen:

$$\{A\} \leftrightarrow \{C\} \qquad \{\{A, B\}\} \leftrightarrow \{\{B, C\}\}$$

Wir können die selben Vereinigungen aber auch dadurch erreichen, dass wir zuerst um die Identifikation $\psi : \mathcal{S}_{\prec A} \rightarrow \mathcal{S}_{\prec C}$ und dann erst um φ erweitern.

Bei der Erweiterung um ψ werden nämlich die Klassen $\{A\}$ und $\{C\}$ vereinigt. Dadurch werden bei der Erweiterung um φ nur noch die Kantenklassen vereinigt. Es handelt sich also in beiden Fällen um primitive Erweiterungen.

Die Zerlegung aus diesem Beispiel lässt sich auch allgemein ausführen.

Bemerkung 2.28. Sei $\mathcal{S}$ ein n -Überlagerungskomplex und $\mathcal{T}$ eine Erweiterung von $\mathcal{S}$. Dann gibt es eine endliche Folge von n -Überlagerungskomplexen

$$\mathcal{S} = \mathcal{S}_0, \mathcal{S}_1, \dots, \mathcal{S}_k = \mathcal{T},$$

so dass $\mathcal{S}_i$ eine primitive Erweiterung von $\mathcal{S}_{i-1}$ für $1 \leq i \leq k$ ist.

Durch diese Zerlegung zerfällt die Faltung von zwei Dreiecken in maximal sieben kleinere Faltungen, da höchstens ein Paar von Flächen, sowie jeweils drei Paare von Kanten und Punkten identifiziert werden. Damit lassen sich die Einschränkungen, denen mögliche Identifikationen unterliegen, genauer verstehen (In Kapitel 3 werden diese Einschränkungen noch deutlicher zu Tage treten).

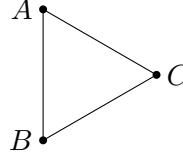
Wir werden jetzt untersuchen, unter welchen Voraussetzungen primitive Erweiterungen der Äquivalenzrelation einen n -Überlagerungskomplex liefern. Wir werden sehen, dass nur in Dimension Null Probleme auftreten können, deshalb behandeln wir diese getrennt von den anderen Dimensionen.

Betrachten wir zunächst primitive Erweiterungen vom Grad 0, bei denen wir zwei Punktäquivalenzklassen identifizieren. Die Abwärtskompatibilität ist bei diesen automatisch erfüllt. Es verbleibt die Überprüfung der Irreduzibilität. Das folgende Beispiel zeigt, dass es bei dieser Komplikationen geben kann.

Beispiel 2.29. Wir betrachten den lokal simplizialen homogenen 1-Komplex $\mathcal{S}$ (der auch ein 1-Überlagerungskomplex ist, wenn man $\sim$ als Gleichheit definiert):

$$\mathcal{K}_0 := \{A, B, C\}$$

$$\mathcal{K}_1 := \{\{A, B\}, \{A, C\}, \{B, C\}\}$$



mit $\prec := \in$.

Die Abbildung $\varphi : \mathcal{S}_{\prec\{A,C\}} \rightarrow \mathcal{S}_{\prec\{B,C\}}$, die vermöge

$$A \mapsto B$$

$$C \mapsto C$$

$$\{A, C\} \mapsto \{B, C\}$$

definiert ist, ist eine Identifikation vom Grad 1 (die höchsten vorkommenden Elemente liegen in $\mathcal{K}_1$), die konstant auf dem Schnitt ist. Aber Hinzufügen von $A \sim_\varphi B$ würde die Irreduzibilität an $\{A, B\}$ brechen.

Wir müssen also bei der Identifikation der Punkte p und q ausschließen, dass es ein $x \in \bigcup_{i=1}^n \mathcal{K}_i$ mit $p \prec x$ und $q \prec x$ gibt. Diese neue Bedingung ist stärker als die Konstanz auf dem Schnitt, da sie diese als Spezialfall enthält.

Lemma 2.30. Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein n -Überlagerungskomplex. Weiter seien $f, g \in \mathcal{K}_0$ mit $f \not\sim g$ gegeben, zusammen mit einer Identifikation $\varphi : \mathcal{S}_{\prec f} \rightarrow \mathcal{S}_{\prec g}$.

Die primitive Erweiterung vom Grad 0 um φ liefert genau dann einen n -Überlagerungskomplex, wenn gilt:

- Es gibt keine $a, b \in \mathcal{K}_0$ mit $a \sim f$ und $b \sim g$, sodass ein $x \in \bigcup_{i=1}^n \mathcal{K}_i$ existiert, das $a \prec x$ und $b \prec x$ erfüllt.

Beweis. Die Abwärtskompatibilität ist stets erfüllt, da es kein $y \in \bigcup_{i=0}^n \mathcal{K}_i$ mit $y \prec f$ gibt. Die einzige Einschränkung ist damit die Irreduzibilität. Diese besagt: Aus $a \sim_\varphi b$ und $a \neq b$ folgt, dass es kein $x \in \bigcup_{i=0}^n \mathcal{K}_i$ mit $a \prec x$ und $b \prec x$ gibt. Dies ist nur für Elemente der $\sim_\varphi$ -Äquivalenzklasse von f bzw. g zu prüfen.

Damit erfüllt die primitive Erweiterung die Irreduzibilität genau dann nicht, wenn es $a \sim f$ und $b \sim g$ (wegen $f \not\sim g$ folgt daraus schon $a \neq b$), sowie $x \in \bigcup_{i=0}^n \mathcal{K}_i$ mit $a \prec x$ und $b \prec x$ gibt. $\square$

Nachdem wir die wichtige Einschränkung bei primitiven Erweiterungen von Grad Null abgehandelt haben, betrachten wir als nächstes primitive Erweiterungen höherer Grade. Bei diesen gibt es glücklicherweise keine Einschränkungen, wie Lemma 2.31 zeigt.

Lemma 2.31. Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein n -Überlagerungskomplex. Zusätzlich sei eine Identifikation φ vom Grad > 0 gegeben. Falls $\sim_\varphi$ eine primitive Erweiterung von $\sim$ ist, dann ist $\mathcal{S}_\varphi$ eine primitive Erweiterung von $\mathcal{S}$.

Beweis. Wir müssen zeigen, dass $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim_\varphi)$ ein n -Überlagerungskomplex ist. Für die Abwärtskompatibilität müssen wir nur die $x, y \in \biguplus_{i=0}^n \mathcal{K}_i$ betrachten, die $x \sim_\varphi y$, aber nicht $x \sim y$ erfüllen. Ohne Einschränkung gelte $\varphi(x) = y$ (ansonsten gehen wir zu den entsprechenden äquivalenten Elementen über). Sei nun $a \in \biguplus_{i=0}^n \mathcal{K}_i$ mit $a \prec x$ gegeben. Da φ primitiv ist, muss $\varphi(a) \sim a$ gelten. Wegen $\varphi(a) \prec \varphi(x) = y$ ist damit die Abwärtskompatibilität erfüllt.

Da die Erweiterung $\sim_\varphi$ primitiv ist, ist die Abwärtskompatibilität erfüllt (alle kleineren Elemente sind schon äquivalent). Wir prüfen noch die Irreduzibilität.

Wäre die Irreduzibilität für $a \sim_\varphi b$ verletzt, dann gäbe es $p, q \in \mathcal{K}_0$ mit $p \prec a$ und $q \prec b$, die $p \sim_\varphi q$ erfüllen und die Irreduzibilität verletzen (da $p \prec x$ und $q \prec x$ aus der Transitivität folgen). Da die Erweiterung primitiv war, gilt $p \sim_\varphi q$ genau dann, wenn $p \sim q$ gilt. Damit wäre $\mathcal{S}$ selbst kein n -Überlagerungskomplex gewesen. Aus diesem Widerspruch folgt, dass die Irreduzibilität stets erfüllt ist. $\square$

Wir haben damit nachgewiesen, unter welchen Voraussetzungen eine primitive Erweiterung der Äquivalenzrelation einen n -Überlagerungskomplex induziert. Zusammen mit Bemerkung 2.28 können wir in Satz 2.32 zusammenfassen, welche Identifikation eine Erweiterung von n -Überlagerungskomplexen erfüllt. Indem man die Bedingung dieses Satzes überprüft, kann man z. B. erkennen, ob zwei Flächen eines Faltzustandes sich zusammenfalten lassen.

Da die Aufspaltung einer Identifikation φ in primitive Erweiterungen stets möglich ist, liefert Lemma 2.30 die Haupteinschränkung für die Erweiterbarkeit. Satz 2.32 fasst die Ergebnisse der Lemmata 2.30 und 2.31 zusammen.

Satz 2.32. *Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein n -Überlagerungskomplex und $\varphi : \mathcal{S}_{\preceq f} \rightarrow \mathcal{S}_{\preceq g}$ eine Identifikation auf $\mathcal{S}$. Dann liefert die Erweiterung von $\sim$ um φ genau dann einen n -Überlagerungskomplex, wenn gilt*

- Für jedes $p \in \mathcal{K}_0$ mit $p \preceq f$ und $p \not\sim \varphi(p)$ gibt es kein $x \in \bigcup_{i=1}^n \mathcal{K}_i$ mit $p \prec x$ und $\varphi(p) \prec x$.

Beweis. Wir zerlegen φ in primitive Erweiterungen. Nach Lemma 2.30 ist die Erweiterung von $\mathcal{S}$ um die primitiven Erweiterungen vom Grad 0 möglich, falls die Bedingung des Satzes gilt. Andernfalls wird ein Widerspruch in der Abwärtskompatibilität erzeugt, der durch keine Erweiterung von $\sim$ aufgelöst werden kann. Alle weiteren primitiven Erweiterungen liefern nach Lemma 2.31 keine Probleme. $\square$

Wir wenden uns jetzt der Frage zu, ob sich jede Teilüberlagerung durch wiederholte Erweiterungen ausdrücken lässt. Dazu zeigen wir in Lemma 2.33 für endliche Komplexe⁵, dass sich jede Teilüberlagerung sogar durch primitive Erweiterungen ausdrücken lässt.

Lemma 2.33. *Seien $\mathcal{S}$ und $\mathcal{T}$ zwei endliche n -Überlagerungskomplexe, wobei $\mathcal{S}$ eine Teilüberlagerung von $\mathcal{T}$ ist. Dann gibt es eine Folge von n -Überlagerungskomplexen*

$$\mathcal{S} = \mathcal{S}_0, \mathcal{S}_1, \dots, \mathcal{S}_k = \mathcal{T},$$

sodass $\mathcal{S}_i$ eine primitive Erweiterung von $\mathcal{S}_{i-1}$ für $1 \leq i \leq k$ ist.

Beweis. Zunächst stellen wir fest, dass es genügt, eine primitive Erweiterung von $\mathcal{S}$ zu konstruieren, die Teilüberlagerung von $\mathcal{T}$ ist. Eine Iteration dieser Konstruktion liefert dann die gewünschte Folge. Da die Anzahl der Äquivalenzklassen stets um 1 abnimmt und die Anzahl

⁵Wenn man daran interessiert ist, kann man diese Aussage mit den Ideen aus Abschnitt 4.2.2 auf den unendlichen Fall übertragen. Diese Möglichkeit verfolgen wir hier aber nicht weiter.

der Äquivalenzklassen von $\sim^{\mathcal{S}}$ und $\sim^{\mathcal{T}}$ endlich ist, terminiert dieses Verfahren bei $\mathcal{S}_k = \mathcal{T}$ für ein $k \geq 0$.

Sei $i \geq 0$ minimal mit $\mathcal{K}_i / \sim^{\mathcal{S}} \neq \mathcal{K}_i / \sim^{\mathcal{T}}$. Wähle eine Klasse von $\sim^{\mathcal{T}}$, die in mehrere $\sim^{\mathcal{S}}$ -Klassen $[x_1], \dots, [x_m]$ zerfällt. Dann ist

$$\varphi : \mathcal{S}_{\preccurlyeq x_1} \rightarrow \mathcal{S}_{\preccurlyeq x_2} : x_1 \mapsto x_2$$

eine Identifikation vom Grad i auf $\mathcal{S}$. Die Erweiterung von $\sim^{\mathcal{S}}$ um φ ist stets primitiv, da alle kleineren Simplizes nach Wahl von i bereits identifiziert sind.

Im Fall $i = 0$ prüfen wir mit Lemma 2.30, ob $\mathcal{S}_\varphi$ ein n -Überlagerungskomplex ist. Gäbe es $y \in \mathcal{K}_1$, $a, b \in \mathcal{K}_0$ mit

$$x_1 \sim^{\mathcal{S}} a \quad x_2 \sim^{\mathcal{S}} b \quad a \prec y \quad b \prec y,$$

dann gäbe es auch (da $\mathcal{S}$ Teilüberlagerung von $\mathcal{T}$ ist)

$$x_1 \sim^{\mathcal{T}} a \quad x_2 \sim^{\mathcal{T}} b \quad a \prec y \quad b \prec y,$$

im Widerspruch zur Irreduzibilität von $\mathcal{T}$. Damit ist $\mathcal{S}_\varphi$ eine primitive Erweiterung von $\mathcal{S}$ und eine Teilüberlagerung von $\mathcal{T}$.

Der Fall $i > 0$ wurde bereits in Lemma 2.31 gezeigt. $\square$

Zusammen mit Bemerkung 2.24 erhalten wir die Äquivalenz der beiden abstrakten Beschreibungen von Faltungen, die wir entwickelt haben: Teilüberlagerungen und Erweiterungen beschreiben (im endlichen Fall) die selben Faltzustände.

Wählt man in Lemma 2.33 die Äquivalenzrelation $\sim$ von $\mathcal{T}$ als Gleichheit, können wir die letzte unserer Fragen beantworten:

Folgerung 2.34. *Man kann den endlichen n -Überlagerungskomplex $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ allein durch primitive Erweiterungen aus dem homogenen n -Komplex $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec)$ erhalten.*

Wir haben damit gezeigt, dass es keinen endlichen n -Überlagerungskomplex gibt, der sich nicht aus einem homogenen n -Komplex erzeugen lässt. Wenn wir einen homogenen 2-Komplex als Papiermuster auffassen, ist ein 2-Überlagerungskomplex einfach eine Wahl von Flächen, die zusammengefasst werden sollen. Dann sagt uns diese Folgerung, dass es möglich ist, jede solche Wahl zu falten. Da dies mit starrem Papier nicht möglich ist, wie z. B. die erzwungenen Faltungen aus den Beispielen 2.3 und 2.4 belegen, handelt es sich hier um eine Schwäche unseres Modells.

Modifikation für lokale Simplizität Wenn wir zusätzlich annehmen, dass unsere n -Überlagerungskomplexe lokal simplizial sind (also dass die Papierfaltung aus Dreiecke besteht), können wir die Aussagen aus Abschnitt 2.3.1 ein wenig verschärfen, um etwas einfachere Bedingungen für das Zusammenfalten von Dreiecken (im Vergleich zu allgemeineren Flächen) zu erhalten.

Um dies auszuführen, müssen wir voraussetzen, dass lokale Simplizität bei Erweiterungen erhalten bleibt. Da diese aber nur vom zugrundeliegenden homogenen n -Komplex abhängt und dieser sich bei Erweiterungen nicht ändert, ist dies der Fall, wie wir in Bemerkung 2.35 festhalten.

Bemerkung 2.35. Sei $\mathcal{S}$ ein lokal simplizialer n -Überlagerungskomplex. Jede Erweiterung von $\mathcal{S}$ ist auch lokal simplizial.

In Folgerung 2.36 verschärfen wir die Aussage von Lemma 2.30. Dort hatten wir untersucht, wann die Erweiterung um eine Identifikation vom Grad Null (also der Punkte p und q) wieder einen n -Überlagerungskomplex liefert. Im Fall von Papier dürfen p und q dafür nicht in einer gemeinsamen Kante oder Fläche (die nicht notwendig ein Dreieck ist) vorkommen. Jetzt, wo wir uns auf Dreiecke beschränken, müssen wir nur überprüfen, dass p und q nicht in einer gemeinsamen Kante vorkommen.

Folgerung 2.36. Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex. Weiter seien $p, q \in \mathcal{K}_0$ mit $p \not\sim q$ gegeben, zusammen mit einer Identifikation $\varphi : \mathcal{S}_{\preceq p} \rightarrow \mathcal{S}_{\preceq q}$. Die primitive Erweiterung vom Grad 0 um φ liefert genau dann einen n -Überlagerungskomplex, wenn gilt:

- Es gibt keine $a, b \in \mathcal{K}_0$ mit $a \sim p$ und $b \sim q$, sodass ein $x \in \mathcal{K}_1$ existiert, das $a \prec x$ und $b \prec x$ erfüllt.

Beweis. Gemäß Lemma 2.30 liefert die primitive Erweiterung vom Grad 0 um φ genau dann einen n -Überlagerungskomplex, wenn gilt:

Es gibt keine $a, b \in \mathcal{K}_0$ mit $a \sim p$ und $b \sim q$, sodass ein $x \in \bigcup_{i=1}^n \mathcal{K}_i$ existiert, das $a \prec x$ und $b \prec x$ erfüllt.

Da $\mathcal{S}_{\preceq x}$ isomorph zu einem Simplex ist und je zwei Punkte eines solchen durch eine Kante verbunden sind, kann ohne Einschränkung $x \in \mathcal{K}_1$ angenommen werden. $\square$

Damit ergibt sich natürlich auch eine Verschärfung von Satz 2.32, der klassifiziert, wann die Erweiterung um eine allgemeine Identifikation einen n -Überlagerungskomplex liefert. Da die einzige Einschränkung dieses Satz auf dem Analogon von Folgerung 2.36 beruhte, überträgt sich deren Vereinfachung auch hier.

Folgerung 2.37. Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex und $\varphi : \mathcal{S}_{\preceq f} \rightarrow \mathcal{S}_{\preceq g}$ eine Identifikation auf $\mathcal{S}$. Dann liefert die Erweiterung von $\sim$ um φ genau dann einen n -Überlagerungskomplex, wenn gilt

- Für jedes $p \in \mathcal{K}_0$ mit $p \preceq f$ und $p \not\sim \varphi(p)$ gibt es kein $x \in \mathcal{K}_1$ mit $p \prec x$ und $\varphi(p) \prec x$.

Insgesamt haben wir also mit Satz 2.32 (im allgemeinen Fall) bzw. Folgerung 2.37 (im lokal simplizialen Fall) klassifiziert, in welchen Fällen eine Identifikation einen n -Überlagerungskomplex liefert. Im Fall von Papier wissen wir also, wann wir in unserem Modell zwei Flächen zusammenfalten können.

2.3.2 Faltung benachbarter Flächen

Wir haben uns in Abschnitt 2.3.1 damit befasst, wann die Erweiterung um eine Identifikation einen n -Überlagerungskomplex liefert. Allerdings ist es noch nicht klar, in welchen Fällen eine

solche Identifikation existiert. In diesem Abschnitt betrachten wir eine Situation, in der wir die Existenz von Erweiterungen häufig allgemein garantieren können.

Dazu schränken wir uns auf das Zusammenfallen benachbarter Flächen ein. Im lokal simplizialen Fall wird eine Identifikation vom Grad 2 durch die Angabe zweier Dreiecke und die Identifikation der Punkte eindeutig festgelegt. Daher gibt es zu zwei Flächen generisch 6 solche Identifikationen, die durch die Konstanz auf dem Schnitt (vgl. Definition 2.22) weiter eingeschränkt werden. Falls die Dreiecke z. B. einen Punkt gemeinsam haben, gibt es nur noch zwei mögliche Identifikationen (man kann noch die beiden anderen Punkte vertauschen).

Sind zwei Flächen aber durch eine Kante verbunden (also benachbart), so gibt es genau eine Identifikation dieser beiden Flächen, die wir Nachbaridentifikation nennen.

Definition 2.38 (Nachbaridentifikation). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein n -Überlagerungskomplex. Eine Identifikation $\varphi : \mathcal{S}_{\preceq f} \rightarrow \mathcal{S}_{\preceq g}$ mit $f, g \in \mathcal{K}_i$ ($0 \leq i \leq n$) heißt **Nachbaridentifikation** (bzgl. $\mathcal{S}$), falls gilt:

- φ ist konstant auf dem Schnitt.
- Es gilt entweder $i = 0$ oder es gibt $k \in \mathcal{K}_{n-1}$ mit $[k] \prec [f]$ und $[k] \prec [g]$.

Wir halten in Bemerkung 2.39 fest, dass es für lokal simpliziale n -Überlagerungskomplexe immer eindeutige Nachbaridentifikationen gibt.

Bemerkung 2.39. Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex. Seien $f, g \in \mathcal{K}_i$ (mit $0 < i \leq n$) und $k \in \mathcal{K}_{i-1}$ gegeben, sodass $[k] \prec [f]$ und $[k] \prec [g]$ gelten.

Dann gibt es genau eine Nachbaridentifikation $\varphi : \mathcal{S}_{\preceq f} \rightarrow \mathcal{S}_{\preceq g}$.

Beweis. Da $\mathcal{S}$ lokal simplizial ist, ist die Identifikation bereits durch die Bilder der Punkte festgelegt. Ein i -Simplex hat $i + 1$ Punkte, von denen durch die Nachbarschaftsbedingung bereits die Bilder von i Punkten festgelegt sind. Da φ bijektiv sein muss, gibt es genau eine mögliche Fortsetzung. $\square$

Für homogene n -Komplexe, die sich komplett auf einen n -Simplex zusammenfallen lassen, können wir sogar garantieren, dass jede Nachbaridentifikation eine Erweiterung ermöglicht.

Satz 2.40. Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex. Es gebe eine Äquivalenzrelation $\sim^\Delta$, sodass gelten:

- $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim^\Delta)$ ist ein n -Überlagerungskomplex.
- $\mathcal{S}$ ist eine Teilüberlagerung von $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim^\Delta)$.
- $\mathcal{K}_n / \sim^\Delta$ ist einelementig.

Seien $\varphi_1, \dots, \varphi_m$ Nachbaridentifikationen. Dann ist $\mathcal{S}_{\varphi_1, \dots, \varphi_m}$ eine Erweiterung von $\mathcal{S}$.

Beweis. Jede Nachbaridentifikation $\varphi : \mathcal{S}_{\preceq f} \rightarrow \mathcal{S}_{\preceq g}$ ist aufgrund von Bemerkung 2.39 eindeutig festgelegt.

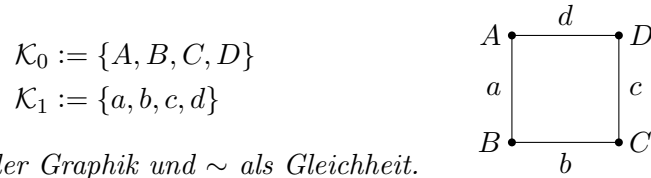
Wir konstruieren nun aus $f \sim^\Delta g$ einen Komplex-Morphismus $\psi : \mathcal{S}_{\preceq f} \rightarrow \mathcal{S}_{\preceq g}$, wobei $x \preceq f$ auf das Element $y \preceq g$ abgebildet wird, das ihm gemäß der starken Abwärtskompatibilität

(Bemerkung 2.18) eindeutig zugeordnet sind. Mit dem selben Argument ist auch φ^{-1} ein Komplex-Morphismus. Insgesamt haben wir also einen Komplex-Isomorphismus bzw. eine Identifikation vorliegen. Diese ist nach Konstruktion konstant auf dem Schnitt.

Folglich müssen φ und ψ übereinstimmen. Insbesondere erfüllt die Erweiterung um φ stets die Irreduzibilitätseigenschaft, ist also wohldefiniert. $\square$

Wir möchten darauf hinweisen, dass die φ_j aus Satz 2.40 nicht notwendig unabhängig sein müssen, z. B. könnte $\mathcal{S}_{\varphi_1, \dots, \varphi_m} = (\mathcal{S}_{\varphi_1, \dots, \varphi_m})_{\varphi_{m+1}}$ eintreten, wie Beispiel 2.41 zeigt.

Beispiel 2.41. Sei $\mathcal{S}$ der lokal simpliziale 1-Überlagerungskomplex



mit $\prec$ wie in der Graphik und $\sim$ als Gleichheit.

Für je zwei benachbarte Elemente aus $\mathcal{K}_1$ legt Satz 2.40 genau eine Identifikation fest. Dies definiert die vier Identifikationen $\varphi_{ab}, \varphi_{bc}, \varphi_{cd}, \varphi_{da}$. Nachdem wir $\mathcal{S}$ um $\varphi_{ab}, \varphi_{bc}$ und φ_{cd} erweitert haben, liefert die Erweiterung um φ_{da} keine Änderung.

Die Situation aus Beispiel 2.41 tritt genau dann auf, wenn es eine Folge

$$f_1, f_2, \dots, f_m = f_1 \quad f_i \in \mathcal{K}_n$$

gibt, sodass f_i und f_{i+1} für jedes $1 \leq i < m$ eine gemeinsame Kante haben⁶.

2.3.3 Kommutativität

Wir haben uns in Abschnitt 2.3 bislang nur mit der Definition und Konstruktion des Faltvorgangs beschäftigt. Wenn man aber versucht, alle möglichen Faltungen zu einem gegebenen n -Überlagerungskomplex zu bestimmen, wächst die Anzahl der Möglichkeiten sehr schnell mit der Größe des Komplexes an. Es ist also sinnvoll, sich mit einer Verringerung dieser Möglichkeiten zu befassen.

In diesem Abschnitt werden wir nachweisen, dass verschiedene Erweiterungen miteinander kommutieren. Es spielt also keine Rolle, in welcher Reihenfolge wir zwei verschiedenen Simplexpaare identifizieren. Dazu zeigen wir für zwei Identifikationen φ und ψ , dass die Erweiterung von $\sim_\varphi$ um ψ gleich der Erweiterung von $\sim_\psi$ um φ ist.

Wir führen diese Aussage auf die Gleichheit der Äquivalenzrelationen $\sim_{\varphi\psi}$ und $\sim_{\psi\varphi}$ zurück, die wir dadurch nachweisen, dass wir eine Darstellung von $\sim_{\varphi\psi}$ angeben, die unabhängig von der Reihenfolge der Identifikationen ist. Um diese zu finden, verwenden wir die präzise Änderung der Äquivalenzrelationen durch eine Erweiterung aus Definition 2.23.

Wir können die Identifikation φ nämlich als **Übergang** zwischen jeweils zwei Äquivalenzklassen verstehen. Die natürliche Verallgemeinerung dieser Anschauung wäre, dass $\sim_{\varphi\psi}$ zwei Übergänge erlaubt, einen mittels φ und einen mittels ψ . Allerdings entspricht das nicht unserer Definition von $\sim_{\varphi\psi}$. In dieser erlauben wir *drei* Übergänge, zuerst mit φ , dann mit

⁶Mit den Bezeichnungen aus Abschnitt 4.1.6 liegt hier also einen geschlossenen Weg vor.

ψ und zuletzt erneut mit φ . Im Allgemeinen kann es passieren, dass man diese alle benötigt, wie Beispiel 2.42 demonstriert.

Beispiel 2.42. Wir betrachten den 1-Überlagerungskomplex $\mathcal{S}$:

$$\begin{aligned} \mathcal{K}_0 &:= \{A, B, C, D\} & \begin{array}{c} A \quad \quad B \\ \bullet \quad \quad \bullet \\ \hline \end{array} \\ \mathcal{K}_1 &:= \{\{A, B\}, \{C, D\}\} & \begin{array}{c} \bullet \quad \quad \bullet \\ C \quad \quad D \\ \hline \end{array} \end{aligned}$$

mit $\prec := \in$, und $\sim$ als Gleichheit.

Nun ist $\varphi : \mathcal{S}_{\preccurlyeq\{A,B\}} \rightarrow \mathcal{S}_{\preccurlyeq\{C,D\}}$ vermöge $\varphi(A) = C$ und $\varphi(B) = D$ eine Identifikation vom Grad 1. Ersichtlich ist $\mathcal{S}_\varphi$ eine Erweiterung von $\mathcal{S}$. Ebenso ist $\psi : \mathcal{S}_{\preccurlyeq B} \rightarrow \mathcal{S}_{\preccurlyeq C}$ mit $\psi(B) = C$ eine Identifikation vom Grad 0. Auch hier ist $\mathcal{S}_\psi$ eine Erweiterung von $\mathcal{S}$.

Allerdings liefert die Erweiterung um beide Identifikationen keinen 1-Überlagerungskomplex, da die Irreduzibilität an $\{A, B\}$ gebrochen wird. Und während $A \sim_{\varphi\psi} D$ wegen

$$A \sim_\varphi C \sim_\psi B \sim_\varphi D$$

gilt, lässt sich das nicht nur mit höchstens zwei Übergängen beschreiben.

Dass die Irreduzibilität bereits bei der Definition der Erweiterung eine entscheidende Rolle spielt, sieht man daran, dass A und D bezüglich $\sim_{\psi\varphi}$ nicht äquivalent sind.

Während uns Beispiel 2.42 Grenzen für die allgemeine Vertauschbarkeit aufzeigt, können wir nachweisen, dass wir nicht so viele Übergänge benötigen, wenn die doppelte Erweiterung einen n -Überlagerungskomplex liefert.

Lemma 2.43. Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein n -Überlagerungskomplex. Seien φ, ψ zwei Identifikationen auf $\mathcal{S}$, sodass eine der beiden doppelten Erweiterungen $(\mathcal{S}_\varphi)_\psi$ oder $(\mathcal{S}_\psi)_\varphi$ existiert. Dann gilt für $x, y \in \bigcup_{i=0}^n \mathcal{K}_i$, dass $x \sim_{\varphi\psi} y$ genau dann zutrifft, wenn mindestens eine der folgenden Bedingungen erfüllt ist:

- (Kein Übergang) $x \sim y$
- (Ein Übergang) Es gibt ein $z \in \bigcup_{i=0}^n \mathcal{K}_i$ mit $x \sim z$ und $\rho(z) \sim y$ für $\rho \in \{\varphi^{\pm 1}, \psi^{\pm 1}\}$
- (Zwei Übergänge) Es gibt $w_1, w_2 \in \bigcup_{i=0}^n \mathcal{K}_i$ mit $x \sim w_1$, $\rho_1(w_1) \sim w_2$ und $\rho_2(w_2) \sim y$ mit $\{\rho_1, \rho_2\} \subseteq \{\varphi^{\pm 1}, \psi^{\pm 1}\}$, wobei ρ_1 und ρ_2 verschiedenfarbig sind.

Beweis. Klar: Falls eine der Bedingungen erfüllt ist, folgt auch $x \sim_{\varphi\psi} y$.

Für die umgekehrte Richtung ist $x \sim_{\varphi\psi} y$ nach Definition 2.23 gleichbedeutend mit einer der folgenden Bedingungen:

1. $x \sim_\varphi y$
2. Es gibt ein $z \in \bigcup_{i=0}^n \mathcal{K}_i$ mit $x \sim_\varphi z$ und $\psi(z) \sim_\varphi y$
3. Es gibt ein $z \in \bigcup_{i=0}^n \mathcal{K}_i$ mit $x \sim_\varphi z$ und $\psi^{-1}(z) \sim_\varphi y$.

Der erste Fall ist unproblematisch. Die beiden anderen Fälle unterscheiden sich nur darin, ob ψ oder ψ^{-1} verwendet wird, lassen sich daher analog behandeln. Im Folgenden werden wir den Fall für ψ weiter ausführen.

Falls eines der $\sim_\varphi$ zu $\sim$ reduziert, sind wir fertig. Falls nicht, so lernen wir aus Definition 2.23, dass der Übergang durch φ oder φ^{-1} vorstatten geht. Da es auch hier nicht darauf ankommt, welche dieser beiden Abbildungen wir φ nennen, können wir ohne Beschränkung der Allgemeinheit davon ausgehen, dass es die erste Abbildung ist. Damit müssen wir noch zwei Fälle untersuchen.

1. Es gebe $w_1, w_2 \in \bigcup_{i=0}^n \mathcal{K}_i$ mit

$$\begin{array}{rcl} x & \sim & w_1 \\ \varphi(w_1) & \sim & z \\ \psi(z) & \sim & w_2 \\ \varphi(w_2) & \sim & y \end{array}$$

Falls $(\mathcal{S}_\varphi)_\psi$ existiert, sind wir fertig, da dann $w_1 \sim_{\varphi\psi} w_2$ gilt. Da w_1 und w_2 im Definitionsbereich von φ liegen, liegen sie in einem gemeinsamen Simplex, sind also schon gleich. Damit können wir $x \sim_\varphi y$ folgern.

Falls $(\mathcal{S}_\psi)_\varphi$ existiert, kann man erkennen, dass $w_1 \sim_\varphi z$ und $\psi(z) \sim w_2$ gilt, weshalb wir $w_1 \sim_{\varphi\psi} w_2$ haben. Wie eben folgt $w_1 = w_2$ und damit $x \sim_\varphi y$.

2. Es gebe $w_1, w_2 \in \bigcup_{i=0}^n \mathcal{K}_i$ mit

$$\begin{array}{rcl} x & \sim & w_1 \\ \varphi(w_1) & \sim & z \qquad w_2 \sim y \\ \psi(z) & \sim & \varphi(w_2) \end{array}$$

Falls $(\mathcal{S}_\varphi)_\psi$ existiert, sind wir fertig, da dann $w_1 \sim_{\psi\varphi} w_2$ gilt. Da w_1 und w_2 im Definitionsbereich von φ liegen, liegen sie in einem gemeinsamen Simplex, sind also schon gleich. Damit können wir $x \sim y$ folgern.

Falls $(\mathcal{S}_\psi)_\varphi$ existiert, kann man erkennen, dass $\varphi(w_1) \sim_\psi \varphi(w_2)$ gilt. Da $\varphi(w_1)$ und $\varphi(w_2)$ in einem gemeinsamen Simplex liegen, sind sie gleich. Daher folgt $w_1 = w_2$ und $x \sim y$. $\square$

Das Lemma 2.43 zeigt, dass die resultierende Äquivalenzrelation aus einer doppelten Erweiterung nicht von der Reihenfolge dieser Erweiterungen abhängt, solange das Ergebnis einen wohldefinierten n -Überlagerungskomplex liefert. Unter dieser Voraussetzung können wir die Kommutativität zweier Erweiterungen zeigen.

Lemma 2.44. *Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein n -Überlagerungskomplex. Seien φ, ψ zwei Identifikationen auf $\mathcal{S}$, sodass die doppelte Erweiterung $(\mathcal{S}_\varphi)_\psi$ ein n -Überlagerungskomplex ist. Dann sind auch $\mathcal{S}_\psi$ und $(\mathcal{S}_\psi)_\varphi$ bereits n -Überlagerungskomplexe und es gilt $(\mathcal{S}_\psi)_\varphi = (\mathcal{S}_\varphi)_\psi$.*

Beweis. Nach Lemma 2.43 sind $\sim_{\varphi\psi}$ und $\sim_{\psi\varphi}$ identisch. Damit ist $(\mathcal{S}_\psi)_\varphi$ genau dann ein n -Überlagerungskomplex, wenn $(\mathcal{S}_\varphi)_\psi$ einer ist. Aus der Gleichheit der Äquivalenzrelationen folgt die Gleichheit der n -Überlagerungskomplexe.

Um nachzuweisen, dass $\mathcal{S}_\psi$ ein n -Überlagerungskomplex ist, verwenden wir Satz 2.32, der die möglichen Erweiterungen charakterisiert. Wäre dieser Satz nicht erfüllt, so gäbe es verschiedene $p, q \in \mathcal{K}_0$ und $z \in \bigsqcup_{i=1}^n \mathcal{K}_i$, sodass $[p]_{\sim\psi} \prec [z]_{\sim\psi}$ und $[q]_{\sim\psi} \prec [z]_{\sim\psi}$, aber auch $p \sim_\psi$ gelten (also ein Gegenbeispiel zur Irreduzibilität von $\mathcal{S}_\psi$). Da $\mathcal{S}_\psi$ eine Teilüberlagerung von $(\mathcal{S}_\psi)_\varphi$ ist, wäre dies dann auch ein Gegenbeispiel zur Irreduzibilität von $(\mathcal{S}_\psi)_\varphi$. Da es sich bei $(\mathcal{S}_\psi)_\varphi$ aber um einen n -Überlagerungskomplex handelt, ist dies ein Widerspruch. Folglich muss Satz 2.32 erfüllt sein und $\mathcal{S}_\psi$ ist ein n -Überlagerungskomplex. $\square$

Das Resultat aus Lemma 2.44, dass zwei verschiedene Erweiterungen in beliebiger Reihenfolge ausgeführt werden können (sofern das Ergebnis wohldefiniert ist), lässt sich durch ein Induktionsargument auf den allgemeinen endlichen Fall übertragen.

Satz 2.45. *Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein n -Überlagerungskomplex. Seien $\varphi_1, \dots, \varphi_m$ Identifikationen auf $\mathcal{S}$, sodass die mehrfache Erweiterung $\mathcal{S}_{\varphi_1 \dots \varphi_m}$ ein n -Überlagerungskomplex ist.*

Für jede Permutation $\pi \in S_m$ ist dann auch $\mathcal{S}_{\varphi_{\pi(1)}, \dots, \varphi_{\pi(m)}}$ ein n -Überlagerungskomplex und $\mathcal{S}_{\varphi_1 \dots \varphi_m} = \mathcal{S}_{\varphi_{\pi(1)} \dots \varphi_{\pi(m)}}$.

Beweis. Wir führen den Beweis per Induktion über m . Der Fall $m = 2$ wurde bereits in Lemma 2.44 bewiesen.

Im Induktionsschritt von m auf $m+1$ unterscheiden wir zwei Fälle. Falls $\pi(m+1) = m+1$, sind wir fertig, da wir nach Induktionsvoraussetzung die übrigen Erweiterungen beliebig vertauschen können.

Andernfalls vertauschen wir φ_m und $\varphi_{\pi(m+1)}$ in $\mathcal{S}_{\varphi_1 \dots \varphi_m}$ (das ist nach Induktionsvoraussetzung möglich). Wir können daher von $\pi(m+1) = m$ ausgehen. Da $\mathcal{S}_{\varphi_1 \dots \varphi_{m+1}}$ eine doppelte Erweiterung von $\mathcal{S}_{\varphi_1 \dots \varphi_{m-1}}$ ist, gilt nach Lemma 2.44 die Gleichheit $\mathcal{S}_{\varphi_1 \dots \varphi_{m-1} \varphi_{m+1} \varphi_m} = \mathcal{S}_{\varphi_1 \dots \varphi_{m+1}}$. Damit sind wir im ersten Fall und können die Induktionsvoraussetzung anwenden. $\square$

Wir haben diesen Abschnitt damit begonnen, die kombinatorische Explosion bei der Bestimmung aller möglichen Faltungen zu beklagen. Mit Satz 2.45 können wir die Anzahl der Möglichkeiten etwas einschränken, da es zur Konstruktion einer Faltung genügt, sich die verwendeten Identifikationen (als Menge) zu merken. Wir bemerken aber, dass die Anzahl der Möglichkeiten auch unter Verwendung dieses Satzes noch immer sehr schnell ansteigt.

2.4 Entfalten

Während wir uns im letzten Abschnitt mit dem Zusammenfalten von n -Überlagerungskomplexen auseinandergesetzt haben, widmet wir uns in diesem Abschnitt ihrer Entfaltung. Unser Hauptresultat (Bemerkung 2.51) wird sein, dass die Menge der möglichen Entfaltungen eine natürliche Verbandsstruktur besitzt, also eine deutlich reichhaltigere Struktur als die Faltungen aus Abschnitt 2.3.

Um dorthin zu gelangen, müssen wir verstehen, was „Entfaltung“ bedeutet. Dabei stellen wir eine Asymmetrie zwischen dem Zusammenfallen und dem Entfallen fest. Denn um zu wissen, welche Äquivalenzklassen *vereinigt* werden sollen (wie beim Zusammenfallen), müssen wir nur je ein Element aus jeder Klasse angeben (daher funktionieren Identifikationen), aber um eine Äquivalenzklasse zu *zerteilen* (wie es beim Entfallen geboten ist), muss die gesamte Partition angegeben werden. Wir können also nicht erwarten, dass das Entfallen ebenso einfach wie das Zusammenfallen beschrieben werden kann.

Bei der Untersuchung von Faltungen haben wir einige Resultate erzielt, indem wir jede Erweiterung in primitive Erweiterungen zerlegt haben. Analog versuchen wir „primitive Entfaltungen“ zu definieren, aus denen jede Entfaltung aufgebaut ist. Die Situation, die mit den primitiven Erweiterungen in größter Analogie steht, ist die Aufspaltung von genau einer Äquivalenzklasse. Solche Entfaltungen bezeichnen wir als primitive Spaltungen.

Definition 2.46 (primitive Spaltung). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim^{\mathcal{S}})$ ein n -Überlagerungskomplex. Sei $f \in \mathcal{K}_j$ (mit $0 \leq j \leq n$). Der n -Überlagerungskomplex $\mathcal{R} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim^{\mathcal{R}})$ heißt **primitive Spaltung (von $\mathcal{S}$) an $[f]$** , falls gilt

- $\mathcal{R}$ ist eine Teilüberlagerung von $\mathcal{S}$
- Für $x \in \bigcup_{i=0}^n \mathcal{K}_i$ mit $x \not\sim^{\mathcal{S}} f$ gilt $[x]_{\mathcal{S}} = [x]_{\mathcal{R}}$. (Nur $[f]_{\mathcal{S}}$ wird aufgespalten.)

Die Partition $\{[x]_{\mathcal{R}} \mid x \sim^{\mathcal{S}} f\}$ heißt **Spaltungspartition von $\mathcal{R}$** .

Definition 2.46 legt besonderen Fokus auf die Spaltungspartition, in die eine Äquivalenzklasse $[f]_{\mathcal{S}}$ zerfällt. Das liegt daran, dass nicht jede Partition von $[f]_{\mathcal{S}}$ eine primitive Spaltung liefert. Um dies zu verstehen, betrachten wir die Einschränkungen von Abwärtskompatibilität und Irreduzibilität, denen n -Überlagerungskomplexe gemäß Definition 2.17 unterliegen.

Während beim Zusammenfallen die Abwärtskompatibilität einfach zu erfüllen ist und die Irreduzibilität für Komplikationen sorgt, ist es es beim Entfallen genau umgekehrt. Die Irreduzibilität kann durch ein *Aufspalten* von Äquivalenzklassen niemals gefährdet werden, die Abwärtskompatibilität aber ständig. Damit zwei Elemente getrennt werden können, dürfen sie nicht durch Elemente höherer Dimension „zusammengehalten“ oder „gebunden“ werden. Diesen Begriff der „Bindung“ formalisieren wir in Definition 2.47.

Definition 2.47 (Bindungsrelation). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein n -Überlagerungskomplex. Seien $f, g \in \mathcal{K}_j$ ($0 \leq j \leq n$) mit $f \sim g$ gegeben.

Wir nennen das Paar (f, g) **gebunden**, falls es $x, y \in \mathcal{K}_j$ (für $i < j \leq n$) mit $f \prec x$, $g \prec y$ und $x \sim y$ gibt. Die dadurch definierte Relation heißt **Bindungsrelation**.

Bemerkung 2.48. Die Bindungsrelation ist eine Äquivalenzrelation auf jeder $\sim$ -Äquivalenzklasse.

Wir führen in Satz 2.49 aus, dass sich primitive Spaltungen mit der Bindungsrelation charakterisieren lassen.

Satz 2.49. Seien $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim^{\mathcal{S}})$ ein n -Überlagerungskomplex und $f \in \mathcal{K}_j$ (mit $0 \leq j \leq n$).

Eine Partition $\uplus_{s=1}^m M_s$ von $[f]_{\mathcal{S}}$ ist genau dann Spaltungspartition einer primitiven Spaltung von $[f]_{\mathcal{S}}$, wenn sie mit der Bindungsrelation verträglich ist (d. h. falls sie gröber als die Partition der Bindungs-Äquivalenzklassen ist).

Beweis. 1. Sei die primitive Spaltung $\mathcal{R}$ gegeben. Wir betrachten die $\sim^{\mathcal{R}}$ -Klasse von x mit $x \sim^{\mathcal{S}} f$. Für jedes y , das mit x gebunden ist, gibt es $u, v \in \mathcal{K}_l$ (für $j < l \leq n$) mit $x \prec u$, $y \prec v$ und $u \sim^{\mathcal{S}} v$, d. h. $u \sim^{\mathcal{R}} v$. Folglich gilt $x \sim^{\mathcal{R}} y$ aufgrund der Abwärtskompatibilität.

2. Sei nun die Partition $\uplus_{j=1}^m M_j$ gegeben, die mit der Bindungsrelation kompatibel ist. Wir müssen zeigen, dass das so konstruierte $\mathcal{R}$ tatsächlich ein n -Überlagerungskomplex ist (dann ist es nach Konstruktion Teilüberlagerung von $\mathcal{S}$).

Da $\sim^{\mathcal{R}}$ eine Verfeinerung von $\sim^{\mathcal{S}}$ ist, ist die Irreduzibilität erfüllt. Zur Überprüfung der Abwärtskompatibilität genügt es, die $u \sim^{\mathcal{R}} v$ mit $u, v \in \mathcal{K}_l$ für $l > j$ zu betrachten.

Da nur $[f]_{\mathcal{S}}$ aufgespalten wird, müssen wir nur $x, y \in [f]_{\mathcal{S}}$ überprüfen. Falls es solche $x \prec u$ und $y \prec v$ gibt, ist (x, y) gebunden, also $x \sim^{\mathcal{R}} y$. Damit ist die Abwärtskompatibilität erfüllt und $\mathcal{R}$ ist ein n -Überlagerungskomplex. $\square$

Mit anderen Worten: Die Äquivalenzklassen der Bindungsrelation bilden die feinste mögliche Aufspaltung einer $\sim$ -Äquivalenzklasse, die die Abwärtskompatibilität erfüllen.

Da die primitiven Spaltungen an $[f]$ durch Partitionen indiziert sind, erhalten sie von diesen eine kanonische Verbandsstruktur. Wir erinnern uns daher an die Definition eines Verbandes.

Definition 2.50 (Verband). Die partielle Ordnung⁷ $(P, \leq)$ heißt **Verband**, falls gilt:

1. Zu je zwei Elementen $a, b \in P$ gibt es ein eindeutiges Supremum, d. h. ein $x \in P$ mit

- $a \leq x$ und $b \leq x$
- Jedes $w \in P$ mit $a \leq w$ und $b \leq w$ erfüllt $x \leq w$.

Dieses Element heißt auch der **join** von a und b und wird mit $a \vee b$ bezeichnet.

2. Zu je zwei Elementen $a, b \in P$ gibt es ein eindeutiges Infimum, d. h. ein $x \in P$ mit

- $x \leq a$ und $x \leq b$
- Jedes $w \in P$ mit $w \leq a$ und $w \leq b$ erfüllt $w \leq x$.

Dieses Element heißt auch der **meet** von a und b und wird mit $a \wedge b$ bezeichnet.

Nachdem wir uns die Definition eines Verbandes ins Gedächtnis gerufen haben, können wir die Verbandsstruktur auf den möglichen Spaltungspartitionen beschreiben. Diese lassen sich über Satz 2.49 auf die primitiven Spaltungen eines n -Überlagerungskomplexe übertragen.

Bemerkung 2.51. Die Menge aller Partitionen, die mit der Bindungsrelation verträglich sind, bildet einen Verband, wobei

⁷also eine binäre Relation, die reflexiv, transitiv und antisymmetrisch ist

- $\{M_1, \dots, M_m\} \leq \{N_1, \dots, N_n\}$ soll genau dann gelten, wenn es zu jedem M_i ein N_j mit $M_i \subseteq N_j$ gibt.
- Der Meet ist definiert als

$$\{M_1, \dots, M_m\} \wedge \{N_1, \dots, N_n\} := \{M_i \cap N_j \mid 1 \leq i \leq m, 1 \leq j \leq n\} \setminus \{\emptyset\}$$

- Der Join ist definiert als

$$\{M_1, \dots, M_m\} \vee \{N_1, \dots, N_n\} := \bigwedge_{P \in \mathcal{P}} P,$$

wobei $\mathcal{P}$ aus allen mit der Bindungsrelation kompatiblen Partitionen $\{P_k\}$ besteht, die $\{M_i\} \leq \{P_k\}$ und $\{N_j\} \leq \{P_k\}$ erfüllen.

Dadurch wird eine Verbandsstruktur auf den primitiven Entfaltungen eines n -Überlagerungskomplexes induziert, dessen partielle Ordnung durch „Teilüberlagerung“ gegeben ist.

Wir haben in Bemerkung 2.51 erkannt, dass die primitiven Spaltungen eines n -Überlagerungskomplexes eine Verbandsstruktur haben und damit die Entfaltung von n -Überlagerungskomplexen beschrieben.

In endlichen Verbänden gibt es ein kleinstes und ein größtes Element. Da das größte Element dieses Verbandes dem ursprünglichen n -Überlagerungskomplex entspricht (also keiner Entfaltung), entspricht das kleinste Element der maximal möglichen Entfaltung. Alle anderen Elemente (außer dem maximalen) entsprechen damit nicht-maximalen Entfaltungen. Um eine eindeutige Entfaltung zu definieren, müssen wir für jede mögliche Äquivalenzklasse ein Element des entsprechenden Verbandes angeben. Von allen nicht-trivialen Entfaltungen lässt sich nur die maximale Entfaltung im Allgemeinen eindeutig identifizieren. Daher ist es plausibel, dieses kleinste Element zu wählen, wenn man eine eindeutige Entfaltung definieren möchte.

Wenn wir diese Überlegung aber auf Papierfaltungen übertragen, stellen wir ein Problem fest: In einem 2-Überlagerungskomplex ist die Bindungsrelation auf den Flächen $\mathcal{K}_2$ trivial. Damit lassen sich die Flächen beliebig aufspalten – im Widerspruch zu unserer Anschauung, in der es möglich ist, dass eine Fläche zwischen zwei anderen „eingequetscht“ ist. Das deutet darauf hin, dass unsere Modellierung an dieser Stelle unvollständig ist. Wir werden uns in Kapitel 3 ausführlich mit dieser Thematik beschäftigen.

2.5 Einbettungen

Im gesamten Kapitel 2 haben wir bisher nur über abstrakte Faltungen gesprochen (es war nie relevant, wie diese Komplexe konkret im $\mathbb{R}^3$ aussehen). Dadurch hat unser Modell Möglichkeiten, die ein Modell auf Basis des $\mathbb{R}^3$ niemals haben kann. Während man in vielen Anwendungen die Faltmuster (also 2-Überlagerungskomplexe) aus dem $\mathbb{R}^3$ nimmt, ist dies nicht die einzige Möglichkeit zur Konstruktion eines 2-Überlagerungskomplexes.

Es ist nämlich möglich, simpliziale Flächen abstrakt zu konstruieren, wie z. B. [BNPS] mithilfe von Gruppentheorie demonstriert. Wenn man mit einem 2-Überlagerungskomplex

aus einer solchen Quelle startet, ist es nicht direkt klar, ob man diesen im $\mathbb{R}^3$ darstellen kann. An dieser Stelle benötigen wir eine abstrakte Beschreibung von Faltungen, da es nicht möglich ist, die „Faltung“ dieses Komplexes im $\mathbb{R}^3$ auszuführen.

Wir haben am Ende des letzten Abschnitts 2.4 zu Entfaltungen zwar eingesehen, dass das Modell dieses Kapitels unvollständig ist. Dennoch können wir die Frage nach möglichen Darstellungen in den $\mathbb{R}^3$ hier schon stellen. Die genauere Analyse, auf die wir in Kapitel 4 eingehen werden, wird auf den Ideen dieses Abschnitts aufbauen.

Die erste Frage ist, was wir mit einer Darstellung oder Einbettung eines 2-Überlagerungskomplexes in den $\mathbb{R}^3$ eigentlich meinen. Da eine Einbettung normalerweise injektiv ist, aber zusammengefaltete Flächen an denselben Teilmengen des $\mathbb{R}^3$ liegen, gibt es hier einen Konflikt. Um diesen aufzulösen, betrachten wir einen lokal simplizialen 2-Komplex mit Ecken $\mathcal{K}_0$, Kanten $\mathcal{K}_1$ und Flächen $\mathcal{K}_2$. Wir überlegen uns zunächst, wie diese Simplizes abgebildet werden müssen und auf Basis dieser Information, wie wir den Konflikt lösen können.

Jeder Ecke soll ein Punkt zugeordnet werden. Eine Kante ist durch ihre beiden Eckpunkte festgelegt⁸. Die Flächen sollen auf 2-Simplizes, also Dreiecke abgebildet werden. Auch diese sind durch ihre Eckpunkte festgelegt. Im Allgemeinen bildet man also einen k -Simplex auf die konvexe Hülle seiner Eckpunkte ab (nachdem man diese eingebettet hat).

An dieser Stelle bekommen wir Schwierigkeiten mit unserer Intuition, falls es zwei verschiedene Simplizes mit den gleichen Eckpunkten gibt. Dies tritt in Fall (4) aus Beispiel 2.8 ein, wo wir zwei verschiedene Kanten mit den gleichen Eckpunkten vorliegen haben. Wenn wir eine Kante also auf die konvexe Hülle ihrer Eckpunkte abbilden, würden diese beiden Kanten identisch eingebettet werden. Um diese Fälle allgemein auszuschließen, beschreiben wir die kritischen Fälle (als Anomalien) gesondert.

Definition 2.52 (Anomalie). Sei $n > 0$ und $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex. Zwei verschiedene Äquivalenzklassen $[f], [g] \in \mathcal{K}_i / \sim$ (für $0 \leq i \leq n$) bilden eine i -**Anomalie**, falls

$$\{[p] \in \mathcal{K}_0 / \sim \mid [p] \prec [f]\} = \{[p] \in \mathcal{K}_0 \mid [p] \prec [g]\}$$

gilt. Die Menge aller Äquivalenzklassen, die paarweise eine i -Anomalie bilden, heißt **Anomalieklasse**.

Falls es zu $\mathcal{S}$ keine Anomalien gibt, heißt $\mathcal{S}$ **anomaliefrei**.

Im Fall von 2-Anomalien sagt diese Definition aus, dass zwei Dreiecke genau dann eine Anomalie bilden, wenn sie dieselben Punkte besitzen. Dadurch liegen sie an der gleichen Position im $\mathbb{R}^3$.

Falls Anomalien auftreten, ist eine Einbettung in den $\mathbb{R}^3$ nicht möglich, da z.B. zwei verschiedene Flächen auf die selbe Teilmenge des $\mathbb{R}^3$ abgebildet würden. Wenn wir diese Fälle ausschließen, können wir die Einbettung eines n -Überlagerungskomplexes allein durch die Bilder seiner Punkte festlegen.

⁸Wenn man krummlinige Einbettungen zulässt, würde man eine Kante auf eine glatte Kurve abbilden, deren Rand durch die Eckpunkte der Kante gegeben sind. Man kann die Definitionen in diese Richtung ausbauen, allerdings werden sie dadurch schwerer zu benutzen. Daher bleiben wir bei den einfachen Definitionen.

Wir müssen aber noch eine Variation der Injektivität garantieren, um den Namen „Einbettung“ zu legitimieren. Die Abbildung sollte „so injektiv wie möglich“ sein, d. h. zwei Simplizes x und y sollten sich in der Einbettung genau dann schneiden, wenn sie es genauso im n -Überlagerungskomplex tun. Damit gelangen wir zur vollständigen Definition der Einbettung, die wir in Definition 2.53 zusammenfassen:

Definition 2.53 (simpliciale Einbettung). *Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein anomaliefreier lokal simplicialer n -Überlagerungskomplex und sei eine Abbildung $\iota : \mathcal{K}_0 / \sim \rightarrow \mathbb{R}^d$ gegeben. Diese induziert eine Abbildung*

$$\bar{\iota} : \text{Pot}(\mathcal{K}_0 / \sim) \rightarrow \text{Pot}(\mathbb{R}^d), M \mapsto \text{ConvexHull}(\{\iota(p) | p \in M\}).$$

Für $x \in \bigsqcup_{i=0}^n \mathcal{K}_i$ definiere

$$\mathcal{K}_0^{\preceq x} := \{[p]_{\sim} \in \mathcal{K}_0 \mid p \preceq x\}.$$

Dann heißt ι (**simpliciale**) **Einbettung** von $\mathcal{S}$ in den $\mathbb{R}^d$, falls für alle $x, y \in \bigsqcup_{i=0}^n \mathcal{K}_i$ gilt, dass

$$\bar{\iota}(\mathcal{K}_0^{\preceq x}) \cap \bar{\iota}(\mathcal{K}_0^{\preceq y}) = \bar{\iota}(\mathcal{K}_0^{\preceq x} \cap \mathcal{K}_0^{\preceq y}).$$

Definition 2.53 beruht im Wesentlichen auf der Idee, dass zwei Dreiecke des 2-Überlagerungskomplexes genau dann eine Kante gemeinsam haben, wenn sie auch in der Einbettung eine Kante gemeinsam haben. Zur Illustration wenden wir diese Bedingung auf die Menge der Punktäquivalenzklassen an, die infolgedessen paarweise verschieden abgebildet werden müssen.

Folgerung 2.54. *Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein anomaliefreier lokal simplicialer n -Überlagerungskomplex und $\iota : \mathcal{K}_0 / \sim \rightarrow \mathbb{R}^d$ eine simpliciale Einbettung. Dann ist ι injektiv.*

Beweis. Seien $x, y \in \mathcal{K}_0$ mit $x \not\sim y$. Dann sind $\mathcal{K}_0^{\preceq x}$ und $\mathcal{K}_0^{\preceq y}$ disjunkt. Folglich müssen auch $\bar{\iota}(\mathcal{K}_0^{\preceq x})$ und $\bar{\iota}(\mathcal{K}_0^{\preceq y})$ disjunkt sein. Wegen $\bar{\iota}(\mathcal{K}_0^{\preceq x}) = \{\iota([x])\}$ und $\bar{\iota}(\mathcal{K}_0^{\preceq y}) = \{\iota([y])\}$ folgt die Behauptung. $\square$

Damit haben wir den Begriff der Einbettung definiert. Wir können jetzt zwei verschiedene Fragestellungen angehen:

1. Kann ein gegebener n -Überlagerungskomplex (gegebenenfalls nach einigen Faltungen) eingebettet werden?
2. Lässt sich die Faltung eines eingebetteten n -Überlagerungskomplexes auch einbetten?

Wir werden einen Teil der ersten Frage in Abschnitt 2.5.1 untersuchen, indem wir die Frage von allgemeinen n -Überlagerungskomplexen auf Simplizes zurückführen. Die zweite Frage ist deutlich komplizierter, sodass wir in Abschnitt 2.5.2 nur einige generelle Einschränkungen angeben können.

2.5.1 Färbbarkeit

In diesem Abschnitt beschäftigen wir uns mit der Frage, ob ein gegebener n -Überlagerungskomplex eingebettet werden kann, falls er zuvor gefaltet werden darf. Wir können uns also auch fragen, ob sich eine der Erweiterungen des n -Überlagerungskomplexes einbetten lässt.

Da Einbettungen für Komplexe mit weniger Flächenklassen einfacher zu bestimmen sind, werden wir die minimalen Faltzustände von n -Überlagerungskomplexen klassifizieren. Wir suchen also die n -Überlagerungskomplexe, die nicht Teilüberlagerung eines anderen n -Überlagerungskomplexes sind. Es wird sich zeigen, dass diese minimalen n -Überlagerungskomplexe sämtlich anomaliefrei sind, da es sich im Wesentlichen um Simplizes handelt.

Ein zentrales Hilfsmittel für diese Klassifikation ist das Konzept der Färbung. Dabei ordnen wir jedem Punkt des lokal simplizialen Komplexes eine Farbe zu, sodass die Eckpunkte jeder Kante verschieden gefärbt sind. Um eine Kante zu färben, benötigen wir zwei Farben, zur Färbung eines Dreiecks drei und für einen Tetraeder vier.

Definition 2.55 (Färbung). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer⁹ n -Überlagerungskomplex. Er heißt **k -färbbar** (mit $k \in \mathbb{N}$), falls es eine Abbildung $f : \mathcal{K}_0 \rightarrow \{1, \dots, k\}$ gibt, die die folgenden Eigenschaften erfüllt:

1. (verträglich mit $\sim$) Für $x, y \in \mathcal{K}_0$ mit $x \sim y$ gilt $f(x) = f(y)$.
2. (verschiedene Farben an den Eckpunkten jeder Kante) Für verschiedene $x, y \in \mathcal{K}_0$, sodass es ein $z \in \mathcal{K}_1$ mit $x \prec z$ und $y \prec z$ gibt, gilt $f(x) \neq f(y)$.

Eine solche Abbildung f heißt **k -Färbung**.

Da in einem k -Simplex je zwei der $k + 1$ Punkte durch eine Kante verbunden sind, ist ein solcher k -Simplex automatisch $(k + 1)$ -färbbar.

Beispiel 2.56. Sei M eine endliche Menge mit $|M| = n + 1$. Dann ist der n -Simplex $((\text{Pot}_{i+1}(M))_{0 \leq i \leq n}, \subsetneq)$, aufgefasst als n -Überlagerungskomplex, $(n + 1)$ -färbbar, da je zwei Elemente aus M in einer Kante vorkommen.

Wir werden zeigen, dass ein $(k - 1)$ -Simplex in einem gewissen Sinne der allgemeinste k -färbbare n -Überlagerungskomplex ist. Das klingt zunächst einmal unrealistisch, da es größere k -färbbare Komplexe gibt. Beispielsweise sind zwei Dreiecke (also 2-Simplizes), die sich eine Kante teilen, auch 3-färbbar. Allerdings lassen sich diese beiden Dreiecke zusammenfallen, sodass nur ein Dreieck übrig bleibt.

Wenn wir Faltungen zulassen wollen, um Färbungen zu beschreiben, müssen wir überprüfen, ob diese beiden Begriffe miteinander verträglich sind. Es ist vergleichsweise leicht zu zeigen, dass die Färbung eines gefalteten Zustandes eine Färbung des entfalteten Zustandes induziert, wie Lemma 2.57 demonstriert.

Lemma 2.57. Seien $\mathcal{S}, \mathcal{T}$ zwei n -Überlagerungskomplexe und $\mathcal{S}$ sei eine Teilüberlagerung von $\mathcal{T}$. Falls $\mathcal{T}$ bereits k -färbbar ist, dann ist es auch $\mathcal{S}$.

⁹Wenn wir diese Bedingung weglassen, sollten wir fordern, dass verschiedene $x, y \in \mathcal{K}_0$, für die ein $z \in \mathcal{K}_n$ mit $x \prec z$ und $y \prec z$ existiert, verschiedene Farben haben.

Beweis. Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim^{\mathcal{S}})$ und $\mathcal{T} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim^{\mathcal{T}})$. Da $\mathcal{S}$ eine Teilüberlagerung von $\mathcal{T}$ ist, sind die Abbildungen

$$\begin{aligned}\alpha_i : \mathcal{K}_i / \sim^{\mathcal{S}} &\rightarrow \mathcal{K}_i / \sim^{\mathcal{T}}, \\ [x]_{\sim^{\mathcal{S}}} &\mapsto [x]_{\sim^{\mathcal{T}}}\end{aligned}$$

wohldefiniert.

Sei $f : \mathcal{K}_0 / \sim^{\mathcal{T}} \rightarrow \{1, \dots, k\}$ eine k -Färbung von $\mathcal{T}$. Definiere die Abbildung $g : \mathcal{K}_0 / \sim^{\mathcal{S}} \rightarrow \{1, \dots, k\}$ durch $g := f \circ \alpha_0$. Wir zeigen, dass g eine k -Färbung von $\mathcal{S}$ ist. Seien $x, y \in \mathcal{K}_0$ durch eine Kante verbunden. Nach Voraussetzung gilt $f([x]_{\sim^{\mathcal{T}}}) \neq f([y]_{\sim^{\mathcal{T}}})$, also auch

$$g([x]_{\sim^{\mathcal{S}}}) = f([x]_{\sim^{\mathcal{T}}}) \neq f([y]_{\sim^{\mathcal{T}}}) = g([y]_{\sim^{\mathcal{S}}}).$$

Damit ist g eine k -Färbung von $\mathcal{S}$. □

Nachdem wir gezeigt haben, dass sich Färbungen gefalteter Zustände auf Färbungen ungefalteter Zustände übertragen, müssen wir noch die umgekehrte Richtung überprüfen, um die Verträglichkeit von Faltung und Färbung nachzuweisen. Hier ist es zwar nicht der Fall, dass jede Faltung die Färbung des n -Überlagerungskomplexes erhält, aber es gibt eine Faltung, die von der gegebenen Färbung ausgezeichnet ist. Bei dieser werden alle gleichfarbigen Punkte zusammengefasst.

Definition 2.58 (Färbungsüberlagerung). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex und f eine Färbung von $\mathcal{S}$. Definiere eine Äquivalenzrelation $\sim^\Delta$ auf $\bigsqcup_{0 \leq i \leq n} \mathcal{K}_i$ wie folgt:

$$x \sim^\Delta y \Leftrightarrow \{f(p) \mid p \in \mathcal{K}_0, p \preceq x\} = \{f(q) \mid q \in \mathcal{K}_0, q \preceq y\}$$

Dann ist $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim^\Delta)$ ein anomaliefreier lokal simplizialer n -Überlagerungskomplex, die sogenannte **Färbungsüberlagerung** (von $\mathcal{S}$ bezüglich f).

Wohldefiniertheit. Wir müssen zeigen, dass ein n -Überlagerungskomplex vorliegt.

Zum Nachweis der Dimensionalität betrachten wir für jedes $x \in \mathcal{K}_i$ die von der Färbung induzierte Abbildung auf den Punkten, die in x liegen:

$$\bar{f} : \{p \in \mathcal{K}_0 \mid p \preceq x\} \rightarrow \{f(p) \mid p \in \mathcal{K}_0, p \preceq x\}$$

Diese Abbildung ist nach Konstruktion surjektiv. Da $\mathcal{S}$ lokal simplizial ist, sind je zwei Elemente aus $\{p \in \mathcal{K}_0 \mid p \preceq x\}$ durch eine Kante verbunden, weshalb ihre Farben verschieden sind. Folglich ist die Abbildung auch injektiv, insgesamt also bijektiv.

Aufgrund der lokalen Simplizität enthält die Menge $\{p \in \mathcal{K}_0 \mid p \preceq x\}$ genau $i + 1$ Elemente, ebenso also die Menge der Farben $\{f(p) \mid p \in \mathcal{K}_0, p \preceq x\}$. Daraus folgt sofort die Dimensionalität.

Wir wollen die starke Abwärtskompatibilität nachweisen (vgl. Bemerkung 2.18). Dazu betrachten wir zwei $x, y \in \bigsqcup_{0 \leq i \leq n} \mathcal{K}_i$ mit $x \sim^\Delta y$. Für jedes $a \prec x$ gilt

$$\{f(p) \mid p \in \mathcal{K}_0, p \preceq a\} \subseteq \{f(p) \mid p \in \mathcal{K}_0, p \preceq x\} = \{f(q) \mid q \in \mathcal{K}_0, q \preceq y\}.$$

Da die Färbung auf einem Simplex eine Bijektion induziert, gibt es genau eine korrespondierende Teilmenge T von $\{q \in \mathcal{K}_0 \mid q \preccurlyeq y\}$. Da $\mathcal{S}$ lokal simplizial ist, gibt es zu dieser Teilmenge T genau ein $b \prec y$, sodass $T = \{q \in \mathcal{K}_0 \mid q \preccurlyeq b\}$ gilt. Nach Konstruktion haben wir $a \sim^\Delta b$. $\square$

Wir haben damit gezeigt, dass Faltung und Färbung von n -Überlagerungskomplexen zum großen Teil miteinander verträglich sind. Daher sind wir legitimiert, Faltungen zur Untersuchung von Färbungen zu verwenden. Wir wollen darauf hinaus, dass ein lokal simplizialer n -Überlagerungskomplex genau dann k -färbbar ist, wenn seine Färbungsüberlagerung in einem $(k-1)$ -Simplex liegt.

An dieser Stelle stoßen wir auf ein Problem, da es nicht klar ist, inwiefern ein 2-Überlagerungskomplex aus z. B. 4 Dreiecken in einem einzelnen Dreieck liegen kann. Um dieses Dilemma zu lösen, betrachten wir Papierfaltungen. Ein Papiermuster aus gleichseitigen Dreiecken ist genau dann 4-färbbar, wenn es sich so zusammenfalten lassen, dass es wie ein Teil eines Tetraeders aussieht.

In dieser Umformulierung erkennen wir das fehlende Puzzlestück: Wir interessieren uns für das *Aussehen* der Faltung, nicht deren innere Struktur. Um diese äußere Form aus einem n -Überlagerungskomplex herauszuholen, müssen wir jede Äquivalenzklasse von zusammengefalteten Elementen als ein Element auffassen. Wir erhalten dadurch den Basiskomplex eines n -Überlagerungskomplexes.

Definition 2.59 (Basiskomplex). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein n -Überlagerungskomplex. Der homogene n -Komplex $((\mathcal{K}_i / \sim)_{0 \leq i \leq n}, \prec)$ mit $\mathcal{K}_i / \sim := \{[x] \mid x \in \mathcal{K}_i\}$ und $\prec$ wie in Definition 2.17 festgelegt, heißt **Basiskomplex von $\mathcal{S}$** , bezeichnet mit $BK[\mathcal{S}]$.

Der Basiskomplex eines Faltzustandes beschreibt ausschließlich, wie dieser Zustand von außen aussieht, wenn man die interne Faltstruktur ignoriert. Damit handelt es sich dabei um eine komplementäre Sichtweise auf n -Überlagerungskomplexe, die wir bislang nur als gefaltete Struktur aufgefasst haben.

Wir können jetzt sagen, dass ein n -Überlagerungskomplex genau dann k -färbbar ist, wenn der Basiskomplex seiner Färbungsüberlagerung in einem $(k-1)$ -Simplex enthalten ist. Dabei stellen wir die „enthalten sein“-Beziehung durch einen Komplex-Monomorphismus (vgl. Definition 2.10) dar. Da die Färbungsüberlagerung nur bei Anwesenheit einer Färbung vorliegt, müssen wir den Satz 2.60 ein wenig allgemeiner formulieren.

Satz 2.60. Sei $\mathcal{S}$ ein lokal simplizialer n -Überlagerungskomplex. Dann ist $\mathcal{S}$ genau dann k -färbbar, wenn es einen n -Überlagerungskomplex $\mathcal{T}$ mit den folgenden Eigenschaften gibt:

- $\mathcal{S}$ ist eine Teilüberlagerung von $\mathcal{T}$.
- Es gibt einen Komplex-Monomorphismus von $BK[\mathcal{T}]$ in einen $(k-1)$ -Simplex.

Beweis. Sei $\mathcal{S}$ k -färbbar. Definiere $\mathcal{T}$ als die Färbungsüberlagerung von $\mathcal{S}$ bezüglich dieser Färbung. Nach Konstruktion hat $\mathcal{T}$ maximal k verschiedene Punktäquivalenzklassen, also hat $BK[\mathcal{T}]$ maximal k verschiedene Punkte. Damit es einen Komplex-Monomorphismus in einen $(k-1)$ -Simplex gibt, müssen wir zeigen, dass $BK[\mathcal{T}]$ anomaliefrei ist (sonst gäbe es

z. B. zwei Kanten mit den selben Punkten). Dies ist genau dann der Fall, wenn $\mathcal{T}$ anomaliefrei ist. Nach Definition 2.58 ist dies aber automatisch garantiert.

Umgekehrt existiere ein solcher n -Überlagerungskomplex $\mathcal{T}$. Nach Beispiel 2.56 ist der $(k-1)$ -Simplex k -färbbar. Da bei einem Komplex-Monomorphismus keine zusätzlichen Kanten in $BK[\mathcal{T}]$ hinzukommen können, ist auch $BK[\mathcal{T}]$ k -färbbar, also auch $\mathcal{T}$ selbst. Schließlich ist nach Lemma 2.57 auch $\mathcal{S}$ k -färbbar. $\square$

Wenden wir diesen Satz auf 3-färbbare Papierfaltungen aus Dreiecken an, so erhalten wir, dass diese sich stets so zusammenfalten lassen, dass sie in einem Dreieck liegen. Allerdings können wir ein Dreieck niemals auf eine Kante falten, demnach sind Papierfaltungen aus Dreiecken genau dann 3-färbbar, wenn man sie auf ein Dreieck zusammenfalten kann. Diese Erkenntnis gilt auch allgemein, wie Folgerung 2.61 demonstriert.

Folgerung 2.61. *Sei $\mathcal{S}$ ein lokal simplizialer n -Überlagerungskomplex. Dann ist $\mathcal{S}$ genau dann $(n+1)$ -färbbar, wenn es einen n -Überlagerungskomplex $\mathcal{T}$ mit den folgenden Eigenschaften gibt:*

- $\mathcal{S}$ ist eine Teilüberlagerung von $\mathcal{T}$.
- $BK[\mathcal{T}]$ ist isomorph zu einem n -Simplex.

Beweis. Sei $\mathcal{S}$ $(n+1)$ -färbbar. Dann gibt es nach Satz 2.60 einen n -Überlagerungskomplex $\mathcal{T} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ und einen Komplex-Monomorphismus von $BK[\mathcal{T}]$ in einen n -Simplex. Da $\mathcal{T}$ ein n -Überlagerungskomplex ist, gibt es mindestens ein Element x in $\mathcal{K}_n$. Damit ist $\mathcal{T}_{\leq x}$ isomorph zu einem n -Simplex und alle anderen Elemente aus $\mathcal{K}_n$ sind äquivalent zu x .

Umgekehrt gebe es ein solches $\mathcal{T}$. Da jeder Komplex-Isomorphismus auch ein Komplex-Monomorphismus ist, folgt die Behauptung aus Satz 2.60. $\square$

Empirisch stellt man fest, dass man viele der 2-Überlagerungskomplexe, auf die man trifft, mit drei oder vier Farben färben kann. Gemäß Satz 2.60 lassen sich diese dann in die Form eines Tetraeders bringen, der sich in den $\mathbb{R}^3$ einbetten lässt. Falls mehr als vier Farben benötigt werden, muss die Situation genauer analysiert werden.

Dieses Ergebnis beruht aber darauf, dass das Modell der 2-Überlagerungskomplexe nicht in der Lage ist, die Reihenfolge von Papierschichten oder deren Überschneidungsfreiheit an den Kanten zu erfassen. Wir werden diese Information in Kapitel 3 hinzufügen und in Kapitel 4 analysieren, welche zusätzlichen Annahmen erfüllt sein müssen, damit das Resultat aus Satz 2.60 in diesem allgemeineren Kontext gilt.

2.5.2 Verträglichkeit mit Faltung

In diesem Abschnitt beschäftigen wir uns mit der Frage, ob die Einbettbarkeit eines n -Überlagerungskomplexes auch die Einbettbarkeit seiner Erweiterungen nach sich zieht (sofern diese anomaliefrei sind).

Wir möchten dabei nicht nur mehrere Einbettungen finden, sondern auch den Faltprozess selbst im $\mathbb{R}^3$ ausführen können. Da bei einem Faltprozess jede Fläche isometrisch transformiert wird, verlangen wir, dass ein Simplex unter allen Einbettungen so abgebildet wird, dass diese Bilder paarweise isometrisch sind.

Definition 2.62 (konsistente Einbettungen). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec)$ ein homogener n -Komplex und $\sim_1, \dots, \sim_m$ seien Äquivalenzrelationen auf $\biguplus_{i=0}^n \mathcal{K}_i \times \mathcal{K}_i$, sodass $(\mathcal{S}, \sim_j)$ anomaliefreie n -Überlagerungskomplexe sind (für $1 \leq j \leq m$).

Eine Familie von simplizialen Einbettungen $\iota_j : \mathcal{K}_0 / \sim_j \rightarrow \mathbb{R}^d$ (mit $1 \leq j \leq m$) heißt **konsistent mit** $\{(\mathcal{S}, \sim_j) | 1 \leq j \leq m\}$, falls für alle $x \in \biguplus_{i=0}^n \mathcal{K}_i$ die Elemente aus

$$\{\bar{\iota}_j(\mathcal{K}_0^{\prec x}) | 1 \leq j \leq m\}$$

isometrisch sind.

Das Konzept der konsistenten Einbettungen beschreibt also, dass jeder Simplex des Faltmusters während des gesamten Faltvorgangs nur isometrisch im Raum bewegt werden kann.

Da die möglichen Erweiterungen stark vom homogenen Grundkomplex abhängen (eingeschränkt von der Definition 2.23 der Erweiterung und Satz 2.32, der die Identifikationen charakterisiert, die zu einer gültigen Erweiterung führen), ist es schwer, allgemeine Aussagen über konsistente Einbettungen zu treffen. Wir haben es mit drei verschiedenen verkomplizierenden Phänomenen zu tun:

1. Nicht alle Identifikationen sind konstant auf dem Schnitt und liefern gültige Erweiterungen (charakterisiert durch Satz 2.32).
2. Nicht alle durch Erweiterungen erhaltenen n -Überlagerungskomplexe sind anomaliefrei.
3. Es ist nicht klar, dass die Einbettung einer Erweiterung überhaupt existiert (aufgrund möglicher globaler Überschneidungen).

Wir werden daher versuchen, möglichst allgemeine Bedingungen zu finden, unter denen konsistenten Einbettungen erlaubt sind.

Würden wir jede mögliche Identifikation zulassen (und für einen Moment die obigen Restriktionen außer Acht lassen), so müssen alle Simplizes gleicher Dimension isometrisch sein, da sie sich aufeinander falten lassen. Da Simplizes schon durch die Forderung, dass alle Kanten die gleiche Länge haben müssen, festgelegt sind (für Flächen nennt man dies „gleichseitige Dreiecke“), sind sie durch die Forderung gleicher Kantenlängen schon eindeutig festgelegt.

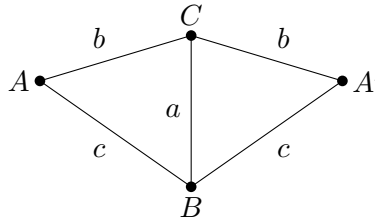
Bemerkung 2.63. Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein anomaliefreier lokal simplizialer n -Überlagerungskomplex und $\mathcal{M}$ die Menge aller anomaliefreien lokal simplizialen n -Überlagerungskomplexe, sodass $\mathcal{S}$ Teilüberlagerung von $\mathcal{T}$ ist.

Es sei für jedes $\mathcal{T} \in \mathcal{M}$ eine Einbettung $\iota_{\mathcal{T}}$ in den $\mathbb{R}^d$ gegeben. Falls es ein $a \in \mathbb{R}$ gibt, sodass für jede Kante $x \in \mathcal{K}_1$ und jede Einbettung $\iota_{\mathcal{T}}$ gilt, dass $\bar{\iota}_{\mathcal{T}}(\mathcal{K}_0^{\prec x})$ Länge a hat, dann ist diese Familie der Einbettungen konsistent bzgl. $\mathcal{M}$.

Folglich ist ein Muster aus gleichseitigen Dreiecken das Muster, das generisch die meisten Faltungen erlaubt. Häufig sind wir aber gar nicht an beliebigen Identifikationen interessiert, sondern z. B. nur an den n -Überlagerungskomplexen, die wir durch Nachbaridentifikationen (vgl. Definition 2.38) erhalten können. Relevant sind natürlich nur die Situationen, in denen kein gleichseitiges Dreiecksmuster gewählt werden muss.

Bei 1-Überlagerungskomplexen ergibt sich keine Änderung – wenn je zwei benachbarte Kanten dieselbe Länge haben und der Komplex zusammenhängend ist, haben alle Kanten dieselbe Länge und wir sind im Fall von Bemerkung 2.63.

Bei 2-Überlagerungskomplexen ist die Lage ein wenig komplizierter. Sobald die Kantenlängen eines Dreiecks festgelegt sind, sind durch die Nachbaridentifikationen auch die Kantenlängen der benachbarten Flächen festgelegt:



Interessant ist der Fall, in dem mehr Freiheiten als bloß gleichseitige Dreiecke bleiben. Um dies zu erhalten, betrachten wir anstelle der Kanten deren Eckpunkte. Da eine Kante (und damit ihre Länge) eindeutig durch ihre Eckpunkte beschrieben ist, genügt eine 3-Färbung der Ecken aus, um zu garantieren, dass wir drei verschiedene Kantenlängen wählen können (die natürlich ein Dreieck bilden müssen).

Lemma 2.64. *Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq 2}, \prec, \sim)$ ein 3-färbbarer, anomaliefreier, lokal simplizialer 2-Überlagerungskomplex und $\mathcal{M}$ die Menge aller anomaliefreien lokal simplizialen 2-Überlagerungskomplexe, die man durch wiederholte Nachbaridentifikationen aus $\mathcal{S}$ gewinnt.*

Es sei für jedes $\mathcal{T} \in \mathcal{M}$ eine Einbettung $\iota_{\mathcal{T}}$ in den $\mathbb{R}^d$ gegeben. Falls für jede Kante $x \in \mathcal{K}_1$ und jede Einbettung $\iota_{\mathcal{T}}$ die Länge von $\overline{\iota_{\mathcal{T}}(x)}$ nur von den Farbäquivalenzklassen ihrer Eckpunkte abhängt, dann ist diese Familie der Einbettungen konsistent bzgl. $\mathcal{M}$.

Beweis. Da 3-Färbbarkeit durch Nachbarschaftserweiterungen erhalten bleibt, folgt die Behauptung. $\square$

Wir haben in Bemerkung 2.63 und Lemma 2.64 allgemeine Bedingungen aufgestellt, unter denen eine möglichst große Klasse von Faltungen möglich ist. In Bemerkung 2.63 lassen wir alle möglichen Faltungen zu, während es in Lemma 2.64 nur Faltungen benachbarter Flächen sind. Diese Bedingungen sind aber weit davon entfernt, Einbettungen zu garantieren. Sie geben nur wieder, welche Einschränkungen allein aus der Forderung der (uneingeschränkten) Faltbarkeit folgen. Jede speziellere Untersuchung muss sich also mit den folgenden Fragen beschäftigen:

- Inwieweit werden Einschränkungen an die möglichen Faltungen gestellt? Da jede mögliche Faltung die Wahl der Einbettung einschränkt, lässt eine Einschränkung der möglichen Faltungen mehr Freiheit in der Wahl der Einbettungen.
- Gibt es zusätzliche Einschränkungen an die Einbettung (z. B. um Überschneidungen zu verhindern)?

3 Faltkomplexe

In Kapitel 2 sind wir dem Ziel, die Faltung von Papier abstrakt zu modellieren, ein wenig näher gekommen. Wir haben dort ein mengentheoretisches Modell entwickelt, das unabhängig von einer speziellen Einbettung der Faltung in den $\mathbb{R}^3$ ist und einige Eigenschaften von Papierfaltungen abbildet. Allerdings haben wir auch festgestellt, dass dieses Modell die Reihenfolge von Papier nicht abbilden kann, was insbesondere bei der Untersuchung der Entfaltung in Abschnitt 2.4 aufgefallen ist.

In diesem Kapitel werden wir diesen Mangel beheben, indem wir das Modell der n -Überlagerungskomplexe so erweitern, dass wir die Reihenfolge von Flächen ebenfalls modellieren können. Dabei stellen wir fest, dass die naive Idee, lineare Ordnungen auf jeder Äquivalenzklasse aus $\mathcal{K}_n / \sim$ zu definieren, fehlschlägt. Sie schlägt deswegen fehl, da sie nicht in der Lage ist, die Reihenfolgen bei einer Faltung richtig zu aktualisieren. Dementsprechend wird sich der gesamte Abschnitt 3.1 mit der Frage auseinandersetzen, wie man solche Reihenfolgen vernünftig definiert.

Nachdem die Frage nach der korrekten Beschreibung von Reihenfolgen in der Definition 3.27 abschließend geklärt ist, wenden wir uns der Untersuchung des Zusammenfaltens zu. Dabei werden wir zuerst untersuchen, wann man zwei Punkte, Kanten oder Flächen zusammenfalten kann, in Analogie zu den primitiven Erweiterungen aus Abschnitt 2.3.1. Nachdem wir diesen konstruktiven Begriff formuliert haben, definieren wir in Abschnitt 3.3 einen abstrakteren Faltungsbegriff, der – ähnlich wie das Konzept der Teilüberlagerungen aus Definition 2.19 – besser dazu geeignet ist, zu erkennen, ob zwei verschiedene Faltzustände durch Faltung auseinander hervorgehen.

Als nächstes untersuchen wir in Abschnitt 3.4, welchen Einfluss die Betrachtung von Reihenfolgen auf den Entfaltvorgang hat. Wir vergewissern uns dabei, dass die unsinnigen Resultate aus Abschnitt 2.4 nicht mehr auftreten. Es verbleibt aber noch immer eine gewisse Asymmetrie zwischen Faltung und Entfaltung.

Um eine klarere Beziehung zwischen Falt- und Entfaltvorgang herzustellen, entwickeln wir in Abschnitt 3.5 einen direkteren Formalismus zur Beschreibung von Faltung und Entfaltung. Dieser beruht darauf, dass es zur Faltung und Entfaltung von Flächen genügt, die beiden Oberflächen zu spezifizieren, die zusammengeführt bzw. getrennt werden sollen.

3.1 Definition des Faltzustandes

Wenn wir Papier falten, dann haben zusammengefaltete Flächen eine feste Reihenfolge. In diesem Abschnitt geht es um die Frage, um welche Informationen wir unser bestehendes Modell der n -Überlagerungskomplexe aus Kapitel 2 erweitern müssen, um mit solchen Reihenfolgen vernünftig umzugehen. Wir werden dazu in mehreren Schritten vorgehen.

Wir beginnen Abschnitt 3.1.1 damit, nachzuweisen, dass es nicht ausreicht, für jede Menge zusammengefalteter Flächen eine Reihenfolge anzugeben. Stattdessen werden sich die Reihenfolgen der Flächen um jede Kante (als zyklische Reihenfolge) als zentral herausstellen, sogar so zentral, dass wir die Reihenfolgen auf den Flächenklassen später aus den Kantenumläufen (und etwas mehr Informationen) rekonstruieren können.

Obwohl die Kantenumläufe dazu ausreichen, die Faltzustände vernünftig zu definieren, sind sie nicht ausreichend, um Faltungen eindeutig beschreiben zu können. Dazu müssen wir darüber sprechen können, welche Seiten der Flächen aufeinandergefaltet werden sollen. In Abschnitt 3.1.2 werden wir uns daher der Beschreibung dieser Oberflächen widmen.

Unter Verwendung der Oberflächen können wir verschiedene Kantenumläufe, die die selbe Fläche enthalten, vergleichen. Damit sind wir dazu in der Lage, die Interaktionen zwischen solchen Kantenumläufen zu beschreiben, was wir in Abschnitt 3.1.3 tun werden. Schließlich werden wir feststellen, dass wir zur vollständigen Rekonstruktion der Reihenfolgen auf den Flächenklassen noch eine Beschreibung der Randstücke brauchen, die wir in Abschnitt 3.1.4 liefern werden.

Diese Überlegungen kulminieren in Abschnitt 3.1.5, in dem wir die endgültige Definition 3.27 unserer Faltzustände präsentieren und auf einige vereinfachende Situationen eingehen.

3.1.1 Definition von Fächern

In diesem Abschnitt weisen wir nach, dass es zur Betrachtung von Reihenfolgen nicht ausreicht, nur die Reihenfolgen auf den zusammengefalteten Flächen zu betrachten. Formal lässt sich das so formulieren, dass die Festlegung einer lineare Ordnung auf jeder Äquivalenzklasse $[f] \in \mathcal{K}_n / \sim$ nicht dazu ausreicht, um die linearen Ordnungen eines gefalteten Zustands aus denen des ungefalteten Zustandes zu berechnen. Das folgende Beispiel 3.1 illustriert dieses Phänomen.

Beispiel 3.1. Wir betrachten den 1-Überlagerungskomplex¹⁰



mit $\prec := \mathcal{K}_0 \times \mathcal{K}_1$ und $\sim$ als Gleichheit.

Jede Kantenklasse besteht aus genau einem Element aus $\mathcal{K}_1$, wodurch die Angabe von linearen Ordnungen auf diesen trivial ist. Würden wir diesen aber wie in der Graphik dargestellt in den $\mathbb{R}^2$ einbetten, könnten wir die Kanten a und c durch eine isometrische Bewegung in der Ebene nicht zu direkten Nachbarn machen. Demnach ist in der Konfiguration der Graphik eine Information enthalten, die von den linearen Ordnungen auf den Kantenklassen nicht erfasst wird, aber dennoch zur Bestimmung gefalteter Zustände essentiell ist.

Während Beispiel 3.1 demonstriert, dass es nicht genügt, die Reihenfolge der zusammengefalteten Kanten festzulegen, suggeriert es auch eine Lösung für dieses Problem. Wenn wir nämlich zu dem zentralen Punkt den Umlauf der Kanten angeben (oder in drei Dimensionen den Umlauf der Flächen um eine Kante), haben wir die fehlende Information repräsentiert. Wir beschreiben einen solchen zyklischen Umlauf durch eine Nachfolgerabbildung ν . Da jedes Element sowohl einen Vorgänger als auch einen Nachfolger hat, muss ν bijektiv sein.

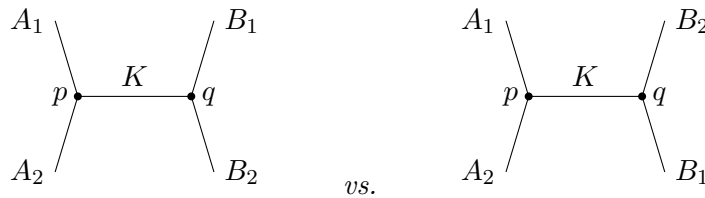
¹⁰Dieser ist zwar nicht lokal simplizial, aber das ändert die konkrete Argumentation nicht.

Definition 3.2 (zyklische Reihung). Sei M eine endliche Menge. Eine bijektive Abbildung $\nu : M \rightarrow M$ heißt **zyklische Reihung auf M** , falls die Gruppe $\langle \nu \rangle \leq \text{Sym}(M)$ transitiv auf M operiert. Die Menge aller zyklischen Reihungen auf M bezeichnen wir mit $\text{Zyk}(M)$.

Eine zyklische Reihung auf M kann also als ein Zykel der Länge $|M|$ in der symmetrischen Gruppe $\text{Sym}(M)$ dargestellt werden.

Wenn wir jedem Punkt eines 1-Überlagerungskomplexes einen eindeutigen Umlauf zuordnen wollen, müssen wir uns damit befassen, dass es immer zwei mögliche Umlaufrichtungen gibt. Dass es sich dabei tatsächlich um ein Problem handelt, zeigt das nächste Beispiel.

Beispiel 3.3. Seien p und q die Endpunkte einer Kante K . Der Kantenumlauf (in Zykelschreibweise) um p sei (K, A_1, A_2) , der um q sei (K, B_1, B_2) . Wenn wir keinen Umlaufsinn vorgeben, haben wir (bis auf Symmetrie) zwei Möglichkeiten, diese Situation in die Ebene einzubetten:



In der linken Darstellung können wir die Kante A_1 durch Falten in der Ebene unmittelbar zwischen K und B_1 bringen. Dies ist in der rechten Darstellung nicht möglich – entweder A_2 oder B_2 kommen dabei in die Quere.

Wir brauchen also eine Möglichkeit, verschiedene Umlaufrichtungen zu unterscheiden. In der Ebene ist das dadurch möglich, dass wir eine universelle Umlaufrichtung festlegen können (den mathematisch positiven Drehsinn bzw. gegen den Uhrzeigersinn). Im dreidimensionalen Raum liefert das Konzept der Orientierung die einfachste Lösung: Wir können den beiden Orientierungen einer Kante die beiden verschiedenen Flächenumläufe zuordnen (z. B. über die Rechte-Hand-Regel aus der Physik). Da wir dies ohne Bezug auf eine spezielle Einbettung in den $\mathbb{R}^3$ lösen wollen, müssen wir auf die simpliziale Definition von Orientierung zurückgreifen.

Definition 3.4 (simpliziale Orientierung). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec)$ ein n -Simplex. Ein Tupel $(a_1, \dots, a_{n+1}) \in \mathcal{K}_0^{n+1}$ heißt **simpliziale Orientierung (von $\mathcal{S}$)**, falls $a_i \neq a_j$ für $i \neq j$ gilt. Die Menge aller simplizialen Orientierungen von $\mathcal{S}$ bezeichnen wir mit $\text{SimplOrient}(\mathcal{S})$.

Gemäß dieser Definition hat ein Dreieck sechs verschiedene simpliziale Orientierungen. Um dies mit unserer Anschauung (die nur von zwei Orientierungen spricht) wieder in Einklang zu bringen, müssen wir erkennen, welche davon wir als gleich ansehen wollen. Es ist sinnvoll, die beiden möglichen Umlaufrichtungen der Punkte mit den Orientierungen des Dreiecks zu identifizieren. Daher sollen die simplizialen Orientierungen (a_1, a_2, a_3) , (a_2, a_3, a_1) und (a_3, a_1, a_2) als äquivalent angesehen werden. Wenn wir diese Äquivalenzklassen als Bahnen einer Gruppenoperation darstellen wollen, sehen wir, dass es sich um die Bahnen der alternierenden Gruppe $A_3 = \langle (1, 2, 3) \rangle$ handelt. Da die Operation der alternierenden Gruppe immer genau zwei Bahnen auf den simplizialen Orientierungen hat, verallgemeinern wir diese Erkenntnis in Bemerkung 3.5.

Bemerkung 3.5. Die symmetrische Gruppe S_{n+1} operiert transitiv auf der Menge der simplizialen Orientierungen eines n -Simplex vermöge

$$S_{n+1} \times \text{SimplOrient}(\mathcal{S}) \rightarrow \text{SimplOrient}(\mathcal{S}) : (\pi, (a_1, \dots, a_{n+1})) \mapsto (a_{\pi(1)}, \dots, a_{\pi(n+1)}).$$

Die Operation der alternierenden Gruppe A_{n+1} hat genau zwei Bahnen auf den simplizialen Orientierungen von $\mathcal{S}$. Diese Bahnen bezeichnen wir als **Orientierungsklassen**.

Wenn man landläufig von „Orientierung“ spricht, meint man eine der Bahnen von A_{n+1} . Für unsere Zwecke ist es allerdings praktischer, mit konkreten Bahnvertretern arbeiten zu können. Dafür müssen wir garantieren, dass alle unsere Definitionen sich unter der Operation von A_{n+1} nicht verändern. Zudem müssen sie konsistent mit der Operation von S_{n+1} sein (negatives Signum führt zu einer Orientierungsumkehr).

Mit dem Konzept der Orientierung in der Hand können wir Kantenumläufe definieren. Dazu müssen wir zuerst beschreiben, aus welchen Flächen ein solcher Umlauf bestehen kann. Es handelt sich dabei um alle Flächen, die an der gegebenen Kantenklasse anliegen.

Definition 3.6 (Corona). Sei $n > 0$ und $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein n -Überlagerungskomplex, sowie $k \in \mathcal{K}_{n-1}$. Dann heißt

$$\text{Cor}([k]) := \text{Cor}_{\mathcal{S}}([k]) := \{f \in \mathcal{K}_n : [k] \prec [f]\}$$

Corona (von $[k]$).

Gemäß der Idee des Kantenumlauf handelt es sich bei diesem um eine zyklische Reihung auf allen anliegenden Flächen, also auf der Corona der Kantenklasse. Dabei müssen wir darauf achten, dass ein Invertieren der Orientierung den Kantenumlauf umdreht, d. h. dass die zyklische Reihung invertiert wird, sobald wir die Orientierungsklasse aus Bemerkung 3.5 wechseln.

Da wir in der Definition des Zusammenfaltens verschiedene solcher Kantenumläufe kombinieren wollen, ist es praktischer, wenn ein Kantenumlauf nicht alle anliegenden Flächen durchlaufen muss.

Definition 3.7 (Fächer). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex. Seien $[k] \in \mathcal{K}_{n-1}/\sim$ und $M \subseteq \text{Cor}_{\mathcal{S}}([k])$ gegeben. Ein **Fächer** (von $[k]$ über M) ist eine Abbildung

$$\nu_{[k]} : \text{SimplOrient}(\mathcal{S}_{\preccurlyeq[k]}) \rightarrow \text{Zyk}(M),$$

sodass für $\pi \in S_{n+1}$ und $a \in \text{SimplOrient}(\mathcal{S}_{\preccurlyeq[k]})$ bereits $\nu_{[k]}(\pi \cdot a) = (\nu_{[k]}(a))^{\text{sign}(\pi)}$ gilt.

Falls $M = \text{Cor}_{\mathcal{S}}([k])$ gilt, bezeichnen wir die Abbildung als **vollen Fächer**.

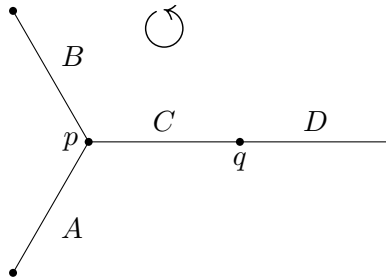
Falls die simpliziale Orientierung a aus dem Kontext klar ist, schreiben wir auch $\hat{\nu}_{[k]}$ für die zyklische Reihung $\nu_{[k]}(a)$.

Dabei beschreiben die vollen Fächer genau die Kantenumläufe, die wir zur Behandlung der Reihenfolge benötigen. Zwei volle Fächer, die eine Fläche gemeinsam haben, können aber nicht völlig unabhängig voneinander gewählt werden. Wir werden uns in Abschnitt 3.1.3 mit dieser Interaktion befassen, aber zunächst müssen wir in Abschnitt 3.1.2 die Mittel zur Behandlung dieser Interaktion bereitstellen.

3.1.2 Definition von Oberflächen

Wir haben im letzten Abschnitt die Kantenumläufe definiert, die wir zur Modellierung von Flächenreihenfolgen benötigen. Bevor wir uns in Abschnitt 3.1.3 mit deren Interaktionen beschäftigen, vervollständigen wir unser Modell an einer anderen Stelle. In Beispiel 3.8 sehen wir, dass wir Oberflächen benötigen, um Faltungen eindeutig festzulegen.

Beispiel 3.8. Wir betrachten den 1-Überlagerungskomplex, der durch folgende Einbettung in den $\mathbb{R}^2$ gegeben ist:

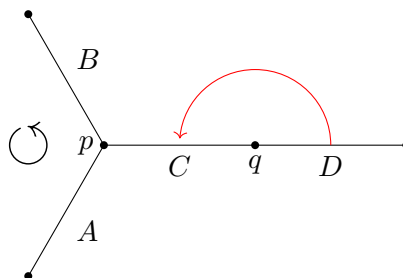


Wenn wir für die Richtung der Punkumläufe um p und q den mathematisch positiven Drehsinn wählen, lauten die zyklischen Reihungen

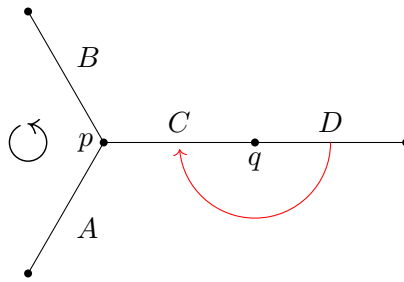
$$\nu_p((p)) = (A, C, B) \qquad \nu_q((q)) = (C, D).$$

Nehmen wir an, dass wir die Kanten C und D zusammenfalten wollen. Dadurch verändert sich die zyklische Reihung bei q nicht. Die zyklische Reihung um p wird hingegen um eine Kante (nämlich D) erweitert. Allerdings ist es nicht eindeutig, wie D in diese zyklische Reihung eingefügt wird.

Wir können die informelle Anweisung „falte C und D zusammen“ nämlich auf zwei verschiedene Weisen interpretieren. Wir können „obenherum“ falten



und die zyklische Reihung (A, C, D, B) an $[p]$ erhalten. Wir können aber auch „untenherum“ falten

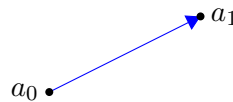


und an die zyklische Reihung (A, D, C, B) gelangen.

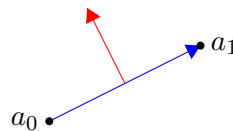
Das Beispiel 3.8 demonstriert, dass wir zur eindeutigen Spezifikation von Faltungen nicht nur sagen müssen, welche Flächen zusammengefaltet werden, sondern auch, wie deren Oberflächen kombiniert werden.

Um unabhängig von einer Einbettung in den $\mathbb{R}^n$ von einer „Oberfläche“ zu sprechen, hilft die Erkenntnis, dass die Begriffe von „Oberfläche“ und „Orientierung“ im $\mathbb{R}^n$ sehr eng miteinander verwoben sind. Bevor wir diese Einsicht allgemein formulieren, betrachten wir die Situation im $\mathbb{R}^2$ und $\mathbb{R}^3$ gesondert, da auf diesen unser Hauptaugenmerk liegt und weil diese einfacher zu visualisieren sind.

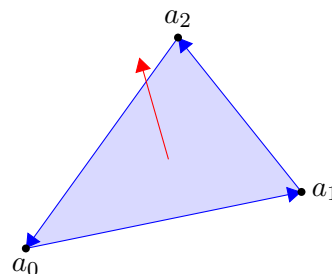
Beispiel 3.9. Im $\mathbb{R}^2$ können wir die simpliziale Orientierung (a_0, a_1) eines eingebetteten 1-Simplex als Richtung interpretieren:



Fasst man diese Richtung als Vektor $a_1 - a_0$ auf, so gibt es genau einen normierten Normalenvektor $\mathbf{n}$, sodass die Basis $(a_1 - a_0, \mathbf{n})$ positiv orientiert ist (d. h. $\det(a_1 - a_0, \mathbf{n}) > 0$). Man kann diesen Normalenvektor als Darstellung der Oberfläche (in diesem Fall die Seite der Kante) verstehen, von der er „wegzeigt“:



Im $\mathbb{R}^3$ können wir analog vorgehen. Zu einem nicht-entarteten Dreieck (a_0, a_1, a_2) gibt es genau einen normierten Vektor $\mathbf{n}$, der orthogonal auf der von (a_0, a_1, a_2) aufgespannten affinen Ebene steht und $\det(a_1 - a_0, a_2 - a_0, \mathbf{n}) > 0$ erfüllt.



In der Abbildung haben wir die Orientierung des Dreiecks angedeutet, die von der simplizialen Orientierung (a_0, a_1, a_2) festgelegt wird.

Dieses Beispiel lässt sich auf beliebige Dimensionen verallgemeinern.

Bemerkung 3.10. Sei $\mathcal{S}$ ein m -Simplex mit simplizialer Orientierung $(a_0, a_1, \dots, a_m)$ und einer simplizialen Einbettung¹¹ $\iota : \{a_0, \dots, a_m\} \rightarrow \mathbb{R}^{m+1}$, sodass die Vektoren $\iota(a_1) - \iota(a_0), \dots, \iota(a_m) - \iota(a_0)$ linear unabhängig sind.

Dann gibt es genau einen Vektor $\mathbf{n} \in \mathbb{R}^{m+1}$ mit den folgenden Eigenschaften:

1. $\mathbf{n}$ ist orthogonal zu $\iota(a_i) - \iota(a_0)$ für jedes $1 \leq i \leq m$.
2. $\mathbf{n}$ hat Länge 1.
3. $\det(\iota(a_1) - \iota(a_0), \dots, \iota(a_m) - \iota(a_0), \mathbf{n}) > 0$.

Falls $\bar{\mathbf{n}} \in \mathbb{R}^{m+1}$ dieser eindeutige Vektor zur simplizialen Orientierung $(a_{\pi(0)}, \dots, a_{\pi(m)})$ (mit $\pi \in \text{Sym}(\{0, 1, \dots, m\})$) ist, gilt

$$\bar{\mathbf{n}} = \text{sign}(\pi) \cdot \mathbf{n}.$$

Beweis. Der Übersichtlichkeit halber werden wir die Punkte a_i vermöge ι als Elemente des $\mathbb{R}^{m+1}$ auffassen.

Da die Vektoren $a_1 - a_0, \dots, a_m - a_0$ linear unabhängig sind, spannen sie einen Teilraum von Dimension m auf. Dessen Orthogonalraum bezüglich $\mathbb{R}^{m+1}$ ist daher eindimensional. Es gibt genau zwei Vektoren dieses Orthogonalraums von Länge 1, die sich durch ihr Vorzeichen unterscheiden. Setze $\mathbf{n}$ als denjenigen davon, der $\det(a_1 - a_0, \dots, a_m - a_0, \mathbf{n}) > 0$ erfüllt (ein Umdrehen des Vorzeichens von $\mathbf{n}$ ändert das Vorzeichen der Determinante). Damit haben wir gezeigt, dass es genau einen Vektor mit den geforderten Eigenschaften gibt.

Sei nun $\pi \in S_{m+1} \cong \text{Sym}(\{0, 1, \dots, m\})$. Dann können wir π in das Produkt $\tau \cdot \rho$ zerlegen, wobei τ die Transposition $(0, \pi(0))$ ist (bzw. die Identität, falls $\pi(0) = 0$ gilt) und $\rho(0) = 0$ gilt. Da sign ein Gruppenhomomorphismus ist und die S_{m+1} auf den simplizialen Orientierungen operiert, genügt es, die Aussage für τ und ρ separat zu zeigen.

Da ρ nur die Argumente der Determinante permutiert, folgt sofort (da die Determinante alternierend ist)

$$\det(a_{\rho(1)} - a_0, \dots, a_{\rho(m)} - a_0, \mathbf{n}) = \text{sign}(\rho) \det(a_1 - a_0, \dots, a_m - a_0, \mathbf{n}).$$

Um $\tau = (0, \pi(0))$ zu behandeln, verwenden wir die Multilinearität der Determinante. Wenn wir die Argumente von $\det(a_{\tau(1)} - a_{\tau(0)}, \dots, a_{\tau(m)} - a_{\tau(0)}, \mathbf{n})$ betrachten, steht an der $\pi(0)$ -ten Stelle $a_{\tau(\pi(0))} - a_{\tau(0)} = a_0 - a_{\tau(0)}$. Durch Subtraktion dieses Eintrages von den anderen ($\leq m$) ändert sich der Wert der Determinante nicht, solange wir die $\pi(0)$ -te Stelle unverändert lassen. Für jede Stelle $j \neq \pi(0)$ erhalten wir dann

$$(a_{\tau(j)} - a_{\tau(0)}) - (a_0 - a_{\tau(0)}) = a_{\tau(j)} - a_0 = a_j - a_0,$$

¹¹Gemäß Folgerung 2.54 handelt es sich also um eine injektive Abbildung.

da $\tau(j) = j$ gilt. Da wir noch $a_0 - a_{\tau(0)}$ an Stelle $\pi(0)$ mit -1 multiplizieren müssen, um auf die Form $\det(a_1 - a_0, \dots, a_m - a_0, \mathbf{n})$ zu kommen, unterscheiden sich diese beiden Determinanten um einen Faktor von -1 , dem Signum von τ .

Die Aussage für $\mathbf{n}$ folgt damit aus unserer Konstruktionsvorschrift und der Linearität der Determinante. $\square$

Bemerkung 3.10 stellt eine allgemeine Beziehung zwischen simplizialen Orientierungen und den „Oberflächen“ eines Simplex her. Wir verwenden sie daher als Legitimation dafür, dass wir die Oberflächen eines Simplex (in geeigneter Dimension) über simpliziale Orientierungen definieren, um den Begriff der Oberflächen in unserem abstrakten Modell zu verstehen.

Das Ziel dieser Abschweifung war es, Faltungen eindeutig über die zusammengefalteten Oberflächen festzulegen. Dazu muss bei jedem Dreieck des Musters eine eindeutige Oberfläche festgelegt sein. Formal bedeutet dies, dass wir bei jedem $f \in \mathcal{K}_n$ eine der Orientierungsklassen aus Bemerkung 3.5 auszeichnen müssen. Dazu geben wir eine Abbildung $\text{SimplOrient}(\mathcal{S}_{\preccurlyeq f}) \rightarrow \{\pm 1\}$ an, die für jedes Element der ausgezeichneten Orientierungsklasse $+1$ liefert und für jedes Element der anderen Orientierungsklasse -1 ergibt.

Definition 3.11 (simpliziale Oberfläche). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec)$ ein lokal simplizialer homogener n -Komplex. Eine **simpliziale Oberfläche von $\mathcal{S}$** ist eine Familie von Abbildungen

$$\sigma_f : \text{SimplOrient}(\mathcal{S}_{\preccurlyeq f}) \rightarrow \{\pm 1\}, \quad f \in \mathcal{K}_n,$$

sodass für $\pi \in S_{n+1}$ und $x \in \text{SimplOrient}(\mathcal{S}_{\preccurlyeq f})$ die Relation $\sigma(\pi \cdot x) = \text{sign}(\pi) \cdot \sigma_f(x)$ erfüllt ist.

Da diese Definition nur vom homogenen n -Komplex abhängig ist, der der Faltung zugrunde liegt, bleibt die Wahl der Oberflächen während des gesamten Faltvorgangs unberührt. Es handelt sich damit um eine vernünftige Grundlage, auf der wir Faltungen definieren können.

Diese Definition von σ_f baut direkt auf den simplizialen Orientierungen auf, wodurch sie in der Formulierung der Theorie hilfreich ist. Zur Angabe von σ_f ist diese Beschreibung aber aufwendiger, als sie sein muss (es gibt schließlich nur zwei mögliche Orientierungsklassen). Unter Verwendung der Gruppenoperation aus Bemerkung 3.5 können wir σ_f durch ihren Wert an einer einzigen simplizialen Orientierung festlegen.

Bemerkung 3.12. Da die S_{n+1} transitiv auf den simplizialen Orientierungen eines n -Simplex operiert, sind die Abbildungen σ_f aus Definition 3.11 bereits durch die Angabe einer einzigen simplizialen Orientierung a mit $\sigma_f(a) = 1$ vollständig festgelegt. Dies kann man auch so interpretieren, dass bei jedem $f \in \mathcal{K}_n$ eine Oberfläche (bzw. Orientierung) willkürlich ausgezeichnet wird.

Diese ausgezeichnete Oberfläche lässt sich graphisch durch einen Pfeil andeuten, der bei einer Kante auf der ausgezeichneten Seite steht:

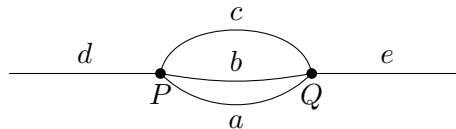


Wir bezeichnen die ausgezeichnete Oberfläche auch mit $(f, +1)$ und die andere mit $(f, -1)$.

3.1.3 Abhängigkeiten zwischen Fächern

Wir haben in Abschnitt 3.1.1 die Kantenumläufe definiert, die wir zur Beschreibung der Reihenfolgen von Faltzuständen benötigen. Da jedes Dreieck aber drei Kanten hat, kommt es in drei verschiedenen Kantenumläufen vor. In diesem Abschnitt beschäftigen wir uns mit dieser Abhängigkeit. Diese wird dann wichtig, wenn wir anfangen, Faltzustände abstrakt zu konstruieren.

Wir betrachten dazu die folgende Situation:



Hier lauten die zyklischen Reihungen der Punkte (bezüglich der Standardorientierung in der Ebene, also gegen den Uhrzeigersinn)

$$\nu_P((P)) = (a, b, c, d) \quad \text{und} \quad \nu_Q((Q)) = (e, c, b, a).$$

Um diese zyklischen Reihungen zu vergleichen, müssen wir uns auf den Schnitt ihrer Coronae (hier die Menge $\{a, b, c\}$) einschränken, da alle übrigen Kanten für den Vergleich irrelevant sind. Dazu betrachten wir z. B. den Punktumlauf um P und entfernen die Kante d . Da wir den Umlauf nicht verändert wollen, wird c dann auf a abgebildet. Wir nennen $\hat{\nu}_P((P)) = (a, b, c)$ dann das Redukt von $\nu_P((P))$.

Definition 3.13 (Redukt). Sei ν eine zyklische Reihung auf der endlichen Menge M und $T \subseteq M$ sei eine Teilmenge. Für jedes $t \in T$ setze

$$n_t := \min\{i \in \mathbb{N} \mid \nu^i(t) \in T\}.$$

Dann ist $\nu|_{T \rightarrow T} : T \rightarrow T, t \mapsto \nu^{n_t}(t)$ eine zyklische Reihung auf T , genannt das **Redukt von ν auf T** .

Wenn wir ν als (zyklische) Nachfolgerabbildung interpretieren, bedeutet die Reduktbildung, dass wir $t \in T$ auf den ersten Nachfolger abbilden, der ebenfalls in T liegt. Wir verdeutlichen dieses Vorgehen an einem Beispiel.

Beispiel 3.14. Sei $M := \{1, \dots, 6\}$ und

$$\nu : M \rightarrow M : x \mapsto \begin{cases} x+1 & x < 6 \\ 1 & x = 6 \end{cases}$$

eine zyklische Reihung auf M , die als 6-Zykel die Form $(1, 2, 3, 4, 5, 6)$ hat. Für

$$T := \{1, 2, 5\} \subseteq M$$

hat das Redukt von ν auf T die Zykeldarstellung $(1, 2, 5)$.

Wir haben das Redukt eingeführt, um die Einschränkung einer zyklischen Reihung beschreiben zu können. Da es häufig notwendig sein wird, diese Konstruktion zu iterieren, weisen wir nach, dass sich jede doppelte Reduktbildung auch durch eine einfache Reduktbildung erhalten lässt.

Lemma 3.15. *Sei ν eine zyklische Reihung auf der endlichen Menge M und $A \subseteq B \subseteq M$ seien Teilmengen. Dann gilt*

$$\nu|_{A \rightarrow A} = (\nu|_{B \rightarrow B})|_{A \rightarrow A}.$$

Beweis. Sei $a \in A$. Dann gilt $\nu|_{A \rightarrow A}(a) = \nu^{n_a}(a)$ mit

$$n_a := \min\{n \in \mathbb{N} | \nu^n(a) \in A\}.$$

Setze $\hat{\nu} := \nu|_{B \rightarrow B}$, dann gilt $\hat{\nu}|_{A \rightarrow A}(a) = \hat{\nu}^p(a)$ mit

$$p := \min\{n \in \mathbb{N} | \hat{\nu}^n(a) \in A\}.$$

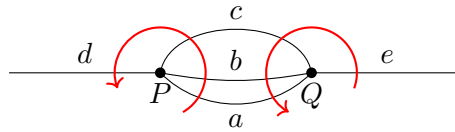
Es gibt $n_1, \dots, n_p \in \mathbb{N}$ mit

$$\begin{aligned} \hat{\nu}(a) &= \nu^{n_1}(a) =: a_1 \\ \hat{\nu}(a_1) &= \nu^{n_2}(a_1) =: a_2 \\ &\vdots \\ \hat{\nu}(a_{p-1}) &= \nu^{n_p}(a_{p-1}) = \hat{\nu}^p(a). \end{aligned}$$

Damit folgt $\nu^{n_1+n_2+\dots+n_p}(a) \in A$.

Gäbe es ein $k \leq n_1 + n_2 + \dots + n_p$ mit $\nu^k(a) \in A \subseteq B$, so gäbe es ein $1 \leq i \leq p-1$, das $a_i \in A$ erfüllt. Dann wäre aber $\hat{\nu}^i(a) \in A$, im Widerspruch zur Minimalität von p . Folglich gilt $n_a = n_1 + n_2 + \dots + n_p$. $\square$

Wir kehren jetzt zu der Situation vom Anfang dieses Abschnitts zurück. Die Redukte von $\nu_P((P))$ und $\nu_Q((Q))$ auf $\{a, b, c\}$ sind dann (a, b, c) und (c, b, a) , also invers zueinander. Wir werden sehen, dass das in der Ebene immer der Fall sein muss. Dazu betrachten wir, welche Umlaufrichtungen von den beiden Punkten auf $\{a, b, c\}$ induziert werden.



Wir erkennen, dass die Menge $\{a, b, c\}$ in zwei verschiedenen Richtungen durchlaufen wird. In der Ebene liegt diese Situation aufgrund der universellen Orientierung immer vor. Damit sind zwei zyklische Reihungen μ, ν über den Mengen M und N in der Ebene genau dann kompatibel, wenn $M \cap N$ leer ist oder wenn

$$\mu|_{M \cap N \rightarrow M \cap N} = (\nu|_{M \cap N \rightarrow M \cap N})^{-1}$$

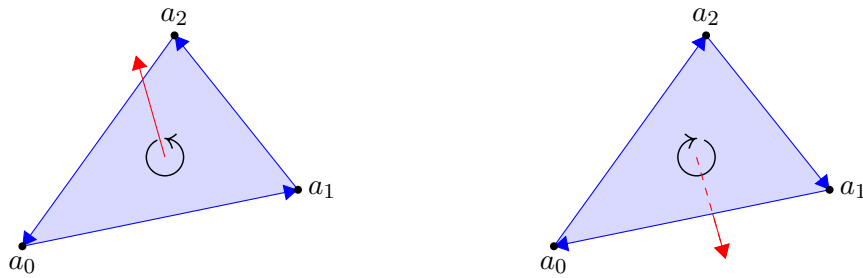
gilt. Um dieses Resultat zu erhalten, genügt es, eine Orientierung der Ebene festzulegen. Das erkennen wir auch daran, dass wir in der obigen Gleichung die zyklischen Reihungen μ und ν durch ihre Inversen ersetzen können (also die Orientierung der Ebene wechseln), wie wir in Bemerkung 3.16 festhalten:

Bemerkung 3.16. *Invertieren ist mit dem Bilden des Redukts verträglich. Genauer gilt:*
Sei $\nu : M \rightarrow M$ eine zyklische Reihung und $T \subseteq M$. Dann gilt

$$(\nu^{-1})|_{T \rightarrow T} = (\nu|_{T \rightarrow T})^{-1}.$$

Wir haben damit die Abhängigkeiten zwischen zwei zyklischen Reihungen in der Ebene, also von Fächern in 1-Überlagerungskomplexen, vollständig beschrieben.

Die Situation verkompliziert sich in drei Dimensionen, da wir die Umlaufrichtungen der Fächer nicht mehr kanonisch wählen können. Wir müssen also mit den simplizialen Orientierungen der Kanten arbeiten. Dabei legt jede solche simpliziale Orientierung eine Oberfläche der anliegenden Dreiecke fest (mit der Konvention aus Beispiel 3.9, die einem orientierten Simplex eine eindeutige Oberfläche zuordnet):



Dieser induzierte Durchlaufsinne lässt sich sehr einfach bestimmen: Ist z. B. (a_1, a_2) die Orientierung der Kante, so ist (a_1, a_2, a_0) die Orientierung des Dreiecks, die der von (a_1, a_2) ausgewählten Oberfläche entspricht.

Diese Zuordnung müssen wir aber noch ein wenig verallgemeinern, da wir in einem Fächer von der simplizialen Orientierung einer Kantenklasse sprechen und nicht von der einer einzelnen Kante. Da die Orientierung einer Kantenklasse aber die Orientierung aller Kanten dieser Klasse festlegt, ist das ebenfalls eindeutig. Wir halten diese Überlegungen in Definition 3.17 fest.

Definition 3.17 (induzierter Durchlaufsinne). *Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex mit simplizialer Oberfläche $\{\sigma_f\}_{f \in \mathcal{K}_n}$. Sei $[k] \in \mathcal{K}_{n-1} / \sim$ gegeben. Dann ist für jedes $f \in \text{Cor}([k])$ der **induzierte Durchlaufsinne** die Abbildung*

$$\begin{aligned} \bar{\sigma}_f : \text{SimplOrient}(\mathcal{S}_{\prec[k]}) &\rightarrow \{\pm 1\} \\ ([a_1], \dots, [a_n]) &\mapsto \sigma_f((b_1, \dots, b_n, b_{n+1})), \end{aligned}$$

mit $b_i \in [a_i]$ für alle $1 \leq i \leq n$ und $(b_1, \dots, b_n, b_{n+1}) \in \text{SimplOrient}(\mathcal{S}_{\prec f})$.

Wohldefiniertheit. Wir müssen nachweisen, dass die Definition des induzierten Durchlaufsinns eindeutig ist. Da f in der Corona von $[k]$ liegt, gibt es genau ein $b_{n+1} \in \mathcal{K}_0$ mit $b_{n+1} \prec f$, sodass $[b_{n+1}] \not\prec [k]$ gilt. Für jeden anderen Punkt $b_i \in \mathcal{K}_0$ mit $b_i \prec f$ gibt es eine Punktklasse $[a_i] \in \mathcal{K}_0 / \sim$ mit $[a_i] \prec [k]$, sodass $b_i \in [a_i]$ gilt.

Lägen in einer Punktklasse $[a_i]$ zwei verschiedene Punkte $b, c \in \mathcal{K}_0$, die in f enthalten sind, dann wäre die Irreduzibilität von $\sim$ verletzt. Folglich ist die Konstruktion aus der Definition eindeutig. $\square$

Damit haben wir den Durchlaufsinns einer Fläche, der von den anliegenden Kanten festgelegt wird, definiert. Kommen wir zum ursprünglichen Problem der Kompatibilität von Fächern zurück. Für zwei zyklische Reihungen müssen wir überprüfen, ob ihre simplizialen Orientierungen den selben Durchlaufsinns induzieren. Falls sie das tun, müssen sie gleich sein, ansonsten invers zueinander. In Definition 3.18 beschreiben wir diese Situation allgemein.

Definition 3.18 (kompatible Fächer). *Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex mit simplizialer Oberfläche $\{\sigma_f\}_{f \in \mathcal{K}_n}$. Seien $\nu_{[k]}, \nu_{[l]}$ Fächer über den Mengen M_k und M_l . Die Fächer heißen **kompatibel**, falls eine der beiden folgenden Bedingungen erfüllt ist.*

- $M_k \cap M_l = \emptyset$.
- *Es gibt ein $f \in M_k \cap M_l$, sowie simpliziale Orientierungen $a \in \text{SimplOrient}(\mathcal{S}_{\prec[k]})$ und $b \in \text{SimplOrient}(\mathcal{S}_{\prec[l]})$, sodass*

$$(\nu_{[k]}(a)|_{M_k \cap M_l \rightarrow M_k \cap M_l})^{\bar{\sigma}_f(a)} = (\nu_{[l]}(b)|_{M_k \cap M_l \rightarrow M_k \cap M_l})^{\bar{\sigma}_f(b)}$$

gilt, wobei $\bar{\sigma}_f$ den induzierten Durchlaufsinns von $[k]$ aus Definition 3.17 bezeichnet.

Wohldefiniertheit. Wir zeigen zuerst, dass die Definition unabhängig von der Wahl der simplizialen Orientierung von $\mathcal{S}_{\prec[k]}$ ist.

Seien $a, c \in \mathcal{S}_{\prec[k]}$. Dann gibt es eine Permutation $\pi \in S_n$ mit $\pi \cdot a = c$. Einsetzen in die Definition des induzierten Durchlaufsinns liefert $\bar{\sigma}_f(\pi \cdot a) = \text{sign}(\pi) \cdot \bar{\sigma}_f(a)$. Zusammen mit der Bemerkung 3.16, die die Verträglichkeit von Redukt und Invertieren sicherstellt, erhalten wir die folgende Gleichungskette:

$$\begin{aligned} (\nu_{[k]}(c)|_{M_k \cap M_l \rightarrow M_k \cap M_l})^{\bar{\sigma}_f(c)} &= (\nu_{[k]}(\pi \cdot a)|_{M_k \cap M_l \rightarrow M_k \cap M_l})^{\bar{\sigma}_f(\pi \cdot a)} \\ &= \left(\left(\nu_{[k]}(a)^{\text{sign}(\pi)} \right)_{|M_k \cap M_l \rightarrow M_k \cap M_l} \right)^{\text{sign}(\pi) \bar{\sigma}_f(a)} \\ &= (\nu_{[k]}(a)|_{M_k \cap M_l \rightarrow M_k \cap M_l})^{\text{sign}(\pi) \cdot \text{sign}(\pi) \cdot \bar{\sigma}_f(a)} \\ &= (\nu_{[k]}(a)|_{M_k \cap M_l \rightarrow M_k \cap M_l})^{\bar{\sigma}_f(a)}. \end{aligned}$$

Als nächstes zeigen wir die Unabhängigkeit von der Wahl von f . Da nur die Exponenten der Gleichung von f beeinflusst werden und nur Werte in $\{\pm 1\}$ annehmen, genügt es zu zeigen, dass $\bar{\sigma}_f(a) \cdot \bar{\sigma}_f(b)$ konstant ist.

Sei $g \in M_k \cap M_l$. Da $\mathcal{S}$ lokal simplizial ist, gibt es gemäß Lemma 2.25 eine Identifikation $\gamma : \mathcal{S}_{\preccurlyeq f} \rightarrow \mathcal{S}_{\preccurlyeq g}$, die konstant auf dem Schnitt ist. Da die S_{n+1} transitiv auf den simplizialen Orientierungen operiert, ist der Term $\sigma_f(x) \cdot \sigma_g(\gamma(x))$ für alle $x \in \text{SimplOrient}(\mathcal{S}_{\preccurlyeq f})$ konstant. Wir bezeichnen seinen Wert mit $\omega \in \{\pm 1\}$.

Betrachten wir nun $\bar{a} := ([a_1], \dots, [a_n]) \in \text{SimplOrient}(\mathcal{S}_{\preccurlyeq[k]})$. Dann gilt

$$\bar{\sigma}_f(\bar{a}) \cdot \bar{\sigma}_g(\bar{a}) = \sigma_f((u_1, \dots, u_n, u_{n+1})) \cdot \sigma_g((v_1, \dots, v_n, v_{n+1}))$$

mit $(u_1, \dots, u_n, u_{n+1}) \in \text{SimplOrient}(\mathcal{S}_{\preccurlyeq f})$ und $(v_1, \dots, v_n, v_{n+1}) \in \text{SimplOrient}(\mathcal{S}_{\preccurlyeq g})$, sowie $u_i, v_i \in [a_i]$ für alle $1 \leq i \leq n$.

Insbesondere gilt $u_i \sim v_i$ für alle $1 \leq i \leq n$. Folglich ist $u_i \mapsto v_i$ (für alle $1 \leq i \leq n+1$) die einzige Möglichkeit für eine Identifikation $\mathcal{S}_{\preccurlyeq f} \rightarrow \mathcal{S}_{\preccurlyeq g}$, die konstant auf dem Schnitt ist. Damit muss es sich um γ handeln. Wir schließen $\bar{\sigma}_f(\bar{a}) \cdot \bar{\sigma}_g(\bar{a}) = \omega$.

Insgesamt haben wir

$$\begin{aligned} \bar{\sigma}_g(a) \cdot \bar{\sigma}_g(b) &= \bar{\sigma}_g(a) \cdot \omega^2 \cdot \bar{\sigma}_g(b) \\ &= \bar{\sigma}_g(a) \cdot \left(\bar{\sigma}_g(a) \cdot \bar{\sigma}_f(a) \right) \cdot \left(\bar{\sigma}_f(b) \cdot \bar{\sigma}_g(b) \right) \cdot \bar{\sigma}_g(b) \\ &= \bar{\sigma}_f(a) \cdot \bar{\sigma}_f(b), \end{aligned}$$

womit die Behauptung gezeigt ist. □

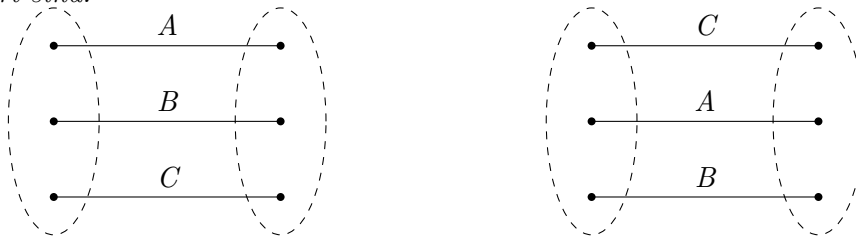
Damit haben wir beschrieben, wie die verschiedenen Kantenumläufe im Formalismus der Fächer miteinander zusammenhängen. Im nächsten Abschnitt werden wir uns mit einer Limitierung der Fächer befassen.

3.1.4 Definition von Randstücken

Während wir uns in den letzten Abschnitten mit den Eigenschaften von Fächern beschäftigt haben, untersuchen wir in diesem Abschnitt, ob Fächer zur Beschreibung von Falzuständen genügen. Um das zu überprüfen, versuchen wir die Reihenfolge zusammengefalteter Flächen aus ihren Fächern zu rekonstruieren.

Leider ist dies nicht im Allgemeinen möglich. Beispiel 3.19 demonstriert eine Situation, in der zwei verschiedene Reihenfolgen den selben Fächer zugeordnet bekommen.

Beispiel 3.19. *Wir betrachten die 1-Überlagerungskomplexe, die in den folgenden Bildern illustriert sind:*

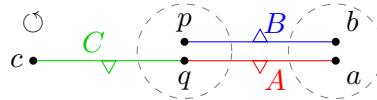


Die gestrichelten Ellipsen deuten an, dass alle Punkte innerhalb einer Ellipse äquivalent sind. Wir wollen annehmen, dass auch die Kanten $\{A, B, C\}$ eine Äquivalenzklasse bilden.

Damit lauten die zyklischen Reihungen der linken Ellipsen (C, B, A) , während die zyklischen Reihungen der rechten Ellipsen die Form (A, B, C) haben. Folglich ist es nicht möglich, die beiden Zustände anhand ihrer Fächer zu unterscheiden.

Um die Situation aus Beispiel 3.19 auszuschließen, nehmen wir die beiden Randstücke jeder Flächenklasse zu unserer Modellierung hinzu. In den meisten Fällen lassen sich diese Randstücke aus den Fächern rekonstruieren, wie z. B. in der Situation aus Beispiel 3.20.

Beispiel 3.20. Betrachten wir den lokal simplizialen 1-Überlagerungskomplex



mit $A \sim B$. Die zyklische Reihung bei $[p]$ ist $\hat{\nu}_{[p]} = (A, B, C)$. Die Kantenorientierungen für die simpliziale Oberfläche (vgl. Bemerkung 3.12) können wir aus den Richtungen der Pfeile in der Graphik und der Konvention aus Lemma 3.10 ablesen:

$$A : (a, q)$$

$$B : (p, b)$$

$$C : (q, c)$$

Um zu erkennen, dass die Oberfläche $(B, +1)$ ein Randstück ist (also die Seite, die durch ein Dreieck markiert ist), genügt die Beobachtung, dass $\hat{\nu}_{[p]}(B) = C \not\sim B$. Um zu erkennen, dass $(A, +1)$ ein Randstück ist, genügt die Beobachtung, dass $\hat{\nu}_{[p]}^{-1}(A) = C \not\sim A$.

Wieso müssen wir einmal $\hat{\nu}_{[p]}$ und einmal $\hat{\nu}_{[p]}^{-1}$ betrachten? Der einzige Unterschied ist die festgelegte Orientierung: $\sigma_B((p, b)) = +1$, aber $\sigma_A((q, a)) = -1$. Daher muss die Orientierung im Exponenten der zyklischen Reihung vorkommen.

Allerdings ist $(B, -1)$ kein Randstück, denn $\hat{\nu}_{[p]}^{-1}(B) = A \sim B$. Da sich $(B, -1)$ nur in der Seitenangabe von $(B, +1)$ unterscheidet, muss auch diese in der Lage sein, $\hat{\nu}_{[p]}$ zu invertieren.

Folglich ist der Exponent von $\hat{\nu}_{[p]}$ das Produkt der Seitenangabe (aus $\{+1, -1\}$) und dem Wert der simplizialen Oberflächenabbildung für diese Kante (wobei die Orientierung der Kante durch den Punkt $[p]$ festgelegt ist).

Wir stoßen damit auf die folgende Charakterisierung: Eine orientierte Seite (X, ε_X) ist genau dann ein Randstück, wenn $\hat{\nu}_{[p]}^{\varepsilon_X \cdot \sigma_X(\tilde{p}, y)}(X) \not\sim X$ gilt. Dabei ist $\tilde{p} \in [p]$ der Vertreter, der $\tilde{p} \prec X$ erfüllt und y ist der andere Eckpunkt dieser Kante.

Wir haben in Beispiel 3.20 gesehen, dass es möglich ist, Randstücke in der Ebene zu erkennen, wenn verschiedene Kantenklassen in einem Fächer vorkommen. Diese Bedingungen lassen sich leicht vom zweidimensionalen Fall auf höhere Dimensionen übertragen. Dabei ist lediglich zu beachten, dass die Kantenorientierungen explizit einbezogen werden müssen, da wir keine globale Umlaufrichtung mehr festlegen können.

Bevor wir diese Verallgemeinerung aber in Definition 3.22 festhalten, sollten wir die Grundlage des Verfahrens aus Beispiel 3.20, mit dem wir die Randstücke erkannt haben, allgemein formulieren. Dabei starten wir mit einer Oberfläche (f, ε_f) im Fächer. Diese legt eine Umlaufrichtung des Fächers fest. Wir erhalten dann die nächste Fläche in diesem Kantenumlauf. Formal beschreiben wir dieses Konzept als ε_f -Nachbarn.

Definition 3.21 (ε_f -Nachbar). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex mit simplizialer Oberfläche $\{\sigma_f\}_{f \in \mathcal{K}_n}$. Seien $f \in \mathcal{K}_n$ und $k \in \mathcal{K}_{n-1}$ mit $k \prec f$, sowie ein Fächer $\nu_{[k]}$ gegeben.

Sei a eine simpliziale Orientierung von $\mathcal{S}_{\leq [k]}$, dann heißt $(\nu_{[k]}(a))^{\varepsilon_f \cdot \bar{\sigma}_f(a)}(f)$ der ε_f -**Nachbar von f** (bzgl. $\nu_{[k]}$). Dabei bezeichnet $\bar{\sigma}_f$ den induzierten Durchlaufsinne von $[k]$ aus Definition 3.17.

Wohldefiniertheit. Wir müssen nachweisen, dass diese Definition unabhängig von der Wahl der simplizialen Orientierung a ist. Sei dazu $\pi \in S_n$. Dann gilt

$$\begin{aligned} (\nu_{[k]}(\pi \cdot a))^{\varepsilon_f \cdot \bar{\sigma}_f(\pi \cdot a)}(f) &= (\nu_{[k]}(a))^{\text{sign}(\pi) \cdot \varepsilon_f \cdot \text{sign}(\pi) \cdot \bar{\sigma}_f(a)}(f) \\ &= (\nu_{[k]}(a))^{\varepsilon_f \cdot \bar{\sigma}_f(a)}(f). \end{aligned}$$

Folglich ist der ε_f -Nachbar von f unabhängig von der Wahl der simplizialen Orientierung. $\square$

Der Begriff des ε -Nachbarn ermöglicht es uns, das Verfahren aus Beispiel 3.20 zur Erkennung von Randstücken wie folgt zu formulieren: Eine orientierte Seite (X, ε_X) ist ein Randstück, falls der ε_X -Nachbar von X nicht zu X äquivalent ist.

In Beispiel 3.19 gab es aber Randstücke, die diese Bedingung nicht erfüllen, da alle Kanten zueinander äquivalent sind. Um die Randstücke, die sich wie in Beispiel 3.20 zu erkennen geben, von allgemeinen Randstücken zu unterscheiden, nennen wir sie Standardränder.

Definition 3.22 (Standardrand). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex mit simplizialer Oberfläche. Sei $(f, \varepsilon_f) \in \mathcal{K}_n \times \{\pm 1\}$ und sei ein Fächer $\nu_{[k]}$ mit $k \in \mathcal{K}_{n-1}$ gegeben, sodass $[k] \prec [f]$ gilt.

Dann heißt (f, ε_f) **Standardrand** (bzgl. $\nu_{[k]}$), falls der ε_f -Nachbar von f nicht zu f äquivalent ist.

Um zu rekapitulieren: Wir haben festgestellt, dass es Situationen wie Beispiel 3.19 gibt, in denen wir allein aus den Fächern nicht rekonstruieren können, wie zusammengefaltete Flächen angeordnet sind. Wir haben erwähnt, dass es möglich ist, diese Unterscheidung dadurch zu treffen, dass wir die Randstücke jeder Flächenklasse in unser Modell aufnehmen. Bis zu diesem Punkt haben wir dann nachgewiesen, dass es sich dabei in den Fällen, in denen wir die linearen Ordnungen rekonstruieren können, um redundante Information handelt.

Als nächstes wollen wir eine Charakterisierung von Randstücken finden, die auch in Beispiel 3.19 anwendbar ist, also wenn der gesamte Faltzustand nur aus einer einzigen Flächenklasse besteht. Eine zentrale Eigenschaft ist, dass sich Randstücke „gegenüber liegen“. Dieses Konzept lässt sich mit diesem Begriff der Nachbarschaft leicht verstehen: Zwei Oberflächen (f, ε_f) und (g, ε_g) liegen sich gegenüber, wenn g der ε_f -Nachbar von f ist und f der ε_g -Nachbar von g ist (natürlich bezüglich desselben Fächers).

Die obige Beschreibung hat aber ein Problem, falls der Fächer nur aus zwei Flächen besteht – dann erfüllt jedes Paar von Oberflächen diese Bedingung, obwohl es eigentlich nur zwei solcher Paare geben sollte. Das liegt daran, dass ein Zykel mindestens drei Elemente enthalten muss, um allein durch seine Abbildungsvorschrift zwischen verschiedenen Orientierungen unterscheiden zu können. Wir können dieses Manko dadurch beheben, dass wir zusätzlich fordern,

dass die Oberflächen entgegengesetzt orientiert sind. Damit meinen wir, dass die induzierten Umlaufrichtungen der beiden Oberflächen verschieden sind.

Definition 3.23 (entgegengesetzt orientiert). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex mit simplizialer Oberfläche $\{\sigma_f\}_{f \in \mathcal{K}_n}$.

Sei $[k] \in \mathcal{K}_{n-1} / \sim$. Die Oberflächen $(f, \varepsilon_f), (g, \varepsilon_g) \in \text{Cor}([k]) \times \{\pm 1\}$ heißen **entgegengesetzt orientiert (bzgl. $[k]$)**, falls

$$\varepsilon_f \cdot \bar{\sigma}_f(a) \neq \varepsilon_g \cdot \bar{\sigma}_g(a)$$

für eine (und damit alle) simplizialen Orientierungen $a \in \text{SimplOrient}(\mathcal{S}_{\preceq[k]})$ gilt, wobei $\bar{\sigma}_f$ und $\bar{\sigma}_g$ die induzierten Durchlaufsinne von $[k]$ aus Definition 3.17 sind.

Nachdem wir festgelegt haben, was es bedeutet, dass zwei Oberflächen verschiedene Umlaufsinne definieren, können wir jetzt das Konzept komplementärer Oberflächen definieren.

Definition 3.24 (komplementäre Oberflächen). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex mit simplizialer Oberfläche $\{\sigma_f\}_{f \in \mathcal{K}_n}$.

Seien $(f, \varepsilon_f), (g, \varepsilon_g) \in \mathcal{K}_n \times \{\pm 1\}$ und sei ein Fächer $\nu_{[k]}$ mit $k \in \mathcal{K}_{n-1}$ gegeben, sodass $[k] \prec [f]$ und $[k] \prec [g]$ erfüllt sind.

Dann heißen (f, ε_f) und (g, ε_g) **komplementär (bzgl. $\nu_{[k]}$)**, falls gilt:

- g ist der ε_f -Nachbar von f bzgl. $\nu_{[k]}$
- f ist der ε_g -Nachbar von g bzgl. $\nu_{[k]}$
- (f, ε_f) und (g, ε_g) sind entgegengesetzt orientiert

Mit Definition 3.24 haben wir beschrieben, wann sich zwei Oberflächen im Fächer „gegenüberliegen“. Damit haben wir eine charakteristische Eigenschaft von Randstücken formuliert. Wir haben die Bedingung der entgegengesetzten Orientierung aber nur in die Definition aufgenommen, um den Spezialfall abzudecken, dass der Fächer genau zwei Flächen enthält. In Bemerkung 3.25 weisen wir nach, dass diese Bedingung in allen anderen Fällen tatsächlich trivial ist.

Bemerkung 3.25. Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex mit simplizialer Oberfläche.

Seien $(f, \varepsilon_f), (g, \varepsilon_g) \in \mathcal{K}_n \times \{\pm 1\}$ und sei ein Fächer $\nu_{[k]}$ mit $k \in \mathcal{K}_{n-1}$ gegeben, sodass $[k] \prec [f]$ und $[k] \prec [g]$ erfüllt sind. Die Corona $\text{Cor}_{\mathcal{F}}([k])$ enthalte mindestens drei Elemente. Weiter gelten:

- g ist der ε_f -Nachbar von f bzgl. $\nu_{[k]}$
- f ist der ε_g -Nachbar von g bzgl. $\nu_{[k]}$

Dann sind (f, ε_f) und (g, ε_g) komplementär (bzgl. $\nu_{[k]}$).

Beweis. Wir verwenden die Notation aus Definition 3.24. Dann ist $\hat{\nu}_{[k]}$ eine zyklische Reihung aus mindestens drei Elementen. Angenommen, $\varepsilon_f \cdot \hat{\sigma}_f = \varepsilon_g \cdot \hat{\sigma}_g$ gilt, also ohne Beschränkung der Allgemeinheit $\hat{\nu}_{[k]}(\textcolor{red}{f}) = \textcolor{blue}{g}$ und $\hat{\nu}_{[k]}(\textcolor{blue}{g}) = \textcolor{red}{f}$. Dann folgt

$$\hat{\nu}_{[k]}^2(\textcolor{red}{f}) = \hat{\nu}_{[k]}(\textcolor{blue}{g}) = \textcolor{red}{f},$$

im Widerspruch dazu, dass $\hat{\nu}$ mindestens Ordnung 3 hat. $\square$

Insgesamt haben wir es damit geschafft, die Randstücke einer Flächenklasse zu charakterisieren. Wir haben uns mit Randstücken auseinandergesetzt, weil wir die Reihenfolgen der zusammengefalteten Flächen nur auf Basis der Fächer nicht in allen Situationen bestimmen konnten. Nachdem wir Randstücke jetzt hinreichend lange beschrieben haben, zeigt Lemma 3.26, dass dieser Ansatz erfolgreich ist.

Lemma 3.26. *Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex mit simplizialer Oberfläche. Sei $\nu_{[k]}$ ein Fächer (für ein $k \in \mathcal{K}_{n-1}$) und sei eine Menge $R \subseteq \text{Cor}([k]) \times \{\pm 1\}$ gegeben, sodass gilt:*

- *Zu jeder Klasse $[f] \in \mathcal{K}_n / \sim$ in $\text{Cor}([k])$ gibt es genau zwei Elemente $(g, \varepsilon_g) \in R$, die $g \in [f]$ erfüllen. (Es gibt genau zwei Randstücke pro Flächenklasse.)*
- *Jeweils zwei Elemente aus R sind komplementär bezüglich $\nu_{[k]}$. (Randstücke sind komplementär.)*
- *Jeder Standardrand bezüglich $\nu_{[k]}$ liegt in R . (Standardränder sind Randstücke.)*

Sei ein $f \in \text{Cor}([k])$ gegeben, dann gilt $[f] = \{f_1, \dots, f_m\}$ und erfüllt folgende Eigenschaften:

1. *$\hat{\nu}(f_i) = f_{i+1}$ für jedes $1 \leq i < m$, wobei $\hat{\nu} := \nu_{[k]}(a)$ für eine simpliziale Orientierung $a \in \text{SimplOrient}(\mathcal{S}_{\prec[k]})$.*
2. *Falls $m = 1$, liegen $(f_1, +1)$ und $(f_1, -1)$ in R .*
3. *Falls $m > 1$, gibt es $\varepsilon_1, \varepsilon_m \in \{\pm 1\}$, sodass $([f] \times \{\pm 1\}) \cap R = \{(f_1, \varepsilon_1), (f_m, \varepsilon_m)\}$ gilt.*

Beweis. Wir unterscheiden zwei Fälle. Falls $\text{Cor}([k]) = [f]$ gilt (für ein $f \in \mathcal{K}_n$), gibt es genau zwei Elemente in R , die wir mit (g, ε_g) und (h, ε_h) bezeichnen. Nach Voraussetzung sind sie komplementär, d. h. $\hat{\nu}(g) = h$ oder $\hat{\nu}(h) = g$ gilt. Ohne Einschränkung gelte die letztere Aussage. Da $\hat{\nu}$ eine zyklische Reihung auf $[f]$, ist, durchläuft

$$g, \hat{\nu}(g), \hat{\nu}^2(g), \dots, \hat{\nu}^{m-1}(g)$$

alle Elemente von $[f]$. Da $\hat{\nu}^m$ die Identität auf $[f]$ ist, folgt $\hat{\nu}^{m-1}(g) = h$.

Im anderen Fall besteht die Corona aus mehreren Äquivalenzklassen. In diesem Fall kann R nur aus Standardrändern bestehen: Es muss nämlich minimale $m_+, m_- \in \mathbb{N}$ geben, die $\hat{\nu}^{m_+}(f) \not\sim f$ und $(\hat{\nu}^{-1})^{m_-}(f) \not\sim f$ erfüllen. Dann liegen bei $\hat{\nu}^{m_+-1}(f)$ und $(\hat{\nu}^{-1})^{m_--1}(f)$ Standardränder vor (falls $[f]$ einelementig ist, haben diese verschiedene zweite Komponenten).

Da dieses Argument für jedes $g \in [f]$ funktioniert und es nur zwei Randstücke mit erster Komponente in $[f]$ gibt, muss jedes Element aus $[f]$ eine der folgenden Formen haben:

$$(\hat{\nu}^{-1})^{m-1}(f), \dots, \hat{\nu}^{-1}(f), f, \hat{\nu}(f), \dots, \hat{\nu}^{m+1}(f) \quad \square$$

In Lemma 3.26 haben wir $[f] = \{f_1, \dots, f_m\}$ mit $\hat{\nu}(f_i) = f_{i+1}$ erhalten. Wir interpretieren diese Situation als die lineare Ordnung $f_1 < \dots < f_m$ auf der Menge $\{f_1, \dots, f_m\}$.

3.1.5 Vollständige Definition

Nachdem wir in den vorangehenden Abschnitten sämtliche Aspekte untersucht und eingeführt haben, die wir zur abstrakten Beschreibung eines Faltzustandes benötigen, sind wir jetzt in der Lage, diesen zu definieren.

Definition 3.27 (Faltkomplex). *Sei $n > 0$. Ein n -Proto-Faltkomplex ist ein lokal simplizialer n -Überlagerungskomplex $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ mit simplizialer Oberfläche $\{\sigma_f\}_{f \in \mathcal{K}_n}$, so dass gilt*

1. *Für jedes $[k] \in \mathcal{K}_{n-1}/\sim$ ist ein voller Fächer $\nu_{[k]} : \text{SimplOrient}(\mathcal{F}_{\prec[k]}) \rightarrow \text{Zyk}(\text{Cor}([k]))$ gegeben.*
2. *Für jedes $[f] \in \mathcal{K}_n/\sim$ ist eine Menge $\partial_{[f]}$ von genau zwei Randstücken (aus $[f] \times \{\pm 1\}$) gegeben.*
3. *Jeder Standardrand $(f, \varepsilon_f) \in \mathcal{K}_n \times \{\pm 1\}$ liegt in $\partial_{[f]}$.*
4. *Bezüglich jedes Fächers $\nu_{[k]}$ sind je zwei der Randstücke aus $\bigsqcup_{[k] \prec [f]} \partial_{[f]}$ komplementär.*

Falls darüber hinaus auch die Bedingung

5. *Je zwei Fächer $\nu_{[k_1]}$ und $\nu_{[k_2]}$ (mit $[k_1], [k_2] \in \mathcal{K}_{n-1}$) sind kompatibel.*

erfüllt ist, sprechen wir von einem n -Faltkomplex.

Wir definieren die n -Proto-Faltkomplexe aus dem folgenden Grund: Während sich die meisten dieser Eigenschaften relativ leicht durch Faltungen übertragen, ist das bei der Kompatibilität der Fächer eher selten der Fall. Der Begriff der n -Proto-Faltkomplexe stellt dieses Phänomen klar heraus, indem er es deutlich macht, wann die Kompatibilität verwendet werden muss.

Wir haben die Randstücke in Abschnitt 3.1.4 nur für den Fall eingeführt, dass genau eine Flächenklasse vorliegt. Demnach sollten wir die Bedingung aus der Definition der n -Faltkomplexe, die die Eigenschaften der Randstücke enthält, in der Regel nicht benötigen. Dazu muss es zu jeder Flächenklasse eine anliegende Kantenklasse geben, die noch an einer weiteren Flächenklasse angrenzt. In dieser Situation sind insbesondere die Voraussetzungen aus Lemma 3.26 stets erfüllt, sodass wir die linearen Ordnungen auf den Flächenklassen immer aus den Fächern rekonstruieren können.

Bemerkung 3.28. Sei $n > 0$ und $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ ein lokal simplizialer n -Überlagerungskomplex mit simplizialer Oberfläche $\{\sigma_f\}_{f \in \mathcal{K}_n}$, der Bedingungen 1, 2, 3 und 5 erfüllt.

Weiter gebe es für jedes $f \in \mathcal{K}_n$ ein $g \in \mathcal{K}_n$ und ein $k \in \mathcal{K}_{n-1}$, sodass $f \not\sim g$, sowie $[k] \prec [f]$ und $[k] \prec [g]$ gelten.

Dann ist $\mathcal{S}$ ein n -Faltkomplex.

Beweis. Sei $k \in \mathcal{K}_{n-1}$ und $\nu_{[k]}$ der korrespondierende Fächer. Falls es verschiedene $[f] \in \mathcal{K}_n / \sim$ mit $[k] \prec [f]$ gibt, ist die Situation einfach: Dann liegen nämlich nur Standardränder bzgl. $\nu_{[k]}$ vor und diese sind nach Konstruktion komplementär zueinander. (Falls wir uns nicht in der Situation von Bemerkung 3.25 wiederfinden, sind alle vier Oberflächen in diesem Fächer Standardränder. Dann sind davon auch jeweils zwei komplementär.)

Falls es nur ein $[f] \in \mathcal{K}_n / \sim$ mit $[k] \prec [f]$ gibt, müssen wir ein wenig mehr arbeiten. Wir wollen nachweisen, dass die beiden Standardränder aus $\partial_{[f]}$ (aufgrund der Annahme sind es Standardränder, aber nicht bezüglich $\nu_{[k]}$) komplementär bezüglich $\nu_{[k]}$ sind. Sei $\nu_{[l]}$ ein solcher Fächer. Wähle Orientierungen so, dass eine zyklische Reihung $\hat{\nu}_{[k]}$ ein Redukt der zyklischen Reihung $\hat{\nu}_{[l]}$ ist. Sind (g, ε_g) und (h, ε_h) die Randstücke aus $\partial_{[f]}$, dann hat $\hat{\nu}_{[l]}$ ohne Einschränkung (sonst vertausche die Rolle der Randstücke) die Form

$$(\underbrace{\dots}_{\notin [f]}, g, \underbrace{\dots}_{\in [f]}, h, \underbrace{\dots}_{\notin [f]}).$$

In der Notation von Definition 3.24 der komplementären Seiten gilt auch $\varepsilon_g \cdot \hat{\sigma}_g = -1$ und $\varepsilon_h \cdot \hat{\sigma}_h = 1$. Das Redukt dieses Zyklus auf $[f]$ ist dann gleich $(g, \dots, h)$, wo die beiden Randstücke komplementär sind. $\square$

3.2 Definition der Faltung

Wir haben in Abschnitt 3.1 das mengentheoretische Modell aus Kapitel 2 so modifiziert, dass es mit Reihenfolgen von Flächen umgehen kann. Nachdem wir damit einen Begriff für Faltzustände entwickelt haben, betrachten wir jetzt Faltungen dieser n -Faltkomplexe.

Für n -Überlagerungskomplexe haben wir den Formalismus der Erweiterungen und primitiven Erweiterungen entwickelt, um diese Frage zu beantworten. Wir haben auch das Konzept der Teilüberlagerung definiert, mit dem wir erkennen konnten, ob zwei verschiedene n -Überlagerungskomplexe durch Erweiterungen auseinander hervorgegangen sind. Dabei haben wir festgestellt, dass diese Beschreibungen äquivalent sind.

In diesem Abschnitt beschäftigen wir uns mit Erweiterungen, während Teilüberlagerungen der Inhalt von Abschnitt 3.3 sind. Wir werden untersuchen, inwiefern primitive Erweiterungen auf den Kontext der n -Faltkomplexe übertragen werden können. Für alle Grade außer $n-1$ ist dies relativ einfach, aber die Behandlung der Fächer in Grad $n-1$ ist deutlich komplizierter. Wir werden in Abschnitt 3.2.1 ein Hilfsmittel zur Behandlung der Fächer konstruieren, mit dem wir in Abschnitt 3.2.2 die primitiven Erweiterungen von Grad $n-1$ beschreiben können.

Wir beginnen unsere Untersuchung der primitiven Erweiterungen mit den kleinen Graden, die kleiner als $n-1$ ist. In diesen Dimensionen unterscheiden sich n -Faltkomplexe nicht von n -Überlagerungskomplexen, weswegen wir die primitiven Erweiterungen dieser Grade leicht übertragen können.

Definition 3.29 (primitive Faltung). Sei $\mathcal{F}$ ein n -Proto-Faltkomplex und $\mathcal{G}$ eine primitive Erweiterung von $\mathcal{F}$ von Grad $i < n - 1$. Dann heißt der n -Proto-Faltkomplex $\mathcal{G}$ mit Randstücken

$$\partial_{[f]}^{\mathcal{G}} := \partial_{[f]}^{\mathcal{F}} \text{ für alle } [f] \in \mathcal{K}_n / \sim^{\mathcal{F}} = \mathcal{K}_n / \sim^{\mathcal{G}}$$

und Fächern

$$\nu_{[k]}^{\mathcal{G}} := \nu_{[k]}^{\mathcal{F}} \text{ für alle } [k] \in \mathcal{K}_{n-1} / \sim^{\mathcal{F}} = \mathcal{K}_{n-1} / \sim^{\mathcal{G}}$$

primitive Faltung von $\mathcal{F}$ von Grad i .

Da sich die Fächer in primitiven Faltungen von kleinerem Grad als $n - 1$ nicht verändert haben, bleiben sie auch kompatibel, was wir in Bemerkung 3.30 festhalten:

Bemerkung 3.30. Sei $\mathcal{F}$ ein n -Proto-Faltkomplex und $\mathcal{G}$ sei eine primitive Faltung von $\mathcal{F}$ vom Grad $i < n - 1$. Dann ist $\mathcal{F}$ genau dann ein n -Faltkomplex, wenn $\mathcal{G}$ einer ist.

Damit haben wir primitive Faltungen mit kleineren Graden als $n - 1$ abgehandelt. Da die primitiven Faltungen von Grad $n - 1$ eine sehr aufwendige Behandlung erfordern, betrachten wir als nächstes die primitiven Faltungen vom Grad n . Wir hatten schon in Beispiel 3.8 eingesehen, dass wir zur Angabe der Faltung die Randstücke angeben müssen, die zusammengeführt werden sollen. Diese müssen aber noch eine weitere Bedingung erfüllen:

Wenn die primitive Erweiterung die Flächenklassen $[f] \neq [g]$ zusammenführt, dann müssen $\mathcal{F}_{\preccurlyeq[f]}$ und $\mathcal{F}_{\preccurlyeq[g]}$ dieselben Kantenklassen haben, da keine Erweiterungen vom Grad $n - 1$ durchgeführt werden kann. Folglich liegen die beiden angegebenen Randstücke in den selben Fächern. Da sie nach der Faltung komplementär sind (im Sinne von Definition 3.24) und sich die Kantenklassen – und somit die Fächer – nicht ändern, müssen sie schon vorher komplementär gewesen sein.

Definition 3.31. Sei $\mathcal{F}$ ein n -Proto-Faltkomplex mit den Randstücken (f, ε_f) und (g, ε_g) (für $f \not\sim g$), die in denselben Fächern vorkommen und dort stets komplementär sind.

Dann heißt der n -Proto-Faltkomplex $\mathcal{G}$, der aus der primitiven Erweiterung von $\mathcal{F}$, die $f \sim_{\mathcal{G}} g$ erfüllt, sowie den Fächern

$$\nu_{[k]}^{\mathcal{G}} := \nu_{[k]}^{\mathcal{F}} \text{ für alle } [k] \in \mathcal{K}_{n-1} / \sim^{\mathcal{F}} = \mathcal{K}_{n-1} / \sim^{\mathcal{G}}$$

und Randstücken

$$\begin{aligned} \partial_{[h]}^{\mathcal{G}} &:= \partial_{[h]}^{\mathcal{F}} \text{ für alle } [h] \in \mathcal{K}_n / \sim^{\mathcal{F}} \text{ mit } [h] \notin \{[f], [g]\} \\ \partial_{[f]}^{\mathcal{G}} &:= \partial_{[f]}^{\mathcal{F}} \uplus \partial_{[g]}^{\mathcal{G}} \setminus \{(f, \varepsilon_f), (g, \varepsilon_g)\} \end{aligned}$$

besteht, die **primitive Faltung von $\mathcal{F}$ vom Grad n .**

Wohldefiniertheit. Wir zeigen zuerst, dass eine solche primitive Erweiterung $\mathcal{G}$ existiert. Dazu ist es hinreichend zu zeigen, dass jedes $k \in \mathcal{K}_{n-1}$ mit $[k] \prec [f]$ auch $[k] \prec [g]$ erfüllt, da nach Lemma 2.31 dann schon $\mathcal{G}$ existiert. Da f im Fächer um $[k]$ liegt, gilt dies nach Annahme auch für g . Dies ist aber nur dann möglich, wenn $[k] \prec [g]$ erfüllt ist.

Da (f, ε_f) und (g, ε_g) Randstücke bezüglich $\mathcal{F}$ sind, ist $\partial_{[f]}^{\mathcal{G}}$ zweielementig. Da sie ein komplementäres Paar in allen Fächern bilden, stört ihre Entfernung die Paarung der übrigen Randstücke nicht. Nach der Konstruktion sind die beiden keine Standardränder mehr, demnach ist auch die Standardrand-Bedingung eines n -Proto-Faltkomplexes erfüllt. $\square$

Da sich auch in der primitiven Faltung vom Grad n die Fächer nicht verändern, bleibt auch hier die Fächerkompatibilität erfüllt, falls $\mathcal{F}$ bereits ein n -Faltkomplex ist.

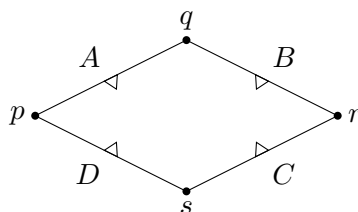
Bemerkung 3.32. Sei $\mathcal{F}$ ein n -Proto-Faltkomplex und $\mathcal{G}$ sei eine primitive Faltung von $\mathcal{F}$ vom Grad n . Dann ist $\mathcal{F}$ genau dann ein n -Faltkomplex, wenn $\mathcal{G}$ einer ist.

Nachdem wir die primitiven Faltungen der Grade ungleich $n - 1$ definiert haben, wenden wir uns jetzt den primitiven Erweiterungen vom Grad $n - 1$ zu. Dazu müssen wir zunächst klären, wie wir verschiedene Fächer kombinieren können. In Abschnitt 3.2.1 werden wir uns umfassend mit dieser Frage beschäftigen. Im Anschluss können wir die primitiven Faltungen vom Grad $n - 1$ in Abschnitt 3.2.2 definieren und untersuchen, wieso es schwierig ist, die Fächerkompatibilität zu garantieren.

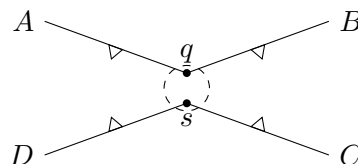
3.2.1 Definition der Fächersumme

Dieser Abschnitt ist der Untersuchung der Fächersumme gewidmet, also der Kombination verschiedener Fächer während einer Faltung. Dass wir diese Betrachtung ausführen müssen, illustriert das folgende Beispiel.

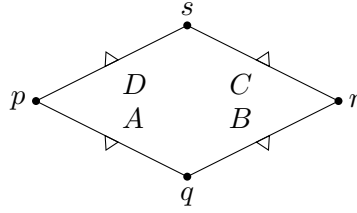
Beispiel 3.33. Wir betrachten den folgenden 1-Faltkomplex, bei dem wir die simpliziale Oberfläche gemäß Bemerkung 3.12 durch einen Pfeil kennzeichnen.



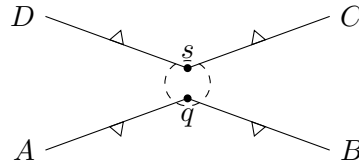
Wir wollen die Punkte q und s zusammenfalten, ohne Kanten zu identifizieren. Das Bild suggeriert, dass die Oberflächen $(A, +1)$ und $(D, +1)$ nach dieser Faltung komplementär bezüglich $\{q, s\}$ sind. In diesem Fall hätte die zyklische Reihung am Punkt $[q] = [s]$ die folgende Gestalt:



Wenn man aber versucht, dies zu formalisieren, stellt man fest, dass uns die Illustration etwas untergejubelt hat, das nicht existierte. Es ist nämlich auch möglich, den Faltkomplex wie folgt darzustellen:



In dieser Darstellung hat man das Gefühl, dass die Oberflächen $(A, -1)$ und $(D, -1)$ nach der Faltung komplementär bezüglich $\{q, s\}$ sind. Die resultierende zyklische Reihung sieht der obigen sehr ähnlich.



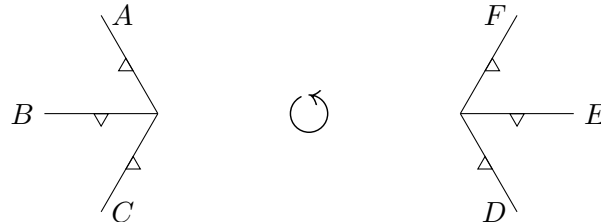
Demnach müssen wir zusätzlich zu den identifizierten Punkten auch angeben, wie die Kantenoberflächen nach der Identifikation liegen sollen. Man muss sich jetzt natürlich fragen, ob auch $(A, +1)$ und $(D, -1)$ komplementär werden können. Das wird aber durch den Fächer an $[p]$ verhindert, da diese beiden Elemente dort nicht komplementär sind und er sich durch Reduktbildung aus dem gefalteten Fächer ergeben müsste.

Das Beispiel demonstriert, dass wir Kanten nicht „einfach so“ zusammenfalten können – wir müssen immer sagen, welche Oberflächen dadurch komplementär werden. Es zeigt aber auch, dass wir uns mit der Kombination von Fächern befassen müssen, denn an diesen können wir die verschiedenen Faltungen des selben Kantenpaars unterscheiden.

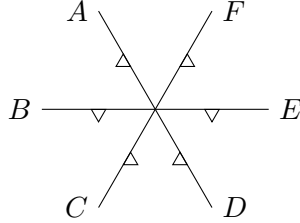
Für zwei zu faltende Kanten zerfällt dieses Problem in zwei Teile, nämlich die Kombination zyklischer Reihungen (bei gleicher Orientierung) und die Behandlung der Orientierung. Um die Orientierungen zu synchronisieren, betrachten wir die Identifikation vom Grad $n - 1$, die die beiden Kanten identifiziert. Diese Identifikation induziert eine bijektive Abbildung zwischen den simplizialen Orientierungen der beiden Kanten. Aufgrund dieser Bijektion können wir bei beiden Kanten die „gleiche“ Orientierung wählen.

Folglich müssen wir uns noch damit beschäftigen, wie wir zwei verschiedene zyklische Reihungen kombinieren können. Das Beispiel 3.34 demonstriert diese Kombination exemplarisch.

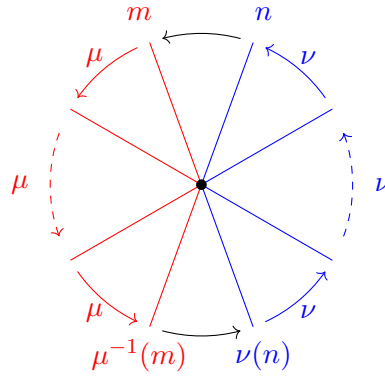
Beispiel 3.34. Wir betrachten die zyklischen Reihungen (A, B, C) und (D, E, F) , die gleich orientiert sind:



Wenn wir die beiden zyklischen Reihungen so an ihren Zentren zusammenführen wollen, dass A und F nebeneinander liegen (und ohne eine davon zu spiegeln), muss entweder $(A, +1)$ an $(F, +1)$ anliegen oder $(A, -1)$ an $(F, -1)$. Mit anderen Worten: Die beiden Oberflächen, die zusammengefügt werden sollen, müssen entgegengesetzt orientiert sein. Ein Aneinanderfügen an $(A, -1)$ und $(F, -1)$ liefert:



In Beispiel 3.34 haben wir zwei zyklische Reihungen zusammengeführt. Um diesen Prozess allgemein zu definieren, starten wir mit zwei zyklischen Reihungen $\mu : M \rightarrow M$ und $\nu : N \rightarrow N$ mit $M \cap N = \emptyset$. Wir wollen diese beiden zyklischen Reihungen an den Elementen $m \in M$ und $n \in N$ kombinieren. Diese Kombination können wir wie folgt illustrieren:



Dabei haben wir nur den Fall abgebildet, in dem n auf m abgebildet wird. Den anderen Fall können wir dadurch erreichen, dass wir die Rollen von μ und ν vertauschen. Da wir mit diesem Konzept also alle möglichen Kombinationen aus Beispiel 3.34 konstruieren können, formalisieren wir diese Konstruktion in der Definition 3.35.

Definition 3.35 (Konfluenz). Seien $\mu : M \rightarrow M$ und $\nu : N \rightarrow N$ zyklische Reihungen mit $M \cap N = \emptyset$, sowie $m \in M, n \in N$. Die **Konfluenz an** (m, n) ist die zyklische Reihung $\mu \overset{m}{\underset{n}{\rightrightarrows}} \nu$ auf $M \uplus N$, die wie folgt definiert ist:

$$\mu \overset{m}{\underset{n}{\rightrightarrows}} \nu : M \uplus N \rightarrow M \uplus N : x \mapsto \begin{cases} \mu(x) & x \in M \setminus \{\mu^{-1}(m)\} \\ \nu(n) & x = \mu(m) \\ \nu(x) & x \in N \setminus \{n\} \\ m & x = n \end{cases}$$

Die Konfluenz ist nicht kommutativ, aber wir können eine Variante beweisen, wenn wir gleichzeitig die zyklischen Reihungen invertieren.

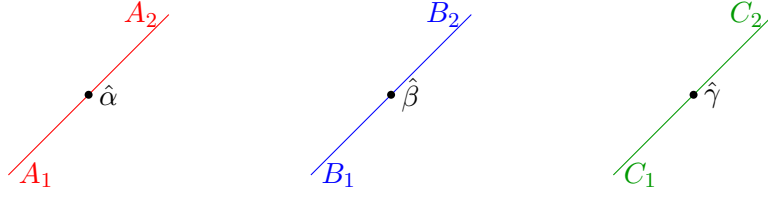
$$\left(\begin{matrix} \textcolor{red}{m} \\ \textcolor{red}{\mu} \rightleftharpoons \textcolor{blue}{\nu} \\ \textcolor{blue}{n} \end{matrix} \right)^{-1} = \textcolor{blue}{\nu}^{-1} \begin{matrix} \textcolor{blue}{n} \\ \textcolor{blue}{\nu} \rightleftharpoons \textcolor{red}{\mu}^{-1} \\ \textcolor{red}{m} \end{matrix}.$$

$$x \mapsto \begin{cases} \mu^{-1}(x) & x \in M \setminus \{f\} \\ \mu^{-1}(f) & x = \nu(g) \\ \nu^{-1}(x) & x \in N \setminus \{\nu(g)\} \\ g & x = f \end{cases}$$

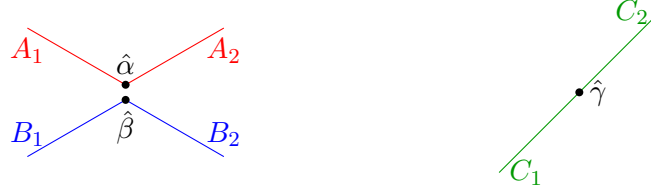
$$x \mapsto \begin{cases} \nu^{-1}(x) & x \in N \setminus \{\nu(g)\} \\ \mu^{-1}(f) & x = \nu(g) \\ \mu^{-1}(x) & x \in M \setminus \{f\} \\ g & x = f \end{cases}$$

The diagram illustrates the mapping $x \mapsto$ for elements x in the sets N and M . The mapping is defined by the piecewise function above. The diagram shows two sets of rays originating from a central point. The left set of rays (blue) represents the image of N under the map, and the right set of rays (red) represents the image of M under the map. The rays are labeled with elements of N and M , and their images under the map. The mapping is defined by the piecewise function above.

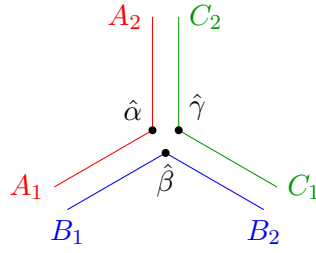
Beispiel 3.37. Wir betrachten zyklische Reihungen (in Zykelschreibweise) $\hat{\alpha} = (A_1, A_2)$, $\hat{\beta} = (B_1, B_2)$ und $\hat{\gamma} = (C_1, C_2)$, die wir wie folgt visualisieren können:



Wir bestimmen zuerst die Konfluenz $\hat{\beta} \xrightarrow[A_1]{B_1} \hat{\alpha}$. Wir erhalten den Zykel (A_1, B_1, B_2, A_2) :



Als nächstes bilden wir die Konfluenz von $\hat{\gamma}$ mit diesem, via $\hat{\gamma} \xrightarrow[B_2]{C_1} \left(\hat{\beta} \xrightarrow[A_1]{B_1} \hat{\alpha} \right)$, und erhalten den Zykel $(A_1, B_1, B_2, C_1, C_2, A_2)$:



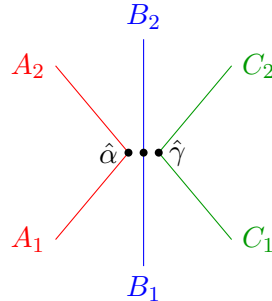
Um dieses Ergebnis zu erhalten, hätten wir die Konfluenzen auch in dieser Reihenfolge berechnen können:

$$\left(\hat{\gamma} \xrightarrow[B_2]{C_1} \hat{\beta} \right) \xrightarrow[A_1]{B_1} \hat{\alpha}$$

Während man in diesem Fall sogar zuerst eine Konfluenz von $\hat{\alpha}$ und $\hat{\gamma}$ hätte ausführen können, ist dies bei

$$\hat{\alpha} \xrightarrow[B_2]{A_1} \left(\hat{\beta} \xrightarrow[C_2]{B_2} \hat{\gamma} \right) = \left(\hat{\alpha} \xrightarrow[B_2]{A_1} \hat{\beta} \right) \xrightarrow[C_2]{B_2} \hat{\gamma}$$

nicht der Fall. Hier enden wir mit $(A_1, B_1, C_1, C_2, B_2, A_2)$, visualisiert wie folgt:

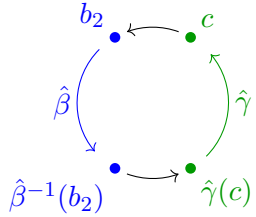


Wir haben in Beispiel 3.37 gesehen, dass die Konfluenz dieser zyklischer Reihungen assoziativ ist. Diese Eigenschaft gilt sogar allgemein, wie wir in Lemma 3.38 nachweisen.

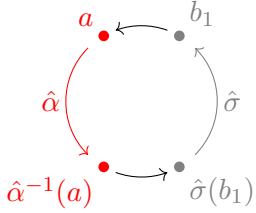
Lemma 3.38. Seien $\hat{\alpha} : A \rightarrow A$, $\hat{\beta} : B \rightarrow B$ und $\hat{\gamma} : C \rightarrow C$ zyklische Reihungen und die Elemente $a \in A$, $b_1, b_2 \in B$ und $c \in C$ gegeben. Dann gilt

$$\hat{\alpha} \xrightarrow[b_1]{a} \left(\hat{\beta} \xrightarrow[c]{b_2} \hat{\gamma} \right) = \left(\hat{\alpha} \xrightarrow[b_1]{a} \hat{\beta} \right) \xrightarrow[c]{b_2} \hat{\gamma}.$$

Beweis. Wir bestimmen $\hat{\alpha} \xrightarrow[b_1]{a} \left(\hat{\beta} \xrightarrow[c]{b_2} \hat{\gamma} \right)$ explizit. Die Konfluenz $\hat{\sigma} := \hat{\beta} \xrightarrow[c]{b_2} \hat{\gamma}$ hat die Form

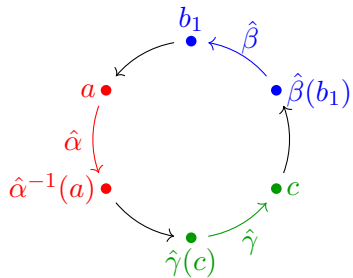
$$x \mapsto \begin{cases} \hat{\beta}(x) & x \in B \setminus \{\hat{\beta}^{-1}(b_2)\} \\ \hat{\gamma}(c) & x = \hat{\beta}^{-1}(b_2) \\ \hat{\gamma}(x) & x \in C \setminus \{c\} \\ b_2 & x = c \end{cases}$$


Nun bestimmen wir $\hat{\alpha} \xrightarrow[b_1]{a} \hat{\sigma}$. Dies hat im Allgemeinen die folgende Gestalt:

$$x \mapsto \begin{cases} \hat{\alpha}(x) & x \in A \setminus \{\hat{\alpha}^{-1}(b_1)\} \\ \hat{\sigma}(b_1) & x = \hat{\alpha}^{-1}(a) \\ \hat{\sigma}(x) & x \in B \uplus C \setminus \{b_1\} \\ a & x = b_1 \end{cases}$$


Um die Abbildung $\hat{\sigma}$ in dieser Abbildungsvorschrift zu ersetzen, müssen wir danach unterscheiden, was $\hat{\sigma}(b_1)$ ist. Hier ergeben sich zwei Fälle, nämlich $b_1 = \hat{\beta}^{-1}(b_2)$ und $b_1 \neq \hat{\beta}^{-1}(b_2)$.

1. Aus $b_1 = \hat{\beta}^{-1}(b_2)$ folgt $\hat{\sigma}(b_1) = \hat{\gamma}(c)$. Einsetzen der Abbildungsvorschriften ineinander ergibt nun:

$$x \mapsto \begin{cases} \hat{\alpha}(x) & x \in A \setminus \{\hat{\alpha}^{-1}(a)\} \\ \hat{\gamma}(c) & x = \hat{\alpha}^{-1}(a) \\ \hat{\gamma}(x) & x \in C \setminus \{c\} \\ \hat{\beta}(b_1) & x = c \\ \hat{\beta}(x) & x \in B \setminus \{b_1\} \\ a & x = b_1 \end{cases}$$


2. Aus $b_1 \neq \hat{\beta}^{-1}(b_2)$ folgt $\hat{\sigma}(b_1) = \hat{\beta}(b_1)$. Einsetzen der Abbildungsvorschriften ineinander ergibt nun:

$$x \mapsto \begin{cases} \hat{\alpha}(x) & x \in A \setminus \{\hat{\alpha}^{-1}(b_1)\} \\ \hat{\beta}(b_1) & x = \hat{\alpha}^{-1}(a) \\ \hat{\beta}(x) & x \in B \setminus \{b_1, \hat{\beta}^{-1}(b_2)\} \\ \hat{\gamma}(c) & x = \hat{\beta}^{-1}(b_2) \\ \hat{\gamma}(x) & x \in C \setminus \{c\} \\ b_2 & x = c \\ a & x = b_1 \end{cases}$$

Als nächstes bestimmen wir eine explizite Form von $\left(\hat{\alpha} \xrightarrow[b_1]{a} \hat{\beta}\right) \xrightarrow[c]{b_2} \hat{\gamma}$. Dazu starten wir mit der Angabe von $\hat{\sigma} := \hat{\alpha} \xrightarrow[b_1]{a} \hat{\beta}$:

$$x \mapsto \begin{cases} \hat{\alpha}(x) & x \in A \setminus \{\hat{\alpha}^{-1}(a)\} \\ \hat{\beta}(b_1) & x = \hat{\alpha}^{-1}(a) \\ \hat{\beta}(x) & x \in B \setminus \{b_1\} \\ a & x = b_1 \end{cases}$$

Dann hat $\hat{\sigma} \xrightarrow[c]{b_2} \hat{\gamma}$ die Gestalt

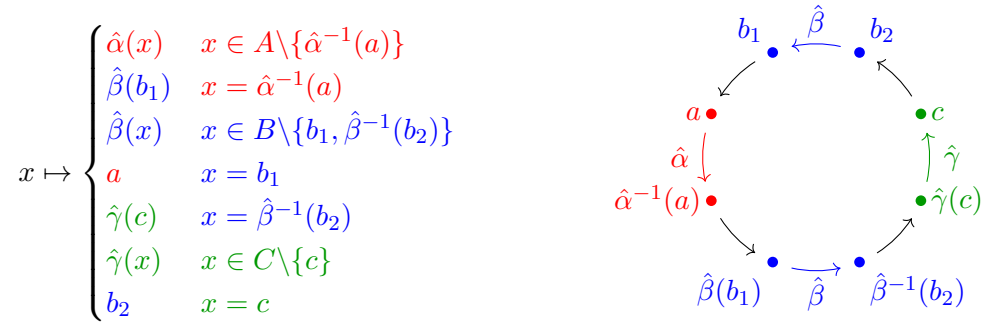
$$x \mapsto \begin{cases} \hat{\sigma}(x) & x \in A \uplus B \setminus \{\hat{\sigma}^{-1}(a)\} \\ \hat{\gamma}(c) & x = \hat{\sigma}^{-1}(b_2) \\ \hat{\gamma}(x) & x \in C \setminus \{c\} \\ b_2 & x = c \end{cases}$$

Wenn man diese beiden Abbildungsvorschriften ineinander einzusetzen versucht, stellt man ebenfalls die Notwendigkeit von zwei Fällen fest.

1. Im Fall $b_2 = \hat{\beta}(b_1)$ folgt $\hat{\sigma}^{-1}(b_2) = \hat{\alpha}^{-1}(a)$. Die Einsetzung ergibt dann

$$x \mapsto \begin{cases} \hat{\alpha}(x) & x \in A \setminus \{\hat{\alpha}^{-1}(a)\} \\ \hat{\beta}(x) & x \in B \setminus \{b_1\} \\ a & x = b_1 \\ \hat{\gamma}(c) & x = \hat{\alpha}^{-1}(a) \\ \hat{\gamma}(x) & x \in C \setminus \{c\} \\ \hat{\beta}(b_1) & x = c \end{cases}$$

2. Im Fall $b_2 \neq \hat{\beta}(b_1)$ gilt $\hat{\sigma}^{-1}(b_2) = \hat{\beta}^{-1}(b_2)$ und die Einsetzung ergibt



Wenn wir die Abbildungsvorschriften in den verschiedenen Fällen vergleichen, sehen wir, dass diese tatsächlich identisch sind. $\square$

Nachdem wir uns ausführlich mit strukturellen Eigenschaften der Konfluenz beschäftigt haben, kehren wir jetzt zur eigentlichen Aufgabenstellung zurück: Wir benötigen die Konfluenz zyklischer Reihungen, um die Kombination von zwei Fächern zu definieren. Dabei müssen wir uns auch überlegen, wie die verschiedenen Orientierungen der beiden Kanten miteinander verglichen werden. Da wir diesen Vergleich über die Identifikation der Kanten definieren, können wir in der Definition der Fächerkombination annehmen, dass diese bereits gleich orientiert sind.

Definition 3.39 (Fächersumme). *Sei $\mathcal{S}$ ein lokal simplizialer n -Überlagerungskomplex mit simplizialer Oberfläche $\{\sigma_f\}_{f \in \mathcal{K}_n}$ und $[k] \in \mathcal{K}_{n-1}/\sim$. Weiter seien $\{M_1, M_2\}$ eine Partition der Corona $\text{Cor}([k])$ und*

$$\begin{aligned} \nu_1 &: \text{SimplOrient}(\mathcal{S}_{\preccurlyeq[k]}) \rightarrow \text{Zyk}(M_1) \\ \nu_2 &: \text{SimplOrient}(\mathcal{S}_{\preccurlyeq[k]}) \rightarrow \text{Zyk}(M_2) \end{aligned}$$

zwei Fächer.

Seien Oberflächen $(f, \varepsilon_f) \in M_1 \times \{\pm 1\}$ und $(g, \varepsilon_g) \in M_2 \times \{\pm 1\}$ entgegengesetzt orientiert (bzgl. $[k]$). Dann definieren wir die **Fächersumme von ν_1 und ν_2 an (f, ε_f) und (g, ε_g) als den Fächer**

$$\rho : \text{SimplOrient}(\mathcal{F}_{\preccurlyeq[k]}) \rightarrow \text{Zyk}(M_1 \uplus M_2),$$

bei dem $\rho(a)$ für $a \in \text{SimplOrient}(\mathcal{F}_{\preccurlyeq[k]})$ wie folgt definiert ist:

- Falls $\varepsilon_f \bar{\sigma}_f(a) = -1$ und $\varepsilon_g \bar{\sigma}_g(a) = +1$, dann handelt es sich um die Konfluenz von $\nu_1(a)$ und $\nu_2(a)$ an f und g .
- Falls $\varepsilon_f \bar{\sigma}_f(a) = +1$ und $\varepsilon_g \bar{\sigma}_g(a) = -1$, dann handelt es sich um die Konfluenz von $\nu_2(a)$ und $\nu_1(a)$ an g und f .

Dabei bezeichnen $\bar{\sigma}_f$ und $\bar{\sigma}_g$ die induzierten Durchlaufsinne bezüglich $[k]$.

Wohldefiniertheit. Wir müssen nachweisen, dass wir auf diese Weise tatsächlich einen Fächer definiert haben. Dazu müssen wir prüfen, ob er mit der Operation von S_n verträglich ist.

Prüfe zuerst, ob die Fallunterscheidung vollständig ist. Sei dazu $\pi \in S_n$, dann gelten

$$\begin{aligned}\bar{\sigma}_f(\pi \cdot a) &= \text{sign}(\pi) \cdot \bar{\sigma}_f(a) \\ \bar{\sigma}_g(\pi \cdot a) &= \text{sign}(\pi) \cdot \bar{\sigma}_g(a),\end{aligned}$$

folglich ist die Fallunterscheidung vollständig.

Wir müssen außerdem zeigen, dass die Inverse der Konfluenz von $\nu_1(a)$ und $\nu_2(a)$ an f und g gleich zur Konfluenz von $\nu_2(a)^{-1}$ und $\nu_1(a)^{-1}$ an g und f ist. Dies ist aber die Aussage aus Bemerkung 3.36. $\square$

Wir führen die Fächersumme an einem minimalen Beispiel durch:

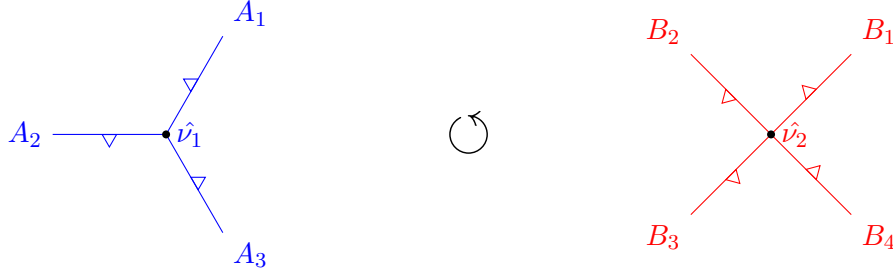
Beispiel 3.40. Wir betrachten einen lokal simplizialen 2-Überlagerungskomplex $\mathcal{S}$ mit

$$\begin{aligned}\text{SimplOrient}(\mathcal{S}_{\leq[k]}) &= \{(p_1, p_2), (p_2, p_1)\} \\ \text{Cor}_{\mathcal{S}}([k]) &= \{A_1, A_2, A_3, B_1, B_2, B_3, B_4\} \\ M_1 &= \{A_1, A_2, A_3\} \\ M_2 &= \{B_1, B_2, B_3, B_4\}\end{aligned}$$

sowie den Fächern ν_1 und ν_2 , die wie folgt definiert sind:

$$\begin{aligned}\hat{\nu}_1 &:= \nu_1(p_1, p_2) = (A_1, A_2, A_3) & \nu_1(p_2, p_1) &= (A_3, A_2, A_1) \\ \hat{\nu}_2 &:= \nu_2(p_1, p_2) = (B_1, B_2, B_3, B_4) & \nu_2(p_2, p_1) &= (B_4, B_3, B_2, B_1)\end{aligned}$$

Für die simpliziale Orientierung (p_1, p_2) können wir diese wie folgt darstellen (bezüglich der Standardorientierung der Ebene):



Dabei haben wir der Einfachheit halber angenommen, dass die induzierten Durchlaufsinne für alle $f \in M_1 \uplus M_2$ gleich definiert sind:

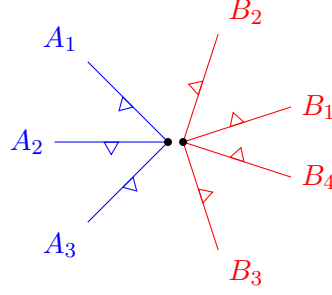
$$\bar{\sigma}_f(p_1, p_2) = +1 \quad \bar{\sigma}_f(p_2, p_1) = -1$$

Folglich sind (f, ε_f) und (g, ε_g) für $f, g \in M_1 \uplus M_2$ genau dann entgegengesetzt orientiert, wenn $\varepsilon_f \neq \varepsilon_g$ gilt.

Wir führen exemplarisch die Faltung an $(A_1, -1)$ und $(B_2, +1)$ aus. Dabei erhalten wir aus Definition 3.39 die Konstruktion eines Fächers

$$\rho : \text{SimplOrient}(\mathcal{S}_{\leq[k]}) \rightarrow \text{Zyk}(M_1 \uplus M_2).$$

$\rho(p_1, p_2)$ ist die Konfluenz von $\nu_2(p_1, p_2)$ und $\nu_1(p_1, p_2)$ an B_2 und A_1 , also der γ -Zykel $(A_1, A_2, A_3, B_3, B_4, B_1, B_2)$, oder graphisch:



$\rho(p_2, p_1)$ ergibt sich als das Inverse dieses 7-Zykels.

3.2.2 Primitive Faltung vom Grad $n - 1$

Nachdem wir uns in Abschnitt 3.2.1 ausführlich mit den Details der Kombination von Fächern beschäftigt haben, wenden wir uns in diesem Abschnitt der Definition der primitiven Faltung vom Grad $n - 1$ zu. Nach unserer Vorarbeit ist deren Definition 3.41 relativ einfach.

Definition 3.41 (primitive Faltung). *Sei $\mathcal{F}$ ein n -Proto-Faltkomplex und $k, l \in \mathcal{K}_{n-1}$ mit $k \not\sim^{\mathcal{F}} l$ und $\mathcal{G}$ eine primitive Erweiterung von $\mathcal{F}$, die $k \sim^{\mathcal{G}} l$ erfüllt. Außerdem seien Randstücke $(f, \varepsilon_f), (g, \varepsilon_g)$ mit $[k] \prec [f]$ und $[l] \prec [g]$ gegeben, die bezüglich $[k]_{\mathcal{G}}$ entgegengesetzt orientiert sind.*

Dann induziert die primitive Erweiterung zwei Bijektionen

$$\begin{aligned} \beta_{[k]_{\mathcal{F}}} &: \text{SimplOrient}(\mathcal{F}_{\prec[k]}) \rightarrow \text{SimplOrient}(\mathcal{G}_{\prec[k]}) \\ \beta_{[l]_{\mathcal{F}}} &: \text{SimplOrient}(\mathcal{F}_{\prec[l]}) \rightarrow \text{SimplOrient}(\mathcal{G}_{\prec[l]}), \end{aligned}$$

sodass $\nu_{[k]} \circ \beta_{[k]}^{-1}$ und $\nu_{[l]} \circ \beta_{[l]}^{-1}$ Fächer sind.

Dann heißt der n -Proto-Faltkomplex $\mathcal{G}$ mit den Randstücken

$$\partial_{[f]}^{\mathcal{G}} := \partial_{[f]}^{\mathcal{F}} \text{ für alle } [f] \in \mathcal{K}_n / \sim^{\mathcal{F}} = \mathcal{K}_n / \sim^{\mathcal{G}}$$

und den Fächern

$$\begin{aligned} \nu_{[t]}^{\mathcal{G}} &:= \nu_{[t]}^{\mathcal{F}} \text{ für alle } t \in \mathcal{K}_{n-1} \text{ mit } t \not\sim^{\mathcal{F}} k \text{ und } t \not\sim^{\mathcal{F}} l \\ \nu_{[k]}^{\mathcal{G}} &\text{ ist die Fächersumme von } \nu_{[k]} \circ \beta_{[k]}^{-1} \text{ und } \nu_{[l]} \circ \beta_{[l]}^{-1} \text{ an } (f, \varepsilon_f) \text{ und } (g, \varepsilon_g) \end{aligned}$$

primitive Faltung von $\mathcal{F}$ vom Grad $n - 1$.

Wohldefiniertheit. Damit wir einen n -Proto-Faltkomplex gemäß Definition 3.27 vorliegen haben, müssen wir die Komplementbedingung und die Standardränder überprüfen.

Die Komplementbedingung muss nur bezüglich $\nu_{[k]}^{\mathcal{G}}$ geprüft werden, da sich die anderen Fächer nicht verändert haben. Nach Voraussetzung waren (f, ε_f) und (g, ε_g) vorher Randstücke, hatten also Komplemente bezüglich $\nu_{[k]}^{\mathcal{F}}$ und $\nu_{[l]}^{\mathcal{G}}$. Nach der Fächersumme sind (f, ε_f) und (g, ε_g) komplementär, weshalb auch ihre ursprünglichen Komplemente nun komplementär zueinander sind.

Untersuche nun, ob die neuen Standardränder bereits Randstücke waren. Da keine Äquivalenzklassen von $\mathcal{K}_n$ vereinigt werden, kann es nur bei $\nu_{[k]}^{\mathcal{G}}$ neue Standardränder geben. Genauer können diese nur bei (f, ε_f) und (g, ε_g) (sowie deren Komplementen in $\mathcal{F}$) auftreten. Alle diese sind aber bereits Randstücke. $\square$

Im Gegensatz zu den primitiven Faltungen von Graden ungleich $n - 1$ ist es sehr schwer, die Fächerkompatibilität allgemein zu garantieren. Wir müssen daher mit einer schwächeren Aussage Vorlieb nehmen.

Bemerkung 3.42. Sei $\mathcal{F}$ ein n -Faltkomplex und $\mathcal{G}$ sei eine primitive Faltung von $\mathcal{F}$ vom Grad $n - 1$. Falls $\mathcal{G}$ ein n -Faltkomplex ist, dann ist auch $\mathcal{F}$ einer.

Beweis. Dass $\mathcal{F}$ ein n -Faltkomplex ist, wenn $\mathcal{G}$ einer ist, folgt daraus, dass man anstatt einer doppelten Reduktbildung nur eine davon ausführen muss, wie Lemma 3.15 zeigt. $\square$

Wir möchten aber eigentlich wissen, wann wir die Fächerkompatibilität garantieren können. Äquivalent können wir fragen, wann $\mathcal{G}$ ein n -Faltkomplex ist, sobald $\mathcal{F}$ einer ist. Da wir kein simples, allgemeines Kriterium zum Garantieren der Fächerkompatibilität haben, muss diese in praktischen Anwendungen nach jeder Faltung manuell überprüft werden.

Im Rest dieses Abschnitts werden wir uns anhand mehrerer Beispiele damit beschäftigen, welche Phänomene bei der Fächersumme auftreten können, die es schwierig machen, die Fächerkompatibilität allgemein zu garantieren. Wir werden anhand dieser Beispiele auch aufzeigen, wieso einige naive Ideen zur Garantie der Fächerkompatibilität fehlschlagen.

Die erste naive Idee beruht darauf, dass (f, ε_f) und (g, ε_g) nach der Faltung komplementär sind. Das gilt natürlich auch für ihre jeweiligen Komplemente. Folglich müssen diese Paare auch vor der Faltung komplementär sein, falls sie im selben Fächer vorkommen. Es genügt aber nicht, diese Eigenschaft zu fordern, um die Fächerkompatibilität allgemein zu garantieren, wie Beispiel 3.43 zeigt.

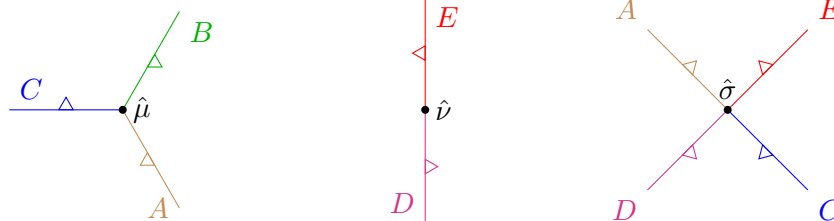
Beispiel 3.43. Wir betrachten die zyklischen Reihungen von drei Fächern, die gleich orientiert sind (daher beschäftigen wir uns nur mit den zyklischen Reihungen):

$$\hat{\mu} := (A, B, C)$$

$$\hat{\nu} := (D, E)$$

$$\hat{\sigma} := (A, D, C, E)$$

Wir visualisieren diese wie folgt (wobei wir die simplizialen Oberflächen kennzeichnen):

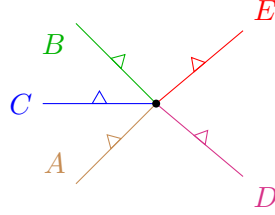


Wir bemerken, dass sowohl $\hat{\mu}$ als auch $\hat{\nu}$ mit $\hat{\sigma}$ kompatibel sind. Als nächstes möchten wir $\hat{\mu}$ und $\hat{\nu}$ an $(A, +1)$ und $(D, +1)$ zusammenfügen, d. h. die Konfluenz von $\hat{\nu}$ und $\hat{\mu}$ an D und A ausführen.

Gemäß unseren Vorüberlegungen soll nach der Konfluenz gelten:

- $(A, +1)$ und $(D, +1)$ sind komplementär.
- $(B, -1)$ und $(E, -1)$ sind komplementär (dies sind die Komplemente von $(A, +1)$ und $(D, +1)$ bezüglich $\hat{\mu}$ und $\hat{\nu}$).

Ersichtlich ist die erste Bedingung auch in $\hat{\sigma}$ erfüllt, die zweite ist schlicht nicht anwendbar. Das Ausführen der Konfluenz liefert (A, D, E, B, C) , oder graphisch:



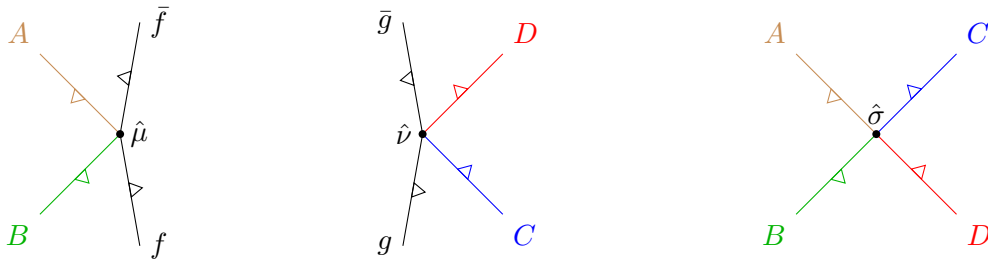
Aber diese zyklische Reihung ist nicht mit $\hat{\sigma}$ kompatibel (genauer: die zugehörigen Fächer sind nicht kompatibel).

Man kann die Situation aus Beispiel 3.43 dadurch verhindern, dass man fordert, dass die Mengen $\{A, C\}$ und $\{D, E\}$ niemals alternieren dürfen. Damit würde man absichern, dass die Coronae von $\hat{\mu}$ und $\hat{\nu}$ überschneidungsfrei in der Konfluenz der beiden sind. Allerdings lässt sich auch dann noch ein Schlupfloch finden, da wir nicht vernünftig mit den Übergängen zwischen den zyklischen Reihungen umgehen können. Das Beispiel 3.44 demonstriert dieses Problem.

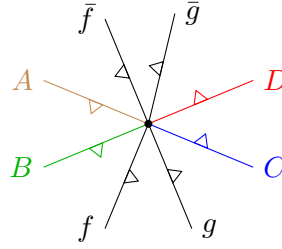
Beispiel 3.44. Wir betrachten die zyklischen Reihungen von drei Fächern, die gleich orientiert sind (daher beschäftigen wir uns nur mit den zyklischen Reihungen):

$$\hat{\mu} := (f, \bar{f}, A, B) \qquad \hat{\nu} := (g, C, D, \bar{g}) \qquad \hat{\sigma} := (A, B, D, C)$$

Wir visualisieren diese wie folgt (wobei wir die simplizialen Oberflächen kennzeichnen):



Wir bemerken, dass sowohl $\hat{\mu}$ als auch $\hat{\nu}$ mit $\hat{\sigma}$ kompatibel sind. Als nächstes möchten wir $\hat{\mu}$ und $\hat{\nu}$ an $(f, +1)$ und $(g, -1)$ zusammenfügen, d. h. die Konfluenz von $\hat{\nu}$ und $\hat{\mu}$ an g und f ausführen. Dies liefert $(A, B, f, g, D, C, \bar{g}, \bar{f})$, oder graphisch:



Aber diese zyklische Reihung ist nicht mit $\hat{\sigma}$ kompatibel (genauer: die zugehörigen Fächer sind nicht kompatibel).

Wir haben in den Beispielen 3.43 und 3.44 gesehen, dass wir die Fächerkompatibilität nicht durch einfache Komplementaritäts- und Überschneidungsfreiheitsbedingungen garantieren können.

Aber selbst wenn zwei Faltungen kompatible Fächer liefern, kann es passieren, dass die Kombination dieser Faltungen dies nicht tut. Das Beispiel 3.45 demonstriert dieses Phänomen.

Beispiel 3.45. Wir betrachten einen 1-Faltkomplex $\mathcal{F}$, bestehend aus

$$\mathcal{K}_0 := \{p, q, r, s\}$$

$$\mathcal{K}_1 := \{A, B, C, D\}$$

mit

$$p \prec A$$

$$p \prec B$$

$$q \prec C$$

$$q \prec D$$

$$r \prec B$$

$$r \prec C$$

$$s \prec A$$

$$s \prec D.$$

und trivialer Äquivalenzrelation $\sim$. Für die simplizialen Orientierungen geben wir jeder Kante eine Richtung:

$$A \leftrightarrow (p, s)$$

$$B \leftrightarrow (p, r)$$

$$C \leftrightarrow (q, r)$$

$$D \leftrightarrow (q, s)$$

Der Einfachheit halber identifizieren wir den Fächer ν_p mit der eindeutigen zyklischen Reihung $\nu_p(p)$. Da an jedem Punkt genau zwei Kanten anliegen, haben wir

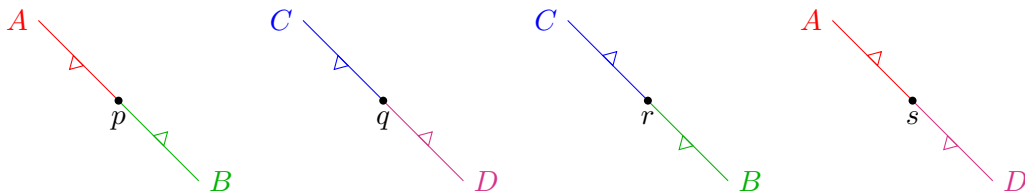
$$\nu_p = (A, B)$$

$$\nu_q = (C, D)$$

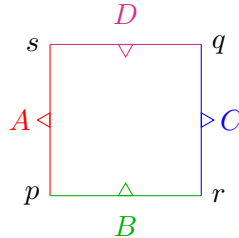
$$\nu_r = (B, C)$$

$$\nu_s = (A, D).$$

Graphisch können wir diese zyklischen Reihungen (inklusive Orientierung) wie folgt darstellen:



Diese Darstellung mag noch nicht allzu aussagekräftig sein, aber wir müssen im Hinterkopf behalten, dass unser Formalismus nicht mehr als das zur Verfügung stehen hat. Außerhalb des Formalismus können wir diese zyklischen Reihungen in eine geschlossene Form bringen:

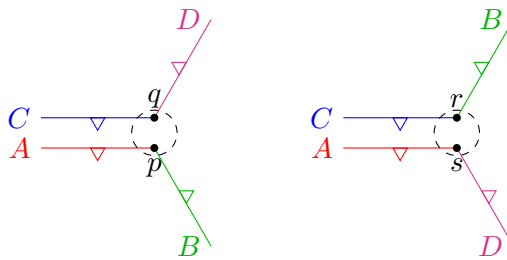


Wir wollen nun mit der Identifikation

$$\varphi : \mathcal{F}_{\preccurlyeq A} \rightarrow \mathcal{F}_{\preccurlyeq C} \quad \begin{cases} p \mapsto q \\ s \mapsto r \\ A \mapsto C \end{cases}$$

die Oberflächen $(A, -1)$ und $(C, +1)$ zusammenfalten.

Führen wir die Fächersummen aus, gelangen wir zu folgenden Resultaten:



Insbesondere liegen die nicht kompatiblen Zyklen (A, B, D, C) und (A, D, B, C) vor. Damit liefert diese Faltung keinen 1-Faltkomplex.

In Beispiel 3.45 sind beide Faltungen von Grad $n - 1$ individuell ausführbar. Erst nachdem eine davon ausgeführt wurde, wird die Problematik der anderen sichtbar. Jede Kriterium, das die Fächerkompatibilität garantiert, muss auch in diesem Beispiel funktionieren.

3.3 Mehrfachfaltung

Nachdem wir uns in Abschnitt 3.2 damit beschäftigt haben, die Faltung eines n -Faltkomplexes konstruktiv zu definieren, wenden wir uns jetzt der Kombination solcher Faltungen zu. Dabei möchten wir die Beziehung zwischen dem Start- und dem Endzustand beschreiben, ohne darauf einzugehen, dass zwischendurch mehrere andere Zustände durchlaufen wurden. Wir möchten ein Konzept erhalten, das analog zur Teilüberlagerung aus Kapitel 2.2 ist.

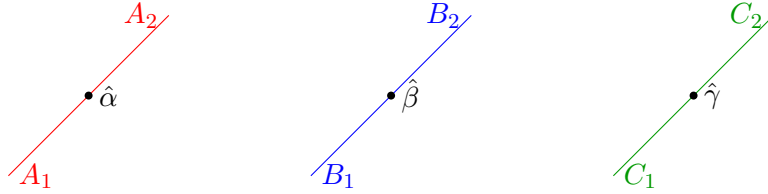
Dazu untersuchen wir, welche Eigenschaften die Faltungen aus Abschnitt 3.2 zwischen ihrem Start- und Endzustand induzieren. Bei Faltungen von Grad n wird die Menge der

Randstücke echt verringert. Bei Faltungen von Grad $n - 1$, in denen ein $[k] \in \mathcal{K}_{n-1}$ auf ein $[l]$ abgebildet wird, müssen wir den ursprünglichen Fächer im Fächer des gefalteten Zustands wiederfinden können. Sobald wir eine simpliziale Orientierung festlegen, ist das Redukt der zyklischen Reihung um $[l]$ auf die Corona von $[k]$ also gleich der zyklischer Reihung um $[k]$.

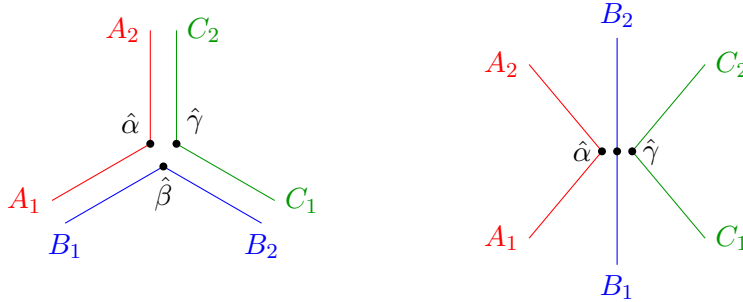
Bei einer Fächersumme müssen wir zwei Dinge beachten: Zum einen bleiben die Blöcke, die einer Flächenklasse entsprechen, fest (d. h. keine Elemente werden dazwischen geschoben). Zum anderen dürfen keine Überschneidungen auftreten. Damit ist gemeint, dass aus den zyklischen Reihungen (A, B) und (C, D) nicht (A, C, B, D) werden kann.

Die erste Eigenschaft ist dadurch garantiert, dass es die Standardränder ändern würde, Elemente zwischen eine Flächenklasse zu schieben. Die zweite Eigenschaft ist noch nicht in einer Form, die durch wiederholte Faltung stabil bleibt. Wir betrachten daher in Beispiel 3.46, was bei wiederholter Konfluenz von zyklischen Reihungen passiert (wir übernehmen dabei die zyklischen Reihungen aus Beispiel 3.37).

Beispiel 3.46. Wir betrachten zyklische Reihungen (in Zykelschreibweise) $\hat{\alpha} = (A_1, A_2)$, $\hat{\beta} = (B_1, B_2)$ und $\hat{\gamma} = (C_1, C_2)$, die wir wie folgt visualisieren können:



Die Konfluenzen $\hat{\gamma} \xrightarrow[B_2]{C_1} \hat{\beta} \xrightarrow[A_1]{B_1} \hat{\alpha}$ und $\hat{\alpha} \xrightarrow[B_2]{A_1} \hat{\beta} \xrightarrow[C_2]{B_1} \hat{\gamma}$ liefern die Zyklen $(A_1, B_1, B_2, C_1, C_2, A_2)$ und $(A_1, B_1, C_1, C_2, B_2, A_2)$ und lassen sich wie folgt visualisieren:



Es fällt auf, dass in beiden Graphiken keine Überschneidungen der verschiedenen zyklischen Reihungen auftreten. Wenn wir das durch die 6-Zykel ausdrücken wollen, müssen wir ausdrücken, dass keine Elemente der ursprünglichen zyklischen Reihungen in alternierender Reihenfolge auftauchen. Betrachten wir exemplarisch $\hat{\alpha}$ und $\hat{\beta}$ Im Zykel $(A_1, B_1, C_1, C_2, B_2, A_2)$ erhalten wir startend von A_1 die Werte B_1 , dann B_2 und schließlich A_2 . Wären die letzten beiden in anderer Reihenfolge erschienen (würden wir also zwischen Elementen in A und B alternieren), würden sich $\hat{\alpha}$ und $\hat{\beta}$ in der Graphik schneiden.

In Beispiel 3.46 haben wir die Überschneidungsfreiheit so beschrieben, dass sie auch für mehrfache Faltungen anwendbar ist. Wir definieren dieses Kriterium allgemein:

Definition 3.47 (überschneidungsfrei). Sei $\nu : M \rightarrow M$ eine zyklische Reihung. Eine Familie disjunkter Teilmengen $\{M_i\}_{i \in I}$ von M heißt **überschneidungsfrei (bzgl. ν)**, falls gilt:

Es gibt kein $f \in M$, sowie natürliche Zahlen $0 < x < y < z < |M|$, sodass $\{f, \nu^y(f)\}$ und $\{\nu^x(f), \nu^z(f)\}$ in verschiedenen M_i liegen.

Nachdem wir das Konzept der Überschneidungsfreiheit allgemein definiert haben, müssen wir nachweisen, dass es unter Fächersummen erhalten bleibt. Der formale Nachweis dieser Aussage ist relativ simpel, aber mit vielen Indizes behaftet:

Lemma 3.48. Seien σ_1, σ_2 zyklische Reihungen über den endlichen, disjunkten Mengen A_1, A_2 . Weiter seien $\{M_s^1 | 1 \leq s \leq t\}$ und $\{M_u^2 | 1 \leq u \leq v\}$ überschneidungsfrei bezüglich σ_1 und σ_2 .

Dann ist $\{M_s^1 | 1 \leq s \leq t\} \cup \{M_u^2 | 1 \leq u \leq v\}$ überschneidungsfrei bezüglich jeder Konfluenz von σ_1 und σ_2 .

Beweis. Dass $\{M_s^1 | 1 \leq s \leq t\} \cup \{M_u^2 | 1 \leq u \leq v\}$ aus disjunkten Teilmengen von $A_1 \uplus A_2$ besteht, ist klar.

Wir betrachten eine Zykelarstellung einer Konfluenz (vgl. Definition 3.35) von σ_1 und σ_2 , in der die Elemente aus A_1 und A_2 separiert sind, d. h.

$$(a_1, \dots, a_{m_1}, a_{m_1+1}, \dots, a_{m_1+m_2}),$$

mit $\{a_1, \dots, a_{m_1}\} = A_1$ und $\{a_{m_1+1}, \dots, a_{m_1+m_2}\} = A_2$.

Angenommen, es gebe eine Überschneidung, also $i < j < k < l$, sodass $\{a_i, a_k\}$ und $\{a_j, a_l\}$ in verschiedenen Mengen der Partition liegen. Da A_1 und A_2 disjunkt sind, muss entweder $k \leq m_1$ oder $i > m_1$ gelten (da sonst $\{a_i, a_k\}$ in keiner Menge der Partition liegen könnte). Ohne Beschränkung der Allgemeinheit nehmen wir $k \leq m_1$ an. Dann folgt $j < m_1$ und demnach muss auch $l \leq m_1$ gelten (selbes Argument für $\{a_j, a_l\}$). Dann könnte aber $\{M_s^1 | 1 \leq s \leq t\}$ nicht überschneidungsfrei bezüglich σ_1 sein. Aus diesem Widerspruch folgt die Behauptung. $\square$

Nachdem wir in Lemma 3.48 nachgewiesen haben, dass Überschneidungsfreiheit ein zentrales Merkmal von Fächersummen ist, erwähnen wir eine Möglichkeit, diese einfacher nachzuweisen. Wir verwenden dabei die Redukte aus Definition 3.13, um die Überschneidungsfreiheit auf eine möglichst kleine zyklische Reihung zurückzuführen.

Bemerkung 3.49. Sei $\nu : M \rightarrow M$ eine zyklische Reihung und $\{M_i\}_{i \in I}$ sei eine Familie disjunkter Teilmengen von M . Dann ist $\{M_i\}_{i \in I}$ genau dann überschneidungsfrei bezüglich ν , falls für je zwei M_i, M_j (mit $i, j \in I$ und $i \neq j$) gilt:

- $\{M_i, M_j\}$ ist überschneidungsfrei bezüglich des Redukts $\nu|_{M_i \uplus M_j \rightarrow M_i \uplus M_j}$.

Nachdem wir den Begriff der Überschneidungsfreiheit hinreichend lange untersucht haben, wenden wir uns wieder dem Ziel dieses Abschnitts zu: Das Konzept der Teilüberlagerung auf die Situation der n -Faltkomplexe zu verallgemeinern. Dazu erweitern wir eine Teilüberlagerung um alle Merkmale, die wir bislang gesammelt haben. Es wird sich in Lemma 3.56 herausstellen, dass diese Merkmale hinreichend für die Äquivalenz zwischen dem neuen Begriff der Teilfaltung und den primitiven Faltungen aus Abschnitt 3.2 ist.

Definition 3.50 (Teilfaltung). *Seien $\mathcal{F}, \mathcal{G}$ zwei n -Proto-Faltkomplexe. Dann heißt $\mathcal{F}$ **Teilfaltung** von $\mathcal{G}$, falls gilt*

1. $\mathcal{F}$ ist Teilüberlagerung von $\mathcal{G}$.
2. Jedes Randstück von $\mathcal{G}$ ist auch eines von $\mathcal{F}$.
3. Für $l \in \mathcal{K}_{n-1}$ sind die Coronae $\{\text{Cor}_{\mathcal{F}}([k]_{\mathcal{F}}) | k \in [l]_{\mathcal{G}}\}$ überschneidungsfrei bezüglich der zyklischen Reihungen $\nu_{[l]_{\mathcal{G}}}(a)$ für eine (und damit jede) simpliziale Orientierung $a \in \text{SimplOrient}(\mathcal{G}_{\preccurlyeq[l]})$.
4. Für $k \in \mathcal{K}_{n-1}$ und eine simpliziale Orientierung $a \in \text{SimplOrient}(\mathcal{F}_{\preccurlyeq k})$ gilt: Das Redukt von $\nu_{[l]_{\mathcal{G}}}([a]_{\mathcal{G}})$ auf die Corona $\text{Cor}_{\mathcal{F}}([k]_{\mathcal{F}})$ ist gleich $\nu_{[l]_{\mathcal{F}}}([a]_{\mathcal{F}})$.

Damit das Konzept der Teilfaltung eine Relation zwischen Start- und Endpunkt einer Faltung darstellen kann, muss es transitiv sein:

Lemma 3.51. *Seien $\mathcal{F}, \mathcal{G}, \mathcal{H}$ drei n -Proto-Faltkomplexe, sodass $\mathcal{F}$ eine Teilfaltung von $\mathcal{G}$ und $\mathcal{G}$ eine Teilfaltung von $\mathcal{H}$ ist. Dann ist $\mathcal{F}$ eine Teilfaltung von $\mathcal{H}$.*

Beweis. Wir prüfen die Eigenschaften aus Definition 3.50 einzeln nach. Dass die Teilüberlagerungs- und Randstückbedingungen transitiv sind, ist klar. Dass eine doppelte Reduktbildung gleich einer einfachen Reduktbildung ist, ist die Aussage von Lemma 3.15.

Wir müssen also noch nachweisen, dass die Überschneidungsfreiheit transitiv ist.

Seien $[k_1], [k_2] \in \mathcal{K}_{n-1} / \sim_{\mathcal{F}}$ verschieden mit $k_1 \sim_{\mathcal{H}} k_2$. Wir wollen zeigen, dass die Coronae von $\text{Cor}_{\mathcal{F}}([k_1])$ und $\text{Cor}_{\mathcal{F}}([k_2])$ überschneidungsfrei bezüglich $\nu_{[k_1]}^{\mathcal{H}}(a)$ für ein $a \in \text{SimplOrient}(\mathcal{H}_{\preccurlyeq[k_1]})$ sind. Fixiere eine solche simpliziale Orientierung a für den Rest des Beweises, sodass wir uns nur um die zyklischen Reihungen kümmern müssen (die Orientierung a induziert auch simpliziale Orientierungen für die anderen betrachteten Fächer).

Wir verwenden die Charakterisierung aus Bemerkung 3.49 und unterscheiden zwei Fälle:

- $k_1 \sim_{\mathcal{G}} k_2$:

Da das Redukt von $\nu_{[k_1]}^{\mathcal{H}}(a)$ auf $\text{Cor}_{\mathcal{G}}([k_1])$ gleich $\nu_{[k_1]}^{\mathcal{G}}(a)$ ist und die Coronae $\text{Cor}_{\mathcal{F}}([k_1])$ und $\text{Cor}_{\mathcal{F}}([k_2])$ überschneidungsfrei sind, sind sie es (aufgrund von Bemerkung 3.49 und Lemma 3.15) auch in $\nu_{[k_1]}^{\mathcal{H}}$.

- $k_1 \not\sim_{\mathcal{G}} k_2$:

In diesem Fall ist jede Corona $\text{Cor}_{\mathcal{F}}([k_i])$ in $\text{Cor}_{\mathcal{G}}([k_i])$ enthalten. Da diese dann überschneidungsfrei bzgl. $\nu_{[k_1]}^{\mathcal{H}}$ sind, gilt das auch für ihre Teilmengen. $\square$

Wir haben bislang den Begriff der Teilfaltung definiert und nachgewiesen, dass dieser transitiv ist und von primitiven Faltungen erfüllt wird. Wir müssen aber auch die umgekehrte Richtung nachweisen, d. h. dass sich jede Teilfaltung in mehrere Einzelfaltungen zerlegen lässt. Dazu bauen wir auf Lemma 2.33 auf, in dem Teilüberlagerungen in primitive Erweiterungen zerlegt werden. Da wir aber gewissen Einschränkungen beim Falten unterliegen, können wir die Folge aus dem Lemma nicht direkt übernehmen. Wir führen den Beweis aber ähnlich. Zu zwei endlichen n -Faltkomplexen $\mathcal{F}$ und $\mathcal{G}$, wobei $\mathcal{F}$ Teilfaltung von $\mathcal{G}$ ist, müssen wir eine primitive Faltung $\mathcal{F}^*$ von $\mathcal{F}$ so konstruieren, dass $\mathcal{F}^*$ eine Teilfaltung von $\mathcal{G}$ ist (der Rest folgt mit Induktion und der Endlichkeit der Komplexe). Für Grade kleiner als $n - 1$ gibt es keine Komplikationen.

Bemerkung 3.52. *Seien $\mathcal{F}, \mathcal{G}$ zwei n -Proto-Faltkomplexe, wobei $\mathcal{F}$ Teilfaltung von $\mathcal{G}$ ist. Es gebe $p_1, p_2 \in \mathcal{K}_i$ mit $i \leq n - 2$, sodass gilt:*

- $p_1 \not\sim_{\mathcal{F}} p_2$
- $p_1 \sim_{\mathcal{G}} p_2$
- Die Vereinigung von $[p_1]$ und $[p_2]$ liefert eine primitive Erweiterung von $\mathcal{F}$

Dann ist die primitive Faltung von $\mathcal{F}$ an p_1 und p_2 eine Teilfaltung von $\mathcal{G}$.

Bei der Identifikation der Kanten (Grad $n - 1$) müssen wir beachten, dass wir die Fächer in der richtigen Reihenfolge zusammenfügen. Dazu nehmen wir zwei Fächer $\nu_{[k_1]}^{\mathcal{F}}$ und $\nu_{[k_2]}^{\mathcal{F}}$ mit $k_1, k_2 \in [l]_{\sim_{\mathcal{G}}}$, deren Coronae in $\nu_{[l]}^{\mathcal{G}}$ benachbart sind, und fügen diese zusammen. Dabei müssen wir sicherstellen, dass diese primitive Faltung noch immer Teilfaltung von $\mathcal{G}$ ist. Die meisten Bedingungen dazu (vgl. Definition 3.50) machen wenig Arbeit, aber bei der Überschneidungsfreiheit müssen wir ein wenig arbeiten. Wir müssen nachweisen, dass die Kombination zweier benachbarter Fächer keine Überschneidungen erzeugt. Dazu verwenden wir das folgende Lemma:

Lemma 3.53. *Die Mengenfamilie $\{M_i | 1 \leq i \leq k\}$ mit $k \geq 2$ sei überschneidungsfrei bezüglich der zyklischen Reihung ν . Es gebe ein $m \in M_1$ mit $\nu(m) \in M_2$.*

Dann ist $\{M_1 \uplus M_2\} \uplus \{M_i | 3 \leq i \leq k\}$ überschneidungsfrei bezüglich ν .

Beweis. Da die Redukte auf alle Mengenpaare außer $M_1 \uplus M_2$ sich nicht verändern, müssen wir nach Bemerkung 3.49 nur die Paare betrachten, in denen $M_1 \uplus M_2$ vorkommt. Wir untersuchen ohne Einschränkung nur $M_1 \uplus M_2$ und M_3 .

Gäbe es eine Überschneidung zwischen diesen beiden Mengen, so gäbe es ein $a \in M_1$ und $0 < \alpha < \beta < \gamma$, sodass $\nu^\beta(a) \in M_2$ und $\{\nu^\alpha(a), \nu^\gamma(a)\} \subseteq M_3$ gilt (a und $\nu^\beta(a)$ können nicht im selben M_i liegen, da wir ansonsten eine Überschneidung in $\{M_1, M_2, M_3\}$ gehabt hätten).

Wir untersuchen, wo das m aus der Voraussetzung liegen kann. Es gibt kein $\alpha < j < \gamma$ mit $\nu^j(a) = m$, da sonst $(a, \nu^\alpha(a), \nu^j(a), \nu^\gamma(a))$ eine Überschneidung von M_1 und M_3 wäre. Damit muss $j < \alpha$ oder $j > \gamma$ gelten. Wegen $\nu^{j+1}(a) \in M_2$ bilden $\{\nu^{j+1}(a), \nu^\beta(a)\} \subseteq M_2$ und $\{\nu^\alpha(a), \nu^\gamma(a)\} \subseteq M_3$ eine Überschneidung. Daher gibt es keinen Platz für dieses m , womit die Annahme der Überschneidung widerlegt ist. $\square$

Mit Lemma 3.53 als Vorbereitung, können wir leicht eine primitive Faltung von Grad $n - 1$ konstruieren, die Teilfaltung von $\mathcal{G}$ bleibt.

Lemma 3.54. *Seien $\mathcal{F}, \mathcal{G}$ zwei n -Proto-Faltkomplexe, sodass $\mathcal{F}$ Teilfaltung von $\mathcal{G}$ ist. Es gebe $k_1, k_2 \in \mathcal{K}_{n-1}$, sodass gilt:*

- $k_1 \not\sim_{\mathcal{F}} k_2$
- $k_1, k_2 \in [l]_{\sim \mathcal{G}} \in \mathcal{K}_{n-1} / \sim^{\mathcal{G}}$
- Vereinigung von $[k_1]$ und $[k_2]$ liefert eine primitive Erweiterung von $\mathcal{F}$

Bezeichne mit M_i die Corona von $\nu_{[k_i]}^{\mathcal{F}}$. Wir nehmen weiter an, dass es $(f, \varepsilon_f) \in M_1 \times \{\pm 1\}$ und $(g, \varepsilon_g) \in M_2 \times \{\pm 1\}$ gibt, sodass (f, ε_f) und (g, ε_g) in $\nu_{[l]}^{\mathcal{G}}$ komplementär sind.

Dann ist die primitive Faltung von $\mathcal{F}$, die k_1 und k_2 an (f, ε_f) und (g, ε_g) zusammenfaltet, eine Teilfaltung von $\mathcal{G}$.

Beweis. Nach Definition 3.41 ist diese Konstruktion wohldefiniert. Wir prüfen die Eigenschaften aus der Definition 3.50 der Teilfaltung einzeln nach.

1. Da $k_1 \sim^{\mathcal{G}} k_2$ gilt, handelt es sich bei der primitiven Faltung von $\mathcal{F}$ um eine Teilüberlagerung von $\mathcal{G}$.
2. Die Randstücke werden nicht berührt, folglich ist auch diese Bedingung erfüllt.
3. Die Überschneidungsfreiheit ist durch Lemma 3.53 gedeckt.
4. Seien $a_1 \in \text{SimplOrient}(\mathcal{F}_{\preccurlyeq k_1})$ und $a_2 \in \text{SimplOrient}(\mathcal{F}_{\preccurlyeq k_2})$ simpliziale Orientierungen, die $[a_1]_{\mathcal{G}} = [a_2]_{\mathcal{G}}$ erfüllen. Definiere die zyklischen Reihungen

$$\begin{aligned}\hat{\nu}_1^{\mathcal{F}} &:= \nu_{[k_1]}^{\mathcal{F}}([a_1]_{\mathcal{F}}) \text{ mit Corona } M_1 \\ \hat{\nu}_2^{\mathcal{F}} &:= \nu_{[k_2]}^{\mathcal{F}}([a_2]_{\mathcal{F}}) \text{ mit Corona } M_2 \\ \hat{\nu}^{\mathcal{G}} &:= \nu_{[l]}^{\mathcal{G}}([a_1]_{\mathcal{G}}),\end{aligned}$$

dann gilt nach Voraussetzung $\hat{\nu}_{|M_1 \rightarrow M_1}^{\mathcal{G}} = \hat{\nu}_1^{\mathcal{F}}$ und $\hat{\nu}_{|M_2 \rightarrow M_2}^{\mathcal{G}} = \hat{\nu}_2^{\mathcal{F}}$. Da (f, ε_f) und (g, ε_g) komplementär sind, nehmen wir ohne Einschränkung an, dass $\hat{\nu}^{\mathcal{G}}(f) = g$ gilt (ansonsten vertausche deren Rollen). Da $\{M_1, M_2\}$ überschneidungsfrei bezüglich $\nu_{[l]}^{\mathcal{G}}$ ist, ist dadurch schon das Redukt $\hat{\nu}_{|M_1 \sqcup M_2 \rightarrow M_1 \sqcup M_2}^{\mathcal{G}}$ festgelegt. Dieses ist gleich der Konfluenz von $\hat{\nu}_2^{\mathcal{F}}$ und $\hat{\nu}_1^{\mathcal{F}}$ an g und f , die auch bei der Fächersumme zum Einsatz kommt. Damit folgt die behauptete Gleichheit. $\square$

Nachdem wir in Bemerkung 3.52 und Lemma 3.54 die primitiven Faltungen aller Grade ungleich n untersucht haben, wenden wir uns jetzt der Identifikationen von Flächen zu (oder allgemeiner Grad n). Hier müssen wir nur darauf achten, dass wir die richtigen Oberflächen zusammenfalten.

Lemma 3.55. *Seien $\mathcal{F}, \mathcal{G}$ zwei n -Proto-Faltkomplexe, sodass $\mathcal{F}$ Teilfaltung von $\mathcal{G}$ ist. Es gebe $f, g \in \mathcal{K}_n$, sodass gilt:*

- $g \not\sim_{\mathcal{F}} f$
- $f \sim_{\mathcal{G}} g$
- Vereinigung von $[f]$ und $[g]$ liefert eine primitive Erweiterung von $\mathcal{F}$

Seien (f, ε_f) und (g, ε_g) (mit $\varepsilon_f, \varepsilon_g \in \{\pm 1\}$) komplementäre Oberflächen¹², die keine Randstücke von $\mathcal{G}$ sind.

Dann ist die primitive Faltung von $\mathcal{F}$, die f und g an (f, ε_f) und (g, ε_g) zusammenfaltet, eine Teilfaltung von $\mathcal{G}$.

Beweis. Nach Definition 3.31 ist diese Konstruktion wohldefiniert. Wir prüfen die Eigenschaften aus der Definition 3.50 der Teilfaltung einzeln nach.

1. Da $f \sim_{\mathcal{G}} g$ gilt, handelt es sich bei der primitiven Faltung von $\mathcal{F}$ um eine Teilüberlagerung von $\mathcal{G}$.
2. Es werden nur Randstücke entfernt, die nicht schon in $\mathcal{G}$ vorkamen, folglich ist auch diese Bedingung erfüllt.
3. Da keine Fächer verändert werden, ist die Überschneidungsfreiheit garantiert.
4. Da keine Fächer verändert werden, ist auch die Redukt-Bedingung trivial. $\square$

Wir fassen unsere Erkenntnisse zusammen.

Lemma 3.56. *Seien $\mathcal{F}, \mathcal{G}$ zwei endliche n -Proto-Faltkomplexe und $\mathcal{F}$ sei eine Teilfaltung von $\mathcal{G}$. Dann gibt es eine Folge von n -Proto-Faltkomplexen*

$$\mathcal{F} = \mathcal{F}_0, \mathcal{F}_1, \dots, \mathcal{F}_k = \mathcal{G},$$

sodass $\mathcal{F}_i$ für alle $1 \leq i \leq k$ eine primitive Faltung von $\mathcal{F}_{i-1}$ ist.

Falls $\mathcal{G}$ ein n -Faltkomplex ist, sind alle $\mathcal{F}_i$ (mit $0 \leq i \leq k$) auch n -Faltkomplexe.

Beweis. Zeige zuerst, dass es eine primitive Faltung von $\mathcal{F}$ gibt, die Teilfaltung von $\mathcal{G}$ ist. Sei dazu $0 \leq i \leq n$ minimal, sodass es $x, y \in \mathcal{K}_i$ mit $[x]_{\sim_{\mathcal{F}}} \uplus [y]_{\sim_{\mathcal{F}}} \subseteq [x]_{\sim_{\mathcal{G}}}$ gibt. Für $i < n - 1$ folgt die Aussage aus Bemerkung 3.52.

Für $i = n - 1$ müssen wir die Voraussetzungen aus Lemma 3.54 prüfen. Wir wissen, dass $[x]_{\sim_{\mathcal{G}}}$ in mehrere $\sim_{\mathcal{F}}$ -Klassen zerfällt. Die Coronae dieser Klassen sind in der Corona $\text{Cor}_{\mathcal{G}}([x])$ enthalten, da $\mathcal{F}$ Teilfaltung von $\mathcal{G}$ ist. Es gibt daher zwei komplementäre Oberflächen aus verschiedenen Coronae, wie in Lemma 3.54 gefordert. Die Vereinigung der entsprechenden Äquivalenzklassen induziert also eine primitive Faltung, die Teilfaltung von $\mathcal{G}$ ist.

¹²Diese Bezeichnung ist eindeutig, da aufgrund der Primitivitätsannahme alle Kanten diese beiden Flächen enthalten.

Für $i = n$ müssen wir die Voraussetzungen von Lemma 3.55 prüfen. Die Klasse $[x]_{\sim \mathcal{G}}$ zerfällt in mehrere $\sim^{\mathcal{F}}$ -Klassen. In $[x]_{\mathcal{G}}$ gibt es komplementäre Oberflächen aus verschiedenen $\sim^{\mathcal{F}}$ -Klassen. Da $\mathcal{F}$ eine Teilfaltung von $\mathcal{G}$ ist und alle Fächer der beiden n -Proto-Faltkomplexe übereinstimmen, sind diese beiden Oberflächen auch in $\mathcal{F}$ komplementär. Die Vereinigung der entsprechenden Äquivalenzklassen liefert als nach Lemma 3.55 eine primitive Faltung, die Teilfaltung von $\mathcal{G}$ ist.

Da die Komplexe endlich sind, lässt sich die obige Konstruktion nur endlich oft anwenden. Sobald es kein solches i mehr gibt, gilt bereits $\mathcal{F} = \mathcal{G}$. Damit ist die Behauptung gezeigt. $\square$

Damit haben wir jede Teilfaltung in primitive Faltungen zerlegt. Folglich sind die Begriffe von Teilfaltung und primitiven Faltungen äquivalent (im endlichen Fall).

3.4 Entfaltungen

Nachdem wir die Faltung von n -Faltkomplexen in den Abschnitten 3.2 und 3.3 ausführlich untersucht haben, wenden wir uns in diesem Abschnitt deren Entfaltung zu. Wir erweitern dazu das Konzept der primitiven Spaltung aus Abschnitt 2.4 von einer Teilüberlagerung auf eine Teilfaltung:

Definition 3.57 (primitive Entfaltung). Seien $\mathcal{F}$ ein n -Faltkomplex und $f \in \bigcup_{j=0}^n \mathcal{K}_j$.

Der n -Faltkomplex $\mathcal{E}$ heißt **primitive Entfaltung (von $\mathcal{F}$) an $[f]$** , falls gilt:

- $\mathcal{E}$ ist eine Teilfaltung von $\mathcal{F}$.
- $\mathcal{E}$ ist eine primitive Spaltung von $\mathcal{F}$ an $[f]$ (als n -Überlagerungskomplexe aufgefasst).

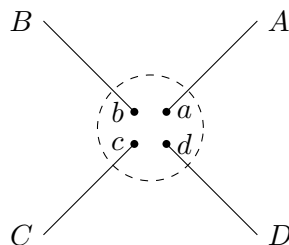
In den meisten Fällen ändert sich nur wenig durch diese Modifikation. Da Faltkomplexe nur bei $\mathcal{K}_{n-1}$ und $\mathcal{K}_n$ ein von Überlagerungskomplexen abweichendes Verhalten zeigen, können wir Satz 2.49 unmittelbar auf primitive Entfaltungen von allen $\mathcal{K}_j$ mit $j < n - 1$ übertragen.

Folgerung 3.58. Sei $\mathcal{F}$ ein n -Proto-Faltkomplex und sei $f \in \mathcal{K}_j$ (mit $0 \leq j < n - 1$).

Eine Partition $\bigsqcup_{s=1}^m M_s$ von $[f]_{\sim}$ ist genau dann Spaltungspartition einer primitiven Entfaltung von $[f]$, wenn sie mit der Bindungsrelation verträglich ist.

Wir betrachten die Klassen $\mathcal{K}_{n-1}/\sim$ als nächstes. Hier ändern die zyklischen Reihungen die Situation. Wenn wir die Definition 3.50 der Teilfaltung betrachten, erkennen wir eine mögliche Komplikation: Die Coronae der Fächer von $\mathcal{E}$ müssen überschneidungsfrei bezüglich jeder zyklische Reihung von $\mathcal{F}$ sein. Dass dies nicht trivial ist, zeigt das folgende Beispiel.

Beispiel 3.59. Betrachten wir den folgenden Ausschnitt eines 1-Faltkomplexes (bei dem wir auf die Kennzeichnung der Kantenoberflächen verzichten haben):



Wir haben eine Punktäquivalenzklasse $\{a, b, c, d\}$ vorliegen, deren zyklische Reihung gleich (A, B, C, D) ist. Wenn alle Kantenklassen einelementig sind, erkennt man aus der Graphik, dass die Bindungsrelation auf $\{a, b, c, d\}$ trivial ist. Nach Satz 2.49 ist die Zerlegung in $\{a, c\}$ und $\{b, d\}$ eine primitive Spaltung. Diese Partition der Punkte induziert die Kantenpartition $\{A, C\} \uplus \{B, D\}$, die nicht überschneidungsfrei bezüglich (A, B, C, D) ist. Demnach handelt es sich bei dieser primitiven Spaltung nicht um eine primitive Entfaltung.

Das Beispiel 3.59 demonstriert, dass die Überschneidungsfreiheit der induzierten Kantenpartition notwendig für das Vorliegen einer primitiven Entfaltung ist. Wir müssen also klären, was mit dem „Induzieren einer Partition“ gemeint ist.

Definition 3.60 (induzierte Fächerpartition). Sei $\mathcal{F}$ ein n -Proto-Faltkomplex und $k \in \mathcal{K}_{n-1}$. Sei $\uplus_{j=1}^m M_j$ eine Partition von $[k]$. Dann definiert $\uplus_{j=1}^m N_j$ mit

$$N_j := \{f \in \mathcal{K}_n \mid \exists l \in M_j : l \prec f\}$$

eine Partition von $\text{Cor}([k])$, die **induzierte Fächerpartition**.

Wohldefiniertheit. Wir müssen zeigen, dass es sich tatsächlich um eine Partition handelt.

- Für $f \in \text{Cor}([k])$ gibt es nach Definition 3.6 ein $l \in \mathcal{K}_{n-1}$ und ein $g \in \mathcal{K}_n$, sodass $k \sim l \prec g \sim f$ erfüllt ist. Aufgrund der Abwärtskompatibilität von $\sim$ gibt es ein $\bar{l} \sim l$ mit $\bar{l} \prec f$.

Da $\uplus_{j=1}^m M_j$ eine Partition von $[k]$ ist, liegt $\bar{l}$ in einem M_j . Dann folgt $f \in N_j$.

- Betrachten wir zwei N_i, N_j . Zu einem $f \in N_i \cap N_j$ gibt es $l_i \in M_i$ und $l_j \in M_j$ mit $l_i \prec f$ und $l_j \prec f$. Wegen $l_i \sim l_j$ folgt $l_i = l_j$ aus der Irreduzibilität von $\sim$. Da $\uplus_{j=1}^m M_j$ eine Partition ist, folgt $i = j$.

Folglich ist $\uplus_{j=1}^m N_j$ eine Partition von $\text{Cor}([k])$. □

Damit lässt sich das Analog von Satz 2.49 für $[k] \in \mathcal{K}_{n-1} / \sim$ formulieren.

Lemma 3.61. Sei $\mathcal{F}$ ein n -Proto-Faltkomplex und sei $k \in \mathcal{K}_{n-1}$.

Eine Partition $\uplus_{j=1}^m M_j$ von $[k]$ ist genau dann Spaltungspartition einer primitiven Entfaltung an $[k]$, wenn gilt:

- Sie ist mit der Bindungsrelation verträglich.
- Die induzierte Fächerpartition ist überschneidungsfrei bzgl. einer (und damit jeder) zyklischen Reihung des Fächers $\nu_{[k]}^{\mathcal{F}}$.

Beweis. • Sei $\mathcal{E}$ eine primitive Entfaltung von $\mathcal{F}$ an $[k]$ und $\uplus_{j=1}^m M_j$ sei die Spaltungspartition von $\mathcal{E}$. Nun gilt:

- Da $\mathcal{E}$ primitive Spaltung von $\mathcal{F}$ ist, ist $\uplus_{j=1}^m M_j$ nach Satz 2.49 mit der Bindungsrelation verträglich.

- Die zyklischen Reihungen von $\mathcal{E}$ auf M_j sind auf der induzierten Fächerpartition von $\biguplus_{j=1}^m M_j$ definiert. Da $\mathcal{E}$ Teilfaltung von $\mathcal{F}$ ist, liefern diese also eine überschneidungsfreie Partition bzgl. der zyklischen Reihungen von $\nu_{[k]}^{\mathcal{F}}$.
- Sei eine solche Partition $\biguplus_{j=1}^m M_j$ von $[k]$ gegeben. Nach Satz 2.49 gibt es einen n -Überlagerungskomplex $\mathcal{E}$, der primitive Spaltung von $\mathcal{F}$ an $[k]$ ist. Wir zeigen, dass sich $\mathcal{E}$ so zu einem n -Proto-Faltkomplex fortsetzen lässt, dass dieser Teilfaltung von $\mathcal{F}$ wird. Dies funktioniert wie folgt:

- Alle Randstücke bleiben gleich, d. h. $\partial_{[f]}^{\mathcal{E}} := \partial_{[f]}^{\mathcal{F}}$ für jedes $[f] \in \mathcal{K}_n / \sim$.
- Für $l \in \mathcal{K}_{n-1}$ mit $l \not\sim^{\mathcal{F}} k$ gilt $\nu_{[l]}^{\mathcal{E}} = \nu_{[l]}^{\mathcal{F}}$.

Für die Äquivalenzklasse M_j müssen wir aus einer zyklischen Reihung $\hat{\nu}_{[k]}^{\mathcal{F}}$ eine über M_j konstruieren. Wegen Bedingung 4 in Definition 3.50 handelt es sich dabei um das Redukt auf M_j . Jede simpliziale Orientierung von $\mathcal{F}_{\preccurlyeq[k]}$ induziert dabei eine simpliziale Orientierung auf $\mathcal{E}_{\preccurlyeq[l]}$ für jedes $[l]_{\sim \mathcal{E}} \subseteq [k]_{\sim \mathcal{F}}$.

- Da keine Klasse aus $\mathcal{K}_n / \sim$ aufgespalten wurde, haben sich die Standardränder nicht verändert (da M_j mit der Bindungsrelation verträglich ist).

Um nachzuweisen, dass $\mathcal{E}$ Teilfaltung von $\mathcal{F}$ ist, genügt es nachzuweisen, dass $\{M_j | 1 \leq j \leq m\}$ überschneidungsfrei bzgl. jeder zyklischen Reihung von $\nu_{[k]}^{\mathcal{F}}$ ist. Das ist aber durch die Einschränkung an $\biguplus_{j=1}^m M_j$ garantiert. $\square$

Damit haben wir die primitiven Entfaltungen aller Grade kleiner als n charakterisiert. Als nächstes betrachten wir primitive Entfaltungen an $[f] \in \mathcal{K}_n / \sim$. Eine solche Entfaltung verändert nur die Randstücke, nicht aber die Fächer. Damit wir einen n -Proto-Faltkomplex erhalten, müssen wir garantieren, dass jede Menge der Spaltungspartition maximal zwei Standardränder induziert. Da wir von primitiven Spaltungen sprechen, haben alle diese Flächen die selben Kanten (vor und nach der Entfaltung). Damit ist diese Bedingung nur von der Spaltungspartition selbst abhängig. Anschaulich ist die nötige Bedingung recht leicht zu formulieren: Wir betrachten die lineare Ordnung

$$f_1 < f_2 < \cdots < f_{m_k} \quad (m_k \in \mathbb{N})$$

auf der Flächenklasse $[f]$ (die wir aus den Fächern und Randstücken mittels Lemma 3.26 rekonstruieren können) und zerlegen diese in mehrere Teilordnungen

$$\begin{aligned} f_1 &< f_2 < \cdots < f_{m_1} \\ f_{m_1+1} &< f_{m_1+2} < \cdots < f_{m_2} \\ &\dots \\ f_{m_{k-1}+1} &< f_{m_{k-1}+2} < \cdots < f_{m_k}. \end{aligned}$$

Da wir diese linearen Ordnungen nicht direkt vorliegen haben, müssen wir die Zerlegung etwas anders formulieren. Dabei müssen wir insbesondere auf den Unterschied zwischen einer linearen und einer zyklischen Ordnung aufpassen. Im Ausgleich dafür brauchen wir uns nicht um die Bindungsrelation zu kümmern, da es keine höhere Dimensionen als n gibt.

Lemma 3.62. Sei $\mathcal{F}$ ein n -Faltkomplex und sei $f \in \mathcal{K}_n$. Sei a eine simpliziale Orientierung von $\mathcal{F}_{\preccurlyeq[k]}$ für ein $k \in \mathcal{K}_{n-1}$, das $[k] \prec [f]$ erfüllt. Wähle die Bezeichnung $\hat{\nu} := \nu_{[k]}(a)$.

Eine Partition $\bigsqcup_{j=1}^m M_j$ von $[f]$ ist genau dann Spaltungspartition einer primitiven Entfaltung an $[f]$, wenn gilt

- Zu jedem M_j gibt es ein $x \in M_j$ und ein $t \in \mathbb{N}$ mit

$$M_j = \{\hat{\nu}^i(x) | 0 \leq i \leq t\}.$$

Beweis. • Sei $\mathcal{E}$ eine primitive Entfaltung von $\mathcal{F}$ an $[f]$ und $\bigsqcup_{j=1}^m M_j$ sei die Spaltungspartition von $\mathcal{E}$. Nun gilt:

- Da $\mathcal{E}$ primitive Spaltung von $\mathcal{F}$ ist, ist $\bigsqcup_{j=1}^m M_j$ nach Satz 2.49 mit der Bindungsrelation verträglich.
- Angenommen, es gebe ein $x \in M_j$, sodass für eine zyklische Reihung $\hat{\nu}$ von $\mathcal{E}$ drei Werte $\alpha, \beta, \gamma \in \mathbb{N}$ mit $0 < \alpha < \beta < \gamma < |\text{Cor}_{\mathcal{E}}([f])|$ existieren, die $\{\hat{\nu}^\alpha(x), \hat{\nu}^\gamma(x)\} \cap M_j = \emptyset$ und $\hat{\nu}^\beta(x) \in M_j$ erfüllen.

Dann gibt es ein $0 \leq a < \alpha$ mit $\hat{\nu}^a(x) \in M_j$ und $\hat{\nu}^{a+1}(x) \notin M_j$, also liegt hier ein Randstück aus $\partial_{M_j}^{\mathcal{E}}$ vor. Analog findet man zwischen α und γ zwei weitere. Da es aber nur genau zwei davon gibt, ist das ein Widerspruch.

- Wähle x als Randstück so, dass $\hat{\nu}(x) \in M_j$. Gäbe es $0 < \alpha < \beta < |\text{Cor}_{\mathcal{E}}([f])|$ mit $\hat{\nu}^\alpha(x) \notin M_j$ und $\hat{\nu}^\beta(x) \in M_j$, dann gäbe es analog zum zweiten Fall mindestens zwei weitere Randstücke. Daher kann es keine Unterbrechung geben, bis das zweite Randstück erreicht ist.

- Sei eine solche Partition $\bigsqcup_{j=1}^m M_j$ von $[f]$ gegeben. Nach Satz 2.49 gibt es einen n -Überlagerungskomplex $\mathcal{E}$, der primitive Spaltung von $\mathcal{F}$ an $[f]$ ist. Wir zeigen, dass sich $\mathcal{E}$ so zu einem n -Faltkomplex fortsetzen lässt, dass dieser Teilfaltung von $\mathcal{F}$ wird.

- Alle Fächer bleiben gleich.
- Die Randstücke ∂_{M_j} sind gegeben durch die Elemente von M_j , die Standardränder sind, außer wenn die Partition einelementig ist. Dann bleiben sie gleich $\partial_{[f]}^{\mathcal{F}}$.

Von allen Bedingungen aus Definition 3.27 müssen wir nur überprüfen, ob es immer genau zwei Randflächen gibt (dann ist das so konstruierte $\mathcal{E}$ eine Teilfaltung von $\mathcal{F}$). Zu jedem M_j gibt es ein $x \in M_j$ und ein $t \in \mathbb{N}$ mit $M_j = \{\hat{\nu}^i(x) | 0 \leq i \leq t\}$ (die zyklischen Reihungen von $\mathcal{E}$ sind gleich denen von $\mathcal{F}$). Dann induzieren x und $\hat{\nu}^t(x)$ Standardränder (zu verschiedenen Richtungen), da $\hat{\nu}^{-1}(x) \notin M_j$ und $\hat{\nu}^{t+1}(x) \notin M_j$ gilt. Da es keine weiteren Standardränder gibt, folgt die Behauptung. $\square$

Wir können die Bedingungen dieses Lemmas mit Hilfe von Lemma 3.26 in Bezug auf eine feste lineare Ordnung formulieren.

Folgerung 3.63. Sei $\mathcal{F}$ ein n -Faltkomplex und sei $f \in \mathcal{K}_n$. Sei $(a_1, \dots, a_n, a_{n+1})$ eine simpliziale Orientierung von $\mathcal{F}_{\preccurlyeq f}$ und $k \in \mathcal{K}_{n-1}$ mit $k \prec f$, sodass $(a_1, \dots, a_n)$ eine simpliziale Orientierung von $\mathcal{F}_{\preccurlyeq k}$ ist. Setze $\hat{\nu} := \nu_{[k]}([a_1], \dots, [a_n])$.

Gemäß Lemma 3.26 können wir die Elemente von $[f]$ linear ordnen:

$$g, \hat{\nu}(g), \hat{\nu}^2(g), \dots, \hat{\nu}^{m-1}(g)$$

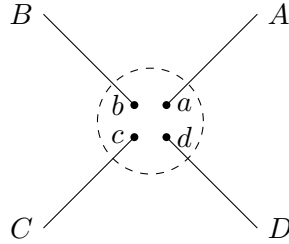
Jede lineare Ordnung $0 < t_1 < t_2 < \dots < t_l = m$ korrespondiert zu einer Spaltungspartition von $[f]$, nämlich

$$\begin{aligned} &\{\hat{\nu}^j(g) | 0 \leq j < t_1\} \\ &\{\hat{\nu}^j(g) | t_1 \leq j < t_2\} \\ &\vdots \\ &\{\hat{\nu}^j(g) | t_{l-1} \leq j < m\} \end{aligned}$$

Insgesamt beschreiben Folgerung 3.58 und die Lemmata 3.61 und 3.62 alle möglichen primitiven Entfaltungen eines n -Faltkomplexes. Als nächstes widmen wir uns der Struktur aller primitiven Entfaltungen.

Wir haben in Bemerkung 2.51 gesehen, dass die primitiven Spaltungen an $[f]$ eine Verbandsstruktur aufweisen, die durch Teilüberlagerung geordnet war. Aufgrund der Einschränkungen, die die Faltkomplex-Struktur vorgibt, überträgt sich das nicht vollständig auf primitive Entfaltungen, wie das nächste Beispiel zeigt:

Beispiel 3.64. Wir betrachten erneut die Situation aus Beispiel 3.59:



Die Partitionen $\{\{a\}, \{b, d\}, \{c\}\}$ und $\{\{a, c\}, \{b\}, \{d\}\}$ von $\{a, b, c, d\}$ induzieren überschneidungsfreie Fächerpartitionen (bzgl. (A, B, C, D)), aber ihr join (im Sinne von Bemerkung 2.51) ist $\{\{a, c\}, \{b, d\}\}$. Die induzierte Fächerpartition $\{\{A, C\}, \{B, D\}\}$ ist nicht überschneidungsfrei bezüglich (A, B, C, D) .

Allerdings ist der meet stets wohldefiniert. Wir benötigen also eine Abschwächung des Verbandsbegriffs. Es ist zwar theoretisch möglich, dass eine andere Definition des join zum Erfolg führt, aber dadurch verlören wir die angenehmen Berechnungseigenschaften.

Definition 3.65 (meet-Halbverband). Eine partielle Ordnung $(P, \leq)$ heißt **meet-Halbverband**, falls zu je zwei $a, b \in P$ ihr meet $a \wedge b$ (im Sinne der Definition 2.50 eines Verbandes) wohldefiniert ist.

Lemma 3.66. Sei $\mathcal{F}$ ein n -Faltkomplex und $x \in \mathcal{K}_i$ ($0 \leq i \leq n$). Die Menge der Spaltungspartitionen primitiver Entfaltungen von $[x]$ bildet für $i \neq n - 1$ einen Verband und für $i = n - 1$ einen meet-Halbverband, der wie in Bemerkung 2.51 definiert ist.

Beweis. Wir betrachten die drei Fälle $i < n - 1$, $i = n - 1$ und $i = n$ separat.

$i < n - 1$ Da nach Folgerung 3.58 primitive Entfaltungen und primitive Spaltungen zusammenfallen, ist nichts zu zeigen.

$i = n - 1$ Wir müssen nur zeigen, dass der meet von zwei überschneidungsfreien Partitionen selbst wieder überschneidungsfrei ist.

Seien dazu $\{M_1, \dots, M_m\}$ und $\{L_1, \dots, L_l\}$ überschneidungsfrei bezüglich der zyklischen Reihung $\hat{\nu}$ (auf N Elementen). Ihr meet sei die Partition $\{T_1, \dots, T_t\}$. Angenommen, die verschiedenen Mengen T_u und T_v überschneiden sich. Dann gibt es $x \in T_u$, sowie $0 < \alpha < \beta < \gamma < N$, sodass $\{x, \hat{\nu}^\beta(x)\} \subseteq T_u$ und $\{\hat{\nu}^\alpha(x), \hat{\nu}^\gamma(x)\} \subseteq T_v$ gelten.

Da $\{M_i | 1 \leq i \leq m\}$ überschneidungsfrei ist und $\{T_k\} \leq \{M_i\}$ gilt, gibt es ein $i \in \{1, \dots, m\}$, sodass $T_u \subseteq M_i$ und $T_v \subseteq M_i$ gelten. Analog gibt es ein $j \in \{1, \dots, l\}$, sodass $T_u \subseteq L_j$ und $T_v \subseteq L_j$ gelten.

Folglich gilt $T_u \subseteq M_i \cap L_j$, also nach Konstruktion des meet: $T_u = M_i \cap L_j$. Ebenso folgt $T_v = M_i \cap L_j$, woraus wir $T_u = T_v$ erhalten, im Widerspruch zur Annahme.

$i = n$ Um die Aussage nachzuweisen, verwenden wir die Darstellung aus Folgerung 3.63. Für die Spaltung von $[f] \in \mathcal{K}_n / \sim$ betrachten wir eine der zyklischen Reihungen $\hat{\nu}$, die $[f]$ enthält. Es gibt ein $x \in [f]$ mit

$$[f] = \{\hat{\nu}^t(x) | 0 \leq t < |[f]|\}.$$

Nach Folgerung 3.63 lässt sich eine primitive Entfaltung durch die Wahl einer Liste $0 < t_1 < t_2 < \dots < t_k \leq |[f]|$ festlegen. Seien nun zwei Spaltungspartitionen durch diese Listen gegeben. Ihr join ist durch den Schnitt dieser Listen gegeben, ihr meet durch deren Vereinigung. Damit liegt eine Verbandsstruktur vor. $\square$

Wir haben in Lemma 3.66 nachgewiesen, dass die Spaltungspartitionen an einer Äquivalenzklasse einen meet-Halbverband bilden. Wir möchten diese Struktur aber auch auf die primitiven Entfaltungen selbst übertragen. Dieser Übergang ist nicht so einfach wie der analoge Fall der n -Überlagerungskomplexe. Dort mussten wir nur die Äquivalenzrelation überprüfen, hier müssen wir auch auf die anderen Eigenschaften von Teilfaltungen aus Definition 3.50 achten. Wir führen die notwendige Argumentation in Satz 3.67 aus.

Satz 3.67. *Sei $\mathcal{F}$ ein n -Faltkomplex und $x \in \mathcal{K}_i$ (mit $0 \leq i \leq n$). Dann bildet die Menge der primitiven Entfaltungen an $[x] \in \mathcal{K}_i / \sim$ (wobei die partielle Ordnung durch Teilfaltung gegeben ist) für $i \neq n - 1$ einen Verband und für $i = n - 1$ einen meet-Halbverband.*

Beweis. Aus Lemma 3.66 lesen wir ab, dass die Menge der primitiven Entfaltungen die geforderte Verbandsstruktur besitzt, falls die partielle Ordnung durch Teilüberlagerung gegeben ist.

Seien $\mathcal{A}$ und $\mathcal{B}$ zwei n -Faltkomplexe, die Teilfaltung von $\mathcal{F}$ sind, sodass $\mathcal{A}$ eine Teilüberlagerung von $\mathcal{B}$ ist. Wir wollen zeigen, dass $\mathcal{A}$ auch eine Teilfaltung von $\mathcal{B}$ ist.

Dazu überprüfen wir die Eigenschaften aus Definition 3.50 (Teilfaltung). Beginnen wir mit der Untersuchung der Randstücke. Diese kommen nur bei $i = n$ ins Spiel und sind gemäß

des Beweises von Lemma 3.62 eindeutig durch die Spaltungspartition festgelegt. Die Ränder kommen genau an den Intervallgrenzen (aus Folgerung 3.63) zustande. Da nach Voraussetzung jede Intervallgrenze von $\mathcal{B}$ auch eine von $\mathcal{A}$ ist, kommt jedes Randstück von $\mathcal{B}$ auch in $\mathcal{A}$ vor.

Betrachten wir als nächstes die Fächer, also $i = n - 1$. Aufgrund der Transitivität der Redukte (Lemma 3.15) ist die Reduktbedingung automatisch erfüllt. Betrachten wir also noch die Überschneidungsfreiheit. Angenommen, es gebe eine Überschneidung beim Übergang von $\mathcal{A}$ zu $\mathcal{B}$. Dann liegt diese komplett in einer der Äquivalenzklassen von $\mathcal{B}$. Das Redukt der entsprechenden zyklischen Reihung von $\mathcal{F}$ auf diese vier Elemente ist aber gleich dem Redukt der zyklischen Reihung in $\mathcal{B}$ auf diese vier Elemente. Folglich läge auch eine Überschneidung von $\mathcal{A}$ in $\mathcal{F}$ vor, im Widerspruch zur Annahme. Daher gilt die Überschneidungsfreiheit. $\square$

Als Konsequenz dieses Satzes ergibt sich eine kanonische Möglichkeit zur Definition der Entfaltung (sollte man dies wollen) – nämlich das kleinste Element des meet-Halbverbandes.

Folgerung 3.68. *Sei $\mathcal{F}$ ein n -Faltkomplex und sei $x \in \mathcal{K}_i$ ($0 \leq i \leq n$) gegeben. Dann gibt es eine primitive Entfaltung $\mathcal{A}$ an $[x]$, sodass für jede primitive Entfaltung $\mathcal{E}$ an $[x]$ gilt, dass $\mathcal{A}$ Teilfaltung von $\mathcal{E}$ ist.*

Wir haben damit die Entfaltung von n -Faltkomplexen vollständig beschrieben. Falls wir die primitive Entfaltung aus Folgerung 3.68 als eindeutige Entfaltung an einer Flächenklasse definieren, trennt jede Entfaltung einer Flächenklasse alle Flächen dieser Klassen voneinander. Wenn wir uns die Entfaltung von Papierfaltungen vorstellen, sollte es aber möglich sein, die Flächenklasse durch Entfaltung nur in zwei kleinere Flächenklassen zu trennen. Wir werden uns im nächsten Abschnitt mit einer alternativen Methode zur Definition einer „eindeutigen Entfaltung“ beschäftigen, die diese Eigenschaft abbildet.

3.5 Faltpläne

Wir haben n -Faltkomplexe bislang analog zu n -Überlagerungskomplexen studiert und die Unterschiede zwischen diesen betont. Aus diesem Grund beruhten alle bisherigen Resultate auf der mengentheoretischen Sprache aus Kapitel 2. Da wir aber über Reihenfolgen und Oberflächen sprechen können, stehen uns weitere Möglichkeiten zur Beschreibung von Faltung und Entfaltung offen. Wir werden in diesem Abschnitt die Idee formalisieren, dass es zur Beschreibung einer Faltung bzw. Entfaltung genügen sollte, wenn man sagt, an welchen Oberflächen dies geschehen soll.

Mit dieser Beschreibung ist es möglich, unabhängig vom Faltzustand über Faltungen und Entfaltungen zu sprechen. Während wir bislang nur sagen konnten, dass bei einem spezifischen n -Faltkomplex gewisse Faltungen oder Entfaltungen möglich sind, war ein Vergleich zwischen den Faltungen verschiedener n -Faltkomplexe sehr schwierig. Die Beschreibung dieses Kapitels erlaubt es, über Faltungen zu sprechen, die unabhängig vom Faltzustand sind und nur vom zugrunde liegenden Muster, also dem homogenen n -Komplex, abhängen.

Der Abschnitt 3.5.1 beschäftigt sich mit der Frage, inwieweit man verschiedene solche Faltungen miteinander vertauschen kann. Im Abschnitt 3.5.2 befassen wir uns genauer mit der Entfaltung und zeichnen eine spezielle Klasse von n -Faltkomplexen aus, für die Faltung und Entfaltung ineinander rückgängig machen.

Während wir bislang alle möglichen Faltungen aller möglichen Dimensionen beschrieben haben, möchten wir uns jetzt auf eine einfachere Situation beschränken: Das Zusammenfalten von zwei Flächen. Dafür benötigen wir zwei Informationen: Um welche Flächen es sich handelt und welche Oberflächen zusammengefasst werden. Da der lokal simpliziale homogene n -Komplex alle diese Informationen bereits enthält, können wir unabhängig vom konkreten Faltzustand über die selben Faltungen reden.

Definition 3.69 (Faltplan). Sei $\mathcal{S} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec)$ ein lokal simplizialer homogener n -Komplex. Ein **Faltplan** (für $\mathcal{S}$) besteht aus:

- Einer Identifikation $\varphi : \mathcal{S}_{\preceq f} \rightarrow \mathcal{S}_{\preceq g}$ mit $f, g \in \mathcal{K}_n$.
- Zwei Oberflächen (f, ε_f) und (g, ε_g) mit $\varepsilon_f, \varepsilon_g \in \{\pm 1\}$.

Wir notieren diesen Faltplan als $(\varphi, (f, \varepsilon_f) \leftrightarrow (g, \varepsilon_g))$.

Bemerkung 3.70. In der Notation von Definition 3.69 ist die Reihenfolge der Oberflächen irrelevant. Daher gilt $(\varphi, (f, \varepsilon_f) \leftrightarrow (g, \varepsilon_g)) = (\varphi, (g, \varepsilon_g) \leftrightarrow (f, \varepsilon_f))$.

Wir sind daran interessiert, wann man einen Faltplan auf einen n -Faltkomplex anwenden kann (dies geht z. B. dann nicht, wenn die angegebenen Oberflächen keine Randstücke sind). Dazu starten wir mit einer informellen Beschreibung, wie wir uns die Ausführung der Faltung vorstellen und überprüfen diese dann auf ihre Durchführbarkeit.

1. Die Flächen sollen zusammengefasst werden, d. h. wir wechseln auf der Ebene der Überlagerungskomplexe von $\mathcal{F}$ zur Erweiterung $\mathcal{F}_\varphi$ (für diesen Schritt müssen wir die Bedingung aus Satz 2.32 überprüfen).
2. Die beiden Oberflächen (f, ε_f) und (g, ε_g) müssen vor der Faltung Randstücke sein und nach der Faltung nicht mehr. Die verbliebenen beiden Randstücke von $[f]$ und $[g]$ bilden die Randstücke der neuen Flächenklasse.
3. Beim Zusammenfalten werden $[f]$ und $[g]$ zu einer gemeinsamen Flächenklasse zusammengefügt. Daher müssen die Fächer, die diese Flächenklassen enthalten, geeignet kombiniert werden.

Die Definitionen 3.31 und 3.41 beschreiben, wann primitive Faltungen der Grade $n - 1$ und n möglich sind. Wenn wir diese kombinieren, erhalten wir ein Kriterium für Faltpläne.

Definition 3.71 (zulässiger Faltplan). Sei $\mathcal{F}$ ein n -Faltkomplex und $(\varphi, (f, \varepsilon_f) \leftrightarrow (g, \varepsilon_g))$ ein Faltplan. Der Faltplan heißt **zulässig** (für $\mathcal{F}$), falls gilt:

1. $\mathcal{F}_\varphi$ ist ein n -Überlagerungskomplex.
2. $(f, \varepsilon_f) \in [f] \times \{\pm 1\}$ und $(g, \varepsilon_g) \in [g] \times \{\pm 1\}$ sind Randstücke von $\mathcal{F}$.
3. Die Randstücke (f, ε_f) und (g, ε_g) sind bezüglich jedes Fächers, in dem sie beide vorkommen, komplementär.

4. Für alle $k, l \in \mathcal{K}_{n-1}$ mit $k \prec f$ und $l \prec g$, sowie $k \sim_\varphi l$ ist die Fächersumme von $\nu_{[k]}^\mathcal{F}$ und $\nu_{[l]}^\mathcal{F}$ an (f, ε_f) und (g, ε_g) wohldefiniert.

Unter diesen Voraussetzungen lässt sich mit einem Faltplan ein n -Proto-Faltkomplex falten. Da wir in Abschnitt 3.2 eingesehen hatten, dass es sehr schwer ist, Fächerkompatibilität zu garantieren, können wir nicht darauf hoffen, im Allgemeinen mehr als einen n -Proto-Faltkomplex zu erhalten.

Definition 3.72 (Anwenden eines Faltplans). Sei $\mathcal{F}$ ein n -Proto-Faltkomplex und $P := (\varphi, (f, \varepsilon_f) \leftrightarrow (g, \varepsilon_g))$ ein zulässiger Faltplan. Dann ist $\mathcal{F}_\varphi$ zusammen mit den folgenden Setzungen ein n -Proto-Faltkomplex.

- Für $h \in \mathcal{K}_n$ mit $h \not\prec_\varphi f$ gilt $\partial_{[h]}^{\mathcal{F}_\varphi} = \partial_{[h]}^\mathcal{F}$. (Die Randstücke aller Flächenklassen, die nicht direkt an der Faltung beteiligt sind, bleiben gleich.)
- Es gilt $\partial_{[f]}^{\mathcal{F}_\varphi} = \partial_{[f]}^\mathcal{F} \uplus \partial_{[g]}^\mathcal{F} \setminus \{(f, \varepsilon_f), (g, \varepsilon_g)\}$. (Die Randstücke der zusammengefalteten Flächenklasse sind genau die beiden vorherigen Randstücke, die nicht zusammengefoldet wurden.)
- Für $k \in \mathcal{K}_{n-1}$ mit einer der folgenden Eigenschaften
 - $[k] \prec [f]$ und $[k] \prec [g]$
 - $[k] \not\prec [f]$ und $[k] \not\prec [g]$
 gilt $\nu_{[k]}^{\mathcal{F}_\varphi} = \nu_{[k]}^\mathcal{F}$. (Alle an der Faltung unbeteiligten Fächer bleiben gleich.)
- Für $k, l \in \mathcal{K}_{n-1}$ mit $[k] \prec [f]$ und $[l] \prec [g]$, sowie $k \not\sim l$ und $k \sim_\varphi l$, ist $\nu_{[k]}^{\mathcal{F}_\varphi}$ als die Fächersumme (vgl. Definition 3.39) von $\nu_{[k]}^\mathcal{F}$ und $\nu_{[l]}^\mathcal{F}$ an (f, ε_f) und (g, ε_g) definiert. (Beteiligte Fächer werden kombiniert.)

Dieser n -Proto-Faltkomplex heißt **Anwendung von P auf $\mathcal{F}$** , geschrieben $P(\mathcal{F})$.

Wohldefiniertheit. Die Wohldefiniertheit dieser Konstruktion folgt daraus, dass wir hier nur primitive Faltungen von Grad $n-1$ (Definition 3.41) und Grad n (Definition 3.31) kombiniert haben. $\square$

Wir haben damit nachgewiesen, dass Faltpläne ein Konzept von Faltung definieren. Als nächstes wollen wir nachweisen, dass dieses Konzept kommutativ ist. Im Anschluss danach beschäftigen wir uns mit der Frage, wie man mit Faltplänen entfalten kann. Es wird sich zeigen, dass die Definition der Entfaltung komplizierter als die Konstruktion des Zusammenfaltens ist.

3.5.1 Kommutativität

Nachdem wir im letzten Abschnitt Faltpläne als neues Faltungskonzept eingeführt haben, untersuchen wir jetzt die strukturellen Eigenschaften dieses Konzepts. Dazu übertragen wir Lemma 2.44, wo die Kommutativität von Erweiterungen nachgewiesen wird, auf Faltpläne.

Falls der Leser nicht am allgemeinen Fall dieser Aussage interessiert ist, endet Abschnitt 3.5.2 mit einem einfacheren Beweis dieses Lemmas (der Folgerung 3.85). Im Ausgleich dafür trifft diese Folgerung nur für spezielle Faltkomplexe zu.

Lemma 3.73. *Seien $\mathcal{F}$ ein n -Proto-Faltkomplex und P, Q Faltpläne, sodass die doppelte Anwendung $P(Q(\mathcal{F}))$ existiert. Dann existiert auch die doppelte Anwendung $Q(P(\mathcal{F}))$ und es gilt die Gleichheit $P(Q(\mathcal{F})) = Q(P(\mathcal{F}))$.*

Beweis. Seien $P = (\varphi, (e, \varepsilon_e) \leftrightarrow (f, \varepsilon_f))$ und $Q = (\psi, (g, \varepsilon_g) \leftrightarrow (h, \varepsilon_h))$ gegeben.

Wir werden die Bedingungen aus Definition 3.71 zur Zulässigkeit von Faltplänen einzeln durchgehen und nachweisen, dass sie für die Anwendung von P auf $\mathcal{F}$ und für die Anwendung von Q auf $P(\mathcal{F})$ zutreffen. Dabei werden für die zweite Situation davon ausgegangen, dass $P(\mathcal{F})$ existiert. Sofern wir nachweisen können, dass $P(\mathcal{F})$ in der ersten Situation existiert, ist diese Annahme gerechtfertigt.

Zuerst betrachten wir die n -Proto-Faltkomplexe als n -Überlagerungskomplexe und weisen nach, dass die Erweiterungen $\mathcal{F}_\varphi$ und $P(\mathcal{F})_\psi = (\mathcal{F}_\varphi)_\psi$ existieren (Bedingung 1).

Da $(\mathcal{F}_\psi)_\varphi$ nach Voraussetzung ein n -Überlagerungskomplex ist, können wir gemäß Lemma 2.44 die Erweiterungen vertauschen, sodass auch $(\mathcal{F}_\varphi)_\psi$ ein n -Überlagerungskomplex ist. Da das Weglassen von Erweiterungen ebenfalls aufgrund von Lemma 2.44 gerechtfertigt ist, handelt es sich auch bei $\mathcal{F}_\varphi$ um einen n -Überlagerungskomplex.

Als nächstes müssen wir nachweisen, dass die Oberflächen der Faltpläne Randstücke sind (Bedingung 2). Wir müssen also zeigen, dass (e, ε_e) und (f, ε_f) Randstücke von $\mathcal{F}$ sind, aber auch, dass (g, ε_g) und (h, ε_h) Randstücke von $P(\mathcal{F})$ sind.

Da P zulässig für $Q(\mathcal{F})$ ist, sind die Oberflächen (e, ε_e) und (f, ε_f) Randstücke von $Q(\mathcal{F})$. Außerdem sind die Oberflächen (g, ε_g) und (h, ε_h) keine Randstücke in $Q(\mathcal{F})$, aber in $\mathcal{F}$. Folglich sind die Oberflächen $(e, \varepsilon_e), (f, \varepsilon_f), (g, \varepsilon_g), (h, \varepsilon_h)$ paarweise verschieden. Damit müssen (e, ε_e) und (f, ε_f) Randstücke von $\mathcal{F}$ sein, weil Randstücke, die nicht gefaltet wurden, ihren Status als Randstücke nicht verlieren. Aus dem selben Grund sind (g, ε_g) und (h, ε_h) Randstücke in $P(\mathcal{F})$, falls diese Anwendung existiert.

Die nächste Bedingung verlangt den Nachweis, dass die Randstücke der Faltpläne in allen Fächern komplementär sind, in denen sie gemeinsam vorkommen (Bedingung 3).

Wir betrachten zunächst (e, ε_e) und (f, ε_f) . Sei $k \in \mathcal{K}_{n-1}$ mit $e, f \in \text{Cor}_{\mathcal{F}}([k]_{\mathcal{F}})$. Für jede simpliziale Orientierung a von $\mathcal{F}_{\leq k}$ ist $\nu_{[k]}^{\mathcal{F}}([a]_{\mathcal{F}})$ das Redukt von $\nu_{[k]}^{Q(\mathcal{F})}([a]_{Q(\mathcal{F})})$ auf $\text{Cor}_{\mathcal{F}}([k]_{\mathcal{F}})$. Da P zulässig für $Q(\mathcal{F})$ ist, sind (e, ε_e) und (f, ε_f) komplementär bezüglich $\nu_{[k]}^{Q(\mathcal{F})}$, folglich auch bezüglich $\nu_{[k]}^{\mathcal{F}}$.

Wenden wir uns als nächstes (g, ε_g) und (h, ε_h) für die Anwendung von Q auf $P(\mathcal{F})$ zu. Sei $k \in \mathcal{K}_{n-1}$ mit $g, h \in \text{Cor}_{P(\mathcal{F})}([k]_{P(\mathcal{F})})$ gegeben. Falls g und h bereits bezüglich $\mathcal{F}$ in einer Corona liegen, müssen sie dort bereits komplementär sein, da ansonsten Q nicht zulässig für $\mathcal{F}$ gewesen wäre. Da die Fächersumme im Schritt von $\mathcal{F}$ zu $P(\mathcal{F})$ diese Komplementarität nicht berührt, bleiben sie bezüglich $\nu_{[k]}^{P(\mathcal{F})}$ komplementär.

Falls g und h bezüglich $\mathcal{F}$ in verschiedenen Coronae liegen, gibt es $k_1, k_2 \in \mathcal{K}_{n-1}$ mit $[k]_{P(\mathcal{F})} = [k_1]_{\mathcal{F}} \uplus [k_2]_{\mathcal{F}}$ (dass $[k]_{P(\mathcal{F})}$ nur in höchstens zwei Äquivalenzklassen zerfallen kann,

folgt aus der Definition 2.23 von Erweiterungen). Zudem ist $\nu_{[k]}^{P(\mathcal{F})}$ die Fächersumme von $\nu_{[k_1]}^{\mathcal{F}}$ und $\nu_{[k_2]}^{\mathcal{F}}$ an (e, ε_e) und (f, ε_f) . Da die Anwendung von Q diese beiden Äquivalenzklassen ebenfalls vereinigt, müssen (g, ε_g) und (h, ε_h) in den Fächern $\nu_{[k_1]}^{\mathcal{F}}$ und $\nu_{[k_2]}^{\mathcal{F}}$ vorkommen. Ohne Einschränkung liegen also (e, ε_e) und (g, ε_g) in $\nu_{[k_1]}^{\mathcal{F}}$, sowie (f, ε_f) und (h, ε_h) in $\nu_{[k_2]}^{\mathcal{F}}$.

Bei der Berechnung von $Q(\mathcal{F})$ wurden diese Fächer an (g, ε_g) und (h, ε_h) summiert. Danach lagen (e, ε_e) und (f, ε_f) komplementär (sonst wäre P nicht zulässig für $Q(\mathcal{F})$ gewesen). Eine Prüfung der Konfluenzdefinition 3.35 ergibt, dass dies nur dann möglich ist, wenn (e, ε_e) und (g, ε_g) , sowie (f, ε_f) und (h, ε_h) komplementär bezüglich $\nu_{[k_1]}^{\mathcal{F}}$ bzw. $\nu_{[k_2]}^{\mathcal{F}}$ waren. Damit sind auch (g, ε_g) und (h, ε_h) komplementär bezüglich $\nu_{[k]}^{P(\mathcal{F})}$. Damit haben wir Bedingung 3 nachgewiesen.

Zuletzt müssen wir nachweisen, dass die notwendigen Fächersummen existieren (Bedingung 4). Aus der Definition 3.39 der Fächersumme geht hervor, dass wir eine involvierte Orientierungsbedingung überprüfen müssen. Diese hängt aber bei genauerer Betrachtung nur von der Wahl der simplizialen Oberfläche ab. Da die Randstück-Paare in $P(Q(\mathcal{F}))$ komplementär sind, ist die Orientierungsbedingung für die simpliziale Oberfläche von $\mathcal{F}$ erfüllt. Demnach ist sie auch erfüllt, wenn wir die Faltpläne in umgekehrter Reihenfolge anwenden.

Wir haben also gezeigt, dass $P(\mathcal{F})$ existiert. Da wir unter dieser Annahme die Existenz von $Q(P(\mathcal{F}))$ nachgewiesen haben, existiert auch $Q(P(\mathcal{F}))$. Wir müssen nur noch nachweisen, dass diese beiden n -Proto-Faltkomplexe identisch sind.

Gemäß Lemma 2.44 sind sie als n -Überlagerungskomplexe gleich. Auch ihre Randstücke sind identisch, da diese sich bei beiden durch die Vereinigung der selben Randstückemengen ergeben, aus denen (e, ε_e) , (f, ε_f) , (g, ε_g) , (h, ε_h) entfernt wurden.

Wir müssen also noch die Gleichheit der Fächer nachweisen. Mit einer geeigneten Wahl der simplizialen Orientierungen lässt sich dies auf die Assoziativität der Konfluenz zurückführen, die wir in Lemma 3.38 nachgewiesen haben. $\square$

Da dieses Lemma die Kommutativität von zwei Faltplänen garantiert, kommt die Frage auf, ob aus der Existenz von $P(Q(\mathcal{F}))$ bereits die Existenz von $Q(\mathcal{F})$ folgt. Dazu benutzen wir eine interessante Konsequenz aus der Definition der Teilfaltung:

Folgerung 3.74. *Sei $\mathcal{F}$ ein n -Proto-Faltkomplex und $\mathcal{G}$ ein n -Faltkomplex, sodass $\mathcal{F}$ eine Teilfaltung von $\mathcal{G}$ ist. Dann ist auch $\mathcal{F}$ ein n -Faltkomplex.*

Beweis. Aufgrund von Lemma 3.15 und Bemerkung 3.16 erhält die Reduktbildung (mit der wir die Fächer von $\mathcal{F}$ aus denen von $\mathcal{G}$ erhalten) die Kompatibilitätsbedingung aus der Definition 3.18 der Fächerkompatibilität. Folglich ist $\mathcal{F}$ ein n -Faltkomplex, sobald $\mathcal{G}$ einer ist. $\square$

Aus Lemma 3.73 und Folgerung 3.74 erhalten wir die folgende Aussage:

Folgerung 3.75. *Sei $\mathcal{F}$ ein n -Proto-Faltkomplex und P, Q Faltpläne, sodass $P(Q(\mathcal{F}))$ existiert und ein n -Faltkomplex ist. Dann sind $\mathcal{F}$, $P(\mathcal{F})$ und $Q(\mathcal{F})$ ebenfalls n -Faltkomplexe.*

3.5.2 Entfaltung

Bis zu diesem Punkt haben wir Faltpläne zur Faltung benutzt. Jetzt wenden wir uns der Entfaltung zu. Dabei lassen wir uns von der Idee leiten, dass es eine eindeutige Entfaltung geben sollte, wenn man zwei aneinanderliegende Oberflächen trennt. Es sollte auch der Fall sein, dass Faltung und Entfaltung am selben Faltplan sich gegenseitig invertieren.

Um eine eindeutige Entfaltung zu definieren, müssen wir überlegen, wie diese aussehen soll. Wir hatten in Folgerung 3.68 zwar festgehalten, dass es eine maximale primitive Entfaltung gibt, aber die maximale primitive Entfaltung einer Flächenklasse würde jede Fläche als eigene Äquivalenzklasse auffassen. Wenn wir an einem Faltplan entfalten, wollen wir nur zwei Oberflächen voneinander trennen, daher sollte die Flächenklasse nur in zwei Äquivalenzklassen aufspalten.

Wir müssen also das Uneindeutigkeitsproblem des Entfaltens auf eine andere Art lösen. Da wir vornehmlich an Flächenfaltungen interessiert sind, sollen die Flächenklassen möglichst wenig modifiziert werden. Die Kanten und Punkten sollen eigentlich in den Hintergrund treten, weswegen wir deren Äquivalenzklassen so weit wie möglich aufspalten dürfen. Faltkomplexe, die – außer an den Flächen – maximal entfaltet sind, nennen wir stabil.

Definition 3.76 (stabiler Faltkomplex). *Sei $\mathcal{F}$ ein n -Faltkomplex. Wir nennen $\mathcal{F}$ **stabil**, falls keine primitiven Entfaltungen von Grad k (mit $k < n$) möglich sind.*

Wir wollen die grundlegende Idee an einem Beispiel illustrieren.

Beispiel 3.77. *Betrachten wir den folgenden lokal simplizialen 1-Überlagerungskomplex¹³:*

$$\begin{array}{ll} \mathcal{K}_0 := \{1, 2, 3\} & \bullet \text{---} \bullet \text{---} \bullet \\ \mathcal{K}_1 := \{\{1, 2\}, \{2, 3\}\} & \begin{array}{ccc} 1 & 2 & 3 \end{array} \end{array}$$

mit $\prec := \in$ und $\sim$ später festgelegt wird.

Falls $\{1, 2\} \sim \{2, 3\}$, also auch $1 \sim 3$ gilt, dann ist dieser Komplex stabil (eine Entfaltung von Grad 0 ist nicht möglich, solange die Kanten äquivalent sind). Falls aber nur $1 \sim 3$ gilt (und $\{1, 2\} \not\sim \{2, 3\}$), dann ist der Komplex nicht stabil. Dann ist nämlich eine primitive Entfaltung an der Punktklasse $\{1, 3\}$ gemäß Lemma 3.61 möglich.

Unser Ziel ist es, die Entfaltung an einem Faltplan wie folgt zu konstruieren:

1. Wir trennen zuerst die im Faltplan angegebenen Oberflächen, d. h. wir führen zuerst die primitive Entfaltung der Flächenklasse durch.
2. Dann trennen wir alle Punkt- und Kantenklassen, die nicht durch Flächen zusammengehalten werden. Formal wollen wir zu einem stabilen n -Faltkomplex übergehen.

Damit diese Konstruktion funktioniert, müssen wir nachweisen, dass ein solcher eindeutiger stabiler n -Faltkomplex stets existiert.

¹³Das grundlegende Prinzip ist schon bei Überlagerungskomplexen vorhanden. Die Betrachtung von Faltkomplexen macht alles ein wenig aufwendiger, ändert dieses Prinzip aber nicht.

Lemma 3.78. *Sei $\mathcal{F}$ ein n -Faltkomplex. Dann gibt es genau einen stabilen n -Faltkomplex $\mathcal{E}$, der Teilfaltung von $\mathcal{F}$ ist.*

Beweis. Wir zeigen zuerst die Eindeutigkeit. Seien $\mathcal{D}$ und $\mathcal{E}$ zwei stabile n -Faltkomplexe, die Teilfaltungen von $\mathcal{F}$ sind. Wähle ein $x \in \mathcal{K}_i$ mit $0 \leq i < n$ maximal, sodass die Aufspaltung von $[x]$ in $\sim^{\mathcal{D}}$ - und $\sim^{\mathcal{E}}$ -Klassen verschieden ist.

Betrachten wir zunächst den Fall $i < n - 1$. In diesem Fall sind alle möglichen Entfaltungen von $[x]$ gemäß Folgerung 3.58 und Satz 2.49 durch die Bindungsrelation festgelegt. Da wir i maximal gewählt haben, sind die Bindungsrelationen von $\mathcal{D}$ und $\mathcal{E}$ auf $[x]$ identisch. Da $\mathcal{D}$ und $\mathcal{E}$ stabil sind, sind ihre Aufspaltungen von $[x]$ maximal nach Folgerung 3.68. Aus der selben Folgerung schließen wir, dass die maximale Aufspaltung eindeutig ist, weswegen die Aufspaltungen bezüglich $\mathcal{D}$ und $\mathcal{E}$ identisch sind.

Im Fall $i = n - 1$ werden primitive Entfaltungen zusätzlich dadurch eingeschränkt, dass die induzierte Fächerpartition überschneidungsfrei bezüglich der ursprünglichen zyklischen Reihungen ist (vgl Lemma 3.61). Da die induzierte Fächerpartition eindeutig durch die Spaltungspartition festgelegt ist, genügt aber auch hier eine Betrachtung der Spaltungspartitionen. Daher führt die obige Argumentation auch hier zum Ziel.

Als nächstes zeigen wir die Existenz. Dazu verwenden wir das Beweisprinzip aus dem Eindeutigkeitsbeweis, um eine explizite Konstruktion anzugeben. Wir geben ein Verfahren an, die verschiedenen Klassen so abzulaufen, dass wir jede Klasse maximal aufspalten können (im Sinne von Folgerung 3.68) und keine Klasse am Ende des Prozesses noch weiter aufgespalten werden kann.

Wir betrachten ein $x \in \mathcal{K}_i$ (mit $0 \leq i < n$ maximal), sodass es eine nicht-triviale Spaltungspartition von $[x]$ gibt (die gemäß Lemma 3.61 oder Folgerung 3.58 charakterisiert ist). Nach Folgerung 3.68 gibt es eine primitive Entfaltung $\mathcal{E}$, bei dem an keinem Element der Spaltungspartition eine primitive Entfaltung möglich ist. Da diese Aufspaltung über die Bindungsrelation nur Einfluss auf Klassen $[y] \in \mathcal{K}_j / \sim$ mit $j < i$ hat, sind die primitiven Entfaltungen an allen $[x] \in \mathcal{K}_i / \sim$ voneinander unabhängig. Daher können wir den n -Faltkomplex bilden, bei dem alle diese primitiven Entfaltungen ausgeführt wurden und dieser ist noch immer eine Teilfaltung von $\mathcal{F}$. Da alle Klassen höherer Dimension bereits maximal aufgespalten wurden, ändert sich die maximale Spaltung in den nachfolgenden Aufspaltungen nicht. Da die Dimension des n -Faltkomplexes endlich ist, terminiert dieser Prozess. $\square$

Damit ist gezeigt, dass wir mit dem Konzept stabiler n -Faltkomplexe eine eindeutige Entfaltung auf Basis eines Faltplans definieren können.

Definition 3.79 (zulässiger Entfaltplan). *Sei $\mathcal{F}$ ein n -Faltkomplex und $P := (\varphi, (f, \varepsilon_f) \leftrightarrow (g, \varepsilon_g))$ ein Faltplan. Wir nennen P eine **zulässige Entfaltung von $\mathcal{F}$** , falls gilt:*

1. *Für alle $x \preceq f$ gilt $x \sim \varphi(x)$. (Die Flächen sind zusammengefasst)*
2. *Die Oberflächen (f, ε_f) und (g, ε_g) sind komplementär und keine Randstücke. (Die Oberflächen liegen in der Faltung aufeinander.)*

Den stabilen n -Faltkomplex, den man durch Entfalten von (f, ε_f) und (g, ε_g) erhält (im Sinne von Lemma 3.78), bezeichnen wir mit $P^{-1}(\mathcal{F})$.

Definition 3.80 (Anwenden eines Entfaltplans). Sei $\mathcal{F}$ ein n -Proto-Faltkomplex und $P := (\varphi, (f, \varepsilon_f) \leftrightarrow (g, \varepsilon_g))$ ein zulässiger Entfaltplan. Im Sinne von Folgerung 3.63 haben wir eine primitive Entfaltung zwischen f und g . Dann bezeichnen wir den eindeutigen stabilen n -Proto-Faltkomplex zu dieser Entfaltung mit $P^{-1}(\mathcal{F})$, genannt die **Anwendung von P^{-1} auf $\mathcal{F}$** .

Wohldefiniertheit. Wir müssen überprüfen, ob wir Folgerung 3.63 anwenden können. Wir wollen die Äquivalenzklasse $[f]$ so in $M_1 \uplus M_2$ aufspalten, dass $f \in M_1$ und $g \in M_2$ liegt (dabei wird $g \sim f$ dadurch garantiert, dass P ein zulässiger Entfaltplan von $\mathcal{F}$ ist). Bezüglich einer beliebigen zyklischen Reihung an einer der Kantenklassen lassen sich die Elemente der Flächenklasse (bis auf Invertieren) linear ordnen, wie Lemma 3.26 zeigt:

$$a < \dots < f < g < \dots < b$$

Dabei sind f und g benachbart, da ihre Oberflächen nach Voraussetzung (Bedingung 2 in der Definition 3.79 der Zulässigkeit) komplementär sind. Dann liefert $M_1 := \{a, \dots, f\}$ und $M_2 := \{g, \dots, b\}$ gemäß Folgerung 3.63 eine gültige Spaltungspartition.

Diese primitive Entfaltung liefert einen n -Proto-Faltkomplex $\mathcal{E}$. Nach Lemma 3.78 gibt es einen eindeutigen stabilen n -Faltkomplex, der Teilfaltung von $\mathcal{E}$ ist. Damit sind wir legitimiert, diesen mit $P^{-1}(\mathcal{F})$ zu bezeichnen. $\square$

Damit haben wir nachgewiesen, dass es eine eindeutige Entfaltung an Faltplänen gibt. Ebenso wie bei der Faltung an Faltplänen weisen wir jetzt nach, dass auch die Entfaltungen miteinander vertauschen.

Lemma 3.81. Sei $\mathcal{F}$ ein n -Proto-Faltkomplex und P, Q Faltpläne, sodass $P^{-1}(Q^{-1}(\mathcal{F}))$ existiert. Dann ist auch $Q^{-1}(P^{-1}(\mathcal{F}))$ wohldefiniert und $P^{-1}(Q^{-1}(\mathcal{F})) = Q^{-1}(P^{-1}(\mathcal{F}))$.

Beweis. Seien $P = (\varphi, (a, \varepsilon_a) \leftrightarrow (b, \varepsilon_b))$ und $Q = (\psi, (c, \varepsilon_c) \leftrightarrow (d, \varepsilon_d))$.

Wir zeigen, dass P im Sinne von Definition 3.79 eine zulässige Entfaltung für $\mathcal{F}$ ist.

1. Da $x \sim^{Q^{-1}(\mathcal{F})} \varphi(x)$ für alle $x \preccurlyeq a$ gilt und $Q^{-1}(\mathcal{F})$ Teilfaltung von $\mathcal{F}$ ist, folgt auch $x \sim^{\mathcal{F}} \varphi(x)$ für alle $x \preccurlyeq a$.
2. Da (a, ε_a) und (b, ε_b) in $Q^{-1}(\mathcal{F})$ komplementär und keine Randstücke sind, sind sie es in $\mathcal{F}$ auch nicht (da $a \sim b$ gilt).

Als nächstes müssen wir nachweisen, dass Q eine zulässige Entfaltung für $P^{-1}(\mathcal{F})$ ist.

1. Da $P^{-1}(Q^{-1}(\mathcal{F}))$ existiert, folgt $\{a, b\} \neq \{c, d\}$, insbesondere trennt P^{-1} die Äquivalenz $c \sim d$ nicht auf. Aufgrund der Abwärtskompatibilität ist damit $x \sim \psi(x)$ für alle $x \preccurlyeq c$ erfüllt.
2. Da in $P^{-1}(\mathcal{F})$ noch immer $c \sim d$ gilt, sind (c, ε_c) und (d, ε_d) noch immer komplementär (die Fächer haben sich lokal nicht verändert). Da $\{(a, \varepsilon_a), (b, \varepsilon_b)\}$ und $\{(c, \varepsilon_c), (d, \varepsilon_d)\}$ disjunkt sind (sonst könnte $P^{-1}(Q^{-1}(\mathcal{F}))$ nicht existieren), sind (c, ε_c) und (d, ε_d) auch keine Randstücke.

Zuletzt müssen wir die Gleichheit zeigen. Da in beiden Fällen die gleichen primitiven Spaltungen von $\mathcal{K}_n / \sim$ vorgenommen wurden, folgt die Gleichheit aus der Eindeutigkeit in Lemma 3.78. $\square$

Wenn das Konzept der Faltpläne sinnvoll definiert ist, dann erwarten wir, dass sich Faltung und Entfaltung gegenseitig aufheben. Aufgrund der asymmetrischen Situation (bei Entfaltungen erhalten wir immer stabile n -Faltkomplexe) gilt dies nicht allgemein. Wenn wir unser Interesse aber ausschließlich auf Faltungen von Flächen eingeschränkt haben, ist es durchaus plausibel, nur stabile n -Faltkomplexe zu betrachten. So eingeschränkt erfüllt sich unsere Erwartung.

Lemma 3.82. *Sei $\mathcal{F}$ ein stabiler n -Faltkomplex und P sei ein Faltpfad. Falls $P(\mathcal{F})$ existiert, dann auch $P^{-1}(P(\mathcal{F}))$ und es gilt $P^{-1}(P(\mathcal{F})) = \mathcal{F}$.*

Beweis. Nach Konstruktion von $P(\mathcal{F})$ (Definition 3.72) ist P eine zulässige Entfaltung von $P(\mathcal{F})$.

Der einzige Unterschied zwischen $\mathcal{F}$ und $P(\mathcal{F})$ in Bezug auf ihre $\mathcal{K}_n$ -Klassen ist die Anwendung des Faltpfades. Die beiden $\mathcal{K}_n$ -Klassen, die in der Faltung vereinigt wurde, werden in der Entfaltung wieder getrennt. Da es zu diesen $\mathcal{K}_n$ -Klassen genau einen stabilen n -Faltkomplex gibt und $\mathcal{F}$ schon stabil ist, folgt die geforderte Gleichheit. $\square$

Der Beweis dieses Lemmas war sehr einfach. Das liegt daran, dass wir Entfaltungen indirekt definiert haben und daher auf eine Eindeutigkeitsaussage zurückgreifen können (Lemma 3.78 garantiert uns einen eindeutigen stabilen n -Faltkomplex). Die umgekehrte Richtung – eine Entfaltung rückgängig zu machen – wird daher komplizierter (da wir nicht genau wissen, wie $P^{-1}(\mathcal{F})$ aussieht). An dieser Stelle ist es also notwendig, dass wir uns genauer damit beschäftigen, welches Resultat die Konstruktion des stabilen n -Faltkomplexes in Lemma 3.78 liefert.

Lemma 3.83. *Sei $\mathcal{F}$ ein stabiler n -Faltkomplex und P sei ein Faltpfad. Falls $P^{-1}(\mathcal{F})$ existiert, dann auch $P(P^{-1}(\mathcal{F}))$ und es gilt $P(P^{-1}(\mathcal{F})) = \mathcal{F}$.*

Beweis. Sei $P = (\varphi, (f, \varepsilon_f) \leftrightarrow (g, \varepsilon_g))$ der Faltpfad. Wir zeigen, dass P eine zulässige Faltung für $P^{-1}(\mathcal{F})$ ist (mit den Kriterien aus Definition 3.71).

1. Wir müssen zeigen, dass die Erweiterung von $P^{-1}(\mathcal{F})$ um φ wohldefiniert ist. Da $P^{-1}(\mathcal{F})$ eine Teilfaltung von $\mathcal{F}$ ist und die formale Erweiterung von $\sim^{P^{-1}(\mathcal{F})}$ um φ in $\sim^{\mathcal{F}}$ enthalten ist, erfüllt diese Erweiterung die Irreduzibilitätsbedingung. Folglich ist die Erweiterung von $P^{-1}(\mathcal{F})$ um φ wohldefiniert.
2. Wir müssen nachweisen, dass (f, ε_f) und (g, ε_g) Randstücke in $P^{-1}(\mathcal{F})$ sind. Dies ist durch die Konstruktion der Entfaltung bereits erfüllt.
3. Um nachzuweisen, dass diese Randstücke in jedem Fächer komplementär sind, in dem sie gemeinsam vorkommen, verwenden wir, dass diese Randstücke nach der primitiven Entfaltung von $\mathcal{F}$ an (f, ε_f) und (g, ε_g) komplementär sind. Da $P^{-1}(\mathcal{F})$ eine Teilfaltung dieser primitiven Entfaltung ist, sind die zyklischen Reihungen von $P^{-1}(\mathcal{F})$ Redukte

von denen der primitiven Entfaltung. Damit kommen f und g entweder nicht im selben Fächer vor oder die Randstücke bleiben komplementär.

4. Zuletzt müssen wir zeigen, dass die Fächersummen an den korrespondierenden Kanten wohldefiniert sind. Da die primitive Entfaltung vom Grad n die Fächer nicht verändert, werden die Fächer nur durch die Anwendung von Lemma 3.78 gespalten. Da es sich bei den Spaltungen der Kantenklassen um primitive Entfaltungen handelte, sind sie insbesondere überschneidungsfrei bezüglich der zyklischen Reihungen von $\mathcal{F}$. Daher lassen sich diese Aufspaltungen durch Fächersumme rückgängig machen. Damit sind die Fächersummen wohldefiniert.

Jetzt ist noch die Gleichheit nachzuweisen. Dazu prüfen wir die Äquivalenzklassen, Randstücke und Fächer auf Gleichheit.

Auf $\mathcal{K}_n$ ist die Anwendung von P^{-1} äquivalent zu einer primitiven Entfaltung in genau zwei Mengen. Der Wechsel zu einem stabilen n -Faltkomplex ändert die Klassen von $\mathcal{K}_n$ nicht. Folglich macht die Anwendung von P diese Trennung rückgängig. Da durch die Anwendung von P^{-1} genau die Randstücke hinzukommen, die durch Anwendung von P wieder entfernt werden, sind auch die Randstücke gleich.

Für die anderen Dimensionen zeigen wir zunächst, dass beim Anwenden von P^{-1} jede Äquivalenzklasse in höchstens zwei Klassen gespalten wird. Da die Zerlegung der Fächer gemäß Lemma 3.61 durch die Spaltungspartition festgelegt ist, genügt es, zum Beweis dieser Aussage nur die Bindungsrelation zu betrachten.

Sei also $x \in \mathcal{K}_j$ mit $j < n$ und $[x] \prec [f]$ (bei allen anderen Äquivalenzklassen ändert sich die Bindungsrelation nicht). Bezüglich der Spaltungspartition $[f] = M_1 \uplus M_2$ definieren wir die Mengen

$$N_i := \{y \in [x] \mid \exists h \in M_i : y \prec h\}, \quad i \in \{1, 2\}.$$

Klar: $[x] = N_1 \cup N_2$. Nach Konstruktion ist jedes N_i in einer Bindungsklasse von $[x]$ enthalten, daher zerfällt $[x]$ maximal in diese beiden Mengen (es zerfällt nicht, falls die Bindungsklasse ganz $[x]$ ist). Damit ist die Behauptung gezeigt, da bei der Konstruktion in Lemma 3.78 jede Äquivalenzklasse nur einmal betrachtet werden muss.

Um nachzuweisen, dass diese Zerlegung durch die Anwendung von P wieder rückgängig gemacht wird, genügt es anzumerken, dass die jeweils zu identifizierenden Elemente nach Konstruktion der N_i in verschiedenen N_i liegen müssen (da f und g in verschiedenen M_i liegen). Damit werden die Äquivalenzklassen bei der Faltung wieder vereinigt. Da keine weiteren Vereinigungen vorkommen, sind damit die Äquivalenzklassen von $\mathcal{F}$ und $P(P^{-1}(\mathcal{F}))$ identisch.

In Bezug auf die Fächer ist zu sagen: In $\mathcal{F}$ sind (f, ε_f) und (g, ε_g) komplementär. Da die Fächersumme der relevanten Fächer an genau diesen Randstücken erfolgt, gilt also auch hier Gleichheit. $\square$

Wir haben also gezeigt, dass die Anwendungen von P und P^{-1} auf stabilen n -Faltkomplexen zueinander invers sind. Damit handelt es sich um eine geeignete Struktur zur Modellierung von Faltungen an spezifischen Oberflächen. Da der Bedingung der Stabilität hier sehr zentral

ist, müssen wir natürlich garantieren, dass diese Bedingung durch Faltung nicht verloren geht. Dass die Anwendung von P^{-1} zu einem stabilen n -Faltkomplex führt, folgt direkt aus deren Definition. Dies ist bei der Anwendung von P , also beim Falten des Komplexes, nicht unbedingt klar. Wir müssen also noch nachweisen, dass die Anwendung von P auf einen stabilen n -Faltkomplex wieder einen solchen liefert.

Lemma 3.84. *Sei $\mathcal{F}$ ein stabiler n -Faltkomplex und P ein zulässiger Faltplan für $\mathcal{F}$. Dann ist $P(\mathcal{F})$ auch stabil.*

Beweis. Angenommen, $P(\mathcal{F})$ wäre nicht stabil. Dann gäbe es eine primitive Entfaltung von $P(\mathcal{F})$, die einen Grad kleiner als n hat. Wir unterscheiden zwei Fälle, da beim Grad $n - 1$ zusätzliche Komplikationen auftreten.

Wir betrachten zunächst den Fall, dass die primitive Entfaltung einen Grad kleiner als $n - 1$ hat. Dann spaltet sie eine Äquivalenzklasse $[x]_{\sim P(\mathcal{F})} \in \mathcal{K}_i / \sim^{P(\mathcal{F})}$ mit $i < n - 1$. Nach Konstruktion der Faltung an P (bei einer Erweiterung werden die Äquivalenzklassen höchstens paarweise identifiziert, wie wir in Definition 2.23 festgehalten haben) ist $[x]_{\sim P(\mathcal{F})}$ die Vereinigung von höchstens zwei $\sim$ -Äquivalenzklassen $[x_1]_{\sim}$ und $[x_2]_{\sim}$.

Da $\mathcal{F}$ stabil ist, sind die Bindungsklassen dieser beiden Äquivalenzklassen identisch zu den Klassen selbst. Andernfalls würde uns die Charakterisierung primitiver Entfaltungen aus Bemerkung 3.58 eine primitive Entfaltung von $\mathcal{F}$ liefern. Da die Vereinigung von Äquivalenzklassen die Bindungsklassen nicht verkleinert, gibt es in $[x]_{\sim P(\mathcal{F})}$ entweder eine oder zwei Bindungsklassen. Der einzige Grund für die Vereinigung von $[x_1]_{\sim}$ und $[x_2]_{\sim}$ ist aber, dass zwei Elemente – je eines aus jeder Klasse – durch das Zusammenfügen von zwei Elementen in $\mathcal{K}_n$ gebunden sind. Folglich kann es nur eine Bindungsklasse geben, also keine primitiven Entfaltungen von Grad kleiner als $n - 1$.

Betrachten wir den Fall, dass die primitive Entfaltungen den Grad $n - 1$ hat. Hier kann es Partitionen geben, die mit der Bindungsrelation verträglich sind, aber keine Spaltungspartitionen sind (aufgrund der zusätzliche Bedingung in der Charakterisierung primitiver Entfaltungen aus Lemma 3.61). Untersuchen wir die Situation, dass die Bindungsklasse von $[x_1]_{\sim}$ nicht gleich der gesamten Klasse ist, genauer. Da $\mathcal{F}$ stabil ist, kann eine Zerlegung, die die Bindungsklassen respektiert, keine primitive Entfaltung liefern. Daher ist die induzierte Fächerpartition dieser Zerlegung nicht überschneidungsfrei bezüglich einer zyklischen Reihung von $\mathcal{F}$. Da die Anwendung von P auf $\mathcal{F}$ die Redukte auf die Corona $\text{Cor}_{\mathcal{F}}([x_1]_{\sim})$ nicht verändert, bleibt diese Überschneidung bestehen. Folglich gibt es keine Spaltungspartition von $[x]_{\sim P(\mathcal{F})}$, in der unterschiedliche Bindungsklassen von $[x_1]_{\sim}$ in verschiedenen Mengen liegen. Da dieses Argument auch für die Bindungsklassen von $[x_2]_{\sim}$ gilt, können wir wie im ersten Fall argumentieren, dass die Spaltung $[x_1]_{\sim} \uplus [x_2]_{\sim}$ nicht möglich ist.

Da wir gezeigt haben, dass es keine primitiven Entfaltungen eines Grades kleiner als n geben kann, ist $P(\mathcal{F})$ stabil. $\square$

Wir haben jetzt insgesamt nachgewiesen, dass wir uns zur Untersuchung von Papierfaltungen (wo wir nur Flächen aufeinanderfalten) allein auf stabile n -Faltkomplexe und Faltpläne beziehen können. Diese sind also ein adäquates Modell zur Beschreibung dieser Faltungen.

Um zu demonstrieren, wie stark die Erleichterungen dieses Modells sind, beweisen wir die Kommutativität von Faltplänen, die wir in Lemma 3.73 mühevoll gezeigt haben, erneut. Dabei schränken wir uns auf stabile n -Faltkomplexe ein. Wir werden den Beweis auf die Kommutativität von Entfaltungen aus Lemma 3.81 zurückführen. Dieser ist deutlich einfacher als die korrespondierende Aussage für Faltungen, da die Eindeutigkeit eines stabilen n -Komplexes die Argumentation stark vereinfacht. Man sollte sich aber nicht davon täuschen lassen – wir haben den Hauptteil der Arbeit in den Beweisen der Lemmata 3.82 und 3.83 geleistet, in denen wir nachgewiesen haben, dass sich Faltung und Entfaltung von Faltplänen gegenseitig aufheben.

Folgerung 3.85. *Sei $\mathcal{F}$ ein stabiler n -Faltkomplex und P, Q Faltpläne, sodass $P(Q(\mathcal{F}))$ existiert. Dann existiert auch $Q(P(\mathcal{F}))$ und es gilt $P(Q(\mathcal{F})) = Q(P(\mathcal{F}))$.*

Beweis. Sei $\mathcal{G} := P(Q(\mathcal{F}))$. Es gelten folgende Äquivalenzen (wobei man Falt- und Entfaltungen gemäß der Lemmata 3.82 und 3.83 anwenden darf):

$$\begin{aligned}
& P(Q(\mathcal{F})) = \mathcal{G} \\
\Leftrightarrow & \quad Q(\mathcal{F}) = P^{-1}(\mathcal{G}) \\
\Leftrightarrow & \quad \mathcal{F} = Q^{-1}(P^{-1}(\mathcal{G})) = P^{-1}(Q^{-1}(\mathcal{G})) \\
\Leftrightarrow & \quad P(\mathcal{F}) = Q^{-1}(\mathcal{G}) \\
\Leftrightarrow & \quad Q(P(\mathcal{F})) = \mathcal{G}
\end{aligned}$$

Die Vertauschung der Entfaltungen im mittleren Schritt ist dabei durch die Kommutativität von Entfaltungen aus Lemma 3.81 gewährleistet. $\square$

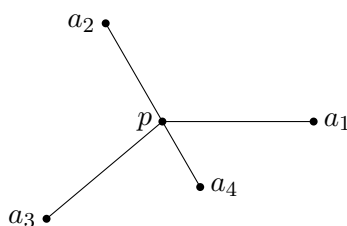
4 Die Einbettungsfrage

Wir haben in dieser Arbeit den Ansatz gewählt, Faltungen unabhängig von ihrer Position im $\mathbb{R}^3$ zu beschreiben. Dazu haben wir in den Kapiteln 2 und 3 ein Modell aufgestellt, das viele Eigenschaften von Faltungen dadurch beschreibt, dass es die Umläufe der Flächen um jede Kante speichert. In diesem Abschnitt möchten wir die Überlegungen zu Einbettungen von n -Überlagerungskomplexen aus Abschnitt 2.5 auf das vollständige Modell erweitern. Zunächst müssen wir den Begriff der Einbettung für n -Faltkomplexe formulieren.

Dazu müssen wir die Fächer geeignet repräsentieren (keine Eigenschaft der Randstücke ist in einer Einbettung sichtbar). Hier ergeben sich zwei Unterprobleme: Wir müssen die Reihenfolge von äquivalenten Flächen und die Reihenfolge von nicht-äquivalenten Flächen korrekt darstellen können. Wir beschreiben zuerst die Reihenfolge der nicht-äquivalenten Flächen und zeigen im Anschluss, dass die Reihenfolge der äquivalenten Flächen keine weiteren Restriktionen liefert.

Wir betrachten die Situation zuerst in zwei Dimensionen, bevor wir sie auf drei (und damit fast unmittelbar auch höhere) Dimensionen erweitern. Dazu betrachten wir exemplarisch einen Fächer ν_p eines 1-Faltkomplexes (wir nehmen an, dass alle Äquivalenzklassen einelementig sind und sparen uns damit die Erinnerung an die Existenz dieser Klassen). Wir separieren die beiden Unterprobleme und nehmen zunächst an, dass alle Kanten, die an den Punkt p angrenzen, paarweise nicht äquivalent sind.

Die einzige simpliziale Orientierung von p ist das 1-Tupel (p) , die zyklische Reihung $\nu_p(p)$ sei $(a_1, \dots, a_m)$. Graphisch liegt die folgende Situation (für den Fall $m = 4$) vor:



Eine naive Beschreibung der Reihenfolge in der Einbettung wird durch Polarkoordinaten (mit Radius r und Winkel φ) geliefert. Dabei legen wir den Winkel φ_1 des Punktes a_1 auf 0 fest (wie in obiger Graphik dargestellt). Damit liegen die Punktklassen $a_1, \dots, a_m$ auf den Winkeln $\varphi_1, \dots, \varphi_m$ (wobei wir jeweils den Vertreter aus $[0, 2\pi)$ wählen). Die Punktklassen liegen genau dann in der richtigen Reihenfolge, wenn $\varphi_1 < \varphi_2 < \dots < \varphi_m$ gilt.

Während diese Methode sehr anschaulich ist, ist ihre praktische Umsetzung sehr umständlich. Für jeden Fächer muss ein eigenes Koordinatensystem definiert werden und im Dreidimensionalen müssen die Punktklassen zuerst auf eine Ebene projiziert werden, die orthogonal zur ursprünglichen Kantenklasse steht.

Wir versuchen daher mit den Differenzen $a_i - p$ zu arbeiten, die unabhängig von der konkreten Lage der Punkte sind. Unser Ansatz ist, uns in mathematisch positiver Drehrichtung (also gegen den Uhrzeigersinn) um den Punkt p herumzubewegen, während wir die Punkte $a_1, a_2, \dots, a_m, a_1$ in dieser Reihenfolge ablaufen. Die Punkte sind genau dann richtig angeordnet, wenn wir den Punkt p genau einmal umlaufen haben.

Wir müssen also die Winkel φ_i bestimmen. In zwei Dimensionen lässt sich dies vergleichsweise einfach dadurch regeln, dass wir sowohl das Skalar- als auch das Kreuzprodukt (vermöge der kanonischen Einbettung $\mathbb{R}^2 \hookrightarrow \mathbb{R}^3$) bestimmen (diese liefern uns dann Kosinus und Sinus des Winkels). In drei Dimensionen werden die Differenzen $a_i - p$ aber uneindeutig, da es dann zwei Eckpunkte p für die zentrale Kante gibt. Wir müssen dann mit den Orthogonalprojektionen auf die zentrale Kante arbeiten. Sobald diese Projektionen sich aber unterscheiden, ergeben sich weitere Komplikationen. Daher versuchen wir, die Winkel indirekt abzulesen. Dazu verwenden wir die Determinante, denn es gilt:

$$\det(a_1 - p, a_2 - p) \text{ ist } \begin{cases} > 0 & \text{falls der orientierten Winkel in } (0, \pi) \text{ liegt} \\ 0 & \text{falls den orientierten Winkel in } \{0, \pi\} \text{ liegt} \\ < 0 & \text{falls den orientierten Winkel in } (\pi, 2\pi) \text{ liegt} \end{cases}$$

Da dies nur eine unzuverlässige Methode zur Winkelbestimmung ist (sie ist in der Regel zweideutig), funktioniert sie nur für drei Punkte a_1, a_2, a_3 .

Lemma 4.1. *Seien $a_1, a_2, a_3 \in \mathbb{R}^2$ mit paarweise verschiedenen Winkeln (bzgl. Polarkoordinaten) $\varphi_1 = 0, \varphi_2, \varphi_3$. Dann gilt $\varphi_1 < \varphi_2 < \varphi_3$ genau dann, wenn höchstens eine der folgenden Determinanten negativ ist:*

$$\det(a_1, a_2) \qquad \det(a_2, a_3) \qquad \det(a_3, a_1)$$

Beweis. Wir bezeichnen den orientierten Winkel zwischen den Vektoren a_i und a_{i+1} (Indizes modulo 3) mit $\Delta_i \in (0, 2\pi)$. Dann gilt $\varphi_1 < \varphi_2 < \varphi_3$ genau dann, wenn $\Delta_1 + \Delta_2 + \Delta_3 = 2\pi$. Zudem ist diese Summe stets ein positives Vielfaches von 2π . Die Bedingung $\det(a_i, a_{i+1}) \geq 0$ ist äquivalent zu $\Delta_i \in (0, \pi]$. Ebenso ist $\det(a_i, a_{i+1}) < 0$ äquivalent zu $\Delta_i \in (\pi, 2\pi)$.

Wir zeigen das Lemma, indem wir die Äquivalenz der Negationen nachweisen.

Sei $\Delta_1 + \Delta_2 + \Delta_3 \neq 2\pi$, d. h. $\Delta_1 + \Delta_2 + \Delta_3 \geq 4\pi$. Angenommen, zwei der Determinanten seien ≥ 0 , dann sind ohne Beschränkung der Allgemeinheit Δ_1 und Δ_2 kleiner oder gleich π . Es folgt $\Delta_3 \geq 2\pi$, im Widerspruch zu $\Delta_3 < 2\pi$.

Seien nun mindestens zwei der Determinanten negativ. Dann sind ohne Einschränkung Δ_1 und Δ_2 größer als π , demnach müsste $\Delta_1 < 0$ sein, um in der Summe 2π zu erreichen. Dies ist ein Widerspruch. $\square$

Obwohl Lemma 4.1 nur eine Aussage über drei Punkte trifft, benutzen wir es in Lemma 4.2, um eine Aussage über beliebig lange zyklische Reihenfolgen zu machen (auch wenn dabei der Rechenaufwand stark ansteigt).

Lemma 4.2. *Seien $a_1, a_2, \dots, a_m \in \mathbb{R}^2$ mit paarweise verschiedenen Winkeln (bzgl. Polarkoordinaten) $\varphi_1 = 0, \varphi_2, \dots, \varphi_m$. Dann gilt $\varphi_1 < \varphi_2 < \dots < \varphi_m$ genau dann, wenn für alle $1 < j < k \leq m$ maximal eine der folgenden Determinanten negativ ist:*

$$\det(a_1, a_j) \qquad \det(a_j, a_k) \qquad \det(a_k, a_1)$$

Beweis. Wie in Lemma 4.1 zeigen wir die Äquivalenz der Negationen.

Seien $1 < j < k \leq m$ mit $a_j > a_k$ gegeben. Dann liefert Lemma 4.1 mit den Punkten a_1, a_j, a_k , dass mindestens zwei der Determinanten

$$\det(a_1, a_j) \qquad \det(a_j, a_k) \qquad \det(a_k, a_1)$$

negativ sind.

Seien umgekehrt $1 < j < k \leq m$ gegeben, sodass mindestens zwei der Determinanten negativ sind. Nach Lemma 4.1 ist dann $0 < \varphi_j < \varphi_k$ falsch, also gilt $\varphi_j > \varphi_k$. $\square$

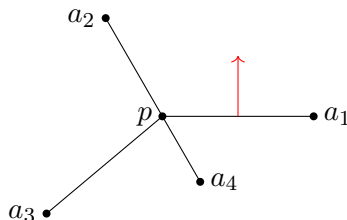
Da die Wahl des Koordinatensystems beliebig ist (schließlich können wir jedem Punkt den Winkel 0 zuordnen), ist die Voraussetzung $\varphi_1 = 0$ in Lemma 4.2 unerheblich und die Wahl dieses Punktes beliebig. Damit ist diese Determinantenbedingung eine adäquate Reformulierung für die zyklische Reihenfolge der Kanten $a_1, \dots, a_m$.

Allerdings müssen wir überprüfen, dass diese Bedingung nicht allein durch unsere vereinfachenden Annahmen entsteht. Wir müssen noch die folgenden Probleme beheben:

1. Die Bedingung muss kompatibel mit der Existenz äquivalenter Punkte sein.
2. Wir haben uns auf Dimension 2 eingeschränkt, möchten aber auch eine Bedingung in Dimension 3 haben.
3. Die Reihenfolge von äquivalenten Flächen kann mit dieser Bedingung nicht erkannt werden.

Die Determinantenbedingung ist bereits kompatibel mit äquivalenten Punkten: wir betrachten dazu die vier Punkte a_1, a_2, a_3, a_4 mit $a_2 = a_3$. Die beiden Mengen $\{a_1, a_2, a_4\}$ und $\{a_1, a_3, a_4\}$ liefern dieselben Determinanten, sodass man nur eine der Mengen untersuchen muss. Betrachten wir als nächstes die Menge $\{a_1, a_2, a_3\}$. Wegen $\det(a_2, a_3) = 0$ ist diese Determinante nicht negativ, weshalb sich das Problem auf die beiden Punkte a_1 und a_2 reduziert. Für zwei Punkte ist die Bedingung aus Lemma 4.2 aber trivial. Folglich ist die Determinantenbedingung mit der Existenz äquivalenter Punkte kompatibel.

Als nächstes untersuchen wir die Verallgemeinerung auf Dimension 3, die uns unmittelbar alle höheren Dimensionen erschließen wird. Dazu müssen wir verstehen, inwiefern die Wahl des mathematisch positiven Drehsinns ausgezeichnet war. Wir können die Drehrichtung als alleinige Konsequenz des Vektors $a_1 - p$ ansehen, indem wir uns auf Bemerkung 3.10 besinnen, die jedem Vektor im $\mathbb{R}^2$ eine eindeutige Oberfläche zuordnet.



Dies suggeriert die korrekte Formulierung: Für eine simpliziale Orientierung (p_1, p_2) der Kante und dazugehörige zyklische Reihung $(a_1, \dots, a_m)$ darf für alle $1 < j < k \leq m$ höchstens eine der folgenden Determinanten negativ sein:

$$\begin{aligned} \det(p_2 - p_1, a_1 - p_1, a_j - p_1) \\ \det(p_2 - p_1, a_j - p_1, a_k - p_1) \\ \det(p_2 - p_1, a_k - p_1, a_1 - p_1) \end{aligned}$$

Damit ist natürlich auch die Verallgemeinerung in alle höheren Dimensionen einsichtig.

Zuletzt müssen wir uns um die Reihenfolge der äquivalenten Flächen kümmern, die der Betrachtung durch diese Methode entfliehen. Dabei fällt uns auf, dass wir an keiner Eigenschaft der Einbettung selbst erkennen können, in welcher „Reihenfolge“ die Flächen liegen, da sie alle auf dieselbe Menge abgebildet werden. Daher müssen wir nur darauf achten, dass die Angaben der verschiedenen Fächer an den Rändern der Fläche miteinander konsistent sind. Dies wird aber bereits durch die Kompatibilität der Fächer garantiert. Somit kommen wir zur endgültigen Formulierung der Einbettung:

Definition 4.3 (Falteinbettung). *Seien $\mathcal{F}$ ein n -Faltkomplex und ι eine simpliziale Einbettung von $\mathcal{F}$ (aufgefasst als n -Überlagerungskomplex). Dann heißt ι (**simpliziale**) **Falteinbettung**, falls für jeden Fächer $\nu_{[k]}$ und jede simpliziale Orientierung $\sigma := ([p_1], \dots, [p_n])$ von $[k]$ mit zyklischer Reihenfolge $\nu_{[k]}(\sigma) = (a_1, \dots, a_m)$ gilt:*

Für alle $1 < j < k \leq m$ ist maximal eine der folgenden Determinanten negativ:

$$\begin{aligned} \det(\iota(p_2) - \iota(p_1), \dots, \iota(p_n) - \iota(p_1), \iota(a_1) - \iota(p_1), \iota(a_j) - \iota(p_1)) \\ \det(\iota(p_2) - \iota(p_1), \dots, \iota(p_n) - \iota(p_1), \iota(a_j) - \iota(p_1), \iota(a_k) - \iota(p_1)) \\ \det(\iota(p_2) - \iota(p_1), \dots, \iota(p_n) - \iota(p_1), \iota(a_k) - \iota(p_1), \iota(a_1) - \iota(p_1)) \end{aligned}$$

Bemerkung 4.4. *Für einen 2-Faltkomplex reduzieren sich die Determinanten aus Definition 4.3 auf*

$$\begin{aligned} \det(\iota(p_2) - \iota(p_1), \iota(a_1) - \iota(p_1), \iota(a_j) - \iota(p_1)) \\ \det(\iota(p_2) - \iota(p_1), \iota(a_j) - \iota(p_1), \iota(a_k) - \iota(p_1)) \\ \det(\iota(p_2) - \iota(p_1), \iota(a_k) - \iota(p_1), \iota(a_1) - \iota(p_1)) \end{aligned}$$

Wir haben bereits in Abschnitt 2.5 gesehen, dass es sehr schwer ist, allgemeine Aussagen über Einbettungen zu treffen und werden uns daher auf die Untersuchung einiger einfacher Komplexe beschränken.

Der einfachste Fall ist der, dass sich der 2-Faltkomplex zu einem einzelnen Dreieck falten lässt. Wir werden diesen Fall ausführlich in Abschnitt 4.1 behandeln. Die Ergebnisse aus diesen Überlegungen können wir als notwendige Bedingungen einiger anderer ebenen Einbettungen interpretieren. Diese werden wir in Abschnitt 4.2 untersuchen.

4.1 Überlagerung eines Dreiecks

In diesem Abschnitt untersuchen wir, welche 2-Faltkomplexe sich auf ein Dreieck zusammenfalten lassen und damit einbettbar sind. Wir haben in Abschnitt 2.5.1 nachgewiesen, dass sich alle 3-färbbaren 2-Überlagerungskomplexe zu einem Dreieck erweitern lassen. Diese Aussage gilt aber nicht mehr für 2-Faltkomplexe, da wir dort auch die Überschneidungsfreiheit an den Kantenklassen überprüfen müssen. Wir suchen also nach den 2-Faltkomplexen $\mathcal{F}$, die Teilfaltung eines Dreiecks sind.

Während es zwar interessant sein kann, diese Frage in voller Allgemeinheit zu beantworten, sind wir primär an solchen Faltkomplexen interessiert, die auch unabhängig von dieser Arbeit von Interesse sind. Wir wählen dazu die Untersuchung geschlossener simplizialer Flächen.

Definition 4.5 (Flächenkomplex). *Sei $((\mathcal{K}_0, \mathcal{K}_1, \mathcal{K}_2), \prec)$ ein lokal simplizialer homogener 2-Komplex. Er heißt **Flächenkomplex**, falls es zu jedem $k \in \mathcal{K}_1$ höchstens zwei $f \in \mathcal{K}_2$ mit $k \prec f$ gibt. Dann heißen diese beiden Elemente **benachbart (an k)**. Der Flächenkomplex heißt **geschlossen**, falls es zu jedem $k \in \mathcal{K}_1$ genau zwei $f \in \mathcal{K}_2$ mit $k \prec f$ gibt.*

Mit den Flächenkomplexen aus Definition 4.5 lässt sich jede Triangulation einer Oberfläche beschreiben. In Bemerkung 4.6 weisen wir zudem nach, dass sie sich fast mühelos zu 2-Faltkomplexen ergänzen lassen.

Bemerkung 4.6. *Sei $((\mathcal{K}_0, \mathcal{K}_1, \mathcal{K}_2), \prec)$ ein Flächenkomplex. Für jede simpliziale Oberfläche lässt sich dieser wie folgt zu einem 2-Faltkomplex erweitern:*

- *Die Äquivalenzrelation $\sim$ ist die Gleichheit.*
- *Für $k \in \mathcal{K}_1$ gibt es zwei Fälle: Falls es nur ein $f \in \mathcal{K}_2$ mit $k \prec f$ gibt, ist der Fächer konstant dem 1-Zykel auf diesem Element. Ansonsten ist er konstant gleich dem 2-Zykel, der aus den beiden $f \in \mathcal{K}_2$ besteht, die $k \prec f$ erfüllen.*
- *Für jedes $f \in \mathcal{K}_2$ sind $(f, +1)$ und $(f, -1)$ Randstücke.*

Beweis. Da Gleichheit die Eigenschaften aus der Definition 2.17 von 2-Überlagerungskomplexen erfüllt, haben wir einen solchen vorliegen. Damit müssen wir nur noch die Eigenschaften aus Definition 3.27 von 2-Faltkomplexen nachweisen.

1. Wir müssen prüfen, ob wir tatsächlich Fächer gemäß Definition 3.7 definiert haben. Da alle Zykel aus einem oder zwei Elementen selbstinvers ist, ist das gegeben.
2. Wir haben für jede Flächenäquivalenzklasse genau zwei Randstücke definiert (da alle Klassen einelementig sind).
3. Da jedes mögliche Randstück auch tatsächlich ein Randstück ist, ist die Bedingung über Standardränder automatisch erfüllt.
4. Auch die Komplementarität von Randstücken folgt daraus, dass jedes mögliche Randstück eines ist.

5. Wir müssen nachweisen, dass die Fächer gemäß Definition 3.18 kompatibel sind. Da sich zwei verschiedene Coronae nur in höchstens zwei Elementen schneiden und 2-Zykel selbstinvers sind, ist dies aber klar.

Folglich liegt ein 2-Faltkomplex vor. $\square$

Zu beachten ist, dass die Konstruktion aus Bemerkung 4.6 unabhängig von der simplizialen Oberfläche ist. Damit sich ein solcher 2-Faltkomplex $\mathcal{F}$ auf ein Dreieck zusammenfalten lässt, sind nach Definition 3.50 (Teilfaltung) die folgenden Bedingungen zu überprüfen:

1. $\mathcal{F}$ muss eine Teilüberlagerung des Dreiecks sein. Nach Folgerung 2.61 ist dies äquivalent zur 3-Färbbarkeit von $\mathcal{F}$.
2. Alle Randstücke des Dreiecks müssen bereits in $\mathcal{F}$ Randstücke sein. Da wir uns mit geschlossenen Flächenkomplexen befassen und bei diesen zu Beginn alle orientierten Seiten auch Randstücke sind, ist diese Bedingung trivialerweise erfüllt.
3. Bezüglich jeder zyklische Reihung einer Kantenklasse $[k]$ des Dreiecks sind die Grundmengen der Fächer von $\mathcal{F}$, die zu einem $l \in [k]$ gehören, überschneidungsfrei.
4. Die Redukte der zyklischen Reihungen des Dreiecks auf die korrespondierenden Flächen von $\mathcal{F}$ sollen identisch zu den zyklischen Reihungen von $\mathcal{F}$ sein. Da jede zyklische Reihung von $\mathcal{F}$ aus genau zwei Flächen besteht (und damit selbstinvers ist), ist auch diese Bedingung unmittelbar einsichtig.

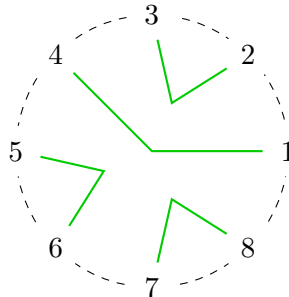
4.1.1 Kreisdarstellung

Wir haben uns im letzten Abschnitt gefragt, welche 3-färbbaren, geschlossenen Flächenkomplexe sich auf ein Dreieck zusammenfalten lassen. Dabei haben wir festgestellt, dass wir nur überprüfen müssen, ob die Grundmengen der Fächer von $\mathcal{F}$ überschneidungsfrei bezüglich der zyklischen Reihungen der Kantenklassen des Dreiecks sind.

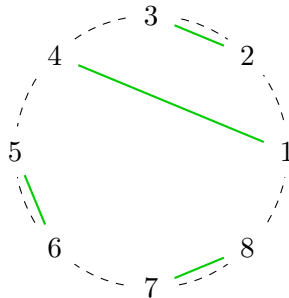
Wir haben also folgende Situation gegeben: Die zweielementigen Mengen M_j sind überschneidungsfrei bezüglich der zyklischen Reihung σ . Wir wollen diese Bedingung etwas handhabbarer (und vor allem algorithmisch angreifbarer) machen. Dazu betrachten wir das Beispiel $\sigma = (1, 2, 3, 4, 5, 6, 7, 8)$ mit den Mengen

$$M_1 = \{1, 4\} \quad M_2 = \{2, 3\} \quad M_3 = \{5, 6\} \quad M_4 = \{7, 8\}$$

Wir stellen uns die Situation graphisch etwa wie folgt vor:



Insbesondere schneiden sich die farbigen Kanten innerhalb des Kreises nicht, also handelt es sich hierbei um einen planaren Graphen (mit Mehrfachkanten). Wir modifizieren die farbigen Kanten durch Homotopien zu Strecken (wobei die gestrichelten Kanten garantieren, dass die farbigen Kanten im Inneren des „Kreises“ bleiben) und erhalten das folgende Bild:



Wir werden solche Darstellungen als Hilfsmittel verwenden, um die Bedingung der Überschneidungsfreiheit in ein gruppentheoretisches Problem zu überführen. Daher müssen wir diese Darstellungen auch formal verstehen.

Definition 4.7 (Kreisdarstellung). *Seien $n \in \mathbb{N}$ und σ ein $2n$ -Zykel auf $\{1, \dots, 2n\}$, sowie $\{M_j\}_{1 \leq j \leq n}$ eine Partition zweielementiger Teilmengen von $\{1, \dots, 2n\}$.*

Sei zudem $\iota : \{1, \dots, 2n\} \rightarrow \{x \in \mathbb{C} \mid \|x\| = 1\}$ mit $\iota(\sigma(m)) = e^{\frac{2\pi i}{2n}} \iota(m)$ für alle $m \in \{1, \dots, 2n\}$ gegeben.

Das Mengensystem $\{S_j \mid 1 \leq j \leq n\} \cup \{Z_k \mid 1 \leq k \leq 2n\}$ mit den Kreisbögen

$$Z_k := \{\iota(k) \cdot e^{\frac{2\pi i}{2n} x} \mid x \in [0, 1]\}$$

und den Strecken

$$S_j := \{\alpha \iota(x_1) + \beta \iota(x_2) \mid x_1, x_2 \in M_j, \alpha, \beta \geq 0, \alpha + \beta = 1\}$$

*heißt **Kreisdarstellung** (von σ bezüglich ι).*

*Der Punkt $\iota(k) \in Z_k$ heißt **Sockel** von Z_k , der Punkt $\iota(\sigma(k)) \in Z_k$ heißt **Spitze** von Z_k .*

Mit der formalen Definition der Kreisdarstellung wollen wir die Überschneidungsfreiheit bezüglich zyklischer Reihungen auf die Überschneidungsfreiheit der Kreisdarstellung zurückführen. Dafür müssen wir zunächst klären, was wir mit der Überschneidungsfreiheit der Kreisdarstellung meinen. Wir meinen damit, dass die Kreisdarstellung ein planarer Graph ist.

Definition 4.8 (überschneidungsfreie Kreisdarstellung). *Sei σ ein $2n$ -Zykel ($n \in \mathbb{N}$) auf $\{1, \dots, 2n\}$ und $\{M_j\}_{1 \leq j \leq n}$ sei eine Partition zweielementiger Teilmengen von $\{1, \dots, 2n\}$.*

Sei zudem $\iota : \{1, \dots, 2n\} \rightarrow \{x \in \mathbb{C} \mid \|x\| = 1\}$ mit $\iota(\sigma(m)) = e^{\frac{2\pi i}{2n}} \iota(m)$ für alle $m \in \{1, \dots, 2n\}$ gegeben.

*Die Kreisdarstellung von σ bezüglich ι heißt **überschneidungsfrei**, wenn je zwei Mengen dieses Mengensystems sich höchstens in $\{\iota(x) \mid x \in \{1, \dots, 2n\}\}$ schneiden.*

Nachdem wir die Überschneidungsfreiheit von Kreisdarstellung auf diese allgemeine Weise definiert haben, zeigen wir in Lemma 4.9, dass die Überprüfung der Strecken genügt, um die Überschneidungsfreiheit sicherzustellen.

Lemma 4.9. *Eine Kreisdarstellung ist genau dann überschneidungsfrei, wenn ihre Strecken paarweise disjunkt sind.*

Beweis. Wir müssen zeigen, dass sich zwei beliebig gewählte Mengen (die nicht beide Strecken sind) höchstens in $\{\iota(x) | x \in \{1, \dots, 2n\}\}$ schneiden.

1. Z_a und Z_b (mit $a \neq b$):

Klar: Es gibt ein $p \in \mathbb{N}_{<2n}$ mit $\sigma^p(a) = b$.

Ein Element aus dem Schnitt $Z_a \cap Z_b$ hat die Form

$$\iota(a) \cdot e^{\frac{\pi i}{n}x} = \iota(b) \cdot e^{\frac{\pi i}{n}y}$$

mit $x, y \in [0, 1]$. Aufgrund der Eigenschaft von ι gilt

$$\begin{aligned} \iota(b) \cdot e^{\frac{\pi i}{n}y} &= \iota(\sigma^p(a)) e^{\frac{\pi i}{n}y} \\ &= e^{\frac{\pi i}{n}p} \iota(a) e^{\frac{\pi i}{n}y}, \end{aligned}$$

d. h. nach Kürzen von $\iota(a)$ bleibt

$$e^{\frac{\pi i}{n}x} = e^{\frac{\pi i}{n}(p+y)}.$$

Folglich gilt $x = p + y$ (modulo $2n$). Für $1 < p < 2n - 1$ gibt es keine Lösung dieser Gleichung, da $x, y \in [0, 1]$. Für $p = 2n - 1$ ist die Gleichung äquivalent zu $x + 1 = y$ (modulo $2n$). Wir betrachten also nur den Fall $x = 1 + y$ (modulo $2n$).

Diese Gleichung hat maximal Lösungen für $x, y \in \{0, 1\}$ mit $x \neq y$, aber damit ist der Schnittpunkt $\iota(a)$ oder $\iota(b)$.

2. Z_a und S_b (mit $1 \leq a \leq 2n$ und $1 \leq b \leq n$):

Da S_b die konvexe Hülle von $\iota(x)$ und $\iota(y)$ mit $M_b = \{x, y\}$ ist, also deren Verbindungsstrecke, und alle Punkte daraus (mit Ausnahme von $\iota(x)$ und $\iota(y)$) eine Norm kleiner als 1 haben, folgt die Behauptung.

Da zwei verschiedene Strecken keine Punkte aus $\{\iota(x) | x \in \{1, \dots, 2n\}\}$ gemeinsam haben, folgt die Behauptung. $\square$

Nachdem wir in Lemma 4.9 nachgewiesen haben, dass die Kreisdarstellungen genau dann planar sind, wenn sich deren Strecken nicht schneiden, können wir im nächsten Lemma 4.10 die Äquivalenz zwischen der Überschneidungsfreiheit bezüglich zyklischen Reihungen und Kreisdarstellungen zeigen.

Lemma 4.10. *Sei σ ein $2n$ -Zykel ($n \in \mathbb{N}$) auf $\{1, \dots, 2n\}$ und $\{M_j\}_{1 \leq j \leq n}$ sei eine Partition zweielementiger Teilmengen von $\{1, \dots, 2n\}$. Dann sind äquivalent*

1. Die M_j sind überschneidungsfrei bezüglich σ .

2. Eine Kreisdarstellung von σ ist überschneidungsfrei.

Beweis. Sei ι wie in Definition 4.7 gegeben. Wir verwenden die Charakterisierung aus Lemma 4.9.

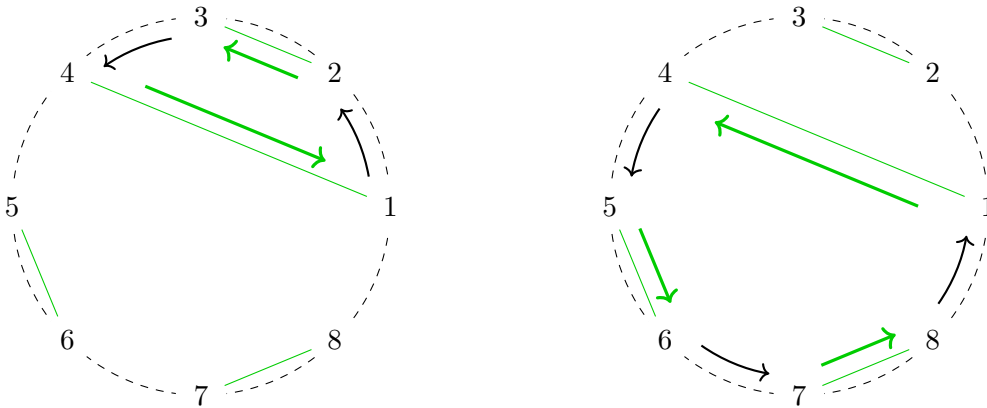
Seien $\overline{\iota(a_1)\iota(a_2)}$ und $\overline{\iota(a_3)\iota(a_4)}$ zwei Strecken. Wir wollen zeigen, dass diese sich genau dann schneiden, wenn $\{a_1, a_2\}$ und $\{a_3, a_4\}$ eine Überschneidung bezüglich σ bilden. Dazu betrachten wir die $\iota(a_i)$ in Polarkoordinaten mit Winkeln φ_i , wobei $\iota(a_1) = 1$, also $\varphi_1 = 0$ gilt. Mit der Eigenschaft von ι folgt, dass $\{a_1, a_2\}$ und $\{a_3, a_4\}$ genau dann eine Überschneidung bezüglich σ bilden, wenn $0 = \varphi_1 < \varphi_3 < \varphi_2 < \varphi_4 < 2\pi$ gilt (dabei sei $\{a_3, a_4\}$ richtig angeordnet).

Der Kreis $\{x \in \mathbb{C} \mid \|x\| = 1\}$ wird von $\overline{\iota(a_1)\iota(a_2)}$ in zwei Zusammenhangskomponenten zerteilt. Deren Schnitte mit dem Kreisrand werden von den Intervallen $[0, \varphi_2]$ und $[\varphi_2, 2\pi]$ induziert. Wir unterscheiden zwei Fälle:

1. Falls $\{\varphi_3, \varphi_4\} \cap [0, \varphi_2]$ einelementig ist, liegen $\iota(a_3)$ und $\iota(a_4)$ in zwei verschiedenen Zusammenhangskomponenten. Da der Kreis konvex ist, liegt ihre Verbindungsstrecke im Inneren des Kreises, muss also die Strecke $\overline{\iota(a_1)\iota(a_2)}$ schneiden.
2. Falls $\{\varphi_3, \varphi_4\} \cap [0, \varphi_2]$ leer oder zweielementig ist, liegen $\iota(a_3)$ und $\iota(a_4)$ in derselben Zusammenhangskomponente. Da jede der Zusammenhangskomponenten konvex ist, liegt auch ihre Verbindungsstrecke komplett in dieser. Insbesondere schneidet sie die Strecke $\overline{\iota(a_1)\iota(a_2)}$ nicht.

Damit ist die Behauptung gezeigt. □

Falls die Kreisdarstellung planar (d. h. überschneidungsfrei) ist, können wir die Facetten (oder Fenster) des Graphen definieren. Abstrakt sind dies einfach die Zusammenhangskomponenten des Komplements in der Ebene $\mathbb{C}$, wir sind aber an einer expliziteren Beschreibung interessiert. Dazu beobachten wir, dass sich in allen beschränkten Facetten die Kantenfarben abwechseln. Wir illustrieren dies in den folgenden beiden Graphiken.



Die schwarzen, gekrümmten Pfeile lassen sich als Anwendung von σ auffassen. Wenn wir die farbigen, geraden Pfeile auch als Anwendung einer Permutation interpretieren wollen,

erhalten wir $(1, 4)(2, 3)(5, 6)(7, 8)$. Es handelt sich um eine fixpunktfreie Involution (d. h. sie ist selbstinvers).

Um die Verbindung zwischen überschneidungsfreien Kreisdarstellungen und fixpunktfreien Involutionen klar zu formulieren, müssen wir uns mit den Eigenschaften des Komplements befassen. Dazu starten wir mit einer überschneidungsfreien Kreisdarstellung und konstruieren die Involution. In der nächsten Bemerkung beginnen wir damit, die Anzahl der Zusammenhangskomponenten des Komplements zu bestimmen.

Bemerkung 4.11. *Sei $\mathcal{M}$ eine überschneidungsfreie Kreisdarstellung des $2n$ -Zykels σ . Dann enthält $\mathbb{C} \setminus \bigcup_{M \in \mathcal{M}} M$ genau $n + 1$ beschränkte Zusammenhangskomponenten.*

Beweis. Wir fassen die Mengen der Kreisdarstellung als Kanten eines planaren Graphen auf. Dies ist wohldefiniert, da sich zwei dieser Kanten nur in den Punkten $\{\iota(k) | 1 \leq k \leq 2n\}$ schneiden können, wobei $\iota : \{1, \dots, 2n\} \rightarrow \mathbb{C}$ die Abbildung ist, bezüglich der die Kreisdarstellung definiert ist.

Aufgrund der Euler-Charakteristik ist die alternierende Summe von Eckenzahl, Kantenzahl und Facettenzahl gleich 1. Es gibt $2n$ Ecken und $2n + n$ Kanten, folglich ist die Anzahl der Facetten (also die Anzahl der beschränkten Zusammenhangskomponenten) gleich $-2n + 3n + 1 = n + 1$. $\square$

Nachdem wir die Anzahl der Zusammenhangskomponenten im Komplement von überschneidungsfreien Kreisdarstellungen bestimmt haben, beschreiben wir diese Komponenten in Lemma 4.12 genauer. Insbesondere wollen wir nachweisen, dass sich die Farben auf dem Rand dieser Komponente immer abwechseln.

Lemma 4.12. *Sei $\mathcal{M}$ eine überschneidungsfreie Kreisdarstellung des $2n$ -Zykels σ bezüglich $\iota : \{1, \dots, 2n\} \rightarrow \mathbb{C}$. Dann gilt für die Zusammenhangskomponenten des Komplements $\mathbb{C} \setminus \bigcup_{M \in \mathcal{M}} M$:*

- *Es gibt genau eine unbeschränkte Zusammenhangskomponente $\{x \in \mathbb{C} | \|x\| > 1\}$ und jedes Kreissegment Z_j liegt auf ihrem Rand.*
- *Es gibt genau $n + 1$ beschränkte Zusammenhangskomponenten. Für jede davon gibt es ein $m \in \mathbb{N}$, sodass ihr Rand aus Mengen*

$$Z_{j_1} \quad S_{k_1} \quad Z_{j_2} \quad S_{k_2} \quad \dots \quad Z_{j_m} \quad S_{k_m}$$

besteht, die die folgenden Eigenschaften erfüllen:

- *Die Schnitte $Z_{j_t} \cap S_{k_t}$ und $S_{k_t} \cap Z_{j_{t+1}}$ sind für alle $1 \leq t \leq m$ einelementig (Indizes sind modulo m zu lesen).*
- *Sei a_t das eindeutige Element aus $S_{k_{t-1}} \cap Z_{j_t}$ und sei b_t das eindeutige Element aus $Z_{j_t} \cap S_{k_t}$, dann gilt $\sigma(a_k) = b_k$.*

Beweis. Da alle Strecken in $\{x \in \mathbb{C} | \|x\| \leq 1\}$ liegen und

$$\bigcup_{j=1}^{2n} Z_j = \{x \in \mathbb{C} | \|x\| = 1\}$$

gilt, ist $\{x \in \mathbb{C} \mid \|x\| > 1\}$ die einzige unbeschränkte Zusammenhangskomponente und jede beschränkte Zusammenhangskomponente ist eine Teilmenge von $\{x \in \mathbb{C} \mid \|x\| < 1\}$.

Sei nun κ eine beschränkte Zusammenhangskomponente. Da die Kreisdarstellung überschneidungsfrei ist, ist ihr Rand eine Vereinigung von Mengen aus der Kreisdarstellung. Wir gehen für den Beweis der weiteren Eigenschaften in den folgenden Schritten vor:

1. Zu jedem $\iota(t)$ aus dem Rand von κ gibt es genau ein Kreissegment und genau eine Strecke aus dem Rand von κ , die $\iota(k)$ enthalten.
2. Dies ermöglicht die alternierende Abfolge der Mengen aus dem Rand, sowie die Schnittbedingung.
3. Wir können die Abfolge so wählen, dass $\sigma(a_1) = b_1$ gilt.
4. Die Gleichung $\sigma(a_t) = b_t$ ist äquivalent zu $\sigma(a_{t+1}) = b_{t+1}$.

Dann folgt die Behauptung per Induktion. Wir gehen die Schritte einzeln durch:

1. Zu jedem $\iota(t)$ gibt es genau drei Mengen der Kreisdarstellung, die $\iota(t)$ enthalten: Z_t , Z_{t+1} und eine Strecke. Jedes Kreissegment liegt im Rand von genau zwei Zusammenhangskomponenten. Da die Menge $\{x \in \mathbb{C} \mid \|x\| \leq 1\}$ von der Strecke in genau zwei Zusammenhangskomponenten zerlegt wird, liegt auch die Strecke im Rand von genau zwei Zusammenhangskomponenten.

Wir wissen, dass Z_t und Z_{t+1} im Rand der unbeschränkten Zusammenhangskomponente liegen. Da die Strecke nicht im Rand der unbeschränkten Komponente liegt, gibt es noch zwei weitere Zusammenhangskomponenten, die die Strecke im Rand haben. Da $\iota(t)$ ein Endpunkt der Strecke ist, müssen sie auch eines der beiden Kreissegmente im Rand haben. Damit kommt jedes Kreissegment in genau zwei Zusammenhangskomponenten vor. Folglich gibt es keine Komponente, für die alle drei Mengen im Rand liegen.

Insbesondere gilt für jede beschränkte Zusammenhangskomponente, dass sie genau ein Kreissegment und genau eine Strecke, die $\iota(t)$ enthalten, im Rand haben.

2. Klar.
3. Entweder gilt $\sigma(a_1) = b_1$ oder $\sigma(b_1) = a_1$. Falls letzteres zutrifft, invertiere die Reihenfolge der Mengen.
4. Es genügt zu zeigen, dass die Endpunkte jeder Strecke jeweils Spitze und Sockel (der angrenzenden Kreissegmente) sind. Jede Strecke zwischen $\iota(x)$ und $\iota(y)$ spaltet die Menge $\{x \in \mathbb{C} \mid \|x\| \leq 1\}$ in zwei Zusammenhangskomponenten, also auch ihren Rand $\{x \in \mathbb{C} \mid \|x\| = 1\}$ in zwei Teiltränder (die durch Hinzufügen der Strecke zum gesamten Rand der Komponente werden).

Parametrisieren wir $\{x \in \mathbb{C} \mid \|x\| = 1\}$ durch

$$\alpha : [0, 1) \rightarrow \{x \in \mathbb{C} \mid \|x\| = 1\}, t \mapsto e^{2\pi i t},$$

so liegen $\iota(x)$ und $\iota(y)$ an den Stellen $t_1 \neq t_2$. Ohne Einschränkung nehmen wir $t_1 < t_2$ an. Dann ist $[t_1, t_2] \subseteq [0, 1)$. In $\alpha([t_1, t_2])$ liegen Kreissegmente, insbesondere solche, die $\iota(x)$ und $\iota(y)$ enthalten. Für jenes, das $\iota(x)$ enthält, ist $\iota(x)$ ein Sockel (da keine kleineren Winkel vorkommen). Dasjenige, das $\iota(y)$ enthält, hat mit analoger Begründung $\iota(y)$ als Spitze.

Folglich (jedes $\iota(t)$ kommt einmal als Spitze und einmal als Sockel vor) gilt die Behauptung auch für die andere Zusammenhangskomponente des Randes.

Zeige nun, dass diese Kreissegmente auch im Rand von κ liegen. Da die Strecke im Rand von κ liegt, muss es noch eine weitere Menge geben, die $\iota(x)$ enthält und im Rand liegt. Da κ in einer der beiden obigen Zusammenhangskomponenten liegt, entfällt das Kreissegment der jeweils anderen Komponente als Möglichkeit. Das zeigt die geforderte Behauptung. $\square$

Nachdem wir in Lemma 4.12 die Zusammenhangskomponenten einer überschneidungsfreien Kreisdarstellung charakterisiert haben, können wir die Korrespondenz zwischen den fixpunktfreien Involutionen und den überschneidungsfreien Kreisdarstellungen nachweisen:

Lemma 4.13. *Sei σ ein $2n$ -Zykel ($n \in \mathbb{N}$) auf $\{1, \dots, 2n\}$ und ρ sei eine fixpunktfreie Involution auf $\{1, \dots, 2n\}$, sodass die Bahnen von ρ auf $\{1, \dots, 2n\}$ eine überschneidungsfreie Partition von $\{1, \dots, 2n\}$ bilden.*

Sei $\mathcal{M}$ eine Kreisdarstellung von σ und BZ bezeichne die Menge der beschränkten Zusammenhangskomponenten von $\mathbb{C} \setminus \bigcup_{M \in \mathcal{M}} M$.

Dann ist die Abbildung

$$BZ \rightarrow \{B \subseteq \{1, \dots, 2n\} \mid B \text{ Bahn von } \rho\sigma\},$$

bei der jede beschränkte Zusammenhangskomponente mit Rand $Z_{j_1}, S_{k_1}, Z_{j_2}, S_{k_2}, \dots, Z_{j_m}, S_{k_m}$ (wie in Lemma 4.12 beschrieben) auf die Menge

$$\bigcup_{t=1}^m Z_{j_t} \cap S_{k_{t-1}}$$

abgebildet wird, eine Bijektion.

Beweis. Die Kreisdarstellung ist aufgrund von Lemma 4.10 überschneidungsfrei. Wir verwenden im Folgenden die Charakterisierung beschränkter Zusammenhangskomponenten aus Lemma 4.12.

Zeige zuerst die Wohldefiniertheit. Ist a_t das eindeutige Element aus $Z_{j_t} \cap S_{k_{t-1}}$, so ist $\sigma(a_t)$ das eindeutige Element aus $Z_{j_t} \cap S_{k_t}$, liegt insbesondere in S_{k_t} . Demnach ist $\rho\sigma(a_t)$ dann der andere Endpunkt von S_{k_t} . Damit liegt die Bahn von a_t in der Menge. Umgekehrt erreichen wir kein Element einer anderen Bahn.

Da jeder Punkt $\iota(t)$ einmal als Spitze und einmal als Sockel eines Kreissegmentes vorkommt und unsere Konstruktion nur Sockel abbildet, wird jeder Punkt höchstens einmal abgebildet. Damit ist die Abbildung injektiv.

Um Surjektivität zu zeigen, betrachten wir eine Bahn B von $\rho\sigma$. Diese enthält ein Element $b \in B$, sodass es ein Kreissegment Z gibt, mit $b \in Z$ als Sockel. Dann wird die eindeutige beschränkte Zusammenhangskomponente, die Z im Rand hat, auf B abgebildet. $\square$

Nachdem wir gezeigt haben, dass zu jeder überschneidungsfreien Kreisdarstellung eine fixpunktfreie Involution korrespondiert, betrachten wir nun die andere Richtung. Wir starten also mit einer fixpunktfreien Involution und zeigen, dass die zugehörige Kreisdarstellung überschneidungsfrei ist. Dies beruht darauf, dass das Produkt $\rho\sigma$ mindestens einen Fixpunkt hat, wie wir in Lemma 4.14 festhalten.

Lemma 4.14. *Sei $\pi \in S_n$ eine Permutation mit mehr als $\frac{n}{2}$ Zykeln. Dann hat π mindestens einen Fixpunkt.*

Beweis. Falls π keinen Fixpunkt hat, enthält jede Bahn von π mindestens zwei Elemente. Da die Bahnen von π auf $\{1, \dots, n\}$ disjunkt sind, folgt

$$n = \sum_{B \text{ Bahn}} |B| \geq \sum_{B \text{ Bahn}} 2 = 2 \cdot \# \text{ Bahnen} > n.$$

Dies ist ein Widerspruch, also muss π einen Fixpunkt haben. $\square$

Dass $\rho\sigma$ mindestens einen Fixpunkt hat, ist deswegen signifikant, da jeder dieser Fixpunkte eine Kante repräsentiert, deren Eckpunkte minimalen Abstand haben. Folglich lässt sich diese Kante aus der Kreisdarstellung entfernen, ohne die Überschneidungsfreiheit der übrigen Kanten zu gefährden. Dass diese Reduktion möglich ist, müssen wir aber formal beweisen. Wir zeigen im folgenden Lemma 4.15 zunächst, dass wir nach Entfernung dieser Kante noch immer einen Fixpunkt vorliegen haben.

Lemma 4.15. *Sei σ ein $2n$ -Zykel auf $\{1, \dots, 2n\}$ und ρ eine fixpunktfreie Involution, sodass $\rho\sigma$ genau $n+1$ Zyklen hat ($n > 1$). Sei $f \in \{1, \dots, n\}$ ein Fixpunkt von $\rho\sigma$ (existiert wegen Lemma 4.14).*

Definiere Permutationen auf $M := \{1, \dots, 2n\} \setminus \{f, \sigma(f)\}$ durch die Redukte

$$\hat{\sigma} := \sigma|_{M \rightarrow M}$$

$$\hat{\rho} := \rho|_{M \rightarrow M}$$

Dann hat $\hat{\rho}\hat{\sigma}$ genau n Zyklen.

Beweis. Um die Aussage nachzuweisen, betrachten wir eine Bahn B von $\rho\sigma$ genauer.

1. $B = \{f\}$:

Dieser Fall entfällt.

2. $\sigma(f) \in B$:

Dann ist $\sigma^{-1}(f) = \sigma^{-1}(\rho^{-1}(\sigma(f))) \in B$. Bestimme nun $\hat{\rho}\hat{\sigma}\sigma^{-1}(f)$. Da $\hat{\sigma}$ das Redukt von σ ist und

$$\begin{aligned}\sigma(\sigma^{-1}(f)) &= f \notin M \\ \sigma^2(\sigma^{-1}(f)) &= \sigma(f) \notin M\end{aligned}$$

gelten, folgt $\hat{\sigma}\sigma^{-1}(f) = \sigma^2(f) \in M$ (da $n > 1$). Schließlich erhalten wir

$$\begin{aligned}\hat{\rho}\hat{\sigma}\sigma^{-1}(f) &= \hat{\rho}\sigma^2(f) \\ \rho\sigma^2(f) &= (\rho\sigma)(\sigma(f)) \in B.\end{aligned}$$

Damit ist $B \setminus \{\sigma(f)\}$ eine Bahn von $\hat{\rho}\hat{\sigma}$.

3. $f, \sigma(f) \notin B$:

Da $\hat{\rho}\hat{\sigma} = \rho\sigma$ auf dieser Bahn, bleibt sie gleich.

Insgesamt verlieren wir eine Bahn. □

Wir haben in Lemma 4.15 gezeigt, dass die Reduktion einer Kante ein Induktionsargument zulässt. Dieses Argument führen wir im folgenden Lemma aus, das die Überschneidungsfreiheit unserer Konstruktion sicherstellt.

Lemma 4.16. *Sei σ ein $2n$ -Zykel auf $\{1, \dots, 2n\}$ und ρ eine fixpunktfreie Involution, sodass $\rho\sigma$ genau $n+1$ Zykel hat. Dann ist jede Kreisdarstellung von σ (wobei die Bahnen von ρ die Partition von $\{1, \dots, 2n\}$ bilden) überschneidungsfrei.*

Beweis. Wir beweisen die Aussage per Induktion nach n .

$n = 1$ Da σ ein 2-Zykel ist und $\rho = \sigma$ gilt, gibt es nur eine einzige Strecke in der Kreisdarstellung. Gemäß Lemma 4.9 ist diese damit überschneidungsfrei.

$n > 1$ Sei $\iota : \{1, \dots, 2n\} \rightarrow \mathbb{C}$ wie in Definition 4.7 gegeben. Wegen Lemma 4.14 hat $\rho\sigma$ einen Fixpunkt f , d. h. $\rho(\sigma(f)) = f$. Ist $\sigma(f) = g$, so gilt $\rho(g) = f$.

Damit kann die Strecke der Kreiseinbettung, die zu dem Zykel (f, g) gehört, keine andere Strecke schneiden (da f und g bezüglich ι benachbart sind). Wenn wir also zeigen können, dass sich die übrigen Strecken nicht schneiden, ist die Kreisdarstellung überschneidungsfrei.

Dazu definieren wir $\hat{\sigma}$ und $\hat{\rho}$ wie in Lemma 4.15. Nach ebensolchem hat ihr Produkt genau n Zykel, weshalb nach Induktionsvoraussetzung deren Kreisdarstellungen überschneidungsfrei sind. gemäß Lemma 4.10 bilden die Bahnen von $\hat{\rho}$ also eine überschneidungsfreie Partition von $\hat{\sigma}$.

Zeige nun, dass eine Überschneidung der Bahnen von ρ bezüglich σ eine Überschneidung der Bahnen von $\hat{\rho}$ bezüglich $\hat{\sigma}$ induziert. Wegen Bemerkung 3.49 genügt es zu zeigen, dass für $M \subseteq \{1, \dots, 2n\} \setminus \{f, g\}$ die Redukte $\sigma|_{M \rightarrow M}$ und $\hat{\sigma}|_{M \rightarrow M}$ identisch sind (da es schon bekannt ist, dass die Strecke zu $\{f, g\}$ keine andere Strecke schneidet). Dies folgt aber sofort aus Lemma 3.15 und der Definition von $\hat{\sigma}$ aus Lemma 4.15. □

Nachdem wir gezeigt haben, dass eine fixpunktfreie Involution eine überschneidungsfreie Kreisdarstellung induziert, können wir das Hauptergebnis dieses Abschnitts in Satz 4.17 verewigen.

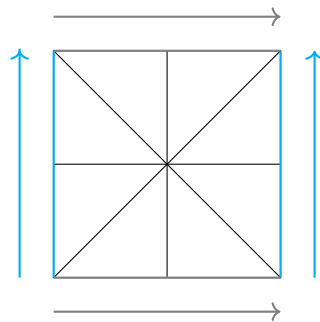
Satz 4.17. *Sei σ ein $2n$ -Zykel ($n \in \mathbb{N}$) auf $\{1, \dots, 2n\}$ und ρ sei eine fixpunktfreie Involution auf $\{1, \dots, 2n\}$. Dann sind äquivalent:*

- Die Bahnen von ρ auf $\{1, \dots, 2n\}$ sind überschneidungsfrei bezüglich σ .
- $\rho\sigma$ hat genau $n + 1$ Bahnen auf $\{1, \dots, 2n\}$.

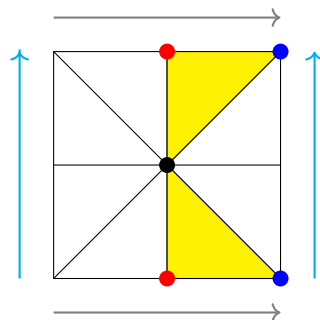
4.1.2 Einfaches Beispiel

Wir haben in Abschnitt 4.1.1 mit Satz 4.17 alle 3-färbbaren geschlossenen simplizialen Flächen charakterisiert, die man auf ein Dreieck zusammenfalten kann. Wir möchten diese abstrakte Charakterisierung in diesem Abschnitt konkret anwenden. Dabei liefert jede Farbklasse von Kanten (bzgl. der 3-Färbung) eine fixpunktfreie Involution ρ . Wir müssen also alle $2n$ -Zykel σ finden, die im Produkt mit jeder Involution genau $n + 1$ Zyklen haben.

Dieses Verfahren führen wir an einem einfachen Beispiel aus, für das wir alle möglichen Faltungen auf ein Dreieck bestimmen werden. Wir betrachten die folgende simpliziale Fläche:



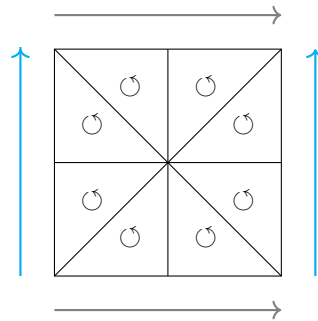
Dabei werden die parallelen Kanten entlang der Pfeilrichtung identifiziert (es handelt sich also um einen topologischen Torus). Diese Fläche kann nicht gemäß Definition 2.53 eingebettet werden. Wir betrachten dazu die markierten Facetten:



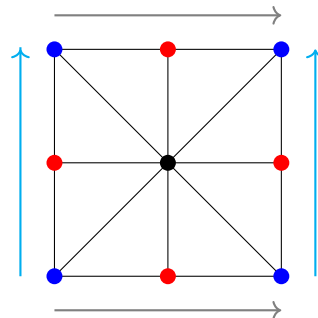
Nach Definition 2.53 müssten die schattierten Facetten auf dieselbe Menge abgebildet werden (da sie dieselben Punkte enthalten). Da die Facetten aber verschieden sind, ist dies unmöglich.

Wenn wir wissen wollen, ob wir sie nach Ausführung von Faltungen einbetten können, überprüfen wir, ob sie sich auf ein Dreieck zusammenfalten lässt. Dazu müssen wir sie als 2-Faltkomplex beschreiben, was wir in Beispiel 4.18 tun.

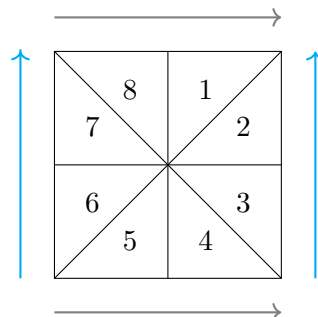
Beispiel 4.18. *Gemäß Bemerkung 4.6 müssen wir nur eine simpliziale Oberfläche angeben, um aus einem Flächenkomplex einen 2-Faltkomplex zu konstruieren. Dazu müssen wir zu jeder Facette eine Oberfläche auszeichnen. Gemäß Bemerkung 3.10 können wir auch eine Umlaufrichtung der Punkte angeben. Wir werden dies durch einen Pfeil andeuten.*



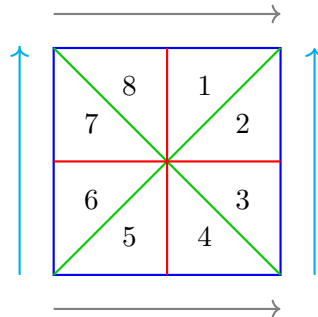
Damit wir diese geschlossene simpliziale Fläche auf ein Dreieck zusammenfalten können, muss sie 3-färbbar sein, wie wir zu Beginn von Kapitel 4.1 festgestellt haben. Eine solche Färbung ist hier bis auf Permutation der Farben eindeutig:



Als nächstes widmen wir uns der Überschneidungsfreiheit. Um die Involutionen konkret angeben zu können, legen wir eine Nummerierung der Facetten fest.

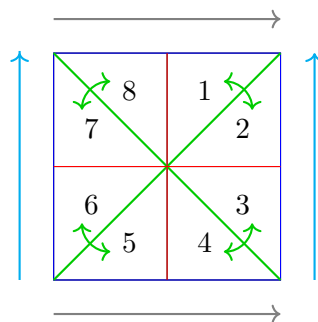


Die Darstellung der verschiedenen Kantenklassen durch eine Punktfärbung ist ein wenig unpraktisch. Daher gehen wir zu der Kantenfärbung über, die von der Punktfärbung induziert wird. Wir erhalten die folgende Struktur:

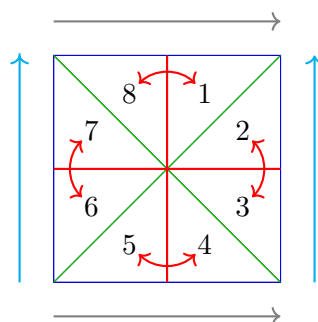


An dieser Stelle merken wir an, dass dieser Übergang zur Kantenfärbung eine MMM-Struktur liefert (vgl. [BNPS] für die genaue Definition und weitere Eigenschaften). Umgekehrt lässt sich aus jeder MMM-Struktur eine gültige 3-Färbung konstruieren.

Betrachten wir nun, welche Flächen entlang der Kantenklasse benachbart sind:



Wir erhalten damit die fixpunktfreie Involution $\rho_1 := (1, 2)(3, 4)(5, 6)(7, 8)$. Aus der Kantenklasse erhalten wir $\rho_2 := (1, 8)(2, 3)(4, 5)(6, 7)$, wie man aus der Graphik



ablesen kann. Schließlich wird die Involution $\rho_3 := (1, 4)(2, 7)(3, 6)(5, 8)$ von der Kantenklasse induziert. Damit liefert Satz 4.17 die folgende Problemstellung:

Finde alle 8-Zykel σ , sodass $\rho_i \cdot \sigma$ für alle $1 \leq i \leq 3$ genau 5 Zykel hat.

Wenn man alle diese 8–Zykel bestimmt (z. B. indem man alle Möglichkeiten ausprobiert), findet man heraus, dass es genau 16 Lösungen gibt, nämlich

(1, 2, 3, 6, 7, 8, 5, 4)	(1, 2, 3, 8, 5, 6, 7, 4)
(1, 2, 7, 6, 3, 4, 5, 8)	(1, 2, 7, 4, 5, 6, 3, 8)
(1, 4, 3, 2, 7, 6, 5, 8)	(1, 4, 3, 2, 5, 8, 7, 6)
(1, 4, 5, 6, 3, 2, 7, 8)	(1, 4, 7, 6, 5, 8, 3, 2)
(1, 4, 5, 8, 7, 6, 3, 2)	(1, 6, 3, 4, 5, 2, 7, 8)
(1, 8, 5, 4, 3, 6, 7, 2)	(1, 6, 7, 8, 5, 2, 3, 4)
(1, 8, 5, 6, 7, 2, 3, 4)	(1, 8, 3, 6, 5, 4, 7, 2)
(1, 8, 7, 2, 3, 6, 5, 4)	(1, 8, 7, 2, 5, 4, 3, 6)

Dabei ist die Anordnung der Permutationen im Kontext der Analyse aus Abschnitt 4.1.3 vor-ausschauend gewählt. Dort werden wir diese Zykel mithilfe einer Gruppenoperation strukturi-
rieren, bevor wir uns in Abschnitt 4.1.4 damit beschäftigen, was genau diese Zykel im Kontext
einer Faltung bedeuten.

4.1.3 Gruppenoperation

Wir haben in Abschnitt 4.1.2 das Beispiel eines Torus betrachtet und alle Möglichkeiten be-
stimmt, wie man diesen auf ein Dreieck zusammenfalten kann. Die 16 Möglichkeiten haben wir
am Ende des Abschnitts aufgelistet. Wir möchten diese Lösungen etwas übersichtlicher struk-
turieren und untersuchen daher einige Operationen, die die Lösungsmenge nicht verändern:

1. Falls $\rho \cdot \sigma$ genau $n + 1$ Zykel hat, dann gilt dies auch für dessen Inverses

$$(\rho\sigma)^{-1} = \sigma^{-1}\rho^{-1} = \sigma^{-1}\rho.$$

Die Konjugation mit ρ ändert die Zykelzahl nicht¹⁴, folglich hat auch

$$\rho(\sigma^{-1}\rho)\rho^{-1} = \rho\sigma^{-1}$$

genau $n + 1$ Zykel. Falls σ eine Lösung ist, ist also auch σ^{-1} eine.

2. Sei g im Zentralisator von ρ , d. h. $g\rho g^{-1} = \rho$. Dann hat

$$g(\rho\sigma)g^{-1} = g\rho g^{-1}g\sigma g^{-1} = \rho(g\sigma g^{-1})$$

ebenfalls genau $n + 1$ Zykel. Damit ist $g\sigma g^{-1}$ eine Lösung, falls σ bereits eine ist.

3. Allgemeiner gilt: Falls die Permutation g die Menge $\{\rho_1, \rho_2, \rho_3\}$ durch Konjugation
stabilisiert, erhält die Konjugation mit g Lösungen.

Wir fassen diese Überlegungen in Bemerkung 4.19 zusammen.

¹⁴Genauer: Der Zykelzähler ist eine Invariante der Konjugation.

Bemerkung 4.19. Seien $\rho_1, \dots, \rho_k \in S_{2n}$ fixpunktfreie Involutionen und

$$\mathcal{L} := \{\sigma \in S_{2n} \mid \sigma \text{ ist } 2n\text{-Zykel}\}$$

die Lösungsmenge. Dann operiert $C_2 \times \text{Stab}_{S_{2n}}(\{\rho_1, \dots, \rho_k\})$ auf $\mathcal{L}$, wobei gilt:

- Die C_2 operiert durch Invertieren.
- $\text{Stab}_{S_{2n}}(\{\rho_1, \dots, \rho_k\})$ ist der mengenweise Stabilisator von $\{\rho_1, \dots, \rho_k\}$ unter Konjugation.

Beweis. Dass die Gruppe auf $\mathcal{L}$ operiert, ist klar (vgl. auch Aufzählung vor der Bemerkung).

Wir zeigen, dass Invertieren und Konjugation miteinander vertauschen. Sei dazu $\sigma \in \mathcal{L}$ und $g \in \text{Stab}_{S_{2n}}(\{\rho_1, \dots, \rho_k\})$, dann gilt

$$(g\sigma g^{-1})^{-1} = g\sigma^{-1}g^{-1},$$

folglich kommutieren Invertieren und Konjugieren. $\square$

Wir möchten die Gruppenoperation aus Bemerkung 4.19 auf unser Beispiel aus Abschnitt 4.1.2 anwenden. Wenn wir dies z. B. mit GAP tun, erkennen wir, dass die beiden Spalten am Ende von Abschnitt 4.1.2 die beiden Bahnen dieser Gruppenoperation sind.

An dieser Stelle (da wir GAP verwenden) ist es sinnvoll, sich kurz über die Berechnung des Stabilisators Gedanken zu machen. Es ist in der Regel einfach, die Gruppe zu bestimmen, die $\{\rho_1, \rho_2, \rho_3\}$ punktweise stabilisiert. Der mengenweise Stabilisator ist aufwendiger zu berechnen. Ein Trick, den man zu Hilfe nehmen kann, ist die Berechnung des *Normalisators*. Wir berechnen also den Stabilisator der Gruppe $\langle \rho_1, \rho_2, \rho_3 \rangle$ in der S_8 und bestimmen den mengenweisen Stabilisator nur in dieser kleineren Gruppe.

Die Gruppe aus Bemerkung 4.19 ist rein algebraisch motiviert und das Nachrechnen der Operation verwendet nur gruppentheoretische Argumente. Da wir die Zyklen aber geometrisch interpretieren (als zyklische Reihungen von Fächern), kann man nach einer geometrischen Interpretation dieser Gruppenoperation fragen.

Dazu starten wir mit dem mengenweisen Stabilisator $\text{Stab}_{S_{2n}}(\{\rho_1, \rho_2, \rho_3\})$. Der Anschauung halber ist es sinnvoll, zuerst eine Untergruppe davon zu untersuchen, nämlich den elementweisen Stabilisator der Menge $\{\rho_1, \rho_2, \rho_3\}$. Wir können diese als Automorphismengruppe der simplizialen Fläche auffassen, die jede Kantenfarbe invariant lässt. Damit ergibt sich auch die Interpretation des mengenweisen Stabilisators: Er enthält auch jene Automorphismen, die Farben vertauschen können.

Wir sollten uns aber darüber im Klaren sein, was wir hier mit „Automorphismen“ meinen, da wir diesen Begriff bislang nicht auf Faltkomplexe angewendet haben. Da wir an späterer Stelle daran interessiert sind, wie ein Automorphismus sich auf der gefalteten Struktur auswirkt, formulieren wir diese so, dass wir die Kenntnis des Zielkomplexes nicht voraussetzen.

Definition 4.20 (induzierter Faltkomplex). Sei

$$\mathcal{F} = ((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim, \{\sigma_f\}_{f \in \mathcal{K}_n}, \{\nu_{[k]}\}_{k \in \mathcal{K}_{n-1}}, \{\partial_{[f]}\}_{f \in \mathcal{K}_n})$$

ein n -Faltkomplex und φ sei ein Komplex-Isomorphismus von $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec)$ in sich selbst. Dann induziert φ einen n -Faltkomplex

$$\varphi(\mathcal{F}) = \left((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim^*, \{\sigma_f^*\}_{f \in \mathcal{K}_n}, \{\nu_{[k]}^*\}_{k \in \mathcal{K}_{n-1}}, \{\partial_{[f]}^*\}_{f \in \mathcal{K}_n} \right)$$

wie folgt:

- Es gilt $x \sim y$ genau dann, wenn $\varphi(x) \sim^* \varphi(y)$ gilt. Wir bezeichnen die $\sim^*$ -Klasse von $x \in \bigcup_{i=0}^n \mathcal{K}_i$ mit $[x]^*$.
- Zu $\sigma_f : \text{SimplOrient}(\mathcal{S}_{\preccurlyeq f}) \rightarrow \{\pm 1\}$ definieren wir

$$\begin{aligned} \sigma_{\varphi(f)}^* : \text{SimplOrient}(\mathcal{S}_{\preccurlyeq \varphi(f)}) &\rightarrow \{\pm 1\} \\ a &\mapsto \sigma_f(\varphi^{-1}(a)). \end{aligned}$$

- Zum Fächer $\nu_{[k]} : \text{SimplOrient}(\mathcal{S}_{\preccurlyeq [k]}) \rightarrow \text{Zyk}(\text{Cor}_{\mathcal{F}}(k))$ definieren wir

$$\begin{aligned} \nu_{[\varphi(k)]^*}^* : \text{SimplOrient}(\mathcal{S}_{\preccurlyeq [\varphi(k)]^*}) &\rightarrow \text{Zyk}(\text{Cor}_{\varphi(\mathcal{F})}(\varphi(k))) \\ t &\mapsto \varphi_n \circ \nu_{[k]}(\varphi^{-1}(t)) \circ \varphi_n^{-1}, \end{aligned}$$

wobei φ_n als Element von $\text{Sym}(\mathcal{K}_n)$ aufzufassen ist.

- Definiere für jedes $f \in \mathcal{K}_n$:

$$\partial_{[\varphi(f)]^*}^* := \{(\varphi(g), \varepsilon_g) \mid (g, \varepsilon_g) \in \partial_{[f]}\}.$$

Wohldefiniertheit. Wir müssen nachweisen, dass wir hier tatsächlich einen n -Faltkomplex gemäß Definition 3.27 konstruiert haben.

Wir zeigen zuerst, dass $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim^*)$ ein n -Überlagerungskomplex ist. Es genügt dazu, die starke Abwärtskompatibilität aus Bemerkung 2.18 nachzuweisen. Seien $a, x, y \in \bigcup_{i=0}^n \mathcal{K}_i$ mit $x \sim^* y$ und $a \preccurlyeq x$ gegeben. Dann gelten $\varphi^{-1}(x) \sim \varphi^{-1}(y)$ und $\varphi^{-1}(a) \preccurlyeq \varphi^{-1}(x)$. Gemäß der starken Abwärtskompatibilität gibt es genau ein $b \in \bigcup_{i=0}^n \mathcal{K}_i$ mit $\varphi^{-1}(a) \sim b$ und $b \preccurlyeq \varphi^{-1}(y)$, folglich hat genau $\varphi(b)$ die Eigenschaften $a \sim^* \varphi(b)$ und $\varphi(b) \preccurlyeq y$.

Als nächstes weisen wir nach, dass die simpliziale Oberfläche und die Fächer wohldefiniert sind. Die Wohldefiniertheit der simplizialen Oberfläche folgt daraus, dass die Operationen von φ und der symmetrischen Gruppe vertauschen. Bei der Wohldefiniertheit der Fächer kommt noch hinzu, dass Invertieren und Konjugation vertauschen.

Um die Eigenschaften eines n -Faltkomplexes aus Definition 3.27 nachzuweisen, benötigen wir eine Hilfsaussage: Ist g der ε_f -Nachbar von f , dann ist $\varphi(g)$ der ε_f -Nachbar von $\varphi(f)$. Seien dazu $f \in \mathcal{K}_n$ und $k \in \mathcal{K}_{n-1}$ mit $k \prec f$ gegeben. Sei $(a_1, \dots, a_{n+1})$ eine simpliziale Orientierung von $\mathcal{S}_{\preccurlyeq f}$, sodass $(a_1, \dots, a_n)$ eine simpliziale Orientierung von $\mathcal{S}_{\preccurlyeq k}$ ist. In der

Notation von Definition 3.21 haben wir:

$$\begin{aligned}
\hat{\nu}_{[\varphi(k)]^*} &:= \hat{\nu}_{[\varphi(k)]^*}^*((\varphi(a_1), \dots, \varphi(a_n))) \\
&= \varphi_n \circ \nu_{[k]}((a_1, \dots, a_n)) \circ \varphi_n^{-1} \\
&= \varphi_n \circ \hat{\nu}_{[k]} \circ \varphi_n^{-1} \\
\hat{\sigma}_{\varphi(f)}^* &:= \sigma_{\varphi(f)}^*((\varphi(a_1), \dots, \varphi(a_{n+1}))) \\
&= \sigma_f((a_1, \dots, a_{n+1})) \\
&= \hat{\sigma}_f
\end{aligned}$$

Folglich gilt

$$\begin{aligned}
\hat{\nu}_{[\varphi(k)]^*}^{\varepsilon_f \cdot \hat{\sigma}_{\varphi(f)}^*}(\varphi(f)) &= (\varphi_n \circ \hat{\nu}_{[k]} \circ \varphi_n^{-1})^{\varepsilon_f \cdot \hat{\sigma}_{\varphi(f)}^*}(\varphi(f)) \\
&= \left(\varphi_n \circ \hat{\nu}_{[k]}^{\varepsilon_f \cdot \hat{\sigma}_{\varphi(f)}^*} \circ \varphi_n^{-1} \right)(\varphi(f)) \\
&= \left(\varphi_n \circ \hat{\nu}_{[k]}^{\varepsilon_f \cdot \hat{\sigma}_f} \right)(f) \\
&= \varphi_n(g).
\end{aligned}$$

Damit folgen die Bedingungen zum Standardrand und zur Komplementarität sofort.

Zum Nachweis der Fächerkompatibilität nutzen wir, dass Konjugation im folgenden Sinne mit der Reduktbildung verträglich ist: Für eine zyklische Reihung $\mu : M \rightarrow M$ und eine Bijektion $\beta : M \rightarrow N$, sowie eine Teilmenge $T \subseteq M$ gilt

$$\beta \circ \mu|_{T \rightarrow T} \circ \beta^{-1} = (\beta \circ \mu \circ \beta^{-1})_{\beta(T) \rightarrow \beta(T)}.$$

Insgesamt ist damit gezeigt, dass $\varphi(\mathcal{F})$ tatsächlich ein n -Faltkomplex ist. $\square$

Da wir in Definition 4.20 beschrieben haben, wie Automorphismen auf n -Faltkomplexe wirken, können wir jetzt untersuchen, wie sich Faltungen unter Automorphismen verhalten.

Bemerkung 4.21. Seien $\mathcal{F}, \mathcal{G}$ zwei n -Faltkomplexe über demselben homogenen n -Komplex $\mathcal{S}$ und $\mathcal{F}$ sei eine Teilfaltung von $\mathcal{G}$. Ist φ ein Komplex-Isomorphismus von $\mathcal{S}$ in sich, dann ist $\varphi(\mathcal{F})$ eine Teilfaltung von $\varphi(\mathcal{G})$.

Gilt insbesondere $\varphi(\mathcal{F}) = \mathcal{F}$ (also falls φ einen Automorphismus induziert), dann ist $\mathcal{F}$ Teilfaltung von $\varphi(\mathcal{G})$.

Beweis. Da sich alle Eigenschaften aus Definition 3.50 der Teilfaltung durch Anwenden von φ übertragen, folgt die Behauptung. $\square$

Wir haben also geklärt, dass wir den Stabilisator aus der Gruppenoperation von Bemerkung 4.19 als Automorphismengruppe des Komplexes interpretieren können. Damit verbleibt die Interpretation des Invertierens. Diese erschließt sich, wenn man sich daran erinnert, dass die 8-Zykel zu Fächern gehören und mit der Definition 3.7 der Fächer vergleicht. Dort ist das Invertieren der zyklischen Reihungen mit einer Änderung der simplizialen Orientierung

verknüpft. Wir können die Invertierungsoperation also als eine Invertierung der „globalen Orientierung des Raums“ auffassen. Diese Perspektive erklärt auch, warum die C_2 immer auf der Lösungsmenge operiert – sie hängt nur vom Fehlen einer kanonischen Orientierung ab, nicht von speziellen Eigenschaften des Komplexes.

4.1.4 Tatsächliche Faltung

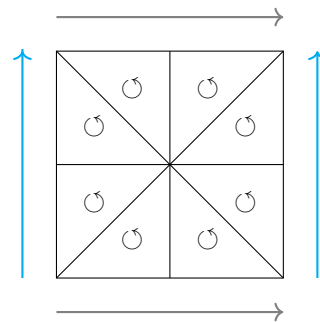
Wir haben in Abschnitt 4.1.2 damit begonnen, alle Faltungen eines konkreten Torus auf ein Dreieck zu bestimmen. In Abschnitt 4.1.3 haben wir die möglichen Faltungen mit einer Gruppenoperation strukturiert und in Bemerkung 4.21 gezeigt, dass wir nur einen Vertreter jeder Bahn untersuchen müssen. In diesem Abschnitt werden wir aus Vertretern der beiden Bahnen, die in unserem Beispiel vorliegen, tatsächliche Faltungen konstruieren. Wir wählen diese beiden Vertreter zur genaueren Untersuchung¹⁵:

$$(1, 4, 3, 2, 7, 6, 5, 8)$$

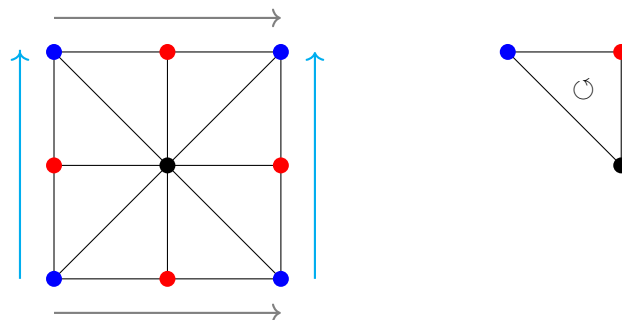
$$(1, 4, 3, 2, 5, 8, 7, 6)$$

Wir wollen aus jedem dieser Zyklen alle 2-Faltkomplexe $\mathcal{G}$ konstruieren, sodass der Torus aus Beispiel 4.18 eine Teilfaltung von $\mathcal{G}$ ist. Dann können wir mit Bemerkung 4.21 alle weiteren möglichen Faltungen durch die Anwendung der Gruppe aus Bemerkung 4.19 bestimmen.

Formulieren wir die Problemstellung präzise: Wir wollen auf Basis des geschlossenen homogenen 2-Komplexes (mit gegebener simplicialer Oberfläche) aus Beispiel 4.18



einen 2-Faltkomplex erzeugen, der nur eine Flächenäquivalenzklasse besitzt und einen der obigen 8-Zyklen als zyklische Reihung hat. Da der Komplex 3-färbbar ist, liefert die Färbungsüberlagerung (Definition 2.58) den gewünschten 2-Überlagerungskomplex.



¹⁵Während die Wahl eigentlich beliebig ist, wählen wir die Vertreter so, dass die resultierenden Faltungen möglichst einfach mit unserer Darstellung des Torus erkennbar sind.

Wir prüfen Definition 3.27 von Faltkomplexen, um herauszufinden, welche Informationen wir noch benötigen, um die Färbungsüberlagerung zu einem 2-Faltkomplex zu ergänzen:

1. Die drei Fächer für die drei Kantenklassen.
2. Die beiden Randstücke.
3. Die Standardränder müssen Randstücke sein.
4. Die Fächer müssen kompatibel sein.

Bearbeiten wir diese Bedingungen Stück für Stück. Die Fächer legen wir fest, indem wir die zyklischen Reihungen bezüglich der induzierten Kantenorientierungen des Dreiecks in der letzten Abbildung angeben. Diese sind alle gleich einem unserer 8-Zykel (womit wir automatisch die Kompatibilität garantiert haben). Dabei spielen die genauen Orientierungen im Torus keine Rolle, da die zyklischen Reihungen dort selbstinvers sind.

Da es nur eine Flächenklasse gibt, haben wir keine Standardränder, weswegen wir bei der Festlegung der beiden Randstücke nur darauf achten müssen, dass diese komplementär sind. Jedes solche Paar liefert dann einen gültigen 2-Faltkomplex $\mathcal{G}$. Da es bei acht Flächen genau 16 Oberflächen gibt und diese paarweise komplementär sind, haben wir acht mögliche Randstück-Paare. Wir möchten diese ungern alle separat berechnen.

Wir zeigen daher, dass wir aus jedem einzelnen solchen 2-Faltkomplex jeden anderen herstellen können. Dazu entfalten wir an den beiden Oberflächen, die später Randstücke werden sollen, und falten die beiden vorherigen Randstücke zusammen. Dass dieser Prozess funktioniert, weisen wir in Lemma 4.22 nach.

Lemma 4.22. *Sei $\mathcal{F}$ ein 2-Faltkomplex mit genau einer Flächenäquivalenzklasse $[f] = \mathcal{K}_2$. Weiter seien die Oberflächen $(a, \varepsilon_a), (b, \varepsilon_b)$ komplementär und keine Randstücke. Dann gilt:*

Die primitive Entfaltung von $\mathcal{F}$ an (a, ε_a) und (b, ε_b) ist ein 2-Faltkomplex, dessen primitive Faltung an $\partial_{[f]}^{\mathcal{F}}$ ein 2-Faltkomplex ist, der sich von $\mathcal{F}$ nur in den Randstücken unterscheidet (dieser hat (a, ε_a) und (b, ε_b) als Randstücke).

Beweis. Da wir von einer primitiven Entfaltung in zwei Mengen sprechen, genügt es, eine einzige Grenzziehung anzugeben, um die Bedingung aus Lemma 3.62 zu erfüllen. Da die beiden neuen Oberflächen komplementär sind, liegen sie in verschiedenen Komponenten. Alle angrenzenden Flächen sind dadurch festgelegt.

Definition 3.31 garantiert dann sofort das Zusammenfallen. □

Gemäß Lemma 4.22 genügt es also, die Faltung für ein einzelnes komplementäres Paar zu bestimmen. Da sich äquivalente Flächen in der Einbettung nicht unterscheiden lassen, haben wir eine Einbettung gefunden, sobald wir ein einzelnes Dreieck eingebettet haben. Damit ist zwar den Formalitäten genüge getan, aber wir würden trotzdem gerne eine Anschauung dafür haben, wie sich die beiden Bahnen voneinander unterscheiden.

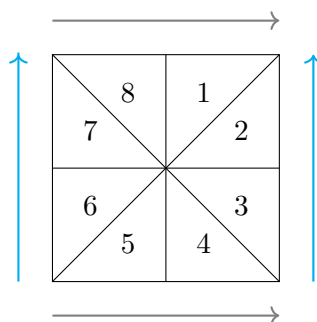
Dazu starten wir mit dem Torusnetz (ohne Identifikationen) und falten dieses. Dabei müssen wir darauf achten, dass am Ende die zu identifizierenden Kanten an denselben Stellen liegen und sich keine Überschneidungen ergeben. Wenn wir aber einen unserer 8-Zykel

nehmen, garantiert uns dieser die Überschneidungsfreiheit. Folglich müssen wir nur eine Faltung konstruieren, die diese beiden Reihenfolgen generiert. Man könnte dazu jedes Dreieck nacheinander falten, verliert aber schnell die Übersicht. Es ist in der Regel sinnvoller, entlang mehrerer Kanten gleichzeitig zu falten. Leider haben wir noch kein allgemeines Verfahren entwickelt, das eine solche Faltfolge bestimmt. Daher begnügen wir uns mit der Angabe der beiden verschiedenen Faltungen, die durch die Zyklen

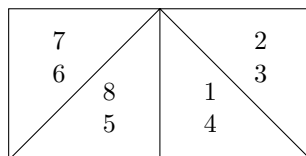
$(1, 4, 3, 2, 7, 6, 5, 8)$

$(1, 4, 3, 2, 5, 8, 7, 6)$

in diesem Beispiel beschrieben werden. Wir starten mit dem Torus

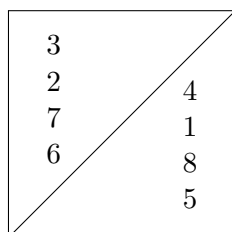


und falten die oberen vier Flächen auf die unteren vier Flächen. Damit erreichen wir den folgenden Zustand,



in dem die äußeren beiden vertikalen Strecken identifiziert werden müssen. Dabei stellen wir uns vor, „von oben“ auf das Muster hinabzuschauen. Die Beschriftungen der Dreiecke geben an, in welcher Reihenfolge die ursprünglichen Dreiecke liegen. Beispielsweise ist der Eintrag des Dreiecks oben links so zu lesen, dass das Dreieck mit der Nummer 7 oben liegt und das mit der Nummer 6 darunter.

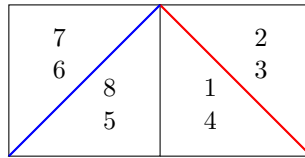
Als nächstes falten wir die rechten beiden Dreiecke auf die linken beiden Dreiecke (die rechten Dreiecke sollen oben liegen). Dabei invertiert sich deren Reihenfolge und wir erhalten



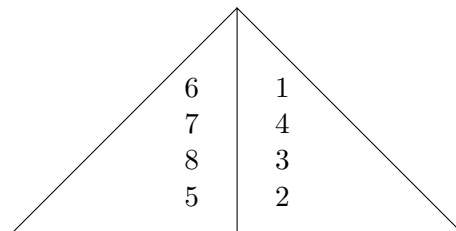
An dieser Stelle sind alle Identifikationen des Torus erfolgt. Es gibt also eine Einbettung dieses Torus mit zwei Dreiecksklassen. Falten wir jetzt das obere linke Dreieck auf das untere rechte

Dreieck, erhalten wir die Flächenabfolge (6, 7, 2, 3, 4, 1, 8, 5), die bis auf Invertieren unserem ersten Zykel entspricht.

Kommen wir nun zum anderen Zykel (1, 4, 3, 2, 5, 8, 7, 6). Wir führen wie vorhin die Faltung an der horizontalen Mittellinie aus und erhalten:



Als nächstes falten wir das obere rechte Dreieck entlang der rechten diagonalen roten Kante *auf* das benachbarte Dreieck und das obere linke Dreieck entlang der linken diagonalen blauen Kante *unter* das benachbarte Dreieck. Damit erhalten wir



Es muss noch eine Identifikation getätigt werden – die ursprünglich vertikalen Kanten liegen beide horizontal (und nicht aufeinander). Wenn wir das linke Dreieck auf das rechte Dreieck falten, erhalten wir genau den geforderten 8–Zykel (5, 8, 7, 6, 1, 4, 3, 2) als Flächenreihenfolge.

Da es genau zwei Bahnen unter der Operation aus Bemerkung 4.19 gibt, sind diese beiden Faltungen (bis auf Automorphie) die einzigen möglichen Faltungen dieses Musters auf ein Dreieck, die eine überschneidungsfreie Identifikation der entsprechenden Kanten ermöglichen.

4.1.5 Maximale Einbettung

Wir haben im letzten Abschnitt gesehen, dass wir den Torus aus Abschnitt 4.1.2 im Wesentlichen auf zwei verschiedene Weisen auf ein Dreieck zusammenfalten können. Eine davon erlaubte sogar eine Einbettung mit zwei Dreiecken. Da allein aufgrund der Randidentifikationen die Flächen aus {1, 4, 5, 8} und {2, 3, 6, 7} paarweise Anomalien bilden, gibt es keine Einbettung mit mehr Dreiecken. Damit liefert die Anzahl der Flächenanomalieklassen eine obere Schranke für die Anzahl der Dreiecke, mit denen man die Struktur einbetten kann. Diese muss natürlich nicht angenommen werden, wie z. B. Faltkomplexe demonstrieren, die sich nicht einbetten lassen. Das einfachste Beispiel hierzu bilden alle anomaliefreien Triangulierungen einer projektiven Ebene. Keine davon kann in den $\mathbb{R}^3$ eingebettet werden, da sie alle topologisch zur projektiven Ebene äquivalent sind.

In diesem Abschnitt skizzieren wir, wie man aus einer gefundenen Einbettung eine mit möglichst vielen Dreiecken konstruiert. Dazu starten wir mit einer bestehenden Einbettung (z. B. aus Abschnitt 4.1.4) und versuchen diese zu entfalten, ohne Anomalien zu erzeugen. Das Verfahren garantiert nicht, dass diese Entfaltung selbst eingebettet werden kann.

Um keine neuen Anomalien einzuführen, lassen wir die Entfaltung eines Faltpans nur dann zu, wenn jede Menge aus der Spaltungspartition der primitiven Entfaltung von Grad 2 eine Vereinigung von Anomalieklassen ist (im Sinne von Definition 2.52).

Betrachten wir exemplarisch die beiden Einbettungen aus Abschnitt 4.1.4. Wir haben die Anomalieklassen $\{1, 4, 5, 8\}$ und $\{2, 3, 6, 7\}$. Wir betrachten die erste Einbettung mit der Flächenreihenfolge

$$5 < 8 < 1 < 4 < 3 < 2 < 7 < 6$$

Man erkennt sofort die mögliche Spaltung in $5 < 8 < 1 < 4$ und $3 < 2 < 7 < 6$, die mit den Anomalieklassen kompatibel ist. Man sieht auch, dass diese Eigenschaft nicht unter der Anwendung von Lemma 4.22 invariant ist, da eine alternative Reihenfolge

$$1 < 4 < 3 < 2 < 7 < 6 < 5 < 8$$

nicht aufgespalten werden kann.

Die andere Einbettung hat die lineare Ordnung

$$2 < 3 < 4 < 1 < 6 < 7 < 8 < 5,$$

und man erkennt, dass sich die Anomalieklassen so abwechseln, dass wir keine Einbettung dieser Äquivalenzklasse entfalten können.

4.1.6 Orientierbarkeit

Wir haben in Abschnitt 4.1.1 herausgefunden, auf welche Weisen man einen geschlossenen Flächenkomplex auf ein Dreieck zusammenfalten kann. Diese abstrakte Charakterisierung haben wir in Abschnitt 4.1.2 am Beispiel eines Torus konkret ausgeführt. Während wir uns in den Abschnitten 4.1.3, 4.1.4 und 4.1.5 mit weiteren Eigenschaften der Lösungen beschäftigt haben, wollen wir uns jetzt mit der Bestimmung dieser Lösungen befassen.

Um die gesuchten Zykel σ zu bestimmen, ist die Gruppentheorie nicht gut ausgestattet. Wir möchten daher die Kreisdarstellung verwenden, um die Berechnung der σ ein wenig zu beschleunigen. Bevor wir unseren Algorithmus in Abschnitt 4.1.7 vorstellen, müssen wir uns aber näher mit der Struktur der Kreisdarstellung auseinandersetzen, da diese zentral für den Algorithmus sein wird. Das Hauptergebnis dieses Abschnitts wird sein, dass nur orientierbare Flächenkomplexe auf ein Dreieck zusammengefaltet werden können.

Die grundlegende Idee ist, alle möglichen überschneidungsfreien Kreisdarstellungen direkt am Kreis zu konstruieren. Dabei fällt auf, dass wir nicht jeden Punkt aus $\{e^{2\pi i \frac{k}{2n}} | 0 \leq k < 2n\}$ mit jedem anderen verbinden können, was wir in Bemerkung 4.23 festhalten.

Bemerkung 4.23. Sei $\overline{xy}$ eine Strecke aus einer überschneidungsfreien Kreisdarstellung bezüglich $\iota : \{1, \dots, 2n\} \rightarrow \mathbb{C}$. Dann gilt $y = e^{2\pi i \frac{k}{2n}} x$ mit ungeradem k .

Beweis. Nach Voraussetzung an ι gibt es ein $k \in \mathbb{Z}$ mit $y = e^{2\pi i \frac{k}{2n}} x$.

Die Strecke zerlegt den Kreis in genau zwei Zusammenhangskomponenten. Auf dem Rand von einer davon liegen genau die folgenden Punkte aus dem Bild von ι :

$$\{e^{2\pi i \frac{l}{2n}} x | 0 \leq l \leq k\}$$

Je zwei davon müssen durch eine Strecke verbunden sein, daher muss die Anzahl der Elemente dieser Menge gerade sein. Daher muss k ungerade sein. $\square$

Damit verbinden Strecken in einer Kreisdarstellung nur Punkte, zwischen denen ein Winkel $\frac{\pi \cdot k}{n}$ mit ungeradem k liegt. Man kann sich fragen, ob jedes beliebige k , das $y = e^{2\pi i \frac{k}{2n} x}$ erfüllt, ungerade sein muss. Da wir stets eine gerade Zahl von Punkten (nämlich $2n$) vorliegen haben, ist das der Fall. Folglich können wir in Definition 4.24 von einem ungeraden Abstand sprechen, auch wenn wir keinen Abstandsbegriff auf dem Kreis definiert haben.

Definition 4.24 ((un)gerader Abstand). Sei $\iota : \{1, \dots, 2n\} \rightarrow \mathbb{C}$ wie in Definition 4.7 gegeben. Für zwei Elemente $a, b \in \{1, \dots, 2n\}$ gibt es ein $k \in \mathbb{Z}$ mit

$$\iota(y) = e^{2\pi i \frac{k}{2n}} \iota(x).$$

Falls k gerade ist, haben a und b **geraden Abstand (bezüglich ι)**, andernfalls **ungeraden Abstand (bezüglich ι)**.

Beweis. Das k ist bis auf Vielfache von $2n$ eindeutig festgelegt. Die Addition mit einer geraden Zahl ändert den Rest modulo 2 nicht. $\square$

Sobald man den Begriff eines geraden Abstands vorliegen hat, lässt sich leicht sehen, dass die Mengen der Elemente mit geradem Abstand zueinander Äquivalenzklassen bilden. Diese Erkenntnis halten wir in Bemerkung 4.25 fest.

Bemerkung 4.25. Sei $\iota : \{1, \dots, 2n\} \rightarrow \mathbb{C}$ wie in Definition 4.7 der Kreisdarstellung gegeben. Dann gibt es eine Äquivalenzrelation mit genau zwei Äquivalenzklassen auf $\{1, \dots, 2n\}$, die wie folgt definiert ist:

- a und b heißen genau dann äquivalent, wenn sie geraden Abstand bezüglich ι haben.

Wir wenden dieser Erkenntnisse in Folgerung 4.26 auf die Situation aus Abschnitt 4.1.1 an, wo wir zu drei fixpunktfreien Involutionen $\{\rho_1, \rho_2, \rho_3\}$ eine Kreisdarstellung suchen, die bezüglich jedes ρ_j überschneidungsfrei ist.

Folgerung 4.26. Sei eine Kreisdarstellung bezüglich $\iota : \{1, \dots, 2n\} \rightarrow \mathbb{C}$ gegeben, die für jede fixpunktfreie Involution in $\{\rho_1, \dots, \rho_m\}$ mit $\rho_j \in S_{2n}$ überschneidungsfrei ist.

Falls für $a, b \in \{1, \dots, 2n\}$ die Beziehung $\rho_j(a) = b$ (für ein $1 \leq j \leq m$) gilt, dann haben a und b ungeraden Abstand bezüglich ι .

Da a und $\rho_j(a)$ niemals in derselben Äquivalenzklasse liegen (für alle $a \in \{1, \dots, 2n\}$ und $1 \leq j \leq m$), lassen sich die Äquivalenzklassen aus den ρ_j rekonstruieren, falls $\langle \rho_1, \dots, \rho_m \rangle$ transitiv auf $\{1, \dots, 2n\}$ operiert. Insbesondere hängen die Äquivalenzklassen dann nur von den gruppentheoretischen Eigenschaften der ρ_j und nicht von der konkreten Abbildung ι ab.

Wenn es also nicht möglich ist, $\{1, \dots, 2n\}$ so in zwei Mengen zu zerlegen, dass a und $\rho_j(a)$ für alle $a \in \{1, \dots, 2n\}$ und $1 \leq j \leq m$ in verschiedenen Äquivalenzklassen liegen (was wir allein mithilfe der ρ_j überprüfen können), kann es keine Kreisdarstellung geben, die bezüglich aller ρ_i überschneidungsfrei ist.

Wir versuchen, diese gruppentheoretische Bedingung geometrisch zu verstehen. Die Voraussetzung dafür, dass die ρ_j die Äquivalenzklassen eindeutig festlegen, ist die Transitivität der Operation von $\langle \rho_1, \dots, \rho_m \rangle$ auf $\{1, \dots, 2n\}$. Wir starten damit, diese Transitivität geometrisch zu interpretieren. Da die ρ_i durch die Kantenklassen der Fläche induziert werden, operiert $\langle \rho_1, \rho_2, \rho_3 \rangle$ genau dann transitiv auf den Flächen, wenn es zu je zwei Flächen f und g eine Folge von Flächen

$$f = f_1, f_2, \dots, f_m = g$$

gibt, sodass f_i und f_{i+1} für alle $1 \leq i < m$ eine Kante gemeinsam haben. Damit stoßen wir auf das Konzept des Wegzusammenhangs.

Definition 4.27 (Wegzusammenhang). *Sei $\mathcal{S}$ ein n -Überlagerungskomplex. Eine Folge*

$$f_1, f_2, \dots, f_m$$

*mit $f_j \in \mathcal{K}_n / \sim$ heißt **Weg (zwischen f_1 und f_m)**, falls $\mathcal{S}_{\preceq f_j}$ und $\mathcal{S}_{\preceq f_{j+1}}$ für alle $1 \leq j < m$ eine $(\mathcal{K}_{n-1} / \sim)$ -Klasse gemeinsam haben.*

*Falls es zu je zwei $f, g \in \mathcal{K}_n / \sim$ einen Weg zwischen f und g gibt, heißt $\mathcal{S}$ **wegzusammenhängend**.*

Wir bemerken, dass zwei Flächen, die nur eine Ecke, aber keine Kante gemeinsam haben, noch nicht automatisch wegzusammenhängend sind. Das Konzept des Wegzusammenhangs unterscheidet sich daher vom üblichen Zusammenhangsbegriff simplizialer Komplexe.

Wir haben es geschafft, die Transitivität der Operation von $\langle \rho_1, \dots, \rho_m \rangle$ auf $\{1, \dots, 2n\}$ geometrisch als Wegzusammenhang des Flächenkomplexes zu interpretieren. Als nächstes möchten wir die Äquivalenzrelation der geraden Abstände aus Bemerkung 4.25 auf geschlossene Flächen übertragen. Wir suchen alle wegzusammenhängenden geschlossenen Flächenkomplexe, deren Flächen in zwei Äquivalenzklassen zerfallen, wobei zwei Elemente derselben Klasse niemals benachbart sind. In Lemma 4.28 beschreiben wir diese Flächenkomplexe durch eine Bedingung an die Wege des Flächenkomplexes.

Lemma 4.28. *Sei $\mathcal{F}$ ein wegzusammenhängender geschlossener Flächenkomplex. Dann sind die folgenden Aussagen äquivalent:*

1. *Es gibt eine Äquivalenzrelation mit genau zwei Äquivalenzklassen auf $\mathcal{K}_2$, sodass benachbarte Flächen nicht äquivalent sind.*
2. *Für je zwei Flächen $f, g \in \mathcal{K}_2 / \sim$ haben die Wege zwischen diesen Flächen entweder alle gerade oder alle ungerade Länge.*

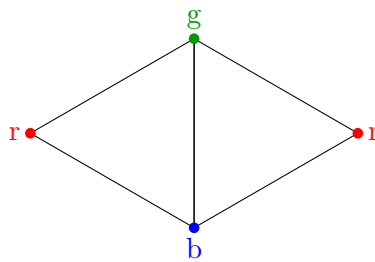
Beweis. Zeige zuerst die Richtung (2) $\Rightarrow$ (1). Definiere die Äquivalenzrelation wie folgt: f und g seien äquivalent, wenn alle Wege zwischen ihnen gerade Länge haben. Da benachbarte Flächenklassen damit in verschiedenen Äquivalenzklassen liegen, gibt es mindestens zwei Äquivalenzklassen. Da ein Weg entweder gerade oder ungerade Länge hat, gibt es maximal zwei Äquivalenzklassen.

Zeige nun die andere Richtung. Seien M_0 und M_1 die beiden Äquivalenzklassen. Wähle ein beliebiges $f \in M_0$. Dann gilt nach Annahme, dass alle $g \in \mathcal{K}_2$, die zu f benachbart sind,

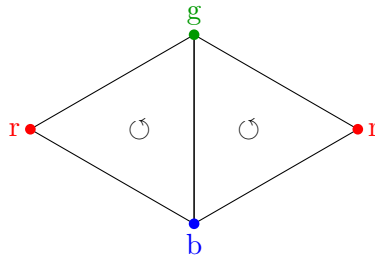
in M_1 liegen (es gibt einen Weg der Länge 1 zwischen ihnen). Iterativ zeigt man, dass ein $h \in \mathcal{K}_2$, zu dem es einen Weg gerader Länge gibt, in M_0 liegt. Wenn der Weg ungerade Länge hat, liegt es in M_1 . Da h nicht in beiden Klassen liegen kann, müssen die Wege zwischen f und h entweder alle gerade oder alle ungerade Länge haben. $\square$

Da die Zerlegung der Flächen in zwei Äquivalenzklassen, bei denen benachbarte Flächen nicht äquivalent sind, eine notwendige Bedingung für die Existenz einer überschneidungsfreien Kreisdarstellung ist, haben wir diese Bedingung mit Lemma 4.28 geometrisch interpretiert. Allerdings ist es sehr aufwendig, diese Bedingung zu überprüfen, da wir *alle* Wege untersuchen müssen. Daher versuchen wir, diese Eigenschaft prägnanter zu formulieren.

Dazu betrachten wir zwei benachbarte Dreiecke in der Ebene, zusammen mit ihrer 3-Färbung.



Wir suchen nach einer Eigenschaft, die auf jedem der Dreiecke genau zwei Werte annehmen kann und bei beiden Dreiecken verschieden ist. Eine solche Eigenschaft ist durch die induzierte Orientierung der Ebene gegeben, die wir innerhalb jedes Dreiecks als zyklische Reihenfolge der Ecken verstehen können.



Die simpliziale Orientierung des linken Dreiecks induziert den 3-Zykel (r, b, g) , während die simpliziale Orientierung des rechten Dreiecks den 3-Zykel (r, g, b) induziert. Diese beiden Zyklen sind invers zueinander. Da es genau zwei mögliche 3-Zyklen in einer Menge von drei Elementen gibt, erfüllt die Wahl dieser Zyklen unsere geforderten Bedingungen.

Wir müssen uns also klar werden, was es bedeutet, die simplizialen Orientierungen der Dreiecke in dieser Weise zu wählen. Wenn die simplizialen Orientierungen von zwei benachbarten Dreiecken immer in dieser Art vorliegen, haben wir die gesamte Fläche orientiert. Den Begriff der Orientierbarkeit führen wir in Definition 4.29 ein.

Definition 4.29 (Orientierbarkeit). Sei $\mathcal{S} = ((\mathcal{K}_0, \mathcal{K}_1, \mathcal{K}_2), \prec)$ ein simplizialer Flächenkomplex. Eine simpliziale Oberfläche $\{\sigma_f\}_{f \in \mathcal{K}_2}$ heißt **Orientierung von $\mathcal{S}$** , falls für jedes Paar benachbarter Flächen $f, g \in \mathcal{K}_2$ und $k \in \mathcal{K}_1$ mit $k \prec f$ und $k \prec g$ folgendes gilt:

- Es gibt simpliziale Orientierungen $(a_1, a_2, a_3) \in \text{SimplOrient}(\mathcal{S}_{\preccurlyeq f})$ und $(b_1, b_2, b_3) \in \text{SimplOrient}(\mathcal{S}_{\preccurlyeq g})$ mit $\sigma_f((a_1, a_2, a_3)) = 1 = \sigma_g((b_1, b_2, b_3))$, sodass (a_1, a_2) und (b_1, b_2) verschiedene simpliziale Orientierungen von $\mathcal{S}_{\preccurlyeq k}$ sind.

Falls eine solche Orientierung existiert, heißt $\mathcal{S}$ **orientierbar**, sonst **nicht orientierbar**.

In dieser Definition setzen wir benachbarte simpliziale Orientierungen dadurch in Relation, dass wir sie auf der verbindenden Kante untersuchen. Man kann sie auch dadurch festlegen, dass man die Oberflächen benachbarter Flächen in Relation setzt. Dass unsere Definition auch diese Situation beschreibt, zeigt Folgerung 4.30:

Folgerung 4.30. Sei $\mathcal{S} = ((\mathcal{K}_0, \mathcal{K}_1, \mathcal{K}_2), \prec)$ ein Flächenkomplex mit Orientierung $\{\sigma_f\}_{f \in \mathcal{K}_2}$. Dann gilt für benachbarte Flächen $f, g \in \mathcal{K}_2$, dass $(f, +1)$ und $(g, +1)$ komplementär sind, ebenso wie $(f, -1)$ und $(g, -1)$ (im Faltkomplex, der in Bemerkung 4.6 bezüglich dieser Orientierung definiert ist).

Beweis. Sei $k \in \mathcal{K}_1$ mit $k \prec f$ und $k \prec g$. Seien simpliziale Orientierungen $(k_1, k_2, a) \in \text{SimplOrient}(\mathcal{S}_{\preccurlyeq f})$ und $(k_1, k_2, b) \in \text{SimplOrient}(\mathcal{S}_{\preccurlyeq g})$ gegeben, sodass (k_1, k_2) eine simpliziale Orientierung von $\mathcal{S}_{\preccurlyeq k}$ ist. Gemäß Definition 3.24 der Komplementarität müssen wir nachweisen, dass $\sigma_f((k_1, k_2, a)) \neq \sigma_g((k_1, k_2, b))$ gilt. Dies ist aber klar: Gälte hier Gleichheit, so müssten nach Definition 4.29 (k_1, k_2) und (k_1, k_2) verschiedene Orientierungen von $\mathcal{S}_{\preccurlyeq k}$ sein. $\square$

Wir haben also festgestellt, dass wir die Weglängenbedingung aus Lemma 4.28 auch dadurch beschreiben können, dass der Flächenkomplex orientierbar ist. Das halten wir in Satz 4.31 fest.

Satz 4.31. Sei $\mathcal{F}$ ein geschlossener Flächenkomplex, der sich auf ein Dreieck zusammenfalten lässt. Dann ist $\mathcal{F}$ orientierbar.

Beweis. Wir betrachten ohne Einschränkung jede Wegzusammenhangskomponente separat. Auf jeder solchen operiert die Gruppe $\langle \rho_1, \rho_2, \rho_3 \rangle$, die von den fixpunktfreien Involutionen der Kanten erzeugt wird, transitiv auf den Flächen der Komponente. Da sich der Flächenkomplex auf ein Dreieck zusammenfalten lässt, gibt es gemäß Folgerung 4.26 eine Äquivalenzrelation mit genau zwei Äquivalenzklassen auf den Flächen, bei der benachbarte Flächen niemals äquivalent sind.

Wir zeigen, dass eine solche Äquivalenzrelation eine Orientierung des Flächenkomplexes induziert. Wir legen für eine Fläche einen 3-Zykel der Farbmenge fest. Jeder Fläche mit geradem Abstand von dieser Fläche ordnen wir denselben Zykel zu, allen anderen dessen Inverses (gemäß Lemma 4.28 ist diese Konstruktion wohldefiniert).

Wir konstruieren aus dieser Zuordnung eine simpliziale Oberfläche, die wir später als Orientierung erkennen werden. Sei dazu $f \in \mathcal{K}_2$ eine Fläche und z der gewählte 3-Zykel für f . Jede simpliziale Orientierung $a \in \text{SimplOrient}(\mathcal{F}_{\preccurlyeq f})$ induziert einen 3-Zykel auf der Eckenmenge $\{p \in \mathcal{K}_0 \mid p \preccurlyeq f\}$, also auch auf der Farbmenge (indem jeder Punkt durch seine Farbe ersetzt wird). Wir definieren die Abbildung σ_f der simplizialen Oberfläche dann wie folgt:

$$\sigma_f : \text{SimplOrient}(\mathcal{F}_{\preccurlyeq f}) \rightarrow \{\pm 1\} \quad a \mapsto \begin{cases} +1 & a \text{ induziert } z \\ -1 & a \text{ induziert } z^{-1} \end{cases}$$

Dadurch ist eine simpliziale Oberfläche $\{\sigma_f\}_{f \in \mathcal{K}_2}$ definiert.

Wir zeigen jetzt, dass diese simpliziale Oberfläche auch eine Orientierung des Flächenkomplexes ist. Seien also zwei Flächen $f, g \in \mathcal{K}_2$ gegeben, die entlang der Kante $k \in \mathcal{K}_1$ benachbart sind. Die Eckpunkte der Kante seien mit den Farben α und β gefärbt, die dritte Farbe sei γ . Ohne Einschränkung sei f der 3-Zykel z zugeordnet, der $z(\alpha) = \beta$ erfüllt. Damit gibt es eine simpliziale Orientierung $a \in \text{SimplOrient}(\mathcal{F}_{\preceq f})$ mit $\sigma_f(a) = 1$, die zum Tripel (α, β, γ) korrespondiert. Da g der 3-Zykel z^{-1} zugeordnet ist, gibt es für diesen ebenso eine simpliziale Orientierung $b \in \text{SimplOrient}(\mathcal{F}_{\preceq g})$ mit $\sigma_g(b) = 1$, die das Tripel (β, α, γ) induziert. Folglich ist $\{\sigma_f\}_{f \in \mathcal{K}_2}$ eine Orientierung des Flächenkomplexes. $\square$

Wir werden die strukturellen Einsichten dieses Abschnitts als Voraussetzungen im Algorithmus des nächsten Abschnitts 4.1.7 verwenden.

4.1.7 Algorithmisches Vorgehen

Wir haben im letzten Abschnitt damit begonnen, strukturelle Eigenschaften der Kreisdarstellung zu bestimmen. Unser Ziel war die Konstruktion eines Algorithmus, der die Lösungen für das Problem aus Abschnitt 4.1.1 auf Basis der Kreisdarstellung bestimmt.

Das zentrale Resultat aus Abschnitt 4.1.6 ist Satz 4.31, der aussagt, dass ein geschlossener Flächenkomplex, der sich auf ein Dreieck falten lässt, orientierbar sein muss. Damit haben wir neben der 3-Färbbarkeit eine weitere notwendige Bedingung, die wir überprüfen können, bevor wir mit der Lösungssuche beginnen.

Damit die Orientierbarkeit mehr als eine geometrische Eigenschaft ist, müssen wir uns genauer mit der Problemformulierung aus Abschnitt 4.1.1 befassen. In dieser sind fixpunktfreie Involutionen $\rho_1, \rho_2, \rho_3 \in S_{2n}$ gegeben, zu denen wir einen $2n$ -Zykel σ suchen, dessen Kreisdarstellung bezüglich jedes ρ_j mit $1 \leq j \leq 3$ überschneidungsfrei ist.

Wir haben im letzten Abschnitt eingesehen, dass die Orientierbarkeit des Flächenkomplexes gleichbedeutend zur Existenz einer Äquivalenzrelation mit genau zwei Äquivalenzklassen auf den Flächen $\{1, \dots, 2n\}$ ist, sodass benachbarte Flächen nie äquivalent sind. Diese Äquivalenzklassen können wir daher allein aus der Kenntnis der ρ_j bestimmen. Da es sich bei dieser Äquivalenzrelation um die Relation „haben geraden Abstand“ aus Bemerkung 4.25 handelt, müssen die Elemente aus den beiden Äquivalenzklassen in jedem $2n$ -Zykel σ alternieren.

Um diese Erkenntnis nutzbar zu machen, verwenden wir die Gruppenoperation auf dem Lösungsraum aus Abschnitt 4.1.3. Dort hatten wir gesehen, dass der mengenweise Stabilisator $\text{Stab}_{S_{2n}}(\{\rho_1, \rho_2, \rho_3\})$ (den wir als eine Automorphismengruppe erkannt haben) auf dem Lösungsraum durch Konjugation operiert. Diese Gruppe operiert auch auf den Flächen $\{1, \dots, 2n\}$ durch Anwenden. Da benachbarte Flächen unter Automorphismen benachbart bleiben, bleiben die beiden Äquivalenzklassen unter dieser Operation fest oder werden vertauscht (sie bilden also ein Blocksystem dieser Operation).

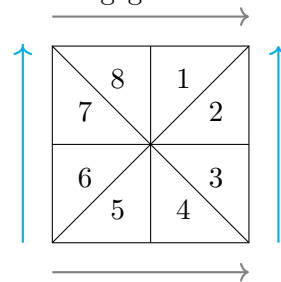
Wir betrachten jetzt den Stabilisator eines Blocks und bestimmen die Bahnen dieses Stabilisators auf den n -Zykeln dieses Blocks (wobei wir auch Invertieren zulassen). Wenn wir einen Vertreter aus einer Bahn wählen, legt dieser bereits die Hälfte der Punkte in der Kreisdarstellung von σ fest. Auf Basis der Überschneidungsfreiheit der Kreisdarstellung bestimmen

wir die übrigen Punkte mit einem Backtrack-Algorithmus, den wir an unserem Beispiel aus Abschnitt 4.1.2 demonstrieren. Dort hatten wir drei Involutionen gegeben:

$$\rho_1 := (1, 2)(3, 4)(5, 6)(7, 8)$$

$$\rho_2 := (1, 4)(2, 7)(3, 6)(5, 8)$$

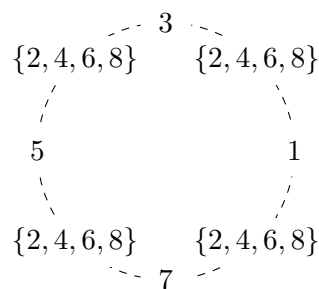
$$\rho_3 := (1, 8)(2, 3)(4, 5)(6, 7)$$



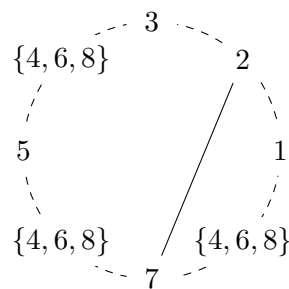
Mit unserer Namensgebung entsprechen die Blöcke den geraden Zahlen $\{2, 4, 6, 8\}$ und den ungeraden Zahlen $\{1, 3, 5, 7\}$. Der Blockstabilisator ist die Diedergruppe mit 8 Elementen, die von den Zykeln $(1, 3, 5, 7)$ und $(3, 7)$ erzeugt wird. Es gibt genau sechs 4-Zykel auf $\{1, 3, 5, 7\}$, die unter der Operation des Blockstabilisators in zwei Bahnen zerfallen:

$$\{(1, 3, 5, 7), (1, 7, 5, 3)\} \quad \{(1, 3, 7, 5), (1, 7, 3, 5), (1, 5, 3, 7), (1, 5, 7, 3)\}$$

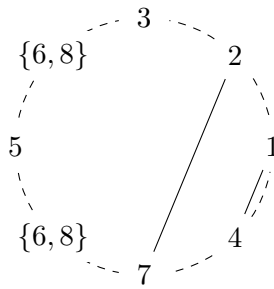
Wir starten mit dem Zykel $(1, 3, 5, 7)$, der die folgende unvollständige Kreisdarstellung liefert:



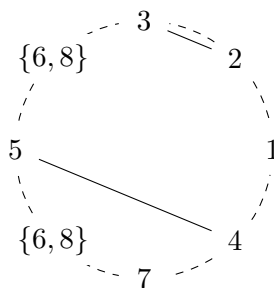
Dabei stellen die Mengen dar, dass an diesen Stellen jedes Element der Menge stehen könnte. Wir betrachten den Fall, dass zwischen 1 und 3 (also im Nord-Osten) eine 2 steht und tragen die Strecken der Involution $(1, 4)(2, 7)(3, 6)(5, 8)$ – soweit es möglich ist – ein:



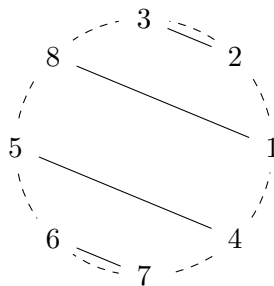
Da der 8-Zykel überschneidungsfrei bezüglich $(1, 4)(2, 7)(3, 6)(5, 8)$ sein muss, dürfen sich die Strecken, die die Involution induziert, gemäß Lemma 4.9 nicht schneiden. Daher muss 1 mit der Position im Süd-Osten, zwischen 1 und 7, verbunden sein. Dort muss also eine 4 stehen. Ebenso können wir schließen, dass die Positionen im Westen nur von den Partnern der Punkte 3 und 5 besetzt sein können, also 6 und 8.



Mit den Strecken der Involution $(1, 8)(2, 3)(4, 5)(6, 7)$ erhalten wir das folgende Bild:

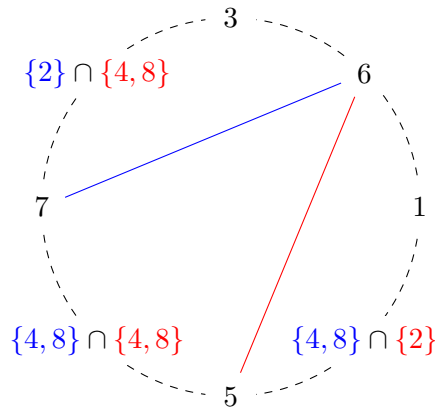


Da die übrigen Strecken die bereits bestehenden nicht schneiden können, muss 1 mit der Position im Nord–Westen verbunden sein, wo also 8 stehen muss. Ebenso argumentiert man, dass die 6 im Süd–Westen stehen muss. Wir erhalten also die vollständige Kreisdarstellung:



Dies liefert den 8–Zykel $(1, 2, 3, 8, 5, 6, 7, 4)$, der eine gültige Lösung ist (dazu muss man überprüfen, ob er bezüglich allen drei Involutionen überschneidungsfrei ist).

Analog (durch Einsetzen verschiedener Zahlen für die Position im Nord–Osten zwischen 1 und 3) erhält man die Zyklen $(1, 4, 3, 2, 5, 8, 7, 6)$, $(1, 6, 3, 4, 5, 2, 7, 8)$ und $(1, 8, 3, 6, 5, 4, 7, 2)$. Dieses Verfahren funktioniert aber nicht immer. Das wird von der folgende Konfiguration des 4–Zykels $(1, 3, 7, 5)$ demonstriert, bei dem wir mit den Involutionen $(1, 2)(3, 4)(5, 6)(7, 8)$ und $(1, 8)(2, 3)(4, 5)(6, 7)$ die Möglichkeiten einschränken.



Es gibt keine Position, an der sich die 2 befinden könnte. Durch solche Phänomene kann es auch vorkommen, dass ein mit dieser Methode berechneter Zykel ungültig ist – da es sinnvoller ist, möglichst wenige Involutionen zu betrachten, könnten wir einen solchen Widerspruch übersehen.

Während dieses Algorithmus die Komplexität des Problems nicht verändert, verringert er die Laufzeit in der Praxis aber erheblich. Insbesondere dadurch, dass wir nur noch alle n -Zykel anstatt aller $(2n)$ -Zykel durchlaufen, sorgt für einen gewaltigen Geschwindigkeitsschub. Es wäre schön, sich auch diese Berechnung zu sparen oder einen Zusammenhang zwischen den Positionen der geraden und der ungeraden Zahlen zu finden. Bei beidem hatten wir bislang aber keinen Erfolg.

4.2 Überlagerung einiger ebener Gitter

Wir haben zu Beginn von Kapitel 4 eingesehen, dass es sehr schwierig ist, allgemeingültige und nicht-triviale Aussagen über das Einbettungsproblem zu treffen. Wir haben uns daher in Abschnitt 4.1 zunächst nur mit der Frage beschäftigt, wann sich ein geschlossener Flächenkomplex auf ein Dreieck zusammenfallen lässt und erkannt, dass bereits in diesem einfachen Problem eine große Menge an Struktur enthalten ist.

Eine etwas allgemeinere Frage als „Kann man auf ein Dreieck zusammenfallen?“ wäre, ob man so falten kann, dass der resultierende 2-Faltkomplex in einer Ebene liegt. Im Allgemeinen ist auch das eine schwierige Frage, die davon abhängt, welche Winkel man für die Dreiecke im 2-Faltkomplex wählt. Wir werden in diesem Abschnitt zwei generische Fälle präsentieren, in denen sich die Frage nach einer solchen ebenen Einbettung auf die Frage nach der Faltbarkeit auf ein Dreieck reduziert.

Im ersten dieser Fälle beschäftigen wir uns mit Mustern die nur aus gleichseitigen Dreiecken bestehen. Die formale Beschreibung dieses Szenarios mit 2-Faltkomplexen liefern wir in Abschnitt 4.2.1. Wir werden feststellen, dass die Betrachtung unendlicher Faltungen für den allgemeinen Nachweis der Faltbarkeit nützlich ist. Aus diesem Grund werden wir uns in Abschnitt 4.2.2 mit diesen beschäftigen und eine unendliche Faltung des hexagonalen Musters konkret ausführen. Unter Verwendung der Automorphismengruppe werden wir in Abschnitt 4.2.3 nachweisen, dass wir damit alle unendlichen Faltungen ausführen können, die

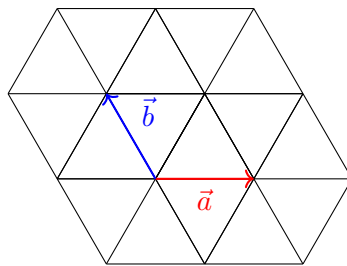
wir benötigen, damit wir jedes ebene Muster aus gleichseitigen Dreiecken auf ein einzelnes Dreieck zusammenfalten können. Schließlich werden wir die Konzepte aus diesen Abschnitten benutzen und in Abschnitt 4.2.4 auf den Fall rechtwinkliger, gleichschenkliger Dreiecke anwenden.

In beiden Fällen verwenden wir dieselbe Technik: Wir zeigen, dass sich jeder endliche Teilbereich der ebenen Muster so falten lässt, dass er eben bleibt und weniger Dreiecke enthält. Dieses Verfahren führen wir solange aus, bis nur ein einzelnes Dreieck übrig bleibt.

4.2.1 Beschreibung des hexagonalen Gitters

In diesem Abschnitt beginnen wir mit der Untersuchung von ebenen Faltmustern, die aus gleichseitigen Dreiecken bestehen. Daher müssen wir uns mit dem hexagonalen Gitter befassen. Damit wir mit diesem arbeiten können (und die „Faltung“ des Gitters einen Sinn bekommt), müssen wir es als 2-Faltkomplex beschreiben.

Wir brauchen also eine Parametrisierung dieses Gitters. Dazu betrachten wir zuerst die Gitterpunkte und erkennen, dass diese ein zweidimensionales $\mathbb{Z}$ -Gitter im $\mathbb{R}^2$ bilden. Eine Gitterbasis davon ermöglicht also eine eindeutige Beschreibung aller Gitterpunkte. Wir wählen diese Gitterbasis:



Für eine Kantenlänge von 1 hätten diese Vektoren bezüglich der Standardbasis $(\vec{e}_1, \vec{e}_2)$ die Form

$$\vec{a} = \vec{e}_1 \qquad \vec{b} = -\frac{1}{2}\vec{e}_1 + \frac{\sqrt{3}}{2}\vec{e}_2.$$

Ab jetzt werden wir alle Punkte bezüglich der Gitterbasis $(\vec{a}, \vec{b})$ angeben, d. h. der Punkt mit den Koordinaten $\begin{pmatrix} \alpha \\ \beta \end{pmatrix}$ bezeichne den Punkt $\alpha\vec{a} + \beta\vec{b}$. Damit ist die Punktmenge des hexagonalen Gitters genau

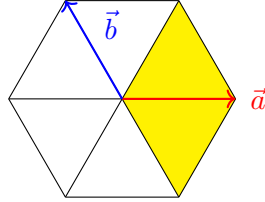
$$\mathcal{K}_0 := \left\{ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \right\} \mid \alpha, \beta \in \mathbb{Z} \right\}.$$

Die Kantenmenge lässt sich aus der Graphik ablesen – die Addition der Basisvektoren $\vec{a}$ und

$\vec{\beta}$, sowie von $\vec{a} + \vec{\beta}$ verbindet zwei Punkte entlang einer Kante.

$$\begin{aligned} \mathcal{K}_1 := & \left\{ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta \end{pmatrix} \right\} \mid \alpha, \beta \in \mathbb{Z} \right\} \\ & \uplus \left\{ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta+1 \end{pmatrix} \right\} \mid \alpha, \beta \in \mathbb{Z} \right\} \\ & \uplus \left\{ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha \\ \beta+1 \end{pmatrix} \right\} \mid \alpha, \beta \in \mathbb{Z} \right\} \end{aligned}$$

Die Flächen dieses Gitters lassen sich in zwei Familien einteilen. Diese erhalten wir dadurch, dass wir zu jedem Punkt die beiden Dreiecke angeben, die an der angrenzenden $\vec{a}$ -Kante anliegen.



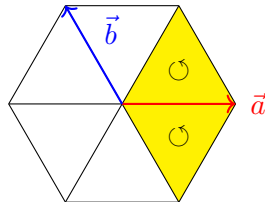
In Formeln handelt es sich um die folgende Menge:

$$\begin{aligned} \mathcal{K}_2 := & \left\{ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta+1 \end{pmatrix} \right\} \mid \alpha, \beta \in \mathbb{Z} \right\} \\ & \uplus \left\{ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha \\ \beta-1 \end{pmatrix} \right\} \mid \alpha, \beta \in \mathbb{Z} \right\} \end{aligned}$$

Damit können wir das hexagonale Gitter formal definieren.

Definition 4.32 (hexagonaler Gitterkomplex). *Seien $\mathcal{K}_0, \mathcal{K}_1, \mathcal{K}_2$ wie oben definiert. Der geschlossene Flächenkomplex $((\mathcal{K}_0, \mathcal{K}_1, \mathcal{K}_2), \subseteq)$ heißt **hexagonaler Gitterkomplex**.*

Wir haben damit das hexagonale Gitter als geschlossenen Flächenkomplex definiert. Damit wir es falten können, müssen wir diesen zu einem 2-Faltkomplex ergänzen. Gemäß Bemerkung 4.6 müssen wir dazu nur eine simpliziale Oberfläche für diesen Flächenkomplex angeben. Dazu verwenden wir die globale Orientierung der Ebene, die wir wie folgt auf die Flächen des hexagonalen Gitterkomplexes übertragen:



Diese simpliziale Oberfläche entspricht der Auszeichnung der folgenden simplizialen Orientierungen, die wir auch als 3-Zykel der Ecken verstehen können:

$$\left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta+1 \end{pmatrix} \right\} \rightsquigarrow \left(\begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta+1 \end{pmatrix} \right)$$

$$\left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha \\ \beta-1 \end{pmatrix} \right\} \rightsquigarrow \left(\begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha \\ \beta-1 \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta \end{pmatrix} \right)$$

Mit dem resultierenden 2-Faltkomplex möchten wir die folgenden Aussagen nachweisen:

1. Jede Kante aus $\mathcal{K}_1$ legt eine eindeutige Gerade fest, die eine Vereinigung von Kanten aus $\mathcal{K}_1$ ist.
2. Diese Gerade ist eine Spiegelungsachse des Musters, an der wir den hexagonalen Gitterkomplex falten können.
3. Der Basiskomplex nach der Faltung ist im ursprünglichen Muster enthalten.

Damit weisen wir nach, dass sich jedes Teilmuster (außer einem einzelnen Dreieck) so falten lässt, dass es weniger Flächenklassen enthält, aber ein Teilmuster bleibt. Für endliche Komplexe ergibt sich durch Iteration dieses Prozesses, dass wir in der Situation aus Abschnitt 4.1 sind.

Da wir noch nicht alle Operationen dieses Prozesses geklärt haben (z.B. verlangt der zweite Schritt die Faltung von unendlich vielen Flächenpaaren), erklären und beweisen wir unser Vorgehen am Beispiel der Kante $\left\{ \begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \right\}$. Danach werden wir alle anderen Kanten durch Automorphismen auf diesen Fall zurückführen.

Diese Kante legt die Gerade $\{x\vec{a} | x \in \mathbb{R}\}$ fest. Diese Gerade lässt sich in $\mathbb{Z}$ -Abschnitte unterteilen:

$$\{x\vec{a} | x \in \mathbb{R}\} = \bigcup_{\alpha \in \mathbb{Z}} \{(\alpha + x)\vec{a} | x \in [0, 1]\}$$

Jeder Abschnitt $\{(\alpha + x)\vec{a} | x \in [0, 1]\}$ entspricht einer Kante $\left\{ \begin{pmatrix} \alpha \\ 0 \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ 0 \end{pmatrix} \right\}$ aus $\mathcal{K}_1$. Dabei ist festzuhalten, dass wir keine Eigenschaften einer konkreten Darstellung im $\mathbb{R}^2$ verwendet haben. Diese Zerlegungseigenschaft beruht allein auf der Definition von $\mathcal{K}_1$ mithilfe der Gitterbasis.

Um nachzuweisen, dass diese Gerade eine Spiegelungsachse ist, müssen wir einen Automorphismus von Grad 2 des homogenen 2-Komplexes angeben, der diese Gerade invariant lässt. Wir wählen die Abbildung

$$\mu : \mathcal{K}_0 \rightarrow \mathcal{K}_0, \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \mapsto \begin{pmatrix} \alpha - \beta \\ -\beta \end{pmatrix},$$

die offenbar selbstinvers und nicht trivial ist, sowie die obige Gerade fest lässt. Um nachzuweisen, dass es sich um einen Automorphismus handelt, müssen wir aber auch die Kanten

und Flächen untersuchen. Die Kanten werden wie folgt abgebildet:

$$\begin{aligned} \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta \end{pmatrix} \right\} &\mapsto \left\{ \begin{pmatrix} \alpha-\beta \\ -\beta \end{pmatrix}, \begin{pmatrix} \alpha-\beta+1 \\ -\beta \end{pmatrix} \right\} \\ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta+1 \end{pmatrix} \right\} &\mapsto \left\{ \begin{pmatrix} \alpha-\beta \\ -\beta \end{pmatrix}, \begin{pmatrix} \alpha-\beta \\ -\beta-1 \end{pmatrix} \right\} \\ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha \\ \beta+1 \end{pmatrix} \right\} &\mapsto \left\{ \begin{pmatrix} \alpha-\beta \\ -\beta \end{pmatrix}, \begin{pmatrix} \alpha-\beta-1 \\ -\beta-1 \end{pmatrix} \right\} \end{aligned}$$

Insbesondere sind alle Kanten auch nach der Anwendung von μ noch Kanten. Auch die Flächen werden durch Anwendung von μ auf Flächen abgebildet:

$$\begin{aligned} \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta+1 \end{pmatrix} \right\} &\mapsto \left\{ \begin{pmatrix} \alpha-\beta \\ -\beta \end{pmatrix}, \begin{pmatrix} \alpha-\beta+1 \\ -\beta \end{pmatrix}, \begin{pmatrix} \alpha-\beta \\ -\beta-1 \end{pmatrix} \right\} \\ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha \\ \beta+1 \end{pmatrix} \right\} &\mapsto \left\{ \begin{pmatrix} \alpha-\beta \\ -\beta \end{pmatrix}, \begin{pmatrix} \alpha-\beta+1 \\ -\beta \end{pmatrix}, \begin{pmatrix} \alpha-\beta+1 \\ -\beta+1 \end{pmatrix} \right\} \end{aligned}$$

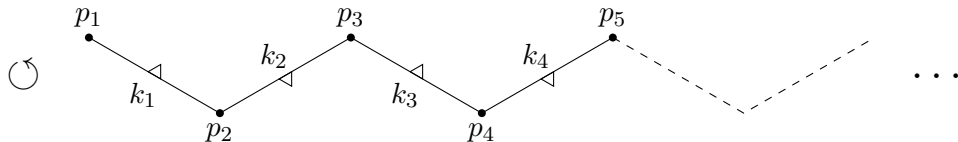
Folglich ist μ eine Spiegelung des Musters, womit die Gerade $\{x\vec{a} | x \in \mathbb{R}\}$ eine Spiegelungsachse ist. Anschaulich sollte es daher möglich sein, die beiden Hälften des Musters aufeinanderzufalten – schließlich wird nur ein Blatt Papier entlang einer Kante gefaltet. Da dieses Blatt aber in unendlich viele Segmente zerlegt wurde, wird der Formalismus ein wenig komplizierter. Wir werden uns im nächsten Abschnitt 4.2.2 damit befassen, wie man solche unendlichen Faltungen definiert. Die Faltung bezüglich μ werden wir direkt im Anschluss an Satz 4.37 ausführen.

4.2.2 Unendliche Faltungen

Wir haben im letzten Abschnitt 4.2.1 erkannt, dass wir eine Faltung mit unendlich vielen Dreiecken ausführen müssen, wenn wir den hexagonalen Gitterkomplex an einer Spiegelungsachse falten wollen. In diesem Abschnitt werden wir das Konzept unendlicher Faltungen formal definieren und in Satz 4.37 einige einfache Kriterien angeben, die hinreichend für die Existenz einer unendlichen Faltung sind. Wir werden dabei den Formalismus der Faltpläne aus Abschnitt 3.5 verwenden.

Während man eine endliche Menge von Faltungen aufgrund der Kommutativität von Faltplänen aus Lemma 3.73 auf wiederholte Anwendung einzelner Faltpläne zurückführen kann, treten bei unendlich vielen Faltungen neue Komplikationen auf. In Beispiel 4.33 zeigen wir eine solche Situation, in der die Ausführung von unendlich vielen Faltplänen nicht einmal einen wohldefinierten Proto-Faltkomplex liefert.

Beispiel 4.33. *Wir betrachten den folgenden unendlichen 1-Faltkomplex $\mathcal{F}$:*



Formal beschreiben wir diesen 1-Faltkomplex wie folgt:

$$\mathcal{K}_0 := \{p_i \mid i \in \mathbb{N}\}$$

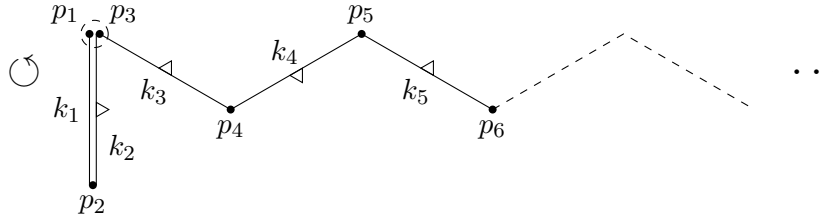
$$\mathcal{K}_1 := \{k_j \mid j \in \mathbb{N}\}$$

mit $p_i \prec k_j$ genau dann, wenn $i \in \{j, j+1\}$ gilt. Die simpliziale Oberfläche für die Kanten ist in der Abbildung gekennzeichnet, zur Orientierung der Fächer wählen wir die Standardorientierung der Ebene.

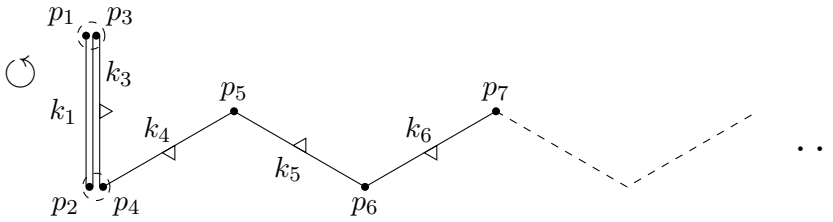
Da die Kanten k_j und k_{j+1} für jedes $j \in \mathbb{N}$ benachbart sind, gibt es gemäß Bemerkung 2.39 eindeutige Identifikationen $\varphi_j : \mathcal{F}_{\preceq k_j} \rightarrow \mathcal{F}_{\preceq k_{j+1}}$. Wir betrachten dann die folgende unendliche Menge von Faltplänen:

$$\mathcal{P} := \{P_j := (\varphi_j, (k_j, +1) \leftrightarrow (k_{j+1}, -1)) \mid j \in \mathbb{N}\}$$

Um eine Intuition darüber zu gewinnen, wie die Anwendung von all diesen Faltplänen auf $\mathcal{F}$ aussehen muss, wenden wir zuerst P_1 auf $\mathcal{F}$ an. Dabei erhalten wir



Dabei ist die zyklische Reihung an $\{p_1, p_3\}$ durch den 3-Zykel (k_1, k_2, k_3) gegeben. Als nächstes wenden wir P_2 an, wodurch $(k_2, +1)$ und $(k_3, -1)$ aufeinandergefaltet werden:



Dadurch ändert sich die zyklische Reihung an $\{p_2, p_4\}$ zu (k_4, k_3, k_2, k_1) . Als nächstes wenden wir P_3 an und erhalten die zyklische Reihung $(k_1, k_2, k_3, k_4, k_5)$ an $\{p_1, p_3, p_5\}$.

Das Muster sollte klar geworden sein: Nachdem wir die Faltpläne $\{P_1, P_2, \dots, P_{2m+1}\}$ angewendet haben ($m \in \mathbb{N}$), lautet die zyklische Reihung an $\{p_1, p_3, \dots, p_{2m+1}\}$ wie folgt:

$$\{k_1, k_2, \dots, k_{2m+1}\} \rightarrow \{k_1, k_2, \dots, k_{2m+1}\}, \quad k_j \mapsto \begin{cases} k_{j+1} & j < 2m+1 \\ k_1 & j = 2m+1 \end{cases}$$

Insbesondere wird das k_j , das zum angewendeten Faltplan P_j mit dem höchsten Index korrespondiert, auf k_1 abgebildet. Wenn wir alle Faltpläne aus $\mathcal{P}$ anwenden, gibt es aber keinen höchsten Index mehr. Die zyklischen Reihungen degenerieren zur Nachfolgerabbildung

$$\{k_j \mid j \in \mathbb{N}\} \rightarrow \{k_j \mid j \in \mathbb{N}\}, \quad k_j \mapsto k_{j+1}.$$

Da diese Abbildung nicht mehr bijektiv ist (k_1 hat kein Urbild), hat der so konstruierte 1-Überlagerungskomplex keine wohldefinierten Fächer, ist also kein 1-Faltkomplex.

Es gibt zwei Möglichkeiten, mit der Einschränkung aus Beispiel 4.33 umzugehen. Wir könnten sie als Manko unseres Formalismus auffassen und das Modell der n -Faltkomplexe so erweitern, dass es mit solchen Situationen umgehen kann. Allerdings führt das von unserer ursprünglichen Fragestellung weg, in der wir im Wesentlichen nur ein Stück Papier an einer Kante falten wollten. Anstatt unser Modell anzupassen, akzeptieren wir dessen Unzulänglichkeit in solchen Situationen und suchen nach den Kriterien, unter denen die Ausführung unendlich vieler Faltpläne wohldefiniert ist.

In Beispiel 4.33 entsteht das Problem dadurch, dass wir unendlich viele Äquivalenzklassen vereinigen. In dem Fall aus Abschnitt 4.2.1, der uns interessiert, vereinigen wir immer nur höchstens zwei Äquivalenzklassen. Wir werden uns also auf solche unendlichen Faltungen einschränken, bei denen lokal nur endlich viele Äquivalenzklassen vereinigt werden.

Definition 4.34 (lokal endlich zulässig). *Sei $\mathcal{F}$ ein n -Faltkomplex und $\mathcal{P}$ eine Menge von Faltplänen. Die Menge $\mathcal{P}$ heißt **lokal endlich zulässig**, falls gilt:*

1. *Für jede endliche Teilmenge $\{P_1, \dots, P_m\} \subseteq \mathcal{P}$ ist $(P_1 \circ \dots \circ P_m)(\mathcal{F})$ ein wohldefinierter n -Faltkomplex.*
2. *Für jedes $x \in \bigsqcup_{i=0}^n \mathcal{K}_i$ existieren endlich viele Faltpläne $P_1, \dots, P_m \in \mathcal{P}$, sodass es kein $P \in \mathcal{P} \setminus \{P_1, \dots, P_m\}$ gibt, das die Äquivalenzklasse von x in $(P_1 \circ \dots \circ P_m)(\mathcal{F})$ durch die formale Erweiterung der Äquivalenzrelation in P modifiziert.*

In Definition 4.34 fordern wir zwei verschiedene Eigenschaften. Wir fordern zunächst, dass jede endliche Teilmenge der Faltpläne angewendet werden kann. Diese Eigenschaft muss von jeder unendlichen Faltung erfüllt sein, ist aber nicht hinreichend für die Existenz der unendlichen Faltung, wie Beispiel 4.33 demonstriert. Daher fordern wir auch die zweite Bedingung, die dafür sorgt, dass wir an jeder Stelle nur eine endliche Zahl von Faltungen ausführen müssen, um zu wissen, wie diese aussieht.

Dass wir lokal endlich zulässige Mengen von Faltplänen tatsächlich ausführen können, zeigen wir in Definition 4.35.

Definition 4.35 (unendliche Faltung). *Sei $\mathcal{F}$ ein n -Faltkomplex mit n -Überlagerungskomplex $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec, \sim)$ und $\mathcal{P}$ eine lokal endlich zulässige Menge von Faltplänen. Dann lässt sich der homogene n -Komplex $((\mathcal{K}_i)_{0 \leq i \leq n}, \prec)$ wie folgt zu einem n -Faltkomplex ergänzen: Für jedes $x \in \bigsqcup_{i=0}^n \mathcal{K}_i$ seien $P_1, \dots, P_m$ die Faltpläne aus Definition 4.34 der lokal endlich zulässigen Faltungen.*

- *Die Äquivalenzklasse von x ist dessen Äquivalenzklasse in $(P_1 \circ \dots \circ P_m)(\mathcal{F})$.*
- *Falls $x \in \mathcal{K}_n$ gilt, ist die Menge dessen Randstücke gleich der Menge der Randstücke in $(P_1 \circ \dots \circ P_m)(\mathcal{F})$.*
- *Falls $x \in \mathcal{K}_{n-1}$ gilt, ist dessen Fächer gleich dem Fächer in $(P_1 \circ \dots \circ P_m)(\mathcal{F})$.*

Wir bezeichnen diesen n -Faltkomplex mit $\mathcal{P}(\mathcal{F})$, der **unendlichen Faltung von $\mathcal{P}$ auf $\mathcal{F}$** .

Wohldefiniertheit. Wir müssen nachweisen, dass die Konstruktion nur von der Menge $\mathcal{P}$ abhängig ist und dass sie tatsächlich einen n -Faltkomplex liefert.

Für den ersten Teil betrachten wir $x \in \biguplus_{i=0}^n \mathcal{K}_i$. Seien $\{P_1, \dots, P_m\}$ und $\{Q_1, \dots, Q_l\}$ zwei Mengen von Faltplänen, die beide die Eigenschaft aus Definition 4.34 in Bezug auf x erfüllen. Wir vergleichen die beiden n -Faltkomplexe

$$\begin{aligned}\mathcal{F}_P &:= (P_1 \circ \dots \circ P_m)(\mathcal{F}) \\ \mathcal{F}_Q &:= (Q_1 \circ \dots \circ Q_l)(\mathcal{F})\end{aligned}$$

mit dem n -Faltkomplex $\mathcal{G}$, der durch Anwenden von $\{P_1, \dots, P_m\} \cup \{Q_1, \dots, Q_l\}$ aus $\mathcal{F}$ hervorgeht. Nach Voraussetzung ändert sich die Äquivalenzklasse von x beim Übergang von $\mathcal{F}_P$ zu $\mathcal{G}$ nicht, ebenso wenig beim Übergang $\mathcal{F}_Q$ zu $\mathcal{G}$. Folglich muss die Äquivalenzklasse von x in $\mathcal{F}_P$ und $\mathcal{F}_Q$ gleich sein.

Wir müssen noch nachweisen, dass sich Randstücke und Fächer nur dann verändern können, wenn sich auch Äquivalenzklassen ändern. Dies folgt aber sofort aus der Konstruktion der primitiven Faltungen von Grad $n - 1$ (Definition 3.41) und n (Definition 3.31).

Um nachzuweisen, dass wir hier tatsächlich einen n -Faltkomplex konstruieren, weisen wir die Eigenschaften aus Definition 3.27 der n -Faltkomplexe nach. Dass es sich um einen n -Überlagerungskomplex handelt, folgt daraus, dass wir dabei gemäß Satz 2.32 nur darauf achten müssen, keine zwei Punkte zu identifizieren, die in einer gemeinsamen Kante vorkommen. Wenn dies der Fall wäre, dann wäre eine der Eigenschaften aus der Definition 4.34 der lokal endlichen Zulässigkeit nicht erfüllt (sobald man alle Faltpläne ausgeführt hat, die eine solche Äquivalenzklasse ändern, läge kein n -Faltkomplex mehr vor).

Da wir alle Randstücke und Fächer aus bestehenden n -Faltkomplexen übernommen haben, sind diese individuell wohldefiniert. Um nachzuweisen, dass alles miteinander verträglich ist, wenden wir alle Faltpläne für eine Kante an, danach alle Faltpläne für die anliegenden Flächen. Dann sind die Fächer und Randstücke dieser Elemente nach Voraussetzung gleich zu denen von $\mathcal{P}(\mathcal{F})$ und (als Teil eines n -Faltkomplexes) miteinander verträglich. $\square$

Wir bemerken, dass diese Definition unendlicher Faltungen im endlichen Fall trivialerweise mit der Definition endlicher Faltungen übereinstimmt. Wir weisen jetzt eine Eigenschaft nach, die definitiv von einer „unendlichen Faltung“ erfüllt werden muss.

Lemma 4.36. *Sei $\mathcal{F}$ ein n -Faltkomplex und $\mathcal{P}$ sei eine lokal endlich zulässige Menge von Faltplänen. Dann ist jede Teilmenge $\mathcal{Q}$ von $\mathcal{P}$ auch lokal endlich zulässig und $\mathcal{Q}(\mathcal{F})$ ist eine Teilfaltung von $\mathcal{P}(\mathcal{F})$.*

Beweis. Klar: Jede endliche Teilmenge von $\mathcal{Q}$ ist anwendbar.

Wir zeigen noch die Lokalität: Sei $x \in \biguplus_{i=0}^n \mathcal{K}_i$ und $P_1, \dots, P_m$ seien die Faltpläne aus Definition 4.34. Da jeder Faltplan jeweils höchstens zwei verschiedene Äquivalenzklassen vereinigt (vgl. den Beweis von Lemma 3.83), wird $[x]_{\mathcal{P}(\mathcal{F})}$ von endlich vielen Äquivalenzklassen $[x_1]_{\mathcal{F}}, \dots, [x_l]_{\mathcal{F}}$ partitioniert.

Für jedes mögliche Paar solcher Äquivalenzklassen wählen wir höchstens einen Faltplan aus $\mathcal{Q}$, der diese Klassen vereinigt. Nach Anwenden dieser endlichen Menge führt kein anderer

Faltplan mehr zu einer weiteren Vereinigung von Äquivalenzklassen, da $[x]_{\mathcal{P}(\mathcal{F})}$ maximal bezüglich $\mathcal{P}$ ist.

Dass es sich um eine Teilfaltung handelt, ist klar: Für $x \in \bigsqcup_{i=0}^n \mathcal{K}_i$ seien $P_1, \dots, P_m$ die Faltpläne aus Definition 4.34 für $\mathcal{P}$ und $Q_1, \dots, Q_l$ diejenigen für $\mathcal{Q}$. Führe alle diese Faltpläne aus, dann ist (lokal) $\mathcal{Q}(\mathcal{F})$ eine Teilfaltung davon. Da die Ausführung von all diesen Faltplänen sich in x nicht von der Ausführung von $P_1, \dots, P_m$ unterscheidet, folgt die Behauptung. $\square$

Damit haben wir eine allgemeine Situation beschrieben, in der unendliche Faltungen vorliegen. Da die Überprüfung der lokal endlichen Zulässigkeit aufwendig ist, suchen wir nach einfacheren Charakterisierungen. Da wir in Abschnitt 4.2.1 an einer unendlichen Faltung interessiert sind, die von einer Spiegelung herrührt, versuchen wir zu charakterisieren, in welchen Fällen eine Spiegelung zu einer validen unendlichen Faltung führt.

Dazu müssen wir uns klar werden, was eine Spiegelung eigentlich ist. Wir wissen, dass es sich dabei um einen selbstinversen Automorphismus des Faltkomplexes handeln muss. Dies lässt aber noch die Identität und einige Drehungen zu. Um diese auszuschließen, benötigen wir die Eigenschaft von Spiegelungen, Orientierungen zu invertieren. Dies ergibt wenig Sinn, wenn wir keine konsistente Orientierung vorliegen haben, demnach schränken wir uns vorerst auf orientierbare Flächenkomplexe ein (vgl. Definition 4.29). Wir fordern dann, dass bei Anwendung der Spiegelung auf eine Fläche dessen Orientierung umgedreht wird.

Bei einer Spiegelung muss man natürlich auch immer auf die Spiegelachse (oder Spiegelfläche) achten, die unter selbiger fest bleibt. Bei der Interaktion dieser Achse mit dem Faltkomplex können unangenehme Dinge passieren: Falls eine Kante durch die Spiegelung fest gelassen wird, ihre Eckpunkte aber vertauscht werden, erreichen wir durch Identifikation der Punkte keinen Überlagerungskomplex. Folglich müssen wir auch diese Möglichkeit ausschließen. Es zeigt sich, dass diese Bedingungen genügen, um eine unendliche Faltung zu garantieren, wie wir in Satz 4.37 nachweisen.

Satz 4.37. *Sei $\mathcal{F} = ((\mathcal{K}_0, \mathcal{K}_1, \mathcal{K}_2), \prec)$ ein orientierbarer Flächenkomplex, den wir mittels Bemerkung 4.6 und einer Orientierung $\{\sigma_f\}_{f \in \mathcal{K}_2}$ von $\mathcal{F}$ als 2-Faltkomplex auffassen. Sei ein Automorphismus μ von $\mathcal{F}$ gegeben, der die folgenden Eigenschaften erfüllt:*

1. μ ist selbstinvers.
2. Für eine simpliziale Orientierung a von $f \in \mathcal{K}_2$ gilt $\sigma_f(a) \neq \sigma_{\mu(f)}(\mu(a))$.
3. Falls μ eine Kante stabilisiert, dann stabilisiert sie auch deren Eckpunkte.

Dann ist die Menge der Faltpläne

$$\{(\mu|_f, (f, +1) \leftrightarrow (\mu(f), +1)) \mid f \in \mathcal{K}_2\}$$

lokal endlich zulässig, wobei $\mu|_f : \mathcal{F}_{\preceq f} \rightarrow \mathcal{F}_{\preceq \mu(f)}, x \mapsto \mu(x)$ die Einschränkung von μ auf die Fläche f ist.

Beweis. Wir bezeichnen die Menge der Faltpläne aus der Formulierung des Satzes mit $\mathcal{P}$.

Wir zeigen zuerst die Lokalität. Dadurch, dass wir die Identifikationen der Faltpläne durch Einschränkung der Abbildung μ definieren und μ selbstinvers ist, wird jedes $x \in \mathcal{K}_0 \sqcup \mathcal{K}_1 \sqcup \mathcal{K}_2$

nur zu $\mu(x)$ äquivalent sein. Die Äquivalenzklasse von x kann also höchstens einmal durch die Anwendung eines Faltplans modifiziert werden.

Als nächstes zeigen wir, dass alle endlichen Teilmengen der Faltpläne wohldefinierte n -Faltkomplexe erzeugen. Dazu müssen wir die Bedingungen zur Zulässigkeit von Faltplänen aus Definition 3.71, sowie die Fächerkompatibilität aus der Definition 3.27 der n -Faltkomplexe überprüfen. Sei dazu $\{P_1, \dots, P_m\} \subseteq \mathcal{P}$ endlich und

$$(\mu|_f, (f, +1) \leftrightarrow (\mu(f), +1)) \in \mathcal{P} \setminus \{P_1, \dots, P_m\}.$$

Wir nehmen an, dass $\mathcal{G} := (P_1 \circ \dots \circ P_m)(\mathcal{F})$ ein n -Faltkomplex ist. Wir wollen zeigen, dass $(\mu|_f, (f, +1) \leftrightarrow (\mu(f), +1))$ zulässig für $\mathcal{G}$ ist und dass seine Anwendung auf $\mathcal{G}$ einen n -Faltkomplex liefert. Daraus folgt die geforderte Behauptung durch ein Induktionsargument.

1. Wir müssen nachweisen, dass $\mathcal{G}_{\mu|_f}$ ein wohldefinierter 2-Überlagerungskomplex ist. Dazu müssen wir überprüfen, ob Satz 2.32 gültig ist. Anstatt uns in den Details der Faltpläne zu verlieren, besinnen wir uns darauf, dass die Identifikationen aller Faltpläne aus $\mathcal{P}$ einen gemeinsamen Ursprung in μ haben. Wenn wir also nachweisen können, dass diese Abbildung keine Widersprüche in der Äquivalenzrelation erzeugt, haben wir die Behauptung gezeigt.

Es gibt genau dann ein Problem, falls es zwei Punkte mit einer Kante zwischen ihnen gibt, die durch μ identifiziert werden. Dies tritt aber genau dann auf, wenn die Kante stabilisiert wird, ihre Eckpunkte jedoch nicht. Diese Situation kann nach Voraussetzung aber nicht auftreten.

2. Um nachzuweisen, dass $(f, +1)$ und $(\mu(f), +1)$ tatsächlich Randstücke in $\mathcal{G}$ sind, müssen wir uns daran erinnern, dass kein anderer Faltplan diese beiden Oberflächen verknüpft. Da sie zu Beginn Randstücke waren (alle Oberflächen sind in dem 2-Faltkomplex aus Bemerkung 4.6 Randstücke) und dieser Faltplan noch nicht ausgeführt wurde, sind sie gemäß der Konstruktion aus Definition 3.72 noch immer Randstücke.
3. Wir wollen nachweisen, dass die beiden Oberflächen komplementär bezüglich aller Fächer sind, in denen sie beide vorkommen. Sei dazu eine solche Kante von f gegeben. Falls es sich um eine Fixkante handelt, ist die Behauptung trivialerweise erfüllt, da der Fächer von zwei Elementen nicht modifiziert wird und die Oberflächen gemäß Folgerung 4.30 komplementär sind. Andernfalls muss es bereits eine Fächersumme gegeben haben. Diese erfolgt entlang zweier positiver Oberflächen. Folglich liegen die beiden anderen positiven Oberflächen (die wir jetzt zusammenfalten wollen) komplementär (gemäß Folgerung 4.30 gibt es bei Komplementarität keine vom Vorzeichen gemischten Paare von Oberflächen).
4. Um zu zeigen, dass die Fächersumme ausgeführt werden kann, müssen wir gemäß der Definition 3.39 der Fächersumme nachweisen, dass die Orientierungen von f und $\mu(f)$ invers zueinander sind. Dies folgt nach Voraussetzung: Eine positive Orientierung von f wird durch μ zu einer negativen Orientierung von $\mu(f)$.

Damit ist der Faltplan zulässig. Als nächstes müssen wir nachweisen, dass die so entstandenen Fächer kompatibel sind. Nach der Identifikation der Spiegelung μ liegen in jeder Flächenklasse genau zwei Flächen. Da Redukte auf höchstens zwei Elemente immer gleich sind, folgt die Kompatibilität sofort. $\square$

Nachdem wir in Satz 4.37 allgemeine Kriterien aufgestellt haben, unter denen eine unendliche Faltung an einer Spiegelung möglich ist, wenden wir diese auf die Spiegelung μ aus Abschnitt 4.2.1 an.

Da μ offenbar selbstinvers ist, müssen wir die anderen beiden Bedingungen von Satz 4.37 überprüfen.

Betrachten wir zuerst die Kantenbedingung. Seien zwei Punkte $\begin{pmatrix} \alpha \\ \beta \end{pmatrix}$ und $\begin{pmatrix} \alpha - \beta \\ -\beta \end{pmatrix}$ mit einer Kante zwischen ihnen gegeben, die durch μ identifiziert werden. Ihre Differenz ist $\beta \vec{a} + 2\beta \vec{b}$, während aus der Definition der Kanten folgt, dass die Differenz der Eckpunkte nur Werte in $\{\pm \vec{a}, \pm \vec{b}, \pm(\vec{a} + \vec{b})\}$ annehmen kann. Da dies nicht möglich ist, kann diese Situation nicht auftreten.

Als nächstes müssen wir nachweisen, dass μ die Orientierung einer Fläche umkehrt. Sei Δ positiv orientiert:

$$\left(\begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha + 1 \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha + 1 \\ \beta + 1 \end{pmatrix} \right)$$

Dann ist $\mu(\Delta)$ negativ orientiert:

$$\left(\begin{pmatrix} \alpha - \beta \\ -\beta \end{pmatrix}, \begin{pmatrix} (\alpha - \beta) + 1 \\ -\beta \end{pmatrix}, \begin{pmatrix} \alpha - \beta \\ -\beta - 1 \end{pmatrix} \right)$$

Damit folgt aus Satz 4.37, dass wir die unendliche Faltung an μ ausführen können. Wir müssen noch zeigen, dass wir damit wieder im ursprünglichen Muster gelandet sind (um dieses Verfahren iterieren zu können). Die folgt aber schon daraus, dass die Spiegelung μ bijektiv zwischen $\left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \mid \beta \geq 0 \right\}$ und $\left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \mid \beta \leq 0 \right\}$ ist. Die Einschränkung des Gitters auf eine der beiden Mengen bleibt Teil des Gitters.

4.2.3 Automorphismen des hexagonalen Gitters

Wir haben im letzten Abschnitt eine Faltung des hexagonalen Gitters ausgeführt. Um Faltungen an den übrigen Kanten auszuführen, könnten wir analog vorgehen. Wir können aber auch auf die strukturellen Eigenschaften des hexagonalen Gitterkomplexes zurückgreifen und dessen Automorphismen verwenden. Dies geht in letzter Instanz auf Bemerkung 4.21 zurück, die aussagt, dass die Teilfaltungsrelation durch Automorphismen erhalten bleibt. Da die Definition 4.35 unendlicher Faltungen auf Teilfaltungen beruht (wir wenden endlich viele Faltpläne an, womit wir eine Teilfaltung erhalten), überträgt sich die Aussage der Bemerkung auch auf unseren Fall.

Wir müssen aber überprüfen, ob sich auch alle unsere Argumente übertragen. Beim Durchlesen der vorherigen Abschnitte bemerken wir die folgenden Einschränkungen:

- Der Automorphismus muss Geraden auf Gerade abbilden. Da alle unsere Automorphismen durch affine Abbildungen der Ebene induziert sein werden, ist dies automatisch erfüllt.
- Wir müssen nach der Faltung wieder im hexagonalen Gitter landen. Es ist aber klar, dass ein Teil des Gitters durch einen Automorphismus auf einen anderen Teil des Gitters abgebildet wird.

Damit wird unsere Analyse durch Anwendung von Automorphismen nicht berührt. Wir können uns also allein darauf konzentrieren, affine Automorphismen zu finden, die jede Kante auf die Kante der bereits ausgeführten Faltung zurückführen. Dazu ist es hinreichend zu zeigen, dass solche Automorphismen transitiv auf den Kanten operieren.

Es genügt, sich die folgenden Automorphismen anzusehen (dass es sich um Automorphismen handelt, ist leicht nachzurechnen):

1. Die Translation um $\vec{a}$:

$$\tau_{\vec{a}} : \mathcal{K}_0 \rightarrow \mathcal{K}_0, \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \mapsto \begin{pmatrix} \alpha + 1 \\ \beta \end{pmatrix}$$

2. Die Translation um $\vec{b}$:

$$\tau_{\vec{b}} : \mathcal{K}_0 \rightarrow \mathcal{K}_0, \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \mapsto \begin{pmatrix} \alpha \\ \beta + 1 \end{pmatrix}$$

3. Die Drehung um 60° :

$$\rho : \mathcal{K}_0 \rightarrow \mathcal{K}_0, \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \mapsto \begin{pmatrix} 1 & -1 \\ 1 & 0 \end{pmatrix} \cdot \begin{pmatrix} \alpha \\ \beta \end{pmatrix} = \begin{pmatrix} \alpha - \beta \\ \alpha \end{pmatrix}$$

Durch Anwendung der Translationen lässt sich jede Kante auf eine der folgenden Formen bringen:

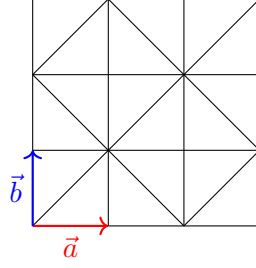
$$\left\{ \begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \right\} \quad \left\{ \begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right\} \quad \left\{ \begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right\}$$

Durch die Verwendung der Drehung ρ erkennt man, dass alle diese in einer Bahn liegen. Folglich überträgt sich die Faltung aus Abschnitt 4.2.2 auf jede beliebige Kante des Musters.

4.2.4 Beschreibung des quadratischen Gitters

Wir haben uns ab Abschnitt 4.2.1 damit beschäftigt, ein ebenes Muster aus gleichseitigen Dreiecken auf ein Dreieck zusammenzufalten. In diesem Abschnitt wollen wir nach demselben Prinzip ebene Muster aus rechtwinkligen, gleichschenkligen Dreiecken untersuchen, die wir quadratische Gitter nennen.

Wir verwenden die Standardbasisvektoren als $\mathbb{Z}$ -Basis $(\vec{a}, \vec{b})$ der Gitterpunkte, bezüglich der wir alle Punkte im $\mathbb{R}^2$ darstellen werden.



Damit lauten die Punkte:

$$\mathcal{K}_0 := \left\{ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \right\} \mid \alpha, \beta \in \mathbb{Z} \right\}.$$

Die Kantenbeschreibung wird ein wenig komplizierter als im hexagonalen Fall. Die Anwendung von $\vec{a}$ oder $\vec{b}$ erzeugt stets eine Kante. Die diagonalen Kanten kommen aber nicht an jeder Stelle vor, sondern nur an jeder zweiten. Daher lautet die Kantenmenge:

$$\begin{aligned} \mathcal{K}_1 := & \left\{ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta \end{pmatrix} \right\} \mid \alpha, \beta \in \mathbb{Z} \right\} \\ & \uplus \left\{ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha \\ \beta+1 \end{pmatrix} \right\} \mid \alpha, \beta \in \mathbb{Z} \right\} \\ & \uplus \left\{ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta+1 \end{pmatrix} \right\} \mid \alpha, \beta \in \mathbb{Z}, \alpha + \beta \text{ gerade} \right\} \\ & \uplus \left\{ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha-1 \\ \beta+1 \end{pmatrix} \right\} \mid \alpha, \beta \in \mathbb{Z}, \alpha + \beta \text{ gerade} \right\} \end{aligned}$$

Zur Beschreibung der Flächen wählen wir zu jedem „geraden“ Punkt alle Dreiecke, die diesen Punkt enthalten und „darüber“ liegen.

$$\begin{aligned} \mathcal{K}_2 := & \left\{ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta+1 \end{pmatrix} \right\} \mid \alpha, \beta \in \mathbb{Z}, \alpha + \beta \text{ gerade} \right\} \\ & \uplus \left\{ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ \beta+1 \end{pmatrix}, \begin{pmatrix} \alpha \\ \beta+1 \end{pmatrix} \right\} \mid \alpha, \beta \in \mathbb{Z}, \alpha + \beta \text{ gerade} \right\} \\ & \uplus \left\{ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha \\ \beta+1 \end{pmatrix}, \begin{pmatrix} \alpha-1 \\ \beta+1 \end{pmatrix} \right\} \mid \alpha, \beta \in \mathbb{Z}, \alpha + \beta \text{ gerade} \right\} \\ & \uplus \left\{ \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix}, \begin{pmatrix} \alpha-1 \\ \beta+1 \end{pmatrix}, \begin{pmatrix} \alpha-1 \\ \beta \end{pmatrix} \right\} \mid \alpha, \beta \in \mathbb{Z}, \alpha + \beta \text{ gerade} \right\} \end{aligned}$$

Die Reihenfolge, mit der wir die Punkte in diese Mengen geschrieben haben, ist auch die Reihenfolge, die wir für die simpliziale Orientierung der entsprechenden Fläche wählen. Da es sich um einen geschlossenen Flächenkomplex handelt, haben wir gemäß Bemerkung 4.6 einen

2-Faltkomplex vorliegen. Wir wollen wie im hexagonalen Fall falten können. Dazu nutzen wir die Erkenntnis aus Abschnitt 4.2.3, dass sich der Aufwand durch geschickte Anwendung der Automorphismengruppe reduzieren lässt. Wir beginnen also damit, einige affine Automorphismen anzugeben und deren Bahnen auf der Kantenmenge zu bestimmen. Wir verwenden die folgenden Automorphismen:

- Die Translation um $\vec{a}$:

$$\tau_{\vec{a}} : \mathcal{K}_0 \rightarrow \mathcal{K}_0, \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \mapsto \begin{pmatrix} \alpha + 1 \\ \beta \end{pmatrix}$$

- Die Translation um $\vec{b}$:

$$\tau_{\vec{b}} : \mathcal{K}_0 \rightarrow \mathcal{K}_0, \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \mapsto \begin{pmatrix} \alpha \\ \beta + 1 \end{pmatrix}$$

- Drehung um 90° :

$$\rho : \mathcal{K}_0 \rightarrow \mathcal{K}_0, \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \mapsto \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix} \cdot \begin{pmatrix} \alpha \\ \beta \end{pmatrix} = \begin{pmatrix} -\beta \\ \alpha \end{pmatrix}$$

Nach Anwendung der beiden Translationen dürfen wir annehmen, dass $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$ in der Kante vorkommt. Der andere Endpunkt liegt dann in der Menge¹⁶

$$\left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \mid \alpha \in \{-1, 0, 1\}, \beta \in \{-1, 0, 1\}, (\alpha, \beta) \neq (0, 0) \right\}.$$

Als nächstes wenden wir die Drehung ρ an, unter der diese Menge in zwei vierelementige Bahnen zerfällt. Es verbleiben also diese beiden Kanten:

$$\left\{ \begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \right\} \qquad \left\{ \begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right\}$$

Da diese beiden Kanten nicht durch Automorphismen auseinander hervorgehen (am Punkt $\begin{pmatrix} 1 \\ 0 \end{pmatrix}$ stoßen vier Kanten zusammen, während an allen anderen Eckpunkten acht Kanten zusammenstoßen), müssen wir zwei Faltungen untersuchen. Es ist nicht allzu schwer, die Voraussetzungen von Satz 4.37 zu überprüfen.

Betrachten wir zuerst die Spiegelung an der Kante $\left\{ \begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 0 \end{pmatrix} \right\}$ mit der Abbildungsvorschrift

$$\mu : \mathcal{K}_0 \rightarrow \mathcal{K}_0, \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \mapsto \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \cdot \begin{pmatrix} \alpha \\ \beta \end{pmatrix} = \begin{pmatrix} \alpha \\ -\beta \end{pmatrix}.$$

Wenn man bedenkt, dass $\alpha + \beta$ durch Addition von -2β seine Parität nicht verändert, ist es leicht nachzuweisen, dass es sich bei μ um einen Automorphismus handelt und dass

¹⁶Man kann diese Menge allein durch Translationen noch weiter verkleinern. Dies ist in Hinblick auf die weitere Argumentation aber nicht notwendig und verkompliziert die Darstellung an dieser Stelle unnötig.

dieser die Orientierungen der Flächen umkehrt. Wir sehen auch leicht, dass μ nur die Kanten $\left\{ \begin{pmatrix} \alpha \\ 0 \end{pmatrix}, \begin{pmatrix} \alpha+1 \\ 0 \end{pmatrix} \right\}$ stabilisiert, deren Eckpunkte selbst invariant unter μ sind. Damit lässt sich die Faltung entlang μ ausführen. Um nachzuweisen, dass wir wieder im quadratischen Gitter landen, beachten wir, dass μ eine Bijektion zwischen den folgenden beiden Mengen induziert:

$$\left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \middle| \alpha, \beta \in \mathbb{Z}, \beta \geq 0 \right\} \qquad \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \middle| \alpha, \beta \in \mathbb{Z}, \beta \leq 0 \right\}$$

Betrachten wir nun die Spiegelung an der anderen Kante $\left\{ \begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right\}$, die die Abbildungsvorschrift

$$\mu : \mathcal{K}_0 \rightarrow \mathcal{K}_0, \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \mapsto \begin{pmatrix} -1 & 0 \\ 0 & -1 \end{pmatrix} \cdot \begin{pmatrix} \alpha \\ \beta \end{pmatrix} = \begin{pmatrix} -\alpha \\ -\beta \end{pmatrix}$$

hat. Der Nachweis der Voraussetzungen von Satz 4.37 ist auch hier simpel. Wir geben noch die beiden Mengen an, zwischen denen μ eine Bijektion induziert:

$$\left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \middle| \alpha, \beta \in \mathbb{Z}, \alpha + \beta \geq 0 \right\} \qquad \left\{ \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \middle| \alpha, \beta \in \mathbb{Z}, \alpha + \beta \leq 0 \right\}$$

Daher lässt sich jede Überlagerung eines Gitterteilbereichs auf ein Dreieck zusammenfallen.

5 Modellerweiterungen

In Kapitel 2 haben wir ein mengentheoretisches Modell von Faltungen geschaffen, das wir in Kapitel 3 um die Reihenfolge der Flächen ergänzt haben. Das Kapitel 4 war ersten Ansätzen in der Einbettungsfrage gewidmet. In diesem Kapitel werden wir auf mögliche Verallgemeinerungen unseres abstrakten Faltmodells eingehen. Wir interessieren uns dabei primär für den Fall $n = 2$, also das Falten von Flächen im Raum, da es dort viele Anwendungen gibt, die ein etwas allgemeineres Modell erfordern.

Wir werden zwei Annahmen der 2-Faltkomplexe auf ihre Verallgemeinerbarkeit untersuchen und jeweils skizzieren, wie eine solche Verallgemeinerung aussehen müsste:

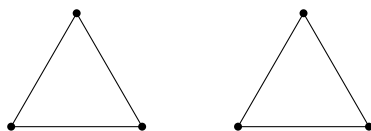
- Das Modell der Faltkomplexe ist nur im lokal simplizialen Fall gültig. Es ist nicht ausformuliert, was z. B. bei viereckigen Flächen passiert. Wir werden in Abschnitt 5.1 genauer auf diesen Fall eingehen.
- Wir haben in der Regel stillschweigend angenommen, dass die Dreiecke unserer Faltkomplexe kongruent sind (ansonsten kann man benachbarte Dreiecke nicht aufeinander falten). In Abschnitt 5.2 geben wir einige Hinweise, wie man 2-Faltkomplexe auf diese Situation anpassen kann.

5.1 Allgemeine m -Ecke als Flächen

In Kapitel 3 haben wir lokale Simplizität in unserer Definition von n -Faltkomplexen vorausgesetzt. Das haben wir primär deswegen getan, um den Begriff der simplizialen Orientierung (und damit des induzierten Durchlaufsinns) definieren zu können. Wir würden aber auch gerne die Faltung von quadratischen Flächen erlauben. Dass dies prinzipiell möglich ist, sehen wir daran, dass der Großteil der Theorie zu n -Überlagerungskomplexen aus Kapitel 2.2 nicht auf der lokalen Simplizität beruht. Wir werden unser Modell in diesem Abschnitt für beliebige m -Ecke formulieren, aber nicht auf die Einbettungsproblematik eingehen.

Zunächst beschreiben wir, wie sich homogene 2-Komplexe auf m -Ecke spezialisieren, ähnlich dazu, wie die lokale Simplizität aus Abschnitt 2.1.1 die Spezialisierung auf Dreiecke beschrieben hat. Dazu müssen wir nur die Bedingung der lokalen Simplizität abändern. Unsere Definition von 2-Komplexen ist dafür abstrakt genug – schließlich legen wir dort nur eine Menge $\mathcal{K}_2$ fest, die auch eine Menge von m -Ecken sein kann. Eine grundlegende Eigenschaft von m -Ecken ist, dass sie genau m Kanten und genau m Punkte haben.

Diese Eigenschaft legt ein m -Eck aber nicht eindeutig fest, da die Kanten und Punkten geeignet interagieren müssen. Nur zu fordern, dass jeweils zwei Kanten einen gemeinsamen Punkt haben, ist nicht ausreichend, wie das Muster



demonstriert. Es handelt sich nämlich um kein Sechseck, obwohl es jeweils sechs Punkte und Kanten hat, sodass jeweils zwei Kanten einen gemeinsamen Punkt haben. Um diese Situation auszuschließen, betrachten wir die Kantenabfolge selbst. Wir legen dazu eine zyklische

Reihenfolge der Kanten fest, also für jedes $f \in \mathcal{K}_2$ eine bijektive Abbildung

$$\kappa_f : (\mathbb{Z}/m\mathbb{Z}, +) \rightarrow \{l \in \mathcal{K}_1 \mid l \prec f\},$$

sodass $\kappa_f(i)$ und $\kappa_f(i+1)$ genau einen Punkt gemeinsam haben¹⁷. Diese Beschreibung lässt sich zwar nicht in höhere Dimensionen übertragen (es gibt nur in der Ebene eine solche eindeutige Umlaufrichtung), genügt aber für unsere Zwecke.

Definition 5.1 (lokal zellulär). *Ein homogener 2-Komplex $((\mathcal{K}_0, \mathcal{K}_1, \mathcal{K}_2), \prec)$ heißt **lokal zellulär**, falls gilt*

1. *Für jedes $y \in \mathcal{K}_1$ gibt es genau zwei $x_1, x_2 \in \mathcal{K}_0$, sodass $x_j \prec y$ für $j \in \{1, 2\}$ gilt.*
2. *Für jedes $z \in \mathcal{K}_2$ gibt es ein $m \in \mathbb{N}_{\geq 3}$ und genau m Elemente $y_1, \dots, y_m \in \mathcal{K}_1$ mit $y_j \prec z$ für $j \in \{1, \dots, m\}$. Zudem existiert eine bijektive Abbildung*

$$\kappa_z : (\mathbb{Z}/m\mathbb{Z}, +) \rightarrow \{y_1, \dots, y_m\}$$

mit der Eigenschaft, dass es zu $\kappa_z(j)$ und $\kappa_z(j+1)$ genau ein $x_j \in \mathcal{K}_0$ gibt, das $x_j \prec y_j$ und $x_j \prec y_{j+1}$ erfüllt (Indizes sind modulo m zu lesen).

Mit dieser Definition kann jede Fläche ein beliebiges m -Eck sein. Wenn jedes $x \in \mathcal{K}_2$ ein Dreieck ist, erhalten wir den Fall von lokaler Simplität. Wir haben damit herausgestellt, wie sich homogene n -Komplexe auf den Fall beliebiger m -Ecke spezialisieren.

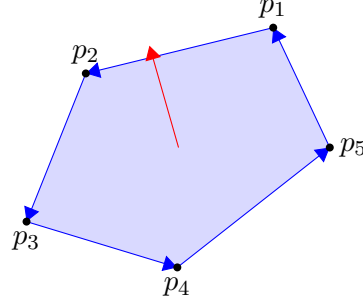
Wir untersuchen zunächst, wie sich diese Verallgemeinerung bei n -Überlagerungskomplexen auswirkt. Die lokale Simplität hatte die Frage vereinfacht, wann eine Identifikation eine gültige Erweiterung liefert. Im generischen Fall in Satz 2.32 wurde eine Erweiterung ausgeschlossen, falls zwei zu identifizierende Punkte p und q in einer gemeinsamen Fläche liegen. Unter der Annahme der lokalen Simplität konnten wir diese Bedingung in Folgerung 2.36 dazu vereinfachen, dass wir nur überprüfen müssen, ob p und q in einer gemeinsamen Kante liegen. Dass diese Vereinfachung bei allgemeinen m -Ecken nicht funktioniert, lässt sich bereit an einem Viereck einsehen. Die gegenüberliegenden Ecken dürfen nicht identifiziert werden, aber sie sind durch keine Kante verbunden.

Wir versuchen als nächstes, die Annahme der lokalen Simplität aus der Definition der 2-Faltkomplexe durch die Annahme der lokalen Zellularität zu ersetzen. Die lokale Simplität ist nur dafür notwendig, simpliziale Orientierungen (vgl. Definition 3.4) und induzierte Durchlaufsinne (vgl. Definition 3.17) zu definieren. Wir müssen diese Begriffe also auf Basis der lokalen Zellularität definieren. Zur Verallgemeinerung der simplizialen Orientierung benötigen wir den Begriff eines m -Ecks:

Definition 5.2 (m -Eck). *Ein lokal zellulärer homogener 2-Komplex $((\mathcal{K}_0, \mathcal{K}_1, \mathcal{K}_2), \prec)$ heißt **m -Eck**, falls $\mathcal{K}_2$ einelementig ist und $\mathcal{K}_0$ genau m Elemente enthält.*

Um die Orientierung eines m -Ecks zu definieren, stellen wir es uns in der Ebene eingebettet vor. Dann lassen sich die Orientierungen der Ebene auf das m -Eck übertragen. Diese Orientierungen sind identisch zu den Punktumlaufrichtungen des m -Ecks.

¹⁷Formal gesehen sollten wir $\kappa_f(i + m\mathbb{Z})$ anstelle von $\kappa_f(i)$ schreiben.



Wir können die Orientierung eines m -Ecks daher koordinatenunabhängig über dessen Punktumläufe angeben.

Definition 5.3. Sei $\mathcal{Z} = ((\mathcal{K}_0, \mathcal{K}_1, \mathcal{K}_2), \prec)$ ein m -Eck. Ein m -Tupel $(p_1, \dots, p_m) \in \mathcal{K}_0^m$ heißt **zelluläre Orientierung (von $\mathcal{Z}$)**, falls gilt:

- Für $i \neq j$ gilt $p_i \neq p_j$
- Die Punkte p_i und p_{i+1} liegen in einer gemeinsamen Kante für alle $1 \leq i \leq m$, wobei wir die Indizes modulo m lesen.

Die Menge aller zellulären Orientierungen schreiben wir als $\text{ZellOrient}(\mathcal{Z})$.

Ebenso wie bei den simplizialen Orientierungen aus Definition 3.4 definieren mehrere dieser zellulären Orientierungen die selbe Umlaufrichtung. Für simpliziale Orientierungen haben wir diese Schwierigkeit in Bemerkung 3.5 dadurch gelöst, dass wir eine Gruppenoperation auf den simplizialen Orientierungen definiert haben. Um diese Idee auch bei zellulären Orientierungen anzuwenden, müssen wir uns damit beschäftigen, welche Gruppe auf diesen operiert.

Betrachten wir eine zelluläre Orientierung $(p_1, p_2, \dots, p_m)$, so sind sowohl $(p_2, \dots, p_m, p_1)$ als auch $(p_m, \dots, p_2, p_1)$ zelluläre Orientierungen. Diese beiden Permutationen erzeugen eine Diedergruppe mit $2m$ Elementen, wie wir in Bemerkung 5.4 nachweisen.

Bemerkung 5.4. Sei $m \in \mathbb{Z}_{\geq 3}$. Die Abbildungen

$$\tilde{z} : \{1, \dots, m\} \rightarrow \{1, \dots, m\}, \quad i \mapsto \begin{cases} i+1 & i < m \\ 1 & i = m \end{cases}$$

und

$$\tilde{s} : \{1, \dots, m\} \rightarrow \{1, \dots, m\}, \quad i \mapsto m+1-i$$

sind Permutationen aus der S_{2m} und erzeugen eine Diedergruppe mit $2m$ Elementen.

Beweis. Wir wollen nachweisen, dass die Gruppe $\langle \tilde{z}, \tilde{s} \rangle \leq S_{2m}$ isomorph zur Diedergruppe mit $2m$ Elementen ist, die durch ihre Präsentation $D_{2m} := \langle a, b \mid a^m, b^2, (ab)^2 \rangle$ definiert ist.

Wir betrachten dazu den Epimorphismus der freien Gruppe mit zwei Erzeugern nach $\langle \tilde{z}, \tilde{s} \rangle$:

$$\eta : F(a, b) \rightarrow \langle \tilde{z}, \tilde{s} \rangle \quad \begin{cases} a \mapsto \tilde{z} \\ b \mapsto \tilde{s} \end{cases}$$

und suchen den Kern dieser Abbildung.

Es ist leicht einzusehen, dass $\tilde{z}^m = Id$ und $\tilde{s}^2 = Id$ gelten. Wir berechnen $\tilde{z}\tilde{s}(i)$ für jedes $i \in \{1, \dots, m\}$:

$$\begin{aligned}\tilde{z}\tilde{s}(i) &= \begin{cases} (m+1-i)+1 & m+1-i < m \\ 1 & m+1-i = m \end{cases} \\ &= \begin{cases} m+2-i & 1 < i \\ 1 & 1 = i \end{cases}\end{aligned}$$

Damit folgt $(\tilde{z}\tilde{s})^2 = Id$, weswegen das Normalteilererzeugnis $\langle a^m, b^2, (ab)^2 \rangle_{F(a,b)} \leq F(a,b)$ im Kern von η liegt. Faktorisieren wir diesen Normalteiler heraus, erhalten wir den Epimorphismus $\hat{\eta}$:

$$\hat{\eta} : \langle a, b | a^m, b^2, (ab)^2 \rangle \rightarrow \langle \tilde{z}, \tilde{s} \rangle \quad \begin{cases} a \mapsto \tilde{z} \\ b \mapsto \tilde{s} \end{cases}$$

Wir müssen noch zeigen, dass $\hat{\eta}$ ein Isomorphismus ist. Da $\hat{\eta}$ bereits ein Epimorphismus ist und $\langle a, b | a^m, b^2, (ab)^2 \rangle$ aus genau $2m$ Elementen besteht, genügt es zu zeigen, dass $\langle \tilde{z}, \tilde{s} \rangle$ mindestens $2m$ Elemente besitzt.

Dazu betrachten wir die zyklische Untergruppe $\langle \tilde{z} \rangle \leq \langle \tilde{z}, \tilde{s} \rangle$. Da aus $(\tilde{z}\tilde{s})^2 = Id$ schon $\tilde{s}\tilde{z}\tilde{s} = \tilde{z}^{-1}$ folgt und $\tilde{s} = \tilde{s}^{-1}$ gilt, handelt es sich bei $\langle \tilde{z} \rangle$ um einen Normalteiler von $\langle \tilde{z}, \tilde{s} \rangle$. Wegen $\langle \tilde{z}, \tilde{s} \rangle / \langle \tilde{z} \rangle = \langle \tilde{s} \rangle$ folgt

$$|\langle \tilde{z}, \tilde{s} \rangle| = |\langle \tilde{z} \rangle| \cdot |\langle \tilde{s} \rangle| = m \cdot 2,$$

weswegen $\hat{\eta}$ ein Isomorphismus ist. □

Die Diedergruppe aus Bemerkung 5.4 operiert auf den zellulären Orientierungen. Die Bahnen unter der zyklischen Untergruppe $\langle \tilde{z} \rangle$ entsprechen dabei den Punktlaufrichtungen, während die Anwendung von $\tilde{s}$ diese umdreht.

Bemerkung 5.5. *Sei $\mathcal{Z}$ ein m -Eck. Die Diedergruppe $\langle \tilde{z}, \tilde{s} \rangle$ aus Bemerkung 5.4 operiert transitiv auf der Menge der zellulären Orientierungen von $\mathcal{Z}$, vermöge*

$$\pi \cdot (p_1, \dots, p_m) = (p_{\pi(1)}, \dots, p_{\pi(m)}).$$

Die zyklische Gruppe $\langle \tilde{z} \rangle \leq \langle \tilde{z}, \tilde{s} \rangle$ hat für $m \geq 3$ genau zwei Bahnen auf $\text{ZellOrient}(\mathcal{Z})$.

Mit Bemerkung 5.5 haben wir das Konzept der simplizialen Orientierung auf den lokal zellulären Fall übertragen. Sobald wir auch den Begriff der induzierten Durchlaufsinne übertragen haben, lassen sich alle Aussagen über 2-Faltkomplexe (insbesondere der Faltplan-Formalismus), die nicht auf der Anzahl der Kanten beruhen, sofort auf den lokal zellulären Fall übertragen.

Bevor wir uns den induzierten Durchlaufsinnen widmen können, müssen wir uns mit einer Verallgemeinerung der simplizialen Oberfläche aus Definition 3.11 beschäftigen. Diese lässt

sich analog definieren, indem man alle zellulären Orientierungen einer Umlaufrichtung auf +1 und die anderen auf -1 abbildet. Wir können die Verträglichkeit mit der Gruppenoperation aber nicht wie im simplizialen Fall durch das Signum definieren, sondern müssen sie in Bezug auf die Erzeuger $\tilde{z}$ und $\tilde{s}$ direkt angeben.

Definition 5.6 (zelluläre Oberfläche). Sei $\mathcal{S} = ((\mathcal{K}_0, \mathcal{K}_1, \mathcal{K}_2), \prec)$ ein lokal zellulärer homogener 2-Komplex. Eine **zelluläre Oberfläche** von $\mathcal{S}$ ist eine Familie von Abbildungen

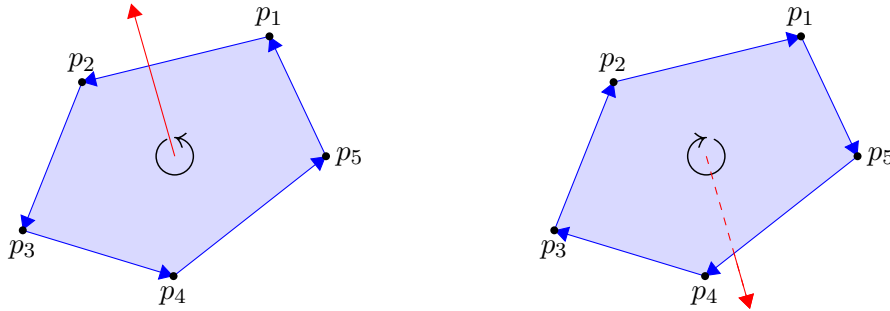
$$\sigma_f : \text{ZellOrient}(\mathcal{S}_{\preccurlyeq f}) \rightarrow \{\pm 1\}, \quad f \in \mathcal{K}_n,$$

sodass für die Operation der Diedergruppe $\langle \tilde{z}, \tilde{s} \rangle$ aus Bemerkung 5.4 die Relationen

$$\begin{aligned} \sigma_f(\tilde{z} \cdot a) &= \sigma_f(a) \\ \sigma_f(\tilde{s} \cdot a) &= -\sigma_f(a), \end{aligned}$$

für jedes $a \in \text{ZellOrient}(\mathcal{S}_{\preccurlyeq f})$ gelten.

Wir können jetzt die induzierten Durchlaufsinne betrachten. Diese beruhen im lokal simplizialen Fall darauf, dass die simpliziale Orientierung einer Kante bereits die simpliziale Orientierung des gesamten Dreiecks festlegt. Aber auch bei m -Ecken ist diese Eigenschaft gegeben, wie die folgenden Graphiken zeigen.



Der Kantenorientierung (p_4, p_5) wird z. B. die zelluläre Orientierung $(p_4, p_5, p_1, p_2, p_3)$ zugewiesen. Dass dieses Vorgehen allgemein funktioniert, zeigen wir in Definition 5.7.

Definition 5.7 (induzierter Durchlaufsinne). Sei $\mathcal{S} = ((\mathcal{K}_0, \mathcal{K}_1, \mathcal{K}_2), \prec, \sim)$ ein lokal zellulärer 2-Überlagerungskomplex mit zellulärer Oberfläche $\{\sigma_f\}_{f \in \mathcal{K}_n}$. Sei $[k] \in \mathcal{K}_{n-1}/\sim$ gegeben. Dann ist für jedes $f \in \text{Cor}([k])$ der **induzierte Durchlaufsinne** definiert als die Abbildung

$$\begin{aligned} \bar{\sigma}_f : \text{SimplOrient}(\mathcal{S}_{\preccurlyeq [k]}) &\rightarrow \{\pm 1\} \\ ([a_1], [a_2]) &\mapsto \sigma_f((b_1, b_2, \dots, b_{m+1})), \end{aligned}$$

mit $b_i \in [a_i]$ für alle $1 \leq i \leq 2$ und $(b_1, b_2, \dots, b_m) \in \text{ZellOrient}(\mathcal{S}_{\preccurlyeq f})$.

Wohldefiniertheit. Wegen $[k] \prec [f]$ enthalten $[a_1]$ und $[a_2]$ jeweils einen Punkt aus $\mathcal{K}_0$, der in f liegt. Gäbe es mehr als einen solchen Punkt in $[a_i]$, wäre die Irreduzibilität von $\sim$ verletzt. Folglich sind b_1 und b_2 eindeutig festgelegt und liegen auf einer gemeinsamen Kante.

Wir müssen noch zeigen, dass es genau eine zelluläre Orientierung $(p_1, \dots, p_m)$ mit $p_1 = b_1$ und $p_2 = b_2$ gibt. Wir betrachten zuerst die Existenz. Wir zeigen dazu, dass wir diese zelluläre Orientierung durch die Gruppenoperation aus Bemerkung 5.5 erreichen können. Sei also die zelluläre Orientierung $(p_1, \dots, p_m)$ beliebig gegeben. Da b_1 und b_2 durch eine Kante verbunden sind, gibt es ein $1 \leq i \leq m$ mit

$$p_i = b_1 \qquad b_2 \in \{p_{i-1}, p_{i+1}\}.$$

Falls $b_2 = p_{i-1}$ gilt, wenden wir eine Spiegelung aus D_{2m} an. Daher können wir ohne Einschränkung von $b_2 = p_{i+1}$ ausgehen. Nach Anwendung einer Drehung erhalten wir

$$p_1 = b_1 \qquad p_2 = b_2.$$

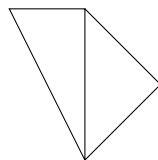
Wir müssen noch die Eindeutigkeit nachweisen. Innerhalb des m -Ecks $\mathcal{S}_{\prec f}$ ist jeder Punkt $p \in \mathcal{K}_0$ in genau zwei Kanten enthalten. Damit sind die beiden „Nachbarn“ von p in einer zellulären Orientierung eindeutig festgelegt. Sobald wir b_1 und b_2 festgelegt haben, sind damit alle weiteren b_i induktiv festgelegt. $\square$

Wir haben die Begriffe der simplizialen Orientierung und des induzierten Durchlaufsinns auf den lokal zellulären Fall verallgemeinert. Da die lokale Simplizität an keiner anderen Stelle (außer den Einbettungen) für 2-Faltkomplexe relevant ist, haben wir damit 2-Faltkomplexe auf den lokal zellulären Fall verallgemeinert.

5.2 Unregelmäßige Flächen

Wir haben uns in dieser Arbeit nur sehr selten mit dem Kongruenztyp von Dreiecken befasst. Diese sind für unsere 2-Faltkomplexe nur dann relevant, wenn wir diese einbetten wollen, wie in Abschnitt 2.5.2. In diesem Abschnitt haben wir die zugelassenen Faltungen mit den erlaubten Kongruenztypen der Dreiecke in Verbindung gebracht. Beispielsweise zeigt Bemerkung 2.63, dass nur gleichseitige Dreiecke die Faltung von zwei beliebigen Flächen erlauben. Selbst wenn wir nur die Faltung benachbarter Flächen erlauben, können nach Lemma 2.64 im generischen Fall höchstens annehmen, dass alle Dreiecke kongruent sein müssen.

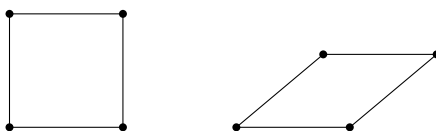
Diese impliziten Annahmen an die Flächen sind aber in vielen Anwendungen nicht erfüllt. Häufig wissen wir bereits, dass unsere Dreiecke nicht gleichseitig oder kongruent sind. In diesem Fall gibt es Identifikationen, die eine gültige Faltung des 2-Faltkomplexes liefern, aber nicht ausführbar sind, wenn man diesen einbettet. Das folgende Bild illustriert eine solche Situation:



Die beiden Dreiecke können nicht deckungsgleich aufeinandergefaltet werden, obwohl sie als 2-Faltkomplexe faltbar sind.

Um diese Modifikation darzustellen, müssen wir die Menge der möglichen Identifikationen einschränken. Für Dreiecke können wir z. B. zu jeder Kante eine Länge speichern und nur solche Identifikationen φ zulassen, bei denen für alle Kanten k sowohl k als auch $\varphi(k)$ die selbe Länge haben.

Da wir aber auch allgemeine m -Ecke als Flächen zulassen wollen (vgl. den vorigen Abschnitt 5.1), ist eine Speicherung der Kantenlängen nicht ausreichend. Die folgenden Vierecke haben zwar die gleichen Kantenlängen, sind aber nicht isometrisch:



Wir benötigen also eine Zusatzinformation, die es uns erlaubt, zwischen verschiedenen Isometrieklassen zu unterscheiden. Um ein m -Eck bis auf Isometrie eindeutig zu rekonstruieren, genügt es, sich zu jeder Kante deren Länge sowie zu jedem Punkt den zugehörigen Innenwinkel zu merken. Wir müssen bei einer Identifikation dann nur darauf achten, dass diese Kennzahlen bei aufeinandergefalteten Elementen identisch sind.

Bei dieser Art der Speicherung ist die Größe der Winkel oder Längen eigentlich unwichtig – es ist nur notwendig, zwei Winkel oder Längen auf Gleichheit zu überprüfen. Daher können wir z. B. die Winkel der Raute in einem Punktzug durch den Zykel (a, b, a, b) ausdrücken, womit wir zum Ausdruck bringen, dass benachbarte Winkel verschieden, aber gegenüberliegende Winkel gleich sind. Relevant ist an dieser Stelle auch der Vergleich zum Quadrat, dessen Winkel wir als (c, c, c, c) schreiben würden. Während der Zykel der Raute die Winkelverteilung *jeder* Raute (mit Ausnahme des Quadrats) beschreibt, beschreibt der Zykel des Quadrats ausschließlich die Winkelverteilung des Quadrats. Wir schließen daraus, dass wir mit unserem Formalismus direkt die Faltung aller möglichen Rauten (außer dem Quadrat) simultan bestimmen können.

Diese Eigenschaft ist für praktische Anwendungen sehr interessant, da es dort häufiger vorkommt, dass man noch nicht genau weiß, mit welchen Winkeln das Faltmuster ausgestattet sein soll, wenn man das Flächenmuster vorliegen hat. Während konventionelle Modelle für jede neue Winkelangabe neue Berechnungen anstellen müssen (da sich dadurch die Flächen im $\mathbb{R}^3$ verändern), ist unser Formalismus gegen solche Probleme immun. Wenn der Formalismus interessante Faltungen findet, muss man allerdings noch immer überprüfen, ob die gefundenen Faltungen auch tatsächlich im $\mathbb{R}^3$ ausführbar sind.

Literatur

- [ABD⁺01] Esther M. Arkin, Michael A. Bender, Erik D. Demaine, Martin L. Demaine, Joseph S. B. Mitchell, Saurabh Sethia, and Steven S. Skiena, *When can you fold a map?*, Algorithms and Data Structures: 7th International Workshop, WADS 2001 Providence, RI, USA, August 8–10, 2001 Proceedings (Berlin, Heidelberg) (Frank Dehne, Jörg-Rüdiger Sack, and Roberto Tamassia, eds.), 2001, pp. 401–413.
- [Bau14] Markus Baumeister, *Periodische Faltmuster mit Singularitäten*, Bachelorarbeit, RWTH Aachen University, September 2014.
- [BNPS] Karl-Heinz Brakhage, Alice Niemeyer, Wilhelm Plesken, and Ansgar Strzelczyk, *Simplicial surfaces controlled by one triangle*, zu erscheinen.
- [COO57] R. J. Wisner C. O. Oakley, *Flexagons*, The American Mathematical Monthly **64** (1957), no. 3, 143–154.
- [DO10] Erik D. Demaine and Joseph O’Rourke, *Geometric Folding Algorithms, Linkages, Origami, Polyhedra*, Cambridge University Press, 32 Avenue of the Americas, New York, 2010.
- [Hul94] Thomas Hull, *On the mathematics of flat origamis*, Congressus Numerantium **100** (1994), 215–224.
- [Lan04] Robert Lang, *Orchestra*, <http://www.langorigami.com/artwork/orchestra>, 2004, abgerufen am 20. September 2016.
- [Lan06] ———, *Miura-ken beauty rose*, <http://www.langorigami.com/artwork/miura-ken-beauty-rose-opus-482>, 2006, abgerufen am 20. September 2016.
- [PH97] Hans Walser Peter Hilton, Jean Pedersen, *The faces of the tri-hexaflexagon*, Mathematics Magazine **70** (1997), no. 4, 243–251.
- [Whe58] Roger F. Wheeler, *The flexagon family*, The Mathematical Gazette **42** (1958), no. 339, 1–6.

Index

- m -Eck, 152
- Überlagerungskomplex, 18
 - Einschränkung, 16
 - wegzusammenhängender, 130
- Überschneidungsfreiheit, 80, 109
- Abstand
 - gerader, 129
 - ungerader, 129
- Abwärtskompatibilität, 18
 - starke, 18
- Anomalie, 37
 - klasse, 37
- anomaliefrei, 37
- Basiskomplex, 41
- benachbart, 107
- Bindungsrelation, 34
- Corona, 48
- Dimensionalität, 18
- Durchlaufsinne
 - induzierter, 55, 155
- Einbettung, 38
 - Falt-, 106
 - konsistente, 43
 - simpliziale, 38
- Einschränkung, 16
- Entfaltung
 - Faltplan, 98
 - primitive, 85
 - zulässige, 97
- entgegengesetzt orientiert, 60
- Epsilon-Nachbar, 59
- Erweiterung, 21
 - formale, 21
 - primitive, 23
- Fächer, 48
 - Corona, 48
 - induzierte Fächerpartition, 86
 - kompatible, 56
 - Summe, 72
 - voller, 48
 - zyklische Reihung, 46
- Fächersumme, 72
- färbbar, 39
- Färbung, 39
- Färbungsüberlagerung, 40
- Falteinbettung, 106
- Faltkomplex, 62
 - Einschränkung, 16
 - Fächer, 48
 - stabiler, 96
 - voller Fächer, 48
 - wegzusammenhängender, 130
- Faltplan, 92
 - Anwendung, 93, 98
 - zulässiger, 92
- Faltung
 - Faltplan, 93
 - primitive, 64, 74
 - unendliche Faltung, 142
- Flächen
 - benachbarte, 107
- Flächenkomplex, 107
 - geschlossener, 107
 - orientierter, 131
- gebunden, 34
- gerader Abstand, 129
- geschlossen, 107
- Gitterkomplex
 - hexagonaler, 138
- hexagonaler Gitterkomplex, 138
- homogener Komplex, 12
 - Einschränkung, 16
 - lokal simplizial, 16
 - lokal zellulär, 152
- Identifikation, 20

- konstant auf dem Schnitt, 21
 - Nachbar-, 29
- induzierte Fächerpartition, 86
- induzierter Durchlaufsin, 55, 155
- Irreduzibilität, 18
- Kompatibilität von Fächern, 56
- Komplementarität, 60
- Komplex
 - Isomorphismus, 15
 - Morphismus, 15
 - Überlagerungs-, 18
 - Basis-, 41
 - Falt-, 62
 - Flächen-, 107
 - homogener, 12
 - lokal simplizialer, 16
 - lokal zellulärer, 152
- Konfluenz, 67
- konstant auf dem Schnitt, 21
- Kreisdarstellung, 109
 - überschneidungsfreie, 109
 - Abstand, 129
- lokal endlich zulässig, 142
- lokal simplizial, 16
- lokal zellulär, 152
- meet–Halbverband, 89
- Nachbar, 107
 - bezüglich eines Fächers, 59
- Nachbaridentifikation, 29
- Oberfläche
 - induzierter Durchlaufsin, 55, 155
 - komplementäre, 60
 - simpliziale, 52
 - zelluläre, 155
- Orientierbarkeit, 131
- Orientierung, 131
 - sklasse, 48
 - entgegengesetzte, 60
 - simpliziale, 47
 - zelluläre, 153
- Orientierungsklasse, 48
- primitive Faltung, 64, 74
- primitive Spaltung, 34
- Redukt, 53
- Simplex, 16
- simplizial
 - lokal simplizial, 16
- simpliziale Oberfläche, 52
- simpliziale Orientierung, 47
- Spaltung
 - primitive, 34
- Spaltungspartition, 34
- stabil, 96
- Standardrand, 59
- starke Abwärtskompatibilität, 18
- starr, 9
- Starrheit, 9
- Teilüberlagerung, 19
- Teilfaltung, 81
- unendliche Faltung, 142
- ungerader Abstand, 129
- Verband, 35
 - meet–Halbverband, 89
- voller Fächer, 48
- Weg, 130
- Wegzusammenhang, 130
- zelluläre Oberfläche, 155
- zelluläre Orientierung, 153
- zulässig
 - lokal endlich, 142
- zulässige Entfaltung, 97
- zulässiger Faltplan, 92
- zyklische Reihung, 46
 - überschneidungsfrei, 80
- Konfluenz, 67
- Redukt, 53

Danksagung

Ich danke Professor Plesken für seine ausführliche Betreuung, sowie dafür, mich stets wieder auf den rechten Pfad zu ziehen, falls ich mich in Details verloren habe.

Ich danke Charlotte Lorenz und Hannah Arndt für das Korrekturlesen meiner Arbeit. Insbesondere dank Hannahs Kritik ist diese Arbeit deutlich besser verständlich geworden.

Erklärung

Hiermit erkläre ich, Markus Baumeister, an Eides statt, dass ich die vorliegende Masterarbeit selbständig verfasst und keine anderen als die im Literaturverzeichnis angegebenen Quellen und Hilfsmittel benutzt habe.

Aachen, den 16. Juni 2019
