

Kapitel 6

Stabile Teilkomplexe und Kommutatoren

6.1 Absolutes Kommutatorschieben

Eine denkbare Strategie, mit der man ein Gegenbeispiel zu (AC) als solches erweisen könnte, hätte wegen Satz 3.2.16 darin bestehen können, verschiedene Präsentationsklassen durch Projektion von Relatoren ihrer Vertreter in geeignete Kommutatorstufen zu unterscheiden. Die Hoffnung würde darin bestehen, entlang höherer Kommutatoren der Relatorengruppe Invarianten für Q^{**} -Präsentationsklassen zu schaffen. Cynthia Hog-Angeloni und Wolfgang Metzler haben 1990 in ihrer Arbeit „Stabilization by free products giving rise to Andrews-Curtis equivalences“, [11] folgenden Satz bewiesen:

Satz 6.1.1. *Sind die kompakten und zusammenhängenden CW-2-Komplexe K und L einfach homotopieäquivalent, so können sie durch geeignete Deformationen der Dimension 3 in Komplexe K' und L' deformiert werden, so dass die zugehörigen abgelesenen Präsentationen $\mathcal{P} = \mathcal{P}_{K'}$ und $\mathcal{Q} = \mathcal{P}_{L'}$ die folgenden Eigenschaften haben:*

$\mathcal{P} = \langle a_1, \dots, a_g | R_1, \dots, R_h \rangle$ und $\mathcal{Q} = \langle a_1, \dots, a_g | S_1, \dots, S_h \rangle$ sind endliche Präsentationen mit gleicher Relatorengruppe $N \subseteq F(a_i)$, und es gilt $R_j S_j^{-1} \in N^{(1)}$ für alle j . $\square$

In ihrer Arbeit „Andrews-Curtis-Operationen und Höhere Kommutatoren der Relatorengruppe“ [10] haben Cynthia Hog-Angeloni und Wolfgang Metzler dann - ebenfalls 1990 - folgenden Satz bewiesen, der die Hoffnung auf eine irgendwo in höheren Kommutatoren lebende Invariante zunichte macht:

Satz 6.1.2. *Seien $\mathcal{P} = \langle a_1, \dots, a_g | R_1, \dots, R_h \rangle$ und $\mathcal{Q} = \langle a_1, \dots, a_g | S_1, \dots, S_h \rangle$ endliche Präsentationen mit gleicher Relatorengruppe $N \subseteq F(a_i)$ und mit jeweils $R_j S_j^{-1} \in N^{(1)}$. Dann läßt sich zu jedem $n \in \mathbb{N}$ durch eine Q -Operation an $\mathcal{P}$ erreichen, dass die transformierten Relatoren $R_j S_j^{-1} \in N^{(n)}$ erfüllen.*

Den Beweis erledigen wir gemeinsam mit dem von Satz 6.2.1, der eine relative Version und damit eine Verschärfung von Satz 6.1.2 darstellt. Die Q -Transformation des Satzes 6.1.2 kann nämlich unter gewissen Zusatzvoraussetzungen besonders „schön“ gewählt werden:

6.2 Relatives Kommutatorschieben

Satz 6.2.1. Seien $\mathcal{P} = \langle a_1, \dots, a_g | R_1, \dots, R_l \rangle$ und $\mathcal{Q} = \langle a_1, \dots, a_g | S_1, \dots, S_l \rangle$ endliche Präsentationen mit gleicher Relatorengruppe $N \subseteq F(a_i)$ und mit jeweils $R_j S_j^{-1} \in N^{(1)}$. Sei $n \in \mathbb{N}$, sei $2 \leq h < l$, und es gelte $R_i S_i^{-1} \in N^{(n)}$ für alle $i \in \{h+1, \dots, l\}$. Dann gibt es eine Q -Transformation an $\mathcal{P}$ relativ $\{R_{h+1}, \dots, R_l\}$, so dass die damit transformierten Relatoren $R_j S_j^{-1} \in N^{(n)}$ für alle $j \in \{1, \dots, h\}$ erfüllen.

Korollar 6.2.2. Es seien

$$\begin{aligned} \mathcal{P} &= \langle a_1, \dots, a_g | R_1, \dots, R_h, R_{h+1}, \dots, R_l \rangle \\ \mathcal{Q} &= \langle a_1, \dots, a_g | S_1, \dots, S_h, R_{h+1}, \dots, R_l \rangle \end{aligned}$$

endliche Präsentationen mit gleicher Relatorengruppe $N \subseteq F(a_i)$ und $R_j S_j^{-1} \in N^{(1)}$ für alle j . Dann läßt sich zu jedem $n \in \mathbb{N}$ durch eine Q -Operation an $\mathcal{P}$ relativ $\{R_{h+1}, \dots, R_l\}$ erreichen, dass die transformierten Relatoren $R_j S_j^{-1} \in N^{(n)}$ für alle $j \in \{1, \dots, h\}$ erfüllen.

Wenn gewisse Relatoren in $\mathcal{P}$ also bereits zu ihren „Partnern“ in $\mathcal{Q}$ kongruent modulo n -ter Kommutatoren sind, also so, wie man sie gerne hätte, und man mindestens zwei Relatoren verändern darf, dann kann man eine Q -Transformation wählen, die die bereits kongruenten Relatoren festläßt. „Festlassen eines Relators“ heißt dabei nicht nur, dass dieser nach der Transformation dieselbe Gestalt wie vorher hat, sondern sogar, dass er während der gesamten Transformation nicht verändert wurde.

Die Voraussetzung $h \geq 2$ ist für den gleich folgenden Beweis wesentlich. Inwieweit und gegebenenfalls mit welchen Mitteln sich dieser „Schönheitsfehler“ beheben läßt, werden wir im Kapitel 6.3 diskutieren.

Natürlich wäre es denkbar und möglich, Satz 6.2.1 etwas suggestiver zu formulieren. Zu dem Vortrag des Satzes in der obigen Fassung auf dem Treffen der AG in Heidenrod stellte Sergeij Matveev die verständliche Frage: „Warum so viele Buchstaben?“ Die Antwort darauf ist: Es erscheint sinnvoll, den Satz 6.2.1 so zu formulieren, dass sein Beweis möglichst wörtlich aus dem für Satz 6.1.2 erreicht werden kann, und genau das wollen wir nun tun:

Beweis. Zunächst haben wir uns um die Situation $h = 1$ im ersten Satz zu kümmern: Ist $h = 1$, so erzeugen die Worte R_1 und S_1 in $\langle a_1, \dots, a_g \rangle$ denselben normalen Abschluß, und dann sind sie nach dem Korollar zum Freiheitssatz von Magnus (Korollar 3.4.2) konjugiert oder invers konjugiert zueinander. Insbesondere sind sie dann aber Q -äquivalent.

Wir können für Satz 6.1.2 also $h \geq 2$ annehmen, im Fall von Satz 6.2.1 ist dies ohnehin Voraussetzung. (Satz 6.1.2 ist in diesem Fall der Spezialfall von Satz 6.2.1 mit der Besetzung der $R_{h+1}, \dots, R_l$ mit dem leeren Wort.) Der gemeinsame Beweis besteht aus einer Induktion nach der Kommutatorstufe n .

Dabei ist die Verankerung, $n = 1$, bereits unmittelbar in der Voraussetzung erledigt. Zum Beweis beider Sätze nehmen wir nun als Induktionshypothese an, die Behauptungen seien für die Kommutatorstufe $n - 1$ mit $n \geq 2$ gezeigt. Der Induktionsschluß besteht nun aus sechs Schritten.

Ziel von *Schritt A* ist das beweistechnische Zurückspielen des relativen Falls auf den absoluten. Wir bezeichnen die bereits transformierten Relatoren in $\mathcal{P}$ so

wie die ursprünglichen mit R_i – das wird nicht zu Missverständnissen führen. Wir haben durch die Induktionshypothese folgende Situation:

$$R_i = S_i \cdot W_i \quad \text{mit } W_i \in N^{(n-1)} \quad \text{für } i = 1, \dots, h \quad (1)$$

$$R_i = S_i \cdot W_i \quad \text{mit } W_i \in N^{(n)} \quad \text{für } i = h+1, \dots, l \quad (2)$$

Dabei ist (1) für sich alleine auch schon als Induktionshypothese in den Bezeichnungen von Satz 6.1.2 zu verstehen, und (2) ist gerade die Zusatzvoraussetzung von Satz 6.2.1. Sie ist aber keineswegs eine Wiederholung oder Selbstverständlichkeit, sondern integraler Bestandteil der Induktionshypothese; denn diese besagt ja gerade, dass die Q -Transformation, die die Äquivalenz modulo $N^{(n-1)}$ hergestellt hat, schon so gewählt werden konnte, dass die Relatoren $R_{h+1}, \dots, R_l$ von ihr nicht verändert wurden.

Die W_i sind nach Lemma 3.2.17 Produkte von Kommutatoren $(n-1)$ -ter Stufe bzw. n -ter Stufe in Konjugaten der $S_j^{\pm 1}$. Wir schreiben $W_i = W_i(S_1, \dots, S_l)$. Eines der wesentlichen Mittel des Induktionsschlusses ist nun folgendes Verfahren: Kommt in R_i , etwa in der eben gewählten Darstellung, das Wort $S_j^{\pm 1}$, $j \neq i$ vor, so erlaubt uns eine Q -Transformation an R_i , dass wir $S_j^{\pm 1}$ durch das Gegenwort $W^{\mp 1}(S_1, \dots, S_l)$ nach (1) und (2) ersetzen. Sei etwa $R_i = S_i X S_j Z$, dann können wir die gewünschte Transformation wie folgt direkt angeben:

$$S_i X S_j Z \rightarrow S_i X S_j Z \cdot (Z^{-1} W_j (S_j W_j)^{-1} W_j^{-1} Z) = S_i X W_j^{-1} Z. \quad (3)$$

Für ein Vorkommen von S_j mit negativem Exponenten verfährt man ähnlich. Natürlich muß für jedes in R_i vorkommende S_j , $i \neq j$, einzeln eine Operation nach (3) durchgeführt werden.

Nach dieser Methode ersetzen wir in den R_i mit $1 \leq i \leq h$ die dort vorkommenden S_j mit $h+1 \leq j \leq l$. So kommen wir nach (3) zu folgender Situation – die R_i bezeichnen wieder die bereits transformierten Relatoren:

$$R_i = S_i \cdot W_i(S_1, \dots, S_h, W_{h+1}^{-1}, \dots, W_l^{-1}) \quad \text{für } i \in \{1, \dots, h\} \quad (4)$$

Bislang haben wir echte Gleichheiten angegeben. Ab dieser Stelle betrachten wir die R_i nun zum ersten Mal nur noch bis auf Kongruenz modulo $N^{(n)}$ und bezeichnen dies durch den Doppelpunkt statt des Gleichheitszeichens. Nach (2) gilt für $j \in \{h+1, \dots, l\}$ schon $W_j \in N^{(n)}$, also sind die W_j kongruent eins modulo n -ter Kommutatoren, und mit (4) haben wir bereits:

$$R_i : S_i \cdot W_i(S_1, \dots, S_h, 1, \dots, 1) \quad \text{für } i \in \{1, \dots, h\} \quad (5)$$

Jede innerste Klammer der Gestalt $[a S a^{-1}, 1]$ ist trivial. Wird ein S_j durch 1 ersetzt, so werden also alle innersten Kommutatorklammern, in denen S_j vorkam, trivial. Wegen Lemma 3.2.17 schreiben wir:

$$R_i = S_i \cdot W_i(S_1, \dots, S_h) \quad \text{für } i \in \{1, \dots, h\} \quad (6)$$

Die W_i sind hier strenggenommen nicht dieselben wie die in (1). Korrekt formuliert wäre: Wir gewinnen durch Q -Operationen an den R_i neue $W_i \in N^{(n)}$,

die nur noch Produkte von Konjugaten der S_j mit $1 \leq j \leq h$ in der Gestalt von Lemma 3.2.17 sind. Damit ist der Beweis für den relativen Fall vollständig auf den für den absoluten zurückgespielt. Wir haben bislang nur an denjenigen R_i mit $i \leq h$ Transformationen vorgenommen. Der weitere Beweisgang folgt fast wörtlich dem in der Originalarbeit [10]. Auch hier werden nur an denjenigen R_i mit $i \leq h$ Transformationen vorgenommen. Wir benötigen nun noch fünf Schritte.

Unser Ziel in *Schritt B* ist nun, Transformationen relativ $\{R_{h+1}, \dots, R_l\}$ zu finden, die die Relatoren $R_1, \dots, R_h$ in die nachfolgende Gestalt überführen. Dabei nimmt i einen festen Wert zwischen 1 und h an, und j durchläuft alle übrigen Werte zwischen 1 und h :

$$\begin{aligned} R_i : S_i & \cdot W_i(W_1^{-1}(1, \dots, 1, S_i, 1, \dots, 1), \dots, \\ & W_{i-1}(1, \dots, 1, S_i, 1, \dots, 1), S_i, \\ & W_{i+1}(1, \dots, 1, S_i, 1, \dots, 1), \dots, \\ & W_h(1, \dots, 1, S_i, 1, \dots, 1)) \\ R_j : S_j & \cdot W_j(S_1, \dots, S_{i-1}, 1, S_{i+1}, \dots, S_h) \quad i \neq j \end{aligned} \quad (7)$$

Unser Ziel ist also, (7) aus (6) zu gewinnen. Wieder ersetzen wir nach dem Muster von (3) die „fremden“ S_j und erhalten:

$$\begin{aligned} R_i : S_i & \cdot W_i(W_1^{-1}(S_1, \dots, S_h), \dots, W_{i-1}^{-1}(S_1, \dots, S_h), S_i, \\ & W_{i+1}^{-1}(S_1, \dots, S_h), \dots, W_h^{-1}(S_1, \dots, S_h)), \\ R_j : S_j & \cdot W_j(S_1, \dots, S_h), \quad j \neq i. \end{aligned} \quad (8)$$

Nun ersetzen wir nach (3) in den R_j alle Vorkommnisse von $S_i^{\pm 1}$ erneut durch deren Gegenfaktor in seiner Gestalt von (8). Wir bekommen damit:

$$\begin{aligned} R_i : S_i & \cdot W_i(W_1^{-1}(S_1, \dots, S_h), \dots, W_{i-1}^{-1}(S_1, \dots, S_h), S_i, \\ & W_{i+1}^{-1}(S_1, \dots, S_h), \dots, W_h^{-1}(S_1, \dots, S_h)), \\ R_j : S_j & \cdot W_j(S_1, \dots, S_{i-1}, W_i^{-1}(W_1^{-1}(S_1, \dots, S_h), \dots, \\ & W_{i-1}^{-1}(S_1, \dots, S_h), \underline{S_i}, W_{i+1}^{-1}(S_1, \dots, S_h), \dots, \\ & W_h^{-1}(S_1, \dots, S_h)), S_{i+1}, \dots, S_h) \quad j \neq i \end{aligned} \quad (9)$$

Danach ersetzen wir das unterstrichene S_i noch einmal und bekommen auf diese Weise:

$$\begin{aligned} R_i : S_i & \cdot W_i(W_1^{-1}(S_1, \dots, S_h), \dots, W_{i-1}^{-1}(S_1, \dots, S_h), S_i, \\ & W_{i+1}^{-1}(S_1, \dots, S_h), \dots, W_h^{-1}(S_1, \dots, S_h)), \\ R_j : S_j & \cdot W_j(S_1, \dots, S_{i-1}, W_i^{-1}(W_1^{-1}(S_1, \dots, S_h), \dots, \\ & W_{i-1}^{-1}(S_1, \dots, S_h), W_i^{-1}(W_1^{-1}(S_1, \dots, S_h), \dots, \\ & W_{i-1}^{-1}(S_1, \dots, S_h), S_i, W_{i+1}^{-1}(S_1, \dots, S_h), \dots, \\ & W_h^{-1}(S_1, \dots, S_h)), W_{i+1}^{-1}(S_1, \dots, S_h), \dots, \\ & W_h^{-1}(S_1, \dots, S_h)), S_{i+1}, \dots, S_h) \quad j \neq i \end{aligned} \quad (10)$$

Diese Darstellung ist nicht wirklich übersichtlich. Auf einen Teil der Argumente in den W_k können wir aber verzichten. Dann haben wir:

$$\begin{aligned}
R_i &: S_i \cdot W_i(W_1^{-1}(S_1, \dots, S_h), \dots, W_{i-1}^{-1}(S_1, \dots, S_h), S_i, \\
&\quad W_{i+1}^{-1}(S_1, \dots, S_h), \dots, W_h^{-1}(S_1, \dots, S_h)), \\
R_j &: S_j \cdot W_j(S_1, \dots, S_{i-1}, W_i^{-1}(W_1^{-1}, \dots, W_{i-1}^{-1}, \\
&\quad W_i^{-1}(W_1^{-1}, \dots, W_{i-1}^{-1}, S_i, W_{i+1}^{-1}, \dots, W_h^{-1}), \\
&\quad W_{i+1}^{-1}, \dots, W_h^{-1}), S_{i+1}, \dots, S_h) \quad j \neq i
\end{aligned} \tag{11}$$

Nun ist aber

$$\begin{aligned}
&W_i^{-1}(W_1^{-1}, \dots, W_{i-1}^{-1}, \\
&W_i^{-1}(W_1^{-1}, \dots, W_{i-1}^{-1}, S_i, W_{i+1}^{-1}, \dots, W_h^{-1}), \\
&W_{i+1}^{-1}, \dots, W_h^{-1})
\end{aligned} \tag{12}$$

ein $(n-1)$ -ter Kommutator in $(n-1)$ -ten Kommutatoren, insbesondere Kommutator in $(n-1)$ -ten Kommutatoren und damit n -ter Kommutator und kongruent eins modulo $N^{(n)}$. Wir wollen die R_k aber nur bis auf Kongruenz betrachten und dürfen deshalb für (12) in (11) eins einsetzen. Damit haben aber die R_j bereits in die Gestalt von (7); denn wir erhalten:

$$\begin{aligned}
R_i &: S_i \cdot W_i(W_1^{-1}(S_1, \dots, S_h), \dots, W_{i-1}^{-1}(S_1, \dots, S_h), S_i, \\
&\quad W_{i+1}^{-1}(S_1, \dots, S_h), \dots, W_h^{-1}(S_1, \dots, S_h)), \\
R_j &: S_j \cdot W_j(S_1, \dots, S_{i-1}, 1, S_{i+1}, \dots, S_h) \quad i \neq j
\end{aligned} \tag{13}$$

Die soeben gewonnene Gestalt der R_j verwenden wir nun zur Einsetzung in R_i . Zunächst wird nach der Methode (3) in R_i jedes $S_j, j \neq i$ durch den Ausdruck

$$W_j^{-1}(S_1, \dots, S_{i-1}, 1, S_{i+1}, \dots, S_h)$$

ersetzt. Nochmaliges Einsetzen bringt $W_j^{-1}(W_1^{-1}, \dots, W_{i-1}^{-1}, 1, W_{i+1}^{-1}, \dots, W_h^{-1})$, wo vorher S_j stand. Dieser Ausdruck ist nun wieder Kommutator in $(n-1)$ -ten Kommutatoren, also n -ter Kommutator und kongruent eins modulo $N^{(n)}$. Also erhalten wir aus (13) durch die vorgenannten Einsetzungen in R_i und Repräsentantenwechsel modulo $N^{(n)}$ die gewünschte Darstellung (7).

Wir haben ab (8) nur mit Repräsentantenwechseln und Einsetzungen gearbeitet. Den Repräsentantenwechseln entspricht gar keine Q -Transformation, den Einsetzungen nach (3) eine Q -Transformation an dem Relator, in dem eingesetzt wurde. Demnach haben wir bei allen bisherigen Transformationen die Relatoren $R_{h+1}, \dots, R_l$ unverändert gelassen, die bisherige Transformation war also tatsächlich *relativ* $\{R_{h+1}, \dots, R_l\}$, wie verlangt. Wichtig ist auch der Hinweis, dass schon in den letzten Schritten die Voraussetzung $h \geq 2$ wesentlich war. Es sind für das hier benutzte Beweisprinzip mindestens zwei veränderbare Relatoren nötig: Der zweite wird mit dem ersten transformiert, und der erste wird wiederum mit dem transformierten zweiten transformiert.

In *Schritt C* wiederholen wir Schritt B sukzessive $h-1$ mal. Bei jedem Durchgang nimmt die innerhalb des Durchgangs feste Variable i einen neuen Wert

an. Hat i nach allen Durchgängen alle Werte von 1 bis h durchlaufen, so hat jeder Relator eine der von R_i in (7) entsprechende Darstellung erhalten, und wir haben:

$$R_j : S_j \cdot W_j(1, \dots, 1, S_j, 1, \dots, 1) \quad (14)$$

Wie schon am Ende von Schritt A bei (6) können wir nun zu einfacheren Worten für die W_j übergehen, diesmal ändern wir aber die Bezeichnung und erhalten:

$$R_j : S_j \cdot V_j(S_j) \quad (15)$$

Im jetzt folgenden *Schritt D* wollen wir für zwei der Relatoren die Darstellung

$$\begin{aligned} R_j &: S_i \cdot V(S_i) \\ R_m &: S_m \end{aligned} \quad (16)$$

erhalten.

Dafür multiplizieren wir zunächst R_i mit dem Konjugat $V_i^{-1}(S_i) \cdot R_m \cdot V_i(S_i)$ und verwenden dabei die Relatoren in ihrer Form nach (15):

$$\begin{aligned} R_i &: S_i \cdot S_m \cdot V_m(S_m) \cdot V_i(S_i), \\ R_m &: S_m \cdot V_m(S_m) \end{aligned} \quad (17)$$

Bis auf n te Kommutatoren können wir (17) transformieren zu:

$$\begin{aligned} R_i &: S_i \cdot S_m \cdot V_m(S_m) \cdot V_i(S_i), \\ R_m &: S_m \cdot V_m(S_m) \cdot V_m^{-1}(S_i \cdot S_m) \end{aligned} \quad (18)$$

Dabei soll $V_m^{-1}(S_i \cdot S_m)$ aus $V_m^{-1}(S_m)$ dadurch hervorgehen, dass jedes S_m in einer innersten Kommutatorklammer durch $S_i \cdot S_m$ ersetzt wird. Den Übergang von (17) nach (18) sieht man, indem man umgekehrt (17) aus (18) erreicht: $S_i \cdot S_m$ kann man durch den Gegenfaktor $V_i^{-1}(S_i) \cdot V_m^{-1}(S_m)$ aus R_i in (18) ersetzen. Der Term $V_m^{-1}(V_i^{-1}(S_i) \cdot V_m^{-1}(S_m))$ liegt aber wieder in $N^{(n)}$ und fällt deshalb in Kongruenz modulo $N^{(n)}$ weg. Für den relativen Fall macht man sich noch einmal mit Lemma 4.3.3 klar, dass es egal ist, ob man mit einer Q -Transformation relativ $\{R_{h+1}, \dots, R_l\}$ von (17) nach (18) oder von (18) nach (17) kommt.

Wir multiplizieren nun in einmal ein Konjugat von R_m aus (18) an R_i . An die Stelle des S_m direkt nach S_i tritt das Gegenwort $(V_m(S_m) \cdot V_m^{-1}(S_i \cdot S_m))^{-1}$ und ergibt

$$\begin{aligned} R_i &: S_i \cdot V_m(S_i \cdot S_m) \cdot V_i(S_i) \\ R_m &: S_m \cdot V_m(S_m) \cdot V_m^{-1}(S_i \cdot S_m) \end{aligned} \quad (19)$$

Wir betrachten nun

$$\begin{aligned} W_i(S_i, S_m) &:= V_m(S_i \cdot S_m) \cdot V_i(S_i) \\ W_m(S_i, S_m) &:= V_m(S_m) \cdot V_m^{-1}(S_i \cdot S_m) \end{aligned}$$

Dann ist $R_i : S_i \cdot W_i(S_i, S_m)$ und $R_m : S_m \cdot W_m(S_i, S_m)$, und wir können wie in Schritt B verfahren, um die zu (7) analoge Situation zu erhalten. Diese lautet dann:

$$\begin{aligned} R_i &: S_i \cdot V_m(S_i \cdot V_m^{-1}(S_i)) \cdot V_i(S_i) \\ R_m &: S_m \end{aligned} \quad (20)$$

Mit $V(S_i) = V_m(S_i \cdot V_m^{-1}(S_i)) \cdot V_i(S_i)$ ist (20) von der Gestalt (16).

Schritt E ist wieder verhältnismäßig kurz. Wir halten das i aus Schritt D weiter fest, wiederholen das Verfahren von Schritt E aber noch $h - 2$ mal und lassen dabei die Variable m alle noch übrigen Werte in $\{1, \dots, h\}$ annehmen. Wir erhalten

$$\begin{aligned} R_i &: S_i \cdot U(S_i) \\ R_j &: S_j \quad \text{für alle } j \neq i \end{aligned} \quad (21)$$

Es folgt zum Abschluß des Beweises nun *Schritt F*: Wir müssen von (21) durch Q -Transformationen zu

$$R_j : S_j \quad 1 \leq j \leq h \quad (22)$$

gelangen.

Wir zeigen umgekehrt, wie man bei beliebig vorgegebenem $U(S_i) \in N^{(n-1)}$ von (22) nach (21) gelangt. Dazu benötigen wir außer R_i noch ein weiteres R_m . Sei also $U(S_i)$ vorgelegt. Dann bilden wir $U^*(S_m^{-1}, S_i)$ dadurch, dass jede innerste Kommutatorklammer $[vS_i^\delta v^{-1}, wS_i^\epsilon w^{-1}]$, (v, w Worte in den a_k ; $\delta, \epsilon \in \{1, -1\}$), durch $[vS_m^{-\delta} v^{-1}, wS_i^\epsilon w^{-1}]$ ersetzt wird.

Dann gilt offenbar $U^*(S_i, S_i) = U(S_i)$, außerdem $U^*(1, S_i) = 1 = U^*(S_m^{-1}, 1)$. Lemma 3.2.18 ergibt direkt, dass $U^*(S_m^{-1}, S_i)$ im normalen Abschluß von S_m liegt. Für $R_m : S_m$ ist also Multiplikation mit $U^*(S_m^{-1}, S_i)$ eine Q -Operation mit R_m , und wir können R_i und R_m aus (22) in

$$\begin{aligned} R_i &: S_i \cdot U^*(S_m^{-1}, S_i) \\ R_m &: S_m \end{aligned} \quad (23)$$

überführen. Danach multiplizieren wir das neue R_i an R_m und erhalten

$$\begin{aligned} R_i &: S_i \cdot U^*(S_m^{-1}, S_i) \\ R_m &: S_m \cdot S_i \cdot U^*(S_m^{-1}, S_i) \end{aligned} \quad (24)$$

Dann ersetzen wir S_m in R_i durch den Gegenfaktor aus R_m :

$$\begin{aligned} R_i &: S_i \cdot U^*(S_i \cdot U^*(S_m^{-1}, S_i), S_i) \\ R_m &: S_m \cdot S_i \cdot U^*(S_m^{-1}, S_i) \end{aligned} \quad (25)$$

Den Faktor S_i in R_m ersetzen wir durch das Gegenwort aus R_i und bekommen

$$\begin{aligned} R_i &: S_i \cdot U^*(S_i \cdot U^*(S_m^{-1}, S_i), S_i) \\ R_m &: S_m \cdot (U^*)^{-1}(S_i \cdot U^*(S_m^{-1}, S_i), S_i) \cdot U^*(S_m^{-1}, S_i) \end{aligned} \quad (26)$$

An dieser Stelle betrachten wir

$$\begin{aligned} W_i(S_i, S_m) &:= U^*(S_i \cdot U^*(S_m^{-1}, S_i), S_i) \\ W_m(S_i, S_m) &:= (U^*)^{-1}(S_i \cdot U^*(S_m^{-1}, S_i), S_i) \cdot U^*(S_m^{-1}, S_i). \end{aligned}$$

Wie schon in (7) können wir S_m in R_i durch 1 ersetzen. Danach können wir mit derselben Methode S_i in R_m durch $W_i^{-1}(1, S_m) = (U^*)^{-1}(1 \cdot U^*(S_m^{-1}, 1), 1) = 1$ ersetzen.

Damit haben wir

$$\begin{aligned} R_i &: S_i \cdot U^*(S_i \cdot 1, S_i) = S_i \cdot U(S_i) \\ R_m &: S_m \cdot (U^*)^{-1}(1 \cdot U^*(S_m^{-1}, 1), 1) \cdot U^*(S_m^{-1}, 1) = S_m \end{aligned}$$

Damit haben wir (21) und den Abschluß von Schritt F erreicht, die Beweise von Satz 6.1.2 und Satz 6.2.1 sind vollständig. $\square$

Das Korollar 6.2.2 folgt direkt aus dem Satz: Gilt $R_j = S_j$, so ist nämlich trivialerweise $R_j S_j^{-1} \in N^{(n)}$ für alle $n \in \mathbb{N}$. $\square$

6.3 Der Fall eines einzigen beweglichen Relators

Es seien nun zwei Präsentationen

$$\begin{aligned} \mathcal{P} &= \langle a_1, \dots, a_k | R, R_1, \dots, R_l \rangle \\ \mathcal{Q} &= \langle a_1, \dots, a_k | S, S_1, \dots, S_l \rangle \end{aligned}$$

gegeben, und es gelte $RS^{-1} \in N^{(1)}$, $R_j S_j^{-1} \in N^n$. Gibt es dann eine Q -Transformation an $\mathcal{P}$ relativ $\{R_1, \dots, R_l\}$, so dass für den einzigen dann transformierten Relator R danach gilt $RS^{-1} \in N^{(n)}$? Diese Frage können wir auch mit Satz 6.2.1 nicht positiv beantworten. Der Beweis von Satz 6.2.1 hat, wie schon entlang des Beweisgangs erläutert, gerade deshalb funktioniert, weil mindestens an zwei Relatoren Veränderungen vorgenommen werden durften. Diese Möglichkeit entfällt nun. Und siehe da, das gewünschte Ergebnis ist in diesem Fall auch nicht zu erzielen:

6.3.1 Ein Q -Gegenbeispiel für den Fall $h = 1$

Satz 6.3.1. (Metzler 2000) Seien für $j \notin \{-1, 0, 1\}$ und $m \in \mathbb{Z}, m \not\equiv 0 \pmod{m^2}$ und

$$\begin{aligned} R &= t^j a t^{-j} [t a t^{-1}, a^m] \\ S &= t^j a t^{-j} \\ T &= [t a t^{-1}, a^m]^m \end{aligned}$$

die Präsentationen

$$\begin{aligned} \mathcal{P} &= \langle a, t | R, T \rangle \\ \mathcal{Q} &= \langle a, t | S, T \rangle \end{aligned}$$

gegeben¹. Dann gibt es keine Q -Operation an $\mathcal{P}$ relativ $\{T\}$, so dass der transformierte Relator R zu S kongruent modulo zweiter Kommutatoren der Relatorengruppe wird.

Man mache sich zunächst klar, dass beispielsweise $j = m = 2$ die Voraussetzungen erfüllen, und dass der Satz dann tatsächlich ein Gegenbeispiel zu „Satz 6.2.1 mit $h = 1$ “ ergibt: Die beiden Präsentationen erfüllen in der Tat die geforderten Relatorkongruenzen. Damit haben wir, umgangssprachlich formuliert, folgendes Ergebnis: Um sicher zu sein, dass relatives Hochschieben in zweite Kommutatoren möglich ist, braucht man mindestens zwei bewegliche Relatoren. Doch hier zunächst der

Beweis von Satz 6.3.1. Wir werden noch etwas Schärferes zeigen, und zwar, dass sich durch Q -Transformationen an R noch nicht einmal eine Kongruenz von R und S modulo des ersten Zentralreihenterms (also Gleichheit im Reidemeisterquotienten der Relatorengruppe) erreichen lässt. Weil zweite Kommutatoren insbesondere zum ersten Zentralreihenterm gehören, folgt damit auch die Behauptung des Satzes.

Mit der Reidemeister-Schreier-Methode² sieht man, dass die Relatorengruppe $N \trianglelefteq F(a, t)$ von den Konjugaten

$$a, \quad b := tat^{-1}, \quad c := t^2at^{-2}, \dots$$

frei erzeugt wird. Wir betrachten fortan die Elemente der Relatorengruppe nur noch bis auf Gleichheit im Reidemeisterquotienten. Dazu erinnern wir an den Reidemeister-Normalformensatz 3.2.13, wenden diesen auf die eben erklärten $a, b, c, \dots$ an und betrachten Normalformen der folgenden besonderen Gestalt (die Reihenfolge der Elementarkommutatoren in der Normalform ist im Reidemeisterquotienten definitionsgemäß beliebig):

$$t^k a^{\pm 1} t^{-k} \prod [t^k at^{-k}, t^i at^{-i}]^{\zeta_{k,i}} \prod_{r,s \neq k} [t^r at^{-r}, t^s at^{-s}]^{\zeta_{r,s}}. \quad (1)$$

Offenbar ist R gerade von dieser Gestalt, und zwar gemäß Voraussetzung des Satzes so, dass rechts neben dem zweiten Produktzeichen mindestens ein Exponent nicht kongruent null modulo m^2 ist³. Diese Eigenschaft bleibt, wie wir jetzt zeigen, bei jeder der für uns zulässigen Q -Transformationen erhalten: Modulo $[N, [N, N]]$ kann jede Q -Transformation an R dargestellt werden als Kompositum der folgenden Operationen:

1. Inversion
2. Konjugation mit t
3. Konjugation mit a
4. Multiplikation mit T

¹ $\mathcal{P}$ und $\mathcal{Q}$ präsentieren übrigens $\mathbb{Z}(t)$. Weil die Whiteheadgruppe von $\mathbb{Z}$ trivial ist (siehe Cohen [5]), sind die zugehörigen Standardkomplexe einfach homotopieäquivalent und in Sachen Andrews-Curtis interessant. Sie bilden das aktuelle „Lieblings-Andrews-Curtis-Gegenbeispiel“ ihrer Erfinder Hog-Angeloni und Metzler.

²ausführlich beschrieben bei Magnus, Karass und Solitar [18] und Zieschang, Vogt und Coldewey [30]

³Das ist für R selbst trivial: Betrachte mit Bemerkung 3.2.11 den Faktor $[tat^{-1}, a]^m$ - wir hatten vorausgesetzt, dass m nicht kongruent null modulo m^2 ist.

Letzteres hakt im Reidemeisterquotienten auch die Multiplikation mit Konjugaten von T ab: Weil T Kommutatorprodukt ist, ist T Reidemeister-äquivalent zu seinen Konjugaten.

Betrachten wir nun die einzelnen Schritte und ihre Wirkung auf eine Normalform nach (1):

Inversion ändert nur die Vorzeichen der ζ und macht aus dem $a^{\pm 1}$ ein $a^{\mp 1}$. Die bei Inversion sonst übliche Reihenfolgenvertauschung in den Erzeugenden fällt im Reidemeisterquotienten wegen der Vertauschbarkeit mit Kommutatoren weg. Inkongruenz eines ζ modulo m^2 bleibt stets erhalten, also bleibt uns eine Normalform in der gewünschten Gestalt erhalten.

Konjugation mit t ist dasselbe wie die Konjugation jedes einzelnen Faktors mit t . Diese bewirkt in (1) eine Veränderung jedes Index k, i, r, s um 1. Die ζ springen einfach einen Index weiter; die Sortierung der ζ nach Gleichheit eines Index mit dem ersten t -Exponenten bleibt unberührt. Die Gestalt der Normalform bleibt erhalten.

Konjugation mit a : Weil a in Reidemeisteräquivalenz mit allen Kommutatoren vertauschbar ist, macht nur das „Vorbeiziehen“ eines a an dem Term $t^k a^{\pm 1} t^{-k}$ eine Veränderung an der Normalform aus. Dieses entspricht, wie wir im Existenzbeweis von Satz 3.2.13 gesehen haben, einer Multiplikation mit einer Potenz des Elementarkommutators $[t^k a t^{-k}, a]$. In (1) werden also nur Exponenten verändert, die *zwischen* den beiden Produktzeichen stehen. Insbesondere bleiben diejenigen rechts des zweiten Produktzeichens ungeändert.

Multiplikation mit T ist im Reidemeisterquotienten Multiplikation mit dem Faktor $[tat^{-1}, a]^{m^2}$. Dieser ist Potenz eines Normalformenfaktors. Bis auf Kongruenz modulo m^2 bleiben also alle Exponenten in (1) erhalten - insbesondere bleibt Inkongruenz modulo m^2 erhalten.

Zusammengefasst heißt das: Welche Q -Transformation wir auch an R vornehmen, immer bleibt in der Reidemeister-Normalform ein Kommutatorfaktor mit echtem Exponenten stehen. Die Normalform tat^{-1} , in der es keinen solchen Faktor gibt, kann also nicht erreicht werden. Damit ist Satz 6.3.1 bewiesen. $\square$