

Theorie der Nebenläufigkeit

Basierend auf einer Mitschrift

der Vorlesung von Sergey Goncharov im Wintersemester 2016
von Christian Bay, Frederik Haselmeier und Julian Jakob,
überarbeitet von Sergey Goncharov und Lutz Schröder

Inhaltsverzeichnis

1	Prozessalgebra	3
1.1	Inhalt der Vorlesung	3
1.2	Literatur	3
1.3	CCS: Motivation	4
1.4	CCS: Konstrukte	5
1.5	Labelled Transition Systems	8
1.6	Syntax von CCS	9
1.7	Semantik von CCS	10
1.8	Value-passing CCS	10
1.9	Ableitungen	12
2	Verhaltensäquivalenz	13
2.1	Trace-Äquivalenz	14
2.2	Bisimilarität und Bisimulation	15
2.3	Schwache Bisimilarität	22
2.4	Bisimulationsspiele	28
2.5	Schwache Bisimulationsspiele	30
2.6	Bisimulation als Fixpunkt	31
2.7	Der Paige-Tarjan-Algorithmus	34
3	Hennessey-Milner-Logik	41
4	Der modale μ-Kalkül	46
4.1	Der (multi-)modale μ -Kalkül	47
4.2	Syntax	47
4.3	Semantik	48
4.4	Temporale Spezifikation im μ -Kalkül	51

5	Erfülltheitsspiele	53
6	Der π-Kalkül	56
6.1	Syntax	57
6.2	Semantik	59
6.3	Beispiele	60

1 Prozessalgebra

Bekannte Prozessalgebren:

- *Calculus of Communicating Systems (CCS)* (Milner)
- *Communicating Sequential Processes (CSP)* (Hoare)
- *Algebra of Communicating Processes (ACP)* (Bergstra/Klop)

Berührungspunkte mit Automatentheorie.

1.1 Inhalt der Vorlesung

- CCS
- Labelled Transition System (LTS): Semantik für *CCS*.
- Hennessy/Millner Logik (HML): Einfache Spezifikationslogik für LTS
- Bisimilarität: Ununterscheidbarkeit von Prozessen
- μ -Kalkül: ausdrucksstarke Erweiterung der HML, Spezifikation temporaler Eigenschaften ('immer', 'letztendlich' etc.)
- π -Kalkül: Erweiterung von *CCS* um dynamische Konfiguration der Prozesskommunikation – Erzeugen und Versenden von Kanalnamen.

1.2 Literatur

- Robin Milner, *Communication and Concurrency*, Prentice-Hall, 1989.
- Luca Aceto, Kim G. Larsen, *An Introduction to Milner's CCS*.
- Luca Aceto, Anna Ingólfssdóttir, Kim Guldstrand Larsen and Jiri Srba: *Reactive Systems*, Cambridge University Press, 2007
- Davide Sangiorgi and David Walker, *The Pi-Calculus - A Theory of Mobile Processes*, Cambridge University Press, 2001.
- Joachim Parrow, *An introduction to the π -calculus*, in J. Bergstra, A. Ponse and S. Smolka, eds., *Handbook of Process Algebra*, Elsevier 2001, pp. 479-543.

Die Entwicklung des Materials in der Vorlesung folgt größtenteils Aceto et al. (2007); von dort stammen insbesondere das Getränkeautomatenbeispiel und viele der Übungsaufgaben. Das Kapitel über den π -Kalkül folgt Parrow (2001), insbesondere auch hinsichtlich des Beispielmaterials.

1.3 CCS: Motivation

Kommunikationsmodelle involvieren typischerweise ein Medium, das die Kommunikation trägt:

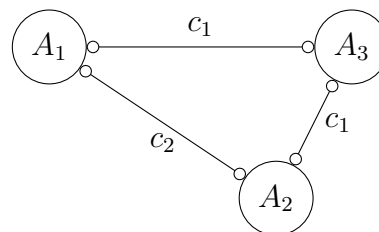
$$\text{Sender} \longrightarrow \text{Medium} \longrightarrow \text{Receiver}$$

Verbreitete Medientypen sind etwa

1. Medium = Ether: Sender kann immer schicken, Receiver kann immer empfangen, Reihenfolge ist nicht festgelegt.
2. Medium = Puffer: Wie in Ether, aber mit definierter Reihenfolge.
3. Medium = Shared Memory: Gemeinsam genutzter Speicher, in dem Reihenfolge und Häufigkeit des Nachrichtenempfangs nicht spezifiziert sind.
4. Weitere Varianten: Begrenzter Speicher, spontaner Verlust von Nachrichten

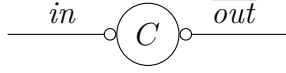
Milners Idee ist nun, dass man nur eine primitive Art von Kommunikation braucht, den sogenannten *Handshake*. Kompliziertere Medien wie Puffer sind dann eigene *Prozesse* oder *Agenten*.

Wir stellen die Kommunikationsstruktur gelegentlich durch *Interfacediagramme* dar; z.B. besagt



dass Agent A_3 über den Kanal c_1 mit den Agenten A_1 und A_3 kommuniziert, und A_1 und A_2 untereinander über Kanal c_2 .

Beispiel 1. Im Diagramm



wird demnach ein einzelner Prozess dargestellt, der über zwei Kanäle $\overline{out}$ und in mit einer unspezifizierten Umgebung kommuniziert. Ein mögliches Verhalten dieses Prozesses wäre, abwechselnd auf beiden Kanälen zu kommunizieren, was man in CCS folgendermaßen notieren würde:

$$\begin{aligned} C &= in(x).C'(x) \\ C'(x) &= \overline{out}(x).C; \end{aligned}$$

äquivalent, als einzelne Gleichung:

$$C = in(x).\overline{out}(x).C$$

1.4 CCS: Konstrukte

Wir führen die syntaktischen Konstrukte von CCS ein, mit zunächst nur informeller Semantik (die formale Semantik folgt später).

1. Deadlock: $\emptyset$ (Keine Aktion ist möglich).
2. Aktionen: Ports bzw. Kanäle, über die Prozesse miteinander kommunizieren können. Formal parametrisieren wir die Syntax über eine Menge $\mathcal{A}$ von *Aktionen*, z.B.

$$\mathcal{A} \ni \text{pushButton}, \text{tick}, \dots$$

Aktionen werden syntaktisch als Präfixe verwendet:

$$a.P \quad (a \in \mathcal{A}, P \text{ Prozess})$$

Informell: “Führe a durch, verhalte dich dann wie P .”

Beispiel 2. Modell eines Streichholzes:

$$\begin{aligned} \text{BadMatch} &= \text{Strike}.\emptyset \\ \text{GoodMatch} &= \text{Strike}.\text{Burn}.\emptyset \end{aligned}$$

3. Rekursion/Prozessdefinition: Wir unterstellen weiterhin eine Menge K von *Prozessnamen*. Eine (*rekursive*) *Definiton* hat dann die Form

$$A = P \quad (A \in K, P \text{ Prozess})$$

Beispiele haben wir oben bereits gesehen, hier noch eins:

$$Clock = tick. Clock$$

Intuitiv impliziert dies übrigens

$$Clock = \underbrace{tick. tick. \dots tick}_{n\text{-mal}}. Clock$$

– wir werden das später formal fassen.

4. (Nichtdeterministische) Auswahl

$$P + Q \qquad (P, Q \text{ Prozesse})$$

Informelle Semantik:

“Wähle einen möglichen Schritt in P oder in Q , und verhalte dich dann weiter wie P bzw Q .”

Beispiel 3.

$$\begin{aligned} B_0 &= in(x).B_1(x) \\ B_1(x) &= in(y).B_2(y, x) + \overline{out}(x).B_0 \\ B_2(x, y) &= \overline{out}(y).B_1(x) \end{aligned}$$

Beispiel 4. Ein eventuell nicht funktionierendes Streichholz lässt sich naiv auf eine der folgenden Arten modellieren:

$$\begin{aligned} Match &= strike.\emptyset + strike.burn.\emptyset \\ Match' &= strike.(\emptyset + burn.\emptyset) \end{aligned}$$

Dies wirft die Frage auf, ob diese beiden Prozesse verhaltensgleich sind. Die Antwort ist **Nein**. Die oben notierte informelle Semantik besagt, dass in $\emptyset + burn.\emptyset$ ein möglicher Schritt aus entweder $\emptyset$) oder $burn.\emptyset$ ausgewählt wird. Der Prozess $\emptyset$ hat aber keine Schritte, daher wird immer $burn.\emptyset$ ausgewählt. Damit gilt in Wirklichkeit also

$$Match' = strike.(\emptyset + burn.\emptyset) = strike.burn.\emptyset = GoodMatch$$

Dagegen wird in $Match$ zwischen zwei nicht-blockierten Prozessen ausgewählt, beide Zweige tragen also zum Verhalten bei, so dass ein Fehlschlagen tatsächlich möglich ist.

Beispiel 5. Eine Uhr, die stehen bleiben kann, wäre demnach wie folgt zu formulieren:

$$Clock = tick. Clock + tick. \emptyset$$

5. Parallelkomposition:

$$P \mid Q \quad (P, Q \in \text{Prozess})$$

Informelle Semantik: “Führe P und Q nebenläufig aus per Interleaving, oder synchronisiert”.

- Interleaving:

Beispiel 6.

$$Clock1 = tick1. Clock1$$

$$Clock2 = tick2. Clock2$$

Welche Aktionsfolgen hat $Clock1 \mid Clock2$?

- Synchronisation:

Zu jedem $a \in \mathcal{A}$ gibt es eine *duale Aktion* (oder *Co-Aktion*) $\bar{a} \in \mathcal{A}$, also $\bar{\cdot} : \mathcal{A} \rightarrow \mathcal{A}$. Diese Abbildung ist eine *Involution*, d.h. es gilt $\bar{\bar{a}} = a$. (Es folgt sofort, dass $\bar{\cdot}$ injektiv ist; wie?)

Die übliche Lesart ist, dass a als Input und $\bar{a}$ als entsprechender Output verstanden wird. Über gleichzeitige Ausführung von a und $\bar{a}$ können parallele Prozesse *synchronisieren*; nach außen sind die Aktionen dann nicht sichtbar. Hierzu wird formal eine gesonderte *stumme Aktion* (*silent action*) τ eingeführt.

6. Restriktion:

$$P \setminus L \quad (L \subseteq \mathcal{A})$$

Informelle Semantik: “Verhalte dich wie P , aber ohne sichtbare Aktionen aus L .” Der Effekt hiervon ist, dass die Aktionen in L nur für interne Synchronisation in P verwendet werden können.

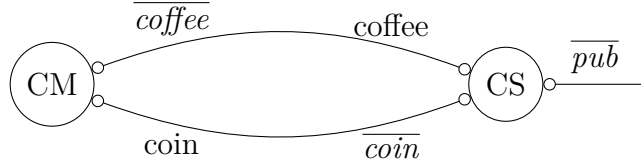
Beispiel 7. Modellierung eines kaffeetrinkenden Informatikers:

$$CoffeeMachine = coin. \overline{coffee}. CoffeeMachine$$

$$ComputerScientist = \overline{pub}. \overline{coin}. coffee. ComputerScientist$$

$$SmUni = (CM \mid CS) \setminus coin \setminus coffee$$

Dann stehen *coin* und *coffee* als Kanäle nicht mehr für die Kommunikation mit der Außenwelt zur Verfügung, wir erhalten also folgendes Kommunikationsdiagramm:



7. Umbenennung: Sei $f : Act \rightarrow Act$ eine Funktion. Dann ist

$$P[f]$$

der Prozess, den man aus P gewinnt, indem man Aktionen gemäß f umbenennt.

Beispiel 8. Seien $P = tick.P$ und $P' = P[tick \mapsto tack]$. Dann gilt effektiv (wiederum in einem später zu präzisierenden Sinn)

$$P' = tack.P'.$$

1.5 Labelled Transition Systems

Die Semantik von CCS wird in Begriffen von *Transitionssystemen* (*labelled transition systems, LTS*) beschrieben, die wir als nächstes definieren.

Definition 9. Ein LTS ist ein Tupel $(Q, \mathcal{A}, (\xrightarrow{\alpha})_{\alpha \in \mathcal{A}})$, bestehend aus

- einer Menge Q von *Zuständen*;
- einer Menge $\mathcal{A}$ von *Aktionen* oder *Labeln* (dies wird bei uns meist die gleiche Menge wie in der CCS-Syntax sein, so dass wir der Einfachheit halber gleich denselben Bezeichner wählen);
- einer *Transitions-* oder *Zustandsübergangsrelation*

$$\rightarrow \subseteq Q \times \mathcal{A} \times Q.$$

Man schreibt

$$\begin{aligned} s &\xrightarrow{\alpha} s', \text{ wenn } (s, \alpha, s') \in \rightarrow; \\ s &\xrightarrow{\alpha}, \text{ wenn } s \xrightarrow{\alpha} s' \text{ für ein } s' \in Q; \\ s &\rightarrow, \text{ wenn } s \xrightarrow{\alpha} \text{ für ein } \alpha \in \mathcal{A}, \end{aligned}$$

und entsprechend negiert ($s \not\rightarrow$ etc.). Ein LTS ist *endlich*, wenn Q endlich ist. Ein LTS ist *endlich verzweigend*, wenn für alle $s \in Q$ und alle $\alpha \in \mathcal{A}$ die Menge $\{s' \in Q \mid s \xrightarrow{\alpha} s'\}$ endlich ist.

Aus einem LTS $T = (\mathcal{Q}, \mathcal{A}, (\xrightarrow{\alpha})_{\alpha \in \mathcal{A}})$ bilden wir ein LTS $T^* = (\mathcal{Q}, \mathcal{A}^*, (\xrightarrow{w})_{w \in \mathcal{A}^*})$, in dem also die Aktionen Wörter über den ursprünglichen Aktionen sind, wie folgt induktiv:

- $s \xrightarrow{\epsilon} s$ stets ($\epsilon \in \mathcal{A}^*$ ist das leere Wort);
- $s \xrightarrow{aw} t$, wenn es ein $s' \in \mathcal{Q}$ gibt, so dass $s \xrightarrow{a} s'$ und $s' \xrightarrow{w} t$.

In Begriffen von Relationskomposition bedeutet die letzte Klausel, dass

$$\xrightarrow{aw} = \xrightarrow{w} \circ \xrightarrow{a} = \xrightarrow{a}; \xrightarrow{w},$$

wobei wir mit $\circ$ applikative Komposition und mit $;$ diagrammatische Komposition bezeichnen (die sich nur durch die Reihenfolge ihrer Argumente unterscheiden).

1.6 Syntax von CCS

Wir legen jetzt noch einmal die Syntax von CCS übersichtlich und vollständig formal fest. Wir fixieren also (erneut) eine abzählbare Menge $\mathcal{A}$ von *Aktionen* mit einer *stummen Aktion* $\tau \in \mathcal{A}$ und einer involutiven Operation $\bar{\cdot}$ (*Komplement*) auf $\mathcal{A} \setminus \{\tau\}$ sowie eine ebenfalls abzählbar unendlichen Menge $\mathcal{K}$ von *Prozessnamen*.

Die Menge $\text{Proc}(\mathcal{K}_0)$ der *Prozesse* P, Q mit Namen aus $\mathcal{K}_0 \subseteq \mathcal{K}$ ist dann durch die Grammatik

$$P, Q ::= K \mid \alpha.P \mid \sum_{i \in I} P_i \mid P \mid Q \mid P[f] \mid P \setminus L.$$

definiert; hierbei sind

- $K \in \mathcal{K}_0$ ein Prozessname;
- $\alpha \in \mathcal{A}$ eine Aktion, und $L \subseteq \mathcal{A} \setminus \{\tau\}$;
- I eine beliebige Indexmenge (die insbesondere auch unendlich sein darf);
- $f : \mathcal{A} \setminus \{\tau\} \rightarrow \mathcal{A} \setminus \{\tau\}$ eine Umbenennungsfunktion mit $f(\bar{a}) = \overline{f(a)}$.

Eine *Spezifikation* besteht dann aus einer Teilmenge $\mathcal{K}_0 \subseteq \mathcal{K}$, einem CCS-Prozess $P \in \text{Proc}(\mathcal{K}_0)$ und für jedes $K \in \mathcal{K}_0$ einer rekursiven Gleichung $K = Q$ mit $Q \in \mathcal{Q}(\mathcal{K}_0)$.

Die Operatorpräzedenz lautet, von höherer zu niedrigerer Priorität,

- Restriktion $\setminus$ und Umbenennung $[]$;

- Aktionspräfix;
- Parallelkomposition $|$;
- Auswahl $+$.

1.7 Semantik von CCS

Sei $\mathcal{S} = (P, K_1 = P_1, \dots, K_n = P_n)$ eine CCS-Spezifikation (die Angabe der vorkommenden Prozessnamen lassen wir meist weg). Die Semantik von $\mathcal{S}$ ist ein ausgezeichneter Zustand ein einem wie folgt induktiv definierten LTS:

- Zustände sind zunächst alle CCS-Prozesse;
- der erwähnte ausgezeichnete Zustand ist P ;
- die Menge der Aktionen ist $\mathcal{A}$;
- die Transitionen sind induktiv durch folgende Regeln definiert (die das Format der *Structural Operational Semantics (SOS)* einhalten).

$$\begin{array}{ll}
(ACT) \frac{}{\alpha.P \xrightarrow{\alpha} P} & (COM3) \frac{P \xrightarrow{a} P' \quad Q \xrightarrow{\bar{a}} Q'}{P \mid Q \xrightarrow{\tau} P' \mid Q'} \\
(SUM_j) \frac{P_j \xrightarrow{\alpha_j} Q}{\sum_{i \in I} P_i \xrightarrow{\alpha_j} Q} & (RES) \frac{P \xrightarrow{\alpha} P'}{P \setminus L \xrightarrow{\alpha} P' \setminus L} \quad (\alpha, \bar{\alpha} \notin L) \\
(COM1) \frac{P \xrightarrow{a} P'}{P \mid Q \xrightarrow{a} P' \mid Q} & (REL) \frac{P \xrightarrow{\alpha} P'}{P[f] \xrightarrow{f(\alpha)} P'[f]} \\
(COM2) \frac{Q \xrightarrow{a} Q'}{P \mid Q \xrightarrow{a} P \mid Q'} & (CON) \frac{P_i \xrightarrow{\alpha} P'}{K_i \xrightarrow{\alpha} P'}
\end{array}$$

Um nicht mit einem unnötig großen LTS zu hantieren, schränkt man dann abschließend noch auf den vom ausgezeichneten Prozess P aus erreichbaren Teil ein, d.h. auf diejenigen Zustände (Prozesse) Q , so dass $P \xrightarrow{w} Q$ für ein Wort $w \in \mathcal{A}^*$.

1.8 Value-passing CCS

Wir führen noch eine syntaktische Erweiterung von CCS ein, in der wir mittels der Aktionen noch Datenwerte transportieren können; wir nennen

diese Erweiterung *value-passing CCS*. Man kann dies direkt auch in CCS kodieren, indem man für jeden möglichen Eingabewert eigene Aktionen und Prozessnamen einführt. So wird zum Beispiel

$$\begin{aligned} C &= in(x). C'(x) \\ C'(x) &= \overline{out}(x). C \end{aligned}$$

kodiert zu:

$$\begin{aligned} C &= in_n. C_n \\ C_n &= \overline{out}_n. C \quad (n \in \mathbb{N}) \end{aligned}$$

Aus Gründen des Komforts definiert man dennoch gesonderte semantische Regeln für value-passing CCS. Man führt dabei wie oben angedeutet Parameter für Aktionen ein, die dann die Form $\alpha(x_1, \dots, x_n)$ haben und mit Werten aus einem vorgegebenen Wertebereich, z.B. den natürlichen Zahlen, belegt werden können (allgemeiner kann man noch Ausdrücke über einer algebraischen Signatur zulassen; dies lassen wir hier der Einfachheit halber weg). Prozess-Gleichungen haben dann die Form $K(x_1, \dots, x_n) = Q$ mit $FV(Q) \subseteq \{x_1, \dots, x_n\}$. Hierbei bezeichnet $FV(Q)$ die Menge freien Variablen in Q , wobei man in diesem Kontext von der bisherigen Symmetrie zwischen Ein- und Ausgabe abweicht: Man betrachtet Aktionen ohne Querstrich als *Eingabeaktionen* und solche mit Querstrich als *Ausgabeaktionen* und fasst die Präfixe $a(x_1, \dots, x_n)$ als Bindungsoperationen auf. Formal wird $FV(Q)$ induktiv wie folgt definiert:

$$\begin{aligned} FV(K(x_1, \dots, x_n)) &= \{x_1, \dots, x_n\} \\ FV(\tau.P) &= FV(P) \\ FV(a(x_1, \dots, x_n).P) &= FV(P) \setminus \{x_1, \dots, x_n\} \\ FV(\bar{a}(x_1, \dots, x_n).P) &= FV(P) \cup \{x_1, \dots, x_n\} \\ FV(\sum_i P_i) &= \bigcup_i FV(P_i) \\ FV(P \mid Q) &= FV(P) \cup FV(Q) \\ FV(P[f]) &= FV(P) \\ FV(P \setminus L) &= FV(P) \end{aligned}$$

(Man beachte die unterschiedliche Behandlung von Input-Aktionen a und Output-Aktionen $\bar{a}$.) In der Semantik sind die Zustände des relevanten LTS nun *Instanzen* der Prozessterme; eine Instanz entsteht aus einem Prozessterm, indem seine *freien* Variablen durch Werte substituiert werden. Wir verlangen in Prozessspezifikationen, dass der ausgezeichnete Prozess *geschlossen* ist, d.h. keine freien Variablen enthält (er ist dann selbst eine Instanz, also ein Zustand des LTS). Die Regeln für Präfix werden wie folgt modifiziert:

$$(ACT_1) \frac{}{\bar{a}(n_1, \dots, n_k).P \xrightarrow{\bar{a}(n_1, \dots, n_k)} P}$$

$$(ACT_2) \frac{}{a(x_1, \dots, x_k).P \xrightarrow{a(n_1, \dots, n_k)} P[n_1/x_1, \dots, n_k/x_k]}$$

Hierbei stehen $n_1, \dots, n_k$ für Werte aus dem vorgegeben Bereich. Man beachte wiederum den Unterschied zwischen Ein- und Ausgaben: Eingaben werden gewissermaßen im Prozess gespeichert, d.h. an die entsprechenden Variablen gebunden, Ausgaben (natürlich) nicht. Mit $[n_1/x_1, \dots, n_k/x_k]$ bezeichnen wir die Substitution von Werten $n_1, \dots, n_k$ für die Variablen $x_1, \dots, x_k$ im Term P ; Achtung: Wie aus anderen Formalismen mit Substitution, wie etwa Prädikatenlogik oder λ -Kalkül, bekannt, werden *nur* freie Vorkommen von Variablen bei der Substitution ersetzt.

1.9 Ableitungen

Definition 10 (α -Ableitung). Sei α eine Aktion und P ein CCS-Prozess. Die α -Ableitung von P ist die Menge

$$D_\alpha(P) = \{P' \mid P \xrightarrow{\alpha} P'\}.$$

Beispiel 11.

$$D_{up}(up.(C \mid \emptyset)) = \{C \mid \emptyset\}$$

Satz 12. Sei $(P, K_1 = P_1, \dots, K_n = P_n)$ eine CCS-Spezifikation. Dann gelten folgende Gleichungen:

$$\begin{aligned} D_\alpha(K_i) &= D_\alpha(P_i) \\ D_\alpha(\alpha.P) &= \{P\}, \quad D_\alpha(\beta.P) = \{\} \quad \text{für } \alpha \neq \beta \\ D_\alpha(\sum_i Q_i) &= \cup_i D_\alpha(Q_i) \\ D_\alpha(P \mid Q) &= \{P' \mid Q \mid P' \in D_\alpha(P)\} \cup \{P \mid Q' \mid Q' \in D_\alpha(Q)\} \quad \text{für } \alpha \neq \tau \\ D_\tau(P \mid Q) &= \{P' \mid Q \mid P' \in D_\tau(P)\} \cup \{P \mid Q' \mid Q' \in D_\tau(Q)\} \cup \\ &\quad \{P' \mid Q' \mid \exists \alpha \in \mathcal{A} \setminus \{\tau\}. P' \in D_\alpha(P), Q' \in D_{\bar{\alpha}}(Q)\} \\ D_\alpha(P \setminus L) &= \begin{cases} \{P' \setminus L \mid P' \in D_\alpha(P)\} & \text{für } \alpha \notin L \cup \bar{L} \\ \emptyset & \text{sonst} \end{cases} \\ D_\alpha(P[f]) &= \cup_{f(\beta)=\alpha} D_\beta(P)[f] \end{aligned}$$

Der Beweis der Gleichungen ergibt sich jeweils unmittelbar aus der Struktur der Semantikregeln.

Die Gleichung $C = a.C + C$ hat mehrere (sogar unendlich viele) Lösungen:

1. $C = a.a.a. \dots$
2. $C = a.a.a. \dots + b.b.b. \dots$

Definition 13 (Bewachte Terme (Guarded terms)). Eine Spezifikation $(P, K_1 = P_1, \dots, K_n = P_n)$ heißt *guarded* bzw. *bewacht*, wenn alle Vorkommnisse von K_i in P_j ausschließlich in Untertermen der Form $\alpha.Q$ vorkommen.

Satz 14. Für jede bewachte Spezifikation $(P, K_1 = P_1, \dots, K_n = P_n)$ terminieren die rekursive Gleichungen für die $D_\alpha(P)$; insbesondere ist ihre Lösung eindeutig.

(Warum?) Da wir aus den Ableitungen wiederum ein LTS rekonstruieren können, ist also das durch die Semantikregeln definierte LTS das einzige, das die Spezifikation erfüllt, sofern die Spezifikation eben bewacht ist. Natürlich definieren die Semantikregeln auch sonst ein eindeutiges LTS; man überlege sich z.B., welches Verhalten sich aus den Regeln für den durch $A = A$ definierten Prozess ergibt. Ohne Bewachtheit ist aber die durch die Semantikregeln definierte Lösung der Spezifikation eben nicht die einzig mögliche.

2 Verhaltensäquivalenz

Wir suchen nunmehr nach geeigneten Begriffen von *Verhaltensäquivalenz* von Prozessen. Solche Begriffe sollen zum einen unsere Intuition davon widerspiegeln, was es heißt, zwei Prozesse durch Beobachtung oder Interaktion unterscheiden zu können (oder eben nicht), und zum anderen gewisse Abschlusseigenschaften einhalten. Insbesondere sollte z.B. eine Verhaltensäquivalenz tatsächlich eine Äquivalenzrelation sein (also reflexiv, symmetrisch und transitiv), und zum anderen sollte sie nach Möglichkeit eine *Kongruenz* bezüglich der Prozessalgebraoperationen sein, d.h. ein Austauschen von Prozessen gegen verhaltensäquivalente Prozesse innerhalb eines Prozessterms sollte wiederum verhaltensäquivalente Prozesse liefern. Wir führen im folgenden verschiedene Begriffe von Verhaltensäquivalenz ein und beurteilen sie vor diesem Hintergrund.

Wir erinnern vorab kurz an einige Standardbegriffe hinsichtlich Eigenschaften von Relationen:

Definition 15 (Äquivalenz, Präordnung, partielle Ordnung). Eine Relation $R \subseteq X \times X$ heißt

- *reflexiv*, wenn $\forall x \in X. (x, x) \in R$;
- *symmetrisch*, wenn $\forall x, y \in X. (x, y) \in R \Rightarrow (y, x) \in R$;

- *transitiv*, wenn $\forall x, y, z \in X. (x, y) \in R \wedge (y, z) \in R \Rightarrow (x, z) \in R$;
- *antisymmetrisch*, wenn $\forall x, y \in X. (x, y) \in R \wedge (y, x) \in R \Rightarrow x = y$;
- eine *Präordnung*, wenn R reflexiv und transitiv ist;
- eine *Äquivalenz*, wenn R reflexiv, transitiv und symmetrisch ist;
- eine (*partielle*) *Ordnung*, wenn R reflexiv, transitiv und antisymmetrisch ist.

Definition 16 (Kongruenz). Eine Äquivalenzrelation $R \subseteq X \times X$ ist eine *Kongruenz* bezüglich einer Operation $f : X^n \rightarrow X$, wenn mit $x_1 R x'_1, \dots, x_n R x'_n$ stets auch $f(x_1, \dots, x_n) R f(x'_1, \dots, x'_n)$ gilt (alternativ reicht der Fall, wo $x_i = x'_i$ für alle i bis auf eins; warum?). Ferner ist R eine *Kongruenz* für eine gegebene Sprache, wie z.B. CCS, wenn R eine Kongruenz bezüglich aller Operationen der Sprache ist, äquivalenterweise wenn $\forall x, y \in X. (x, y) \in R \Rightarrow (C(x), C(y)) \in R$ für alle Kontexte $C(\cdot)$. (Wir erinnern daran, dass ein Kontext ein Term mit genau einem Freiplatz $(\cdot)$ ist).

2.1 Trace-Äquivalenz

Unser erster Begriff von Verhaltensäquivalenz ist rein beobachtend, d.h. orientiert sich nur daran, welche Aktionsfolgen ein Prozess ausführen kann:

Definition 17 (Spuräquivalenz (Trace Equivalence)). Sei $T = (\mathcal{Q}, \mathcal{A}, (\xrightarrow{\alpha})_{\alpha \in \mathcal{A}})$ ein LTS. Ein *Trace* eines Zustands $x \in \mathcal{Q}$ ist eine Wort $\alpha_1 \dots \alpha_k \in \mathcal{A}^*$ ($k \geq 0$), so dass eine Folge

$$x = x_0 \xrightarrow{\alpha_1} x_1 \dots \xrightarrow{\alpha_k} x_k$$

von Transitionen in T existiert. Mit $Traces(x)$ bezeichnen wir die Menge aller Traces von x ; in früher eingeführter Notation:

$$Traces(x) = \{w \in \mathcal{A}^* \mid x \xrightarrow{w}\}.$$

Zustände x, y in (womöglich verschiedenen) LTS heißen *spuräquivalent* oder *trace-äquivalent*, wenn $Traces(x) = Traces(y)$.

Beispiel 18 (Trace-Äquivalenz). Wir haben

$$\begin{aligned} Traces(a. (b. \emptyset + c. \emptyset)) &= \{\epsilon, a, ab, ac\} \\ Traces(a.b. \emptyset + a.c. \emptyset) &= \{\epsilon, a, ab, ac\}, \end{aligned}$$

die Prozesse $a.(b.\emptyset + c.\emptyset)$ und $a.b.\emptyset + a.c.\emptyset$ sind also trace-äquivalent. Das ist aber durchaus problematisch; z.B. betrachte man folgende sehr ähnliche Prozesse:

$$\begin{aligned} CM &= \text{coin}.\overline{\text{coffee}}.CM + \overline{\text{tea}}.CM \\ CM' &= \text{coin}.\overline{\text{coffee}}.CM' + \text{coin}.\overline{\text{tea}}.CM' \\ CA &= \overline{\text{coin}}.\text{coffee}.CA \end{aligned}$$

und bilde dann folgende zusammengesetzten Prozesse:

$$\begin{aligned} P_1 &= (CA \mid CM) \setminus \{\text{coin}, \text{coffee}, \text{tea}\} \\ P_2 &= (CA \mid CM') \setminus \{\text{coin}, \text{coffee}, \text{tea}\} \end{aligned}$$

Nach Spuräquivalenz sind diese Prozesse gleich. Man denkt zunächst, dass es sich hier um ein Kongruenzproblem handelt, das durch die Interaktion mit CA entsteht; wir werden aber gleich sehen, dass das gerade nicht der Fall ist. Vielmehr ist das unterliegende Problem, dass Trace-Äquivalenz offenbar die Gleichung

$$a.(P + Q) \neq a.P + a.Q$$

erfüllt, die offenbar nicht immer wünschenswert ist.

Satz 19. *Trace-Äquivalenz ist eine Kongruenz für CCS.*

Beweis. Wir können eine Berechnungsvorschrift für $\text{Traces}(P)$ durch strukturelle Rekursion über Terme angeben; Details erarbeiten wir in der Veranstaltung gemeinsam. Warum folgt daraus die Kongruenzeigenschaft? $\square$

2.2 Bisimilarität und Bisimulation

Wir führen nun eine *feinergranulare* (oder einfach *feinere*) Form von Verhaltensäquivalenz ein, die also *mehr* Prozesse unterscheidet und auf diesem Wege das im vorigen Abschnitt angedeutete Problem der Trace-Äquivalenz eliminiert. Dies ist der Begriff der *starken Bisimilarität*. (Es gibt zahlreiche weitere Äquivalenzen zwischen den beiden Begriffen; in den Übungen lernen wir noch *completed-Trace-Äquivalenz* oder *starke Trace-Äquivalenz* kennen.)

Definition 20 (Bisimulation/Bisimilarität). Seien T_1 und T_2 LTS über demselben Alphabet $\mathcal{A}$ von Aktionen mit Zustandsmengen $\mathcal{Q}_1$ bzw. $\mathcal{Q}_2$. Eine binäre Relation $R \subseteq \mathcal{Q}_1 \times \mathcal{Q}_2$ heißt eine (*starke*) *Simulation*, wenn für alle $x \in \mathcal{Q}_1$, $y \in \mathcal{Q}_2$ mit xRy gilt:

Forth: Für $x \xrightarrow{\alpha} x'$ existiert stets y' mit $y \xrightarrow{\alpha} y'$ und $x'Ry'$.

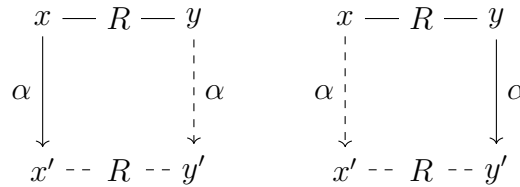
Ferner heißt R eine (*starke*) *Bisimulation*, wenn sowohl R als auch seine Umkehrrelation

$$R^- = \{(y, x) \mid (x, y) \in R\} \subseteq Q_2 \times Q_1$$

Bisimulationen sind; die entsprechende Bedingung für R^- übersetzt sich dann in Begriffen von R statt R^- , d.h. für xRy , als

Back: Für $y \xrightarrow{\alpha} y'$ existiert stets x' mit $x \xrightarrow{\alpha} x'$ und $x'Ry'$.

In diagrammatischer Darstellung:



(*Starke*) *Bisimilarität* ist die durch

$$x \sim y \iff \text{es gibt eine Bisimulation } R \text{ mit } xRy$$

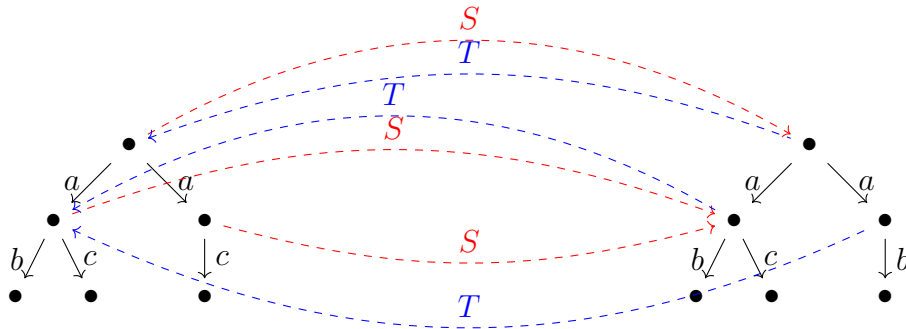
definierte Relation $\sim$ zwischen Zuständen von LTS (mit gemeinsamer Aktionsmenge).

Zusammenfassend haben wir

$$\frac{\text{Beweis}}{\text{Gültigkeit}} = \frac{\text{Bisimulation}}{\text{Bisimilarität}},$$

d.h. die Angabe einer Bisimulation zwischen zwei Zuständen beweist ihre Bisimilarität.

Beispiel 21. Bisimilarität ist stärker als gegenseitige Simulation:



Diese beiden LTS simulieren sich gegenseitig, sind aber nicht bisimilar.

Beispiel 22. Wir definieren ein LTS $(\mathcal{Q}, \mathcal{A}, (\xrightarrow{\alpha})_{\alpha \in \mathcal{A}})$ durch

$$\begin{aligned}\mathcal{Q} &= \{s, s_1, s_2, t_0, t_1\} \\ \mathcal{A} &= \{a, b\} \\ \xrightarrow{a} &= \{(s, s_1), (s, s_2), (t, t_1)\} \\ \xrightarrow{b} &= \{(s_1, s_2), (s_2, s_2), (t_1, t_1)\}\end{aligned}$$

Wir haben dann eine Bisimulation $R = \{(s, t), (s_1, t_1), (s_2, t_1)\}$:

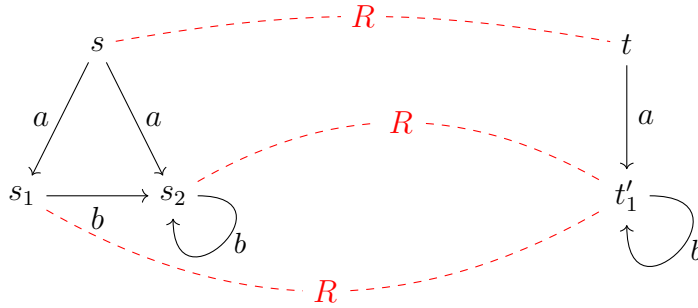


Abbildung 1: Bisimulation

Beispiel 23. Wir definieren ein LTS $(\mathcal{Q}, \mathcal{A}, (\xrightarrow{\alpha})_{\alpha \in \mathcal{A}})$ durch

$$\begin{aligned}\mathcal{Q} &= \{s_i \mid i \geq 1\} \cup \{t\} \\ \mathcal{A} &= \{a\} \\ \xrightarrow{a} &= \{(s_i, s_{i+1}) \mid i \geq 1\} \cup \{(t, t)\}.\end{aligned}$$

Dann haben wir $s \sim t$, weil $R = \{(s_i, t) \mid i \geq 1\}$ eine Bisimulation ist:

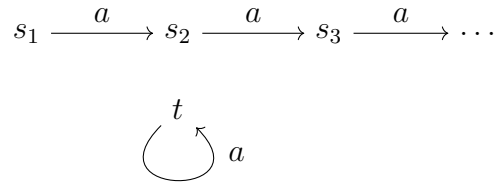
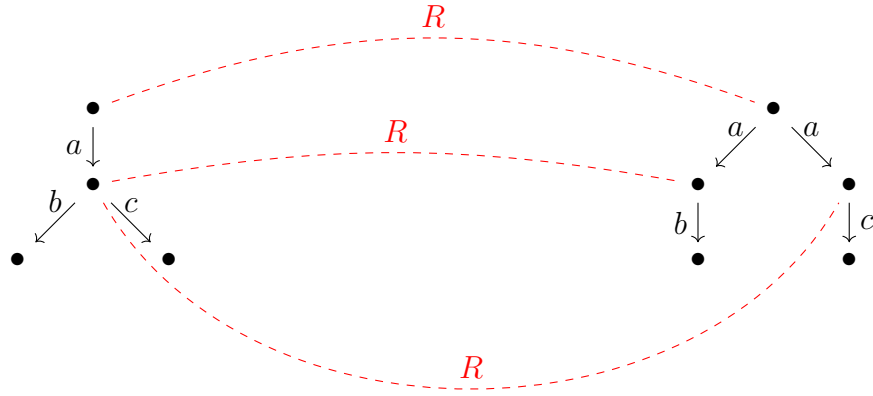


Abbildung 2: Beispiel 2: Bisimulation

Beispiel 24. Gegenbeispiel:

$$a.(b.0 + c.0) \not\sim a.b.0 + a.c.0$$

Der Versuch, eine Bisimulation zwischen den beiden Prozessen zu konstruieren, scheitert wie folgt:



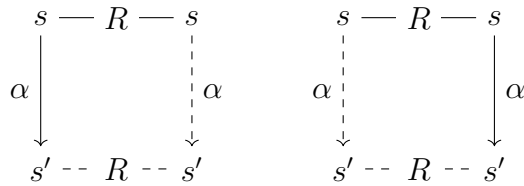
Satz 25. Es gilt Folgendes:

1. Bisimilarität ist eine Äquivalenzrelation
2. Bisimilarität ist die größte Bisimulation
3. Für Zustände s, t gilt

$$s \sim t \iff \begin{aligned} &\forall s \xrightarrow{\alpha} s'. \exists t \xrightarrow{\alpha} t'. s' \sim t' \text{ und} \\ &\forall t \xrightarrow{\alpha} t'. \exists s \xrightarrow{\alpha} s'. s' \sim t'. \end{aligned}$$

Beweis. 1. $\sim$ ist eine Äquivalenz:

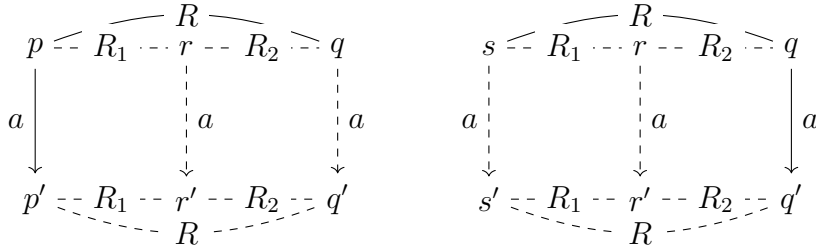
- Reflexivität: für alle s gilt $s \sim s$, weil $R = \{(s, s) \mid s \in \mathcal{Q}\}$ eine Bisimulation ist:



- Symmetrie: Sei $s \sim t$; zu zeigen ist $t \sim s$. Nach Voraussetzung gibt es eine Bisimulation R mit $s R t$. Dann ist R^- wieder eine Bisimulation, denn weil R eine Bisimulation ist, sind sowohl R^- als auch $(R^-)^- = R$ Simulationen. Natürlich gilt $t R^- s$, und damit $t \sim s$.
- Transitivität: Sei $s \sim t$ und $t \sim u$; zu zeigen ist $s \sim u$. Nach Voraussetzung existieren Bisimulationen R_1 und R_2 mit $s R_1 t$ und $t R_2 u$. Wir setzen

$$R = R_2 \circ R_1 = \{(p, q) \mid \exists r. (p, r) \in R_1, (r, q) \in R_2\}.$$

Dann gilt nach Voraussetzung $s R u$. Ferner ist R eine Bisimulation:



Das zeigt $s \sim t$.

2. Da Bisimilarität $\sim$ per Definition alle Bisimulationen enthält, ist nur noch zu zeigen, dass $\sim$ selbst eine Bisimulation ist. Seien also $s \sim t$ und $s \xrightarrow{\alpha} s'$. Nach Voraussetzung existiert dann eine Bisimulation R mit $s R t$. Dann haben wir

$$\begin{array}{ccc} s & \text{---} R \text{---} & t \\ \alpha \downarrow & & \downarrow \alpha \\ s' & \text{---} R \text{---} & t' \end{array}$$

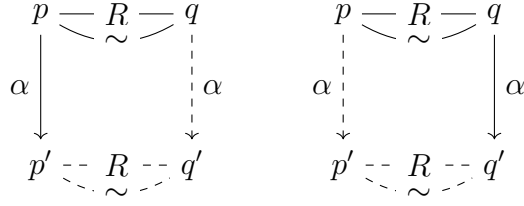
und somit auch $s' \sim t'$. Der Beweis Back-Bedingung ist symmetrisch.

3. Die Implikation „ $\Rightarrow$ “ gilt nach dem vorherigen Punkt; wir zeigen „ $\Leftarrow$ “. Dazu zeigen wir, dass unter den gegebenen Annahmen

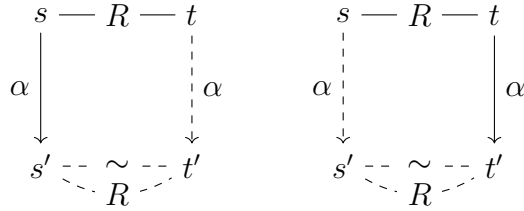
$$R = \sim \cup \{(s, t)\}$$

eine Bisimulation ist. Wir unterscheiden für $p R q$ die beiden Fälle in der Definition von R :

1. **Fall:** $(p, q) \in \sim$, dann:



2. Fall: $(p, q) = (s, t)$; dann haben wir per Annahme



□

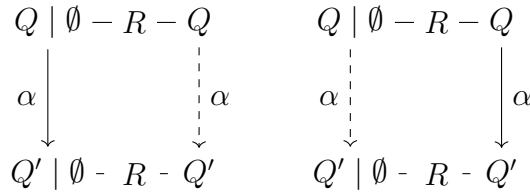
Satz 26. *Es gilt:*

$$\begin{aligned}
P + Q &\sim Q + P \\
P + P &\sim P \\
P + 0 &\sim P \\
P + (Q + R) &\sim (P + Q) + R \\
P \mid Q &\sim Q \mid P \\
P \mid \emptyset &\sim P \\
(P \mid Q) \mid R &\sim P \mid (Q \mid R)
\end{aligned}$$

Beweis. Wir greifen zwei Beispiele heraus:

$P + Q = Q + P$: Dies ist eine leichte Anwendung von Theorem 25.3: Nach der Semantik von CCS haben $P + Q$ und $Q + P$ unter allen Aktionen die gleichen Nachfolger (nämlich die von P und die von Q), so dass wir in der Bedingung in Theorem 25.3 sogar $=$ statt nur $\sim$ haben. Die anderen Aussagen über $+$ zeigt man genauso.

$P \mid \emptyset = P$: Wir definieren $R = \{(P \mid \emptyset, P) \mid P \text{ ein CCS-Prozess}\}$ und zeigen, dass R eine Bisimulation ist:



□

Beispiel 27. Im allgemeinen gilt *nicht* $(P + Q) \mid S \sim (P \mid S) + (Q \mid S)$; z.B. $(a.\emptyset + b.\emptyset) \mid c.\emptyset \not\sim (a.\emptyset \mid c.\emptyset) + (b.\emptyset \mid c.\emptyset)$:



Satz 28. $\sim$ ist eine Kongruenz für CCS

Beweis. Wir beweisen den interessantesten Fall, die Parallelkomposition $\mid$:
Sei

$$R = \{(P \mid S, Q \mid S) \mid P, Q, S \text{ CCS-Prozesse}, P \sim Q\};$$

Es reicht zu zeigen, dass R eine Bisimulation ist. Da R symmetrisch ist, genügt die Forth-Bedingung. Wir unterscheiden Fälle über die Nachfolger von $P \mid S$ gemäß der Semantikregeln:

(COM1): Sei $P \xrightarrow{a} P'$; da $P \sim Q$, existiert Q' mit $Q \xrightarrow{a} Q'$ und $P' \sim Q'$.
Dann haben wir

$$\begin{array}{ccc} P \mid S & \xrightarrow{\quad} R & \xrightarrow{\quad} Q \mid S \\ \alpha \downarrow & & \downarrow \alpha \\ P' \mid S & \text{----} R & \text{----} Q' \mid S \end{array}$$

(COM2): Sei $S \xrightarrow{a} S'$; dann haben wir

$$\begin{array}{ccc} P \mid S & \xrightarrow{\quad} R & \xrightarrow{\quad} Q \mid S \\ \alpha \downarrow & & \downarrow \alpha \\ P \mid S' & \text{----} R & \text{----} Q \mid S' \end{array}$$

(COM3): Sei $P \xrightarrow{a} P'$ und $S \xrightarrow{\bar{a}} S'$. Da $P \sim Q$, existiert Q' mit $Q \xrightarrow{a} Q'$ und $P' \sim Q'$; damit haben wir

$$\begin{array}{ccc} P \mid S & \xrightarrow{\quad} R & \xrightarrow{\quad} Q \mid S \\ \tau \downarrow & & \downarrow \tau \\ P \mid S' & \text{----} R & \text{----} Q' \mid S' \end{array}$$

□

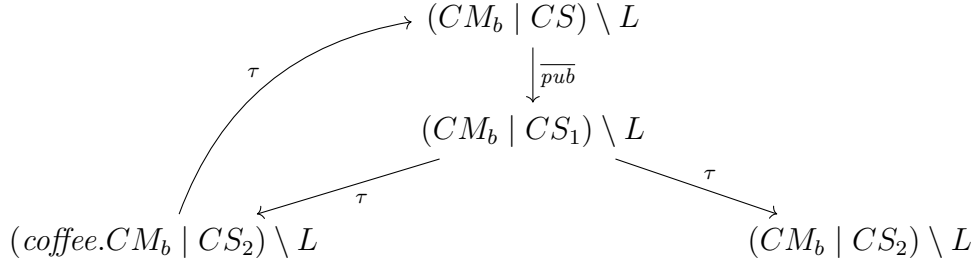
2.3 Schwache Bisimilarität

Die bisherigen Begriffe von Prozessäquivalenz berücksichtigen die besondere Rolle der stummen Aktion nicht; z.B. gilt natürlich

$$a.\emptyset \approx a.\tau.\emptyset,$$

in der Tat sind die beiden Prozesse noch nicht einmal trace-äquivalent. Intuitiv sollte die Äquivalenz aber letztlich gelten, weil ja die zusätzliche τ -Aktion auf der rechten Seite gerade nicht beobachtbar sein soll. Ähnliches gilt unter den geübten Definitionen von CS und CM für $(CM \mid CS) \setminus L$ mit $L = \{coffee, coin\}$ einerseits und $SmUni = \overline{pub}.SmUni$ andererseits. Die folgende alternative Kaffeemaschine CM_b wollen wir dagegen in diesem Kontext sicher weiterhin unterscheiden:

$$CM_b = coin.\overline{coffee}.CM_b + coin.CM_b$$



D.h. eine womögliche Idee, durch τ -Aktionen verbundene Zustände einfach zu identifizieren, tut offenbar nicht das Richtige. Stattdessen modifizieren wir die Transitionsrelation:

Definition 29 (Double-Arrow-Konstruktion). Sei $(\mathcal{Q}, \mathcal{A}, (\overset{\alpha}{\rightarrow})_{\alpha \in \mathcal{A}})$ ein LTS und $\tau \in \mathcal{A}$. Die *schwache Transitionsrelation* $\overset{\alpha}{\Rightarrow} \subseteq \mathcal{Q} \times \mathcal{Q}$ ist wie folgt definiert: $P \overset{\alpha}{\Rightarrow} Q$ bedeutet, dass

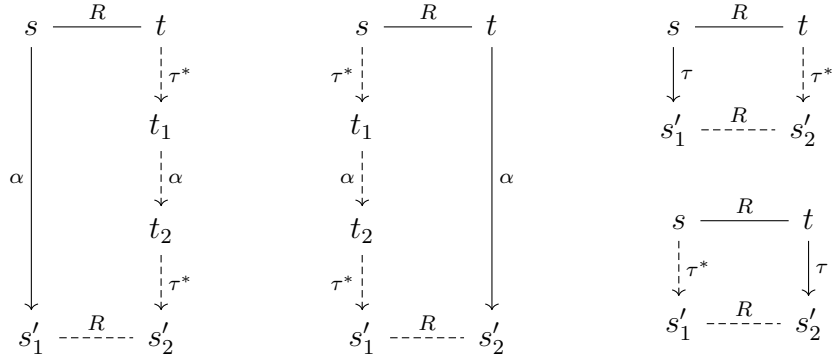
- $P \xrightarrow{\tau}^* P' \xrightarrow{\alpha} Q' \xrightarrow{\tau}^* Q$ mit geeignetem P', Q' wenn $\alpha \neq \tau$, und
- $P \xrightarrow{\tau}^* Q$, wenn $\alpha = \tau$.

(dabei bezeichnet $\xrightarrow{\tau^*}$ oder auch $\xrightarrow{\tau^*}$ den reflexiv-transitiven Abschluss von $\xrightarrow{\tau}$, d.h. $P \xrightarrow{\tau^*} Q$ genau dann, wenn es $w = \tau \dots \tau$ gibt, so dass $P \xrightarrow{w} Q$).

Definition 30 (Schwache Bisimulation/Bisimilarität). Eine binäre Relation $R \in \mathcal{Q} \times \mathcal{Q}$ ist eine *schwache Bisimulation*, wenn aus sRt stets folgt, dass

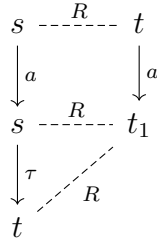
- $s \xrightarrow{\alpha} s' \implies \exists t'. t \xrightarrow{\alpha} t' \wedge s'Rt'$ (*forth*)
- $t \xrightarrow{\alpha} t' \implies \exists s'. s \xrightarrow{\alpha} s' \wedge s'Rt'$ (*back*).

Graphisch:



Wir schreiben $P \approx Q$, wenn es eine schwache Bisimulation R gibt, so dass $P R Q$. Die Relation $\approx$ heißt *schwache Bisimilarität*.

Beispiel 31. $a.\emptyset \approx a.\tau.\emptyset$, aber $a.\emptyset \not\approx a.\tau.\emptyset$:



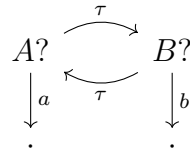
Beispiel 32. $(CM \mid CS) \setminus L \approx SmUni$ mit $SmUni = \overline{pub}.SmUni$:

Beispiel 35. Man betrachte die folgende ‘Polling-Schleife’:

$$A? = a. \emptyset + \tau.B?$$

$$B? = b. \emptyset + \tau.A?$$

Als LTS:



Für

$$P = a. \emptyset + b. \emptyset$$

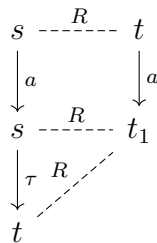
haben wir dann

$$A? \approx P,$$

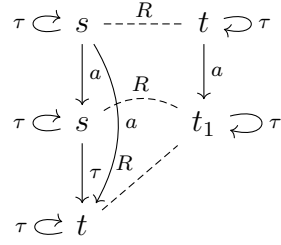
obwohl $A?$ hat einen Livelock hat, P dagegen nicht. Diese Phänomen ist unter dem Namen *fair abstraction from divergence* bekannt.

Satz 36. Seien $T_i = (\mathcal{Q}_i, \mathcal{A}, (\xrightarrow{\alpha})_{\alpha \in \mathcal{A}})$, $i = 1, 2$, LTS mit $\tau \in \mathcal{A}$, und seien $T'_i = (\mathcal{Q}_i, \mathcal{A}, (\xRightarrow{\alpha})_{\alpha \in \mathcal{A}})$. Dann ist $R \subseteq \mathcal{Q}_1 \times \mathcal{Q}_2$ eine schwache Bisimulation zwischen T_1 und T_2 genau dann, wenn R eine starke Bisimulation zwischen T'_1 und T'_2 ist.

Beispiel 37. T :



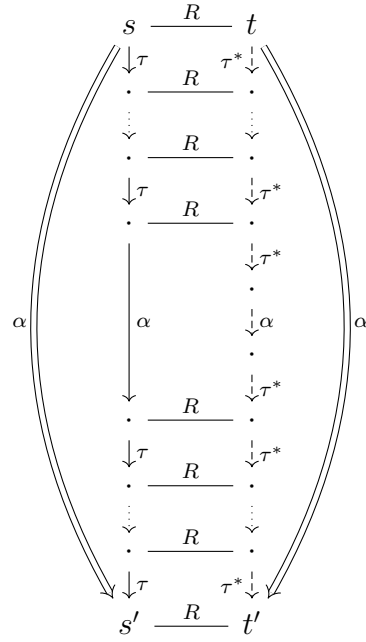
T' :



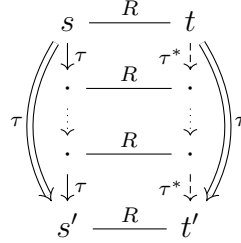
Satz 36. “ $\Leftarrow$ ”: Klar durch Vergleich der Definitionen, da aus $s \xrightarrow{\alpha} s'$ stets $s \xRightarrow{\alpha} s'$ folgt.

“ $\Rightarrow$ ”: Sei sRt und $s \xRightarrow{\alpha} s'$. Wir unterscheiden Fälle über α :

$\alpha \neq \tau$:



$\alpha = \tau$:



□

Korollar 38.

1. $\approx$ ist eine Äquivalenz.
2. $\approx$ ist die größte schwache Bisimulation.
3. Für Zustände s, t gilt $s \approx t$ genau dann, wenn für alle $\alpha \in \mathcal{A}$ gilt:
 - $s \xrightarrow{\alpha} s' \Rightarrow \exists t'. t \xRightarrow{\alpha} t' \wedge s' \approx t'$
 - $t \xrightarrow{\alpha} t' \Rightarrow \exists s'. s \xRightarrow{\alpha} s' \wedge s' \approx t'$

Beweis. Aus Satz 36 und Satz 25 erhalten wir sofort alle drei Behauptungen, allerdings mit $\xRightarrow{\alpha}$ statt $\xrightarrow{\alpha}$ auf den linken Seiten der Forth- und Back-Bedingungen in Behauptung 3. Man zeigt wie in Satz 36, dass die Bedingungen in dieser Form zu den in der Behauptung angegebenen äquivalent sind. □

Beispiel 39.

$$\emptyset \approx \tau. \emptyset, \quad \text{aber} \\ a. \emptyset + \emptyset \not\approx a. \emptyset + \tau. \emptyset$$

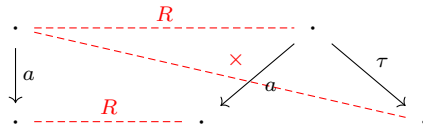


Abbildung 3:

Satz 40. $\approx$ ist bezüglich aller Operationen außer $+$ eine Kongruenz.

2.4 Bisimulationsspiele

Wir formalisieren als nächstes eine spielorientierte Sichtweise auf Bisimulation, d.h. wir geben ein Spiel an, in dem ein Spieler versucht, die Bisimilarität zweier Zustände zu zeigen, und einer, sie zu widerlegen; das Spiel ist so gebaut, dass der richtige Spieler gewinnt.

Die Spieler nennen wir

- S für *Spoiler* (auch *Attacker*, *Opponent*)
- D für *Duplicator* (auch *Defender*, *Proponent*, *Prover*)

Wie man anhand der Namen vielleicht errät, spielt D für und S gegen Bisimilarität.

Das Spiel wird in Runden gespielt; jede Runde besteht aus einem Zug von S und einem anschließenden Zug von D . Das Spiel dauert so lange, bis einer der Spieler nicht mehr ziehen kann, und kann insbesondere auch unendlich lange dauern.

Seien nun $T_1 = (Q_1, \mathcal{A}, \rightarrow)$ und $T_2 = (Q_2, \mathcal{A}, \rightarrow)$ LTS, mit Anfangszuständen $s_0 \in Q_1$, $t_0 \in Q_2$. Wir beschreiben das *Bisimulationsspiel* $\mathcal{G}(s_0, t_0)$. Die Spielpositionen nennen wir *Positionen*; eine Position ist ein Paar (s, t) von Zuständen $s \in Q_1$, $t \in Q_2$. Die Spielregeln sind dann wie folgt beschrieben:

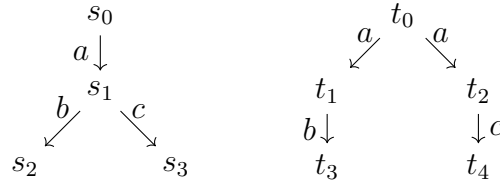
- *Anfangsposition*: (s_0, t_0)
- *Züge* bei aktueller Position (s, t) :
 - S wählt eines der beiden LTS, z.B. T_2 , und darin eine Transition vom aktuellen Zustand, hier $t \xrightarrow{\alpha} t'$.
 - D wählt daraufhin im anderen LTS, im Beispiel also T_1 , eine Transition $t \xrightarrow{\alpha} t'$.

Die nach der Runde erreichte neue Position ist dann (s', t') .

- *Gewinnbedingung*: Wer am Zug ist und keinen Zug hat (d.h. von einer Transition von einem Deadlock aus wählen müsste), verliert. Unendlich lange Partien gewinnt D .

Satz 41. *Gegeben Daten wie oben gilt $s_0 \sim t_0$ genau dann, wenn D in $\mathcal{G}(s_0, t_0)$ eine Gewinnstrategie besitzt.*

Beispiel 42. In der folgenden Situation gewinnt S :

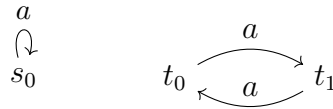


Eine Gewinnstrategie (von mehreren) für S ist wie folgt:

- Ziehe in der ersten Runde $t_0 \xrightarrow{a} t_1$. Die einzige Antwort von D ist dann $s_0 \xrightarrow{a} s_1$, so dass die Position (s_1, t_1) erreicht wird. (Es ist in der Tat egal, was S in der ersten Runde tut, durch diese Wahl vermeiden wir nur Fallunterscheidungen in der Spielanalyse.)
- Ziehe $s_1 \xrightarrow{c} s_3$. Spieler D kann nicht antworten und verliert.

Laut obigem Satz sind also s_0, t_0 nicht bisimilar, was wir ja auch schon wussten.

Beispiel 43. Im folgenden Fall gewinnt offenbar D :



– S zieht in jeder Runde eine a -Transition auf einer der beiden Seiten, die D dann mit der (jeweils einzigen) verfügbaren a -Transition auf der anderen Seite beantwortet. Da das Spiel damit nie abbricht, gewinnt D . Nach obigem Satz sind die Prozesse also bismilar.

Bemerkung 44. Es ist bei unendlichen Spielen im allgemeinen nicht klar, dass sie *determiniert* sind, d.h. in jedem Fall einer der Spieler eine Gewinnstrategie hat. Bisimulationsspiele gehören aber zu einer Klasse von Spielen, wo dies der Fall ist; ein Beweis führt für unsere gegenwärtigen Zwecke zu weit ab vom Thema. Wir werden diese Tatsache im Beweis von Satz 41 auch nicht verwenden.

Beweis (Satz 41). ‘ $\Rightarrow$ ’: Sei R eine Bisimulation mit $s_0 R t_0$. Wir verwenden R als Gewinnstrategie, genauer gesagt als *Invariante*, deren Durchsetzung D den Gewinn garantiert. Allgemein werden Gewinnstrategien in Spielen gerne durch Invarianten beschrieben. Eine Invariante I ist dabei eine Eigenschaft von Positionen; damit sie tatsächlich eine Gewinnstrategie liefert, muss sie folgende Eigenschaften erfüllen:

- I muss für die *Anfangsposition* gelten;

- I muss tatsächlich *durchsetzbar* sein, d.h. wenn I vor einer Runde erfüllt ist, muss D erzwingen können, dass I auch nach der Runde erfüllt ist;
- Partien, in denen durchgängig die Invariante gilt, müssen von D *gewonnen* werden.

Wir prüfen der Vollständigkeit halber alle Bedingungen:

- Für die Anfangsposition gilt s_0Rt_0 per Annahme.
- Wenn vor einer Runde sRt gilt und S o.E. eine Transition $t \xrightarrow{a} t'$ in T_2 spielt, dann existiert per *Back*-Bedingung eine Transition $s \xrightarrow{a} s'$ mit $s'Rt'$. Spieler D spielt diese Transition und hat die Invariante damit durchgesetzt.
- Gegeben sei nun eine Partie, in dem für alle Positionen sRt gilt. Wir unterscheiden zwei Fälle:
 - Das Spiel bricht ab, weil einer der Spieler nicht ziehen kann. Wir haben gesehen, dass D immer ziehen kann, wenn S ziehen kann; in diesem Fall kann also S nicht ziehen, und D gewinnt.
 - Das Spiel bricht nicht ab; dann gewinnt D .

‘ $\Leftarrow$ ’: Wir zeigen, dass

$$R = \{(s, t) \in Q_1 \times Q_2 \mid D \text{ gewinnt } \mathcal{G}(s, t)\}$$

eine Bisimulation ist. Per Symmetrie der Spielregeln genügt es, die Forth-Bedingung zu zeigen. Sei dazu sRt und $s \xrightarrow{a} s'$. Dann kann S in der ersten Runde von $\mathcal{G}(s, t)$ eben diese Transition $s \xrightarrow{a} s'$ ziehen; sei $t \xrightarrow{a} t'$ die Antwort von D gemäß seiner Gewinnstrategie. Da dies eben eine Gewinnstrategie ist, gewinnt D das verbleibende Spiel $\mathcal{G}(s', t')$ (da dessen Gewinnbedingungen aus denen von $\mathcal{G}(s, t)$ folgen), d.h. es gilt wie verlangt $s'Rt'$. $\square$

2.5 Schwache Bisimulationsspiele

Völlig analog erhält man ein Spiel für schwache Bisimilarität, das sich von dem für starke Bisimilarität nur dadurch unterscheidet, dass D auch mit einer schwachen Transition $\xRightarrow{a}$ antworten kann. Man zeigt wie oben:

Satz 45. *Für Zustände s', t' in LTS gilt $s \approx t$ genau dann, wenn D eine Gewinnstrategie in $\mathcal{G}(s, t)$ hat.*

2.6 Bisimulation als Fixpunkt

Eine weitere besonders wichtige Perspektive auf Bisimilarität liefert die Fixpunkttheorie; insbesondere erhält man durch eine Fixpunktcharakterisierung von Bisimilarität mehr oder weniger sofort einen Algorithmus zur Berechnung der Bisimilaritätsrelation auf einem gegebenen LTS. Wir führen zunächst in die benötigten Aspekte der Fixpunkttheorie auf geordneten Mengen ein.

Definition 46. Sei $(X, \sqsubseteq)$ eine (partiell) geordnete Menge (Poset). Eine Funktion $f : X \rightarrow X$ ist *monoton*, wenn $\forall x, y \in X. x \sqsubseteq y \implies f(x) \sqsubseteq f(y)$. Ein Element $x \in X$ ist *obere Schranke* einer Teilmenge $S \subseteq X$ (Schreibweise: $S \sqsubseteq x$), wenn $\forall y \in S. y \sqsubseteq x$, und *Supremum* von S ($x = \bigsqcup S$), wenn x *kleinste* obere Schranke von S ; Letzteres bedeutet im einzelnen, dass x erstens selbst eine obere Schranke von S ist, und zweitens kleiner oder gleich jeder oberen Schranke von S (insgesamt also $S \sqsubseteq z$ gdw. $x \sqsubseteq z$).

Die Begriffe *untere Schranke* und *Infimum* werden in der offensichtlichen Weise dual hierzu definiert; insbesondere ist x Infimum von S ($x = \bigsqcap S$), wenn x die größte untere Schranke von S ist.

Das Poset $(X, \sqsubseteq)$ ist ein *vollständiger Verband*, wenn jede Teilmenge von X ein Supremum besitzt.

In einem vollständigen Verband $(X, \sqsubseteq)$ ist $\perp = \bigsqcup \emptyset$ das kleinste Element und $\top = \bigsqcup X$ das größte Element. Es gilt dann natürlich auch $\perp = \bigsqcap X$ und $\top = \bigsqcap \emptyset$. Allgemeiner bekommen wir aus der Existenz aller Suprema auch die Existenz aller Infima:

Satz 47. *In einem vollständigen Verband hat jede Teilmenge ein Infimum.*

(Natürlich gilt auch die duale Aussage, dass die Existenz aller Infima die Existenz aller Suprema impliziert, so dass also vollständige Verbände äquivalenterweise auch durch die Existenz aller Infima charakterisiert sind.)

Beweis. Sei $(X, \sqsubseteq)$ ein vollständiger Verband und $S \subseteq X$. Setze

$$\bar{S} = \{y \in X \mid y \sqsubseteq S\}.$$

Wir behaupten, dass $\bigsqcup \bar{S}$ Infimum von S ist. Nach Konstruktion ist $\bigsqcup \bar{S}$ größer gleich allen unteren Schranken von S ; zu zeigen ist also nur noch $\bigsqcup \bar{S} \sqsubseteq S$. Sei also $x \in S$. Dann $\bar{S} \sqsubseteq x$: Für $y \in \bar{S}$ gilt per Definition $y \sqsubseteq S$, insbesondere also $y \sqsubseteq x$. Da aber $\bigsqcup \bar{S}$ die kleinste obere Schranke von $\bar{S}$ ist, folgt wie verlangt $\bigsqcup \bar{S} \sqsubseteq x$. $\square$

Beispiel 48. Beispiele von vollständigen Verbänden:

1. $([0, 1], \leq)$ ist ein vollständiger Verband.
2. $(\mathcal{P}X, \subseteq)$ ist ein vollständiger Verband, wobei $\mathcal{P}X = \{Y \mid Y \subseteq X\}$ die Potenzmenge von X ist.
3. $(\mathbb{N}, \leq)$ ist *kein* vollständiger Verband, da $\bigsqcup \mathbb{N}$ nicht existiert. Dagegen ist $(\mathbb{N} \cup \{\infty\}, \leq)$ ein vollständiger Verband.

Beispiel 49. Sei $(X, \sqsubseteq)$ ein vollständiger Verband; dann ist $(X, \sqsubseteq)^{op} = (X, \sqsupseteq^{-1})$ ebenfalls ein vollständiger Verband.

Definition 50. Sei $(X, \sqsubseteq)$ eine geordnete Menge und $f: X \rightarrow X$ eine monotone Funktion. Wir definieren die folgenden Mengen:

$$\begin{aligned} \text{Fix}(f) &= \{x \in X \mid fx = x\} && \text{(Fixpunkte von } f) \\ \text{Fix}^\mu(f) &= \{x \in X \mid fx \sqsubseteq x\} && \text{(Präfixpunkte von } f) \\ \text{Fix}^\nu(f) &= \{x \in X \mid x \sqsubseteq fx\} && \text{(Postfixpunkte von } f) \end{aligned}$$

Wir bezeichnen mit μf das kleinste Element von $\text{Fix}^\mu(f)$ und mit νf das größte Element von $\text{Fix}^\nu(f)$ (also den kleinsten Präfixpunkt bzw. den größten Postfixpunkt von f), wenn solche Elemente jeweils existieren.

Lemma 51. Sei $(X, \sqsubseteq)$ eine geordnete Menge und $f: X \rightarrow X$ monoton, so dass μf (bzw. νf) existiert. Dann ist μf (bzw. νf) ein Fixpunkt von f .

Beweis. Per Dualität reicht es, die Aussage für μf zu zeigen.

Es bleibt nur noch zu zeigen, dass $f(\mu f) \sqsubseteq \mu f$. Nach Annahme gilt $f(\mu f) \sqsubseteq \mu f$; per Monotonie von f folgt $f(f(\mu f)) \sqsubseteq f(\mu f)$, d.h. $f(\mu f) \in \text{Fix}^\mu(f)$. Da μf kleinstes Element von $\text{Fix}^\mu(f)$ ist, folgt $\mu f \sqsubseteq f(\mu f)$. $\square$

Satz 52 (Knaster-Tarski). Sei $(X, \sqsubseteq)$ ein vollständiger Verband und $f: X \rightarrow X$ monoton. Dann gilt

$$\mu f = \bigcap \text{Fix}^\mu(f) \quad \nu f = \bigcup \text{Fix}^\nu(f).$$

Beweis. Per Dualität reicht es, die Behauptung für μf zu zeigen. Setze $m = \bigcap \text{Fix}^\mu f$. Damit liegt m unter allen Präfixpunkten von f ; nach Lemma 51 genügt also $m \in \text{Fix}^\mu(f)$, d.h. $f(m) \sqsubseteq m$. Das wiederum ist äquivalent zu $f(m) \sqsubseteq \text{Fix}^\mu(f)$. Sei also $x \in \text{Fix}^\mu(f)$. Per Definition von m gilt dann $m \sqsubseteq x$; per Monotonie von f folgt wie verlangt $f(m) \sqsubseteq f(x) \sqsubseteq x$, wobei die letzte Ungleichung einfach die definierende Eigenschaft von x als Präfixpunkt von f ist. $\square$

Satz 53 (Fixpunktsatz von Kleene für endliche Posets). *Sei $(X, \sqsubseteq)$ eine endliche geordnete Menge und $f: X \rightarrow X$ monoton. Dann bilden die $f^i(\perp)$ eine aufsteigende Folge, und es gibt ein n , ab dem die Folge bei $f^n(\perp)$ stationär wird (d.h. $f^k(\perp) = f^n(\perp)$ für alle $k \geq n$); dual bilden die $f^i(\top)$ eine absteigende Folge, und es gibt ein l , ab dem die Folge bei $f^l(\top)$ stationär wird. Es gilt dann $\mu f = f^n(\perp)$, $\nu f = f^l(\top)$.*

Dabei definieren wir f^n wie üblich induktiv durch

$$\begin{aligned} f^0(x) &= x \\ f^{i+1}(x) &= f(f^i(x)). \end{aligned}$$

Beweis. Per Dualität reicht es, die Behauptung für μf zu beweisen. Man zeigt induktiv, dass die $f^i(\perp)$ eine aufsteigende Folge bilden: Der Induktionsanfang $f^0(\perp) \sqsubseteq f^1(\perp)$ folgt sofort aus der Eigenschaft von $f^0(\perp) = \perp$ als kleinstem Element von X . Im Induktionsschritt folgt $f^{i+1}(\perp) = f(f^i(\perp)) \sqsubseteq f(f^{i+1}(\perp)) = f^{i+2}(\perp)$ per Monotonie von f aus der Induktionsannahme $f^i(\perp) \sqsubseteq f^{i+1}(\perp)$.

Da X endlich ist, muss es daher ein n mit $m := f^n(\perp) = f^{n+1}(\perp)$ geben, und dann bleibt die Folge ab dieser Stelle stationär im Sinne des Satzes. Damit ist m insbesondere ein (Prä-)fixpunkt von f . Sei nun $x \in \text{Fix}^\mu(f)$; zu zeigen bleibt $m \sqsubseteq x$. Allgemeiner zeigen wir per Induktion über i , dass $f^i(\perp) \sqsubseteq x$ für alle i . Der Induktionsanfang ist klar, da $f^0(\perp) = \perp$. Im Induktionsschritt erhalten wir aus der Induktionsannahme $f^i(\perp) \sqsubseteq x$ per Monotonie von f

$$f^{i+1}(\perp) = f(f^i(\perp)) \sqsubseteq f(x) \sqsubseteq x,$$

wiederum unter Ausnutzung der Annahme $x \in \text{Fix}^\mu(f)$ im letzten Schritt. $\square$

Wir charakterisieren nun Bisimulationen als Postfixpunkte und (damit) Bisimilarität als größten (Post-)Fixpunkt eines monotonen Operators auf Relationen. Sei $T = (\mathcal{Q}, \mathcal{A}, \{\xrightarrow{\alpha} \mid \alpha \in \mathcal{A}\})$ ein LTS und $R \subseteq \mathcal{Q} \times \mathcal{Q}$ eine binäre Relation auf $\mathcal{Q}$. Wir definieren $\mathcal{F}(R) \subseteq \mathcal{Q} \times \mathcal{Q}$ wie folgt:

$(s, t) \in \mathcal{F}(R)$ und $s, t \in \mathcal{Q}$ genau dann, wenn

- $s \xrightarrow{\alpha} s'$ impliziert, dass t' existiert, so dass $t \xrightarrow{\alpha} t'$ und $(s', t') \in R$ (*Forth*);
- $t \xrightarrow{\alpha} t'$ impliziert, dass s' existiert, so dass $s \xrightarrow{\alpha} s'$ und $(s', t') \in R$ (*Back*).

Dann gilt:

- Eine Relation $R \subseteq \mathcal{Q} \times \mathcal{Q}$ ist genau dann eine Bisimulation, wenn $R \subseteq \mathcal{F}(R)$, d.h. wenn R Postfixpunkt von $\mathcal{F}$ ist; und somit

$$\bullet \sim = \sqcup \{R \in \mathcal{Q} \times \mathcal{Q} \mid R \subseteq \mathcal{F}(R)\} = \nu \mathcal{F}$$

Aus dem Satz von Kleene ergibt sich nun unmittelbar ein Algorithmus zur Berechnung der Bisimilaritätsrelation auf einem gegebenen LTS $T = (\mathcal{Q}, \mathcal{A}, \rightarrow)$:

1. Initialisiere R auf $R = \mathcal{Q} \times \mathcal{Q}$
2. Wenn $R = \mathcal{F}(R)$, terminiere und gib R zurück.
3. Setze $R := \mathcal{F}(R)$.
4. Gehe zu Schritt 2.

Was ist die Laufzeit dieses Algorithmus?

Wichtig bei der Analyse ist unter anderem folgende Feststellung:

Lemma 54. *Im obigen Algorithmus ist R zu jedem Zeitpunkt eine Äquivalenzrelation.*

Beweis. Wir müssen zeigen, dass das Funktional $\mathcal{F}$ Äquivalenzrelationen bewahrt. Sei also R eine Äquivalenzrelation; wir zeigen, dass $\mathcal{F}(R)$ eine Äquivalenzrelation ist.

Reflexivität: Sei $s \in \mathcal{Q}$. Zu zeigen ist $(s, s) \in \mathcal{F}(R)$. Sei also $s \xrightarrow{a} s'$; wir wählen auf der „anderen“ Seite ebenfalls die Transition $s \xrightarrow{a} s'$, und da R eine Äquivalenzrelation ist, gilt wie verlangt $s'Rs'$. Die Back-Bedingung zeigt man symmetrisch.

Symmetrie: Sei $(s, t) \in \mathcal{F}(R)$; zu zeigen ist $(t, s) \in \mathcal{F}(R)$. Wir beweisen o.E. nur die Forth-Bedingung. Sei also $t \xrightarrow{a} t'$. Dann existiert gemäß der Back-Bedingung für $(s, t) \in \mathcal{F}(R)$ ein s' mit $s \xrightarrow{a} s'$ und $s'Rt'$; da R symmetrisch ist, gilt wie verlangt $t'Rs'$.

Transitivität: Seien $(s, t), (t, u) \in \mathcal{F}(R)$; zu zeigen ist $(s, u) \in \mathcal{F}(R)$. Wir zeigen o.E. nur die Forth-Bedingung. Sei also $s \xrightarrow{a} s'$. Da $(s, t) \in \mathcal{F}(R)$, existiert $t \xrightarrow{a} t'$, so dass $s'Rt'$. Da $(t, u) \in \mathcal{F}(R)$, existiert $u \xrightarrow{a} u'$, so dass $t'Ru'$. Da R transitiv ist, folgt wie verlangt $s'Ru'$. $\square$

Dementsprechend heißen Algorithmen des obigen Typs *Partition-Refinement-Algorithmen*. Wie viele Iterationen benötigt also der obige Algorithmus schlimmstenfalls?

2.7 Der Paige-Tarjan-Algorithmus

Allgemein arbeitet man bei Partition-Refinement-Algorithmen statt mit Äquivalenzrelationen äquivalenterweise mit *Partitionen* $P = \{B_1, \dots, B_b\}$, d.h.

disjunkten Zerlegungen $\mathcal{Q} = B_1 \cup \dots \cup B_n$ der Zustandsmenge in Äquivalenzklassen B_i , die man hier kurz als *Blöcke* bezeichnet. Der Einfachheit halber nehmen wir an, es gebe nur eine Aktion, die wir dann nicht mehr erwähnen. Wir schreiben für $S \subseteq \mathcal{Q}$ im folgenden

$$\rightarrow [S] = \{s' \mid \exists s \in S. s \rightarrow s'\}$$

und

$$\leftarrow [S] = \{s \mid \exists s' \in S. s \rightarrow s'\}$$

für die Menge der Nachfolger bzw. Vorgänger von Elementen von S . Wir nennen eine Partition P wie oben *stabil* bezüglich $S \subseteq \mathcal{Q}$, wenn für alle $B \in P$ entweder

$$B \subseteq \leftarrow [S] \quad \text{oder} \quad B \cap \leftarrow [S] = \emptyset$$

gilt, wenn also entweder alle oder keine Zustände in B Transitionen nach S haben. Eine Partition P heißt *stabil* bezüglich einer Partition R , wenn P stabil bezüglich aller Blöcke von R ist, und *stabil* schlechthin, wenn P stabil bezüglich P ist. Man überzeugt sich leicht, dass eine Partition genau dann stabil ist, wenn die zugehörige Äquivalenzrelation eine Bisimulation ist. Man arbeitet bei Partition-Refinement-Algorithmen typischerweise mit einer *Initialpartition* P_0 und sucht dann die *größte stabile Verfeinerung* der Initialpartition, wobei eine Partition P *feiner* ist als eine Partition R , wenn jeder Block von R Vereinigung von Blöcken von P ist (was in Begriffen von Äquivalenzrelationen bedeutet, dass die von P repräsentierte Äquivalenzrelation in der von R repräsentierten enthalten ist). Wir interessieren uns hier nur für die Berechnung von Bisimilarität und unterstellen deswegen, dass P_0 stets aus (maximal) zwei Blöcken besteht, den Deadlocks und allen anderen Zuständen (einer dieser Blöcke kann natürlich leer sein und entfällt dann).

Die Definition von Stabilität motiviert eine Funktion *split*: Gegeben eine Partition P und $S \subseteq \mathcal{Q}$ ist $\text{split}(S, P)$ die Partition, die man aus P erhält, indem man jeden Block $B \in P$ in $B' = B \cap \leftarrow [S]$ und $B'' = B \setminus B'$ aufspaltet. Wir nennen S einen *Spalter* von P , wenn $\text{split}(S, P) \neq P$, wenn also P nicht stabil bezüglich S ist. In diesen Begriffen besteht eine Berechnung von $\mathcal{F}(R)$ wie im vorigen Abschnitt definiert im Wesentlichen in einer sukzessiven Anwendung von *split* auf alle Blöcke von R .

Man kann den Algorithmus dann liberalisieren, indem man *split* wiederholt auf beliebige Spalter anwendet. Wenn man als Spalter stets Blöcke der aktuellen Partition verwendet, erreicht man so bei geschickter Datenehaltung und guter Implementierung von *split* eine Laufzeit von $\mathcal{O}(m)$ pro Anwendung von *split* und damit eine Gesamtlaufzeit von $\mathcal{O}(mn)$, wobei m die Anzahl Transitionen und n die Anzahl Zustände des Systems ist. Der

Paige-Tarjan-Algorithmus geht noch etwas subtiler vor und erreicht damit eine Gesamtlaufzeit $\mathcal{O}(m \log n)$.

Der Algorithmus arbeitet mit *zwei* Partitionen P und R und hält die folgende Invariante aufrecht:

P ist feiner als R und stabil bezüglich allen Blöcken von R . $(*)$

Ein Block $B \in R$ heißt *einfach*, wenn $B \in P$, und sonst *zusammengesetzt*. Der Algorithmus verläuft im Groben wie folgt:

1. Initialisiere P auf P_0 und R auf $\{Q\}$.
2. Wiederhole folgenden Schritt *Refine* so lange, bis $P = R$:
 - (a) Wähle einen zusammengesetzten Block $S \in R$.
 - (b) Wähle $B \in P$ mit $B \subseteq S$ und $|B| \leq |S|/2$.
 - (c) Ersetze in R S durch B und $S \setminus B$.
 - (d) Setze $P := \text{split}(S \setminus B, \text{split}(B, P))$.

Im letzten Schritt von *Refine* werden effektiv alle Blöcke $D \in P$ in drei Teile (sofern jeweils nichtleer) unterteilt, nämlich

$$D \cap \leftarrow[B] \cap \leftarrow[S \setminus B] \quad D \cap \leftarrow[B] \setminus \leftarrow[S \setminus B] \quad D \setminus \leftarrow[B]$$

(der dritte Teil würde nach Definition von *split* in die Mengen $(D \setminus \leftarrow[B]) \cap \leftarrow[S \setminus B]$ und $(D \setminus \leftarrow[B]) \setminus \leftarrow[S \setminus B]$ zerfallen, von denen aber, weil P stabil bezüglich S ist, entweder die erste leer ist, nämlich wenn $D \cap \leftarrow[S] = \emptyset$, oder die zweite, nämlich wenn $D \subseteq \leftarrow[S]$). Man beachte, dass die mittlere der drei Mengen gleich

$$D \cap (\leftarrow[B] \setminus \leftarrow[S \setminus B])$$

ist.

Wir zeigen zunächst, dass der skizzierte Algorithmus korrekt ist:

Lemma 55. *Der Paige-Tarjan-Algorithmus berechnet die Bisimilaritätsrelation auf Q (als Partition).*

Beweis. Wir zeigen, dass

P vergrößert Bisimilarität $(+)$

eine (weitere) Invariante des Algorithmus ist; mit anderen Worten, P trennt nur nicht-bisimilare Zustände. Nach Initialisierung gilt die Invariante für P nach unserer festen Wahl von P_0 ; wir müssen zeigen, dass *Refine* die Invariante

erhält: Sei $s \sim t$, und sei $D \in P$ der Block mit $s, t \in D$; wir müssen zeigen, dass s und t nach Aufspaltung von D im gleichen Block verbleiben. Dazu reicht es zu zeigen, dass für jeden Block $B \in P$ gilt

$$s \in \leftarrow[B] \iff t \in \leftarrow[B].$$

Per Symmetrie reicht eine Implikation. Sei also $s \in \leftarrow[B]$, d.h. wir haben $s \rightarrow s'$ mit $s' \in B$. Wir erhalten t' mit

$$\begin{array}{ccc} s & \xrightarrow{\sim} & t \\ \downarrow & & \downarrow \\ s' & \xrightarrow{\sim} & t' \end{array}$$

Da P Bisimilarität vergrößert, folgt aus $s' \in B$, dass auch $t' \in B$ und damit wie verlangt $t \in \leftarrow[B]$.

Ferner ist klar, dass der Algorithmus tatsächlich den nichttrivialen Teil der Invarianten (*)

P ist stabil bezüglich R

einhält: Nach Initialisierung gilt die Invariante nach unserer festen Wahl von P_0 , und die beiden *split*-Anwendungen am Ende eines *Refine*-Schritts stellen gerade sicher, dass P stabil bezüglich des neuen R wird.

Da nach Terminierung P stabil ist (nämlich stabil bezüglich R , was jetzt aber gleich P ist), induziert P eine Bisimulation, verfeinert also Bisimilarität. Nach der obigen Invarianten (+) vergrößert P außerdem Bisimilarität, fällt also insgesamt gerade mit Bisimilarität zusammen. $\square$

Die Herausforderung liegt dann in den Implementierungsdetails von *Refine*. Der Algorithmus verwendet die folgenden Datenstrukturen:

- *Zustände*
- *Transitionen*
- *Blöcke* in R
- *Blöcke* in P
- *Zusammengesetzte Blöcke*
- *Zähler* $\text{count}(s, S) = |\rightarrow[\{s\}] \cap S|$ für $s \in \mathcal{Q}$, $S \in R$ (d.h. $\text{count}(s, S)$ gibt die Anzahl in S liegender Nachfolger von s an)

Diese Daten sind wie folgt verlinkt:

- Transitionen $s \rightarrow s'$ auf s ;
- Zustände s' auf eine Liste aller Transitionen $s \rightarrow s'$;
- Zustände s auf temporäre Counter $count(s, -)$;
- Block $B \in P$ hat Integer $|B|$ und zeigt auf seine Elemente als doppelt verkettete Liste (daher Entfernung von Elementen in $\mathcal{O}(1)$);
- Zustand s auf Block $B \in P$ mit $s \in B$;
- Block $S \in R$ auf Blöcke $B \in P$ mit $B \subseteq S$ als doppelt verkettete Liste;
- Block $B \in P$ auf Block $S \in R$ mit $B \subseteq S$;
- Transition $s \rightarrow s'$ auf $count(s, S)$ für $s' \in S \in R$;
- die zusammengesetzten Blöcke bilden eine Menge C .

Diese Datenstrukturen werden beim Einlesen des Eingabe-LTS geeignet initialisiert, insbesondere muss also die Größe der Blöcke in P ermittelt werden.

Die Implementierung von *Refine* verfährt dann im einzelnen wie folgt:

1. (*Spalter auswählen*) Entferne einen Block S aus C . Vergleiche die ersten beiden P -Blöcke in S und wähle den kleineren als B .
2. (*R aktualisieren*) Entferne B aus S und erzeuge einen neuen Block S' in R , der nur aus B besteht. Wenn S noch zusammengesetzt ist, füge es wieder zu C hinzu.
3. (*Berechne $\leftarrow[B]$*) Kopiere die Elemente von B in einen temporären Block B' . Initialisiere $\leftarrow[B]$ auf leer. Gehe durch $s' \in B$ und eingehende Kanten $s \rightarrow s'$ (verlinkt von s' aus!) und füge die aufgefundenen s zu $\leftarrow[B]$ hinzu und markiere sie zur Vermeidung von Duplikaten. Berechne gleichzeitig einen neuen Zähler $count(s, B)$ für jedes s , und verlinke diesen mit s .
4. (*Berechne $P' = split(B, P)$*) Spalte jeden Block $D \in P$ mit $D \cap \leftarrow[B] \neq \emptyset$ in $D_1 = D \cap \leftarrow[B]$ und $D_2 = D \setminus D_1$. Gehe dazu durch die Elemente $s \in \leftarrow[B]$. Finde für jedes s das $D \in P$ mit $s \in D$, erzeuge zu D einen neuen als zu D gehörig markierten und mit D entsprechend verlinkten Block D' in S (wenn nicht schon vorher geschehen), und verschiebe s von D nach D' . (D.h. D' wird in den obigen Bezeichnungen am Ende D_1 , während die Elemente von D_2 in D verbleiben.) Dabei muss über die Anzahl Elemente der neu entstehenden Blöcke buchgeführt werden.

Erzeuge währenddessen eine Liste der aufgespaltenen Blöcke D , d.h. derjenigen Blöcke, für die ein zugehöriges D' erzeugt worden ist. Verarbeite diese Liste anschließend wie folgt:

- (a) Hebe die Zugehörigkeit von D' zu D auf (so dass im nächsten Durchgang alle Blöcke als noch nicht gespalten gelten).
 - (b) Entferne D , wenn jetzt $D = \emptyset$.
 - (c) Sonst: Füge den Block $S' \in R$ mit $D \subseteq S'$ zu C hinzu, wenn S' jetzt genau zwei Blöcke enthält (nämlich D und D'). (Dann war nämlich S' vorher nicht zusammengesetzt, jetzt aber schon.)
5. (*Berechne $\leftarrow[B] \setminus \leftarrow[S \setminus B]$*) Gehe wie in Schritt 3 durch die Transitionen $s \rightarrow s'$ mit $s' \in B'$ (wir verwenden jetzt die Kopie B' von B , weil B im vorigen Schritt eventuell schon aufgespalten worden sein könnte). Wenn $\text{count}(s, B) = \text{count}(s, S)$ (verlinkt mit s bzw. $s \rightarrow s'$), füge s zu $\leftarrow[B] \setminus \leftarrow[S \setminus B]$ hinzu (wenn nicht schon drin, s.o.).
6. (*Berechne $\text{split}(S \setminus B, P')$*) Wie für B , aber mit $\leftarrow[B] \setminus \leftarrow[S \setminus B]$ statt $\leftarrow[B]$.
7. (*Aktualisiere Zähler*) Gehe wie gehabt durch die Transitionen $s \rightarrow s'$ mit $s' \in B'$. Dekrementiere $\text{count}(s, S)$ (verlinkt mit $s \rightarrow s'$). Wenn dies 0 wird, lösche den Zähler $\text{count}(s, S)$. Ersetze den an $s \rightarrow s'$ hängenden Zähler $\text{count}(s, S)$ (in jedem Fall) durch $\text{count}(s, B)$ (verlinkt mit s). Lösche abschließend B' , hebe alle Markierungen auf (wie?).

(Das eventuelle Löschen von $\text{count}(s, S)$ im letzten Schritt dient gewissermaßen nur dem Aufräumen; wenn $\text{count}(s, S)$ nämlich 0 wird, ist der Zähler $\text{count}(s, S)$ von keiner Transition $s \rightarrow s'$ mehr verlinkt, da alle Nachfolger von s im Block B liegen, der ja zu diesem Zeitpunkt bereits aus S entfernt ist. Die Endwerte der Zähler $\text{count}(s, S)$ sind natürlich vorhersagbar, sie betragen am Ende des Gesamtschritts jeweils $\text{count}(s, B)$ weniger als vorher, aber man müsste bei direkter Subtraktion zur Vermeidung von wiederholter Subtraktion die schon abgehandelten s wieder aufwändig markieren. Die schrittweise Dekrementierung bedeutet dagegen insofern keinen Effizienzverlust, als ja ohnehin alle Transitionen $s \rightarrow s'$ mit $s' \in B'$ einzeln durchgegangen werden müssen.)

Die Laufzeitanalyse gestaltet sich dann wie folgt. Wir schreiben für $A \subseteq Q$

$$\text{Trans}(A) = \{(s, s') \mid s \rightarrow s', s' \in A\},$$

d.h. $\text{Trans}(A)$ ist die Menge der Transitionen nach A .

Lemma 56. *Ein Refine-Schritt mit Spalter B benötigt Zeit*

$$\mathcal{O}(|B| + \sum_{s' \in B} |\text{Trans}(\{s'\})|).$$

(Die Summe in der obigen Formel ist gerade $|\text{Trans}(B)|$.)

Beweis. Wir schätzen die Zeit für die einzelnen Phasen von *Refine* ab; es reicht natürlich, dass jede Phase die angegebene asymptotische Schranke einhält.

1. (*Spalter auswählen*) $\mathcal{O}(1)$.
2. (*R aktualisieren*) $\mathcal{O}(1)$.
3. (*Berechne $\leftarrow[B]$*) Diese Phase enthält eine Schleife über B , mit Schleifenkörper in $\mathcal{O}(1)$, und eine Schleife über $s' \in B$ und $s \rightarrow s'$, mit Schleifenkörper in $\mathcal{O}(1)$, was gerade Zeit $\mathcal{O}(|B| + \sum_{s' \in B} |\text{Trans}(\{s'\})|)$ benötigt.
4. (*Berechne $P' = \text{split}(B, P)$*) Diese Phase enthält eine Schleife über $s \in \leftarrow[B]$ mit Schleifenkörper in $\mathcal{O}(1)$, was also Zeit $\mathcal{O}(|\leftarrow[B]|) \leq \mathcal{O}(\sum_{s' \in B} |\text{Trans}(\{s'\})|)$ benötigt. Die zweite Schleife über die Liste der aufgespaltenen Blöcke hat höchstens so viele Durchläufe wie die erste.
5. (*Berechne $\leftarrow[B] \setminus \leftarrow[S \setminus B]$*) Wie Phase 3. (Diese harmlose Bemerkung ist einer der Schlüssel zur Effizienz des Algorithmus: Dank des Tricks mit den Zählern muss man hier eben nicht über $S \setminus B$, sondern nur wieder über B bzw. seine Kopie B' laufen.)
6. (*Berechne $\text{split}(S \setminus B, P')$*) Wie Phase 4 (mit $\leftarrow[B] \setminus \leftarrow[S \setminus B]$ statt $\leftarrow[B]$).
7. (*Aktualisiere Zähler*) Diese Phase enthält eine Schleife über $\leftarrow[B]$, mit Schleifenkörper in $\mathcal{O}(1)$, d.h. die Laufzeit ist wie in Phase 4. Das abschließende Löschen der Markierungen hält dieselbe Zeitschranke ein.

□

Satz 57. *Der Paige/Tarjan-Algorithmus läuft in Zeit $\mathcal{O}((m+n) \log n)$, wobei m die Anzahl Transitionen des Systems bezeichnet und n die Anzahl Zustände.*

Beweis. In einem Lauf des Algorithmus taucht ein gegebener Zustand s' in höchstens $\log(n) + 1$ als Spalter verwendeten Blöcken auf, da jeder s' enthaltende Spalter immer höchstens halb so groß wie der vorige solche Spalter ist. (Denn: Wenn s' in einem Spalter B vorkommt, wird B anschließend ein Block in R ; sobald also s erneut in einem Spalter B' vorkommt, ist der entsprechende zusammengesetzte Block $S \supseteq B'$ eine Teilmenge von B , und B' höchstens halb so groß wie S .) Wenn $B_1, \dots, B_k$ die vom Algorithmus verwendeten Spalter sind, ist die Gesamtlaufzeit nach dem obigen Lemma also, wenn c die Konstante in der Laufzeitabschätzung des *Refine*-Schritts bezeichnet, höchstens

$$\begin{aligned}
& \sum_{i=1}^k c \cdot (|B_i| + |\text{Trans}(B_i)|) \\
&= c \cdot \sum_{i=1}^k \sum_{s' \in B_i} (1 + |\text{Trans}(\{s'\})|) \\
&= c \cdot \sum_{s' \in \mathcal{Q}} \sum_{i | s' \in B_i} (1 + |\text{Trans}(\{s'\})|) \\
&\leq c \cdot \sum_{s' \in \mathcal{Q}} (1 + \log n) (1 + |\text{Trans}(\{s'\})|) \\
&= c \cdot (1 + \log n) \cdot \sum_{s' \in \mathcal{Q}} (1 + |\text{Trans}(\{s'\})|) \\
&= c \cdot (1 + \log n) \cdot (n + m) \\
&= \mathcal{O}((m + n) \log n)
\end{aligned}$$

□

Wir bemerken noch, dass außer in wenig sinnvollen Randfällen tendenziell $m \geq n$ gilt, so dass die Schranke üblicherweise leicht verkürzt als $\mathcal{O}(m \log n)$ angegeben wird.

3 Hennessy-Milner-Logik

Hennessy-Milner-Logik (HML) macht Aussagen über die mögliche zukünftige Entwicklung eines LTS, vom aktuellen Zustand aus gesehen. Sie ermöglicht z.B. die Formalisierung von Aussagen wie

- CS möchte (momentan) keinen Tee trinken;
- CS möchte (momentan) weder Tee noch Kaffee trinken;
- CS trinkt gleich Kaffee und veröffentlicht dann.

Hennessy-Milner-Logik ist eine Form von *Modallogik*, und erweitert als solche die übliche Aussagenlogik um *Modaloperatoren*:

Definition 58 (HML, Syntax). Formeln in HML über der Menge $\mathcal{A}$ von Aktionen sind durch die Grammatik

$$\phi, \psi ::= \perp \mid \neg\phi \mid \phi \wedge \psi \mid [a]\psi \quad (a \in \mathcal{A})$$

definiert, und *Negationsnormalformen* durch

$$\phi, \psi ::= \top \mid \perp \mid \phi \wedge \psi \mid \phi \vee \psi \mid \langle a \rangle \phi \mid [a]\psi \quad (a \in \mathcal{A}).$$

In Formeln werden $\top, \vee, \rightarrow$ und $\leftrightarrow$ wie üblich in $\perp, \neg, \wedge$ dekodiert, und $\langle a \rangle \phi$ als $\neg[a]\neg\phi$ (dagegen sind $\top, \vee$ und $\langle a \rangle$ in Negationsnormalformen First-Class Citizens). Für $A = \{a_1, \dots, a_n\}$ vereinbaren wir ferner die Abkürzungen

$$\begin{aligned} \langle A \rangle \phi &= \langle a_1 \rangle \phi \vee \dots \vee \langle a_n \rangle \phi \\ [A] \phi &= [a_1] \phi \wedge \dots \wedge [a_n] \phi \end{aligned}$$

(mit Randfällen $\langle \emptyset \rangle \phi = \perp, [\emptyset] \phi = \top$).

Definition 59 (HML, Semantik). Sei $T = (\mathcal{Q}, \mathcal{A}, (\xrightarrow{a})_{a \in \mathcal{A}})$ ein LTS. Wir definieren *Erfülltheit* $\models$ von Formeln ϕ in Zuständen $s \in \mathcal{Q}$ (mit $s \models \phi$ gelesen als ‘ s erfüllt ϕ ’) rekursiv durch

$$\begin{aligned} s &\not\models \perp \\ s &\models \phi \wedge \psi && \text{gdw. } s \models \phi \text{ und } s \models \psi \\ s &\models [a]\phi && \text{gdw. } s' \models \phi \text{ für alle } s' \text{ mit } s \xrightarrow{a} s'. \end{aligned}$$

Mit der angegebenen Kodierung folgt dann z.B. ferner

$$\begin{aligned} s &\models \top && \text{stets} \\ s &\models \phi \vee \psi && \text{gdw. } s \models \phi \text{ oder } s \models \psi \\ s &\models \langle a \rangle \phi && \text{gdw. } s' \text{ existiert mit } s \xrightarrow{a} s' \text{ und } s' \models \phi. \end{aligned}$$

Wir definieren die *Extension* $\llbracket \phi \rrbracket$ einer Formel ϕ in T als

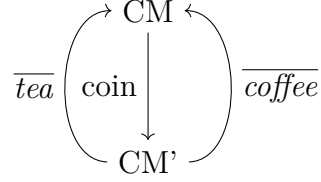
$$\llbracket \phi \rrbracket = \{s \in \mathcal{Q} \mid s \models \phi\}.$$

Damit gilt

$$\begin{aligned} \llbracket \perp \rrbracket &= \emptyset \\ \llbracket \neg\phi \rrbracket &= \mathcal{Q} \setminus \llbracket \phi \rrbracket \\ \llbracket \phi \wedge \psi \rrbracket &= \llbracket \phi \rrbracket \cap \llbracket \psi \rrbracket \\ \llbracket \langle a \rangle \phi \rrbracket &= \{s \in \mathcal{Q} \mid D_a(s) \cap \llbracket \phi \rrbracket \neq \emptyset\} \end{aligned}$$

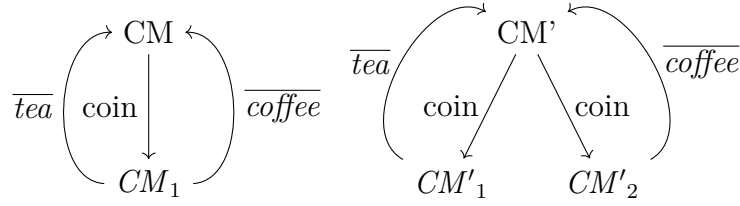
wobei $D_a(s) = \{s' \mid s \xrightarrow{a} s'\}$.

Beispiel 60.



$$\begin{aligned}
s \models [coin]\top &\iff (\forall s'. s \xrightarrow{coin} s' \Rightarrow s' \models \top) \iff s \in \{CM, CM'\} \\
s \models \langle coin \rangle \top &\iff (\exists s'. s \xrightarrow{coin} s') \iff s \in \{CM\} \\
s \models [coin]\langle \overline{coffee} \rangle \top &\iff (\forall s'. s \xrightarrow{coin} s' \Rightarrow s' \models \langle \overline{coffee} \rangle \top) \\
&\iff (\forall s'. s \xrightarrow{coin} s' \Rightarrow s' \in \{CM'\}) \iff s \in \{CM, CM'\}
\end{aligned}$$

Beispiel 61.



Es gibt eine Formel ϕ , so dass $CM \models \phi$, aber $CM' \not\models \phi$, und zwar

$$\phi = [coin]\langle \overline{coffee} \rangle \top,$$

weil $\llbracket \phi \rrbracket = \{CM, CM_1, CM'_1, CM'_2\} \not\ni CM'$.

Wir klären noch kurz, warum wir in der Darstellung von HML via Negationsnormalformen ohne Negation auskommen:

Definition 62. Sei ϕ eine Negationsnormalform in HML. Wir definieren eine Negationsnormalform $\overline{\phi}$ induktiv durch

$$\begin{aligned}
\overline{\top} &= \perp \\
\overline{\perp} &= \top \\
\overline{\phi \wedge \psi} &= \overline{\phi} \vee \overline{\psi} \\
\overline{\phi \vee \psi} &= \overline{\phi} \wedge \overline{\psi} \\
\overline{[a]\phi} &= \langle a \rangle \overline{\phi} \\
\overline{\langle a \rangle \phi} &= [a] \overline{\phi}
\end{aligned}$$

Lemma 63. Für jede Negationsnormalform ϕ in HML sind $\neg\phi$ und $\bar{\phi}$ logisch äquivalent.

Beweis. Induktion über ϕ . □

Wir erinnern daran, dass ein LTS $(\mathcal{Q}, \mathcal{A}, \rightarrow)$ endlich verzweigend ist, wenn $D_\alpha(s) = \{s' \mid s \xrightarrow{\alpha} s'\}$ für alle $s \in \mathcal{Q}$ und alle $\alpha \in \mathcal{A}$ endlich ist.

Beispiel 64 (Unendlich verzweigendes LTS). $P = a. \emptyset \mid (P + a.\emptyset)$ ist nicht endlich verzweigend. (Wie sieht das LTS aus?)

Definition 65. Zustände s, t sind (in HML) logisch ununterscheidbar, wenn s und t dieselben HML-Formeln erfüllen, d.h. wenn die Äquivalenz

$$P \models \phi \iff Q \models \phi$$

für alle HML-Formeln ϕ gilt.

Satz 66 (Bisimulationsinvarianz von HML). Bisimilare Zustände sind in HML logisch ununterscheidbar.

Beweis. Seien s, t Zustände in LTS, so dass $s \sim t$, und sei ϕ eine HML-Formel. Wir zeigen per Induktion über ϕ , dass $s \models \phi$ genau dann, wenn $t \models \phi$. Die Booleschen Fälle $\perp, \neg, \wedge$ sind, wie meistens in Invarianzbeweisen, leicht; zur Illustration führen wir den Schritt für Negation hier explizit durch:

$$\begin{aligned} s \models \neg\phi & \\ \iff s \not\models \phi & \\ \iff t \not\models \phi & \quad \text{(IV)} \\ \iff t \models \neg\phi. & \end{aligned}$$

Der einzig interessante Fall ist der für die modale Box; per Symmetrie von $\sim$ reicht es, eine der beiden Implikationen zu zeigen. Sei also $s \models [a]\phi$; zu zeigen ist $t \models [a]\phi$. Sei also $t \xrightarrow{a} t'$; dann existiert s' mit $s \xrightarrow{a} s'$ und $s' \sim t'$. Da $s \models [a]\phi$, gilt $s' \models \phi$, nach IV also auch $t' \models \phi$, wie verlangt. □

In endlich verzweigenden Systemen gilt auch die Umkehrung:

Satz 67 (Hennessy/Milner). In HML logisch ununterscheidbare Zustände in endlich verzweigenden LTS sind bisimilar.

Beweis. Wir zeigen, dass die Relation

$$R = \{(s, t) \mid s, t \text{ logisch ununterscheidbar}\}$$

eine Bisimulation ist, d.h. dass sich jedes partielle Viereck

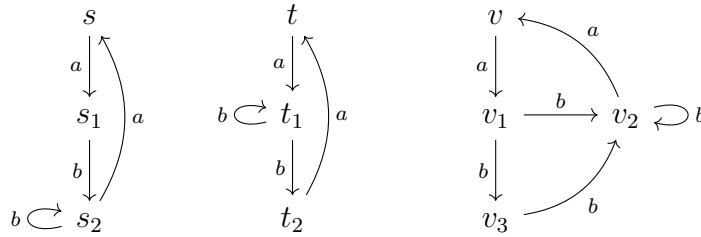
$$\begin{array}{ccc} s & \xrightarrow{R} & t \\ \downarrow \alpha & & \downarrow \alpha \\ s' & \xrightarrow{R} & t' \end{array}$$

wie angedeutet vervollständigen lässt. Wir zeigen die Existenz von t' per Widerspruch. Wenn $t_1, \dots, t_n$ die (endlich vielen!) α -Nachfolger von t sind, nehmen wir also an, dass für jedes $i = 1, \dots, n$ eine Formel ϕ_i existiert, die s' und t_i unterscheidet; da die Logik Negation hat, können wir dann annehmen, dass $s \models \neg\phi_i$ und $t_i \models \phi_i$. Dann gilt $t_i \models \phi_1 \vee \dots \vee \phi_n$ für alle i , aber $s' \not\models \phi_1 \vee \dots \vee \phi_n$, und damit

$$\begin{aligned} t &\models [\alpha](\phi_1 \vee \dots \vee \phi_n) \\ s &\not\models [\alpha](\phi_1 \vee \dots \vee \phi_n), \end{aligned}$$

im Widerspruch zu sRt . □

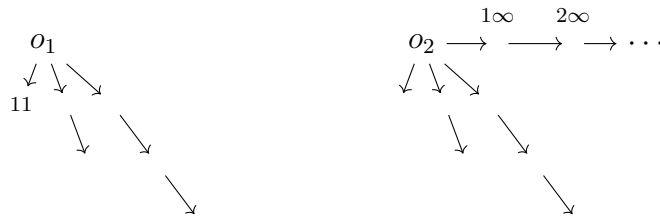
Beispiel 68.



Hier gibt es für s, t, v paarweise unterscheidende Formeln. (Welche?)

In unendlich verzweigende LTS müssen dagegen logisch ununterscheidbare Zustände im allgemeinen nicht bisimilar sein:

Beispiel 69. Man betrachte die beiden Kripkemodelle



Die beiden Wurzelzustände o_1, o_2 sind *nicht* bisimilar. Dies sieht man leicht z.B. mittels des Bisimulationsspiels (wie?). Die Zustände o_1 und o_2 erfüllen aber die gleichen Formeln, da für jede modale Formel mit Schachtelungstiefe n der Modaloperatoren nur Zustände bis zur Entfernung n von der Wurzel relevant sind.

4 Der modale μ -Kalkül

Motivation:

Man betrachte das LTS

$$a \hookrightarrow p$$

$$a \hookrightarrow q \xrightarrow{a} r$$

Die Wurzelzustände p und q sind offenbar nicht bisimilar, und wir finden unmittelbar eine unterscheidende HML-Formel:

$$\begin{aligned} p &\models [a]\langle a \rangle \top \\ q &\not\models [a]\langle a \rangle \top \end{aligned}$$

Das folgende System ist bis auf die Anzahl a -Aktionen, die zum Deadlock führen, sehr ähnlich wie q oben:

$$a \hookrightarrow q_1^n \xrightarrow{a} q_2^n \xrightarrow{a} \dots \xrightarrow{a} q_n^n$$

Wir unterscheiden es in HML wie folgt von p :

$$\begin{aligned} p &\models [a]^n \langle a \rangle \top \\ q_1^n &\not\models [a]^n \langle a \rangle \top \end{aligned}$$

Um Unterschiede wie die zwischen p und q_1^n *allgemein* zu fassen, brauchen wir einen Begriff von *Invarianz* bzw. von *Erreichbarkeit* von Eigenschaften: Wir schreiben *Inv* für *invariant* und *Pos* für *possible*, und setzen informell

$$\begin{aligned} \text{Inv}(\langle a \rangle \top) &= \langle a \rangle \top \wedge [a]\langle a \rangle \top \wedge [a][a]\langle a \rangle \top \wedge \dots \\ \text{Pos}([a]\perp) &= [a]\perp \vee \langle a \rangle[a]\perp \vee \langle a \rangle \langle a \rangle[a]\perp \vee \dots \end{aligned}$$

Solche unendlichen Konjunktionen und Disjunktionen sind aber nicht Teil der logischen Syntax, und es wäre zwar theoretisch möglich, aber praktisch eher ungünstig, sie einzuführen. Stattdessen beobachten wir, dass *Inv*(ϕ) die folgende rekursive Gleichung erfüllt:

$$\text{Inv}(\phi) = \phi \wedge [\mathcal{A}]\text{Inv}(\phi)$$

Ähnlich erfüllt *Pos*(ϕ) die rekursive Gleichung

$$\text{Pos}(\phi) = \phi \vee \langle \mathcal{A} \rangle \text{Pos}(\phi).$$

Wir können diese Gleichungen als Definitionen ansehen, aber auf sehr unterschiedliche Weise. Beide Gleichungen sind Fixpunktgleichungen; die erste besagt z.B., dass $Inv(\phi)$ grob gesagt Fixpunkt einer Abbildung der Art $X \mapsto (\phi \wedge [\mathcal{A}]X)$ ist (formal lässt sich das im Moment noch nicht fassen, da wir keine Notation für Mengen von Zuständen in Formeln haben), ähnlich in der zweiten Gleichung. Wir können dann die zweite Gleichung im Geiste unserer bereits entwickelten Fixpunkttheorie in partiellen Ordnungen als einen *kleinsten* Fixpunkt begreifen, was am ehesten unserem normalen Begriff von Rekursion entspricht – wenn wir uns Pos als rekursiv programmierte Funktion vorstellen, besagt die Gleichung, dass $Pos(\phi)$ in einem Zustand s gilt, also ein ϕ erfüllender Zustand von s aus erreichbar ist, wenn entweder ϕ in s bereits gilt oder s einen Nachfolger hat, von dem aus ein ϕ erfüllender Zustand erreichbar ist; die Antwort „ja“ akzeptieren wir nur bei Terminierung des Verfahrens, nicht abbrechende Suchen nach einem ϕ erfüllenden Nachfolger gelten als fehlgeschlagen, liefern also „nein“ als Antwort.

Dagegen macht eine Sicht auf Inv als kleinsten Fixpunkt keinen Sinn; der kleinste Fixpunkt der rekursiven Gleichung im obigen LTS ist z.B. $\emptyset$ (warum?). Korrekt ist hier stattdessen, $Inv(\phi)$ als *größten* Fixpunkt zu lesen: Wir überprüfen gewissermaßen alle von einem Zustand aus erreichbaren Zustände auf Erfülltheit von ϕ ; wenn wir immer weiter suchen, ohne eine Verletzung von ϕ zu finden, ist die Antwort „ja“. (Was ist der größte Fixpunkt der Definition von Pos ?)

4.1 Der (multi-)modale μ -Kalkül

Wie in HML verwenden wir Modalitäten (mehrere, daher „multi“) $[\alpha]$ und $\langle \alpha \rangle$ mit $\alpha \in \mathcal{A}$.

4.2 Syntax

Wir definieren Formeln im μ -Kalkül durch die Grammatik

$$\phi, \psi ::= \neg x \mid x \mid \phi \wedge \psi \mid \phi \vee \psi \mid [\alpha]\phi \mid \langle \alpha \rangle \psi \mid \mu x. \phi \mid \nu x. \phi$$

wobei x aus einer abzählbaren Menge $Vars$ von Variablen kommt, $\alpha \in \mathcal{A}$ und $\neg$ nur vor freien Variablen x im unten definierten Sinn vorkommt. Wir verwenden gelegentlich $\eta \in \{\mu, \nu\}$ als Metavariablen in Ausdrücken $\eta x. \phi$. Wir können $\top$ und $\perp$ definieren per

$$\begin{aligned} \top &= \nu x. x \\ \perp &= \mu x. x \end{aligned}$$

Beispiel 70. Der folgende Ausdruck ist eine Formel im μ -Kalkül:

$$\nu z. \mu y. \nu x. ([a](\langle b \rangle x \wedge z) \vee [\mathcal{A}]y)$$

Wir sehen μ und ν als *Variablenbinder* an und definieren demgemäß die Menge $FV(\phi)$ der *freien Variablen* in einer Formel ϕ durch

- $FV(x) = x$;
- $FV(\neg x) = x$;
- $FV(\phi \wedge \psi) = FV(\phi \vee \psi) = FV(\phi) \cup FV(\psi)$;
- $FV([a]\phi) = FV(\langle a \rangle \phi) = FV(\phi)$
- $FV(\mu x. \phi) = FV(\nu x. \phi) = FV(\phi) \setminus \{x\}$

Eine Formel ϕ ist *geschlossen*, wenn $FV(\phi) = \emptyset$. Nach obiger Einschränkung der Syntax kommt insbesondere in geschlossenen Formeln keine Negation vor.

4.3 Semantik

Allgemein definieren wir die Semantik von Formeln ϕ über einem LTS $T = (\mathcal{Q}, \mathcal{A}, \rightarrow)$ relativ zu einer *Belegung*

$$\sigma : \text{Vars} \rightarrow \mathcal{P}(\mathcal{Q}),$$

die also eine Interpretation der Variablen als Mengen von Zuständen vorgibt. Die *Extension* $\llbracket \phi \rrbracket \sigma \subseteq \mathcal{Q}$ einer Formel ϕ ist dann rekursiv definiert durch

Definition 71. Sei σ eine Belegung. Wir definieren die Semantik $\llbracket \phi \rrbracket \sigma$ von Formeln induktiv wie folgt:

- $\llbracket x \rrbracket \sigma = \sigma(x)$;
- $\llbracket \neg x \rrbracket \sigma = \mathcal{Q} \setminus \sigma(x)$;
- $\llbracket \phi \vee \psi \rrbracket \sigma = \llbracket \phi \rrbracket \sigma \cup \llbracket \psi \rrbracket \sigma$;
- $\llbracket \phi \wedge \psi \rrbracket \sigma = \llbracket \phi \rrbracket \sigma \cap \llbracket \psi \rrbracket \sigma$;
- $\llbracket \langle a \rangle \phi \rrbracket \sigma = \{s \in \mathcal{Q} \mid \exists s'. s \xrightarrow{a} s' \wedge s' \in \llbracket \phi \rrbracket \sigma\}$;
- $\llbracket [a]\phi \rrbracket \sigma = \{s \in \mathcal{Q} \mid \forall s'. s \xrightarrow{a} s' \implies s' \in \llbracket \phi \rrbracket \sigma\}$;
- $\llbracket \mu x. \phi \rrbracket \sigma = \cap \{S \subseteq \mathcal{Q} \mid \llbracket \phi \rrbracket \sigma[x \mapsto S] \subseteq S\}$;

- $\llbracket \nu x. \phi \rrbracket \sigma = \bigcup \{S \subseteq \mathcal{Q} \mid S \subseteq \llbracket \phi \rrbracket \sigma[x \mapsto S]\};$

wobei $\sigma[x \mapsto S](y) = \sigma(y)$ für $y \neq x$, und $\sigma[x \mapsto S](x) = S$. Wir lassen σ in $\llbracket \phi \rrbracket \sigma$ weg, wenn ϕ geschlossen ist. Wir schreiben in diesem Fall $s \models \phi$ für $s \in \llbracket \phi \rrbracket$.

Beispiel 72.

$$\begin{aligned}\llbracket \top \rrbracket &= \llbracket \nu x. x \rrbracket = \bigcup \{S \subseteq \mathcal{Q} \mid S \subseteq \llbracket x \rrbracket [x \mapsto S]\} = \mathcal{Q}, \\ \llbracket \perp \rrbracket &= \llbracket \mu x. x \rrbracket = \bigcap \{S \subseteq \mathcal{Q} \mid \llbracket x \rrbracket [x \mapsto S] \subseteq S\} = \emptyset.\end{aligned}$$

Wie bereits für HML können wir Negation definieren mittels

$$\begin{aligned}\neg(\neg x) &= x, \\ \neg(\mu x. \phi) &= \nu x. \neg\phi[\neg x/x], \\ \neg(\nu x. \phi) &= \mu x. \neg\phi[\neg x/x]\end{aligned}$$

und weiteren Klauseln wie in HML. (Man mache sich klar, dass die syntaktischen Beschränkungen des μ -Kalküls hier in der Tat eingehalten werden.)

Beispiel 73. Negation einer μ -Kalkül-Formel:

$$\begin{aligned}&\neg(\mu z. ([tick] \perp \vee [\mathcal{A}]z)) \\ &= \nu z. (\langle tick \rangle \top \wedge \langle \mathcal{A} \rangle \neg z) \\ &= \nu z. (\langle tick \rangle \top \wedge \langle \mathcal{A} \rangle z).\end{aligned}$$

Negation wie oben definiert agiert wirklich wie Negation:

Lemma 74. In einem LTS $(\mathcal{Q}, \rightarrow)$ gilt für jede Formel ϕ und jede Belegung σ

$$\llbracket \neg\phi \rrbracket \sigma = \mathcal{Q} \setminus \llbracket \phi \rrbracket \sigma$$

Beweis. Induktion über ϕ ; alle Fälle sind klar außer den Fixpunktfällen. Per Dualität reicht einer, z.B. für $\mu x. \phi$. Wir rechnen direkt mit der obigen Semantik:

$$\begin{aligned}&\mathcal{Q} \setminus \llbracket \mu x. \phi \rrbracket \sigma \\ &= \mathcal{Q} \setminus \bigcap \{S \mid \llbracket \phi \rrbracket \sigma[x \mapsto S] \subseteq S\} \\ &= \bigcup \{\mathcal{Q} \setminus S \mid \llbracket \phi \rrbracket \sigma[x \mapsto S] \subseteq S\} \\ &= \bigcup \{\mathcal{Q} \setminus S \mid \mathcal{Q} \setminus S \subseteq \mathcal{Q} \setminus \llbracket \phi \rrbracket \sigma[x \mapsto S]\} \\ &= \bigcup \{\mathcal{Q} \setminus S \mid \mathcal{Q} \setminus S \subseteq \llbracket \neg\phi \rrbracket \sigma[x \mapsto S]\} \tag{IV} \\ &= \bigcup \{\mathcal{Q} \setminus S \mid \mathcal{Q} \setminus S \subseteq \llbracket \neg\phi[\neg x/x] \rrbracket \sigma[x \mapsto \mathcal{Q} \setminus S]\} \\ &= \bigcup \{S' \mid S' \subseteq \llbracket \neg\phi[\neg x/x] \rrbracket \sigma[x \mapsto S']\} \\ &= \llbracket \nu x. \neg\phi[\neg x/x] \rrbracket \sigma.\end{aligned}$$

Dabei nutzen wir im vorletzten Schritt aus, dass alle Teilmengen von $\mathcal{Q}$ von der Form $\mathcal{Q} \setminus S$ sind. $\square$

Lemma 75. *Sei $\eta x. \phi$ eine Formel und σ eine Belegung. Dann ist die Funktion*

$$\begin{aligned} f : \mathcal{P}(\mathcal{Q}) &\rightarrow \mathcal{P}(\mathcal{Q}) \\ S &\mapsto \llbracket \phi \rrbracket \sigma[x \mapsto S] \end{aligned}$$

monoton, und $\llbracket \mu x. \phi \rrbracket \sigma = \mu f$, $\llbracket \nu x. \phi \rrbracket \sigma = \nu f$

Beweis. Die zweite Aussage folgt nach dem Fixpunktsatz von Knaster und Tarski (Satz 52) sofort aus der ersten. Wir beweisen die erste Aussage per Induktion über ψ . Man beachte, dass der Fall $\phi = \neg x$ nicht vorkommt. Wir beweisen als weitere Beispiele nur die Fälle für $\psi = \langle a \rangle \phi$ und $\psi = \mu y. \phi$ mit $x \neq y$ (der Fall $x = y$ ist trivial, da x dann nicht frei in ψ ist).

$\psi = \langle a \rangle \phi$: Sei $S_1 \subseteq S_2 \subseteq \mathcal{Q}$. Zu zeigen ist

$$\llbracket \langle a \rangle \phi \rrbracket \sigma[x \mapsto S_1] \subseteq \llbracket \langle a \rangle \phi \rrbracket \sigma[x \mapsto S_2].$$

Sei also $s \in \llbracket \langle a \rangle \phi \rrbracket \sigma[x \mapsto S_1]$. Dann existiert s' , so dass $s \xrightarrow{\alpha} s'$ und $s' \in \llbracket \phi \rrbracket \sigma[x \mapsto S_1]$, woraus nach Induktionsvoraussetzung $s' \in \llbracket \phi \rrbracket \sigma[x \mapsto S_2]$ folgt. Per Definition der Semantik gilt also wie verlangt $s \in \llbracket \langle a \rangle \phi \rrbracket \sigma[x \mapsto S_2]$.

$\psi = \mu y. \phi$: Sei $S_1 \subseteq S_2 \subseteq \mathcal{Q}$. Zu zeigen ist

$$\llbracket \mu y. \phi \rrbracket \sigma[x \mapsto S_1] \subseteq \llbracket \mu y. \phi \rrbracket \sigma[x \mapsto S_2].$$

Schreibe U für die rechte Seite dieser Ungleichung. Per Definition der Semantik reicht es zu zeigen, dass

$$\llbracket \phi \rrbracket \sigma[x \mapsto S_1][y \mapsto U] \subseteq U.$$

Man beachte, dass wegen $x \neq y$

$$\sigma[x \mapsto S_i][y \mapsto U] = \sigma[y \mapsto U][x \mapsto S_i] \tag{1}$$

für $i = 1, 2$. Wir rechnen wie folgt:

$$\begin{aligned} &\llbracket \phi \rrbracket \sigma[x \mapsto S_1][y \mapsto U] \\ &\subseteq \llbracket \phi \rrbracket \sigma[x \mapsto S_2][y \mapsto U] && \text{(IV, (1))} \\ &= U, \end{aligned}$$

wobei wir im letzten Schritt die IV in Gestalt der zweiten Aussage des Lemmas verwenden. $\square$

Für endliche LTS gilt damit nach dem Fixpunktsatz von Kleene für endliche Verbände (Satz 53):

$$\begin{aligned}\llbracket \mu x. \phi \rrbracket \sigma &= f^n(\emptyset) \quad \text{für } n \text{ mit } f^n(\emptyset) = f^{n+1}(\emptyset) \\ \llbracket \nu x. \phi \rrbracket \sigma &= f^m(\mathcal{Q}) \quad \text{für } m \text{ mit } f^m(\mathcal{Q}) = f^{n+1}(\mathcal{Q}).\end{aligned}$$

Beispiel 76. Man betrachte das LTS

$$\begin{array}{c} \text{tick} \\ \curvearrowright \\ \text{Clock} \xrightarrow{\text{tick}} \text{tick}.\emptyset \xrightarrow{\text{tick}} \emptyset \end{array}$$

Die Formel $\nu z. \langle \text{tick} \rangle z$ besagt, dass unendlich oft die Aktion *tick* durchgeführt werden kann. Wir berechnen $\llbracket \nu z. \langle \text{tick} \rangle z \rrbracket$ nach der oben angedeuteten iterativen Methode: Sei $g(S) = \llbracket \langle \text{tick} \rangle z \rrbracket [z \mapsto S]$. Dann gilt per Definition

$$\begin{aligned}g^0(\mathcal{Q}) &= \mathcal{Q}, \\ g^1(\mathcal{Q}) &= \llbracket \langle \text{tick} \rangle z \rrbracket [z \mapsto \mathcal{Q}] = \{\text{Clock}, \text{tick}.\emptyset\}, \\ g^2(\mathcal{Q}) &= \llbracket \langle \text{tick} \rangle z \rrbracket [z \mapsto \{\text{Clock}, \text{tick}.\emptyset\}] = \{\text{Clock}\}.\end{aligned}$$

Dagegen besagt die Formel $\mu z. [\text{tick}] \perp \vee \langle \mathcal{A} \rangle z$, dass ein Zustand erreicht werden kann, der die Aktion *tick* verweigert. Wir berechnen $\llbracket \mu z. [\text{tick}] \perp \vee \langle \mathcal{A} \rangle z \rrbracket$ iterativ: Sei $f(S) = \llbracket [\text{tick}] \perp \vee \langle \mathcal{A} \rangle z \rrbracket [S/z]$. Dann

$$\begin{aligned}f^0(\emptyset) &= \emptyset \\ f^1(\emptyset) &= \llbracket [\text{tick}] \perp \vee \langle \mathcal{A} \rangle z \rrbracket [z \mapsto \emptyset] = \{\emptyset\} \\ f^2(\emptyset) &= \llbracket [\text{tick}] \perp \vee \langle \mathcal{A} \rangle z \rrbracket [z \mapsto \{\emptyset\}] = \{\emptyset, \text{tick}.\emptyset\} \\ f^3(\emptyset) &= \llbracket [\text{tick}] \perp \vee \langle \mathcal{A} \rangle z \rrbracket [z \mapsto \{\emptyset, \text{tick}.\emptyset\}] = \{\emptyset, \text{tick}.\emptyset, \text{Clock}\}\end{aligned}$$

Was besagt die Formel $\mu z. [\text{tick}] \perp \vee [\mathcal{A}]z$? Was die Formel $\nu z. [\text{tick}] \perp \vee [\mathcal{A}]z$?

4.4 Temporale Spezifikation im μ -Kalkül

Wir haben bereits folgende Formelschemata kennengelernt:

$$\begin{aligned}Pos(\phi) &= \mu x. (\phi \vee \langle \mathcal{A} \rangle x) & - \text{ weak liveness} \\ Inv(\phi) &= \nu x. (\phi \wedge [\mathcal{A}]x) & - \text{ strong safety}\end{aligned}$$

Die Operatoren *Pos* und *Inv* sind dual, d.h. $\neg Pos(\phi) = Inv(\neg \phi)$, denn

$$\begin{aligned}
& \neg Pos(\phi) \\
&= \neg \mu x. (\phi \vee \langle \mathcal{A} \rangle x) \\
&= \nu x. (\neg \phi \wedge [\mathcal{A}]x) \\
&= Inv(\neg \phi)
\end{aligned}$$

Ein weiteres oft benötigtes Formelschema ist *strong liveness*:

$$\mu x. (\phi \vee (\langle \mathcal{A} \rangle \top \wedge [\mathcal{A}]x)),$$

und dual dazu weak safety:

$$\nu x. (\phi \wedge ([\mathcal{A}] \perp \vee \langle \mathcal{A} \rangle x))$$

Was bedeuten diese Formeln?

Computation Tree Logic (CTL) ist ein Fragment des μ -Kalküls, allerdings üblicherweise einer Form des Kalküls, in der zum einen nur eine Aktion angenommen wird (die zudem meist als nie blockiert vorausgesetzt wird) und zum anderen die Zustände mit Wahrheitsbelegungen für propositionale Atome ausgestattet sind; mit anderen Worten sitzt in CTL die atomare Information üblicherweise an den Zuständen, nicht an den Transitionen. Wir können die CTL-Operatoren dennoch sinnvoll für unseren Kontext anpassen: Wir schreiben kurz $\Diamond = \langle \mathcal{A} \rangle$, $\Box = [\mathcal{A}]$, und setzen

- $EX\phi = \Diamond\phi$
- $A[\phi U \psi] = \mu x. \psi \vee (\phi \wedge \Box x)$
- $E[\phi U \psi] = \mu x. \psi \vee (\phi \wedge \Diamond x)$

Sinn der Notation ist hierbei, dass A und E sogenannte *Pfadquantoren* sind und besagen, dass eine Formel auf allen maximal langen (ggf. unendlichen) Pfaden vom aktuellen Zustand aus gilt bzw. auf mindestens einem Pfad. Die *temporalen Operatoren* X und U dagegen sind als auf einem bereits gegebenen Pfad interpretiert zu verstehen. Dabei besagt $X\phi$, dass ϕ im nächsten Zustand gilt, und $[\phi U \psi]$, dass auf dem Pfad irgendwann ein Zustand erreicht wird, in dem ψ gilt, und bis dahin ϕ gilt. Man mache sich klar, dass obige Definitionen tatsächlich diese Lesart implementieren.

In CTL werden A/E und X/U aber nur in Kombination wie oben verwendet, die dann jeweils als ein Operator angesehen werden (bezeichnet als EX , AU bzw. EU). (In der erweiterten Logik CTL* können Pfadquantoren und temporale Operatoren unabhängig verwendet werden; diese Logik ist allerdings wesentlich schwerer in den μ -Kalkül einzubetten als CTL.)

In CTL werden folgende weitere Operatoren verwendet:

	X (next)	F (finally)	G (globally)	U (until)	W (weak until)
A (all)	AX	AF	AG	AU	AW
E (exists)	EX	EF	EG	EU	EW

Hierbei ist $F\phi$ zu lesen als „ ϕ gilt irgendwann auf dem Pfad“, $G\phi$ als „ ϕ gilt auf dem gesamten Pfad“, und $\phi W \psi$ als „ ϕ gilt auf dem Pfad (mindestens), bis ψ gilt“ (was aber nicht eintreten muss). Die zusätzlichen Operatoren sind aber durch EX , AU und EU ausdrückbar:

$$\begin{aligned}
AX\phi &= \neg EX(\neg\phi) = \neg\Diamond\neg\phi = \Box\phi, \\
AF\phi &= A[\top U\phi] = \mu x. \phi \vee (\top \wedge \Box x) = \mu x. \phi \vee \Box x, \\
EF\phi &= E[\top U\phi] = \mu x. \phi \vee (\top \wedge \Diamond x) = \mu x. \phi \vee \Diamond x, \\
AG\phi &= \neg EF(\neg\phi) = \neg\mu x. \neg\phi \vee \Diamond x = \nu x. \phi \wedge \Box x, \\
EG\phi &= \neg AF(\neg\phi) = \neg\mu x. \neg\phi \vee \Box x = \nu x. \phi \wedge \Diamond x, \\
A[\phi W \psi] &= \neg E[\neg\psi U \neg(\phi \vee \psi)] = \nu x. (\psi \vee (\phi \wedge \Box x)), \\
E[\phi W \psi] &= \neg A[\neg\psi U \neg(\phi \vee \psi)] = \nu x. (\psi \vee (\phi \wedge \Diamond x)).
\end{aligned}$$

5 Erfülltheitsspiele

Man kann die Erfülltheit von Formeln in Zuständen von LTS spieltheoretisch beschreiben. Die Behandlung des vollen μ -Kalküls in dieser Form führt für unsere gegenwärtigen Zwecke zu weit; man kann aber relativ zwanglos die Erfülltheit von Formeln charakterisieren, die Fixpunkte von nur einer Form (μ oder ν) beinhalten.

Wir entwickeln das Spiel zunächst für den einfachen Fall von geschlossenen Formeln der Form $\mu x. \phi_0$ mit ϕ_0 fixpunktfrei. Um die Ausdrucksstärke etwas zu erhöhen, lassen wir $\top$ und $\perp$ als First-Class Citizens in ϕ_0 zu, mit der erwarteten Semantik. Erfülltheit einer solchen Formel in einem Zustand s_0 eines LTS $T = (\mathcal{Q}, \mathcal{A}, \rightarrow)$ charakterisieren wir durch folgendes Spiel:

- *Spieler*: $\forall$ (*Abaelard*) und $\exists$ (*Eloise*); $\exists$ spielt für Erfülltheit, $\forall$ dagegen.

- *Positionen*: Paare

$$(s, \psi)$$

aus einem Zustand $s \in \mathcal{Q}$ und einer Unterformel ψ von ϕ_0 (inklusive ϕ selbst und x).

- *Initialposition*: (s_0, x) .

- *Züge*: Abhängig von der Formel in der aktuellen Position:
 - $(s, \phi \wedge \psi)$: $\forall$ ist am Zug und zieht wahlweise zu (s, ϕ) oder (s, ψ) .
 - $(s, \phi \vee \psi)$: $\exists$ ist am Zug und zieht wahlweise zu (s, ϕ) oder (s, ψ) .
 - $(s, \perp)$: $\exists$ ist am Zug, hat aber keinen Zug.
 - $(s, \top)$: $\forall$ ist am Zug, hat aber keinen Zug.
 - $(s, \langle a \rangle \phi)$: $\exists$ ist am Zug und zieht zu einem (s', ϕ) mit $s \xrightarrow{a} s'$.
 - $(s, [a]\phi)$: $\forall$ ist am Zug und zieht zu einem (s', ϕ) mit $s \xrightarrow{a} s'$.
 - (s, x) : Es wird ein Schiedsrichterzug (der zur Vereinheitlichung beliebig $\forall$ oder $\exists$ zugerechnet werden kann) zu (s, ϕ_0) ausgeführt.
- *Spielende*: Nur dann, wenn der am Zug befindliche Spieler nicht ziehen kann (d.h. bei $(s, \top)$ oder $(s, \perp)$, sowie bei $(s, \langle a \rangle \phi)$ oder $(s, [a]\phi)$, wenn s keine a -Transition hat)
- *Gewinnbedingung*: Wer nicht ziehen kann, verliert. Unendliche Spiele gewinnt $\forall$.

Das Spiel für $\nu x. \phi_0$ verläuft genauso, mit dem einzigen Unterschied, dass $\exists$ unendliche Spiele gewinnt.

Satz 77. Für einen Zustand s_0 in einem LTS $T = (\mathcal{Q}, \mathcal{A}, \rightarrow)$ und eine μ -Kalkül-Formel $\mu x. \phi$ mit ϕ fixpunktfrei und $FV(\phi) \subseteq \{x\}$ gilt $s \models \mu x. \phi$ genau dann, wenn $\exists$ das zugehörige Erfüllungsspiel gewinnt; Entsprechendes gilt für $\nu x. \phi$.

Beweis. Wir beweisen zunächst die Aussage für $\nu x. \phi_0$, die wir in zwei Implikationen zerlegen:

„ $\implies$ “: Sei $s_0 \models \nu x. \phi_0$; wir müssen eine Gewinnstrategie für $\exists$ im Erfüllungsspiel für $\nu x. \phi_0$ angeben, die wir wieder durch eine Invariante auf Positionen (s, ϕ) beschreiben, nämlich gerade

$$s \models \phi[\nu x. \phi_0/x].$$

Wir prüfen die relevanten Bedingungen:

- Die Invariante gilt nach Annahme in der Anfangsposition (s_0, x) .
- $\exists$ kann die Invariante durchsetzen: Wir nehmen an, dass die Invariante in (s, ϕ) gilt, und unterscheiden Fälle nach ϕ .

- $\phi = \chi \vee \psi$: $\exists$ ist am Zug. Per Annahme gilt $s \models \chi[\nu x. \phi_0/x]$ oder $s \models \psi[\nu x. \phi_0/x]$; o.E. Letzteres. Dann zieht s nach (s, ψ) , so dass weiter die Invariante erfüllt ist.
 - $\phi = \chi \wedge \psi$: $\forall$ ist am Zug; beide möglichen Nachfolgepositionen erfüllen die Invariante.
 - $\phi = \langle a \rangle \chi$: $\exists$ ist am Zug. Per Annahme existiert $s \xrightarrow{a} s'$ mit $s' \models \chi[\nu x. \phi_0/x]$, und $\exists$ zieht nach (s', χ) .
 - $\phi = [a] \chi$: $\forall$ ist am Zug und zieht zu (s', χ) mit $s \xrightarrow{a} s'$; nach Annahme gilt dann weiter die Invariante.
 - $\phi = x$: Der Schiedsrichter zieht zu (s, ϕ_0) ; nach der Fixpunkteigenschaft ist $\phi_0[\nu x. \phi_0]$ äquivalent zu $\nu x. \phi_0$, so dass weiterhin die Invariante gilt.
- Die Invariante gewinnt: Wenn (s, ϕ) die Invariante erfüllt und $\exists$ am Zug ist, dann kann $\exists$ ziehen – dies steht überhaupt nur dann in Frage, wenn ϕ die Form $\phi = \langle a \rangle \chi$ hat, und ist in diesem Fall aber gerade durch $s \models \langle a \rangle \dots$ garantiert. Da $\exists$ unendliche Spiele gewinnt, reicht dies aus.

„ $\Leftarrow$ “: Wir schreiben kurz $\llbracket \phi \rrbracket(S) = \llbracket \phi \rrbracket \sigma[x \mapsto S]$; dann ist $\nu x. \phi_0$ der größte Postfixpunkt von $\llbracket \phi \rrbracket$. Es reicht also zu zeigen, dass

$$W_{\exists}^{\nu} = \{s \in \mathcal{Q} \mid \exists \text{ gewinnt } (s, x)\}$$

ein Postfixpunkt von $\llbracket \phi_0 \rrbracket$ ist, d.h.

$$W_{\exists}^{\nu} \subseteq \llbracket \phi \rrbracket(W_{\exists}^{\nu}). \quad (2)$$

Wir zeigen durch Induktion über Unterformeln ϕ von ϕ_0 allgemeiner, dass für $s \in \mathcal{Q}$

$$\exists \text{ gewinnt } (s, \phi) \implies s \in \llbracket \phi \rrbracket(W_{\exists}^{\nu}). \quad (3)$$

- $\phi = x$: Per Definition gilt $s \in W_{\exists}^{\nu} = \llbracket x \rrbracket(W_{\exists}^{\nu})$.
- $\phi = \chi \wedge \psi$: Nach den Regeln des Spiels gewinnt $\exists$ sowohl (s, χ) als auch (s, ψ) , per Induktionsvoraussetzung also $s \in \llbracket \chi \rrbracket(W_{\exists}^{\nu})$ und $s \in \llbracket \psi \rrbracket(W_{\exists}^{\nu})$, also $s \in \llbracket \chi \wedge \psi \rrbracket(W_{\exists}^{\nu})$.
- $\phi = \chi \vee \psi$: O.E. gewinnt $\exists$ (s, χ) ; per Induktionsvoraussetzung weiter wie im vorigen Fall.
- $\phi = \langle a \rangle \chi$: Es existiert $s \xrightarrow{a} s'$, so dass $\exists$ (s', χ) gewinnt. Nach Induktionsvoraussetzung folgt $s' \in \llbracket \chi \rrbracket(W_{\exists}^{\nu})$ und damit per Semantik $s \in \llbracket \langle a \rangle \chi \rrbracket(W_{\exists}^{\nu})$.

- $\phi = [a]\chi$: Sei $s \xrightarrow{a} s'$. Nach den Regeln des Spiels gewinnt $\exists (s', \chi)$; per Induktionsvoraussetzung weiter wie oben.

Damit folgt die Postfixpunkteigenschaft (2): Wenn $\exists (s, x)$ gewinnt, gewinnt nach den Spielregeln $\exists$ auch (s, ϕ_0) ; per (5) folgt $s \in \llbracket \phi_0 \rrbracket (W_{\exists}^{\nu})$ wie verlangt.

Die Aussage für $\mu x. \phi_0$ zeigen wir dann wie folgt.

Für „ $\implies$ “ zeigen wir ähnlich wie in der „ $\Leftarrow$ “-Richtung des ν -Falls, dass

$$W_{\exists}^{\mu} = \{s \mid \exists \text{ gewinnt } (s, x)\} \quad (4)$$

(hier bezieht sich „gewinnt“ natürlich auf das Spiel für $\mu x. \phi_0$!) ein Präfixpunkt von $\llbracket \phi_0 \rrbracket$ ist, bzw. allgemeiner

$$s \in \llbracket \phi \rrbracket (W_{\exists}^{\mu}) \implies \exists \text{ gewinnt } (s, \phi) \quad (5)$$

per Induktion über Unterformeln ϕ von ϕ_0 . Wir behandeln nur ausgewählte Fälle:

- $\phi = x$: Unmittelbar aus der Definition von $W_{\exists}^{\mu}$.
- $\phi = \chi \wedge \psi$: Aus $s \in \llbracket \phi \rrbracket (W_{\exists}^{\mu})$ erhalten wir per Semantik $s \in \llbracket \chi \rrbracket (W_{\exists}^{\mu})$ und $s \in \llbracket \psi \rrbracket (W_{\exists}^{\mu})$, so dass per IV $\exists$ die Positionen (s, χ) und (s, ψ) gewinnt. Damit gewinnt $\exists$ nach den Spielregeln auch $(s, \phi) = (s, \chi \wedge \psi)$.
- $\phi = \langle a \rangle \psi$: Per Semantik haben wir $s \xrightarrow{a} s'$ mit $s' \in \llbracket \psi \rrbracket (W_{\exists}^{\mu})$. Nach IV gewinnt $\exists (s', \psi)$, und damit nach den Spielregeln auch (s, ϕ) , nämlich per Zug nach (s', ψ) .

Wiederum folgt nun die Präfixpunkteigenschaft (4), da $\exists$ nach den Spielregeln mit (s, ϕ_0) auch (s, x) gewinnt.

Für „ $\Leftarrow$ “ gehen wir per Kontraposition vor, und verwenden den schon bewiesenen Teil für ν . Sei also $s \not\models \mu x. \phi_0$. Dann $s \models \nu x. \bar{\phi}_0$, wobei $\bar{\phi}_0 = \neg \phi_0[\neg x/x]$ die Dualisierung von ϕ_0 ist. Nach der Aussage für ν folgt, dass $\exists$ das Spiel für $\nu x. \bar{\phi}_0$ gewinnt. Dieses Spiel erhält man aber gerade, indem man das Spiel für $\mu x. \phi_0$ nimmt und die Rollen von $\exists$ und $\forall$ vertauscht; d.h. $\forall$ gewinnt das Spiel für $\mu x. \phi_0$, woraus unmittelbar folgt, dass $\exists$ dieses Spiel nicht gewinnt, wie verlangt. $\square$

6 Der π -Kalkül

Wir erinnern uns daran, dass CCS-Prozesse sich über gemeinsame Kanäle synchronisieren; wir haben bereits eine rein syntaktische Erweiterung des Kalküls kennengelernt, in der über die Kanäle Werte kommuniziert werden.

Der π -Kalkül ist mit dem Ziel entwickelt worden, *Mobilität* von Prozessen in diesem Paradigma zu modellieren. Dazu werden zwei neue Ideen in den Kalkül eingeführt:

- Über die Kanäle können auch (bzw. in der Grundversion ausschließlich) Kanalnamen kommuniziert werden, die dann im empfangenden Prozess für weitere Kommunikation genutzt werden. Z.B. in

$$\bar{a}b.S \mid a(x).\bar{x}c.R \xrightarrow{\tau} S \mid \bar{b}c.R$$

wird der Kanalname b über Kanal a vom linken zum mittleren Prozess gesendet, der anschließend über b mit weiteren Prozessen, z.B. dem rechten, kommunizieren kann.

- Dies betrifft auch restringierte Namen, also private Kanäle, die damit ihren Scope vergrößern können. Für die Restriktion $P \setminus \{a\}$ schreibt man im π -Kalkül $(\nu a)(P)$ (nicht zu verwechseln mit dem gleichnamigen Fixpunktoperator; das ν spricht sich auf Englisch wie „new“). Z.B. klappt obiger Transfer von b auch dann, wenn b restringiert ist:

$$(\nu b)(\bar{a}b.S \mid P) \mid a(x).\bar{x}c.R \xrightarrow{\tau} (\nu b)(S \mid P \mid \bar{b}c.R)$$

Hier war der Name b zunächst privat für die linken beiden Prozesse, wird aber dann über den öffentlichen Kanal a an den rechten Prozess kommuniziert, und ist danach privat für alle drei Prozesse.

Die im zweiten Punkt angesprochene *scope extrusion* ist die Hauptursache für die im Vergleich zu CCS relativ hohe technische Komplexität des π -Kalküls.

6.1 Syntax

Die Syntax des π -Kalküls unterscheidet sich wie schon angedeutet nur marginal von der von value-passing CCS. Explizit: Wir unterstellen abzählbar unendliche Vorräte von (Kanal-)namen $a, b, \dots$, die wir als Eingabekanäle $a, b, \dots$ oder Ausgabekanäle $\bar{a}, \bar{b}, \dots$ verwenden, sowie Prozessnamen $A, B, \dots$. Wir unterscheiden nicht zwischen Namen und Variablen. Prozessnamen haben eine gegebene Arität und sind dann über entsprechend viele Namen parametrisiert, d.h. wenn A Arität k hat, verwenden wir A in der Form $A(x_1, \dots, x_k)$. Kanalnamen sind der Einfachheit halber stets einstellig; wir haben Präfixe

$$\begin{array}{ll} \bar{a}b & \text{Ausgabe von } b \text{ über Kanal } a \\ a(x) & \text{Empfang von } x \text{ über Kanal } a \\ \tau & \text{Stumme Aktion} \end{array}$$

– Achtung: Wie in den obigen Beispielen gesehen, spielt der Name x in $a(x)$ die Rolle einer Eingabevariablen, während in $\bar{a}b$ der Name b selbst über a gesendet wird. Prozesse $P, Q, \dots$ definieren wir über die Grammatik

$$P, Q ::= A(x_1, \dots, x_k) \mid \emptyset \mid \alpha.P \mid P + Q \mid P|Q \mid (\nu x)P.$$

Wie schon in value-passing CCS *bindet* $a(x)$ seine Variable; Gleiches gilt für die Restriktion ν . Wir schreiben $FN(\cdot)$ für die Menge der freien Namen eines syntaktischen Objekts und $BN(\cdot)$ für die Menge der gebundenen Namen.

Wie bisher besteht eine *Prozessspezifikation* aus einer Auswahl von Prozessnamen und einer definierenden Gleichung $A(x_1, \dots, x_k) = P$ für jeden ausgewählten Prozessnamen.

Strukturelle Kongruenz Die Definition der Semantik des π -Kalküls wird zur Einsparung von Regeln typischerweise über einer vorab deklarierten Gleichheit von Prozesstermen definiert. Diese ist zu unterscheiden von später eingeführten Bisimilaritätsbegriffen; sie identifiziert nur Prozesse, die „offensichtlich“ das Gleiche tun. Im einzelnen beinhaltet strukturelle Kongruenz $\equiv$ folgende Gleichungen:

- α -Äquivalenz, d.h. Umbenennung gebundener Namen;
- Assoziativität und Kommutativität sowie Neutralität von $\emptyset$ für $+$ und $|$;
- Auffaltung von Prozessdefinitionen; und
- Scope-Erweiterung:

$$\begin{aligned} \emptyset &\equiv (\nu x)\emptyset \\ P \mid (\nu x)Q &\equiv (\nu x)(P \mid Q) & (x \notin FV(P)) \\ P + (\nu x)Q &\equiv (\nu x)(P + Q) & (x \notin FV(P)) \\ (\nu y)(\nu x)P &\equiv (\nu x)(\nu y)P. \end{aligned}$$

Wie der Name sagt, ist $\equiv$ per Definition eine Kongruenz, d.h. die obigen Gleichungen können in Folge und in umgekehrter Richtung sowie tief in Termen verwendet werden.

Insbesondere können wir eine nicht unterhalb eines Präfixes stehende Restriktion (νx) stets ganz nach vorn ziehen, ggf. unter geeigneter Umbenennung. Ferner können wir für $x \notin FV(P)$ die Gleichung $(\nu x)P \equiv P$ herleiten:

$$P \equiv P \mid \emptyset \equiv P \mid (\nu x)\emptyset \equiv (\nu x)(P \mid \emptyset) \equiv (\nu x)P.$$

6.2 Semantik

Wie schon für CCS geben wir die Semantik des π -Kalküls mittels eines LTS an. Dieses LTS benötigt mehr Aktionen als zunächst vielleicht erwartet: Die Prozesse

$$(\nu x)\bar{a}x.P \quad \text{und} \quad \bar{a}x.P$$

müssen klarerweise verschiedene Interpretationen bekommen: Das Präfix $\bar{a}x$ bewirkt im linken Prozess die Extrusion eines *gebundenen* Namens x , während rechts x als *freier* Name kommuniziert wird. Wir führen zur Definition des LTS daher folgende Aktionen ein:

- τ (stumme Aktion)
- $\bar{a}x$ (*freie* Ausgabe)
- $a(x)$ (Eingabe)
- $\bar{a}\nu x$ mit $a \neq x$ (*gebundene* Ausgabe).

Dabei ist x in $a(x)$ und in $\bar{a}\nu x$ gebunden, dagegen in $\bar{a}x$ frei. Wir schreiben $FN(\alpha)$ und $BN(\alpha)$ für die Mengen der freien bzw. gebundenen Namen in einer Aktion α . Die Regeln der Semantik lauten dann wie folgt:

$$\begin{aligned}
(STRUCT) \quad & \frac{P' \equiv P \quad P \xrightarrow{\alpha} Q \quad Q \equiv Q'}{P' \xrightarrow{\alpha} Q'} \\
(ACT) \quad & \frac{}{\alpha.P \xrightarrow{\alpha} P} \quad (SUM) \quad \frac{P \xrightarrow{\alpha} P'}{P + Q \xrightarrow{\alpha} P'} \\
(COM_1) \quad & \frac{P \xrightarrow{\alpha} P'}{P \mid Q \xrightarrow{\alpha} P' \mid Q} \quad (BN(\alpha) \cap FN(Q) = \emptyset) \\
(COM_2) \quad & \frac{P \xrightarrow{a(x)} P' \quad Q \xrightarrow{\bar{a}u} Q'}{P \mid Q \xrightarrow{\tau} P'[u/x] \mid Q'} \\
(RES) \quad & \frac{P \xrightarrow{\alpha} P'}{(\nu x)P \xrightarrow{\alpha} (\nu x)P'} \quad (x \notin FN(\alpha) \cup BN(\alpha)) \\
(OPEN) \quad & \frac{P \xrightarrow{\bar{a}x} P'}{(\nu x)P \xrightarrow{\bar{a}\nu x} P'} \quad (a \neq x)
\end{aligned}$$

Scope Extrusion funktioniert in diesem System wie folgt: Für $u \notin FN(P)$ haben wir

$$\begin{array}{c} (RES) \frac{a(x).P \mid \bar{a}u.Q \xrightarrow{\tau} P[u/x] \mid Q}{(\nu u)(a(x).P \mid \bar{a}u.Q) \xrightarrow{\tau} (\nu u)(P[u/x] \mid Q)} \\ (STRUCT) \frac{a(x).P \mid (\nu u)\bar{a}u.Q \xrightarrow{\tau} (\nu u)(P[u/x] \mid Q)}{a(x).P \mid (\nu u)\bar{a}u.Q \xrightarrow{\tau} (\nu u)(P[u/x] \mid Q)} \end{array}$$

(Man sieht hier, dass man gebundene Ausgabe-Transitionen nicht wirklich für die Herleitung anderer, insbesondere stummer, Transitionen braucht; sie sind aber essentiell für die relevanten Varianten von Bisimilarität, die wir hier allerdings leider nicht mehr behandeln.)

6.3 Beispiele

- Wir können auch in den Scope einer Restriktion Namen hineinkommunizieren:

$$(\nu b)(a(x).P) \mid \bar{a}b.Q \xrightarrow{\tau} ((\nu b')P[b'/b][b/x]) \mid Q$$

– wie wird dies im System hergeleitet?

- *Ressourcenmanagement*: Wir nehmen Prozesse R (*Ressource*, z.B. ein Drucker), S (*Server*) und C (*Client*) an. Der Server kontrolliert den Drucker mittels eines nur ihm zur Verfügung stehenden Triggers e :

$$(\nu e)(\bar{c}e.S \mid e.R).$$

Der Client fordert den Trigger vom Server über einen Kanal c an:

$$c(x).\bar{x}.C.$$

Bei $e \notin FN(C)$ haben wir dann insgesamt Transitionen

$$\begin{aligned} c(x).\bar{x}.C \mid (\nu e)(\bar{c}e.S \mid e.R) \\ \xrightarrow{\tau} (\nu e)(\bar{c}.C \mid S \mid e.R) \\ \xrightarrow{\tau} (\nu e)(C \mid S \mid R). \end{aligned}$$

- *Koordination von Messages*: Wir nehmen an, dass es mehrere Clients gibt, und der Server aber nacheinander zwei Namen d, e an *denselben* Client kommunizieren möchte. Dies stellt er durch Einrichten eines gemeinsamen Kanals mit einem Client sicher:

$$(\nu p)\bar{c}p.\bar{p}d.\bar{p}e.S.$$

Ein Client

$$c(z).z(x).z(y).C$$

empfängt nach der Synchronisierung dann beide Namen nacheinander auf p ; die Nachrichten können an keinen zweiten Client gehen, da p privat für S und C ist. Wie reduziert die Parallelkomposition von Client und Server? (Achtung: Was ist, wenn $p \in FN(C)$?)

- Wenn der Server darauf wartet, dass der Client mit einer Aufgabe (z.B. eben drucken) fertig ist, verabredet er ein Signal r mit dem Client, der dieses Signal dann sendet, wenn er fertig ist:

$$((\nu r)\bar{e}r.r.S) \mid e(x).\bar{p}.\bar{x}.C$$

(irrelevante Argumente von Präfixen lassen wir hier weg). Wie verhält sich dieser Prozess laut Semantik? (Man nehme zur Vereinfachung $r \notin FN(C)$ an.)