

Materiały do wykładu

Rachunek lambda

Paweł Urzyczyn
urzy@mimuw.edu.pl
12 sierpnia 2008

Rachunek lambda i logika kombinatoryczna powstały w latach trzydziestych dwudziestego wieku. Początkowo miały stanowić alternatywne wobec teorii mnogości podejście do podstaw matematyki, w którym funkcja (rozumiana intensjonalnie jako definicja) była pojęciem pierwotnym. Ten zamiar się nie powiódł, ale szybko okazało się, że rachunek lambda jest niezwykle użytecznym aparatem w teorii obliczeń. Definicję funkcji obliczalnej sformułowano wcześniej za pomocą rachunku lambda, niż w języku maszyn Turinga.

Później, w miarę rozwoju informatyki, język rachunku lambda okazał się niezastąpionym narzędziem w teorii języków programowania. A wreszcie i w praktyce, gdy pojawiły się języki programowania funkcyjnego. Dzisiaj znajomość podstaw rachunku lambda jest niezbędnym elementem wykształcenia nowoczesnego informatyka.

Poprzez związki z logiką intuicjonistyczną, rachunek lambda odzyskał też należne sobie miejsce w podstawach matematyki, zwłaszcza w teorii dowodu. Dlatego ta tematyka powinna być interesująca także dla studentów matematyki zainteresowanych logiką.

Te notatki obejmują najważniejsze wiadomości z rachunku lambda bez typów i podstawowe informacje o systemach z typami. Zrozumienie ich nie wymaga w zasadzie przygotowania wykraczającego poza materiał wykładany na I roku. Jednak wcześniejsze wysłuchanie wykładu z języków i automatów pozwoli na pewne zagadnienia z szerszej perspektywy.

1 Składnia rachunku lambda

W rachunku lambda rozważamy obiekty zwane *lambda-termami*. Tradycyjnie lambda-termi definiowane są jako formalne wyrażenia pewnego języka, podobnie jak zwykłe termy w algebrze i logice pierwszego rzędu. Istotna różnica polega na tym, że lambda-termi pozostające w pewnej relacji równoważności (alfa-konwersji) są ze sobą utożsamiane (uważane za identyczne). A więc lambda-termi to w istocie nie wyrażenia, ale klasy abstrakcji *lambda-wyrażeń* (nazywanych też *pre-termami*), które teraz zdefiniujemy.

Przyjmijmy, że mamy pewien przeliczalny nieskończony zbiór *zmiennych przedmiotowych*.

- Zmienne przedmiotowe są lambda-wyrażeniami;
- Jeśli M i N są lambda-wyrażeniami, to (MN) też;
- Jeśli M jest lambda-wyrażeniem i x jest zmienną, to $(\lambda x M)$ jest lambda-wyrażeniem.

Wyrażenie postaci (MN) nazywamy *aplikacją*, a wyrażenie postaci $(\lambda x M)$ to λ -*abstrakcja*. Intuicyjny sens aplikacji (MN) to zastosowanie operacji M do argumentu N . Abstrakcję $(\lambda x M)$ interpretujemy natomiast jako definicję operacji (funkcji), która argumentowi x przypisuje M . Oczywiście x może występować w M , tj. M zależy od x . Narzuca się analogia z procedurą (funkcją) o parametrze formalnym x i treści M .

Stosujemy następujące konwencje notacyjne:

- opuszczamy zewnętrzne nawiasy;
- aplikacja wiąże w lewo, tj. MNP oznacza $(MN)P$;
- piszemy $\lambda x_1 \dots x_n. M$ zamiast $\lambda x_1 (\dots (\lambda x_n M) \dots)$.

Kropkę w wyrażeniu $\lambda x_1 \dots x_n. M$ można uważać za lewy nawias, którego zasięg rozciąga się do końca wyrażenia M .

Operator lambda-abstrakcji λ wiąże zmienne, tj. wszystkie wystąpienia x w wyrażeniu $\lambda x M$ uważa się za *związane* (lokalne). Zmienne *wolne* (globalne) definiuje się tak:

- $FV(x) = \{x\}$;
- $FV(MN) = FV(M) \cup FV(N)$;
- $FV(\lambda x M) = FV(M) - \{x\}$.

Alfa-konwersja

Wybór zmiennych używanych w wyrażeniu jako związane jest sprawą drugorzędną. Takie wyrażenia, jak na przykład $\lambda x. xy$ i $\lambda z. zy$, intuicyjnie definiują tę samą operację („zaaplikuj dany argument do y ”) więc z naszego punktu widzenia powinny być uważane za takie same. Dlatego lambda-wyrażenia różniące się tylko zmiennymi związanymi chcemy ze sobą utożsamiać. Utożsamienie to nazywa się *alfa-konwersją*.

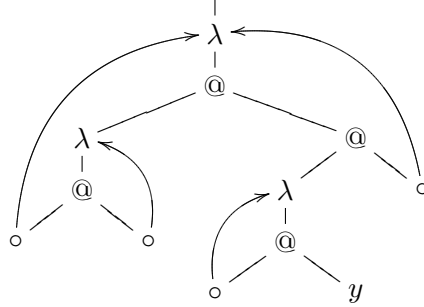
Pojęcie alfa-konwersji można ściśle zdefiniować na wiele sposobów, my wybierzemy interpretację grafową.

Lambda-grafy

Przez *drzewo* rozumiemy poniżej skończone drzewo etykietowane, w którym występować mogą następujące rodzaje wierzchołków:

- wierzchołki o etykiecie @, które mają zawsze dwa następniki: lewy i prawy;
- wierzchołki o etykiecie λ , które mają zawsze jeden następnik;
- liście etykietowane zmiennymi przedmiotowymi;
- liście bez etykiet.

Jeśli z każdego wierzchołka bez etykiety poprowadzimy nową krawędź do któregoś z poprzedników tego wierzchołka o etykiecie λ , to otrzymany graf skierowany nazwiemy *lambda-drzewem*. Liśćmi lambda-drzewa nazywamy jego wierzchołki końcowe (ich etykietami są zmienne). Na przykład taki graf jest lambda-drzewem (wierzchołki bez etykiet oznaczono przez $\circ$).



Teraz dowolnemu lambda-wyrażeniu M przypiszemy lambda-drzewo $G(M)$.

- Jeśli x jest zmienną to $G(x)$ składa się tylko z jednego wierzchołka o etykiecie x .
- Graf $G(PQ)$ ma korzeń o etykiecie $@$. Jego lewym następnikiem jest korzeń grafu $G(P)$ a prawym korzeń grafu $G(Q)$.
- Graf $G(\lambda x P)$ jest otrzymany z $G(P)$ poprzez
 - Dodanie nowego korzenia o etykiecie λ .
 - Połączenie go krawędzią z korzeniem grafu $G(P)$.
 - Dodanie krawędzi od wszystkich wierzchołków z etykietą x do nowego korzenia.
 - Usunięcie etykiet x .

Na przykład $G(\lambda z.(\lambda u.zu)((\lambda u.uy)z))$ to lambda-drzewo na rysunku powyżej. Łatwo widzieć, że etykiety liści w $G(M)$ to dokładnie zmienne wolne M . Na dodatek mamy:

Fakt 1.1 *Jeśli T jest lambda-drzewem, to $T = G(M)$ dla pewnego lambda-wyrażenia M .*

Dowód: Indukcja ze względu na liczbę wierzchołków grafu T . Jeśli jest tylko jeden wierzchołek, to jego etykietą jest pewna zmienna x i mamy $T = G(x)$. Jeśli etykietą korzenia jest $@$, to T składa się z korzenia i dwóch rozłącznych podgrafów T_1 i T_2 , do których stosujemy założenie indukcyjne. Mamy więc $T_1 = G(M_1)$ oraz $T_2 = G(M_2)$. Oczywiście wtedy $T = G(M_1M_2)$. Pozostaje przypadek gdy etykietą korzenia jest λ . Wtedy T składa się z korzenia, podgrafu T_1 zaczepionego w następniku korzenia i pewnej liczby krawędzi prowadzących do korzenia z wierzchołków bez etykiet. Przypisując tym wierzchołkom etykietę z , która nie występuje w grafie T_1 , otrzymujemy z T_1 lambda-drzewo T'_1 . Z założenia indukcyjnego $T'_1 = G(M_1)$ dla pewnego lambda-wyrażenia M_1 i możemy napisać $T = G(\lambda z.M_1)$. ■

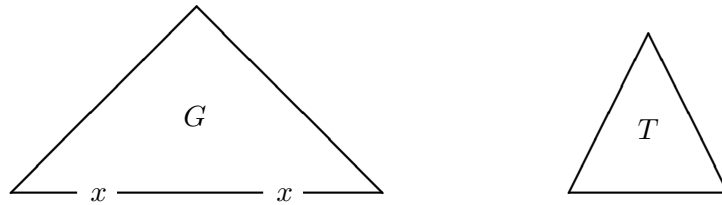
Relację alfa-konwersji (oznaczaną przez $=_\alpha$) definiujemy tak:

$$M =_\alpha N \quad \text{wtedy i tylko wtedy, gdy} \quad G(M) = G(N).$$

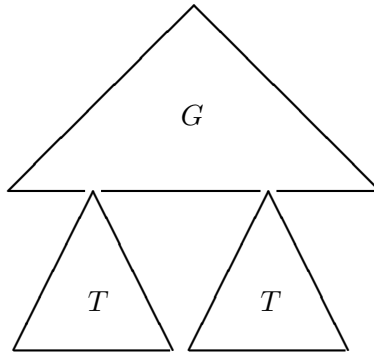
Od tej pory przez *lambda-term* albo po prostu *term* rozumiemy „lambda-wyrażenia z dokładnością do alfa-konwersji”, czyli w istocie klasy abstrakcji relacji $=_\alpha$. Każda taka klasa jest wyznaczona przez pewne lambda-drzewo, możemy więc uważać, że lambda-term to tak naprawdę lambda-drzewo. Lambda-wyrażenia są tylko ich syntaktyczną reprezentacją, przy czym dowolne dwie alfa-równoważne reprezentacje tego samego termu (lambda-drzewa) są tak samo dobre. Na przykład $(\lambda x.x(\lambda x.xy))(\lambda y.yx)$ to to samo co $(\lambda u.u(\lambda x.xy))(\lambda v.vx)$, ale nie to samo co $(\lambda y.y(\lambda y.yx))(\lambda x.xy)$.

Podstawienie

Dla danych lambda-grafów G i T przez *podstawienie* $G[x := T]$ rozumiemy graf otrzymany z G przez zamianę każdego liścia z etykietą x na korzeń odrębnej kopii grafu T . Ilustracja jest chyba prostsza i bardziej zrozumiała niż jakakolwiek formalna definicja. Jeśli mamy grafy



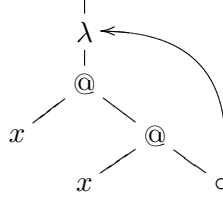
to graf $G[x := T]$ wygląda tak jak na Obrazku 1 (etykiety x już nie ma):



Obrazek 1: Podstawienie $G[x := T]$

Jeśli $G = G(M)$ i $T = G(N)$ to graf $G[x := T]$ reprezentuje wynik podstawienia wyrażenia N na miejsce wszystkich wolnych wystąpień zmiennej x w M . Ale ta interpretacja nie zawsze jest poprawna. Problem powstaje wtedy, gdy jakaś zmienna z wolna w N przy „naiwnym” tekstowym podstawieniu znalazłaby się w zasięgu wiązania λz . Na przykład wynikiem podstawienia z na x w termie $\lambda z.x(xz)$ nie powinno być $\lambda z.z(zz)$, ale $\lambda y.z(zy)$ (przy czym wybór

zmiennej y jest nieistotny). Mamy bowiem podstawić z na x nie w słowie ale w grafie



w którym w ogóle nie ma żadnego z . Tekstowe podstawienie jest poprawne tylko wtedy, gdy nie występuje konflikt nazw wolnych i związanych (globalnych i lokalnych). Dlatego wcześniej należy dokonać odpowiedniej wymiany zmiennych związanych.

Jeśli M i N są lambda-wyrażeniami to notacja $M[x := N]$ oznacza lambda-term o grafie $G(M)[x := G(N)]$. Napiszemy więc na przykład

$$(\lambda z.x(xz))[x := z] = \lambda y.z(zy),$$

pamiętając, że wybór zmiennej związanej jest nieistotny. Następujący lemat stanowi syntaktyczną definicję podstawienia.

Lemat 1.2 *Równość $M[x := N] = R$ ma miejsce wtedy i tylko wtedy, gdy zachodzi jeden z następujących przypadków:*

1. $M = x$, oraz $R = N$;
2. M jest zmienną różną od x , oraz $R = M$;
3. $M = PQ$, oraz $R = P[x := N]Q[x := N]$;
4. $M = \lambda y P$, gdzie $y \notin \{x\} \cup \text{FV}(N)$, oraz $R = \lambda y.P[x := N]$.

Dowód: Każdy term M jest albo zmienną albo aplikacją albo abstrakcją. Ta ostatnia możliwość to jedyny nieoczywisty przypadek, niech więc M będzie abstrakcją. Jeśli z lambda-drzewa termu M odrzucimy korzeń wraz z prowadzącymi do niego krawędziami, to otrzymamy graf M' , w którym niektóre wierzchołki końcowe nie mają etykiet. Nadając tym wierzchołkom „nową” etykietę y otrzymamy taki term P , że $M = \lambda y P$. Ponieważ y nie występuje jako etykieta w lambda-drzewie N , więc nie ma znaczenia czy najpierw w termie P zamienimy x na N , a potem dodamy lambdę wiążącą y , czy postąpimy w odwrotnej kolejności. ■

Zapiszmy treść lematu 1.2 w nieco prostszej postaci:

- $x[x := N] = N$;
- $y[x := N] = y$, gdy y jest zmienną różną od x ;
- $(PQ)[x := N] = P[x := N]Q[x := N]$;
- $(\lambda y P)[x := N] = \lambda y.P[x := N]$, gdy $y \neq x$ oraz $y \notin \text{FV}(N)$.

Powyższe reguły można stosować w dowodach indukcyjnych.

Lemat 1.3

1. $\text{FV}(M[x := N]) = \begin{cases} (\text{FV}(M) - \{x\}) \cup \text{FV}(N), & \text{jeśli } x \in \text{FV}(M); \\ \text{FV}(M), & \text{w przeciwnym przypadku.} \end{cases}$
2. Jeśli $x \notin \text{FV}(M)$ to $M[x := N] = M$.
3. $M[x := x] = M$.
4. Jeśli $\lambda x.P = \lambda y.Q$ to $P[x := y] = Q$ i $Q[y := x] = P$.

Łatwy dowód tego lematu pomijamy, udowodnimy za to następny.

Lemat 1.4 (o podstawieniu) Jeśli $x \neq y$ oraz albo $x \notin \text{FV}(R)$ albo $y \notin \text{FV}(M)$, to

$$M[x := N][y := R] = M[y := R][x := N[y := R]].$$

Dowód: Dowód jest przez indukcję ze względu na rozmiar M . Rozważamy różne przypadki, w zależności od postaci tego termu.

Przypadek 1: Jeżeli $M = x$, to $M[x := N][y := R] = x[x := N][y := R] = N[y := R] = x[x := N[y := R]] = x[y := R][x := N[y := R]]$.

Przypadek 2: Jeżeli $M = y$, to $M[x := N][y := R] = y[x := N][y := R] = y[y := R] = R = R[x := N[y := R]] = y[y := R][x := N[y := R]]$, bo wtedy $x \notin \text{FV}(R)$.

Przypadek 3: Jeżeli M jest jakąś inną zmienną z , to zarówno $M[x := N][y := R]$ jak też $M[y := R][x := N[y := R]]$ jest równe z .

Przypadek 4: Gdy $M = PQ$, to $M[x := N][y := R] = P[x := N][y := R]Q[x := N][y := R] = P[y := R][x := N[y := R]]Q[y := R][x := N[y := R]] = (PQ)[y := R][x := N[y := R]]$, na mocy założenia indukcyjnego.

Przypadek 5: Jeżeli M jest abstrakcją to można zakładać, że $M = \lambda z.Q$, gdzie $z \neq x, y$. Po lewej stronie mamy wtedy $\lambda z.Q[x := N][y := R]$ i z założenia indukcyjnego otrzymamy $M[x := N][y := R] = \lambda z.Q[y := R][x := N[y := R]] = (\lambda z.Q)[y := R][x := N[y := R]]$. ■

Wniosek 1.5 Jeśli $y \notin \text{FV}(M)$ to $M[x := y][y := R] = M[x := R]$.

Ćwiczenia

1. Jaki błąd popełniono w tej definicji lambda-termu:

$$M ::= x \mid MN \mid \lambda x.M \text{ ?}$$

2. Udowodnić lemat 1.3.

3. Napis $M(a|b)$ oznacza wyrażenie powstałe z M przez jednoczesną zamianę wszystkich wystąpień symbolu a w M na wystąpienia b i odwrotnie. Przez $\equiv_\alpha$ oznaczmy najmniejszą relacją równoważności w zbiorze lambda-wyrażeń, spełniającą następujące warunki:

- $\lambda x.M \equiv_\alpha \lambda y.M(x|y)$, dla $y \notin \text{FV}(M)$;
- jeżeli $M \equiv_\alpha M'$, to $\lambda x.M \equiv_\alpha \lambda x.M'$;
- jeżeli $M =_\alpha M'$ i $N =_\alpha N'$ to $MN =_\alpha M'N'$.

Udowodnić, że $M \equiv_\alpha n$ wtedy i tylko wtedy, gdy $M =_\alpha N$.

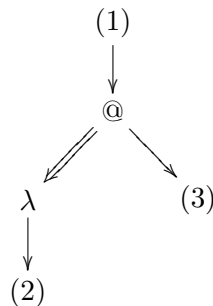
2 Redukcje

Relacja *beta-redukcji* to najmniejsza relacja w zbiorze lambda-termów, spełniająca warunki:

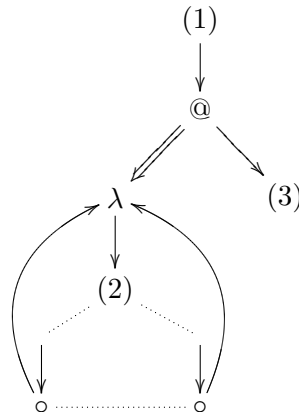
- $(\lambda x P)Q \rightarrow_\beta P[x := Q]$;
- jeśli $M \rightarrow_\beta M'$, to $MN \rightarrow_\beta M'N$, $NM \rightarrow_\beta NM'$ oraz $\lambda x M \rightarrow_\beta \lambda x M'$.

Inaczej mówiąc, $M \rightarrow_\beta M'$ zachodzi gdy podterm termu M postaci $(\lambda x P)Q$, czyli *β -redeks*, zostaje zamieniony na $P[x := Q]$. Znakiem $\rightarrow_\beta$ oznaczamy domknięcie przechodnio-zwrotne relacji $\rightarrow_\beta$, a znakiem $=_\beta$ najmniejszą relację równoważności zawierającą $\rightarrow_\beta$. Tę ostatnią nazywamy relacją *beta-równości*.

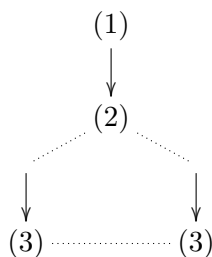
Z punktu widzenia grafowego, beta-redeks to krawędź od @ do λ :



Beta-redukcja polega na usunięciu tej krawędzi, tj. na „wyprostowaniu” drogi od (1) do (2). Każda krawędź wchodząca do λ (reprezentująca wystąpienie zmiennej związanej) zostaje przy tym skierowana do osobnej kopii termu zaczepionego w (3). Tj. podgraf postaci



zostaje zastąpiony podgrafem postaci:



Uwaga: może się zdarzyć, że lambda nie wiąże żadnego wystąpienia zmiennej (nie ma żadnej krawędzi prowadzącej do λ). Wtedy podgraf (3) znika. Odpowiada to redukcji postaci $(\lambda x.M)N \rightarrow_\beta M$, gdzie $x \notin \text{FV}(M)$.

Następujący łatwy lemat stwierdza, że β -redukcja jest zgodna z operacją podstawienia.

Lemat 2.1 *Jeśli $M \rightarrow_\beta M'$ i $N \rightarrow_\beta N'$ to $M[x := N] \rightarrow_\beta M'[x := N']$.*

Dowód: Należy pokazać dwa prostsze stwierdzenia:

- (1) Jeśli $M \rightarrow_\beta M'$, to $M[x := N] \rightarrow_\beta M'[x := N]$;
- (2) Jeśli $N \rightarrow_\beta N'$, to $M[x := N] \rightarrow_\beta M[x := N']$,

a następnie zastosować indukcję ze względu na liczbę kroków redukcji. Zarówno (1) jak (2) można udowodnić przez indukcję ze względu na długość termu M . Istotny przypadek w dowodzie części (1) zachodzi wtedy, gdy $M = (\lambda y.P)Q$ i $M' = P[y:=Q]$, dla pewnych P, Q . Przyjmując $y \notin \text{FV}(N)$, mamy $M[x:=N] = (\lambda y.P[x:=N])Q[x:=N] \rightarrow_\beta P[x:=N][y:=Q[x:=N]] = P[y:=Q][x:=N] = M'[x:=N]$, na mocy lematu o podstawieniu 1.4. Szczegóły pozostawiamy jako ćwiczenie. ■

Teoria λ

Rachunek lambda bez typów można przedstawić jako teorię równościową o aksjomatach i regułach jak na Obrazku 2. Jeśli w tym systemie można wyprowadzić równość $M = N$, to napiszemy $\lambda \vdash M = N$. Łatwo sprawdzić równoważność:

$$\lambda \vdash M = N \text{ wtedy i tylko wtedy, gdy } M =_\beta N.$$

$(\beta) (\lambda x.M)N = M[x := N]$	$x = x$
$\frac{M = N}{MP = NP}$	$\frac{M = N}{PM = PN}$
$(\xi) \frac{M = N}{\lambda x.M = \lambda x.N}$	
$\frac{M = N}{N = M}$	$\frac{M = N, N = P}{M = P}$

Obrazek 2: Rachunek lambda jako teoria równościowa

Jednakże dla nas ważne są też własności beta-redukcji, a nie tylko beta-równości. Term postaci $(\lambda x.P)Q$ nazywamy *beta-redeksem*. Jeśli w termie M nie występuje żaden beta-redek, to dla żadnego N nie zachodzi $M \rightarrow_\beta N$. Mówimy wtedy, że M jest *w postaci β -normalnej*, lub po prostu *w postaci normalnej*. Nietrudno zauważyć, że postaci normalne to dokładnie termy kształtu $\lambda \vec{x}.y\vec{N}$, gdzie $\vec{N}$ są normalne.

Przykład: Następujące termy są w postaci normalnej:

$$\begin{aligned}\mathbf{I} &= \lambda x.x \\ \mathbf{K} &= \lambda xy.x \\ \mathbf{S} &= \lambda xyz.xz(yz) \\ \omega &= \lambda x.xx\end{aligned}$$

Natomiast term $\Omega = \omega\omega$ nie jest w postaci normalnej, bo $\Omega \rightarrow_\beta \Omega$.

Jeśli $M \rightarrow_\beta N$ i N jest w postaci normalnej, to mówimy, że M ma postać normalną, i że N jest postacią normalną termu M . Mówimy, że term M ma własność *silnej normalizacji* (jest *silnie normalizowalny*) jeżeli nie istnieje nieskończony ciąg redukcji $M = M_0 \rightarrow_\beta M_1 \rightarrow_\beta M_2 \rightarrow_\beta \dots$. Oczywiście term silnie normalizowalny musi mieć postać normalną.

Jeszcze jedno przydatne czasem oznaczenie. Piszemy $M \downarrow_\beta N$, gdy istnieje taki term P , że $M \rightarrow_\beta P_\beta \leftarrow N$.

Głównym przedmiotem naszego zainteresowania jest beta-redukcja. Dlatego np. symbole $\rightarrow$ i $\rightarrow$ domyślnie¹ oznaczają $\rightarrow_\beta$ i $\rightarrow_\beta$. Ale czasami rozważamy także relację *eta-redukcji*. Jest to najmniejsza relacja $\rightarrow_\eta$, spełniająca warunki:

- $\lambda x.Mx \rightarrow_\eta M$, gdy $x \notin \text{FV}(M)$;
- jeśli $M \rightarrow_\eta M'$, to $MN \rightarrow_\eta M'N$, $NM \rightarrow_\eta NM'$ oraz $\lambda xM \rightarrow_\eta \lambda xM'$.

Symbol $\rightarrow_{\beta\eta}$ oznacza sumę tych relacji. Używamy odpowiednio notacji np. $\rightarrow_{\beta\eta}$, $=_{\beta\eta}$, mówimy o „postaciach η -normalnych” itd.

Jeżeli do aksjomatów i reguł równościowych rozważanych powyżej dodamy nowy aksjomat

$$(\eta) \quad \lambda x.Mx = M,$$

dla $x \notin \text{FV}(M)$, to oczywiście otrzymamy system wnioskowania, pozwalający na wyprowadzenie równości $M = N$ wtedy i tylko wtedy gdy $M =_{\beta\eta} N$. Mniej oczywiste jest to, że aksjomat (η) można równoważnie zastąpić przez następującą *regulę ekstensjonalności*:

$$(\text{ext}) \quad \frac{Mx = Nx}{M = N} \quad (x \notin \text{FV}(M) \cup \text{FV}(N))$$

Reguła ekstensjonalności wyraża taką myśl: funkcje, które dla tych samych argumentów dają te same wyniki, są równe.

Dygresja: Zauważmy, że beta- i eta-redukcja są w pewnym sensie dualne do siebie. Beta-redeks to term, w którym najpierw utworzyliśmy funkcję, tj. *wprowadziliśmy* abstrakcję (czyli zapytanie o argument, prompt), a potem ją *eliminujemy* przez dostarczenie argumentu. Eta-redeks odpowiada odwrotnej kolejności: najpierw eliminujemy prompt, dostarczając zmiennej jako argumentu, a potem wprowadzamy abstrakcję ze względu na tę zmienną.

¹W odniesieniu do znaku równości nie stosujemy tej konwencji (ale spotyka się to często w literaturze).

Ćwiczenia

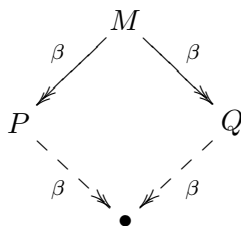
1. Napisać program, który drukuje własny tekst. (Nie wolno odwoływać się do nazwy pliku, miejsca w pamięci itp.) Zadanie łatwe w Lispie, trudne w Pascalu.
2. Udowodnić, że jeśli $(\lambda xP)Q = (\lambda xP')Q'$ to $P[x := Q] = P'[x := Q']$. *Wskazówka:* Użyć lematu 1.3 i wniosku 1.5.
3. Udowodnić, że aksjomat (η) jest równoważny regule (ext) , tj. te same równości można wyprowadzić w systemie z Obrazka 2 rozszerzonym o (η) i w systemie rozszerzonym o regułę (ext) .

3 Twierdzenie Churcha-Rossera

W jednym termie może występować więcej niż jeden redeks. Można więc go redukować na różne sposoby. Niedobrze jednak byłoby, gdyby te różne ciągi redukcji prowadziły do nieodwracalnie różnych rezultatów. Na szczęście mamy *twierdzenie Churcha-Rossera*.

Twierdzenie 3.1 *Jeśli $P \xleftarrow{\beta} M \xrightarrow{\beta} Q$, to $P \downarrow_{\beta} Q$.*

Treść twierdzenia Churcha-Rossera przedstawiamy za pomocą diagramu (Obrazek 3).



Obrazek 3: Twierdzenie Churcha-Rossera

Inaczej mówiąc, twierdzenie to mówi, że relacja $\rightarrow_{\beta}$ ma tzw. *własność rombu*, która niestety nie zachodzi dla relacji $\rightarrow_{\beta}$. Jako przykład można podać term $M = (\lambda x.xx)((\lambda x.x)y)$.

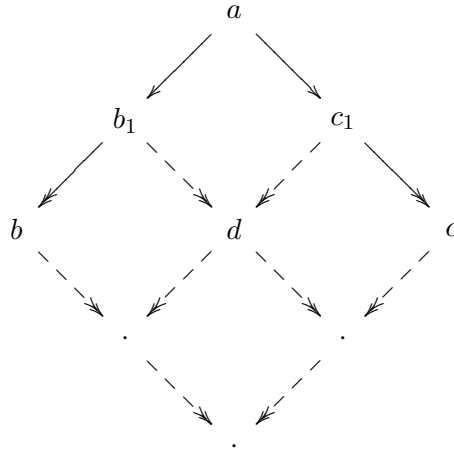
Zanim udowodnimy twierdzenie Churcha-Rossera, zauważmy, że zjawisko, o którym mówimy, może dotyczyć każdej relacji dwuargumentowej. Poniżej przyjmujemy taką konwencję: podwójna strzałka zawsze oznacza domknięcie przechodnio-zwrotne relacji oznaczonej strzałką pojedynczą.

Mówimy, że relacja $\rightarrow$ w zbiorze A ma *własność Churcha-Rossera* (CR), jeżeli dla dowolnych elementów $a, b, c \in A$, takich że $b \leftarrow a \rightarrow c$, istnieje takie d , że $b \rightarrow d \leftarrow c$. Relacja $\rightarrow$ ma *słabą własność Churcha-Rossera* (WCR), jeżeli $b \leftarrow a \rightarrow c$ implikuje $b \rightarrow d \leftarrow c$ dla pewnego d . Oczywiście własność rombu implikuje CR (ale nie na odwrót), a własność CR implikuje WCR. Mniej oczywiste jest to, że własność CR nie wynika z WCR. Najprostszy przykład jest taki:

$$\bullet \leftarrow \bullet \longleftrightarrow \bullet \rightarrow \bullet$$

Jeśli jednak relacja $\rightarrow$ ma *własność silnej normalizacji* (SN), tj. nie istnieje nieskończony ciąg postaci $a_0 \rightarrow a_1 \rightarrow a_2 \rightarrow \dots$, to już jest dobrze.

Fakt 3.2 (Lemat Newmana) *Jednoczesne zachodzenie WCR i SN implikuje CR.*



Obrazek 4: Rysunek do lematu Newmana

Dowód: Jeśli relacja $\rightarrow$ ma własność SN to relacja $\leftarrow$ jest dobrym ufundowaniem zbioru A . Przez indukcję ze względu na porządek $\leftarrow$ dowodzimy, że każdy element $a \in A$ ma własność:

Dla dowolnych b, c , jeśli $b \leftarrow a \rightarrow c$, to $b \downarrow c$.

Jeśli $a = b$ lub $a = c$ to teza jest oczywista. Załóżmy więc, że $b \leftarrow b_1 \leftarrow a \rightarrow c_1 \rightarrow c$.

Na mocy WCR jest takie d , że $b_1 \rightarrow d \leftarrow c_1$. Z założenia indukcyjnego, zastosowanego do b_1 i c_1 mamy więc $b \downarrow d \downarrow c_1$ i możemy zastosować założenie indukcyjne dla d . ■

Dowód twierdzenia Churcha-Rossera

Dla dowodu zdefiniujemy pewną pomocniczą relację. Będzie to relacja $\xrightarrow{1}$, określona jako najmniejsza relacja na termach, która spełnia następujące warunki:

- $x \xrightarrow{1} x$, gdy x jest zmienną;
- jeśli $M \xrightarrow{1} M'$, to $\lambda x M \xrightarrow{1} \lambda x M'$;
- jeśli $M \xrightarrow{1} M'$ i $N \xrightarrow{1} N'$, to $MN \xrightarrow{1} M'N'$, oraz $(\lambda x M)N \xrightarrow{1} M'[x := N']$.

Relacja $\xrightarrow{1}$ odpowiada jednoczesnej redukcji kilku beta-redeksów.

Lemat 3.3

(1) Dla dowolnego M zachodzi $M \xrightarrow{1} M$.

(2) Jeśli $M \xrightarrow{1} M'$ i $N \xrightarrow{1} N'$, to $M[x := N] \xrightarrow{1} M'[x := N']$.

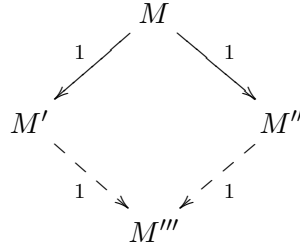
Dowód:

Dowód części (1) jest przez indukcję ze względu na budowę M , a dowód części (2) przez indukcję ze względu na wyprowadzenie $M \xrightarrow{1} M'$. W dowodzie punktu (2) nieoczywisty jest

przypadek, gdy $M = (\lambda y P)Q \xrightarrow{1} P'[y := Q']$, gdzie $P \xrightarrow{1} P'$ oraz $Q \xrightarrow{1} Q'$. Zakładając, że zmienna związana y jest tak wybrana, że $y \notin FV(N)$, i korzystając z założenia indukcyjnego o $P \xrightarrow{1} P'$ i $Q \xrightarrow{1} Q'$, mamy wtedy na mocy lematu o podstawieniu:

$$M[x := N] = (\lambda y P[x := N])Q[x := N] \xrightarrow{1} P'[x := N'][y := Q'[x := N']] = P'[y := Q'] [x := N']. \quad \blacksquare$$

Uwaga: Użyty w powyższym dowodzie zwrot „dowód przez indukcję ze względu na wyprowadzenie $M \xrightarrow{1} M'$ ” należy ściśle rozumieć tak: Zbiór par termów $\{(M, M') \mid M \xrightarrow{1} M'\}$ przedstawiamy jako sumę wstępującego ciągu zbiorów X_i , gdzie $X_0 = \{(x, x) \mid x \text{ jest zmienną}\}$, oraz X_{i+1} powstaje z X_i przez dodanie wszystkich par $(MN, M'N')$, $((\lambda x M)N, M'[x := N'])$ oraz $(\lambda x M, \lambda x M')$, dla dowolnych (M, M') i (N, N') , które są już w X_i . Indukcję tak naprawdę prowadzimy ze względu na najmniejsze takie i , że $(M, M') \in X_i$.

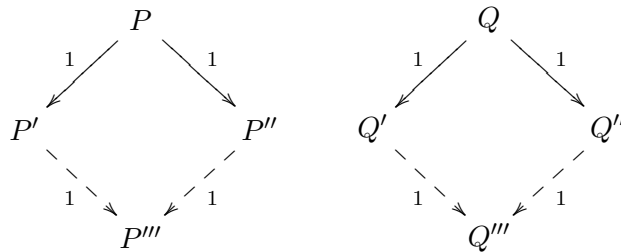


Obrazek 5: Lemat 3.4

Lemat 3.4 Relacja $\xrightarrow{1}$ ma własność rombu, tj. jeśli $M \xrightarrow{1} M'$ i $M \xrightarrow{1} M''$, to dla pewnego termu M''' zachodzi $M' \xrightarrow{1} M''' \xleftarrow{1} M''$.

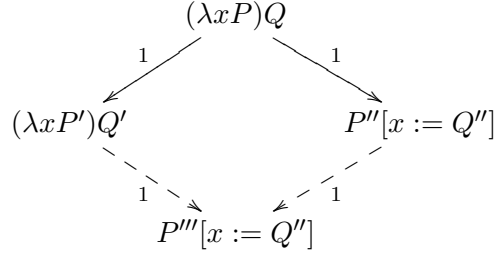
Dowód: Dowód jest przez indukcję ze względu na wyprowadzenie $M \xrightarrow{1} M'$. Większość przypadków jest łatwa. Niech na przykład $M = \lambda x P \xrightarrow{1} \lambda x P' = M'$, gdzie $P \xrightarrow{1} P'$. Wtedy term M'' musi być postaci $\lambda x P''$ i należy skorzystać z założenia indukcyjnego o $P \xrightarrow{1} P'$.

Mniej trywialny jest np. przypadek gdy $M = (\lambda x P)Q$ oraz $M' = (\lambda x P')Q'$ a $M'' = P''[x := Q'']$. Wtedy $P \xrightarrow{1} P'$, P'' i $Q \xrightarrow{1} Q'$, Q'' . Z założenia indukcyjnego dla $P \xrightarrow{1} P'$ i $Q \xrightarrow{1} Q'$ dostajemy



Aby z tego otrzymać tezę, należy skorzystać z lematu 3.3(2), jak na Obrazku 6. $\blacksquare$

Dowód twierdzenia 3.1: Ponieważ relacja $\xrightarrow{1}$ ma własność rombu, więc tym bardziej jej domknięcie przechodnio-zwrotne $\xrightarrow{1}^*$ ma własność rombu.



Obrazek 6: Dowód lematu 3.4.

Zauważmy teraz, że mamy następujące zawierania pomiędzy relacjami:

$$\rightarrow_{\beta} \subseteq \xrightarrow{1} \subseteq \rightarrow_{\beta}.$$

Stąd łatwo wywnioskować, że relacje $\xrightarrow{1}$ i $\rightarrow_{\beta}$ są równe, a więc $\rightarrow_{\beta}$ ma własność rombu. ■

Skoro udowodniliśmy twierdzenie Churcha-Rossera, zobaczmy jakie ma ono konsekwencje.

Wniosek 3.5

0) Jeśli $M \downarrow_{\beta} N \downarrow_{\beta} Q$, to $M \downarrow_{\beta} Q$.

1) Jeśli $M =_{\beta} N$, to $M \downarrow_{\beta} N$.

2) Każdy term ma co najwyżej jedną postać normalną.

3) Rachunek lambda jest niesprzeczny jako teoria równościowa: nie można np. wyprowadzić równości $x = y$, ponieważ $x \neq_{\beta} y$.

Dowód: Jedyna nieoczywista część to (1). Jeśli $M =_{\beta} N$, to mamy taki ciąg:

$$M = M_0 \leftarrow Q_1 \rightarrow M_1 \leftarrow Q_2 \rightarrow M_2 \leftarrow \dots \rightarrow M_n = N$$

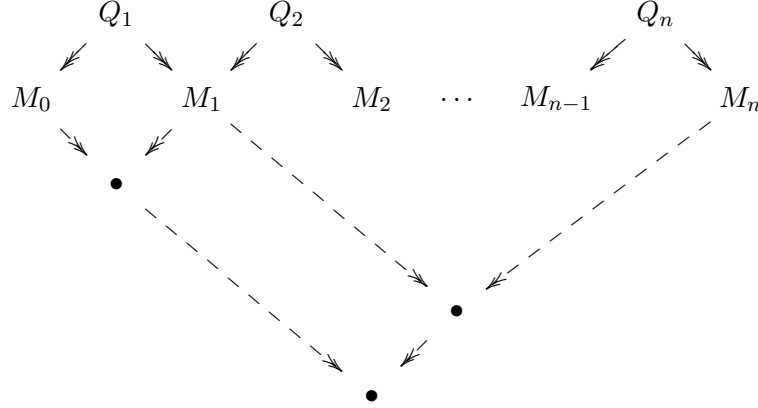
Dowodzimy $M \downarrow N$, przez indukcję ze względu na n . Jeśli $n = 0$, to sprawa jest oczywista, w przeciwnym razie z założenia indukcyjnego mamy $M_1 \downarrow_{\beta} N$. Ale $M_0 \downarrow_{\beta} M_1$ na mocy CR, więc $M \downarrow_{\beta} N$ wynika z części (0). Zob. Obrazek 7. ■

A zatem sens twierdzenia Churcha-Rossera można wyrazić tak. Chociaż term może być redukowany na wiele sposobów, to wszystkie te sposoby są zgodne. Każde obliczenie zakończone sukcesem (uzyskaniem postaci normalnej) daje ten sam wynik.

Ćwiczenia

1. Udowodnić następujące wzmocnienie lematu 3.4:

Dla dowolnego M istnieje taki term $M^{\bullet}$, że z warunku $M \xrightarrow{1} M'$ wynika $M \xrightarrow{1} M^{\bullet}$.



Obrazek 7: Dowód wniosku 3.5(1)

4 Eta-redukcja

Każdy krok η -redukcji $\lambda x. Mx \rightarrow_\eta M$ zmniejsza długość termu, zatem zachodzi

Fakt 4.1 *Relacja η -redukcji ma własność silnej normalizacji.*

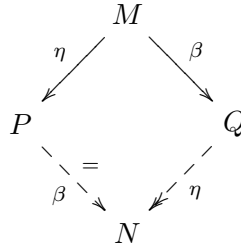
Aby wywnioskować, że $\rightarrow_\eta$ ma własność Churcha-Rossera, wystarczy więc pokazać WCR. W istocie mamy prawie² własność rombu, co można łatwo pokazać, badając możliwe przypadki wzajemnego położenia redeksów.

Lemat 4.2 *Domknięcie zwrotne $\rightarrow_\eta^=$ relacji $\rightarrow_\eta$ ma własność rombu.*

Wniosek 4.3 *Relacja η -redukcji ma własność Churcha-Rossera.*

Pokażemy teraz, że własność Churcha-Rossera zachodzi także dla relacji $\rightarrow_{\beta\eta}$.

Lemat 4.4 *Jeśli $P \eta \leftarrow M \rightarrow_\beta Q$ to istnieje takie N , że $P \rightarrow_\beta^= N \eta \leftarrow Q$.*

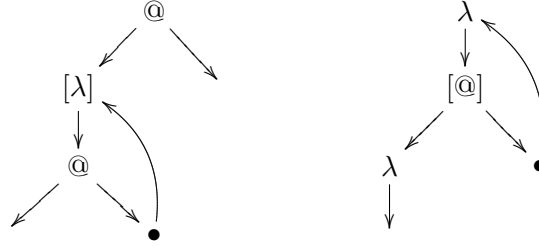


Dowód: Badając wzajemne położenie beta- i eta-redeksów w grafie termu M (Obrazek 8), widzimy że kolejność, w której redukujemy te redeksy, na ogół nie ma znaczenia i rezultat



Obrazek 8: Eta-redeks i beta-redeks

jest taki sam. Pierwszy wyjątek, to sytuacja gdy eta-redeks znajduje się w argumentie beta-redeksu i wskutek redukcji może zostać usunięty lub powielony. Dlatego w treści lematu po prawej stronie mamy strzałkę podwójną $\eta \leftarrow$ a nie pojedynczą. Pozostałe dwa wyjątki zachodzą wtedy, gdy nasze dwa redeksy wykorzystują ten sam wierzchołek grafu (lambdę lub aplikację). Ma to miejsce (Obrazek 9) w podtermach postaci $(\lambda x.Mx)N$, gdzie $x \notin \text{FV}(M)$ oraz postaci $\lambda x(\lambda y M)x$, gdzie $x \notin \text{FV}(\lambda y M)$. Szczęśliwie w obu przypadkach redukcja

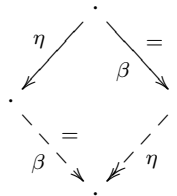


Obrazek 9: Pary krytyczne dla beta-eta-redukcji

każdego z konkurujących redeksów daje ten sam rezultat. ■

Sytuacja zilustrowana na Obrazku 9 może dotyczyć także innych systemów redukcyjnych. Jeśli redukcja każdego z dwóch redeksów wymaga zużycia tego samego „zasobu”, to mówimy, że redeksy te tworzą *parę krytyczną*. Jeśli każdą parę krytyczną można „uzgodnić” to dana relacja redukcji ma słabą własność Churcha-Rossera.³

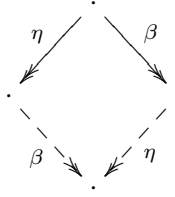
Z lematu 4.4 wynika łatwo, że mamy następujący diagram:



²Własność rombu nie zachodzi bo $y \eta \leftarrow \lambda x. yx \rightarrow_{\eta} y$, tymczasem nie ma termu, do którego y redukowałby się w jednym kroku.

³Dla systemów o własności SN i bez par krytycznych mamy więc ogólną metodę dowodzenia własności Churcha-Rossera: zastosować lemat Newmana 3.2.

Z tego diagramu otrzymamy następny. Mówi on, że relacje $\rightarrow_\beta$ i $\rightarrow_\eta$ są *przemienne*:



Ponieważ obie relacje mają własność Churcha-Rossera, nietrudno teraz pokazać

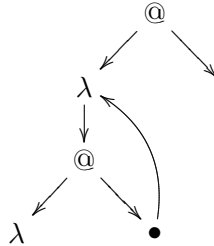
Twierdzenie 4.5 *Relacja $\rightarrow_{\beta\eta}$ ma własność Churcha-Rossera.*

Beta-eta-redukcja

Własności $\beta\eta$ -redukcji są często podobne do własności β -redukcji. Wymienimy tu kilka z nich, pomijając większość dowodów. Zaczniemy od technicznego lematu.

Lemat 4.6 *Niech $M \rightarrow_\eta N \rightarrow_\beta P$ ale $M \not\rightarrow_\beta N$. Wtedy istnieje takie Q , że $M \rightarrow_\beta Q \rightarrow_\eta P$.*

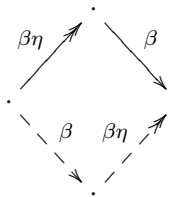
Dowód: Zauważmy, że jedyna sytuacja, w której eta-redukcja może wykreować nowy beta-redeks wygląda tak, jak na Obrazku 10. To jest akurat sytuacja zabroniona, bo zamiast eta-



Obrazek 10: Eta-redukcja tworzy beta-redeks

redeksu możemy z tym samym skutkiem zredukować beta-redeks. W pozostałych przypadkach beta-redeks redukowany w drugim kroku już istnieje w termie M i może zostać zredukowany zanim wykonana będzie eta-redukcja. Potem dopiero wykonamy tyle eta-redukcji ile kopii naszego redeksu (być może zero) występuje w Q . ■

Wniosek 4.7 *Jeśli $M \rightarrow_{\beta\eta} P \rightarrow_\beta N$ to $M \rightarrow_\beta Q \rightarrow_{\beta\eta} N$, dla pewnego Q . Graficznie:*



Wniosek 4.8 *Term ma nieskończoną β -redukcję wtedy i tylko wtedy gdy ma nieskończoną $\beta\eta$ -redukcję. (Inaczej: term ma własność β -SN wtedy i tylko wtedy, gdy ma własność $\beta\eta$ -SN.)*

Dowód: Implikacja z lewej do prawej jest oczywiście oczywista. Załóżmy więc, że mamy nieskończoną $\beta\eta$ -redukcję. Jeśli beta-kroki występują w tej redukcji nieskończenie wiele razy, to wniosek 4.7 pozwala z niej uzyskać nieskończoną β -redukcję (poprzez „przesuwanie” beta-kroków w lewo). A nie może być tak, że prawie wszystkie kroki są typu eta. ■

Pokażemy teraz, że każdy ciąg beta-eta redukcji można „posortować” tak aby eta-redukcje były na końcu.

Twierdzenie 4.9 (o odkładaniu η -redukcji) *Jeśli $M \twoheadrightarrow_{\beta\eta} N$ to $M \twoheadrightarrow_{\beta} P \twoheadrightarrow_{\eta} N$, dla pewnego P .*

Dowód: Niech $\succ$ oznacza najmniejszą relację w zbiorze termów, spełniającą warunki:

- $x \succ x$, dla dowolnej zmiennej x ;
- jeśli $A \succ A'$ oraz $x \notin \text{FV}(A)$, to $\lambda x.Ax \succ A'$;
- jeśli $A \succ A'$, to $\lambda x.A \succ \lambda x.A'$;
- jeśli $A \succ A'$ oraz $B \succ B'$, to $AB \succ A'B'$.

Relacja $\succ$ mniej więcej tak się ma do eta-redukcji, jak relacja $\xrightarrow{1}$ do beta-redukcji. Jej najważniejsza własność jest taka:

$$\text{Jeśli } A \succ B \rightarrow_{\beta} C \text{ to istnieje taki term } D, \text{ że } A \twoheadrightarrow_{\beta} D \succ C. \quad (*)$$

Niech teraz $M \twoheadrightarrow_{\beta\eta} N$. Każdy krok eta-redukcji można uważać za krok typu $\succ$ i „przesuwać” w prawo z pomocą (*). Szczegóły pozostawiamy jako ćwiczenie 3. ■

Twierdzenie 4.10 *Term ma postać beta-normalną wtedy i tylko wtedy, gdy ma postać beta-eta-normalną.*

Dowód: Łatwiejsza jest implikacja z lewej do prawej. Wystarczy zauważyć, że eta-redukowanie postaci beta-normalnej nie tworzy nowych beta-redeksów (ćwiczenie 2), musi więc w końcu dać postać $\beta\eta$ -normalną.

W dowodzie implikacji odwrotnej korzysta się z twierdzenia o odkładaniu eta-redukcji. Szczegóły tego dowodu pozostawiamy jako ćwiczenie 4. ■

Ćwiczenia

1. W jaki sposób beta-redukcja może spowodować powstanie nowego beta-redeksu? Eta-redeksu?
2. Pokazać, że jeśli M jest w postaci β -normalnej oraz $M \rightarrow_{\eta} N$ to N jest w postaci β -normalnej.

3. Udowodnić następujące własności relacji $\succ$, określonej w dowodzie twierdzenia 4.9:

- (a) $\rightarrow_\eta \not\subseteq \succ \not\subseteq \rightarrow_\eta$;
- (b) Jeśli $A \succ A'$ oraz $B \succ B'$, to $A[x := B] \succ A'[x := B']$.

Wynioskować z nich własność (*) o której mowa w dowodzie twierdzenia i uzupełnić dowód twierdzenia.

4. Udowodnić część ($\Leftarrow$) twierdzenia 4.10. *Wskazówka: Pokazać najpierw, że jeśli $M \succ N$ oraz N ma postać beta-normalną to M też ma (indukcja ze względu na liczbę kroków potrzebną do zredukowania N do postaci normalnej). Następnie skorzystać z twierdzenia 4.9 i z ćwiczenia 3(a).*

5 Standaryzacja

W jednym termie może występować więcej niż jeden redeks. Możliwe są więc różne *strategie redukcji*. Jedną z nich polega na wybieraniu zawsze tego redeksu, który zaczyna się najbar-dziej na lewo. Mówimy wtedy o redukcji *lewostronnej*. Okazuje się, że strategia redukcji lewostronnej jest strategią *normalizującą*, tj. tą metodą zawsze dojdziemy do postaci normalnej, jeśli ona istnieje. Popatrzmy na przykład na term $\mathbf{K}z\Omega$. W zależności od wyboru strategii redukcji możemy dojść do postaci normalnej lub zredukować ten term w nieskończoność.

W istocie prawdziwe jest nieco silniejsze twierdzenie, zwane *twierdzeniem o standaryzacji*. Powiemy, że wyprowadzenie

$$M_0 \rightarrow M_1 \rightarrow \dots \rightarrow M_n$$

jest *standardowe*, jeżeli dla dowolnego $i = 0, \dots, n-2$, w kroku $M_i \rightarrow M_{i+1}$ redukujemy redeks położony nie dalej od początku termu niż redeks, który redukujemy w następnym kroku $M_{i+1} \rightarrow M_{i+2}$. Na przykład ta redukcja jest standardowa:

$$\begin{aligned} (\lambda x.xx)((\lambda zy.z(zy))(\lambda u.u)(\lambda w.w)) &\rightarrow (\lambda x.xx)((\lambda y.(\lambda u.u)((\lambda u.u)y))(\lambda w.w)) \rightarrow \\ &\rightarrow (\lambda x.xx)((\lambda u.u)((\lambda u.u)(\lambda w.w))) \rightarrow (\lambda x.xx)((\lambda u.u)(\lambda w.w)) \rightarrow (\lambda x.xx)(\lambda w.w). \end{aligned}$$

a ta nie jest:

$$\begin{aligned} (\lambda x.xx)((\lambda zy.z(zy))(\lambda u.u)(\lambda w.w)) &\rightarrow (\lambda x.xx)((\lambda y.(\lambda u.u)((\lambda u.u)y))(\lambda w.w)) \rightarrow \\ &(\lambda x.xx)((\lambda y.(\lambda u.u)y)(\lambda w.w)) \rightarrow (\lambda x.xx)((\lambda y.y)(\lambda w.w)) \rightarrow (\lambda x.xx)(\lambda w.w). \end{aligned}$$

Twierdzenie 5.1 *Jeśli $M \rightarrow_\beta N$, to istnieje standardowa redukcja z M do N .*

Zajmiemy się teraz dowodem twierdzenia 5.1. Zaczniemy od tego, że każdy term M jest jednej z dwóch możliwych postaci:

- 1) $\lambda \vec{x}.z\vec{R}$;
- 2) $\lambda \vec{x}.(\lambda y.P)Q\vec{R}$,

gdzie długość wektora $\vec{x}$ może być zerem, a z może występować wśród zmiennych $\vec{x}$ lub nie. W przypadku (1) mówimy, że term jest w *czołowej postaci normalnej*⁴. Natomiast w przypadku (2) mówimy, że term ma *redeks czołowy* $(\lambda y.P)Q$. Krok redukcji

$$M = \lambda \vec{x}.(\lambda y.P)Q\vec{R} \rightarrow_\beta \lambda \vec{x}.P[y := Q]\vec{R} = N$$

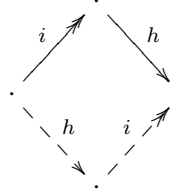
⁴Uwaga: czołowa postać normalna nie jest na ogół postacią normalną.

nazywamy wtedy *redukcją czołową* i zapisujemy $M \xrightarrow{h} N$. Inne kroki redukcji nazywamy *wewnętrzными*, a w takim przypadku używamy symbolu $\xrightarrow{i}$. Kluczowy lemat mówi, że redukcje czołowe można wykonać przed wewnętrznymi:

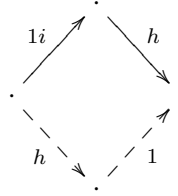
Lemat 5.2 *Jeżeli $M \twoheadrightarrow_{\beta} N$, to $M \xrightarrow{h} P \xrightarrow{i} N$, dla pewnego P .*

Jeśli udowodnimy ten lemat, to pozostaje tylko zauważyć, że po wykonaniu wszystkich redukcji czołowych kształt termu się częściowo „ustala” tj. mamy albo $M \xrightarrow{h} \lambda \vec{x}.z.\vec{R}$ albo $M \xrightarrow{h} \lambda \vec{x}.(\lambda y.P)Q\vec{R}$ i dalsze redukcje odbywają się wewnątrz P , Q i $\vec{R}$. Stosując indukcję do termów P , Q i $\vec{R}$, otrzymujemy standardowe redukcje, które wykonujemy „od lewej” tj. najpierw w P , potem w Q i potem w kolejnych wyrazach ciągu $\vec{R}$.

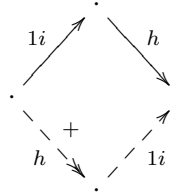
Naiwna próba dowodu lematu 5.2 mogłaby być taka:



Niestety redukcja $(\lambda x P)Q \xrightarrow{i} (\lambda x P')Q' \xrightarrow{h} P'[x := Q']$ może nie dać się zastąpić przez $(\lambda x P)Q \xrightarrow{h} P[x := Q] \xrightarrow{i} P'[x := Q']$, bo kroki wewnątrz P lub wewnątrz Q mogą być czołowe. Idea dowodu polega na użyciu redukcji postaci $\xrightarrow{1}$, które łatwiej rozbić na część czołową i wewnętrzną. Zaczniemy od takiego (łatwego) diagramu



gdzie $\xrightarrow{1i}$ oznacza ciąg redukcji wewnętrznych, które jednocześnie stanowią redukcję postaci $\xrightarrow{1}$. Jeśli teraz uda nam się rozbić $\xrightarrow{1}$ na fazę czołową i fazę postaci $\xrightarrow{1i}$, to dostaniemy diagram



A zatem chcemy udowodnić, że

$$\text{Jeśli } M \xrightarrow{1} N \text{ to } M \xrightarrow{h} L \xrightarrow{1i} N, \text{ dla pewnego } L, \quad (*)$$

i natychmiast mamy kłopot, gdy M jest redeksem $(\lambda x P)Q$. Chcemy bowiem wywnioskować, że $P[x := Q] \xrightarrow{\text{li}} P'[x := Q']$, a wiemy tylko, że $P \xrightarrow{1} P'$ i $Q \xrightarrow{1} Q'$, skąd mamy zaledwie $P[x := Q] \xrightarrow{1} P'[x := Q']$. Potrzebna jest relacja $\Rightarrow$, która ma taką własność:

$$\text{Jeśli } M \Rightarrow M' \text{ oraz } N \Rightarrow N' \text{ to } M[x := N] \Rightarrow M'[x := N'], \quad (**)$$

rozkłada się na $\xrightarrow{h}$ i $\xrightarrow{\text{li}}$, i na dodatek zawiera relację $\xrightarrow{1}$. Definiujemy ją tak: $M \Rightarrow N$ zachodzi wtedy gdy $M \xrightarrow{h} P \xrightarrow{\text{li}} N$, oraz każdy term Q występujący w redukcji $M \xrightarrow{h} P$ spełnia warunek $Q \xrightarrow{1} N$.

Należy teraz udowodnić własność (**). Z jej pomocą przez indukcję ze względu na definicję relacji $\xrightarrow{1}$ dowodzi się, że $\xrightarrow{1} \subseteq \Rightarrow$ a teraz (*) jest już oczywiste.

Udowodniliśmy więc twierdzenie o standaryzacji. A oto najważniejszy wniosek z niego.

Wniosek 5.3 *Jeśli M ma postać normalną N , to istnieje lewostronna redukcja z M do N .*

Powyższy wniosek⁵ można odczytać tak: Jeśli istnieje nieskończona redukcja lewostronna zaczynająca się od M , to term M nie ma postaci normalnej. Można to stwierdzenie wzmocnić w taki sposób. Powiemy, że nieskończony ciąg redukcji jest *quasi-lewostronny*, jeśli nieskończenie wiele razy następuje w tym ciągu redukcja redeksu położonego najbardziej na lewo.

Fakt 5.4 *Jeśli istnieje nieskończony quasi-lewostronny ciąg redukcji zaczynający się od M , to term M nie ma postaci normalnej.*

Dowód: Na potrzeby tego dowodu powiedzmy, że term jest *wytrzymały*, gdy ma nieskończoną redukcję quasi-lewostronną, w której występuje nieskończenie wiele kroków postaci $\xrightarrow{h}$. Zauważmy, że jeśli N jest wytrzymały, to $N \xrightarrow{h}^+ N'$ dla pewnego wytrzymałego termu N' (ćwiczenie 4).

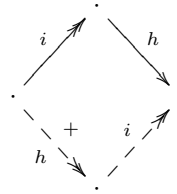
Założmy teraz, że M ma postać normalną. Przez indukcję ze względu na jej długość udowodnimy, że M nie może mieć nieskończonej redukcji quasi-lewostronnej. Przypuśćmy najpierw, że M jest termem wytrzymałym. Używając ćwiczenia 4 można wtedy zdefiniować ciąg wytrzymałych termów N_i , takich że $N_0 = M$ oraz $N_i \xrightarrow{h}^+ N_{i+1}$ dla dowolnego $i \in \mathbb{N}$. Inaczej mówiąc, „wyciągając na początek” redukcje czołowe otrzymamy nieskończoną redukcję czołową (a więc lewostronną) termu M . Nie może więc istnieć postać normalna.

Zatem jeśli $M = M_0 \rightarrow M_1 \rightarrow M_2 \rightarrow \dots$ jest nieskończoną redukcją quasi-lewostronną, to od pewnego miejsca mamy tylko redukcje wewnętrzne postaci $\lambda \vec{z}.y\vec{P} \xrightarrow{i} \lambda \vec{z}.y\vec{P}' \xrightarrow{i} \lambda \vec{z}.y\vec{P}'' \xrightarrow{i} \dots$. Spośród redukcji lewostronnych występujących w tym ciągu, nieskończenie wiele dotyczy tej samej składowej wektora $\vec{P}$, więc wystarczy użyć założenia indukcyjnego. ■

⁵Ten fakt bywa nazywany *twierdzeniem o normalizacji*. Obecnie rzadziej używa się tej nazwy, bo „normalizacja” zwykle oznacza istnienie postaci normalnych.

Ćwiczenia

1. Jaki jest parametr indukcji w dowodzie twierdzenia 5.1?
2. Zdefiniować poprawnie relację $\xrightarrow{1i}$. Dlaczego definicja $\xrightarrow{1i}$ jako iloczynu $\xrightarrow{1}$ i $\xrightarrow{i}$ nie jest poprawna?
3. Zdefiniować poprawnie relację $\Rightarrow$ i uzupełnić brakujące fragmenty dowodu twierdzenia 5.1.
4. Term nazwiemy *wytrzymałym*, gdy ma nieskończoną redukcję quasi-lewostronną, w której występuje nieskończenie wiele kroków postaci $\xrightarrow{h}$. Udowodnić, że jeśli N jest wytrzymały, to $N \xrightarrow{h}^+ N'$ dla pewnego wytrzymałego termu N' . *Wskazówka: Korzystając z twierdzenia o standaryzacji pokazać poprawność diagramu*



6 Rachunek $\lambda\mathbf{I}$

Term M jest nazywany $\lambda\mathbf{I}$ -termem, jeżeli nie występuje w nim jako podterm żadna abstrakcja postaci $\lambda x.N$, gdzie $x \notin \text{FV}(N)$. To znaczy, że każda abstrakcja wiąże przynajmniej jedno wystąpienie zmiennej. Nietrudno zauważyć, że jeśli M jest $\lambda\mathbf{I}$ -termem, oraz $M \rightarrow_\beta M'$, to M' też jest $\lambda\mathbf{I}$ -termem (choć nie na odwrót). Dlatego można mówić o *rachunku $\lambda\mathbf{I}$* , w którym rozważa się tylko takie termy. Pełny rachunek lambda nazywany też bywa rachunkiem $\lambda\mathbf{K}$. Ta terminologia bierze się z tradycyjnych oznaczeń:

$$\begin{aligned}\mathbf{I} &= \lambda x.x \\ \mathbf{K} &= \lambda xy.x\end{aligned}$$

Najciekawsza własność $\lambda\mathbf{I}$ -termów jest taka:

Twierdzenie 6.1 *Jeśli $\lambda\mathbf{I}$ -term ma postać normalną, to ma własność silnej normalizacji.*

Dowód: Ćwiczenie 1. ■

Ćwiczenia

1. Niech $\Delta = (\lambda x.P)Q$ będzie redeksem lewostronnym w termie M i niech $Q \in \text{SN}$. Pokazać, że jeśli $N \in \text{SN}$ powstaje z M przez redukcję redeksu Δ to także $M \in \text{SN}$. Uwaga: założenie, że Δ jest redeksem lewostronnym jest istotne. Przykład: $(\lambda y.(\lambda x.z)(y\omega))\omega$.
2. Udowodnić twierdzenie 6.1. *Wskazówka: Użyć zadania 1 i wniosku 5.3.*
3. W pewnej książce wniosek 5.3 wyprowadza się z twierdzenia 6.1 w taki sposób:
Załóżmy, że term M ma postać normalną ale ma też nieskończoną redukcję. Z Twierdzenia 6.1 wynika, że przyczyną tego musi być jakiś podterm właściwy B termu M , który wcale nie ma

postaci normalnej i który jest „wycinany” przez pewien ciąg redukcji. Wybierzmy takie B możliwie najbardziej na lewo. Skoro ten term jest „wycinany” przez jakiś ciąg redukcji, to tym bardziej przez redukcję lewostronną.

Gdzie jest istotny błąd w tym rozumowaniu?

7 Siła wyrazu

Rachunek lambda jest systemem pozornie bardzo prostym. Abstrakcja i aplikacja wydają się trywialnymi operacjami, i może się zdawać, że niczego ciekawego nie da się z nich uzyskać. Okazuje się jednak, że siła wyrazu tego rachunku jest tak duża jak moc obliczeniowa maszyn Turinga.

Kombinator(y) punktu stałego

Zacniemy od dość paradoksalnej własności rachunku lambda. Każde równanie stałopunktowe ma w nim rozwiązanie. Odpowiedzialny za to jest na przykład taki term:

$$\mathbf{Y} := \lambda f.(\lambda x.f(xx))(\lambda x.f(xx)).$$

Główna własność termu $\mathbf{Y}$ jest taka: Dla dowolnego termu F zachodzi beta-równość:

$$(*) \quad F(\mathbf{Y}(F)) =_{\beta} \mathbf{Y}(F).$$

Rzeczywiście, $\mathbf{Y}(F) =_{\beta} (\lambda x.F(xx))(\lambda x.F(xx)) =_{\beta} F((\lambda x.F(xx))(\lambda x.F(xx))) =_{\beta} F(\mathbf{Y}(F))$. Termy o własności $(*)$ nazywamy *kombinatorami punktu stałego*.

Fakt 7.1 Dla dowolnego termu M istnieje taki term N , że $N =_{\beta} M[x := N]$.

Dowód: Wystarczy przyjąć $N = \mathbf{Y}(\lambda x.M)$. ■

Przykład 7.2

- Istnieje taki term M , że $Mxy =_{\beta} MyxM$. Można przyjąć $M = \mathbf{Y}(\lambda mxy.myxm)$. Wtedy $M =_{\beta} \lambda xy.MyxM$, skąd także $Mxy =_{\beta} MyxM$.
- Istnieje taki term M , że dla dowolnego N zachodzi $MN =_{\beta} MNN$, na przykład $M = \mathbf{Y}(\lambda mx.mxx)$.

Proste rzeczy

Teraz pokażemy jak w rachunku lambda można zinterpretować pewne proste konstrukcje. Zacniemy od wartości logicznych

$$\mathbf{true} = \lambda xy.x$$

$$\mathbf{false} = \lambda xy.y$$

i instrukcji warunkowej

$$\text{if } P \text{ then } Q \text{ else } R = PQR.$$

Łatwo sprawdzić, że $\text{if true then } Q \text{ else } R \rightarrow_\beta Q$ oraz $\text{if false then } Q \text{ else } R \rightarrow_\beta R$.

Następna konstrukcja to para uporządkowana. Przyjmiemy

$$\begin{aligned} \langle M, N \rangle &= \lambda x. xMN; \\ \pi_i &= \lambda x_1 x_2. x_i \text{ dla } i = 1, 2. \end{aligned}$$

Jak należy się spodziewać, mamy $\langle M_1, M_2 \rangle \pi_i \rightarrow_\beta M_i$. Zauważmy jednak, że nie zachodzi równość $\langle M\pi_1, M\pi_2 \rangle =_\beta M$, tj. nasza para uporządkowana nie jest *surjektywna*.

Uwaga: Dodanie do rachunku lambda surjektywnej pary uporządkowanej powoduje utratę własności Churcha-Rossera. Oznacza to w szczególności, że w (beztypowym) rachunku lambda nie można takiej operacji zdefiniować.

Liczby naturalne reprezentujemy w rachunku lambda jako tzw. *liczebniki Churcha*:

$$c_n = \lambda f x. f^n(x),$$

gdzie notacja $f^n(x)$ oznacza oczywiście term $f(f(\dots(x)\dots))$, w którym f występuje n razy. Zwykle zamiast c_n będziemy pisać $\mathbf{n}$, co jest mniej precyzyjne ale wygodne. Więc na przykład:

$$\begin{aligned} \mathbf{0} &= \lambda f x. x; \\ \mathbf{1} &= \lambda f x. f x; \\ \mathbf{2} &= \lambda f x. f(f x), \end{aligned}$$

i tak dalej. Zauważmy, że $\mathbf{1} =_{\beta\eta} \mathbf{I}$, ale $\mathbf{1} \neq_\beta \mathbf{I}$.

Powiemy, że funkcja częściowa $f : \mathbb{N}^k \rightarrow \mathbb{N}$ jest *definiowalna* w beztypowym rachunku lambda (λ -*definiowalna*) jeżeli istnieje term zamknięty F , spełniający dla dowolnych liczebników $n_1, \dots, n_k \in \mathbb{N}$ następujące warunki:

- Jeżeli $f(n_1, \dots, n_k) = m$, to $F\mathbf{n}_1 \dots \mathbf{n}_k =_\beta \mathbf{m}$;
- Jeżeli $f(n_1, \dots, n_k)$ jest nieokreślone, to $F\mathbf{n}_1 \dots \mathbf{n}_k$ nie ma postaci normalnej.

Mówimy oczywiście, że F *definiuje* lub *reprezentuje* funkcję f w rachunku lambda. W przypadku, gdy dodatkowo dla wszystkich $n_1, \dots, n_k \in \mathbb{N}$ spełniony jest warunek

- Jeżeli $f(n_1, \dots, n_k)$ jest określone, to $F\mathbf{n}_1 \dots \mathbf{n}_k$ ma własność silnej normalizacji,

to powiemy, że F *dobrze* definiuje funkcję f , którą nazwiemy *dobrze* definiowalną. Kilka przykładów termów definiujących pewne funkcje mamy poniżej.

Przykład 7.3

- Następnik: $\text{succ} \equiv \lambda n f x. f(n f x)$;
- Dodawanie: $\text{add} \equiv \lambda m n f x. m f(n f x)$;

- Mnożenie: **mult** $\equiv \lambda mnfx.m(nf)x$;
- Potegowanie: **exp** $\equiv \lambda mnfx.mnfx$;
- Test na zero: **zero** $\equiv \lambda m.m(\lambda y.\mathbf{false})\mathbf{true}$;
- Funkcja k -argumentowa stale równa zeru: **Z_k** $= \lambda m_1 \dots m_k.\mathbf{0}$;
- Rzut k -argumentowy na i -tą współrzędną: **Π_kⁱ** $= \lambda m_1 \dots m_k.m_i$.

Reprezentowanie funkcji częściowo rekurencyjnych

Przypomnijmy, że funkcje częściowo rekurencyjne tworzą najmniejszą klasę funkcji częściowych nad $\mathbb{N}$ zawierającą *funkcje bazowe*:

- następnik, $\text{succ}(n) = n + 1$;
- rzuty, $\Pi_k^i(n_1, \dots, n_k) = n_i$;
- funkcje stale równe zeru, $Z_k(n_1, \dots, n_k) = 0$,

i zamkniętą ze względu na składanie, rekursję prostą i operację minimum. Na wszelki wypadek, przypomnijmy też znaczenie tych terminów.

- Funkcja $f : \mathbb{N}^k \multimap \mathbb{N}$ powstaje przez *składanie* funkcji $h : \mathbb{N}^\ell \multimap \mathbb{N}$ z funkcjami $g_1, \dots, g_\ell : \mathbb{N}^k \multimap \mathbb{N}$, gdy
 - $\text{Dom}(f) = \{\vec{n} \mid \vec{n} \in \text{Dom}(g_i) \text{ dla wszystkich } i, \text{ oraz } (g_1(\vec{n}), \dots, g_\ell(\vec{n})) \in \text{Dom}(h)\}$;
 - $f(\vec{n}) = h(g_1(\vec{n}), \dots, g_\ell(\vec{n}))$, dla dowolnego $\vec{n} \in \text{Dom}(f)$.
- Funkcja $f : \mathbb{N}^{k+1} \multimap \mathbb{N}$ powstaje przez *rekursję prostą* z funkcji $h : \mathbb{N}^{k+2} \multimap \mathbb{N}$ i funkcji $g : \mathbb{N}^k \multimap \mathbb{N}$, jeżeli dla dowolnego $n \in \mathbb{N}$ i dowolnego $\vec{n} \in \mathbb{N}^k$ spełnione są równania
 - $f(0, \vec{n}) = g(\vec{n})$;
 - $f(n+1, \vec{n}) = h(f(n, \vec{n}), n, \vec{n})$,

przy czym równanie uważamy za spełnione, gdy obie strony są określone i równe, lub obie strony są nieokreślone. Wartość wyrażenia $h(f(n, \vec{n}), n, \vec{n})$ jest określona wtedy i tylko wtedy gdy $(n, \vec{n}) \in \text{Dom}(f)$ oraz $(f(n, \vec{n}), n, \vec{n}) \in \text{Dom}(h)$.

- Funkcja $f : \mathbb{N}^k \multimap \mathbb{N}$ powstaje z funkcji $h : \mathbb{N}^{k+1} \multimap \mathbb{N}$ przez zastosowanie *operacji minimum*, co zapisujemy $f(\vec{n}) = \mu y[h(\vec{n}, y) = 0]$, gdy dla dowolnego $\vec{n} \in \mathbb{N}^k$,
 - Jeśli istnieje takie m , że $h(\vec{n}, m) = 0$ oraz wszystkie wartości $h(\vec{n}, i)$ dla $i < m$ są określone i różne od zera, to $f(\vec{n}) = m$.
 - Jeśli takiego m nie ma, to $f(\vec{n})$ jest nieokreślone.

Klasa funkcji *pierwotnie rekurencyjnych* to najmniejsza klasa funkcji całkowitych nad $\mathbb{N}$ zawierająca funkcje bazowe i zamknięta ze względu na składanie i rekursję prostą. Szczególne funkcje pierwotnie rekurencyjne to *funkcja pary*

$$p(m, n) = \frac{(m+n)(m+n+1)}{2} + m,$$

i dwie funkcje *left* i *right*, takie że dla dowolnego n ,

$$n = p(\text{left}(n), \text{right}(n)),$$

czyli funkcje odwrotne do funkcji pary. Mamy następujące

Twierdzenie 7.4 (o postaci normalnej Kleene’go) *Każdą funkcję częściowo rekurencyjną można przedstawić w postaci*

$$f(\vec{n}) = \text{left}(\mu y [g(\vec{n}, y) = 0]),$$

gdzie g jest funkcją pierwotnie rekurencyjną.⁶

Udowodnimy teraz, że każda funkcja częściowo rekurencyjna jest definiowalna w rachunku lambda. Zaczynamy od najłatwiejszego.

Lemat 7.5 *Funkcje bazowe są lambda-definiowalne.*

Dowód: Przykład 7.3. ■

Lemat 7.6 *Jeśli funkcja całkowita f powstaje przez składanie lambda-definiowalnych funkcji całkowitych, to też jest lambda-definiowalna.*

Dowód: Oczywiście. Ale tylko dlatego, że mowa o funkcjach całkowitych. ■

Lemat 7.7 *Jeśli funkcja całkowita f powstaje przez rekursję prostą z lambda-definiowalnych funkcji całkowitych, to też jest lambda-definiowalna.*

Dowód: To już nie jest oczywiste. Załóżmy, że f jest zdefiniowana równaniami

$$\begin{aligned} f(0, \vec{n}) &= g(\vec{n}); \\ f(n+1, \vec{n}) &= h(f(n, \vec{n}), n, \vec{n}), \end{aligned}$$

i że funkcje g i h są definiowalne odpowiednio za pomocą termów G i H . Zdefiniujmy pomocnicze termy

$$\begin{aligned} \text{Step} &\equiv \lambda p. \langle \text{succ}(p\pi_1), H(p\pi_2)(p\pi_1)x_1 \dots x_m \rangle; \\ \text{Init} &\equiv \langle 0, Gx_1 \dots x_m \rangle. \end{aligned}$$

⁶Tak naprawdę to jest nawet funkcja elementarnie rekurencyjna.

Funkcja f jest wtedy definiowalna termem

$$F \equiv \lambda x x_1 \dots x_m. x \text{ Step Init } \pi_2.$$

Ta definicja wyraża następujący algorytm obliczania wartości funkcji f : generujemy ciąg par

$$(0, a_0), (1, a_1), \dots, (n, a_n),$$

gdzie $a_0 = g(n_1, \dots, n_m)$, każde a_{i+1} to $h(a_i, i, n_1, \dots, n_m)$ oraz $a_n = f(n, n_1, \dots, n_m)$. ■

Wniosek 7.8 *Funkcje pierwotnie rekurencyjne są lambda-definiowalne.*

Ponieważ mamy twierdzenie Kleene’go, więc pozostaje już (prawie) tylko pokazać lambda-definiowalność funkcji określonych przez minimum. Można do tego wykorzystać kombinatory punktu stałego **Y**, rozwiązując odpowiednie równanie stałopunktowe. My to zrobimy trochę delikatniej, „opóźniając” działanie kombinatora, dzięki czemu uzyskamy silniejszy wynik.

Twierdzenie 7.9 *Wszystkie funkcje częściowo rekurencyjne są lambda-definiowalne.*

Dowód: Niech $f(\vec{n}) = \text{left}(\mu y[g(\vec{n}, y) = 0])$, gdzie g jest funkcją pierwotnie rekurencyjną. Na mocy wniosku 7.8, zarówno g jak left są lambda-definiowalne pewnymi termami G i L . Określmy pomocniczy term

$$W \equiv \lambda y. \text{if zero}(G\vec{x}y) \text{ then } \lambda w. Ly \text{ else } \lambda w. w(\text{succ } y)w.$$

Funkcja f jest definiowana termem

$$F \equiv \lambda \vec{x}. W0W.$$

Rzeczywiście, przypuśćmy najpierw, że $\mu y[g(\vec{n}, y) = 0] = n$, oraz $\text{left}(n) = r$. Wtedy mamy taki ciąg redukcji:

$$F\vec{n} \rightarrow_\beta \overline{W0W} \rightarrow_\beta \overline{W1W} \rightarrow_\beta \dots \rightarrow_\beta \overline{WnW} \rightarrow_\beta L\mathbf{n} \rightarrow_\beta \mathbf{r}, \quad (1)$$

gdzie $\overline{W}$ oznacza wynik podstawienia $\vec{n}$ na $\vec{x}$ w termie W .

Jeśli minimum jest nieokreślone, to znaczy, że wszystkie wartości $g(\vec{n}, m)$ są określone i różne od zera, bo funkcja g jest całkowita. Wtedy mamy nieskończony ciąg redukcji

$$F\vec{n} \rightarrow_\beta \overline{W0W} \rightarrow_\beta \overline{W1W} \rightarrow_\beta \dots \rightarrow_\beta \overline{WnW} \rightarrow_\beta \dots$$

Ten ciąg jest quasi-lewostronny, a zatem na mocy Faktu 5.4 postaci normalnej nie ma. ■

W istocie prawdziwy jest silniejszy fakt:

Fakt 7.10 *Wszystkie funkcje częściowo rekurencyjne są dobrze lambda-definiowalne.*

Dowód Faktu 7.10 opiera się na tej samej konstrukcji co dowód twierdzenia 7.9, ale wymaga dodatkowo dowodu silnej normalizacji, który opuszczamy.

Wniosek 7.11 *Następujące problemy są nierozstrzygalne:*

- Dane termy M i N . Czy $M \rightarrow_{\beta} N$?
- Dane termy M i N . Czy $M =_{\beta} N$?
- Dany term M . Czy M ma postać normalną?
- Dany term M . Czy M ma własność silnej normalizacji?

Ćwiczenia

1. Pokazać, że istnieje nieskończenie wiele kombinatorów punktu stałego.
2. I że żaden z nich nie ma postaci normalnej.
3. Niech $G = \mathbf{SI} =_{\beta} \lambda y f.f(yf)$. Term jest kombinatorem punktu stałego wtedy i tylko wtedy, gdy jest punktem stałym G .
4. Zdefiniować w rachunku lambda poprzednik dla liczb naturalnych.
5. Pokazać, że w definicji funkcji definiowalnej można żądać $F\mathbf{n}_1 \dots \mathbf{n}_k \rightarrow_{\beta} \mathbf{f}(\mathbf{n}_1, \dots, \mathbf{n}_k)$, i że nic się nie zmieni, jeśli dopuścić $FV(F) \neq \emptyset$.
6. Nic się też nie zmieni, jeśli w definicji funkcji λ -definiowalnej zamiast $=_{\beta}$ użyjemy $=_{\beta\eta}$.

8 Lambda-teorie i drzewa Böhma

Jak już mówiliśmy, zwykły rachunek lambda można uważać (w uproszczeniu) za pewną teorię równościową. Inną teorię otrzymamy, jeśli do aksjomatów beta-konwersji dołożymy ekstensjonalność w postaci reguły (ext) lub aksjomatu (η). A co jeśli dodać jako nowy aksjomat na przykład równość $\mathbf{K} = \mathbf{S}$? Wtedy dla dowolnego termu M wywnioskujemy

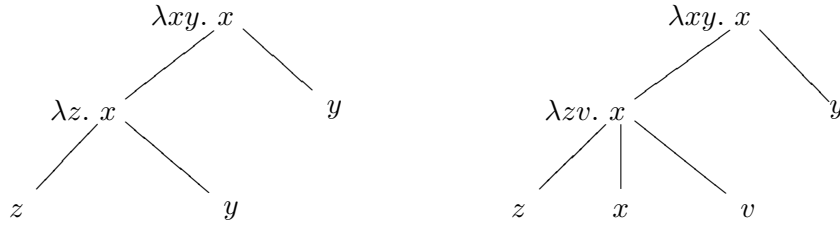
$$M = \mathbf{SI}(\mathbf{KM})\mathbf{I} = \mathbf{KI}(\mathbf{KM})\mathbf{I} = \mathbf{I},$$

a więc takie rozszerzenie rachunku lambda byłoby sprzeczne. Termy $\mathbf{K}$ i $\mathbf{S}$ to tylko przykład. Okazuje się, że sprzeczna musi być każda teoria, która skleja ze sobą dwie różne postaci $\beta\eta$ -normalne. Wynika to z następującego twierdzenia Böhma:

Twierdzenie 8.1 (C. Böhm, 1968) *Niech M i N będą kombinatorami w postaci β -normalnej i niech $M \neq_{\beta\eta} N$. Wtedy istnieje taki ciąg zamkniętych termów $\vec{P}$, że $M\vec{P} =_{\beta} \mathbf{true}$ oraz $N\vec{P} =_{\beta} \mathbf{false}$.*

Pełny dowód twierdzenia Böhma jest za długi jak na nasze możliwości. Ale warto, choćby na przykładzie, zobaczyć jak można odróżnić jeden term od drugiego. Weźmy na przykład $M = \lambda xy.x(\lambda z.xzy)y$ i $N = \lambda xy.x(\lambda zv.xzxv)y$. Wygodnie jest przedstawić oba termy za pomocą drzew (Obrazek 11).

Istotna różnica pomiędzy tymi drzewami to liście odpowiadające zmiennej y w termie M i zmiennej x w termie N . (Wystąpienie v w termie N jest nieistotne, bo $\lambda zv.xzxv =_{\eta} \lambda z.xzx$.) Aby wykorzystać tę różnicę do „odseparowania” termów M i N należy zaaplikować je do takich



Obrazek 11: Drzewa dla $M = \lambda xy.x(\lambda z.xzy)y$ i $N = \lambda xy.x(\lambda zv.xzxv)y$

argumentów, które wyciągną z nich „na wierzch” właśnie te liście. Posłużymy się trickiem, zwanym „*Böhm-out technique*”. Niech $P = \lambda uw.\langle u, w \rangle$. Wtedy dla dowolnego Q zachodzi $MPQ =_\beta \langle \lambda z.\langle z, Q \rangle, Q \rangle$. Dalej $MPQ \mathbf{true} R \mathbf{false} =_\beta Q$ dla dowolnego R . Tymczasem $NPQ \mathbf{true} R \mathbf{false} =_\beta P$. Wybierzmy więc jakiegokolwiek R oraz $Q = \lambda uvv.\mathbf{true}$ i niech $\vec{P}$ oznacza ciąg $(P, Q, \mathbf{true}, R, \mathbf{false}, \mathbf{false}, \mathbf{I}, \mathbf{true})$. Wtedy $M\vec{P} =_\beta \mathbf{true}$ oraz $N\vec{P} =_\beta \mathbf{false}$.

Rozwiązalność

Pierwsza intuicja związana z rachunkiem lambda (rozumianym jako abstrakcyjny język programowania) jest taka: „wartością” termu jest jego postać normalna. A zatem term, który nie ma postaci normalnej jest pozbawiony wartości. Każdy taki term jest równie dobry, albo raczej równie zły, i nie powinno być przeszkód w utożsamieniu ich wszystkich ze sobą. Niestety, utożsamienie na przykład termów $\lambda x.x\mathbf{K}\Omega$ i $\lambda x.x\mathbf{S}\Omega$ daje w wyniku sprzeczną teorię — wystarczy oba zaaplikować do $\mathbf{K}$.

Term bez postaci normalnej może więc czasem objawiać dobrze określone działanie, czyli w pewnym sensie mieć „wartość”. Można uważać, że tą wartością jest czołowa postać normalna, o ile istnieje. Przypomnijmy, że czołowa postać normalna to każdy term postaci $\lambda \vec{x}.z\vec{R}$. Term M ma czołową postać normalną N gdy $M =_\beta N$ i N jest w czołowej postaci normalnej.

Powiemy, że term zamknięty M jest *rozwiązalny* jeżeli $M\vec{P} =_\beta \mathbf{I}$ dla pewnych kombinatorów $\vec{P}$. Term ze zmiennymi wolnymi $\vec{x}$ jest *rozwiązalny* jeżeli rozwiązalne jest jego *domknięcie*, czyli kombinator $\lambda \vec{x}.M$.

Lemat 8.2 *Term jest rozwiązalny wtedy i tylko wtedy, gdy ma czołową postać normalną.*

Dowód: Bez straty ogólności można założyć, że mowa o termie zamkniętym.

Jeśli $M =_\beta \lambda x_1 x_2 \dots x_n. x_i R_1 \dots R_m$, to oczywiście $MP_1 P_2 \dots P_n =_\beta \mathbf{I}$, gdzie $P_i = \lambda y_1 \dots y_m. \mathbf{I}$.

Na odwrót, jeżeli $M\vec{P} =_\beta \mathbf{I}$, to w istocie $M\vec{P} \rightarrow_\beta \mathbf{I}$. Skoro więc $M\vec{P}$ ma czołową postać normalną $\mathbf{I}$, to i M musi mieć czołową postać normalną (ćwiczenie 3). ■

Jeśli termy nierozwiązalne uznamy za pozbawione wartości, to znaczy, że postulujemy teorię równościową, w której te wszystkie termy są równe. Jakie rozwiązalne termy powinny być równe w tej teorii? Nie możemy identyfikować termów za pomocą ich czołowych postaci

normalnych, bo każdy term rozwiązalny ma ich nieskończenie wiele. Nie możemy się też posłużyć beta-równością, bo chcemy aby $x\Omega$ było w naszej teorii równe $x(\Omega y)$. Pozostaje obserwować zachowanie termów w różnych sytuacjach. Jeśli jest ono zawsze takie samo, to powinniśmy dane termy utożsamić.

Niech M i N będą dowolnymi termami i niech $FV(M) \cup FV(N) = \vec{x}$. Powiemy, że M i N są *obserwacyjnie równoważne*, i napiszemy $M \equiv N$, gdy dla dowolnego kombinatora P ,

$$P(\lambda\vec{x}.M) \text{ jest rozwiązalny} \iff P(\lambda\vec{x}.N) \text{ jest rozwiązalny}.$$

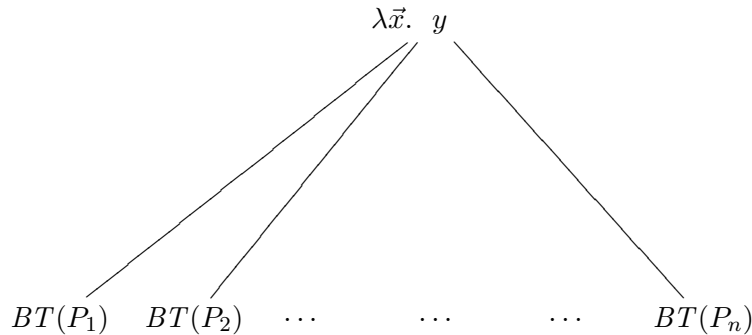
Nietrudno zauważyć, że $\equiv$ jest rozszerzeniem ekstensjonalnej teorii równościowej $\lambda\beta\eta$ i że $M \equiv N$ dla dowolnych nierozwiązalnych M i N (ćwiczenie 4). Ponadto, z twierdzenia Böhma wynika, że dla termów w postaci normalnej zachodzi równoważność

$$M \equiv N \iff M =_{\beta\eta} N.$$

Drzewa Böhma

Drzewo Böhma termu M , oznaczane przez $BT(M)$, definiujemy przez indukcję.

- Jeśli term M jest nierozwiązalny, to $BT(M)$ składa się tylko z jednego wierzchołka, oznaczonego przez Ω .
- Jeśli term M ma czołową postać normalną $\lambda\vec{x}.yP_1 \dots P_n$, to drzewo $BT(M)$ jest zbudowane jak na Obrazku 12.



Obrazek 12: Drzewo $BT(\lambda\vec{x}.yP_1 \dots P_n)$

Drzewo Böhma to uogólnienie postaci normalnej. Zwykła postać normalna to skończone drzewo Böhma bez wystąpienia Ω .

Chcemy teraz uogólnić pojęcie eta-równości na drzewa Böhma. Jasne, co powinno znaczyć $B \rightarrow_\eta B'$, ale w przypadku nieskończonych drzew musimy pozwolić na nieskończenie wiele różnic.

Oczywiście drzewo Böhma można uważać za „granicę” ciągu termów skończonych, w których może występować stała Ω . Ścisłej, powiemy, że ciąg drzew T_n jest *zbieżny* do drzewa B

jeżeli dla dowolnego n istnieje takie m , że wszystkie drzewa $T_m, T_{m+1}, \dots$ są do głębokości n identyczne z drzewem B . Ciąg skończony jest zaś zbieżny do swojego ostatniego elementu.

Powiemy, że drzewa Böhma B i B' są η -równoważne ($B \approx_\eta B'$), gdy istnieją (skończone lub nieskończone) ciągi *eta-ekspansji*:

$$B = B_0 \xrightarrow{\eta} B_1 \xrightarrow{\eta} B_2 \xrightarrow{\eta} B_3 \xrightarrow{\eta} \dots \quad \text{oraz} \quad B' = B'_0 \xrightarrow{\eta} B'_1 \xrightarrow{\eta} B'_2 \xrightarrow{\eta} B'_3 \xrightarrow{\eta} \dots$$

zbieżne do tego samego drzewa.

Przebrnąwszy przez tę ostatnią definicję, możemy sformułować twierdzenie Wadswortha:

Twierdzenie 8.3 (Wadsworth, 1975) *Termy M i N są obserwacyjnie równoważne wtedy i tylko wtedy, gdy $BT(M) \approx_\eta BT(N)$.*

Dowód z lewej do prawej polega na zastosowaniu techniki „Böhm-out”. W przeciwną stronę stosuje się metody semantyczne.

Ćwiczenia

1. Skonstruować takie $\mathbf{X}$, że $\mathbf{X}(\lambda yz.y(yz(yy))z) \rightarrow_{\beta\eta} \mathbf{K}$ oraz $\mathbf{X}(\lambda yz.y(yz(yz))z) \rightarrow_{\beta\eta} \mathbf{S}$.
Rozwiązanie: $\mathbf{X} = \lambda x.x(\lambda uv.\langle u, v \rangle)(\lambda uvw.\mathbf{S})\pi_1\pi_2\mathbf{I}\pi_1\mathbf{K}\mathbf{I}\pi_1$.
2. Udowodnić, że jeśli M nie ma czołowej postaci normalnej (tj. ma nieskończoną redukcję czołową) to żaden term postaci $M[x := P]$ nie ma czołowej postaci normalnej. *Wskazówka:* Jeśli $M \xrightarrow{h} N$ to $M[x := P] \xrightarrow{h} N[x := P]$.
3. Udowodnić, że jeśli M nie ma czołowej postaci normalnej, to MN też nie ma czołowej postaci normalnej. *Wskazówka:* Użyć poprzedniego ćwiczenia.
4. Udowodnić, że jeśli M nie ma czołowej postaci normalnej, ale PM ma czołową postać normalną, to każdy term postaci PN ma czołową postać normalną. *Wskazówka:* Jak wygląda czołowa postać normalna termu P ?
5. Załóżmy, że kombinatory M i N są obserwacyjnie równoważne i że $PM =_{\beta\eta} \lambda x_1 \dots x_n. x_i A_1 \dots A_k$, dla pewnych $A_1, \dots, A_k$. Udowodnić, że $PN =_{\beta\eta} \lambda x_1 \dots x_n. x_i B_1 \dots B_k$ dla pewnych $B_1, \dots, B_k$.
6. Znaleźć term, którego drzewo Böhma to jedna nieskończona gałąź postaci $\lambda x_1.x_1(\lambda x_2.x_2(\lambda x_3 \dots$.
Wskazówka: użyć kombinatora $\mathbf{Y}$.
7. Skonstruować term M o takiej własności: drzewo Böhma $BT(M)$ jest nieskończone i każdy jego wierzchołek ma więcej synów niż miał jego ojciec.
Rozwiązanie: $M = \mathbf{Y}F\mathbf{1}$, gdzie $F = \lambda V\lambda n\lambda x.n(\lambda y.y(V(\mathbf{succ}\ n))x)$.
8. Czy zawsze istnieje taki term M , że drzewo $BT(M)$ ma żądany kształt?

9 Rachunek kombinatorów

Zdefiniujemy teraz system CL (beztypowego) rachunku kombinatorów, zwany także (z przyczyn historycznych) *logiką kombinatoryczną*, chociaż ta druga nazwa jest nieco myląca. Zbiór termów $\mathcal{C}$ definiujemy jako najmniejszy zbiór o własnościach:

- Zmienne przedmiotowe należą do $\mathcal{C}$;

- Stałe K i S należą do $\mathcal{C}$;
- Jeśli $F, G \in \mathcal{C}$, to $(FG) \in \mathcal{C}$.

Konwencje nawiasowe są takie same jak dla termów rachunku lambda. Podstawienie w CL definiujemy podobnie jak dla rachunku lambda, ale łatwiej, bo nie ma zmiennych związanych.

- $x[x := F] = F$ oraz $y[x := F] = y$ dla $y \neq x$;
- $K[x := F] = K$ oraz $S[x := F] = S$;
- $GH[x := F] = G[x := F]H[x := F]$.

Relację *słabej redukcji* w zbiorze $\mathcal{C}$ definiujemy jako najmniejszą relację $\rightarrow_w$ o własnościach:

- $KAB \rightarrow_w A$ dla dowolnych A, B ;
- $SABC \rightarrow_w AC(BC)$ dla dowolnych A, B, C ;
- Jeśli $A \rightarrow_w B$, to $AC \rightarrow_w BC$ oraz $CA \rightarrow_w CB$, dla dowolnych A, B, C .

Jak zwykle, symbol $\rightarrow_w$ oznacza domknięcie przechodnio-zwrotne relacji $\rightarrow_w$, a $=_w$ (słaba równość), to najmniejsza relacja równoważności zawierająca $\rightarrow_w$. Termy rachunku kombinatorów nazywamy oczywiście *kombinatorami*. Zamknięte termy rachunku lambda przejęły tę nazwę od swoich odpowiedników w $\mathcal{C}$. Oto kilka ważnych kombinatorów i związane z nimi redukcje:

- Niech $I = SKK$. Wtedy $IF \rightarrow_w KF(KF) \rightarrow_w F$ dla dowolnego F .
- Term $\Omega = SII(SII)$ redukuje się sam do siebie.
- Niech $W = SS(KI)$. Wtedy $WFG \rightarrow_w FGG$ dla dowolnych F, G .
- Niech $B = S(KS)K$. Wtedy $BFGH \rightarrow_w F(GH)$, dla dowolnych F, G, H .
- Niech $C = S(BBS)(KK)$. Wtedy $CFGH \rightarrow_w FHG$ dla dowolnych F, G, H .
- Niech $B' = CB$. Wtedy $B'FGH \rightarrow_w G(FH)$.
- Niech $2 = SB(SB(KI))$. Wtedy $2FG \rightarrow_w F(FG)$.

Uwaga: Termy $K, S, KS, SK, I, B, 2, W$ są w postaci w -normalnej, a termy B', C nie są. Czasem wszystkie te kombinatory są jednak traktowane jak stałe.

System CL, jako system redukcyjny, ma własności podobne do rachunku lambda. Na przykład prawdziwe jest twierdzenie Churcha-Rossera, a takie własności jak normalizacja, silna normalizacja, czy równość, są nierozstrzygalne. System CL ma bowiem podobną siłę wyrazu — można w nim na przykład zdefiniować liczebniki i reprezentować funkcje częściowo rekurencyjne. W istocie, logika kombinatoryczna została wynaleziona przez Curry'ego i Schönfinkela niezależnie od rachunku lambda i mniej więcej w tym samym czasie, ale w gruncie rzeczy w podobnym celu.

Podobieństwo pomiędzy CL i rachunkiem lambda można wyrazić definiując translacje

$$(\)_{\Lambda} : \mathcal{C} \rightarrow \Lambda \qquad (\)_{\mathcal{C}} : \Lambda \rightarrow \mathcal{C}$$

Pierwsza z nich jest łatwa:

- $(x)_{\Lambda} = x$ dla $x \in V$;
- $(K)_{\Lambda} = \lambda xy.x$;
- $(S)_{\Lambda} = \lambda xyz.xz(yz)$;
- $(FG)_{\Lambda} = (F)_{\Lambda}(G)_{\Lambda}$.

Następujący łatwy fakt wyraża poprawność tej translacji ze względu na redukcje.

Fakt 9.1

- Jeśli $F \rightarrow_w G$, to $(F)_{\Lambda} \rightarrow_{\beta} (G)_{\Lambda}$.
- Jeśli $F =_w G$, to $(F)_{\Lambda} =_{\beta} (G)_{\Lambda}$.

Fakt 9.1 można odczytać tak: translację $(\)_{\Lambda}$ można uważać za przekształcenie z $\mathcal{C}/=_w$ do $\Lambda/=_\beta$, czyli za morfizm z teorii równościowej CL do teorii równościowej Λ .

Aby określić translację $(\)_{\mathcal{C}} : \Lambda \rightarrow \mathcal{C}$ musimy przede wszystkim zdefiniować w CL konstrukcję (zwaną *kombinatoryczną abstrakcją*), która zachowuje się tak jak abstrakcja. Można to zrobić na kilka sposobów, na przykład tak:

- $\lambda^*x.F = KF$, gdy $x \notin \text{FV}(F)$;
- $\lambda^*x.x = I$;
- $\lambda^*x.FG = S(\lambda^*x.F)(\lambda^*x.G)$, w przeciwnym przypadku.

Dowód poniższego faktu pozostawiamy jako ćwiczenie.

Fakt 9.2 $(\lambda^*x.F)G \rightarrow_w F[x := G]$.

Wniosek 9.3 (kombinatoryczna zupełność) Dla dowolnych x i F istnieje takie H , że dla dowolnego G zachodzi $HG \rightarrow_w F[x := G]$.

Teraz możemy zdefiniować translację $(\)_{\mathcal{C}} : \Lambda \rightarrow \mathcal{C}$.

- $(x)_{\mathcal{C}} = x$ dla $x \in V$;
- $(MN)_{\mathcal{C}} = (M)_{\mathcal{C}}(N)_{\mathcal{C}}$;
- $(\lambda x.M)_{\mathcal{C}} = \lambda^*x.(M)_{\mathcal{C}}$.

Fakt 9.4 Dla dowolnego $M \in \Lambda$ zachodzi $((M)_{\mathcal{C}})_{\Lambda} =_{\beta} M$.

Dowód: Indukcja ze względu na M . ■

A zatem translacja $(\)_\Lambda$ wyznacza przekształcenie „na” z teorii $\mathcal{C}/=_w$ do $\Lambda/=_\beta$. Poniższy wniosek stwierdza, że termy $\mathbf{K}$ i $\mathbf{S}$ stanowią bazę dla rachunku lambda.

Wniosek 9.5 *Każdy zamknięty lambda-term można otrzymać (z dokładnością do $=_\beta$) z termów $\mathbf{K}$ i $\mathbf{S}$ przez stosowanie aplikacji.*

Dowód: Natychmiast z Faktu 9.4. ■

Niestety, operacja $(\)_\mathcal{C}$ nie ma tak dobrych własności jak translacja $(\)_\Lambda$:

Fakt 9.6

- Złożenie $((\)_\Lambda)_\mathcal{C}$ nie jest identycznością. Na przykład $((\mathbf{K}_\Lambda)_\mathcal{C})\mathbf{I} \neq_w \mathbf{K}$.
- Z równości $M =_\beta N$ nie wynika $(M)_\mathcal{C} =_w (N)_\mathcal{C}$. Na przykład $\lambda x.\mathbf{K}\mathbf{I}x \rightarrow_\beta \lambda x.\mathbf{I}$, ale $(\lambda x.\mathbf{K}\mathbf{I}x)_\mathcal{C} = \mathbf{S}(\mathbf{K}(\mathbf{S}(\mathbf{K}\mathbf{K})\mathbf{I}))\mathbf{I} =_w \mathbf{S}(\mathbf{K}(\mathbf{K}\mathbf{I}))\mathbf{I} \neq_w \mathbf{K}\mathbf{I} = (\lambda x.\mathbf{I})_\mathcal{C}$.

A zatem „słaba” równość jest w istocie *silniejsza* niż $=_\beta$. Przyczyną tego jest to, że kombinatoryczna abstrakcja λ^* nie spełnia warunku słabej ekstensjonalności ze względu na $=_w$:

$$\frac{M = N}{\lambda x.M = \lambda x.N}, \quad (\xi)$$

W istocie

$$\lambda^*x.\mathbf{K}\mathbf{I}x = \mathbf{S}(\mathbf{K}(\mathbf{K}\mathbf{I}))\mathbf{I} \neq_w \mathbf{K}\mathbf{I} = \lambda^*x.\mathbf{I},$$

choć $\mathbf{K}\mathbf{I}x \rightarrow_w \mathbf{I}$. Można to poprawić kosztem wprowadzenia ekstensjonalności. Niech $=_{ext}$ będzie najmniejszą relacją równoważności w $\mathcal{C}$, taką że:

- Jeśli $G =_w H$, to $G =_{ext} H$;
- Jeśli $Gx =_{ext} Hx$ i $x \notin \text{FV}(G) \cup \text{FV}(H)$, to $G =_{ext} H$;
- Jeśli $G =_{ext} G'$, to $GH =_{ext} G'H$ i $HG =_{ext} HG'$.

Fakt 9.7

- Dla dowolnych $G, H \in \mathcal{C}$, warunki $G =_{ext} H$ i $(G)_\Lambda =_{\beta\eta} (H)_\Lambda$ są równoważne.
- Dla dowolnych $M, N \in \Lambda$, warunki $M =_{\beta\eta} N$ i $(M)_\mathcal{C} =_{ext} (N)_\mathcal{C}$ są równoważne.
- Równanie $((G)_\Lambda)_\mathcal{C} =_{ext} G$ zachodzi dla dowolnego $G \in \mathcal{C}$.

Ćwiczenia

1. Udowodnić Fakt 9.2.
2. Niech $\mathbf{B} = \lambda xyz.x(yz)$ i niech $\mathbf{C} = \lambda xyz.xzy$. Pokazać, że każdy $\lambda\mathbf{I}$ -term można otrzymać (z dokładnością do $=_\beta$) z termów $\mathbf{S}$, $\mathbf{B}$, $\mathbf{C}$, $\mathbf{I}$, poprzez stosowanie aplikacji. *Wskazówka: zdefiniować konstrukcję $\lambda^\circ x.F$ w ten sposób, że:*
 - $\lambda^\circ x.FG = \mathbf{C}(\lambda^\circ x.F)G$, gdy $x \in \text{FV}(F)$ oraz $x \notin \text{FV}(G)$;
 - $\lambda^\circ x.FG = \mathbf{B}F(\lambda^\circ x.G)$, gdy $x \notin \text{FV}(F)$ oraz $x \in \text{FV}(G)$.
3. Term nazwiemy *afinicznym*, gdy każdy jego podterm postaci $\lambda x.M$ zawiera co najwyżej jedno wolne wystąpienie zmiennej x . Pokazać, że każdy afiniczny lambda-term można otrzymać (z dokładnością do $=_\beta$) z termów $\mathbf{B}$, $\mathbf{C}$, $\mathbf{K}$ poprzez stosowanie aplikacji.
4. Term nazwiemy *liniowym*, gdy każdy jego podterm postaci $\lambda x.M$ zawiera dokładnie jedno wolne wystąpienie zmiennej x . Pokazać, że każdy liniowy lambda-term można otrzymać (z dokładnością do $=_\beta$) z termów $\mathbf{B}$, $\mathbf{C}$, $\mathbf{I}$ poprzez stosowanie aplikacji.

Problem otwarty:

Zdefiniować taką (obliczalną) translację $T : \Lambda \rightarrow \mathcal{C}$, że

$$M =_\beta N \iff T(M) =_w T(N).$$

10 Modele

Zacniemy od pewnej ogólnej teorii, potrzebnej dla uściślenia pojęcia modelu dla rachunku lambda. Naszym zamiarem jest interpretacja każdego lambda-termu M w pewnej konkretnej dziedzinie A , jako pewnego elementu $\llbracket M \rrbracket \in A$. Ponieważ mamy w termach zmienne wolne, więc nasza interpretacja może zależeć od *wartościowania* $v : V \rightarrow A$ (gdzie V oznacza zbiór wszystkich zmiennych), co zaznaczymy pisząc $\llbracket M \rrbracket_v$. Indeks v pomijamy, gdy v jest znane, lub M nie ma zmiennych wolnych. Poniższa definicja formułuje nasze podstawowe oczekiwania w stosunku do takiej interpretacji.

Niech $\cdot$ będzie dwuargumentową operacją w zbiorze A i niech $\llbracket \cdot \rrbracket : \Lambda \times A^V \rightarrow A$. Strukturę $\mathcal{A} = \langle A, \cdot, \llbracket \cdot \rrbracket \rangle$ nazywamy *lambda-interpretacją*, jeżeli spełnione są warunki (a) – (d) poniżej. Stosujemy notację $\llbracket M \rrbracket_v$ na oznaczenie wartości $\llbracket \cdot \rrbracket(M, v)$. Symbol $v[x \mapsto a]$ oznacza wartościowanie różniące się od v tylko tym, że $v[x \mapsto a](x) = a$.

- (a) Jeśli x jest zmienną, to $\llbracket x \rrbracket_v = v(x)$;
- (b) $\llbracket PQ \rrbracket_v = \llbracket P \rrbracket_v \cdot \llbracket Q \rrbracket_v$;
- (c) $\llbracket \lambda x.P \rrbracket_v \cdot a = \llbracket P \rrbracket_{v[x \mapsto a]}$, dla dowolnego $a \in A$;
- (d) Jeśli $v|_{\text{FV}(P)} = u|_{\text{FV}(P)}$, to $\llbracket P \rrbracket_v = \llbracket P \rrbracket_u$.

Na pierwszy rzut oka może się wydawać, że warunki (a) – (d) stanowią indukcyjną *definicję* funkcji $\llbracket \cdot \rrbracket$. Tak nie jest, bo warunek (c) nie określa *którym elementem* modelu jest $\llbracket \lambda x.P \rrbracket_v$

a postuluje tylko jego zachowanie przy aplikacji. (Istnienie elementu, który tak się zachowuje, nie jest zresztą oczywiste.)

Jeśli $a, b \in A$, to napiszemy $a \approx b$, gdy dla dowolnego $c \in A$ zachodzi równość $a \cdot c = b \cdot c$. Oznacza to, że zachowanie a i b jako funkcji jest takie samo. Interpretacja jest *ekstensjonalna* wtedy i tylko wtedy, gdy z warunku $a \approx b$ wynika $a = b$.

Lambda-interpretację $\mathcal{A} = \langle A, \cdot, \llbracket \cdot \rrbracket \rangle$ nazywamy

- *lambda-algebrą*, gdy warunek $M =_{\beta} N$ implikuje $\llbracket M \rrbracket_v = \llbracket N \rrbracket_v$ dla dowolnego v (co zapiszemy $\mathcal{A} \models M = N$);
- *lambda-modelem*, gdy warunek $\llbracket \lambda x.M \rrbracket_v \approx \llbracket \lambda x.N \rrbracket_v$ implikuje $\llbracket \lambda x.M \rrbracket_v = \llbracket \lambda x.N \rrbracket_v$.

A więc lambda-algebra to poprawna interpretacja beta-równości. Natomiast lambda-model to interpretacja słabo ekstensjonalna, tj. taka, w której ekstensjonalność zachodzi dla abstrakcji. Aby zrozumieć o co chodzi, zauważmy tu, że $\llbracket \lambda x.M \rrbracket_v \approx \llbracket \lambda x.N \rrbracket_v$ oznacza tyle samo co $\llbracket M \rrbracket_{v[x \mapsto a]} = \llbracket N \rrbracket_{v[x \mapsto a]}$ dla wszystkich a . A zatem definicja lambda-modelu mówi, że znaczenie abstrakcji $\llbracket \lambda x.M \rrbracket$ jest zdeterminowane przez funkcję, która każdemu a przypisuje element $\llbracket M \rrbracket_{v[x \mapsto a]}$. Oznacza to tyle, że abstrakcja jest operacją semantyczną (ma sens w modelu, niezależny od składni) tak samo jak aplikacja.

Przykład 10.1 Szczególnym przypadkiem interpretacji jest interpretacja zbudowana z samych termów. Niech $\mathfrak{M}(\lambda) = \langle \Lambda / \equiv_{\beta}, \cdot, \llbracket \cdot \rrbracket \rangle$, gdzie $[M]_{\beta} \cdot [N]_{\beta} = [MN]_{\beta}$ oraz $\llbracket M \rrbracket_v$ jest wynikiem *jednoczesnego* podstawienia termu $v(x)$ na każdą zmienną wolną x termu M . Przez $\mathfrak{M}^0(\lambda)$ oznaczmy analogiczną interpretację, której dziedzinę tworzą klasy termów zamkniętych. Podobnie $\mathfrak{M}(\lambda\eta)$ i $\mathfrak{M}^0(\lambda\eta)$ oznaczają struktury w których zbiór termów (odp. zbiór kombinatorów) podzielono przez $\equiv_{\beta\eta}$. Wszystkie te struktury są lambda-algebrami, ale $\mathfrak{M}^0(\lambda)$ i $\mathfrak{M}^0(\lambda\eta)$ nie są lambda-modelami.

Następujący lemat mówi o własności, której oczekujemy od każdej sensownej semantyki. Znaczenie podstawienia $M[x := N]$ powinno być takie samo jak znaczenie termu M przy odpowiednio zmienionym wartościowaniu.

Lemat 10.2 W dowolnym lambda-modelu zachodzi tożsamość $\llbracket M[x := N] \rrbracket_v = \llbracket M \rrbracket_{v[x \mapsto \llbracket N \rrbracket_v]}$.

Dowód: Udowodnimy najpierw, że teza zachodzi przy dodatkowym założeniu $x \notin \text{FV}(N)$. Postępujemy przez indukcję ze względu na M . Przypadki zmiennej i aplikacji są łatwe, niech więc $M = \lambda y.P$. Przyjmijmy, że $y \notin \text{FV}(N)$. Dla dowolnego a , korzystając z założenia indukcyjnego o P , dostaniemy

$$\begin{aligned} \llbracket (\lambda y.P)[x := N] \rrbracket_v \cdot a &= \llbracket \lambda y.P[x := N] \rrbracket_v \cdot a = \llbracket P[x := N] \rrbracket_{v[y \mapsto a]} = \\ &= \llbracket P \rrbracket_{v[y \mapsto a][x \mapsto \llbracket N \rrbracket_{\bar{v}}]} = \llbracket P \rrbracket_{v[y \mapsto a][x \mapsto \llbracket N \rrbracket_v]} = \llbracket (\lambda y.P) \rrbracket_{v[x \mapsto \llbracket N \rrbracket_v]} \cdot a, \end{aligned}$$

gdzie $\bar{v} = v[y \mapsto a]$. Z tego wynika

$$\llbracket (\lambda y.P)[x := N] \rrbracket_v = \llbracket (\lambda y.P)[x := N] \rrbracket_{v[x \mapsto \llbracket N \rrbracket_v]} = \llbracket (\lambda y.P) \rrbracket_{v[x \mapsto \llbracket N \rrbracket_v]},$$

bo mamy do czynienia z lambda-modelem.

Niech teraz $x \in \text{FV}(N)$. Ustalmy nową zmienną z i niech w oznacza $v[z \mapsto \llbracket x \rrbracket_v]$. Zauważmy, że $w[x \mapsto \llbracket z \rrbracket_w] = w$. A zatem

$$\llbracket M[x := N][x := z] \rrbracket_w = \llbracket M[x := N] \rrbracket_w = \llbracket M[x := N] \rrbracket_v,$$

bo $x \notin \text{FV}(z)$. Niech $N' = N[x := z]$. Wtedy $M[x := N][x := z] = M[x := N']$, więc

$$\llbracket M[x := N][x := z] \rrbracket_w = \llbracket M[x := N'] \rrbracket_w = \llbracket M \rrbracket_{w[x \mapsto \llbracket N' \rrbracket_w]} = \llbracket M \rrbracket_{v[x \mapsto \llbracket N' \rrbracket_v]},$$

bo $x \notin \text{FV}(N')$. Ale $\llbracket N' \rrbracket_w = \llbracket N[x := z] \rrbracket_w = \llbracket N \rrbracket_{w[x \mapsto \llbracket z \rrbracket_w]} = \llbracket N \rrbracket_w = \llbracket N \rrbracket_v$, ponieważ $x \notin \text{FV}(z)$. Stąd ostatecznie $\llbracket M[x := N] \rrbracket_v = \llbracket M[x := N][x := z] \rrbracket_w = \llbracket M \rrbracket_{v[x \mapsto \llbracket N \rrbracket_v]}$. ■

Fakt 10.3 *Każdy lambda-model jest lambda-algebrą.*

Dowód: Przez indukcję ze względu na definicję beta-równości pokazujemy że jeśli $M =_\beta N$, to $\llbracket M \rrbracket_v = \llbracket N \rrbracket_v$.

Najpierw zauważmy, że $\llbracket (\lambda x.P)Q \rrbracket_v = \llbracket \lambda x.P \rrbracket_v \cdot \llbracket Q \rrbracket_v = \llbracket P \rrbracket_{v[x \mapsto \llbracket Q \rrbracket_v]} = \llbracket P[x := Q] \rrbracket_v$ na mocy lematu 10.2.

Niech teraz $M = \lambda x.P$ i $N = \lambda x.Q$, gdzie $P =_\beta Q$. Z założenia indukcyjnego wynika, że $\llbracket P \rrbracket_u = \llbracket Q \rrbracket_u$, przy każdym u , w szczególności gdy $u = v[x \mapsto a]$. A zatem $\llbracket M \rrbracket_v \cdot a = \llbracket P \rrbracket_{v[x \mapsto a]} = \llbracket Q \rrbracket_{v[x \mapsto a]} = \llbracket N \rrbracket_v \cdot a$ dla dowolnego a . Stąd $\llbracket M \rrbracket_v = \llbracket N \rrbracket_v$.

Pozostałe przypadki są rutynowe. ■

Twierdzenie 10.4 (o pełności) *Następujące warunki są równoważne:*

- 1) $M =_\beta N$;
- 2) $\mathcal{A} \models M = N$ dla dowolnej lambda-algebry $\mathcal{A}$;
- 3) $\mathcal{A} \models M = N$ dla dowolnego lambda-modelu $\mathcal{A}$.

Dowód: Implikacja (1) $\Rightarrow$ (2) wynika wprost z definicji, a implikacja (2) $\Rightarrow$ (3) z Lematu 10.3. Natomiast implikacja (3) $\Rightarrow$ (1) wynika stąd, że $\mathfrak{M}(\lambda)$ jest lambda-modelem, a termy równe w $\mathfrak{M}(\lambda)$ muszą być β -równe (rozpatrzmy trywialne wartościowanie $v(x) = x$). ■

Modele częściowo uporządkowane

Na początek przypomnijmy kilka definicji. Niech $\langle A, \leq \rangle$ będzie porządkiem częściowym.

- Podzbiór B zbioru A jest *skierowany* wtedy i tylko wtedy, gdy dla dowolnych $a, b \in B$ istnieje takie $c \in B$, że $a, b \leq c$.
- Zbiór A jest *zupełnym* porządkiem częściowym (cpo) wtedy i tylko wtedy, gdy każdy jego skierowany podzbiór ma kres górny.

Oczywiście każdy łańcuch jest zbiorem skierowanym. W szczególności elementy dowolnego ciągu wstępującego $a_0 \leq a_1 \leq a_2 \leq \dots$ tworzą zbiór skierowany. Także zbiór pusty jest zbiorem skierowanym. Z definicji porządku zupełnego wynika więc istnienie elementu najmniejszego $\sup \emptyset$, tradycyjnie oznaczanego przez $\perp$.

Niech teraz $\langle A, \leq \rangle$ i $\langle B, \leq \rangle$ będą porządkami częściowymi.

- Funkcja $f : A \rightarrow B$ jest *monotoniczna* wtedy i tylko wtedy, gdy dla dowolnych $x, y \in A$ nierówność $x \leq y$ implikuje $f(x) \leq f(y)$.
- Jeśli $\langle A, \leq \rangle$ i $\langle B, \leq \rangle$ są zupełnymi porządkami częściowymi to funkcja $f : A \rightarrow B$ jest *ciągła* wtedy i tylko wtedy, gdy f zachowuje kresy górne niepustych zbiorów skierowanych, tj. dla dowolnego skierowanego i niepustego podzbioru $X \subseteq A$ istnieje $\sup f(X)$ i zachodzi równość $f(\sup X) = \sup f(X)$.

Fakt 10.5 Każda funkcja ciągła jest monotoniczna. ■

Zbiór wszystkich funkcji ciągłych z A do B oznaczamy przez $[A \rightarrow B]$. Ten zbiór, uporządkowany warunkiem:

$$f \leq g \quad \text{wtedy i tylko wtedy, gdy} \quad \forall a \in A (f(a) \leq g(a)),$$

jest zupełnym porządkiem częściowym. Podobnie można nadać strukturę cpo zbiorowi $A \times B$ za pomocą zwykłego uporządkowania „po współrzędnych”:

$$\langle a, b \rangle \leq \langle a', b' \rangle \quad \text{wtedy i tylko wtedy, gdy} \quad a \leq a' \text{ i } b \leq b'.$$

Lemat 10.6 Funkcja $f : D_1 \times D_2 \rightarrow D$ jest ciągła wtedy i tylko wtedy gdy jest ciągła ze względu na obie współrzędne, tj. dla dowolnych $a \in D_1$ i $b \in D_2$ oraz dowolnych $A \subseteq D_1$ i $B \subseteq D_2$ zachodzi $f(\sup A, b) = \sup f(A, b)$ oraz $f(a, \sup B) = \sup f(a, B)$.

Dowód: Implikacja z lewej do prawej jest oczywista. Dla dowodu implikacji odwrotnej rozpatrzmy skierowany zbiór $X \subseteq D_1 \times D_2$ i niech $A = \pi_1(X)$ oraz $B = \pi_2(X)$. Oczywiście $X \subseteq A \times B$, chociaż niekoniecznie na odwrót. Ale jeśli $\langle a, b \rangle \in A \times B$, to zawsze istnieje takie $\langle a', b' \rangle \in X$, że $\langle a, b \rangle \leq \langle a', b' \rangle$. Istotnie, mamy wtedy pary $\langle a, b'' \rangle, \langle a'', b \rangle \in X$. Skoro zbiór X jest skierowany, to $\langle a, b'' \rangle, \langle a'', b \rangle \leq \langle a', b' \rangle$ dla pewnego $\langle a', b' \rangle \in X$. Porządek jest po współrzędnych, więc $\langle a, b \rangle \leq \langle a', b' \rangle$.

Z powyższej obserwacji wynika, że $\sup X = \sup(A \times B) = \langle \sup A, \sup B \rangle$. Niech $\langle a_0, b_0 \rangle$ będzie tym kresem. Należy sprawdzić, że $f(a_0, b_0) = \sup f(X)$. Oczywiście $f(a_0, b_0) \geq \sup f(X)$, przypuśćmy więc, że $c \geq f(a, b)$ dla wszystkich $\langle a, b \rangle \in X$. Wtedy także $c \geq f(a, b)$ dla wszystkich $\langle a, b \rangle \in A \times B$. Jeśli ustalimy dowolne $b \in B$, to na mocy ciągłości ze względu na pierwszą współrzędną, otrzymamy $c \geq f(a_0, b)$. Dalej z ciągłości ze względu na drugą współrzędną wynika także $c \geq f(a_0, b_0)$. ■

Lemat 10.7 Operacja aplikacji $Ap : D_1 \times [D_1 \rightarrow D_2] \rightarrow D_2$, określona przez $Ap(d, f) = f(d)$, jest funkcją ciągłą.

Dowód: Na mocy lematu 10.6 wystarczy sprawdzić ciągłość ze względu na współrzędne. Dla skierowanego $A \subseteq D_1$ i ustalonego f mamy $\sup Ap(A, f) = \sup f(A) = f(\sup A) = Ap(\sup A, f)$, bo f jest ciągła. A jeśli $F \subseteq [D_1 \rightarrow D_2]$ jest skierowany oraz $a \in D_1$ to $\sup Ap(a, F) = \sup\{f(a) \mid f \in F\} = (\sup F)(a) = Ap(a, \sup F)$ z definicji $\sup F$. ■

Lemat 10.8 *Operacja abstrakcji $Abs : [(D_1 \times D_2) \rightarrow D] \rightarrow [D_1 \rightarrow [D_2 \rightarrow D]]$, dana przez $Abs(f)(d)(e) = f(d, e)$, jest dobrze określona i ciągła.*

Dowód: Dla dowolnego d , funkcja $Abs(f)(d)$ jest złożeniem funkcji ciągłej f z funkcją ciągłą $g(e) = \langle d, e \rangle$, zatem jest ciągła.

Dla $A \subseteq D_1$ mamy $Abs(f)(\sup A)(e) = f(\sup A, e) = \sup f(A, e) = \sup Abs(f)(A)(e)$ czyli $Abs(f)(\sup A) = \sup Abs(f)(A)$, co oznacza, że funkcja $Abs(f)$ jest ciągła.

Wreszcie gdy F jest skierowanym podzbiorem $[(D_1 \times D_2) \rightarrow D]$, to $Abs(\sup F)(d)(e) = (\sup F)(d, e) = \sup\{f(d, e) \mid f \in F\} = \sup\{Abs(f)(d)(e) \mid f \in F\}$, czyli sama funkcja Abs też jest ciągła. ■

Refleksywne cpo

Kluczowa definicja jest taka. Zupełny porządek częściowy D jest *refleksywny*, gdy przestrzeń funkcyjna $[D \rightarrow D]$ jest retraktem D , tj. gdy istnieją przekształcenia ciągłe $F : D \rightarrow [D \rightarrow D]$ i $G : [D \rightarrow D] \rightarrow D$, takie, że $F \circ G = \text{id}_{[D \rightarrow D]}$. Zauważmy, że wtedy $[D \rightarrow D]$ jest izomorficzne z pewnym podzbiorem D (obrazem funkcji G).

Na takim cpo możemy określić lambda-interpretację $\mathcal{D} = \langle D, \cdot, \llbracket \cdot \rrbracket \rangle$, gdzie dla $a, b \in D$

$$a \cdot b = F(a)(b),$$

a znaczenie termów jest określone tak:

- Jeśli x jest zmienną, to $\llbracket x \rrbracket_v = v(x)$;
- $\llbracket PQ \rrbracket_v = \llbracket P \rrbracket_v \cdot \llbracket Q \rrbracket_v$;
- $\llbracket \lambda x. P \rrbracket_v = G(\xi_P^{v, x})$.

Powyżej, symbol $\xi_P^{v, x}$ oznacza funkcję określoną równaniem $\xi_P^{v, x}(a) = \llbracket P \rrbracket_{v[x \mapsto a]}$, którą nieformalnie możemy zapisać tak $\Lambda a. \llbracket P \rrbracket_{v[x \mapsto a]}$.

Nietrudno zauważyć, że wtedy zachodzą następujące warunki:

- $\llbracket \lambda x. P \rrbracket_v \cdot a = \llbracket P \rrbracket_{v[x \mapsto a]}$, dla dowolnego $a \in D$;
- Jeśli $v|_{\text{FV}(P)} = u|_{\text{FV}(P)}$, to $\llbracket P \rrbracket_v = \llbracket P \rrbracket_u$,

a zatem faktycznie mamy lambda-interpretację. Nietrudno też przeoczyć pewien drobiazg: poprawność powyższej definicji nie jest wcale oczywista i wymaga dowodu. Operacja G jest bowiem określona tylko dla funkcji ciągłych.

Lemat 10.9 Dla dowolnych M i v , funkcja $\xi_M^{v,x}$ (czyli $\Lambda a. \llbracket M \rrbracket_{v[x \mapsto a]}$) jest ciągła.

Dowód: Indukcja ze względu na M . Dla $M = PQ$, funkcja $\xi_M^{v,x} = \Lambda a. \llbracket M \rrbracket_{v[x \mapsto a]}$ jest złożeniem postaci $Ap \circ h$, gdzie $h(a) = \langle F(\xi_P^{v,x}(a)), \xi_Q^{v,x}(a) \rangle$ jest ciągła. Jeśli $M = \lambda y N$ to z założenia indukcyjnego funkcja f , która parze $\langle a, b \rangle$ przypisuje $\llbracket N \rrbracket_{v[x \mapsto a][y \mapsto b]}$ jest ciągła ze względu na obie współrzędne, a zatem ciągła, na mocy lematu 10.6. Natomiast funkcja $\xi_M^{v,x} = \Lambda a. G(\Lambda b. \llbracket N \rrbracket_{v[x \mapsto a][y \mapsto b]})$ jest złożeniem $G \circ Abs(f)$. ■

Twierdzenie 10.10 Jeśli D jest refleksywnym cpo, to lambda-interpretacja $\mathcal{D} = \langle D, \cdot, \llbracket \cdot \rrbracket \rangle$ jest lambda-modelem.

Dowód: Pozostaje sprawdzić, że mamy słabą ekstensjonalność. Wystarczy w tym celu zauważyć, że warunek $\llbracket \lambda x. M \rrbracket_v \approx \llbracket \lambda x. N \rrbracket_v$ implikuje $\xi_M^{v,x} = \xi_N^{v,x}$. ■

Uwaga: Nasz model jest ekstensjonalny wtedy i tylko wtedy, gdy $G \circ F = \text{id}_D$.

Ćwiczenia

1. Pokazać, że $\mathfrak{M}^0(\lambda)$, $\mathfrak{M}(\lambda)$, $\mathfrak{M}(\lambda\eta)$ i $\mathfrak{M}^0(\lambda\eta)$ są lambda-algebrami.
2. Pokazać, że $\mathfrak{M}(\lambda)$ i $\mathfrak{M}(\lambda\eta)$ są lambda-modelami.
3. Znaleźć w książce Barendregta, że $\mathfrak{M}^0(\lambda)$ ani $\mathfrak{M}^0(\lambda\eta)$ nie są lambda-modelami.
4. (Meyer) Niech $\mathbf{1} = \llbracket \mathbf{1} \rrbracket$. Pokazać, że lambda-algebra jest lambda-modelem wtedy i tylko wtedy, gdy z warunku $a \approx b$ wynika $\mathbf{1} \cdot a = \mathbf{1} \cdot b$.
5. Czy z tego, że $\mathcal{A} \models \lambda x. Mx = M$ wynika ekstensjonalność interpretacji $\mathcal{A}$?
Wskazówka: Użyć zadania 3.
6. Pokazać, że jeśli D jest refleksywnym cpo i $f : D \rightarrow D$ jest ciągła, to istnieje takie $a \in D$, że $f(x) = a \cdot x$ dla dowolnego x .
7. Czy funkcja odwrotna do ciągłej bijekcji musi być ciągła?

11 Model $\mathcal{D}_\infty$

Skonstruujemy teraz konkretny przykład modelu — pewne refleksywne cpo, zwane modelem Scotta $\mathcal{D}_\infty$. Zaczniemy od prostej definicji.

Niech A i B będą cpo. *Projekcją* z B na A nazywamy parę funkcji ciągłych

$$\varphi : A \rightarrow B \quad \text{i} \quad \psi : B \rightarrow A,$$

spełniającą warunki

$$\psi \circ \varphi = \text{id}_A \quad \text{oraz} \quad \varphi \circ \psi \leq \text{id}_B.$$

Zauważmy, że wtedy $\varphi(\perp_A) = \perp_B$, bo $\varphi(\perp_A) \leq \varphi(\psi(\perp_B)) \leq \perp_B$. Gdy mamy taką projekcję, to w szczególności A jest retraktem B . Identyfikując element $d \in A$ z jego obrazem $\varphi(d)$, możemy wtedy uważać zbiór A za podzbiór B . Przykładem projekcji jest para funkcji

$$\varphi_0 : D \rightarrow [D \rightarrow D] \quad \text{ i } \quad \psi_0 : [D \rightarrow D] \rightarrow D,$$

określona tak:

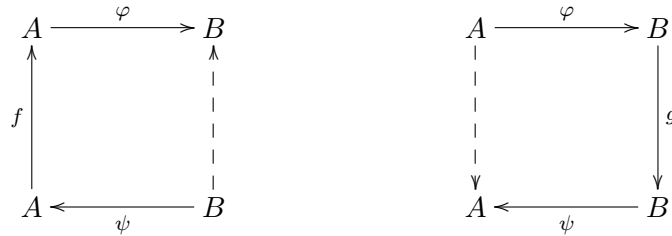
$$\varphi_0(d)(a) = d, \quad \text{ oraz } \quad \psi_0(f) = f(\perp)$$

dla dowolnych $a, d \in D$ i $f \in [D \rightarrow D]$. Mamy tu zanurzenie zbioru D w przestrzeń funkcyjną, przy którym każdy element D jest utożsamiany z funkcją stałą.

Jeśli mamy projekcję (φ, ψ) z B na A , to możemy „podnieść” tę projekcję na poziom funkcji, czyli określić nową projekcję (φ^*, ψ^*) z przestrzeni $[B \rightarrow B]$ na $[A \rightarrow A]$. Dla $f : A \rightarrow A$ i $g : B \rightarrow B$ przyjmujemy

$$\varphi^*(f) = \varphi \circ f \circ \psi \quad \text{ oraz } \quad \psi^*(g) = \psi \circ g \circ \varphi,$$

co ilustruje obrazek:



Lemat 11.1 *Jeśli (φ, ψ) jest projekcją z B na A , to para (φ^*, ψ^*) jest projekcją z $[B \rightarrow B]$ na $[A \rightarrow A]$.*

Dowód: Ćwiczenie. ■

Przypuśćmy teraz, że mamy jakiejkolwiek cpo D . Określimy ciąg zbiorów D_n zaczynając od $D_0 = D$ i indukcyjnie przyjmując $D_{n+1} = [D_n \rightarrow D_n]$. Oczywiście wszystkie D_n są cpo.

Dalej definiujemy przez indukcję ciąg projekcji (φ_n, ψ_n) z D_{n+1} na D_n , przyjmując

$$(\varphi_{n+1}, \psi_{n+1}) = (\varphi_n^*, \psi_n^*).$$

Mamy więc transmisję w obie strony:

$$\begin{aligned} D_0 &\xrightarrow{\varphi_0} D_1 \xrightarrow{\varphi_1} D_2 \xrightarrow{\varphi_2} \dots \\ D_0 &\xleftarrow{\psi_0} D_1 \xleftarrow{\psi_1} D_2 \xleftarrow{\psi_2} \dots \end{aligned}$$

Każdy ciąg $(x_n)_{n \in \mathbb{N}}$ złożony z elementów $x_n \in D_n$ i spełniający dla dowolnego n warunek $x_n = \psi_n(x_{n+1})$ nazywamy *nicią*. Tak wygląda nić:

$$x_0 \xleftarrow{\psi_0} x_1 \xleftarrow{\psi_1} x_2 \xleftarrow{\psi_2} \dots$$

Dla nici stosujemy notację $x = (x_n)_{n \in \mathbb{N}}$.

Zbiór wszystkich nici oznaczamy przez D_∞ . Można go uporządkować po współrzędnych:

$$x \leq y \quad \text{ wtedy i tylko wtedy, gdy } \quad \forall n \in \mathbb{N} (x_n \leq y_n).$$

Lemat 11.2 *Zbiór D_∞ jest zupełnym porządkiem częściowym.*

Dowód: Jeśli $X \subseteq D_\infty$ jest skierowany, to każdy ze zbiorów $X_n = \{x_n \mid x \in X\}$ jest skierowany, zatem istnieją kresy $y_n = \sup X_n$. Wtedy dla dowolnego n :

$$\psi_n(y_{n+1}) = \psi_n(\sup X_{n+1}) = \sup \psi_n(X_{n+1}) = \sup X_n = y_n,$$

bo funkcja ψ_n jest ciągła oraz $\psi_n(X_{n+1}) = X_n$. A zatem ciąg $y = (y_n)_{n \in \mathbb{N}}$ jest nicią (należy do D_∞). Pozostaje łatwe sprawdzenie, że $y = \sup X$. ■

Dla dowolnych $n \leq m$, mamy oczywiście projekcję $(\varphi_{nm}, \psi_{nm})$ z D_m na D_n , czyli złożenie

$$\varphi_{nm} = \varphi_{m-1} \circ \cdots \circ \varphi_n, \quad \psi_{nm} = \psi_n \circ \cdots \circ \psi_{m-1}.$$

Lemat 11.3 *Dla $n \leq m$ zachodzi $\varphi_{nm}^* = \varphi_{n+1, m+1}$ oraz $\psi_{nm}^* = \psi_{n+1, m+1}$.*

Można też określić projekcje $(\varphi_{n\infty}, \psi_{n\infty})$ z D_∞ na D_n . Dla $x \in D_n$ i $y \in D_\infty$:

$$(\varphi_{n\infty}(x))_i = \begin{cases} \psi_{in}(x), & \text{gdy } i < n; \\ x, & \text{gdy } i = n; \\ \varphi_{ni}(x), & \text{gdy } i > n. \end{cases} \quad \psi_{n\infty}(y) = y_n.$$

Od tej pory możemy uważać, że zachodzą inkluzje

$$D_0 \subseteq D_1 \subseteq D_2 \subseteq \cdots \subseteq D_\infty,$$

zadane przez powyższe projekcje. W szczególności każdy element $x \in D_n$ utożsamiamy z nicią prawie wszędzie równą x :

$$x_0 \xleftarrow{\psi_0} x_1 \xleftarrow{\psi_1} x_2 \xleftarrow{\psi_2} \cdots \xleftarrow{\psi_{n-2}} x_{n-1} \xleftarrow{\psi_{n-1}} x \xleftarrow{\psi_n} x \xleftarrow{\psi_{n+1}} x \xleftarrow{\psi_{n+2}} \cdots$$

bo przecież x uważamy za element wszystkich D_m dla $m \geq n$. Zauważmy, że mamy tu nierówności $x_0 \leq x_1 \leq x_2 \leq \cdots \leq x_{n-1} \leq x_n = x$.

Lemat 11.4

1. Każda nica x jest ciągiem wstępującym ($x_0 \leq x_1 \leq x_2 \leq \cdots$) i przy tym $x = \sup x_n$;
2. Najmniejszy element jest zawsze ten sam: $\perp_{D_0} = \perp_{D_n} = \perp_{D_\infty}$.

Dowód: Ćwiczenie. ■

Lemat 11.5 *Aplikacja w D_∞ określona wzorem*

$$x \cdot y = \sup\{x_{n+1}(y_n) \mid n \in \mathbb{N}\}$$

jest operacją ciągłą.

Dowód: Najpierw trzeba pokazać, że kres istnieje, elementy $x_{n+1}(y_n) \in D_n$ nie muszą bowiem tworzyć nici. Jest to jednak ciąg wstępujący, co wynika z takich związków pomiędzy

elementami zbioru D_n :

$$\begin{aligned}\varphi_{n-1}(x_n(y_{n-1})) &= \varphi_{n-1}(\psi_{n-1}(x_{n+1}(\varphi_{n-1}(y_{n-1}))) \leq \\ &x_{n+1}(\varphi_{n-1}(y_{n-1})) \leq x_{n+1}(\varphi_{n-1}(\psi_{n-1}(y_n))) \leq x_{n+1}(y_n)\end{aligned}$$

Dowód ciągłości korzysta z lematu 10.6. Ciągłość ze względu na x liczymy tak: Z jednej strony $\sup_{x \in X} x \cdot y = \sup_{x \in X} \sup_{n \in \mathbb{N}} x_{n+1}(y_n)$, a z drugiej strony

$$(\sup X) \cdot y = \sup_{n \in \mathbb{N}} (\sup X)_{n+1} \cdot y_n = \sup_{n \in \mathbb{N}} (\sup X_{n+1}) \cdot y_n = \sup_{n \in \mathbb{N}} \sup_{x \in X} x_{n+1}(y_n),$$

gdzie $X_{n+1} = \{x_{n+1} \mid x \in X\}$. A przecież to jest to samo. Ciągłość ze względu na y jest nawet łatwiejsza. Mamy bowiem $x \cdot \sup Y = \sup_n (x_{n+1}(\sup Y)) = \sup_n \sup_y x_{n+1}(y)$ oraz $\sup(x \cdot Y) = \sup_y \sup_n x_{n+1}(y)$. ■

Teraz trochę o tym, jak działa aplikacja.

Lemat 11.6

1. Jeśli $x \in D_{n+1}$ oraz $y \in D_n$ to $x \cdot y = x(y)$;
2. Jeśli $y \in D_n$ to $(x \cdot y)_n = x_{n+1}(y)$
3. Zawsze $(x \cdot \perp)_0 = x_0$.

Dowód: 1. W części pierwszej należy pokazać, że $x \cdot y = \sup_m x_{m+1}(y_m) = x_{n+1}(y_n)$. W tym celu zauważmy, że dla $m < n$ mamy

$$x_{m+1}(y_m) = \psi_{m+1,n+1}(x)(\psi_{mn}(y)) = \psi_{mn}^*(x_{n+1})(\psi_{mn}(y_n)) \leq \psi_{mn}(x_{n+1}(y_n)),$$

bo $\psi_{mn}^*(x_{n+1}) = \psi_{mn} \circ x_{n+1} \circ \varphi_{mn}$. Natomiast dla $m \geq n$

$$x_{m+1}(y_m) = \varphi_{n+1,m+1}(x)(\varphi_{nm}(y)) = \varphi_{nm}^*(x)(\varphi_{nm}(y)) = \varphi_{nm}(x(y)),$$

bo $\varphi_{nm}^*(x) = \varphi_{nm} \circ x \circ \psi_{nm}$.

2. W części drugiej, dla dowolnego $i \geq n$ mamy $y_{i+1} = \varphi(y_i)$, więc $x_{i+1}(y_i) = \psi_i(x_{i+2}(y_{i+1}))$, co implikuje $(x_{i+2}(y_{i+1}))_n = (x_{i+1}(y_i))_n = x_{n+1}(y)$. A wszystko to widać na takim obrazku:

$$\begin{array}{ccccc} y_i \in D_i & \xleftarrow{\psi_i} & & \xrightarrow{\varphi_i} & y_{i+1} \in D_{i+1} \\ \downarrow x_{i+1} & & & & \downarrow x_{i+2} \\ D_n & \xleftarrow{\psi_{ni}} & D_i & \xleftarrow{\psi_i} & D_{i+1} \end{array}$$

3. Część trzecia wynika z drugiej, dla $y = \perp \in D_0$. ■

Lemat 11.7 Aplikacja w D_∞ jest ekstensjonalna: jeśli $x \cdot z = y \cdot z$ zachodzi dla dowolnego z , to elementy x i y są równe.

Dowód: W istocie zachodzi implikacja

$$\text{jeśli } \forall z \in D_\infty (x \cdot z \leq y \cdot z) \text{ to } x \leq y.$$

Nierówność $x_0 \leq y_0$ wynika z lematu 11.6(3), bo mamy $x \cdot \perp \leq y \cdot \perp$. Dalej, z lematu 11.6(2),

$$x_{n+1}(z) = (x \cdot z)_n \leq (y \cdot z)_n = y_{n+1}(z),$$

dla dowolnego n i dowolnego $z \in D_n$, a więc $x_{n+1} \leq y_{n+1}$. ■

Lemat 11.8 *Niech $f \in D_{n+1}$ i $g \in D_{n+2}$ będą takie, że $\varphi_n(f(a)) \leq g(\varphi_n(a))$ dla $a \in D_n$. Wtedy $f \leq g$ w zbiorze D_∞ .*

Dowód: Należy sprawdzić, jakie nici wyznaczają f i g . Ponieważ $\varphi_n(f(a)) \leq g(\varphi_n(a))$ więc $f(a) = \psi_n(\varphi_n(f(a))) \leq \psi_n(g(\varphi_n(a))) = \psi_{n+1}(g)(a)$, czyli $f \leq \psi_{n+1}(g)$ w zbiorze D_{n+1} . Inaczej mówiąc $(f)_{n+1} \leq (g)_{n+1}$, a reszta wynika z monotoniczności funkcji φ_i i ψ_i . ■

Twierdzenie 11.9 *Zbiór D_∞ jest refleksywnym cpo, co więcej, jest on izomorficzny z przestrzenią funkcyjną $[D_\infty \rightarrow D_\infty]$.*

Dowód: Funkcja $F : D_\infty \rightarrow [D_\infty \rightarrow D_\infty]$ jest określona w oczywisty sposób:

$$F(x)(y) = x \cdot y.$$

Z lematów 11.5 i 11.7 wynika, że F jest ciągła i różnowartościowa. Aby sprawdzić, że jest to surjekcja, weźmy dowolną funkcję $f \in [D_\infty \rightarrow D_\infty]$ i niech $f^{(n)} : D_n \rightarrow D_n$ będzie jej „aproxymacją” do D_n , tj. niech $f^{(n)}(y) = f(y)_n$ dla $y \in D_n$.

Ciąg $f^{(n)}$ nie musi być nicią, ale jest wstępujący. Istotnie, zauważmy że elementy $y \in D_n$ oraz $\varphi_n(y) \in D_{n+1}$ wyznaczają tę samą nić, zatem z punktu widzenia D_∞ są po prostu równe. A więc $f(y) = f(\varphi_n(y))$ w D_∞ skąd $f^{(n)}(y) = f(y)_n \leq f(\varphi_n(y))_{n+1} = f^{(n+1)}(\varphi_n(y))$ zachodzi dla $y \in D_n$. Ściślej, mamy nierówność $\varphi_n(f^{(n)}(y)) \leq f^{(n+1)}(\varphi_n(y))$ pomiędzy elementami zbioru D_n , co na mocy lematu 11.8 oznacza $f^{(n)} \leq f^{(n+1)}$.

Skoro ciąg $f^{(n)}$ jest wstępujący, to ma kres $a = \sup_n f^{(n)}$. Trzeba teraz sprawdzić, że $F(a) = f$. Weźmy dowolne $y \in D_\infty$. Ponieważ ciągi $f^{(n)}$ i y_n są wstępujące, więc

$$F(a)(y) = a \cdot y = (\sup_n f^{(n)}) \cdot y = \sup_n f^{(n)}(y) = \sup_n f(y)_n = f(y).$$

Jeśli $f, g \in [D_\infty \rightarrow D_\infty]$ oraz $f \leq g$ to oczywiście także $\sup_n f^{(n)} \leq \sup_n g^{(n)}$. Oznacza to, że funkcja $G : [D_\infty \rightarrow D_\infty] \rightarrow D_\infty$, odwrotna do F jest monotoniczna. A zatem funkcja F jest izomorfizmem porządków, w szczególności F i G są ciągłe. ■

Wniosek 11.10 *Zbiór D_∞ wyznacza ekstensjonalny lambda-model $\mathcal{D}_\infty = \langle D_\infty, \cdot, \llbracket _ \rrbracket \rangle$.*

Twierdzenie 11.11 (Hyland, Wadsworth) *Termy M i N są obserwacyjnie równoważne wtedy i tylko wtedy, gdy $\mathcal{D}_\infty \models M = N$.*

Dowód: Można go znaleźć w książce Barendregta. ■

12 Model $\mathcal{P}_\omega$

Model $\mathcal{P}_\omega$ też jest pomysłu Scotta. Symbol P_ω oznacza po prostu zbiór $\mathbf{P}(\mathbb{N})$ uporządkowany przez zwykłą inkluzję. Konstrukcja tego modelu opiera się na tym, że krata $\langle P_\omega, \subseteq \rangle$ jest *algebraiczna* tj. każdy element jest sumą swoich skończonych podzbiorów. Konsekwencją tego jest to że funkcja ciągła nad P_ω jest zdeterminowana przez swoje zachowanie na zbiorach skończonych.

Lemat 12.1 *Funkcja $f : P_\omega \rightarrow P_\omega$ jest ciągła wtedy i tylko wtedy, gdy*

$$f(a) = \bigcup \{f(e) \mid e \text{ skończony oraz } e \subseteq a\},$$

dla dowolnego $a \in P_\omega$.

Ponieważ zbiorów skończonych jest przeliczalnie wiele, funkcji ciągłych jest continuum i można informację o takiej funkcji reprezentować za pomocą jednego elementu P_ω . W tym celu potrzebna jest tylko funkcja pary, na przykład:

$$(m, n) = \frac{(n + m)(n + m + 1)}{2} + m,$$

i kodowanie zbiorów skończonych, na przykład:

$$e_n = \{k_0, k_1, \dots, k_{r-1}\}, \quad \text{dla } n = \sum_{i < r} 2^{k_i}.$$

Teraz można określić funkcje

$$\text{graph} : [P_\omega \rightarrow P_\omega] \rightarrow P_\omega, \quad \text{fun} : P_\omega \rightarrow [P_\omega \rightarrow P_\omega]$$

następująco. Dla $f \in [P_\omega \rightarrow P_\omega]$ i $a, b \in P_\omega$ przyjmujemy

$$\begin{aligned} \text{graph}(f) &= \{(n, m) \mid m \in f(e_n)\}; \\ \text{fun}(a)(x) &= \{m \mid \exists n \in \mathbb{N}(e_n \subseteq x \wedge (n, m) \in a)\}. \end{aligned}$$

Lemat 12.2 *Funkcje graph i fun są ciągłe, a w dodatku $\text{fun} \circ \text{graph} = \text{id}_{[P_\omega \rightarrow P_\omega]}$.*

Dowód: Ćwiczenie. ■

A zatem P_ω jest refleksywnym cpo.

Wniosek 12.3 *Struktura $\mathcal{P}_\omega = \langle P_\omega, \cdot, \ll \rangle$, w której $x \cdot y = \text{fun}(x)(y)$, jest lambda-modelem.*

Fakt 12.4 *Model $\mathcal{P}_\omega$ nie jest ekstensjonalny.*

Dowód: Wystarczy jeśli udowodnimy, że $\mathcal{P}_\omega \not\models x = \lambda y. xy$. Niech $v(x) = a$. Wtedy $\llbracket x \rrbracket_v = a$. Natomiast $\llbracket \lambda y. xy \rrbracket_v = \text{graph}(g)$, gdzie $g(b) = \llbracket xy \rrbracket_{v[y \mapsto b]} = a \cdot b = \text{fun}(a)(b) =$

$\{m \mid \exists k \in \mathbb{N}(e_k \subseteq b \wedge (k, m) \in a)\}$. A więc $\llbracket \lambda y.xy \rrbracket_{[x \mapsto a]} = \{(n, m) \mid m \in g(e_n)\} = \{(n, m) \mid \exists k \in \mathbb{N}(e_k \subseteq e_n \wedge (k, m) \in a)\}$.

Jeśli a jest niepustym zbiorem skończonym, to $\llbracket \lambda y.xy \rrbracket_v$, jest zbiorem nieskończonym, w szczególności $\llbracket \lambda y.xy \rrbracket_v \neq a$. Istotnie, niech $(k, m) \in a$. Wtedy każda para (n, m) gdzie $e_k \subseteq e_n$ należy do $\llbracket \lambda y.xy \rrbracket_v$. ■

Okazuje się, że teoria modelu $\mathcal{P}_\omega$ to dokładnie teoria drzew Böhma.

Twierdzenie 12.5 (Hyland) *Dla dowolnych termów M, N*

$$\mathcal{P}_\omega \models M = N \iff BT(M) = BT(N)$$

Dowód: Można go znaleźć w książce Barendregta. ■

Ćwiczenia

1. Uzupełnić szczegóły dowodów.
2. Jakie nici są interpretacjami termów $\mathbf{I}, \mathbf{K}, \omega$, itd. w modelu $\mathcal{D}_\infty$?
3. Czy $m \in \text{fun}(a)(e_n)$ wtedy i tylko wtedy, gdy $(n, m) \in a$?
4. Jakie zbiory są interpretacjami termów $\mathbf{I}, \mathbf{K}, \omega$, itd. w modelu $\mathcal{P}_\omega$?
5. Pokazać, że $\mathcal{P}_\omega \not\models \mathbf{1} = \mathbf{I}$.
6. Funkcja ciągła z $\mathcal{P}_\omega$ w $\mathcal{P}_\omega$ jest jednoznacznie wyznaczona przez swoje wartości na zbiorach skończonych. Czy jest jednoznacznie wyznaczona przez swoje wartości na zbiorze pustym i singletonach?

13 Typy proste

Zaczynamy od składni typów. Przyjmijmy, że mamy pewien niepusty zbiór *typów atomowych*, który oznaczmy przez TV . Typy atomowe nazywamy też *zmiennymi typowymi*.⁷ Typy definiujemy indukcyjnie:

- Typy atomowe $p, q, r, \dots$ są typami;
- Jeśli σ i τ są typami, to $\sigma \rightarrow \tau$ jest typem.

Zbiór wszystkich typów oznaczmy przez T . Stosujemy taką konwencję, że strzałka jest łączna w prawo, tj. napis $\sigma \rightarrow \tau \rightarrow \rho$ oznacza $\sigma \rightarrow (\tau \rightarrow \rho)$. Zauważmy, że każdy typ można zapisać jako $\sigma_1 \rightarrow \dots \rightarrow \sigma_n \rightarrow p$, gdzie p jest typem atomowym.

Otoczenie typowe to częściowa funkcja ze zbioru zmiennych przedmiotowych w zbiór typów. Taką funkcję utożsamiamy ze zbiorem par postaci $(x : \tau)$ gdzie x jest zmienną przedmiotową

⁷Czasami zakłada się, że jest tylko jeden atom $\mathbf{0}$. My zwykle przyjmujemy, że może ich być dowolnie wiele.

a τ jest typem. Napis $\Gamma(x : \tau)$ oznacza to samo co $\Gamma[x \mapsto \tau]$, a więc $\Gamma(x : \tau)$ powstaje z Γ przez dodanie deklaracji $x : \tau$ poprzedzone usunięciem deklaracji $(x : \Gamma(x))$ jeśli była w Γ .

Przez *rachunek lambda z typami prostymi* (w wersji Curry'ego) rozumiemy system przypisania typów $\lambda \rightarrow$ o następujących regułach

$$\begin{array}{c} \text{(Var)} \quad \Gamma(x : \sigma) \vdash x : \sigma \\ \\ \text{(Abs)} \quad \frac{\Gamma(x : \sigma) \vdash M : \tau}{\Gamma \vdash (\lambda x M) : \sigma \rightarrow \tau} \qquad \text{(App)} \quad \frac{\Gamma \vdash M : \sigma \rightarrow \tau \quad \Gamma \vdash N : \sigma}{\Gamma \vdash (MN) : \tau} \end{array}$$

Aby dany term mógł mieć przypisany typ w otoczeniu Γ , konieczne jest jednoznaczne przypisanie typów wszystkim jego podtermom. Dlatego dobrze otypowany term można sobie wyobrażać jako term z adnotacjami typowymi „w stylu Churcha”. Takie adnotacje stosuje się czasem nieformalnie w formie górnych indeksów. Na przykład napiszemy $(\lambda x^\sigma. N^\tau)P^\sigma : \tau$, żeby zaznaczyć jaki typ użyty został dla x w wyprowadzeniu osądu $(\lambda x. N)P : \tau$.

Poniższe łatwe fakty pozostawiamy bez dowodu.

Lemat 13.1 (o generowaniu)

1. Jeśli $\Gamma \vdash MN : \sigma$, to $\Gamma \vdash M : \tau \rightarrow \sigma$ i $\Gamma \vdash N : \tau$ dla pewnego τ .
2. Jeśli $\Gamma \vdash x : \sigma$, gdzie x jest zmienną, to $\sigma = \Gamma(x)$.
3. Jeśli $\Gamma \vdash \lambda x M : \sigma$, oraz $x \notin \text{Dom}(\Gamma)$, to $\sigma = \rho \rightarrow \tau$ dla pewnych ρ i τ , takich że $\Gamma(x : \rho) \vdash M : \tau$.

Lemat 13.2 Jeśli $\Gamma, x : \sigma \vdash M : \tau$ oraz $\Gamma \vdash N : \sigma$, to $\Gamma \vdash M[x := N] : \tau$.

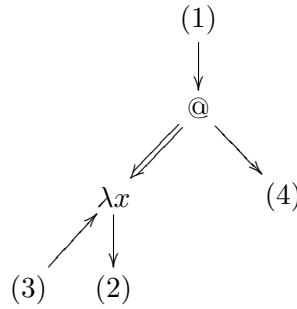
Wniosek 13.3 (poprawność redukcji) Jeśli $\Gamma \vdash M : \tau$ oraz $M \rightarrow_{\beta\eta} N$, to $\Gamma \vdash N : \tau$.

13.1 Normalizacja

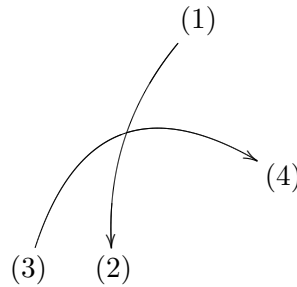
Udowodnimy teraz twierdzenie o normalizacji dla termów z typami: każdy typowalny term ma postać normalną. Zaczniemy od rozwiązania ćwiczenia 1 z rozdziału 4: jakie nowe β -redeksy mogą powstać w termie na skutek wykonania jednego kroku β -redukcji? Jeśli term przedstawimy w postaci grafu, to każdy redeks jest wyznaczony przez krawędź skierowaną w lewo od @ do λ (obrazek 13). Usunięcie tej krawędzi powoduje powstanie bezpośrednich połączeń od wierzchołka (1) do (2) oraz od (3) do (4), jak widać na obrazku 14. Uwaga: Obrazki 13 i 14 są nieco uproszczone. W istocie zamiast jednego wierzchołka (3) mamy ich tyle ile jest wystąpień zmiennej x w podtermie zaczepionym w punkcie (2).

Powstanie nowego redeksu jest możliwe, gdy na skutek skrócenia połączeń pomiędzy (1) i (2) oraz (3) i (4) powstanie krawędź łącząca bezpośrednio wierzchołek typu @ z wierzchołkiem typu λ . Może się to zdarzyć na 3 sposoby:

1. W pozycji (1) jest wierzchołek typu @, a w pozycji (2) jest wierzchołek typu λ . Mamy wtedy podterm postaci $(\lambda x \lambda y Q)NP$, który redukuje się do $(\lambda y Q[x := N])P$.



Obrazek 13: Eta-redeks przed redukcją...



Obrazek 14: ...i po redukcji.

2. W pozycji (3) jest wierzchołek typu @, a w pozycji (4) jest wierzchołek typu λ , tj. podterm $(\lambda x \dots x P \dots)(\lambda y Q)$ redukuje się do $\dots (\lambda y Q) P \dots$.
3. Wierzchołki (2) i (3) są sklejone w jeden, w pozycji (1) jest wierzchołek typu @, a w pozycji (4) jest wierzchołek typu λ . Tak jest w przypadku podtermu $(\lambda x x)(\lambda y Q)P$, który redukuje się do $(\lambda y Q)P$.

Oczywiście na skutek redukcji może także dojść do zwielokrotnienia redeksów już istniejących.

Twierdzenie 13.4 *Jeśli $\Gamma \vdash M : \tau$ to M ma postać normalną.*

Dowód: Rozważamy pewne ustalone wyprowadzenie osądu $\Gamma \vdash M : \tau$. To wyprowadzenie jednoznacznie określa przypisanie typów dla wszystkich termów, które można otrzymać z M w drodze redukcji. A zatem możemy bez obawy nieporozumienia stosować adnotacje typowe w formie indeksów. *Rangą* redeksu $(\lambda x^\sigma. P^\tau)Q^\sigma$ nazwiemy długość typu $\sigma \rightarrow \tau$, czyli typu przypisanego abstrakcji $(\lambda x. P)$. *Ranga* termu to maksymalna ranga występujących w nim redeksów.

Dowód twierdzenia przebiega przez indukcję ze względu na dwa parametry (n, m) , gdzie n jest rangą termu M a m jest liczbą redeksów rangi n występujących w M . Wystarczy udowodnić,

że na skutek redukcji odpowiednio wybranego redeksu w M para (n, m) musi się (leksyko-
graficznie) zmniejszyć. Wybrać należy redeks rangi n położony (zaczynający się) najdalej jak
to możliwe na prawo. Przy takiej redukcji zwielokrotnieniu mogą ulec tylko redeksy mniejszej
rangi (położone na prawo). Pozostaje zauważyć, że nowe redeksy powstające w wyniku naszej
redukcji też muszą być rangi mniejszej od n . Mamy bowiem trzy wcześniej wspomniane możli-
wości:

1. $(\lambda x^\alpha \lambda y^\beta. Q^\gamma) N^\alpha P^\beta \rightarrow (\lambda y^\beta Q[x := N]^\gamma) P^\beta$.
2. $(\lambda x^{\alpha \rightarrow \gamma} \dots x P^\alpha \dots)^\beta (\lambda y^\alpha Q^\gamma) \rightarrow \dots (\lambda y^\alpha Q^\gamma) P \dots$
3. $(\lambda x^{\alpha \rightarrow \beta} x^{\alpha \rightarrow \beta}) (\lambda y^\alpha Q^\beta) P^\alpha \rightarrow (\lambda y^\alpha Q^\beta) P^\alpha$.

W każdym z przypadków nowy redeks jest mniejszej rangi niż redeks wyeliminowany:

1. Abstrakcję typu $\alpha \rightarrow \beta \rightarrow \gamma$ zastąpiono przez abstrakcję typu $\beta \rightarrow \gamma$.
2. Abstrakcję typu $(\alpha \rightarrow \gamma) \rightarrow \beta$ zastąpiono przez abstrakcję typu $\alpha \rightarrow \gamma$.
3. Abstrakcję typu $(\alpha \rightarrow \beta) \rightarrow (\alpha \rightarrow \beta)$ zastąpiono przez abstrakcję typu $(\alpha \rightarrow \beta)$. ■

Silna normalizacja

Twierdzenie o normalizacji można wzmocnić: term rachunku lambda z typami redukuje się
zawsze do postaci normalnej, niezależnie od przyjętej strategii. Inaczej mówiąc, wszystkie
termy z typami mają własność silnej normalizacji (SN). Moralny sens tego faktu jest taki:
statyczna kontrola typów gwarantuje terminację obliczenia niezależnie od kolejności ewaluacji.
Program z typami może się zapętlić tylko wtedy, gdy jakiś jego składnik jawnie na to pozwala
(np. zawiera on pętlę, rekursję itp.).

Metoda dowodu silnej normalizacji (*computability method*), którą zastosujemy, pochodzi od
William'a Taita i bywa stosowana także dla dowodu innych własności języków z typami. W is-
tocie polega ona na konstrukcji pewnego syntaktycznego modelu dla typów prostych. Dla
dowolnego typu τ , określimy pewną interpretację typu τ w zbiorze termów, a mianowicie
zbiór termów *stabilnych*⁸ dla typu τ , który oznaczmy przez $\llbracket \tau \rrbracket$.

- $\llbracket p \rrbracket := \text{SN}$, gdy p jest atomem;
- $\llbracket \tau \rightarrow \sigma \rrbracket := \{M : \tau \rightarrow \sigma \mid \forall N (N \in \llbracket \tau \rrbracket \Rightarrow MN \in \llbracket \sigma \rrbracket)\}$.

Lemat 13.5 *Dla dowolnego typu τ :*

- 1) $\llbracket \tau \rrbracket \subseteq \text{SN}$;
- 2) *Jeśli $N_1, \dots, N_k \in \text{SN}$ to $xN_1 \dots N_k \in \llbracket \tau \rrbracket$. W szczególności zmienne są stabilne.*

⁸Inna, często spotykana terminologia: termy *obliczalne* lub *redukowalne*.

Dowód: Indukcja ze względu na τ . Niech na początek τ będzie atomem. Warunek (1) wynika wprost z definicji. Warunek (2) właściwie też, bo oczywiście $xN_1 \dots N_k \in \text{SN}$.

Teraz niech $\tau = \sigma \rightarrow \rho$ i niech $M \in \llbracket \sigma \rightarrow \rho \rrbracket$. Weźmy dowolną zmienną x . Z założenia indukcyjnego (2) mamy $x \in \llbracket \sigma \rrbracket$, więc z definicji $Mx \in \llbracket \rho \rrbracket$. A więc $Mx \in \text{SN}$, z założenia indukcyjnego (1). Tym bardziej $M \in \text{SN}$. Pokazaliśmy (1).

W punkcie (2) mamy pokazać, że $xN_1 \dots N_k P \in \llbracket \rho \rrbracket$ dla każdego $P \in \llbracket \sigma \rrbracket$. Ale $P \in \text{SN}$ z założenia indukcyjnego (1), więc teza wynika z założenia indukcyjnego (2) dla ρ . ■

Lemat 13.6 *Jeśli $M[x:=N_0]N_1 \dots N_k \in \text{SN}$ oraz $N_0 \in \text{SN}$, to także $(\lambda x.M)N_0N_1 \dots N_k \in \text{SN}$.*

Dowód: Przypuśćmy, że term $P_0 = (\lambda x.M)N_0N_1 \dots N_k$ ma nieskończony ciąg redukcji $P_0 \rightarrow P_1 \rightarrow P_2 \rightarrow \dots$. Ponieważ wszystkie termy $M, N_0, \dots, N_k$ są silnie normalizowalne, więc prędzej czy później będzie zredukowany redeks czołowy. Ściślej, dla pewnego n ,

$$P_n = (\lambda x.M')N'_0N'_1 \dots N'_k \rightarrow M'[x := N'_0]N'_1 \dots N'_k = P_{n+1},$$

gdzie $M \rightarrow M'$ oraz $N_i \rightarrow N'_i$ dla $i \leq k$. Ponieważ $M[x := N_0]N_1 \dots N_k \twoheadrightarrow P_{n+1}$, więc mamy sprzeczność z założeniem. ■

Lemat 13.7 *Jeśli $M[x:=N_0]N_1 \dots N_k \in \llbracket \tau \rrbracket$ oraz $N_0 \in \text{SN}$, to także $(\lambda x.M)N_0N_1 \dots N_k \in \llbracket \tau \rrbracket$.*

Dowód: Indukcja ze względu na τ . Przypadek typu atomowego to dokładnie lemat 13.6. Niech więc $\tau = \sigma \rightarrow \rho$, i niech $M[x := N_0]N_1 \dots N_k \in \llbracket \tau \rrbracket$. Mamy sprawdzić czy dla $P \in \llbracket \sigma \rrbracket$ zachodzi $(\lambda x.M)N_0N_1 \dots N_k P \in \llbracket \rho \rrbracket$. Ale skoro $M[x := N_0]N_1 \dots N_k \in \llbracket \tau \rrbracket = \llbracket \sigma \rightarrow \rho \rrbracket$ więc $M[x := N_0]N_1 \dots N_k P \in \llbracket \rho \rrbracket$ i wystarczy zastosować założenie indukcyjne dla ρ (przyjmując $N_{k+1} = P$). ■

Następujący lemat można uważać za twierdzenie o poprawności dla semantyki zadanej przez zbiory termów stabilnych.

Lemat 13.8 *Jeśli $\Gamma \vdash M : \tau$ oraz $N_i \in \llbracket \Gamma(x_i) \rrbracket$ dla $i \leq n$, to $M[x_1:=N_1, \dots, x_n:=N_n] \in \llbracket \tau \rrbracket$.*

Dowód: Indukcja ze względu na M . Jeśli M jest zmienną x_i , to teza wynika z lematu o generowaniu. Jeśli M jest inną zmienną, to korzystamy z lematu 13.5(2).

W przypadku gdy $M = PQ$, z lematu o generowaniu mamy $\Gamma \vdash P : \zeta \rightarrow \tau$ i $\Gamma \vdash Q : \zeta$ dla pewnego typu ζ . Z założenia indukcyjnego $P \in \llbracket \zeta \rightarrow \tau \rrbracket$ i $Q \in \llbracket \zeta \rrbracket$ więc $M = PQ \in \llbracket \tau \rrbracket$ wprost z definicji $\llbracket \zeta \rightarrow \tau \rrbracket$.

Niech teraz $M = \lambda y.P$. Typ τ jest wtedy postaci $\sigma \rightarrow \rho$, i do tego wiemy jeszcze, że $\Gamma(y : \sigma) \vdash P : \rho$. Wystarczy oczywiście pokazać, że $M[\vec{x} := \vec{N}] \in \llbracket \sigma \rightarrow \rho \rrbracket$. Weźmy więc jakiś term $Q \in \llbracket \sigma \rrbracket$ i rozpatrzmy aplikację $M[\vec{x} := \vec{N}]Q = (\lambda y.P)[\vec{x} := \vec{N}]Q = (\lambda y.P[\vec{x} := \vec{N}])Q$. Przyjeliśmy tu oczywiście, że zmienna y nie jest wolna w żadnym z termów $\vec{N}$. Oznacza to w szczególności, że $P[\vec{x} := \vec{N}][y := Q] = P[\vec{x} := \vec{N}, y := Q]$. Ten ostatni term jest stabilny dla ρ , co wiemy z założenia indukcyjnego. Zatem z lematu 13.7 wynika, że $M[\vec{x} := \vec{N}]Q$ też jest stabilny dla ρ . Pokazaliśmy więc ostatecznie, że $M[\vec{x} := \vec{N}]$ jest stabilny dla $\sigma \rightarrow \rho$. ■

Twierdzenie 13.9 (silna normalizacja) *Jeśli $\Gamma \vdash M : \tau$ dla pewnych Γ i τ to term M jest silnie normalizowalny.*

Dowód: Wystarczy przyjąć $n = 0$ w treści lematu 13.8. ■

14 Izomorfizm Curry’ego-Howarda

Typy proste można utożsamiać z formułami zdaniowymi zbudowanymi z atomów (zmiennych zdaniowych) za pomocą samej implikacji. Ta analogia nie jest przypadkowa, bo reguły wyprowadzania typów odpowiadają regułom wnioskowania:

$$\begin{array}{c} \text{(Ax)} \Gamma, \sigma \vdash \sigma \\ \\ \text{(W}\rightarrow\text{)} \frac{\Gamma, \sigma \vdash \tau}{\Gamma \vdash \sigma \rightarrow \tau} \qquad \text{(E}\rightarrow\text{)} \frac{\Gamma \vdash \sigma \rightarrow \tau \quad \Gamma \vdash \sigma}{\Gamma \vdash \tau} \end{array}$$

Reguły te otrzymaliśmy przez *wytarcie* z reguł systemu $\lambda_{\rightarrow}$ wszystkiego co dotyczy termów i pozostawienie tylko informacji o typach. Oczywiście symbol Γ , występujący w regułach wnioskowania, oznacza po prostu zbiór formuł. Logika określona tymi regułami to *minimalna logika implikacyjna*.

Jeśli Γ jest otoczeniem typowym, to przez $|\Gamma|$ oznaczmy zbiór formuł $\{\Gamma(x) \mid x \in \text{Dom}(\Gamma)\}$. Następujące twierdzenie wyraża w uproszczeniu odpowiedniość nazywaną często *izomorfizmem Curry’ego-Howarda*.

Twierdzenie 14.10

- Jeśli $\Gamma \vdash M : \tau$ w systemie $\lambda_{\rightarrow}$, to $|\Gamma| \vdash \tau$ w logice minimalnej.
- Jeśli $\Gamma \vdash \tau$ w logice minimalnej, to istnieje takie otoczenie typowe E , że $|E| = \Gamma$ oraz $E \vdash M : \tau$ dla pewnego M w systemie $\lambda_{\rightarrow}$.

Dowód: Łatwa indukcja. ■

W sensie ogólnym, izomorfizm Curry’ego-Howarda to właśnie ścisła odpowiedniość pomiędzy

- formułami i typami;
- dowodami i termami (programami).

Powstaje pytanie czy logika minimalna to to samo co implikacyjny fragment logiki klasycznej. Otóż nie. Następująca klasyczna tautologia, zwana *prawem Peirce’a*, nie jest twierdzeniem logiki minimalnej:

$$\pi = ((p \rightarrow q) \rightarrow p) \rightarrow p$$

Można to łatwo stwierdzić, korzystając z prostej obserwacji: jeśli istnieje term M typu τ , to istnieje też taki term w postaci normalnej. Pozostaje więc zbadać, jak może wyglądać zamknięty term w postaci normalnej typu π i przekonać się, że nijak.

A zatem na czym polega różnica? Logika minimalna jest *konstruktywna*. Implikacja w tej logice jest traktowana jak typ pewnej funkcji. Dowód formuły postaci $\alpha \rightarrow \beta$ musi polegać na wskazaniu metody przekształcenia każdego potencjalnego dowodu założenia α w poprawny dowód formuły β . Nie ma innego kryterium prawdy oprócz dowodu, w szczególności nie wystarczy odwołanie do dwuwartościowej semantyki.

Kombinatory z typami

Jak zauważyliśmy, termy rachunku lambda z typami prostymi odpowiadają ściśle dowodom w minimalnej logice implikacyjnej. Chodzi tu o dowody w systemie *naturalnej dedukcji*. Pojęcie dowodu formalnego często jednak kojarzy się nam nie z naturalną dedukcją, ale z podejściem Hilberta. To podejście jest bowiem zwykle prezentowane w podręcznikach logiki. Dla klasycznego implikacyjnego rachunku zdań, hilbertowski system dowodzenia można oprzeć na jednej regule wnioskowania (*modus ponens*)

$$\frac{\alpha, \alpha \rightarrow \beta}{\beta}$$

i trzech następujących schematach aksjomatów.

- $\alpha \rightarrow \beta \rightarrow \alpha$;
- $(\alpha \rightarrow \beta \rightarrow \gamma) \rightarrow (\alpha \rightarrow \beta) \rightarrow \alpha \rightarrow \gamma$;
- $((\alpha \rightarrow \beta) \rightarrow \alpha) \rightarrow \alpha$.

Trzeci z tych aksjomatów, prawo Peirce’a, jak już wiemy, nie jest twierdzeniem logiki minimalnej, ale pierwsze dwa tak. Są to mianowicie typy znanych nam termów zamkniętych:

- $\mathbf{K} = \lambda xy.x$;
- $\mathbf{S} = \lambda xyz.xz(yz)$.

Okazuje się, że jeśli z powyższego systemu odrzucimy trzeci aksjomat a pozostawimy dwa pierwsze, to otrzymamy system równoważny logice minimalnej. Piszemy $\Gamma \vdash_H \varphi$ gdy formuła φ ma dowód w takim systemie Hilberta ze zbioru dodatkowych założeń Γ . Równoważność naszego systemu Hilberta i naturalnej dedukcji wynika z następującego twierdzenia:

Twierdzenie 14.11 (o dedukcji) *Warunki $\Gamma \vdash_H \varphi \rightarrow \psi$ i $\Gamma, \varphi \vdash_H \psi$ są równoważne.*

Dowód: Oczywiście wszyscy znają dowód twierdzenia o dedukcji, ale przypomnijmy go⁹ z powodów metodologicznych. Dowód ten w swojej trudniejszej części — z prawej do lewej — przebiega przez indukcję ze względu na długość dowodu formuły ψ ze zbioru założeń Γ, φ .

⁹W istocie nasz dowód nieznacznie różni się od tego, który występuje w większości podręczników.

Rozpatrzmy najpierw kilka łatwych przypadków.

Pierwszy łatwy przypadek mamy wtedy, gdy w dowodzie nie użyto w ogóle założenia φ . Tak jest w szczególności wtedy, kiedy ψ jest aksjomatem lub $\psi \in \Gamma$, a dowód składa się tylko z tej jednej formuły. Inaczej mówiąc mamy w istocie dowód $\Gamma \vdash_H \psi$. Wtedy dowód formuły $\varphi \rightarrow \psi$ jest otrzymany przez oderwanie ψ od aksjomatu $\psi \rightarrow (\varphi \rightarrow \psi)$.

Jeśli $\psi = \varphi$, to wystarczy udowodnić, że $\vdash_H \varphi \rightarrow \varphi$, co zostawiamy jako ćwiczenie.

Pozostaje główny krok indukcyjny, gdy dowód formuły ψ ze zbioru Γ, φ otrzymano przez odrywanie. Znaczy to, że $\Gamma, \varphi \vdash \vartheta \rightarrow \psi$ i $\Gamma, \varphi \vdash \vartheta$ dla pewnej formuły ϑ , i że odpowiednie dowody są krótsze. Z założenia indukcyjnego otrzymujemy, że formuły $\varphi \rightarrow (\vartheta \rightarrow \psi)$ i $\varphi \rightarrow \vartheta$ mają dowody ze zbioru Γ . Aby otrzymać dowód dla $\varphi \rightarrow \psi$, należy te dwie formuły kolejno oderwać od drugiego aksjomatu w postaci $(\varphi \rightarrow (\vartheta \rightarrow \psi)) \rightarrow ((\varphi \rightarrow \vartheta) \rightarrow (\varphi \rightarrow \psi))$. ■

Wniosek 14.12 *Asercja $\Gamma \vdash \varphi$ jest wyprowadzalna w systemie naturalnej dedukcji wtedy i tylko wtedy, gdy $\Gamma \vdash_H \varphi$.*

Dowód: W obie strony mamy prostą indukcję. Twierdzenie 14.11 jest nam potrzebne w dowodzie części $(\Rightarrow)$ w przypadku wprowadzania implikacji. ■

Dla wnioskowania w stylu Hilberta też można mówić o odpowiedniości Curry’ego-Howarda, ale nie chodzi już teraz o rachunek lambda. Dowody w stylu Hilberta otrzymujemy wyłącznie przez stosowanie reguły *modus ponens*, która jak wiemy odpowiada aplikacji. Aksjomaty logiczne systemu hilbertowskiego można uważać za stałe odpowiednich typów. Dokładniej, dla dowolnych typów α, β, γ możemy postulować stałe $K_{\alpha\beta}$ i $S_{\alpha\beta\gamma}$, którym przypisujemy w dowolnym otoczeniu typy:

- $\vdash K_{\alpha\beta} : \alpha \rightarrow \beta \rightarrow \alpha$;
- $\vdash S_{\alpha\beta\gamma} : (\alpha \rightarrow \beta \rightarrow \gamma) \rightarrow (\alpha \rightarrow \beta) \rightarrow \alpha \rightarrow \gamma$.

Inna możliwość, to założyć, że w systemie są tylko dwie stałe K i S , którym można (w dowolnym otoczeniu) przypisać każdy z typów odpowiedniej postaci. W ten sposób otrzymamy *rachunek kombinatorów z typami prostymi*, a wniosek 14.12 będziemy teraz mogli odczytać tak:

(*) *Typy niepuste w rachunku lambda i rachunku kombinatorów z typami prostymi są takie same.*

Uderzające jest podobieństwo pomiędzy dowodem twierdzenia 14.11 i definicją kombinatorycznej abstrakcji λ^* (ćwiczenie 8). Ciekawe jest to, że ta analogia sięga znacznie dalej niż tylko zgodność typów. Można powiedzieć, że twierdzenie o dedukcji ma sens wykraczający poza język formuł logicznych.

Typologia ogólna: spójniki logiczne

Odpowiedniość pomiędzy formułami i typami byłaby niepełna, gdyby nie rozciągała się na spójniki logiczne inne niż implikacja. W istocie nietrudno jest zinterpretować w tym duchu

zarówno koniunkcję jak i alternatywę. Zaczniemy od przypomnienia reguł naturalnej dedukcji dla tych spójników.

$$\begin{array}{c} \frac{\Gamma \vdash \varphi \quad \Gamma \vdash \psi}{\Gamma \vdash \varphi \wedge \psi} (W\wedge) \qquad \frac{\Gamma \vdash \varphi \wedge \psi}{\Gamma \vdash \varphi} (E\wedge) \quad \frac{\Gamma \vdash \varphi \wedge \psi}{\Gamma \vdash \psi} \\[10pt] \frac{\Gamma \vdash \varphi}{\Gamma \vdash \varphi \vee \psi} (W\vee) \quad \frac{\Gamma \vdash \psi}{\Gamma \vdash \varphi \vee \psi} \qquad \frac{\Gamma \vdash \varphi \vee \psi \quad \Gamma, \varphi \vdash \vartheta \quad \Gamma, \psi \vdash \vartheta}{\Gamma \vdash \vartheta} (E\vee) \end{array}$$

Reguły te możemy objaśniać w myśl tzw. *interpretacji BHK* (od nazwisk: Brouwer, Heyting, Kołmogorow).

- Dowód koniunkcji składa się z dowodów jej składowych, jest więc (uporządkowaną) parą dowodów;
- Dowód alternatywy to dowód jednego z jej członów (ze wskazaniem którego).

A zatem koniunkcja odpowiada iloczynowi kartezjańskiemu (rekordowi) a alternatywa sumie prostej (wariantowi). Wprowadzanie koniunkcji to tworzenie pary, a eliminacja koniunkcji to rzutowanie.

$$\frac{\Gamma \vdash M : \varphi \quad \Gamma \vdash N : \psi}{\Gamma \vdash \langle M, N \rangle : \varphi \wedge \psi} (W\wedge) \qquad \frac{\Gamma \vdash M : \varphi \wedge \psi}{\Gamma \vdash \pi_1(M) : \varphi} (E\wedge) \quad \frac{\Gamma \vdash M : \varphi \wedge \psi}{\Gamma \vdash \pi_2(M) : \psi}$$

Wprowadzanie alternatywy odpowiada tworzeniu obiektu wariantowego. Eliminacja alternatywy to instrukcja wyboru.

$$\frac{\Gamma \vdash M : \varphi}{\Gamma \vdash \mathbf{inl}(M) : \varphi \vee \psi} (W\vee) \quad \frac{\Gamma \vdash M : \psi}{\Gamma \vdash \mathbf{inr}(M) : \varphi \vee \psi} \quad \frac{\Gamma \vdash M : \varphi \vee \psi \quad \Gamma, x : \varphi \vdash P : \vartheta \quad \Gamma, y : \psi \vdash Q : \vartheta}{\Gamma \vdash \mathbf{case } M \mathbf{ of } [x]P, [y]Q : \vartheta} (E\vee)$$

Dla tak rozszerzonego języka możemy teraz postulować następujące redukcje.

Beta redukcje:

$$\begin{aligned} \pi_1(\langle M, N \rangle) &\rightarrow M, \quad \pi_2(\langle M, N \rangle) \rightarrow N; \\ \mathbf{case } \mathbf{inl}(P) \mathbf{ of } [x]M, [y]N &\rightarrow M[x := P]; \\ \mathbf{case } \mathbf{inr}(Q) \mathbf{ of } [x]M, [y]N &\rightarrow N[y := Q]. \end{aligned}$$

Eta redukcje:

$$\langle \pi_1(M), \pi_2(M) \rangle \rightarrow M; \qquad (\mathbf{case } M \mathbf{ of } [x]\mathbf{inl } x, [y]\mathbf{inr } y) \rightarrow M.$$

Mówimy tu o beta i eta redukcjach przez analogię z redukcjami dla implikacji. Beta redukcja odpowiada wprowadzeniu spójnika, po którym natychmiast następuje jego eliminacja. Natomiast postulat eta-redukcji odpowiada założeniu, że każde wyrażenie danego typu opisuje pewien *kanoniczny obiekt* tego typu (funkcję, parę, wariant — ogólnie rezultat operacji wprowadzenia).

Pozostaje sprawa negacji, którą jednak możemy zinterpretować jako specyficzną implikację:

$$\neg\alpha = \alpha \rightarrow \perp$$

Symbol $\perp$ oznacza fałsz, który oczywiście nie może mieć dowodu. Nie ma więc kanonicznych obiektów typu $\perp$ — jest to typ pusty. Mamy jednak regułę eliminacji fałszu:

$$\frac{\Gamma \vdash M : \perp}{\Gamma \vdash \varepsilon_\varphi(M) : \varphi} (E\perp)$$

Symbol ε_φ oznacza cud typu φ . Skoro mamy *rzecz, której nie ma*, to znaczy, że możemy z niej zrobić co zechcemy.

Ćwiczenia

1. Udowodnić twierdzenie 14.10.
2. Udowodnić że nie istnieje kombinator typu $\pi = ((p \rightarrow q) \rightarrow p) \rightarrow p$.
3. Pokazać, że **K** jest jedynym termem zamkniętym w postaci normalnej, który ma typ $s \rightarrow t \rightarrow s$ (gdzie s i t są zmiennymi typowymi). Pokazać analogiczną własność dla **S**.
4. Wyprowadzić $\vdash_H \alpha \rightarrow \alpha$. *Wskazówka: Zbudować ze statych K i S term takiego typu.*
5. Sformułować reguły przypisania typów dla rachunku kombinatorów z typami prostymi i udowodnić stwierdzenie (*).
6. Udowodnić silną normalizację dla rachunku kombinatorów z typami prostymi.
7. Znaleźć typy dla kombinatorów W, B, B' i C.
8. Udowodnić, że w rachunku kombinatorów z typami prostymi warunek $\Gamma, x : \tau \vdash M : \sigma$ implikuje $\Gamma \vdash \lambda^*x.M : \tau \rightarrow \sigma$.

15 Problemy decyzyjne

Zgodnie z izomorfizmem Curry'ego-Howarda, twierdzenia intuicjonistycznej logiki implikacyjnej to dokładnie typy zamkniętych termów rachunku lambda (lub logiki kombinatorycznej). Ponieważ każdy term typowalny ma postać normalną, więc dla ustalenia, czy istnieje term danego typu wystarczy szukać takiego termu w postaci normalnej. Prowadzi to do następującego *algorytmu Ben-Yellesa*, który rozwiązuje nieco ogólniej postawione zadanie:

- Dane są otoczenie Γ i typ τ ;
- Szukamy takiego termu M , że $\Gamma \vdash M : \tau$;

Zadanie to można rozbić na dwa przypadki:

- (1) Jeśli $\tau = \tau_1 \rightarrow \tau_2$, to można zakładać, że M musi być abstrakcją postaci $\lambda x.M'$. (Jeśli term M nie jest abstrakcją, to zamiast M można wziąć term $\lambda x.Mx$, gdzie x jest nowe.) Należy więc znaleźć taki term M' aby $\Gamma, x : \tau_1 \vdash M' : \tau_2$.

- (2) Jeśli τ jest zmienną typową s , to term M musi być aplikacją postaci $xM_1 \dots M_k$, gdzie $\Gamma(x) = \sigma_1 \rightarrow \dots \rightarrow \sigma_k \rightarrow s$. Należy więc znaleźć termy $M_1, \dots, M_k$ spełniające warunki $\Gamma \vdash M_1 : \sigma_1, \dots, \Gamma \vdash M_k : \sigma_k$.

Ponieważ w przypadku (2) możemy mieć do wyboru więcej niż jedną zmienną o typie kończącym się na s , więc nasz algorytm jest w istocie niedeterministycznym algorytmem rekurencyjnym, lub jak kto woli — algorytmem alternującym.¹⁰

Oczywiście, jeśli „inhabitant” typu τ istnieje, to prędzej czy później znajdziemy go tą metodą. Ale może być tak, że nasz algorytm się zapętli — weźmy na przykład typ $(s \rightarrow s) \rightarrow s$, albo $((s \rightarrow t) \rightarrow t) \rightarrow s \rightarrow t$. Jeśli więc rozwiązanie zadania dla danych Γ i τ prowadzi ponownie do tego samego zadania, przerywamy obliczenie z wynikiem negatywnym. Aby uniknąć sytuacji, w której pojawia się nieskończenie wiele różnych zadań, musimy też nieco poprawić działanie algorytmu w przypadku (1).

- (1a) Jeśli $\tau = \tau_1 \rightarrow \tau_2$, oraz w otoczeniu Γ nie ma zmiennej typu τ_1 , to szukamy takiego M' aby $\Gamma, x : \tau_1 \vdash M' : \tau_2$.
- (1b) Jeśli $\tau = \tau_1 \rightarrow \tau_2$, oraz w otoczeniu Γ jest zmienna typu τ_1 , to szukamy takiego M' aby $\Gamma \vdash M' : \tau_2$.

Poprawność przypadku (1b) wynika z następującej prostej obserwacji:

$$\text{Jeśli } \Gamma(y) = \tau_1 \text{ oraz } \Gamma, x : \tau_1 \vdash M' : \tau_2, \text{ to } \Gamma \vdash M'[x := y] : \tau_2.$$

Inaczej mówiąc, wystarczy po jednej zmiennej każdego typu. Liczba typów, które mogą się pojawić w obliczeniu jest proporcjonalna do rozmiaru zadania, a otoczenie stale rośnie. Zatem głębokość rekursji jest wielomianowa (kwadratowa), a stąd wynika, że cały algorytm jest w klasie PSPACE. I nie da się go istotnie ulepszyć.

Twierdzenie 15.1 (Statman) *Implikacyjna logika intuicjonistyczna jest PSPACE-zupełna. A zatem problem niepustości dla typów prostych jest też PSPACE-zupełny.*

Dowód: Redukujemy problem kwantyfikowanych formuł Boole’owskich (QBF), czyli klasyczną logikę zdaniową drugiego rzędu, do intuicjonistycznego implikacyjnego rachunku zdań.¹¹

Niech Φ będzie formułą w języku QBF. Bez straty ogólności możemy założyć, że jest to zdanie i że negacja występuje w Φ tylko w kontekstach postaci $\neg p$, gdzie p jest zmienną zdaniową. Dla porządku założymy jeszcze, że wszystkie zmienne związane w Φ są różne.

Dla dowolnej zmiennej zdaniowej p , która występuje w formule Φ , niech s_p i $s_{\neg p}$ będą nowymi zmiennymi typowymi. Podobnie, dla każdej podformuły φ formuły Φ wybierzmy nowe zmienne typowe s_φ i $s_{\neg\varphi}$. Przez Γ_Φ oznaczmy zbiór złożony z następujących formuł:

- $(s_p \rightarrow s_\psi) \rightarrow (s_{\neg p} \rightarrow s_\psi) \rightarrow s_\varphi$, dla każdej podformuły φ postaci $\forall p\psi$;

¹⁰Krok egzystencjalny to wybór zmiennej x , krok uniwersalny to wybór jednego z termów M_i .

¹¹Jak widać, logika minimalna wcale nie jest taka minimalna.

- $(s_p \rightarrow s_\psi) \rightarrow s_\varphi$ i $(s_{\neg p} \rightarrow s_\psi) \rightarrow s_\varphi$, dla każdej podformuły φ postaci $\exists p\psi$;
- $s_\psi \rightarrow s_\vartheta \rightarrow s_\varphi$, dla każdej podformuły φ postaci $\psi \wedge \vartheta$;
- $s_\psi \rightarrow s_\varphi$ i $s_\vartheta \rightarrow s_\varphi$, dla każdej podformuły φ postaci $\psi \vee \vartheta$.

Jeśli teraz v jest wartościowaniem zerojedynkowym, to Γ_v oznacza Γ_Φ rozszerzone o zmienne

- s_p , gdy $v(p) = 1$;
- $s_{\neg p}$, gdy $v(p) = 0$,

dla wszystkich p . Teraz przez łatwą indukcję ze względu na długość podformuły φ , dowodzimy, że dla dowolnego v zachodzi równoważność

$$\Gamma_v \vdash s_\varphi \quad \text{wtedy i tylko wtedy, gdy} \quad v(\varphi) = 1.$$

W szczególności Φ jest prawdziwa wtedy i tylko wtedy, gdy $\Gamma_\Phi \vdash s_\Phi$. Pozostaje zauważyć, że warunek $\Gamma_\Phi \vdash s_\Phi$ to to samo co $\gamma_1 \rightarrow \dots \rightarrow \gamma_n \rightarrow s_\Phi$ dla odpowiednich γ_i , oraz, że całą konstrukcję można przeprowadzić w pamięci logarytmicznej. ■

Rekonstrukcja typu

Oczywiście nie każdy term jest typowalny w systemie $\lambda_{\rightarrow}$, nawet jeśli jest w postaci normalnej. Przykładem jest choćby $\lambda x.xx$. Problem typowości jest równoważny (ze względu na redukcje w LOGSPACE) problemowi unifikacji dla języka pierwszego rzędu z jedną operacją binarną $\rightarrow$, i dlatego mamy następujący fakt:

Twierdzenie 15.2 *Problem typowości dla $\lambda_{\rightarrow}$ jest w klasie P.* ■

Problem sprawdzenia typu dla $\lambda_{\rightarrow}$ też jest w klasie P. Inny dowód rozstrzygalności obu problemów polega na sprowadzeniu ich do zadania znajdowania cyklu w skończonym grafie.

Równość

Rozpatrzmy term $\mathbf{2} \cdots \mathbf{2}xy$, w którym występuje n dwójek. Nietrudno się przekonać, że postacią normalną tego termu jest $x(x(\cdots(xy)\cdots))$, gdzie liczbą wystąpień zmiennej x jest

$$2^{2^{\cdots^2}} \Big\}_n$$

Najprostszy algorytm sprawdzający, czy dwa termy tego samego typu są beta-równe, polega na obliczeniu i porównaniu ich postaci normalnych. Jak widać z powyższego przykładu rozmiar postaci normalnej może być bardzo duży w stosunku do rozmiaru danego termu. Nawet jeśli uwzględnimy rozmiary użytych *typów* sytuacja zmieni się niewiele. Zauważmy bowiem, że typy, których potrzebujemy aby wywnioskować

$$x : t \rightarrow t, y : t \vdash \mathbf{2} \cdots \mathbf{2}xy : t$$

są „zaledwie” wykładniczego rozmiaru (każda kolejna dwójka ma dwa razy dłuższy typ). A zatem nasz naiwny algorytm ma nieelementarną złożoność. Co gorsza, nie można go istotnie poprawić.

Twierdzenie 15.3 (Statman) *Problem równości termów w rachunku lambda z typami prostymi nie jest problemem elementarnym (nie istnieje algorytm rozwiązujący ten problem w czasie $2^{\cdot^{2 \cdot^2}} \bigg\}_k$, dla żadnego ustalonego k).* ■

Co innego jednak porównać dwa termy, a co innego znaleźć term, lub termy spełniające zadane równanie. Rozwiązywanie równań z niewiadomymi dowolnych typów skończonych nazywamy *unifikacją wyższego rzędu*. Aby ściśle sformułować problem unifikacji przyjmijmy, że dane jest otoczenie Γ i para termów (M, N) , które w tym otoczeniu mają ten sam typ. Wśród zmiennych deklarowanych w Γ wyróżnimy *niewiadome* $x_1, \dots, x_k$ a wszystkie pozostałe zmienne nazwijmy *parametrami*. *Rozwiązaniem* równania $M = N$ o niewiadomych $x_1, \dots, x_k$ nazywamy dowolne termy $P_1, \dots, P_k$ spełniające warunki

- $\Gamma \vdash P_i : \Gamma(x_i)$ dla $i = 1, \dots, k$;
- $M[x_1 := P_1, \dots, x_k := P_k] =_{\beta\eta} N[x_1 := P_1, \dots, x_k := P_k]$.

Problem unifikacji wyższego rzędu to następujący problem decyzyjny: czy istnieje rozwiązanie danego równania? Zwykle zakłada się, że wszystkie typy występujące w zadaniu są zbudowane z jednego atomu. Nie zmienia to istotnie stopnie trudności zadania. Jeśli typy wszystkich niewiadomych są postaci $t \rightarrow t \rightarrow \cdots t \rightarrow t$ lub t , to mówimy o unifikacji drugiego rzędu. Zwykłą unifikację nazywamy też unifikacją pierwszego rzędu (niewiadome mają typ atomowy).

Przykład 15.4

Oto dwa przykłady unifikacji drugiego rzędu z parametrami $f : t \rightarrow t \rightarrow t$, $a : t$, $b : t$ i niewiadomymi $G : t \rightarrow t \rightarrow t$, $F : t \rightarrow t$.

- $Ga(fba) = fa(Gab)$;
- $fa(Fa) = F(faa)$.

Pierwszy przykład nie ma rozwiązania, drugi ma ich nieskończenie wiele. Rozwiązaniami są wszystkie termy postaci $\lambda x.f a(f a(f a(\cdots (f a x) \cdots)))$.

Twierdzenie 15.5 (Goldfarb) *Unifikacja drugiego rzędu jest nierozstrzygalna.* ■

Dowód twierdzenia Goldfarba polega na redukcji Dziesiątego Problemu Hilberta do unifikacji drugiego rzędu. Dla dowolnego wielomianu konstruuje się takie równanie unifikacyjne, które ma rozwiązanie wtedy i tylko wtedy, gdy dany wielomian ma zero całkowite.

Szczególnym przypadkiem unifikacji jest *dopasowanie* (matching). *Problem dopasowania wyższego rzędu* polega na rozwiązaniu równania, w którym niewiadome występują tylko po lewej stronie. Wiadomo, że dopasowanie rzędu czwartego jest rozstrzygalne. Problem jest otwarty dla wyższych rzędów.¹² Jeśli w zadaniu dopasowania zmienić $=_{\beta\eta}$ na $=_{\beta}$, to problem okazuje się być nierozstrzygalny. Udowodnił to kilka lat temu Ralph Loader.

Funkcje obliczalne

Jak pamiętamy, w beztypowym rachunku lambda można reprezentować wszystkie funkcje rekurencyjne, a nawet częściowo rekurencyjne (twierdzenie 7.9). W rachunku lambda z typami prostymi nie należy się spodziewać podobnego wyniku, bo równość termów jest rozstrzygalna. W istocie klasa funkcji definiowalnych w rachunku $\lambda_{\rightarrow}$ jest dość uboga. Ale musimy zacząć od odpowiedniej definicji.

Niech σ będzie dowolnym typem. Przez ω_{σ} oznaczamy typ $(\sigma \rightarrow \sigma) \rightarrow \sigma \rightarrow \sigma$. Jeśli σ jest nieistotne (zwykle jest to ustalony atom), to piszemy po prostu ω . Oczywiście wszystkie liczebники Churcha mają typ ω , więc naturalna jest następująca definicja:

Funkcja $f : \mathbb{N}^k \rightarrow \mathbb{N}$ jest β -definiowalna (lub po prostu *definiowalna*) w rachunku lambda z typami prostymi, w typie ω_{σ} , jeżeli istnieje term zamknięty F , spełniający następujące warunki:

- $\vdash F : \omega_{\sigma} \rightarrow \dots \rightarrow \omega_{\sigma} \rightarrow \omega_{\sigma}$;
- Dla dowolnych $n_1, \dots, n_k \in \mathbb{N}$, jeżeli $f(n_1, \dots, n_k) = m$ to $F \mathbf{n}_1 \dots \mathbf{n}_k =_{\beta} \mathbf{m}$.

Zamieniając w powyższej definicji znak $=_{\beta}$ na $=_{\beta\eta}$ otrzymujemy klasę funkcji $\beta\eta$ -definiowalnych.

Następnik, dodawanie i mnożenie są przykładami funkcji definiowalnych w $\lambda_{\rightarrow}$ (zob. przykład 7.3). Możemy też zdefiniować funkcję warunkową

$$\text{ifzero}(p, q, r) = \begin{cases} q, & \text{jeśli } p = 0; \\ r, & \text{w przeciwnym przypadku.} \end{cases}$$

za pomocą termu

$$\text{ifzero} = \lambda pqr \lambda fx. p(\lambda y. qfx)(rfx).$$

Jeśli zostaniemy przy równości β to więcej zdefiniować nam się nie uda.

Twierdzenie 15.6 (Schwichtenberg) *Dla dowolnego σ klasa funkcji β -definiowalnych w $\lambda_{\rightarrow}$ w typie ω_{σ} to dokładnie klasa wielomianów warunkowych, tj. najmniejsza klasa funkcji nad $\mathbb{N}$, zawierająca*

- rzutowania;
- dodawanie;
- mnożenie;

¹²Ostatnio C. Stirling ogłosił, że problem jest rozstrzygalny.

- funkcje stałe 0 i 1;
- funkcję warunkową *ifzero*,

i zamknięta ze względu na składanie. ■

Z twierdzenia Schwichtenberga wynika, w szczególności, że wybór typu σ nie ma znaczenia. Dotąd uważano, że tak samo jest w przypadku $\beta\eta$ -definiowalności, ale Mateusz Zakrzewski niedawno pokazał, że np. funkcja

$$\text{ifeven}(p, q, r) = \begin{cases} q, & \text{jeśli } p \text{ jest parzyste;} \\ r, & \text{w przeciwnym przypadku,} \end{cases}$$

która nie jest wielomianem warunkowym, jest $\beta\eta$ -definiowalna (gdy liczebniki interpretujemy w odpowiednio dobranym typie ω_σ).

Nawet jednak z użyciem $\beta\eta$ -konwersji, tak proste funkcje jak poprzednik, odejmowanie i potęgowanie nie są definiowalne. Poprzednik i potęgę uda nam się zdefiniować jeśli jeszcze bardziej osłabimy nasze wymagania. Powiemy, że funkcja $f : \mathbb{N}^k \rightarrow \mathbb{N}$ jest *niejednostajnie definiowalna* w $\lambda_\rightarrow$, jeśli istnieje term F spełniający warunki:

- $\vdash F : \omega_{\sigma_1} \rightarrow \dots \rightarrow \omega_{\sigma_k} \rightarrow \omega_\sigma$;
- Dla dowolnych $n_1, \dots, n_k \in \mathbb{N}$, jeżeli $f(n_1, \dots, n_k) = m$ to $F\mathbf{n}_1 \dots \mathbf{n}_k =_\beta \mathbf{m}$.

Okazuje się, że poprzednik i potęgowanie są definiowalne niejednostajnie, ale już odejmowanie nie jest. Klasę funkcji arytmetycznych definiowalnych w skończonych typach może istotnie powiększyć dopiero dodanie do systemu „prawdziwego” operatora iteracji. W ten sposób otrzymujemy tzw. System **T** Gödla.

Ćwiczenia

1. Udowodnić twierdzenie 15.2.
2. Niech τ będzie typem, w którym występują tylko zmienne typowe $s_1, \dots, s_n$ i niech Γ będzie takim otoczeniem, że $\Gamma(x_i) = s_i$ dla $i = 1, \dots, n$. Skonstruować taki term M_τ , że $\Gamma \vdash M : \tau$.
3. Niech $E = \{\tau_1 = \sigma_1, \dots, \tau_k = \sigma_k\}$ będzie instancją problemu unifikacji dla języka pierwszego rzędu z jedną operacją binarną $\rightarrow$ o niewiadomych $s_1, \dots, s_n$. Skonstruować (w pamięci logarytmicznej) taki term M , który jest typowalny wtedy i tylko wtedy, gdy E ma rozwiązanie. *Wskazówka:* Skorzystać z poprzedniego zadania.
4. Zbadać rozwiązalność unifikacji z przykładu 15.4.
5. (Łatwe) Niech $c_n = \lambda x. fa(fa(fa(\dots(fax)\dots)))$, gdzie f występuje n razy. Skonstruować takie równanie drugiego rzędu, którego rozwiązaniami są dokładnie trójki termów postaci c_n, c_m, c_{n+m} .
6. (Trudne) Skonstruować takie równanie drugiego rzędu, którego rozwiązaniami są dokładnie trójki termów postaci $c_n, c_m, c_{n \cdot m}$.
7. (Łatwe, jeśli umiemy rozwiązać poprzednie dwa zadania.) Udowodnić twierdzenie 15.5.
8. Udowodnić, że poprzednik i potęgowanie są niejednostajnie definiowalne w skończonych typach. Dlaczego są kłopoty z odejmowaniem?

16 Semantyka typów prostych

W teorii typów prostych często przyjmuje się dwa założenia, które istotnie upraszczają pewne konstrukcje. Po pierwsze, zamiast równości $=_\beta$ rozważa się ekstensjonalną równość $=_{\beta\eta}$, po drugie zakłada się, że wszystkie typy są zbudowane z jednego tylko typu bazowego. My też teraz przyjmijmy te założenia. A więc typy definiujemy tak:

- Stała 0 jest typem;
- Jeśli σ i τ są typami, to $\sigma \rightarrow \tau$ jest typem.

Niech T oznacza zbiór wszystkich takich typów. Jeśli teraz $\{D_\sigma\}_{\sigma \in T}$ jest rodziną niepustych zbiorów, to możemy rozważać *wartościowania*, które zmiennym typu σ przypisują elementy D_σ . (Wygodnie jest zakładać, że nasze termy są w ortodoksyjnym stylu Churcha.) Rozważamy wielosortowe struktury postaci $\mathcal{A} = \langle \{D_\sigma\}_{\sigma \in T}, \{\cdot_{\sigma\tau}\}_{\sigma, \tau \in T}, \llbracket \cdot \rrbracket^{\mathcal{A}} \rangle$, gdzie dla dowolnych $\sigma, \tau \in T$:

- D_σ jest niepustym zbiorem;
- $d \cdot_{\sigma\tau} e \in D_\tau$ dla $d \in D_{\sigma \rightarrow \tau}$, $e \in D_\sigma$;
- $\llbracket M \rrbracket_v^{\mathcal{A}} \in D_\sigma$, dla dowolnego M typu σ .

Oczywiście zamiast $\cdot_{\sigma\tau}$ będziemy pisać zwykłą kropkę, a zamiast $\llbracket \cdot \rrbracket^{\mathcal{A}}$ napiszemy $\llbracket \cdot \rrbracket$. Taka struktura jest (ekstensjonalnym) *modelem*, gdy spełnia następujące warunki:

- Jeśli $d, d' \in D_{\sigma \rightarrow \tau}$ oraz $d \cdot e = d' \cdot e$ dla dowolnego $e \in D_\sigma$, to $d = d'$.
- Jeśli x jest zmienną, to $\llbracket x \rrbracket_v = v(x)$;
- $\llbracket PQ \rrbracket_v = \llbracket P \rrbracket_v \cdot \llbracket Q \rrbracket_v$;
- $\llbracket \lambda x^\sigma P \rrbracket_v \cdot a = \llbracket P \rrbracket_{v[x \mapsto a]}$, dla dowolnego $a \in D_\sigma$.

Założenie ekstensjonalności powoduje, że powyższe warunki w istocie jednoznacznie definiują funkcję $\llbracket \cdot \rrbracket$ dla danego $\langle \{D_\sigma\}_{\sigma \in T}, \{\cdot_{\sigma\tau}\}_{\sigma, \tau \in T} \rangle$, o ile taka funkcja istnieje (tj. istnieją wszystkie elementy potrzebne do zinterpretowania abstrakcji). Z ekstensjonalności wynika też przez łatwą indukcję pominięty w definicji¹³ warunek:

- Jeśli $v|_{\text{FV}(P)} = u|_{\text{FV}(P)}$, to $\llbracket P \rrbracket_v = \llbracket P \rrbracket_u$.

Następujący lemat jest odpowiednikiem lematu 10.2. Dowodzimy go przez indukcję ze względu na budowę termów

Lemat 16.1 *W dowolnym modelu zachodzi tożsamość $\llbracket M[x := N] \rrbracket_v = \llbracket M \rrbracket_{v[x \mapsto \llbracket N \rrbracket_v]}$.*

Piszemy $\mathcal{A}, v \models M = N$, gdy $\llbracket M \rrbracket_v = \llbracket N \rrbracket_v$. Dalej, $\mathcal{A} \models M = N$ oznacza, że $\mathcal{A}, v \models M = N$ dla dowolnego v , natomiast napis $\models M = N$ mówi, że jest tak w każdym modelu. Zaczynamy od łatwego twierdzenia o poprawności.

¹³Por. analogiczną definicję dla modeli bez typów.

Fakt 16.2 *Jeśli $M =_{\beta\eta} N$ to $\models M = N$.*

Dowód: Indukcja ze względu na definicję $=_{\beta\eta}$, z użyciem lematu 16.1. Istotne równości:

- $\llbracket (\lambda x.P)Q \rrbracket_v = \llbracket \lambda x.P \rrbracket_v \cdot \llbracket Q \rrbracket_v = \llbracket P \rrbracket_{v[x \mapsto \llbracket Q \rrbracket_v]} = \llbracket P[x := Q] \rrbracket_v$.
- $\llbracket \lambda x.Px \rrbracket_v \cdot a = \llbracket Px \rrbracket_{v[x \mapsto a]} = \llbracket P \rrbracket_{v[x \mapsto a]} \cdot a = \llbracket P \rrbracket_v \cdot a$, gdy $x \notin \text{FV}(P)$. ■

Nas, tak naprawdę, interesują tylko modele, w których $D_{\sigma \rightarrow \tau}$ to po prostu zbiór wszystkich funkcji z D_σ do D_τ . Jeśli $D_0 = A$, to taki model oznaczamy przez $\mathfrak{M}(A)$. Dla dowodu twierdzenia o pełności potrzebujemy jednak jeszcze jednego przykładu. Przez $\mathfrak{M}_\eta$ oznaczmy model, w którym dziedzina D_σ składa się ze wszystkich termów typu σ , branych z dokładnością do $\beta\eta$ -konwersji (tj. w istocie z klas abstrakcji). W modelu $\mathfrak{M}_\eta$ znaczenie termów jest określone przez podstawienie, tj. dla $\text{FV}(M) = \{x_1, \dots, x_n\}$ mamy

$$\llbracket M \rrbracket_v = M[x_1, \dots, x_n := v(x_1), \dots, v(x_n)].$$

Nietrudno teraz pokazać twierdzenie odwrotne do Faktu 16.2, czyli najłatwiejszą wersję twierdzenia o pełności.

Fakt 16.3 *Jeśli $\mathfrak{M}_\eta \models M = N$ to $M =_{\beta\eta} N$. Zatem $\models M = N$ implikuje $M =_{\beta\eta} N$.*

W dalszym ciągu pokażemy twierdzenie o pełności dla modeli postaci $\mathfrak{M}(A)$. Wymaga to jednak pewnych przygotowań.

Definicja 16.4 *Częściowy epimorfizm z modelu $\mathcal{A} = \langle \{D_\sigma\}_{\sigma \in T}, \{\cdot_{\sigma\tau}\}_{\sigma, \tau \in T}, \llbracket \cdot \rrbracket \rangle$ do modelu $\mathcal{B} = \langle \{E_\sigma\}_{\sigma \in T}, \{\cdot_{\sigma\tau}\}_{\sigma, \tau \in T}, \llbracket \cdot \rrbracket \rangle$, to rodzina częściowych surjekcji $\phi_\sigma : D_\sigma \xrightarrow{\text{na}} E_\sigma$, zachowująca aplikację, tj. spełniająca warunek $\phi_\tau(d \cdot e) = \phi_{\sigma \rightarrow \tau}(d) \cdot \phi_\sigma(e)$. Ścisłej, żądamy aby wartość $\phi_{\sigma \rightarrow \tau}(d)$ była określona i równa h wtedy i tylko wtedy, gdy dla każdego $e \in \text{Dom}(\phi_\sigma)$ określone jest $\phi_\tau(d \cdot e)$ i zachodzi równość¹⁴ $\phi_\tau(d \cdot e) = h \cdot \phi_\sigma(e)$. Poniżej zamiast $\phi_\sigma(d)$ często będziemy pisać $\bar{d}$. Na przykład zamiast $\phi_\tau(d \cdot e) = \phi_{\sigma \rightarrow \tau}(d) \cdot \phi_\sigma(e)$ napiszemy $\bar{d} \cdot \bar{e} = \bar{d} \cdot \bar{e}$.*

Lemat 16.5 *Jeśli $\{\phi_\sigma\}_{\sigma \in T}$ jest częściowym epimorfizmem z $\mathcal{A}$ do $\mathcal{B}$ oraz $\bar{v}(x) = \overline{v(x)}$ dla dowolnego x , to dla każdego M zachodzi $\llbracket M \rrbracket_v^{\mathcal{B}} = \llbracket M \rrbracket_v^{\mathcal{A}}$, w szczególności $\llbracket M \rrbracket_v^{\mathcal{A}}$ jest określone.*

Dowód: Indukcja ze względu na M . Dla zmiennych teza wynika natychmiast z definicji $\bar{v}$. Dla aplikacji mamy $\llbracket PQ \rrbracket_v^{\mathcal{B}} = \llbracket P \rrbracket_v^{\mathcal{B}} \cdot \llbracket Q \rrbracket_v^{\mathcal{B}} = \llbracket P \rrbracket_v^{\mathcal{A}} \cdot \llbracket Q \rrbracket_v^{\mathcal{A}} = \llbracket P \rrbracket_v^{\mathcal{A}} \cdot \llbracket Q \rrbracket_v^{\mathcal{A}} = \llbracket PQ \rrbracket_v^{\mathcal{A}}$. W przypadku abstrakcji pamiętajmy, że każdy element modelu $\mathcal{B}$ jest postaci $\bar{a}$. Mamy teraz

$$\llbracket \lambda x.P \rrbracket_v^{\mathcal{B}} \cdot \bar{a} = \llbracket P \rrbracket_{v[x \mapsto \bar{a}]}^{\mathcal{B}} = \llbracket P \rrbracket_{v[x \mapsto a]}^{\mathcal{B}} = \overline{\llbracket P \rrbracket_{v[x \mapsto a]}^{\mathcal{A}}} = \overline{\llbracket \lambda x.P \rrbracket_v^{\mathcal{A}} \cdot a} = \llbracket \lambda x.P \rrbracket_v^{\mathcal{A}} \cdot \bar{a},$$

dla dowolnego $\bar{a}$, i stosujemy ekstensjonalność. Czytelnikowi pozostawiamy sprawdzenie, że $\llbracket M \rrbracket_v^{\mathcal{A}}$ jest zawsze określone. ■

¹⁴Jedyność h wynika z ekstensjonalności.

Wniosek 16.6 *Jeśli $\mathcal{A} \models M = N$ i istnieje częściowy epimorfizm z $\mathcal{A}$ do $\mathcal{B}$ to $\mathcal{B} \models M = N$.*

Dowód: Bo każde wartościowanie w $\mathcal{B}$ ma postać $\bar{v}$. ■

Lemat 16.7 *Jeśli w modelu $\mathcal{B} = \langle \{E_\sigma\}_{\sigma \in T}, \{\cdot_{\sigma\tau}\}_{\sigma, \tau \in T}, [\![\]\!] \rangle$ dziedzina E_0 jest przeliczalna, to dowolną surjekcję z $\mathbb{N}$ do E_0 można rozszerzyć do częściowego epimorfizmu z $\mathfrak{M}(\mathbb{N})$ do $\mathcal{B}$.*

Dowód: Oznaczmy przez D_σ dziedziny w modelu $\mathfrak{M}(\mathbb{N})$. Definiujemy $\phi_\sigma : D_\sigma \xrightarrow{\text{na}} E_\sigma$ przez indukcję ze względu na σ , przyjmując daną surjekcję z $\mathbb{N}$ na E_0 jako ϕ_0 . Załóżmy, że ϕ_σ i ϕ_τ są określone. Dla dowolnego $h \in E_{\sigma \rightarrow \tau}$, istnieją funkcje częściowe $d : D_\sigma \rightarrow D_\tau$ o własności $\phi_\tau(d(e)) = h \cdot \phi_\sigma(e)$ (inaczej $\overline{d(e)} = h \cdot \bar{e}$) dla dowolnego $e \in D_\sigma$. Dla każdej takiej funkcji d należy przyjąć $\bar{d} = h$. ■

Twierdzenie 16.8 (H. Friedman, 1975) *Warunki $M =_{\beta\eta} N$ i $\mathfrak{M}(\mathbb{N}) \models M = N$ są równoważne.*

Dowód: Na mocy lematu 16.7 mamy częściowy epimorfizm z $\mathfrak{M}(\mathbb{N})$ do modelu $\mathfrak{M}_\eta$. Jeśli więc $\mathfrak{M}(\mathbb{N}) \models M = N$ to $\mathfrak{M}_\eta \models M = N$ (lemat 16.5) a stąd $M =_{\beta\eta} N$. ■

Twierdzenie Statmana

Pokażemy teraz twierdzenie o pełności dla modeli postaci $\mathfrak{M}(A)$, gdzie A jest skończone. Przypomnijmy, że typ bazowy 0 jest rzędu 0, a typy postaci $0 \rightarrow \dots \rightarrow 0 \rightarrow 0$ są rzędu 1. O zmiennej typu τ powiemy, że jest rzędu 0 lub rzędu 1, gdy takiego rzędu jest typ τ .

Lemat 16.9 *Założmy, że term $Q : 0$ jest w postaci normalnej i że wszystkie jego zmienne wolne są rzędu co najwyżej 1. Wtedy istnieje taka stała k , że dla dowolnego $N : 0$ zachodzi:*

$$\text{Jeśli } \mathfrak{M}(k) \models Q = N \text{ to } Q =_{\beta\eta} N.$$

Dowód: Zauważmy najpierw, że term Q nie może zawierać żadnych abstrakcji (jest zbudowany jak zwykły term „algebraiczny” w którym zmienne rzędu 1 odgrywają rolę symboli funkcyjnych).

Ponumerujemy wszystkie podtermy termu Q , które są typu 0, ustawiając je w ciąg skończony $t_1, t_2, t_3, \dots, t_{k-1}$. Możemy teraz określić wartościowanie v w modelu $\mathfrak{M}(\mathbb{N})$ w ten sposób, że $\llbracket t_i \rrbracket_v = i$ dla każdego $i = 1, \dots, k-1$, oraz $\llbracket N \rrbracket_v = 0$ dla każdego innego termu $N : 0$ w postaci normalnej. Wtedy wartość $\llbracket Q \rrbracket_v$ jest jedną z liczb $1, \dots, k-1$, powiedzmy $k-1$. Przyporządkowanie

$$\phi_0(i) = \begin{cases} i, & \text{jeśli } i = 1, \dots, k-1; \\ 0, & \text{w przeciwnym przypadku,} \end{cases}$$

jest surjekcją z $\mathbb{N}$ do $k = \{0, \dots, k-1\}$, a zatem rozszerza się do częściowego epimorfizmu z modelu $\mathcal{A} = \mathfrak{M}(\mathbb{N})$ do modelu $\mathcal{B} = \mathfrak{M}(k)$. Mamy teraz $\overline{[Q]_v^{\mathcal{A}}} = \overline{k-1} = k-1$, i jeśli $N \neq_{\beta\eta} Q$, to z lematu 16.5 wynika $[N]_v^{\mathcal{B}} = \overline{[N]_v^{\mathcal{A}}} \neq k-1 = \overline{[Q]_v^{\mathcal{A}}} = [Q]_v^{\mathcal{B}}$, czyli $\mathfrak{M}(k), v \not\models Q = N$. ■

Uogólnienie lematu 16.9 na typy wyższych rzędów wymaga dwóch lematów o charakterze syntaktycznym. Mówimy, że term $M : \sigma_1 \rightarrow \dots \rightarrow \sigma_n \rightarrow 0$ jest w postaci Statmana, gdy

$$M = \lambda x_1^{\sigma_1} \dots x_n^{\sigma_n} . x(M_1 x_1 \dots x_n) \dots (M_m x_1 \dots x_n),$$

gdzie $M_1, \dots, M_m$ są w postaci Statmana, oraz $x_1, \dots, x_n \notin \text{FV}(M_i)$. Łatwo widzieć, że każdy term jest beta-eta-równy termowi w postaci Statmana.

Lemat 16.10 *Dla wszystkich typów σ istnieją termy U^σ typu σ , których zmienne wolne są rzędu co najwyżej jeden, i które mają taką własność: Jeśli dwa termy zamknięte M, N typu $\sigma_1 \rightarrow \dots \rightarrow \sigma_n \rightarrow 0$ w postaci Statmana są różnej długości, to $MU^{\sigma_1} \dots U^{\sigma_n} \neq_{\beta\eta} NU^{\sigma_1} \dots U^{\sigma_n}$.*

Dowód: Termy U^σ definiujemy przez indukcję, wraz z termami $V^\sigma : \sigma \rightarrow 0$. Najpierw bierzemy $U^0 = z_0$ (gdzie z_0 jest nową zmienną) oraz $V^0 = \lambda x^0 z_0$. Dla typu σ postaci $\sigma = \sigma_1 \rightarrow \dots \rightarrow \sigma_n \rightarrow 0$ przyjmujemy

$$\begin{aligned} U^\sigma &= \lambda x_1^{\sigma_1} \dots x_n^{\sigma_n} . z_\sigma(V^{\sigma_1} x_1) \dots (V^{\sigma_n} x_n); \\ V^\sigma &= \lambda x^\sigma . x U^{\sigma_1} \dots U^{\sigma_n}, \end{aligned}$$

gdzie zmienna z_σ jest znowu świeża (inna dla różnych σ).

Przypuśćmy teraz, że $M, N : \sigma_1 \rightarrow \dots \rightarrow \sigma_n \rightarrow 0$ oraz $MU^{\sigma_1} \dots U^{\sigma_n} =_{\beta\eta} NU^{\sigma_1} \dots U^{\sigma_n}$. Jeśli $M = \lambda x_1^{\sigma_1} \dots x_n^{\sigma_n} . x_i(M_1 x_1 \dots x_n) \dots (M_m x_1 \dots x_n)$, to lewa strona równości jest $\beta\eta$ -równa termowi $U^{\sigma_i}(M_1 U^{\sigma_1} \dots U^{\sigma_n}) \dots (M_m U^{\sigma_1} \dots U^{\sigma_n})$. Przyjmując $\sigma_i = \tau_1 \rightarrow \dots \rightarrow \tau_m \rightarrow 0$, mamy dalej $MU^{\sigma_1} \dots U^{\sigma_n} =_{\beta\eta} z_{\sigma_i}(V^{\tau_1}(M_1 U^{\sigma_1} \dots U^{\sigma_n})) \dots (V^{\tau_m}(M_m U^{\sigma_1} \dots U^{\sigma_n}))$, co z kolei jest $\beta\eta$ -równe wyrażeniu $z_{\sigma_i}(M_1 U^{\sigma_1} \dots U^{\sigma_n} \vec{U}_1) \dots (M_m U^{\sigma_1} \dots U^{\sigma_n} \vec{U}_m)$, w którym wektory $\vec{U}_1, \dots, \vec{U}_m$ są takie jak w termach $V^{\tau_1}, \dots, V^{\tau_m}$.

Podobnie, jeśli $N = \lambda x_1^{\sigma_1} \dots x_n^{\sigma_n} . x_j(N_1 x_1 \dots x_n) \dots (N_r x_1 \dots x_n)$, to $NU^{\sigma_1} \dots U^{\sigma_n}$ zredukuje się do termu $z_{\sigma_j}(N_1 U^{\sigma_1} \dots U^{\sigma_n} \vec{U}^1) \dots (N_r U^{\sigma_1} \dots U^{\sigma_n} \vec{U}^r)$, gdzie $\sigma_j = \rho_1 \rightarrow \dots \rightarrow \rho_r \rightarrow 0$.

Skoro te termy są $\beta\eta$ -równe, to przede wszystkim musi się zgadzać zmienna czołowa $z_{\sigma_j} = z_{\sigma_i}$. A więc $\sigma_j = \sigma_i$ skąd $m = r$. Dalej $M_l U^{\sigma_1} \dots U^{\sigma_n} \vec{U}_l =_{\beta\eta} N_l U^{\sigma_1} \dots U^{\sigma_n} \vec{U}_l$, dla $l \leq m$ i możemy zastosować indukcję, bo termy M_i są krótsze niż M . (Uwaga: teraz już wiemy, że wektor $\vec{U}_l$ jest po obu stronach taki sam.) ■

Lemat 16.11 *Jeśli M, N są zamkniętymi termami typu $\sigma = \sigma_1 \rightarrow \dots \rightarrow \sigma_n \rightarrow 0$, oraz $M \neq_{\beta\eta} N$, to istnieją termy $V_1^{\sigma_1}, \dots, V_n^{\sigma_n}$, których zmienne wolne są rzędu co najwyżej jeden, takie że $MV_1 \dots V_n \neq_{\beta\eta} NV_1 \dots V_n$.*

Dowód: Indukcja ze względu na M . Można założyć, że M i N są w postaci Statmana:

$$M = \lambda x_1^{\sigma_1} \dots x_n^{\sigma_n} . x_i(M_1 x_1 \dots x_n) \dots (M_m x_1 \dots x_n);$$

$$N = \lambda x_1^{\sigma_1} \dots x_n^{\sigma_n} . x_j (N_1 x_1 \dots x_n) \dots (N_r x_1 \dots x_n).$$

Jeśli $i \neq j$, to sprawa jest prosta: wystarczy wybrać $V_i = \lambda y_1 \dots y_m . z_1$ i $V_j = \lambda y_1 \dots y_r . z_2$, gdzie z_1 i z_2 to dwie różne zmienne. Niech więc $i = j$ (oraz $m = r$). Skoro $M \neq_{\beta\eta} N$, to $M_\ell \neq_{\beta\eta} N_\ell$ dla pewnego $\ell \leq m$. Niech $\sigma_i = \tau_1 \rightarrow \dots \rightarrow \tau_m \rightarrow 0$, oraz $\tau_\ell = \rho_1 \rightarrow \dots \rightarrow \rho_d \rightarrow 0$. Z założenia indukcyjnego są termy $W_1^{\sigma_1}, \dots, W_n^{\sigma_n}, W_{n+1}^{\rho_1}, \dots, W_{n+d}^{\rho_d}$, o zmiennych wolnych rzędu co najwyżej 1, takie że $M_\ell W_1 \dots W_{n+d} \neq_{\beta\eta} N_\ell W_1 \dots W_{n+d}$.

Dla $k \neq i$ przyjmujemy $V_k = W_k$. Aby określić V_i , przyjmijmy, że $W_i = \lambda y_1 \dots y_m . W$. Wtedy

$$V_i = \lambda y_1 \dots y_m . z W(y_\ell W_{n+1} \dots W_{n+d}),$$

gdzie z jest nową zmienną typu $0 \rightarrow 0 \rightarrow 0$. Wówczas

$$M V_1 \dots V_n =_{\beta\eta} V_i (M_1 V_1 \dots V_n) \dots (M_m V_1 \dots V_n) =_{\beta\eta} z W' (M_\ell V_1 \dots V_n W_{n+1} \dots W_{n+d}),$$

gdzie $W' = W[y_1, \dots, y_m := M_1 V_1 \dots V_n, \dots, M_m V_1 \dots V_n]$. Podobnie zredukuje się term $N V_1 \dots V_n$, jeśli więc $M V_1 \dots V_n =_{\beta\eta} N V_1 \dots V_n$, to w szczególności

$$M_\ell V_1 \dots V_n W_{n+1} \dots W_{n+d} =_{\beta\eta} N_\ell V_1 \dots V_n W_{n+1} \dots W_{n+d}.$$

Wektor $V_1 \dots V_n$ różni się od wektora $W_1 \dots W_n$ tylko na i -tej współrzędnej, a wolna zmienna z występuje tylko w termie V_i . Podstawiając w termie V_i na miejsce zmiennej z kombinatory $\mathbf{K}$ otrzymamy więc fałszywą równość $M_\ell W_1 \dots W_{n+d} =_{\beta\eta} N_\ell W_1 \dots W_{n+d}$. ■

Twierdzenie 16.12 (R. Statman, 1982) *Dla dowolnego termu M istnieje taka liczba k (efektywnie obliczalna z M), że jeśli N jest dowolnym termem, to:*

$$M =_{\beta\eta} N \quad \text{wtedy i tylko wtedy gdy} \quad \mathfrak{M}(k) \models M = N.$$

Dowód: Załóżmy, że M jest typu σ . Z lematów 16.10 i 16.11 wynika, że istnieją takie termy $L_1, \dots, L_m : \sigma \rightarrow 0$, że dla dowolnego N mamy

$$\text{Jeśli } M \neq_{\beta\eta} N \text{ to } L_i M \neq_{\beta\eta} L_i N, \text{ dla pewnego } i = 1, \dots, m.$$

Ponadto typy zmiennych wolnych w $L_1, \dots, L_m$ są rzędu co najwyżej 1. Niech Q będzie postacią normalną termu $z(L_1 M) \dots (L_m M)$, gdzie z jest nową zmienną. Do termu Q można zastosować lemat 16.9, weźmy więc odpowiednie k . Jeśli $\mathfrak{M}(k) \models M = N$, to wtedy także $\mathfrak{M}(k) \models Q = z(L_1 N) \dots (L_m N)$, skąd $Q =_{\beta\eta} z(L_1 N) \dots (L_m N)$ i dalej $L_i M =_{\beta\eta} L_i N$ dla wszystkich i . W konsekwencji $M =_{\beta\eta} N$. ■

Wniosek 16.13 (S. Sołowiow, 1981) *Termy M i N są beta-eta-równe wtedy i tylko wtedy, gdy są równe we wszystkich modelach skończonych.*

Istotne w twierdzeniu Statmana jest to, że liczba k nie zależy od M . Przykładem zastosowania tw. Statmana jest niemożność niejednostajnego reprezentowania równości liczb naturalnych w rachunku z typami prostymi. (Podobnie można pokazać, że nie można niejednostajnie reprezentować odejmowania.) Nie jest mi znany syntaktyczny dowód tego faktu.

Wniosek 16.14 *Nie istnieje term $E : \omega_\tau \rightarrow \omega_\sigma \rightarrow \omega_\rho$, taki, że dla dowolnych $p, q \in \mathbb{N}$:*

$$E p^{\omega_\tau} q^{\omega_\sigma} =_{\beta\eta} 0^{\omega_\rho} \quad \text{wtedy i tylko wtedy, gdy} \quad p = q.$$

Dowód: Dobieramy k do $N = \mathbf{0}^{\omega_\rho}$ z twierdzenia Statmana. Dziedzina D_τ w modelu $\mathfrak{M}(k)$ jest skończona, więc istnieją takie liczby $p \neq q$, że $\llbracket \mathbf{p}^{\omega_\tau} \rrbracket = \llbracket \mathbf{q}^{\omega_\tau} \rrbracket$. Wtedy $\llbracket E\mathbf{p}^{\omega_\tau}\mathbf{q}^{\omega_\sigma} \rrbracket = \llbracket E \rrbracket \llbracket \mathbf{p}^{\omega_\tau} \rrbracket \llbracket \mathbf{q}^{\omega_\sigma} \rrbracket = \llbracket E \rrbracket \llbracket \mathbf{q}^{\omega_\tau} \rrbracket \llbracket \mathbf{q}^{\omega_\sigma} \rrbracket = \llbracket E\mathbf{q}^{\omega_\tau}\mathbf{q}^{\omega_\sigma} \rrbracket = \llbracket \mathbf{0}^{\omega_\rho} \rrbracket$, czyli $\mathfrak{M}(k) \models E\mathbf{p}^{\omega_\tau}\mathbf{q}^{\omega_\sigma} = \mathbf{0}^{\omega_\rho}$. Zatem $p = q$ wbrew założeniu. ■

Przez pewien czas otwartym problemem była tzw. hipoteza Plotkina o rozstrzygalności definiowalności w skończonych modelach. Hipoteza okazała się nieprawdziwa.

Twierdzenie 16.15 (R. Loader, 1993) *Następujący problem jest nierozstrzygalny:*

*Dany jest skończony zbiór X , typ τ i element $d \in D_\tau(X)$.
Czy istnieje taki kombinatory M typu τ , że $\llbracket M \rrbracket = d$?*

Pracę Loadera można znaleźć pod adresem

<http://homepages.ihug.co.nz/~suckfish/papers/Church.pdf>

Ćwiczenia

1. Pokazać, że $\mathfrak{M}_\eta$ jest modelem, i że $\mathfrak{M}_\eta \models M = N$ zachodzi wtedy i tylko wtedy gdy $M =_{\beta_\eta} N$.
2. Rodzinę relacji $\{\sim_\sigma\}_{\sigma \in T}$, odpowiednio w $\{D_\sigma\}_{\sigma \in T}$, nazywamy *relacją logiczną*, jeżeli dla dowolnych typów σ i τ i dowolnych $d, d' \in D_{\sigma \rightarrow \tau}$ zachodzi równoważność (zamiast $\sim_\sigma$ piszemy $\sim$):
$$d \sim d' \quad \text{wtedy i tylko wtedy, gdy} \quad \forall e, e' \in D_\sigma (e \sim e' \rightarrow d \cdot e \sim d' \cdot e').$$
Udowodnić, że jeśli $v(x) \sim w(x)$ dla dowolnego $x \in \text{FV}(M)$, to $\llbracket M \rrbracket_v \sim \llbracket M \rrbracket_w$.
3. Jeśli $\{\phi_\sigma\}_{\sigma \in T}$ jest częściowym epimorfizmem, oraz
$$d \sim d' \quad \text{wtedy i tylko wtedy, gdy} \quad \bar{d} = \bar{d}'$$
to $\{\sim_\sigma\}_{\sigma \in T}$ jest relacją logiczną.¹⁵
4. Zdefiniować wartościowanie v , o którym mowa w dowodzie lematu 16.9.
5. Uzasadnić istnienie termów L_i , o których mowa w dowodzie twierdzenia 16.12.
6. Zdefiniować pojęcie modelu dla rachunku z wieloma atomami typowymi. Czy twierdzenie Statmana pozostaje prawdziwe?

17 Semantyka w stylu Curry'ego

Zajmiemy się teraz semantyczną interpretacją asercji typowych w stylu Curry'ego. Musimy się w tym celu odwołać do semantyki beztypowej, bo mamy przecież do czynienia z termami „czystego” rachunku lambda. Przypuśćmy więc, że mamy lambda-model $\mathcal{D} = \langle D, \cdot, \llbracket \rrbracket \rangle$. Podzbiory tego modelu reprezentują własności jego elementów, a także własności termów. Mówimy więc, że element $a \in D$ ma własność $\sigma \subseteq D$, gdy $a \in \sigma$. Podobnie, powiemy że term zamknięty M ma własność $\sigma \subseteq D$, gdy $\llbracket M \rrbracket \in \sigma$. W przypadku termów ze zmiennymi wolnymi będzie to oczywiście zależało od wartościowania. Naturalna jest następująca definicja:

$$\sigma \Rightarrow \tau := \{a \in D \mid \forall b \in D (b \in \sigma \rightarrow a \cdot b \in \tau)\}.$$

¹⁵Uwaga: Relacja, o której tu mowa, jest *częściową relacją równoważności* tj. jest symetryczna i przechodnia, ale nie musi być zwrotna. W ogólności, relacja logiczna nie musi też być symetryczna ani przechodnia.

Własność $\sigma \Rightarrow \tau$ to poprawność ze względu na prewarunek σ i postwarunek τ . Oczywista jest tu analogia z typem funkcyjnym. Nadamy jej ścisły sens, interpretując typy jako podzbiory modelu. Funkcję $\xi : TV \rightarrow \mathbf{P}(D)$ nazwiemy *wartościowaniem typowym*, a znaczenie $\llbracket \tau \rrbracket_\xi$ typu τ przy wartościowaniu ξ zdefiniujemy tak:

- $\llbracket s \rrbracket_\xi = \xi(s)$, gdy s jest zmienną typową;
- $\llbracket \sigma \rightarrow \tau \rrbracket_\xi = \llbracket \sigma \rrbracket_\xi \Rightarrow \llbracket \tau \rrbracket_\xi$.

Piszemy $\mathcal{D}, v, \xi \models M : \sigma$, gdy $\llbracket M \rrbracket_v \in \llbracket \sigma \rrbracket_\xi$. Powiemy wtedy, że term M ma typ σ w modelu $\mathcal{D}$ przy wartościowaniu v i wartościowaniu typowym ξ .

Podobnie, napis $\mathcal{D}, v, \xi \models \Gamma$ oznacza, że dla wszystkich $(x : \tau) \in \Gamma$ zachodzi $v(x) \in \llbracket \tau \rrbracket_\xi$.

Semantycznym odpowiednikiem asercji $\Gamma \vdash M : \sigma$ jest więc następująca własność:

Dla dowolnego modelu $\mathcal{D}$ i dowolnych v, ξ , warunek $\mathcal{D}, v, \xi \models \Gamma$ implikuje $\mathcal{D}, v, \xi \models M : \sigma$, zapisywana $\Gamma \models M : \sigma$. Zachodzi teraz taki fakt:

Twierdzenie 17.1 (o poprawności) *Jeśli $\Gamma \vdash M : \sigma$, to $\Gamma \models M : \sigma$.*

Dowód: Indukcja ze względu na wyprowadzenie $\Gamma \vdash M : \sigma$. Rozważamy kilka przypadków, w zależności od tego jakiej reguły użyto w tym wyprowadzeniu jako ostatniej. Rozpatrzmy przypadek wyprowadzenia kończącego się zastosowaniem reguły (Abs). Wtedy σ jest postaci $\tau \rightarrow \rho$, a term M jest abstrakcją $\lambda x.N$. Natomiast konkluzję $\Gamma \vdash \lambda x.N : \tau \rightarrow \rho$ otrzymano z przesłanki $\Gamma, x : \tau \vdash N : \rho$.

Przypuśćmy, że $\mathcal{D}, v, \xi \models \Gamma$. Mamy pokazać, że $\mathcal{D}, v, \xi \models \lambda x.N : \tau \rightarrow \rho$, innymi słowy, że $\llbracket \lambda x.N \rrbracket_v \in \llbracket \tau \rightarrow \rho \rrbracket_\xi = \llbracket \tau \rrbracket_\xi \Rightarrow \llbracket \rho \rrbracket_\xi$. Niech więc $a \in \llbracket \tau \rrbracket_\xi$. Wtedy $\mathcal{D}, v[x \mapsto a], \xi \models \Gamma, x : \tau$, więc $\llbracket \lambda x.N \rrbracket_v \cdot a = \llbracket N \rrbracket_{v[x \mapsto a]} \in \llbracket \rho \rrbracket_\xi$, bo z założenia indukcyjnego $\Gamma, x : \tau \models N : \rho$. Pozostałe przypadki są natychmiastowe. ■

Twierdzenia odwrotnego (o pełności) nie udowodnimy z oczywistych powodów: równość $=_\beta$, a więc też równość w modelu, nie zachowuje typów. Na przykład $\models \lambda x. \mathbf{K}x(\lambda y yy) : p \rightarrow p$ oraz $\models \lambda x. (\lambda y yy) \mathbf{I}x : p \rightarrow p$ chociaż $\not\models \mathbf{K}x(\lambda y yy) : p \rightarrow p$ oraz $\not\models \lambda x. (\lambda y yy) \mathbf{I}x : p \rightarrow p$. Twierdzenie o pełności zachodzi jednak dla bardziej „elastycznego” systemu przypisania typów.

Ćwiczenia

1. Oznaczmy przez ω całą dziedzinę D . Udowodnić, że dla dowolnych podzbiorów modelu D zachodzą związki:

$$\sigma \subseteq \omega, \quad \omega \subseteq \omega \Rightarrow \omega, \quad \sigma \cap \tau \subseteq \sigma, \quad \sigma \cap \tau \subseteq \tau, \quad \sigma \subseteq \sigma \cap \sigma$$

$$(\sigma \Rightarrow \tau) \cap (\sigma \Rightarrow \rho) \subseteq \sigma \Rightarrow \tau \cap \rho$$

Jeśli $\sigma \subseteq \sigma'$ i $\tau \subseteq \tau'$, to $\sigma \cap \tau \subseteq \sigma' \cap \tau'$ oraz $\sigma' \Rightarrow \tau \subseteq \sigma \Rightarrow \tau'$.

2. Jeśli model $\mathcal{D}$ jest częściowo uporządkowany, to naturalne jest żądanie, aby własności były zbiorami zamkniętymi w górę (jeśli $a \in \sigma$ i $a \leq b$, to $b \in \sigma$) lub były postaci $\uparrow a = \{b \mid a \leq b\}$. Pokazać, że wtedy $\sigma \Rightarrow \tau$ jest zamknięte w górę, oraz że $(\uparrow a \Rightarrow \uparrow b) = \uparrow f$ dla pewnego f .
3. Jeśli każde z σ_i i τ_i jest postaci $\uparrow a$ to z warunku $\bigcap \{\sigma_i \Rightarrow \tau_i \mid i \in I\} \subseteq \sigma \Rightarrow \tau$ wynika $\bigcap \{\tau_i \mid \sigma \subseteq \sigma_i\} \subseteq \tau$.

18 Typy iloczynowe

Zdefiniujemy teraz rachunek lambda z typami *iloczynowymi*,¹⁶ które w odróżnieniu od typów prostych, zbudowane są z pomocą dwóch konstruktorów $\rightarrow$ i $\cap$. Zaczynamy od składni typów.

- Stała ω jest typem;
- Typy atomowe są typami;
- Jeśli σ i τ są typami, to $\sigma \rightarrow \tau$ i $\sigma \cap \tau$ są typami.

Zbiór wszystkich typów oznaczmy przez $\mathcal{T}_{\cap\omega}$. Przyjmujemy konwencję, że iloczyn ma wyższy priorytet niż strzałka. Iloczyn, jak się zaraz okaże, można w gruncie rzeczy uważać za łączny i pomijać nawiasy.

W zbiorze $\mathcal{T}_{\cap\omega}$ określamy relację $\leq$ jako najmniejszy quasiporządek¹⁷ spełniający takie warunki:

$$\sigma \leq \omega, \quad \omega \leq \omega \rightarrow \omega, \quad \sigma \cap \tau \leq \sigma, \quad \sigma \cap \tau \leq \tau, \quad \sigma \leq \sigma \cap \sigma$$

$$(\sigma \rightarrow \tau) \cap (\sigma \rightarrow \rho) \leq \sigma \rightarrow \tau \cap \rho$$

$$\text{Jeśli } \sigma \leq \sigma' \text{ i } \tau \leq \tau', \text{ to } \sigma \cap \tau \leq \sigma' \cap \tau' \text{ oraz } \sigma' \rightarrow \tau \leq \sigma \rightarrow \tau'.$$

Definicja powyżej motywowana jest własnościami operacji $\Rightarrow$ i $\cap$ (ćwiczenie 1 do rozdziału 17).

Określimy teraz reguły przypisania (wyprowadzania) typów iloczynowych dla termów rachunku lambda. Reguły te tworzą system wnioskowania o typach, który nazwiemy *systemem BCD* od nazwisk Barendregt, Coppo i Dezani. System BCD ma następujące aksjomaty i reguły.

$$(\text{Var}) \quad \Gamma, x:\sigma \vdash x:\sigma \qquad (\omega) \quad \Gamma \vdash M:\omega$$

$$(\text{Abs}) \quad \frac{\Gamma, x:\sigma \vdash M:\tau}{\Gamma \vdash (\lambda x.M):\sigma \rightarrow \tau} \qquad (\text{App}) \quad \frac{\Gamma \vdash M:\sigma \rightarrow \tau \quad \Gamma \vdash N:\sigma}{\Gamma \vdash (MN):\tau}$$

$$(\cap I) \quad \frac{\Gamma \vdash M:\sigma \quad \Gamma \vdash M:\tau}{\Gamma \vdash M:\sigma \cap \tau} \qquad (\leq) \quad \frac{\Gamma \vdash M:\sigma}{\Gamma \vdash M:\tau} \quad (\sigma \leq \tau)$$

Reguła $(\cap I)$ jest zwana regułą *wprowadzania iloczynu*. Reguła (Abs) jest często nazywana regułą *wprowadzania strzałki*, a reguła (App) regułą *eliminacji strzałki*. *Eliminacja iloczynu* to taki szczególny przypadek reguły $(\leq)$:

$$(\cap E) \quad \frac{\Gamma \vdash M:\sigma_1 \cap \sigma_2}{\Gamma \vdash M:\sigma_i}$$

¹⁶Dawniej używane określenie „typy koniunkcyjne” wyszło z użycia, bo jest mylące. Obecnie uważamy je za niepoprawne. Odpowiednikiem logicznej koniunkcji nie jest iloczyn typów, ale produkt kartezjański.

¹⁷Quasiporządek to relacja zwrotna i przechodnia.

Jeśli asercja $\Gamma \vdash M : \tau$ ma wyprowadzenie w systemie BCD, to piszemy $\Gamma \vdash_{BCD} M : \tau$ lub po prostu $\Gamma \vdash M : \tau$. Jeśli $\Gamma = \emptyset$, to zamiast $\Gamma \vdash M : \tau$ piszemy $\vdash M : \tau$, lub wręcz $M : \tau$.

Przykład: Następujące asercje są wyprowadzalne w systemie BCD:¹⁸

- $\vdash \lambda x.xx : t \cap (t \rightarrow s) \rightarrow s$;
- $\vdash \mathbf{2} : (t \rightarrow s) \cap (s \rightarrow r) \rightarrow (t \rightarrow r)$;
- $\vdash \mathbf{K} : t \rightarrow s \rightarrow t$;
- $\vdash \mathbf{S} : (t' \rightarrow s \rightarrow r) \rightarrow (t'' \rightarrow s) \rightarrow (t' \cap t'') \rightarrow r$.

Zauważmy, że typowanie termu zależy tylko od jego zmiennych wolnych.

Lemat 18.1 *Niech $\Gamma \vdash M : \tau$ i niech $\Gamma|_{\text{FV}(M)} = \{(x : \Gamma(x)) \mid x \in \text{FV}(M) \text{ i } \Gamma(x) \text{ określone}\}$. Wtedy $\Gamma|_{\text{FV}(M)} \vdash M : \tau$.*

Dowód: Dowód jest przez łatwą indukcję ze względu na wyprowadzenie $\Gamma \vdash M : \tau$. ■

Poprawność redukcji

Naszym celem jest teraz pokazanie, że typy termów zachowują się przy redukcjach. Niestety musimy zacząć od nieprzyjemnych lematów. Notacja $\bigcap S$, gdzie $S = \{\rho_1, \dots, \rho_n\}$ oznacza iloczyn $\rho_1 \cap \dots \cap \rho_n$.

Lemat 18.2 *Jeżeli $\omega \not\leq \rho$ oraz $\bigcap \{\sigma_i \rightarrow \tau_i \mid i \in I\} \leq \rho$ to $\rho = \bigcap \{\xi_j \rightarrow \zeta_j \mid j \in J\}$ dla pewnego J i pewnych ξ_j, ζ_j .*

Dowód: Indukcja ze względu na definicję nierówności. ■

Lemat 18.3 *Jeżeli $\bigcap \{\sigma_i \rightarrow \tau_i \mid i \in I\} \leq \sigma \rightarrow \tau$, oraz $\sigma \rightarrow \tau \not\leq \omega$, to zbiór $\{\tau_i \mid \sigma \leq \sigma_i\}$ jest niepusty, oraz $\bigcap \{\tau_i \mid \sigma \leq \sigma_i\} \leq \tau$.*

Dowód: Przez indukcję ze względu na definicję nierówności pokazujemy ogólniejszy fakt: jeżeli zachodzi nierówność $\bigcap \{\sigma_i \rightarrow \tau_i \mid i \in I\} \cap \bigcap \{p_h \mid h \in H\} \leq (\sigma \rightarrow \tau)$, gdzie p_h są atomami, to zachodzi także nierówność $\bigcap \{\tau_i \mid \sigma \leq \sigma_i\} \leq \tau$. Szczegóły zostawiamy jako ćwiczenie. ■

Lemat 18.4 *Jeśli $\Gamma \vdash M : \sigma$ i $\tau \leq \Gamma(x)$, to $\Gamma(x \mapsto \tau) \vdash M : \sigma$.*

¹⁸Także wtedy, gdy zamiast reguły ($\leq$) używamy tylko słabszej reguły ($\cap E$).

Dowód: Łatwa indukcja ze względu na wyprowadzenie $\Gamma \vdash M : \sigma$. ■

Lemat 18.5 (o generowaniu)

1. Jeśli $\Gamma \vdash MN : \sigma$, to $\Gamma \vdash M : \tau \rightarrow \sigma$ i $\Gamma \vdash N : \tau$ dla pewnego τ .
2. Jeśli $\Gamma \vdash x : \sigma$, gdzie x jest zmienną, to $\Gamma(x) \leq \sigma$.
3. Jeśli $\Gamma \vdash \lambda x.M : \sigma$ i $\sigma \neq \omega$ to $\sigma \equiv \bigcap \{\sigma_i \rightarrow \tau_i \mid i \in I\}$ dla pewnego I i pewnych σ_i, τ_i , przy czym $\tau_i \neq \omega$.
4. Jeśli $\Gamma \vdash \lambda x.M : \sigma \rightarrow \tau$, oraz $x \notin \text{Dom}(\Gamma)$, to $\Gamma, x : \sigma \vdash M : \tau$.

Dowód: (1) Dowód jest przez indukcję ze względu na rozmiar wyprowadzenia $\Gamma \vdash MN : \sigma$ (liczbę wierzchołków w dowodzie). Mamy kilka przypadków w zależności od tego, która reguła została użyta w wyprowadzeniu jako ostatnia. Jeśli jest to reguła (App), to teza jest oczywista. Przypuśćmy, że jest to reguła ($\leq$). To znaczy, że konkluzję $\Gamma \vdash MN : \sigma$ otrzymaliśmy z wcześniej wyprowadzonej asercji $\Gamma \vdash MN : \sigma'$. Dowód tej ostatniej jest mniejszy, więc możemy zastosować założenie indukcyjne. Otrzymujemy $\Gamma \vdash M : \tau \rightarrow \sigma'$ i $\Gamma \vdash N : \tau$. Ponieważ $\tau \rightarrow \sigma' \leq \tau \rightarrow \sigma$, więc $\Gamma \vdash M : \tau \rightarrow \sigma$ tak jak chcemy. Jeśli ostatnią zastosowaną regułą była reguła ($\cap I$), to mamy $\sigma = \sigma_1 \cap \sigma_2$ oraz $\Gamma \vdash MN : \sigma_1$ i $\Gamma \vdash MN : \sigma_2$, przy czym rozmiary tych wyprowadzeń są mniejsze. Z założenia indukcyjnego $\Gamma \vdash M : \tau_1 \rightarrow \sigma_1$, i $\Gamma \vdash M : \tau_2 \rightarrow \sigma_2$ oraz $\Gamma \vdash N : \tau_1$ i $\Gamma \vdash N : \tau_2$. Stąd wnioskujemy, że $\Gamma \vdash N : \tau_1 \cap \tau_2$ oraz $\Gamma \vdash M : (\tau_1 \rightarrow \sigma_1) \cap (\tau_2 \rightarrow \sigma_2)$. Ponieważ

$$(\tau_1 \rightarrow \sigma_1) \cap (\tau_2 \rightarrow \sigma_2) \leq (\tau_1 \cap \tau_2 \rightarrow \sigma_1) \cap (\tau_1 \cap \tau_2 \rightarrow \sigma_2) \leq (\tau_1 \cap \tau_2 \rightarrow \sigma_1 \cap \sigma_2),$$

wiec ostatecznie $\Gamma \vdash M : \tau_1 \cap \tau_2 \rightarrow \sigma_1 \cap \sigma_2$ i też jest dobrze.

Ostatnią regułą nie może być ani reguła (Abs) ani (Var). Pozostaje więc tylko ta możliwość, że całe wyprowadzenie polega na przywołaniu aksjomatu (ω). Ale wtedy $\Gamma \vdash N : \omega$ i $\Gamma \vdash M : \omega$. Ponieważ $\omega \leq \omega \rightarrow \omega$, więc mamy też $\Gamma \vdash M : \omega \rightarrow \omega$.

(2) Dowód jest podobny, a nawet prostszy.

(3) Postępujemy przez podobną indukcję, korzystając z lematu 18.2. Zauważmy tu tylko, że warunek $\tau_i \neq \omega$ wynika stąd, że $\sigma \rightarrow \omega \equiv \omega$ dla dowolnego σ . „Niedobre” człony iloczynu można więc pomijać.

(4) Pokażemy nieco silniejszą tezę: Jeśli $\Gamma \vdash \lambda x.M : (\sigma \rightarrow \tau) \cap \zeta$ to $\Gamma, x : \sigma \vdash M : \tau$. Dowód jest znowu przez indukcję ze względu na wyprowadzenie. Nieoczywisty przypadek jest tylko jeden: gdy konkluzję $\Gamma \vdash \lambda x.M : (\sigma \rightarrow \tau) \cap \zeta$ otrzymano z $\Gamma \vdash \lambda x.M : \rho$ za pomocą reguły ($\leq$). Z części 3 wynika, że typ ρ jest postaci $\bigcap \{\sigma_i \rightarrow \tau_i \mid i \in I\}$. Możemy zastosować założenie indukcyjne do każdego członu tego iloczynu, otrzymując $\Gamma, x : \sigma_i \vdash M : \tau_i$ dla wszystkich $i \in I$. Na mocy lematu 18.4 zachodzi też $\Gamma, x : \sigma \vdash M : \tau_i$ dla wszystkich $i \in I$, dla których $\sigma \leq \sigma_i$. Stąd dalej wnioskujemy, że $\Gamma, x : \sigma \vdash M : \bigcap \{\tau_i \mid \sigma \leq \sigma_i\}$ i do pełni szczęścia brakuje nam tylko nierówności $\bigcap \{\tau_i \mid \sigma \leq \sigma_i\} \leq \tau$. Ponieważ jednak $\rho \leq (\sigma \rightarrow \tau) \cap \zeta \leq (\sigma \rightarrow \tau)$, więc nierówność ta wynika z lematu 18.3. ■

Lemat 18.6 Jeśli $\Gamma, x : \sigma \vdash M : \tau$ oraz $\Gamma \vdash N : \sigma$, to $\Gamma \vdash M[x := N] : \tau$.

Dowód: Dowód jest przez indukcję ze względu na budowę termu M i wykorzystuje lemat o generowaniu. Na przykład jeśli $M = PQ$ to $\Gamma, x : \sigma \vdash M : \tau$ implikuje $\Gamma, x : \sigma \vdash P : \rho \rightarrow \tau$ i $\Gamma, x : \sigma \vdash Q : \rho$. Można więc zastosować założenie indukcyjne do P i Q , otrzymując $\Gamma \vdash P[x := N] : \rho \rightarrow \tau$ i $\Gamma \vdash Q[x := N] : \rho$, skąd łatwo wynika teza. Pozostałe przypadki są podobne (ćwiczenie). ■

Twierdzenie 18.7 (poprawność redukcji)¹⁹ *Jeśli $\Gamma \vdash M : \tau$ oraz $M \rightarrow_{\beta\eta} N$, to $\Gamma \vdash N : \tau$.*

Dowód: Dowód jest przez indukcję ze względu na definicję redukcji. Przypadki nietrywialne mają miejsce gdy M jest beta- lub eta-redeksem.

Jeśli $M = (\lambda x.P)Q \rightarrow_{\beta} P[x := Q] = N$, to z lematu 18.5 otrzymujemy $\Gamma \vdash \lambda x.P : \sigma \rightarrow \tau$ i $\Gamma \vdash Q : \sigma$. Wtedy $\Gamma, x : \sigma \vdash P : \tau$ (znowu z lematu o generowaniu) i pozostaje użyć lematu 18.6.

Niech więc $M = \lambda x.Nx \rightarrow_{\eta} N$ (gdzie $x \notin \text{FV}(M)$). Bez straty ogólności można założyć, że $\tau \neq \omega$. Z lematu o generowaniu typ τ jest więc iloczynem postaci $\bigcap \{\sigma_i \rightarrow \tau_i \mid i \in I\}$ i na dodatek $\Gamma, x : \sigma_i \vdash Nx : \tau_i$ dla wszystkich $i \in I$. Stosując dalej lemat o generowaniu otrzymamy $\Gamma, x : \sigma_i \vdash N : \xi_i \rightarrow \tau_i$ oraz $\Gamma, x : \sigma_i \vdash x : \xi_i$. Wtedy $\sigma_i \leq \xi_i$, więc $\Gamma, x : \sigma_i \vdash N : \sigma_i \rightarrow \tau_i$. Z lematu 18.1 wnioskujemy, że $\Gamma \vdash N : \sigma_i \rightarrow \tau_i$, a skoro tak jest dla wszystkich i , to wreszcie $\Gamma \vdash N : \tau$. ■

System typów iloczynowych BCD jest w tym wyjątkowy, że zachodzi następujące twierdzenie (nieprawdziwe dla większości innych systemów):

Twierdzenie 18.8 *Jeżeli $M \rightarrow_{\beta} N$, to warunki $\Gamma \vdash M : \tau$ i $\Gamma \vdash N : \tau$ są równoważne.*²⁰

Dowód: Implikacja z lewej do prawej to oczywiście twierdzenie 18.7. Dla dowodu implikacji odwrotnej potrzebna nam taka własność:

Jeśli $\Gamma \vdash M[x := N] : \tau$, to istnieje taki typ σ , że $\Gamma \vdash N : \sigma$ i $\Gamma, x : \sigma \vdash M : \tau$,

czyli twierdzenie odwrotne do lematu 18.6. Dowód tej własności przebiega przez indukcję ze względu na budowę termu M , z pomocą lematu o generowaniu (lemat 18.5). Przypadek $M = y \neq x$ wymaga odwołania się do stałej ω . ■

Wniosek 18.9 *Jeżeli $M =_{\beta} N$, to termy M i N mają te same typy w każdym otoczeniu.*

Ćwiczenia

1. Czy twierdzenie 18.7 pozostanie prawdziwe, gdy z definicji $\leq$ usuniemy warunek $\omega \leq \omega \rightarrow \omega$?

¹⁹ *Subject reduction property*

²⁰ *Subject conversion property*

2. Wskazać przykłady termów, które nie mają żadnego typu prostego, ale mają typy iloczynowe, nie zawierające ω .
3. Czy z tego, że term M ma typ w otoczeniu Γ wynika, że $\Gamma(x)$ jest określone dla wszystkich $x \in \text{FV}(M)$?
4. Czy twierdzenie 18.8 uogólnia się dla η -redukcji?

19 Model z filtrów

Wracamy do semantyki typów, ale tym razem iloczynowych. Znaczenie typów przy wartościowaniu typowym ξ definiujemy jak poprzednio, dodając jeszcze dwie klauzule.

- $\llbracket s \rrbracket_\xi = \xi(s)$, gdy s jest zmienną typową;
- $\llbracket \omega \rrbracket_\xi = D$;
- $\llbracket \sigma \cap \tau \rrbracket_\xi = \llbracket \sigma \rrbracket_\xi \cap \llbracket \tau \rrbracket_\xi$;
- $\llbracket \sigma \rightarrow \tau \rrbracket_\xi = \llbracket \sigma \rrbracket_\xi \Rightarrow \llbracket \tau \rrbracket_\xi$.

Terminologia i notacja pozostaje taka jak dla typów prostych. Twierdzenie o poprawności pozostaje prawdziwe.

Lemat 19.1 *Jeśli $\sigma \leq \tau$, to $\llbracket \sigma \rrbracket_\xi \subseteq \llbracket \tau \rrbracket_\xi$ dla dowolnego ξ .*

Dowód: Ćwiczenie. ■

Twierdzenie 19.2 (o poprawności) *Jeśli $\Gamma \vdash M : \sigma$, to $\Gamma \models M : \sigma$.*

Dowód: Podobny do dowodu twierdzenia 17.1. Przypadek reguły $(\leq)$ wynika wprost z lematu 19.1. ■

Teraz zabieramy się za dowód pełności. Podzbiór $F \subseteq \mathcal{T}$ nazywamy *filtrem*, jeżeli spełnione są takie warunki:

- F jest niepusty;
- Jeżeli $\sigma, \tau \in F$, to $\sigma \cap \tau \in F$;
- Jeżeli $\sigma \in F$ i $\sigma \leq \tau$, to $\tau \in F$.

Przykładem filtru jest każdy zbiór postaci $\sigma \uparrow = \{\tau \mid \sigma \leq \tau\}$, nazywany *filtrem głównym*. W zbiorze $\mathbf{F}$ wszystkich filtrów określamy operację aplikacji:

$$F_1 \cdot F_2 = \{\tau \mid \sigma \rightarrow \tau \in F_1 \text{ dla pewnego } \sigma \in F_2\}.$$

Wartościowanie $v : V \rightarrow \mathbf{F}$ przypisuje każdej zmiennej x pewien filtr $v(x)$. Nieformalnie, jest to filtr złożony z typów „dozwolonych” dla x . Mówimy, że otoczenie Γ jest *zgodne* z wartościowaniem v , gdy $\Gamma(x) \in v(x)$ dla wszystkich x , dla których $\Gamma(x)$ jest określone.

Definiujemy teraz znaczenie termu M w zbiorze $\mathbf{F}$ (przy wartościowaniu v), jako zbiór typów „dozwolonych” dla M :

$$\llbracket M \rrbracket_v = \{\sigma \mid \Gamma \vdash M : \sigma, \text{ dla pewnego } \Gamma \text{ zgodnego z } v\}.$$

W następnym dowodzie (i nie tylko) przyda się następująca definicja. Jeśli Γ_1 i Γ_2 są otoczeniami, to przez $\Gamma_1 \& \Gamma_2$ oznaczamy otoczenie, którego dziedzina jest sumą dziedzin Γ_1 i Γ_2 , i które jest określone tak:

$$(\Gamma_1 \& \Gamma_2) = \begin{cases} \Gamma_1(x) \cap \Gamma_2(x), & \text{jeśli } \Gamma_1(x) \text{ i } \Gamma_2(x) \text{ są określone;} \\ \Gamma_1(x), & \text{jeśli tylko } \Gamma_1(x) \text{ jest określone;} \\ \Gamma_2(x), & \text{jeśli tylko } \Gamma_2(x) \text{ jest określone.} \end{cases}$$

Lemat 19.3 *Struktura $\mathcal{F} = \langle \mathbf{F}, \cdot, \llbracket \cdot \rrbracket \rangle$ jest lambda-modelem.*

Dowód: Naszym zadaniem jest sprawdzenie następujących warunków:

- (a) Jeśli x jest zmienną, to $\llbracket x \rrbracket_v = v(x)$;
- (b) $\llbracket PQ \rrbracket_v = \llbracket P \rrbracket_v \llbracket Q \rrbracket_v$;
- (c) $\llbracket \lambda x. P \rrbracket_v \cdot F = \llbracket P \rrbracket_{v[x \mapsto F]}$, dla dowolnego $F \in \mathbf{F}$;
- (d) Jeśli $v|_{\mathbf{FV}(P)} = u|_{\mathbf{FV}(P)}$, to $\llbracket P \rrbracket_v = \llbracket P \rrbracket_u$.
- (e) Jeśli $\llbracket \lambda x. M \rrbracket_v \approx \llbracket \lambda x. N \rrbracket_v$ to $\llbracket \lambda x. M \rrbracket_v = \llbracket \lambda x. N \rrbracket_v$.

Zaczynamy od (a). Przypuśćmy, że $\sigma \in \llbracket x \rrbracket_v$. Z definicji oznacza to, że $\Gamma \vdash x : \sigma$ dla pewnego Γ zgodnego z v , czyli takiego, że $\Gamma(x) \in v(x)$. Z lematu o generowaniu wiemy, że wtedy $\sigma \geq \Gamma(x)$, więc tym bardziej $\sigma \in v(x)$. Mamy więc inkluzję $\llbracket x \rrbracket_v \subseteq v(x)$. Inkluzja odwrotna jest jeszcze łatwiejsza, bo dla $\sigma \in v(x)$ otoczenie $\{x : \sigma\}$ jest zgodne z v .

W punkcie (b) pokażemy inkluzję $\llbracket P \rrbracket_v \llbracket Q \rrbracket_v \subseteq \llbracket PQ \rrbracket_v$, która jest mniej oczywista. Załóżmy, że $\sigma \in \llbracket P \rrbracket_v \llbracket Q \rrbracket_v$. Jest więc takie $\zeta \in \llbracket Q \rrbracket_v$, że $\zeta \rightarrow \sigma \in \llbracket P \rrbracket_v$. Istnieją zatem otoczenia Γ_1 i Γ_2 zgodne z v i takie, że $\Gamma_1 \vdash P : \zeta \rightarrow \sigma$ oraz $\Gamma_2 \vdash Q : \zeta$. Nietrudno sprawdzić, że $\Gamma_0 = \Gamma_1 \& \Gamma_2$ jest zgodne z v , oraz że $\Gamma_0 \vdash P : \zeta \rightarrow \sigma$ i $\Gamma_0 \vdash Q : \zeta$. Stąd $\Gamma_0 \vdash PQ : \sigma$ i mamy $\sigma \in \llbracket PQ \rrbracket_v$.

W części (c) „trudniejsza” jest inkluzja z lewej do prawej. Niech $\sigma \in \llbracket \lambda x. P \rrbracket_v \cdot F$. Oznacza to, że $\Gamma \vdash \lambda x. P : \tau \rightarrow \sigma$ i $\tau \in F$, gdzie Γ jest zgodne z v . Na mocy lematu o generowaniu $\Gamma, x : \tau \vdash P : \sigma$. Ponieważ $\Gamma, x : \tau$ jest zgodne z $v[x \mapsto F]$, więc $\sigma \in \llbracket P \rrbracket_{v[x \mapsto F]}$ i dobrze.

Warunek (d) wynika wprost z lematu 18.1.

Możemy teraz już stwierdzić, że struktura filtrów jest lambda-interpretacją. Pozostaje sprawdzenie warunku (e). Niech więc $\llbracket \lambda x. M \rrbracket_v \approx \llbracket \lambda x. N \rrbracket_v$, czyli $\llbracket \lambda x. M \rrbracket_v \cdot F = \llbracket \lambda x. N \rrbracket_v \cdot F$ dla dowolnego filtru F . Przypuśćmy, że $\sigma \in \llbracket \lambda x. M \rrbracket_v$. Wtedy, na mocy lematu o generowaniu, $\sigma \equiv \bigcap \{\sigma_i \rightarrow \tau_i \mid i \in I\}$ oraz $\Gamma, x : \sigma_i \vdash M : \tau_i$ dla $i \in I$, dla pewnego Γ , zgodnego z v .

Ponieważ otoczenie $\Gamma, x : \sigma_i$ jest zgodne z wartościowaniem $v[x \mapsto \sigma_i \uparrow]$, więc typ τ_i należy do filtru $\llbracket M \rrbracket_{v[x \mapsto \sigma_i \uparrow]} = \llbracket \lambda x M \rrbracket_v \cdot \sigma_i \uparrow = \llbracket \lambda x N \rrbracket_v \cdot \sigma_i \uparrow = \llbracket N \rrbracket_{v[x \mapsto \sigma_i \uparrow]}$. A więc mamy $\Gamma'_i \vdash N : \tau_i$, gdzie Γ'_i jest pewnym otoczeniem zgodnym z $v[x \mapsto \sigma_i \uparrow]$. Ale wtedy $\Gamma'_i(x) \geq \sigma_i$, więc tym bardziej $\Gamma'_i(x \mapsto \sigma_i) \vdash N : \tau_i$, czyli $\Gamma'_i \vdash \lambda x.N : \sigma_i \rightarrow \tau_i$. Ponieważ Γ'_i jest także zgodne z v , więc $\sigma_i \rightarrow \tau_i \in \llbracket \lambda x.N \rrbracket$. Ostatecznie $\sigma \in \llbracket \lambda x.N \rrbracket$, bo σ to iloczyn wszystkich $\sigma_i \rightarrow \tau_i$. ■

Lemat 19.4 W modelu $\mathcal{F} = \langle \mathbf{F}, \cdot, \llbracket \cdot \rrbracket \rangle$ określamy wartościowanie typowe $\xi(p) = \{F \mid p \in F\}$. Wtedy $\llbracket \tau \rrbracket_\xi = \{F \mid \tau \in F\}$, dla dowolnego typu τ .

Dowód: Dowód jest przez indukcję ze względu na τ . Przypadek iloczynu jest łatwy, bo zachodzi równoważność

$$(\sigma \in F) \wedge (\rho \in F) \iff \sigma \cap \rho \in F.$$

Niech więc $\tau = \rho \rightarrow \sigma$. Mamy wykazać, że typ $\rho \rightarrow \sigma$ należy do filtru F wtedy i tylko wtedy, gdy $F \in \llbracket \rho \rrbracket_\xi \Rightarrow \llbracket \sigma \rrbracket_\xi$.

Implikacja z lewej do prawej jest łatwa. Jeśli bowiem $G \in \llbracket \rho \rrbracket_\xi$, to z założenia indukcyjnego dla ρ wynika $\rho \in G$, z zatem $\sigma \in F \cdot G$ wprost z definicji aplikacji $F \cdot G$.

Na odwrót, przypuśćmy, że $F \in \llbracket \rho \rrbracket_\xi \Rightarrow \llbracket \sigma \rrbracket_\xi$. Ponieważ $\rho \in \rho \uparrow$ więc z założenia indukcyjnego dla ρ mamy $\rho \uparrow \in \llbracket \rho \rrbracket_\xi$, a stąd $F \cdot \rho \uparrow \in \llbracket \sigma \rrbracket_\xi$, czyli (na mocy założenia indukcyjnego dla σ) typ σ należy do $F \cdot \rho \uparrow$. Istnieje więc takie $\mu \in \rho \uparrow$, że $\mu \rightarrow \sigma \in F$. Wtedy jednak $\mu \geq \rho$, więc $\rho \rightarrow \sigma \geq \mu \rightarrow \sigma$, i typ $\rho \rightarrow \sigma$ też należy do F . ■

Twierdzenie 19.5 (o pełności) Warunki $\Gamma \vdash M : \sigma$ i $\Gamma \models M : \sigma$ są równoważne.

Dowód: Przypuśćmy, że $\Gamma \vdash M : \sigma$. W modelu $\mathcal{F} = \langle \mathbf{F}, \cdot, \llbracket \cdot \rrbracket \rangle$ definiujemy wartościowanie v , przyjmując $v(x) = \tau \uparrow$ dla dowolnego $(x : \tau) \in \Gamma$. Ponieważ wtedy $\tau \in v(x)$, więc $v(x) \in \llbracket \tau \rrbracket_\xi$ (lemat 19.4). Zatem $\mathcal{F}, v, \xi \models \Gamma$, a wobec tego $\mathcal{F}, v, \xi \models M : \sigma$, gdzie ξ jest takie jak w Lemacie 19.4. To oznacza, że $\llbracket M \rrbracket_v \in \llbracket \sigma \rrbracket_\xi$, czyli $\sigma \in \llbracket M \rrbracket_v$. Jest więc otoczenie Γ' zgodne z v , takie że $\Gamma' \vdash M : \sigma$. Ponieważ Γ' jest zgodne z v , więc $\Gamma'(x) \in v(x)$, czyli $\Gamma(x) \leq \Gamma'(x)$ dla wszystkich x . A więc tym bardziej $\Gamma \vdash M : \sigma$. ■

Ćwiczenia

1. Udowodnić, że jeśli F_1 i F_2 są filtrami, to $F_1 \cdot F_2$ jest filtrem.
2. Udowodnić, że $\llbracket M \rrbracket_v$ zawsze jest filtrem.
3. Czy model z filtrów jest ekstensjonalny?

20 Typy iloczynowe bez omegi

System BCD, o którym była mowa do tej pory, przypisywał każdemu bez wyjątku termowi „trywialny” typ ω . Jeżeli z tego systemu usuniemy aksjomat

$$\Gamma \vdash M : \omega,$$

to już nie każdy term będzie miał typ. Nabiera więc sensu następująca definicja. Term M jest *typowalny* wtedy i tylko wtedy, gdy $\Gamma \vdash M : \tau$, dla pewnych Γ i τ . W istocie, jak się okaże poniżej, termy typowalne za pomocą typów iloczynowych bez omegi, to dokładnie termy silnie normalizowalne.

Dla porządku zaczniemy od definicji. Zbiór typów $\mathcal{T}_\cap$ definiujemy tak samo jak zbiór $\mathcal{T}_{\cap\omega}$, ale z pominięciem omegi:

- Typy atomowe należą do $\mathcal{T}_\cap$;
- Jeśli σ i τ są w $\mathcal{T}_\cap$, to $\sigma \rightarrow \tau$ i $\sigma \cap \tau$ też należą do $\mathcal{T}_\cap$.

W zbiorze $\mathcal{T}_\cap$ określamy relację $\leq$ jako najmniejszy quasiporządek $\leq$ spełniający warunki

$$\sigma \cap \tau \leq \sigma, \quad \sigma \cap \tau \leq \tau, \quad \sigma \leq \sigma \cap \sigma$$

$$(\sigma \rightarrow \tau) \cap (\sigma \rightarrow \rho) \leq \sigma \rightarrow \tau \cap \rho$$

$$\text{Jeśli } \sigma \leq \sigma' \text{ i } \tau \leq \tau', \text{ to } \sigma \cap \tau \leq \sigma' \cap \tau' \text{ oraz } \sigma' \rightarrow \tau \leq \sigma \rightarrow \tau'.$$

Będziemy teraz przypisywać termom typy za pomocą tych samych reguł co poprzednio, ale z wyłączeniem aksjomatu (ω) . Dla odróżnienia od systemu BCD ten nowy system nazwijmy *systemem* $\lambda_{\cap\leq}$. Oto aksjomat i reguły systemu $\lambda_{\cap\leq}$.

$$(\text{Var}) \quad \Gamma, x : \sigma \vdash x : \sigma$$

$$(\text{Abs}) \quad \frac{\Gamma, x : \sigma \vdash M : \tau}{\Gamma \vdash (\lambda x.M) : \sigma \rightarrow \tau} \qquad (\text{App}) \quad \frac{\Gamma \vdash M : \sigma \rightarrow \tau \quad \Gamma \vdash N : \sigma}{\Gamma \vdash (MN) : \tau}$$

$$(\cap\text{I}) \quad \frac{\Gamma \vdash M : \sigma \quad \Gamma \vdash M : \tau}{\Gamma \vdash M : \sigma \cap \tau} \qquad (\leq) \quad \frac{\Gamma \vdash M : \sigma}{\Gamma \vdash M : \tau} \quad (\sigma \leq \tau)$$

Asercje wyprowadzalne w systemie $\lambda_{\cap\leq}$ możemy zapisywać $\Gamma \vdash_\cap M : \sigma$, ale zwykle zamiast $\vdash_\cap$ napiszemy po prostu $\vdash$. O którym systemie mowa, powinno być wiadomo z kontekstu.

Dla systemu $\lambda_{\cap\leq}$ zachodzi następujący wariant lematu 18.1.

Lemat 20.1 *Jeśli $\Gamma \vdash M : \tau$, to $\Gamma(x)$ jest określone dla każdego $x \in \text{FV}(M)$. Ponadto $\Gamma|_{\text{FV}(M)} \vdash M : \tau$, gdzie $\Gamma|_{\text{FV}(M)} = \{(x : \Gamma(x)) \mid x \in \text{FV}(M)\}$.*

Dowód: Łatwa indukcja ze względu na rozmiar wyprowadzenia. ■

Różnica pomiędzy treścią lematu 20.1 i lematu 18.1 bierze się oczywiście z braku aksjomatu (ω) . Teraz, aby przypisać typ całemu termowi, trzeba zadeklarować typy wszystkich jego zmiennych wolnych.

Podstawowe własności systemu $\lambda_{\cap\leq}$ są podobne do własności systemu BCD. W szczególności pozostają prawdziwe lematy 18.2, 18.3, a także lemat 18.5 o generowaniu (z pominięciem

tego, co odnosi się do omegi). Podobnie jak poprzednio, wyprowadzamy stąd lemat o podstawieniu (odpowiednik lematu 18.6) a następnie twierdzenie o poprawności redukcji. Dowody tych faktów pozostają niezmienione (ale prostsze, bo niepotrzebne są rozważania dotyczące omegi).

Dla systemu $\lambda_{\cap\leq}$ nie zachodzi jednak twierdzenie 18.8. Na przykład nietypowalny term $\mathbf{KK}\Omega$ redukuje się do typowalnego $\mathbf{K}$.

Silna normalizacja

Udowodnimy teraz twierdzenie o silnej normalizacji dla systemu $\lambda_{\cap\leq}$. Nasz dowód będzie niewielką adaptacją dowodu twierdzenia 13.9 o silnej normalizacji rachunku z typami prostymi. Definicja $\llbracket \tau \rrbracket$ ma teraz o jeden przypadek więcej:

- $\llbracket p \rrbracket := \text{SN}$, gdy p jest atomem;
- $\llbracket \tau \rightarrow \sigma \rrbracket := \{M : \tau \rightarrow \sigma \mid \forall N (N \in \llbracket \tau \rrbracket \Rightarrow MN \in \llbracket \sigma \rrbracket)\}$.
- $\llbracket \tau \cap \sigma \rrbracket := \llbracket \tau \rrbracket \cap \llbracket \sigma \rrbracket$.

Lematy 13.5 i 13.7 pozostają prawdziwe dla tej definicji, a ich dowody pozostają praktycznie bez zmian (przypadek iloczynu to natychmiastowe zastosowanie założenia indukcyjnego). lemat 13.6 w ogóle nie wymaga żadnej adaptacji. Dowód lematu 13.8 ulega tylko dwu drobnym zmianom. Po pierwsze, w przypadku zmiennej musimy się odwołać do takiej dodatkowej własności:

Lemat 20.2 *Jeśli $\tau \leq \sigma$, to $\llbracket \tau \rrbracket \subseteq \llbracket \sigma \rrbracket$.*

Dowód: Indukcja ze względu na definicję relacji $\leq$. ■

Po drugie, w przypadku abstrakcji zamiast zwykłego $\sigma \rightarrow \rho$ mamy iloczyn $\bigcap \{\sigma_i \rightarrow \rho_i \mid i \in I\}$, gdzie $\Gamma, y : \sigma_i \vdash P : \rho_i$ dla $i \in I$, i musimy pokazać, że $M[\vec{N}/\vec{x}]$ jest stabilny dla wszystkich typów $\sigma_i \rightarrow \rho_i$. Trzeba więc cierpliwie podpisywać wszędzie indeks i . Otrzymamy w końcu

Twierdzenie 20.3 (o silnej normalizacji) *System $\lambda_{\cap\leq}$ ma własność silnej normalizacji: każdy term typowalny jest SN.*

Ćwiczenia

1. Sprawdzić, że lematy 18.3–18.6 i twierdzenie 18.7 dla systemu BCD przenoszą się na system $\lambda_{\cap\leq}$.
2. Udowodnić lemat 20.2.
3. Rozpatrzmy system $\lambda_{\cap}$ powstały z systemu $\lambda_{\cap\leq}$ przez zamianę reguły $(\leq)$ na regułę

$$(\cap E) \frac{\Gamma \vdash M : \sigma_1 \cap \sigma_2}{\Gamma \vdash M : \sigma_i}$$

Pokazać, że termy typowalne w obu systemach są takie same.

4. Czy w systemie $\lambda_{\cap}$ zachodzi twierdzenie o poprawności eta-redukcji?
5. System $\lambda_{\cap\eta}$ to rozszerzenie systemu $\lambda_{\cap}$ o dodatkową regułę

$$(\eta) \frac{\Gamma \vdash M : \sigma, \quad M \rightarrow_{\eta} N}{\Gamma \vdash N : \sigma}$$

Pokazać, że w systemach $\lambda_{\cap\eta}$ i $\lambda_{\cap\leq}$ wyprowadzalne są dokładnie te same asercje typowe.

Twierdzenie Pottingera

Twierdzenie o silnej normalizacji zachodzi dla wielu systemów przypisania typów. Jest to wręcz uważane za dość podstawową własność takich systemów, a system BCD z regułą (ω) stanowi raczej wyjątek, uzasadniony swoimi semantycznymi korzeniami.

Spośród systemów o własności SN, typy iloczynowe stanowią mechanizm najsilniejszy w tym sensie, że twierdzenie o silnej normalizacji jest dla tego systemu odwracalne: każdy term o własności SN jest typowalny. Pierwszy dowód tego faktu opublikował w 1980 roku G. Pottinger. Ten dowód i inne późniejsze dowody zawierały różne usterki. Pierwszy poprawny dowód (niepublikowany) podała w latach dziewięćdziesiątych B. Venneri. Poniższy dowód, mój nadzieję, też jest poprawny.

Naturalna idea dowodu jest oczywiście taka: Zacząć od typowania postaci normalnych i dowieść twierdzenia przez indukcję ze względu na długość redukcji termu do postaci normalnej. Ale zauważyliśmy już, że dla systemu $\lambda_{\cap\leq}$ nie zachodzi twierdzenie 18.8. W szczególności nie zachodzi twierdzenie odwrotne do lematu 18.6. Mamy jednak nieco słabszą własność.

Lemat 20.4 *Jeśli $\Gamma \vdash M[x := N] : \tau$ oraz $\Gamma \vdash N : \rho$, to istnieje taki typ σ , że $\Gamma \vdash N : \sigma$ oraz $\Gamma, x : \sigma \vdash M : \tau$,*

Dowód: Ćwiczenie. ■

Lemat 20.5 *Każdy term w postaci normalnej jest typowalny.*

Dowód: Indukcja ze względu na budowę termu (ćwiczenie). ■

Twierdzenie 20.6 (Pottinger) *Term jest typowalny w systemie $\lambda_{\cap\leq}$ wtedy i tylko wtedy, gdy jest silnie normalizowalny.*

Dowód: Implikacja z lewej do prawej to twierdzenie 20.3 o silnej normalizacji. Dowód implikacji odwrotnej przebiega przez indukcję ze względu na maksymalną możliwą długość ciągu redukcji danego termu M , którą oznaczmy przez $\delta(M)$. Jeśli $\delta(M) = 0$, czyli mamy do czynienia z postacią normalną, to stosujemy lemat 20.5. W kroku indukcyjnym zakładamy, że każdy term N o własności $\delta(N) < n$ jest typowalny i dowodzimy, że jeśli $\delta(M) = n$ to M jest typowalny. Dowód jest przez „lokalną” indukcję ze względu na budowę termu M .

Niech $M \rightarrow_L M'$. Wtedy $\delta(M') < n$, zatem M' jest typowalny. Przypuśćmy, że $\Gamma' \vdash M' : \tau$.

Przypadek 1: $M = \lambda x.N \rightarrow_L \lambda x.N' = M'$. Z lematu o generowaniu typ τ termu M' musi być postaci $\bigcap \{\sigma_i \rightarrow \tau_i \mid i \in I\}$, gdzie $\Gamma, x : \sigma_i \vdash N' : \tau_i$. Z „lokalnego” założenia indukcyjnego term N jest typowalny, więc abstrakcja $M = \lambda x.N$ też jest typowalna.

Przypadek 2: $M = xN_1 \dots N_i \dots N_k \rightarrow_L xN_1 \dots N'_i \dots N_k = M'$, gdzie $N_i \rightarrow_L N'_i$. Skoro term M' ma typ w otoczeniu Γ' , to także termy $N_1, \dots, N'_i, \dots, N_k$ mają pewne typy $\rho_1, \dots, \rho'_i, \dots, \rho_k$ w tym otoczeniu. Z założenia indukcyjnego wiemy, że $\Gamma \vdash N_i : \rho_i$ dla pewnych Γ, ρ_i . Jeśli teraz przyjmiemy $\Gamma_0 = \Gamma' \& \Gamma \& (x : \rho_1 \rightarrow \dots \rightarrow \rho_i \rightarrow \dots \rightarrow \rho_k \rightarrow s)$, gdzie s jest typem atomowym, to otrzymamy $\Gamma_0 : M : s$.

Przypadek 3: Przyszedł czas na przypadek redeksu czołowego:

$$M = (\lambda x.P)QN_1 \dots N_k \rightarrow_L P[x := Q]N_1 \dots N_k = M'.$$

Stosując k -krotnie pierwszą część lematu o generowaniu, wnioskujemy, że $\Gamma' \vdash N_i : \rho_i$ oraz $\Gamma' \vdash P[x := Q] : \rho_1 \rightarrow \dots \rightarrow \rho_k \rightarrow \tau$ dla pewnych ρ_i i σ . Ponadto term Q jest silnie normalizowalny, jako podterm silnie normalizowalnego termu M . Ponieważ M ma redeks czołowy, więc najdłuższa możliwa redukcja termu M jest dłuższa przynajmniej o jeden krok od każdej redukcji termu Q . Zatem $\delta(Q) < \delta(M)$, skąd Q jest termem typowalnym, np. $\Gamma'' \vdash Q : \zeta$. Niech $\Gamma = \Gamma' \& \Gamma''$. Wtedy nadal $\Gamma \vdash P[x := Q] : \rho_1 \rightarrow \dots \rightarrow \rho_k \rightarrow \rho$ oraz Q ma typ w Γ . Z lematu 20.4 dostajemy $\Gamma \vdash Q : \sigma$ oraz $\Gamma, x : \sigma \vdash P : \rho_1 \rightarrow \dots \rightarrow \rho_k \rightarrow \rho$, dla pewnego typu σ . Stąd $\Gamma \vdash \lambda x.P : \sigma \rightarrow \rho_1 \rightarrow \dots \rightarrow \rho_k \rightarrow \rho$, i dalej $\Gamma \vdash M : \rho$. ■

Problemem *typowalności* (typability) (albo problemem *rekonstrukcji typu*) dla danego systemu przypisania typów nazywamy taki problem decyzyjny:

Czy dany term M jest typowalny?

Zbliżonym zagadnieniem jest problem *sprawdzenia typu* (type checking):

Dane są M, Γ i τ . Czy $\Gamma \vdash M : \tau$?

Wniosek 20.7 *Problem typowalności dla typów iloczynowych jest nierozstrzygalny.*

Dowód: Wynika to wprost z twierdzenia 20.6 i z nierozstrzygalności silnej normalizacji (wniosek 7.11). ■

Wniosek 20.8 *Problem sprawdzenia typu dla typów iloczynowych też jest nierozstrzygalny.*

Dowód: Redukujemy problem typowalności do problemu sprawdzenia typu. Niech $\vec{x}$ będą wszystkimi zmiennymi wolnymi termu M . Term M jest typowalny wtedy i tylko wtedy, gdy $x : s \vdash \mathbf{K}x(\lambda \vec{x}M) : s$. ■

Ćwiczenia

1. Udowodnić lemat 20.5. *Wskazówka:* term M jest w postaci normalnej, wtedy i tylko wtedy, gdy zachodzi jedna z trzech możliwości: (a) M jest zmienną, (b) $M = \lambda x.N$, gdzie N jest w postaci normalnej, (c) $M = x\vec{N}$, gdzie $\vec{N}$ są w postaci normalnej.

2. Czy w dowodzie twierdzenia 20.6 można wzmocnić założenie indukcyjne do warunku:
Jeśli $\delta(M) = n$ oraz $M \rightarrow_L M'$ i przy tym $\Gamma \vdash M' : \tau$, to $\Gamma \vdash M : \tau$?
3. Czy w dowodzie twierdzenia 20.6, w założeniu indukcyjnym można zastąpić warunek $M \rightarrow_L M'$ przez $M \rightarrow M'$?

21 Problem niepustości typu

Dualny do problemu typowalności jest *problem niepustości typu*, nazywany też problemem *inhabitacji*, i formułowany tak:

Czy istnieje term zamknięty danego typu?

Nieco ogólniejsze (ale równoważne) sformułowanie jest takie:

Czy dla danych Γ i τ istnieje term M , który ma typ τ w otoczeniu Γ ?

Pokażemy szkic dowodu, że problem inhabitacji dla typów iloczynowych jest nierozstrzygalny. Poniżej mowa o systemie $\lambda_\cap$, ale takie samo rozumowanie stosuje się dla innych wariantów. Zaczniemy od wyspecjalizowanego automatu.

Automaty wierszowe

Definicja 21.1 *Automat wierszowy* to krotka postaci

$$\mathcal{A} = \langle \Sigma, Q, q_0, F, \varrho \rangle,$$

gdzie Σ jest skończonym alfabetem, Q to skończony zbiór stanów, $q_0 \in Q$ jest stanem początkowym i $F \subseteq Q$ jest zbiorem stanów końcowych. Ostatnia składowa to funkcja przejścia $\varrho : (Q - F) \times (\Sigma \cup \{\varepsilon\}) \rightarrow Q \times (\Sigma \cup \{\varepsilon\})$, która musi spełniać warunek: jeśli $\varrho(q, a) = (p, b)$ to albo $a, b \in \Sigma$ albo $a = b = \varepsilon$.

Konfiguracją automatu jest trójka $\langle w, q, v \rangle$, gdzie q jest stanem, $wv \in \Sigma^*$ jest słowem (reprezentującym zawartość taśmy). Interpretacja jest taka, że głowica maszyny „widzi” pierwszy symbol słowa v (lub ε , gdy $v = \varepsilon$). Zmianę konfiguracji w jednym kroku wyraża funkcja $\bar{\varrho}$:

- $\bar{\varrho}(\langle w, q, av \rangle) = \langle wb, p, v \rangle$, jeśli $a \neq \varepsilon$ oraz $\varrho(q, a) = (p, b)$;
- $\bar{\varrho}(\langle w, q, \varepsilon \rangle) = \langle \varepsilon, p, w \rangle$, jeśli $\varrho(q, \varepsilon) = (p, \varepsilon)$.

Język *akceptowany* przez $\mathcal{A}$ definiujemy tak:

$$L(\mathcal{A}) = \{w \in \Sigma^* \mid \bar{\varrho}^k(\langle \varepsilon, q_0, w \rangle) = \langle u, q, v \rangle, \text{ dla pewnego } k \in \mathbb{N} \text{ i pewnego } q \in F\}.$$

Nieformalnie: maszyna czyta słowo od lewej do prawej, przy czym może pisać i zmieniać stan, ale nie może się cofać. Dopiero po dojściu do końca słowa cofa się na sam początek i znowu pracuje od lewej do prawej. Zatrzymuje się po wejściu w stan końcowy.

Fakt 21.2 *Problem niepustości dla automatów wierszowych (czy $L(\mathcal{A}) \neq \emptyset$ dla danego $\mathcal{A}$?) jest nierozstrzygalny.*

Dowód: Wiadomo, że nierozstrzygalny jest problem stopu dla automatów kolejkowych, przy czym można zakładać, że dany automat rozpoczyna pracę z pustą kolejką. Dla danego automatu kolejkowego $\mathcal{K}$ można zbudować automat wierszowy $\mathcal{A}$ o tej własności, że $\mathcal{K}$ zatrzymuje się wtedy i tylko wtedy, gdy $\mathcal{A}$ akceptuje pewne słowo. Słowo to jest postaci $\langle \rangle \# \dots \# \#$ i reprezentuje pustą kolejkę początkową. Automat $\mathcal{A}$ naśladuje działanie automatu $\mathcal{K}$ w ten sposób, że np. kolejka o zawartości 011100010 (początek z lewej) jest przedstawiona na taśmie zawsze jako słowo postaci $\$ \$ \dots \$ \langle 011100010 \rangle \# \dots \# \#$. Instrukcja *insert*(0) jest realizowana przez zamianę $\rangle \#$ na $0 \rangle$, a instrukcja *remove* przez zamianę $\langle 0$ na $\$ \langle$. W ten sposób kolejka przesuwa się wciąż w prawo, a symulacja jest poprawna pod warunkiem, że liczba krzyżyków w słowie wejściowym była dostatecznie duża. ■

Gry na drzewach

Opiszemy teraz gry (dla pojedynczego gracza) polegające na budowaniu drzewa za pomocą dostępnych *reguł*. Niech Σ będzie skończonym alfabetem. Wtedy:

1. *Reguła lokalna* (nad Σ) to skończony niepusty zbiór X par elementów Σ .
2. *Reguła globalna* (nad Σ) to skończony niepusty zbiór G trójek postaci

$$\langle \langle X, b \rangle, \langle Y, c \rangle, d \rangle,$$

gdzie $b, c, d \in \Sigma$ i X, Y są regułami lokalnymi.

3. *Gra* (nad Σ) jest trójką postaci

$$T = \langle a, F, \{G_1, \dots, G_n\} \rangle,$$

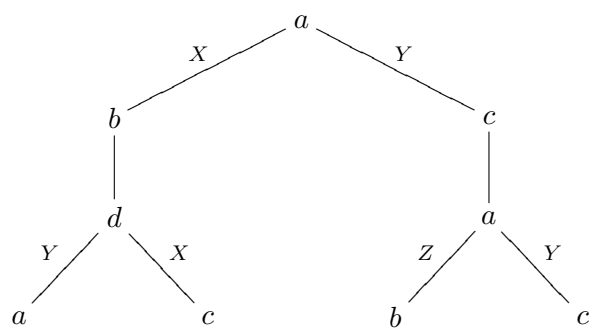
gdzie $a \in \Sigma$ jest *etykietą początkową*, $F \subseteq \Sigma$ jest zbiorem *etykiet końcowych*, a $G_1, \dots, G_n$ są regułami globalnymi.

Pozycję w grze T jest skończone drzewo, którego wierzchołki są etykietowane elementami alfabetu Σ , o takich własnościach:

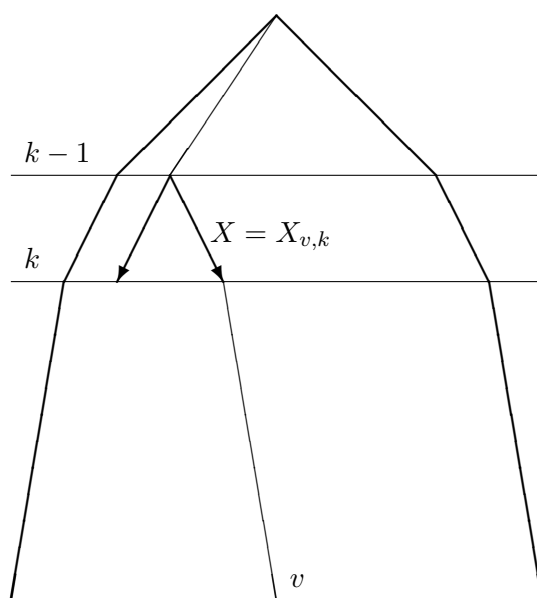
- Korzeń ma etykietę a (etykietę początkową).
- Każdy wierzchołek ma co najwyżej dwóch synów.
- Wierzchołki na tym samym poziomie mają tę samą liczbę synów. W szczególności wszystkie liście są na tym samym poziomie.
- Jeśli wierzchołek v ma dwóch synów v' i v'' , to krawędzie $\langle v, v' \rangle$ and $\langle v, v'' \rangle$ są etykietowane lokalnymi regułami.

Pozycja *początkowa* w grze T składa się tylko z korzenia. Pozycja jest *końcowa*, gdy wszystkie etykiety liści należą do F .

Niech $\mathcal{C}$ będzie pozycją w grze $T = \langle a, F, \{G_1, \dots, G_n\} \rangle$ i niech v będzie liściem $\mathcal{C}$. Przypuśćmy, że dla pewnego k , wszystkie wierzchołki na poziomie $k - 1$ mają dwóch synów



Obrazek 15: Przykładowa pozycja w grze



Obrazek 16: Reguła lokalna X dla wierzchołka v

(Obrazek 16). Pewien wierzchołek u z poziomu $k - 1$ jest przodkiem liścia v , podobnie jak jeden z synów wierzchołka u , powiedzmy u' . Niech X będzie etykietą krawędzi $\langle u, u' \rangle$. Mówimy wtedy, że X jest k -tą regułą lokalną dla v i piszemy $X = X_{v,k}$.

Opiszemy teraz jak można zmienić pozycję. Niech $T = \langle a, F, \{G_1, \dots, G_n\} \rangle$ i niech $\mathcal{C}$ będzie pozycją. Mamy dwie możliwości otrzymania nowej pozycji.

1. Można wykonać krok „globalny”, wybierając jedną z reguł globalnych G_i i wykonując dla każdego liścia v drzewa $\mathcal{C}$ takie czynności:
 - Wybrać taką trójkę $\langle \langle X, b \rangle, \langle Y, c \rangle, d \rangle \in G_i$, że d jest etykietą v ;
 - Utworzyć dwóch synów v , powiedzmy v' i v'' , i przypisać im etykiety b i c ;
 - Oznaczyć krawędź $\langle v, v' \rangle$ przez X a krawędź $\langle v, v'' \rangle$ przez Y .
2. Można też wykonać krok „lokalny”. Zaczynamy go od wybrania poziomu k w $\mathcal{C}$, tak aby każdy wierzchołek na poziomie $k - 1$ miał dwóch synów. Następnie, dla każdego liścia v w drzewie $\mathcal{C}$:
 - Wybieramy parę $\langle c, b \rangle \in X_{v,k}$, gdzie b jest etykietą v ;
 - Tworzymy pojedynczego syna v o etykiecie c .

Oczywiście interesuje nas problem, czy daną grę można *wygrać*, tj. z pozycji początkowej dojść do końcowej. Nazwijmy go *problemem gier na drzewach*. Na następującym przykładzie widać jak gry odpowiadają pewnym przypadkom problemu inhabitacji. Niech $T_0 = \langle 1, \{c\}, \{G_1, G_2\} \rangle$ będzie grą nad $\Sigma = \{a, b, c, 1, 2\}$, gdzie:

- $G_1 = \{ \langle \langle \{(a, a)\}, 1 \rangle, \langle \{(b, a)\}, 2 \rangle, 1 \rangle, \langle \langle \{(a, b)\}, 1 \rangle, \langle \{(b, b)\}, 2 \rangle, 2 \rangle \rangle \}$;
- $G_2 = \{ \langle \langle \{(c, a)\}, a \rangle, \langle \{(c, a)\}, a \rangle, 1 \rangle, \langle \langle \{(c, b)\}, a \rangle, \langle \{(c, b)\}, a \rangle, 2 \rangle \rangle \}$.

Założmy, że elementy alfabetu Σ można uważać za zmienne typowe. Pozycje końcowe w tej grze odpowiadają termom, które mają typ 1 w otoczeniu Γ składającym się z deklaracji:

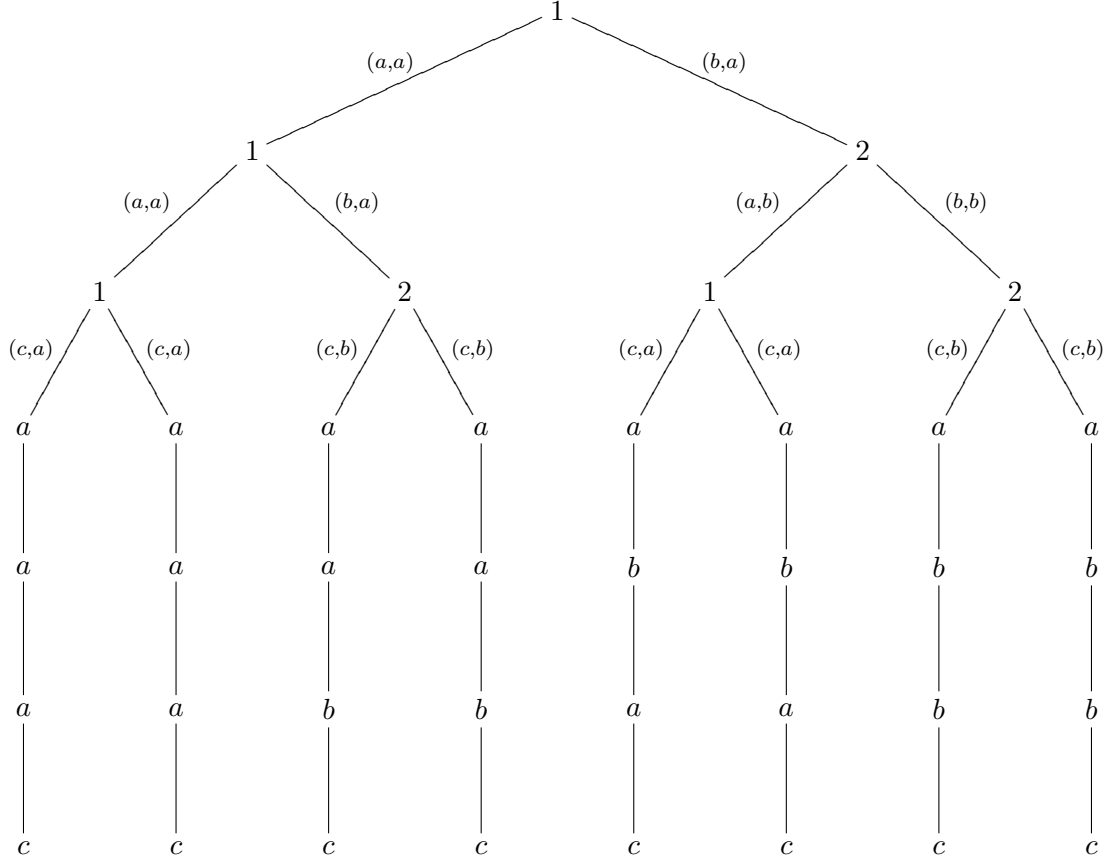
- $x_0 : c$;
- $x_1 : [((a \rightarrow a) \rightarrow 1) \wedge ((b \rightarrow a) \rightarrow 2) \rightarrow 1] \wedge [((a \rightarrow b) \rightarrow 1) \wedge ((b \rightarrow b) \rightarrow 2) \rightarrow 2]$;
- $x_2 : [((c \rightarrow a) \rightarrow a) \wedge ((c \rightarrow a) \rightarrow a) \rightarrow 1] \wedge [((c \rightarrow b) \rightarrow a) \wedge ((c \rightarrow b) \rightarrow a) \rightarrow 2]$.

Na przykład pozycja z Obrazka 17 odpowiada termowi

$$x_1(\lambda y_1(x_1 \lambda y_2(x_2 \lambda y_3. y_1(y_2(y_3 x_0))))).$$

Powyższą obserwację można uogólnić tak:

Fakt 21.3 *Dla dowolnej gry T można efektywnie skonstruować otoczenie Γ i typ a , tak że grę T można wygrać wtedy i tylko wtedy, gdy $\Gamma \vdash M : \tau$ dla pewnego M . A zatem nierozstrzygalność problemu gier na drzewach implikuje nierozstrzygalność inhabitacji.*



Obrazek 17: Pozycja końcowa dla termu $x_1(\lambda y_1(x_1\lambda y_2(x_2\lambda y_3.y_1(y_2(y_3x_0))))))$

Gra cykliczna

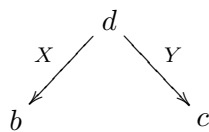
Przykład z Obrazka 17 jest dość pouczający i warto go jeszcze raz obejrzyć. Pozycja na obrazku jest otrzymana w 6 krokach, które można zapisać tak:

$$\mathcal{C}_0 \Rightarrow^{G_1} \mathcal{C}_1 \Rightarrow^{G_1} \mathcal{C}_2 \Rightarrow^{G_2} \mathcal{C}_3 \Rightarrow^1 \mathcal{C}_4 \Rightarrow^2 \mathcal{C}_5 \Rightarrow^3 \mathcal{C}_6 = \mathcal{C},$$

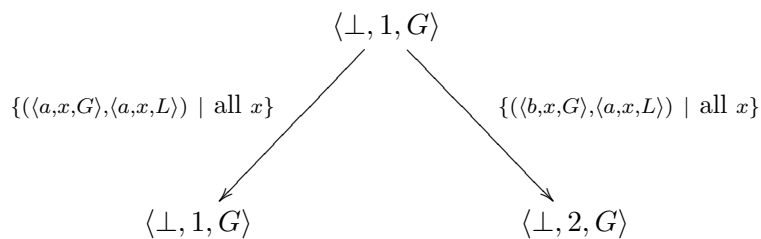
Zauważmy, że każda udana rozgrywka musi przebiegać podobnie ($\mathcal{C}_{2n}$ jest końcowa).

$$\mathcal{C}_0 \Rightarrow^{G_1} \mathcal{C}_1 \Rightarrow^{G_1} \dots \Rightarrow^{G_1} \mathcal{C}_{n-1} \Rightarrow^{G_2} \mathcal{C}_n \Rightarrow^1 \mathcal{C}_{n+1} \Rightarrow^2 \dots \Rightarrow^n \mathcal{C}_{2n}.$$

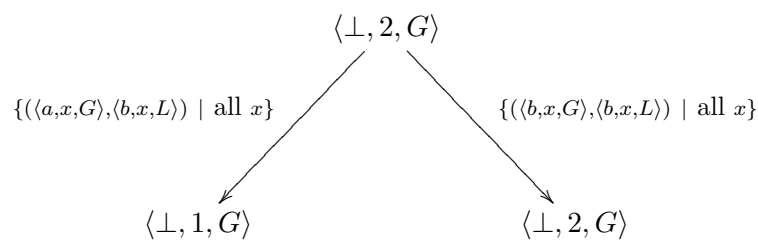
Pierwsza faza gry determinuje drugą. Po wykonaniu G_2 używać musimy kolejno pierwszej, drugiej, trzeciej reguły lokalnej, i tak dalej. Zdefiniujemy teraz nieco bardziej złożoną grę, o podobnym, ale cyklicznym zachowaniu. Poniżej, trójki $\langle\langle X, b \rangle, \langle Y, c \rangle, d \rangle$ reprezentujemy graficznie w taki sposób:



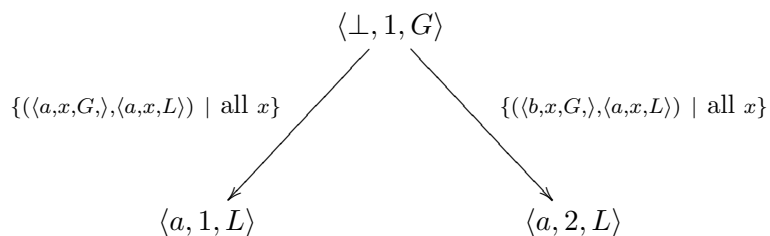
Zaczniemy od „trójwymiarowego” alfabetu $\Sigma_1 = \{\perp, a, b\} \times \{\perp, 1, 2\} \times \{L, G\}$ i określimy grę $T_1 = \langle \langle \perp, 1, G \rangle, \emptyset, \{G_1, G_2, G_3\} \rangle$. (Na razie chcemy tylko grać, nie musimy wygrywać, dlatego nie ma etykiet końcowych.) Reguła globalna G_1 składa się z trójek postaci:



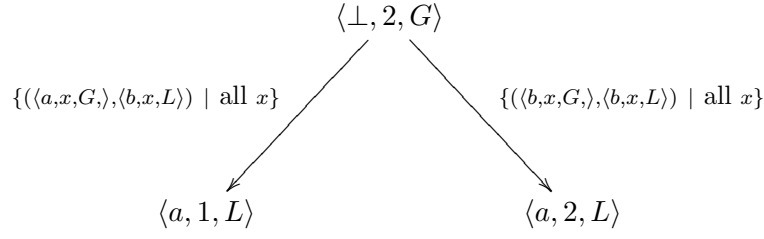
oraz postaci:



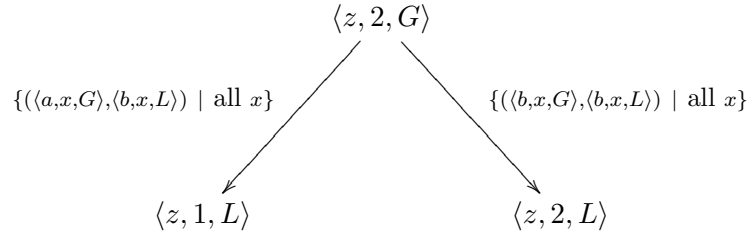
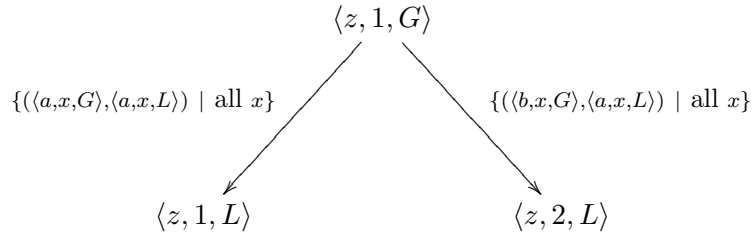
Zauważmy, że G_1 jest bardzo podobna do reguły G_1 z poprzedniego przykładu (por. a i b w pierwszej składowej i 1, 2 w drugiej). Reguła globalna G_2 składa się z dwóch trójek:



oraz:



Dalej, dla dowolnego $z \in \{a, b\}$, tj. $z \neq \perp$, mamy w G_3 trójki postaci:



Lemat 21.4 Każda „rozgrywka” w grze T_1 jest postaci

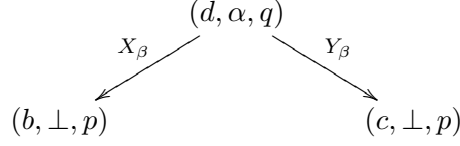
$$\begin{aligned}
 & \mathcal{C}_0 \Rightarrow^{G_1} \mathcal{C}_1 \Rightarrow^{G_1} \dots \Rightarrow^{G_1} \mathcal{C}_i \Rightarrow^{G_1} \mathcal{C}_{i+1} \Rightarrow^{G_1} \mathcal{C}_{i+2} \Rightarrow^{G_1} \dots \Rightarrow^{G_1} \mathcal{C}_{n-1} \Rightarrow^{G_2} \\
 & \mathcal{C}_n \Rightarrow^1 \mathcal{C}_{n+1} \Rightarrow^{G_3} \Rightarrow^2 \dots \Rightarrow^{G_3} \mathcal{C}_{n+2i} \Rightarrow^{i+1} \mathcal{C}_{n+2i+1} \Rightarrow^{G_3} \mathcal{C}_{n+2i+2} \Rightarrow^{i+2} \dots \\
 & \dots \Rightarrow^{G_3} \mathcal{C}_{3n-2} \Rightarrow^n \mathcal{C}_{3n-1} \Rightarrow^{G_3} \mathcal{C}_{3n} \Rightarrow^{n+2} \mathcal{C}_{3n+1} \Rightarrow^{G_3} \mathcal{C}_{3n+2} \Rightarrow^{n+4} \mathcal{C}_{3n+3} \Rightarrow^{G_3} \dots \\
 & \dots \Rightarrow^{G_3} \mathcal{C}_{2i} \Rightarrow^{2i-n+2} \mathcal{C}_{2i+1} \Rightarrow^{G_3} \dots
 \end{aligned}$$

Innymi słowy, zaczynając od pozycji n -tej, kroki lokalne i globalne następują na zmianę, przy czym deklarowane reguły lokalne są używane dokładnie po kolejnych $n-2$ krokach.

Gra z automatu

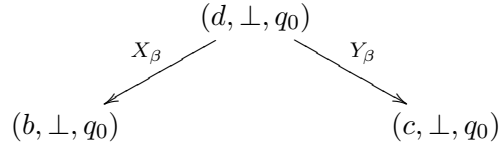
Pokażemy teraz jak zredukować problem niepustości dla automatów wierszowych do problemu gier na drzewach. Niech T_1 będzie grą omawianą powyżej i niech $\mathcal{A} = \langle \Sigma^{\mathcal{A}}, Q, q_0, F^{\mathcal{A}}, \varrho \rangle$ będzie automatem wierszowym. Definiujemy grę $T^{\mathcal{A}} = \langle a, F, \{G_1^{\mathcal{A}}(\beta) \mid \beta \in \Sigma_{\mathcal{A}}\} \cup \{G_2^{\mathcal{A}}, G_3^{\mathcal{A}}\} \rangle$. Jej alfabet jest iloczynem kartezjańskim $\Sigma_1 \times (\Sigma^{\mathcal{A}} \cup \{\perp, \varepsilon\}) \times Q$, a reguły definiujemy tak:

- Dla dowolnej reguły lokalnej X gry T_1 i każdego $\beta \in \Sigma^A \cup \{\varepsilon\}$, definiujemy regułę lokalną $X_\beta = \{(\langle a, \beta, q \rangle, \langle b, \perp, q \rangle) \mid q \in Q \text{ oraz } \langle a, b \rangle \in X\}$.
- Dla każdej trójki $\Delta = \langle \langle X, b \rangle, \langle Y, c \rangle, d \rangle \in G_3$, definiujemy zbiór Δ^A złożony z trójek



gdzie $\varrho(\alpha, q) = (\beta, p)$.

- Jeśli $\Delta = \langle \langle X, b \rangle, \langle Y, c \rangle, d \rangle \in G_1 \cup G_2$, to Δ_β jest trójką



- Dla $\beta \in \Sigma_A$, przyjmujemy $G_1^A(\beta) = \{\Delta_\beta \mid \Delta \in G_1\}$;
- Dalej $G_2^A = \{\Delta_\varepsilon \mid \Delta \in G_2\}$ i $G_3^A = \bigcup \{\Delta^A \mid \Delta \in G_3\}$.

Symbol początkowy gry T^A to $a = \langle a_0, \perp, q_0 \rangle$, gdzie a_0 jest symbolem początkowym w T_1 . Wreszcie $F = \Sigma_1 \times (\Sigma_A \cup \{\perp, \varepsilon\}) \times F^A$, i to cała definicja.

Nasza gra zachowuje się jak „iloczyn kartezjański” gry T_1 (pierwsza współrzędna etykiet) i automatu $\mathcal{A}$. W szczególności lemat 21.4 pozostaje prawdziwy, jeśli G_1, G_2, G_3 zamienimy odpowiednio na $G_1^A(\beta), G_2^A, G_3^A$. Każda rozgrywka wygląda tak:

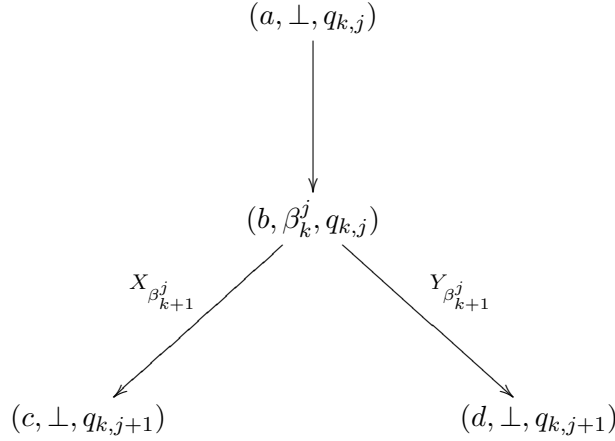
$$\begin{aligned}
 & \mathcal{C}_0 \Rightarrow^{G_1^A(\beta_1)} \mathcal{C}_1 \Rightarrow^{G_1^A(\beta_2)} \dots \Rightarrow^{G_1^A(\beta_{n-1})} \mathcal{C}_{n-1} \Rightarrow^{G_2^A} \\
 & \mathcal{C}_n \Rightarrow^1 \mathcal{C}_{n+1} \Rightarrow^{G_3^A} \Rightarrow^2 \dots \Rightarrow^{G_3^A} \mathcal{C}_{n+2i} \Rightarrow^{i+1} \mathcal{C}_{n+2i+1} \Rightarrow^{G_3^A} \mathcal{C}_{n+2i+2} \Rightarrow^{i+2} \dots \\
 & \dots \Rightarrow^{G_3^A} \mathcal{C}_{3n-2} \Rightarrow^n \mathcal{C}_{3n-1} \Rightarrow^{G_3^A} \mathcal{C}_{3n} \Rightarrow^{n+2} \mathcal{C}_{3n+1} \Rightarrow^{G_3^A} \mathcal{C}_{3n+2} \Rightarrow^{n+4} \mathcal{C}_{3n+3} \Rightarrow^{G_3^A} \dots \\
 & \dots \Rightarrow^{G_3^A} \mathcal{C}_{2i} \Rightarrow^{2i-n+2} \mathcal{C}_{2i+1} \Rightarrow^{G_3^A} \dots
 \end{aligned}$$

Jeśli pozycja $\mathcal{C}$ jest osiągalna z pozycji początkowej, to etykiety wszystkich wierzchołków tego samego poziomu mają zawsze takie same drugie i trzecie współrzędne.

Ciąg pozycji $\mathcal{C}_0, \mathcal{C}_1, \mathcal{C}_2, \dots$ jak wyżej jest całkowicie zdeterminowany przez n oraz $\beta^1, \dots, \beta^{n-1}$. Zauważmy, że β^i , dla $i = 1, \dots, n-1$, są drugimi współrzędnymi etykiet wszystkich liści $\mathcal{C}_{n+2i-1}$. Niech w oznacza słowo $\beta^1 \beta^2 \dots \beta^{n-1}$. Twierdzimy, że

Automat $\mathcal{A}$ akceptuje w wtedy i tylko wtedy, gdy $\mathcal{C}_0, \mathcal{C}_1, \mathcal{C}_2, \dots$ prowadzi do pozycji końcowej.

Aby to udowodnić, założmy, że β_k^j jest symbolem zapisanym w j -tej klatce taśmy automatu, po tym jak maszyna zakończyła $k-1$ fazę pracy (przeczytała całą taśmę $k-1$ razy). Inaczej mówiąc, β_k^j jest symbolem, który ma być czytany w k -tym przejeździe głowicy. Oczywiście $\beta_1^j = \beta^j$. Dla porządku, niech jeszcze $\beta_k^n = \varepsilon$. Dalej przyjmijmy, że $q_{k,j}$ jest stanem maszyny tuż przed przeczytaniem j -tej klatki taśmy po raz k -ty (po $k-1$ pełnych fazach). Pozostawiamy czytelnikowi udowodnienie, że:



Obrazek 18: Symulacja pojedynczego kroku maszyny

1. Drugą współrzędną etykiet wszystkich liści $\mathcal{C}_{(2k-1)n+2j-1}$ jest β_k^j .
2. Trzecią współrzędną etykiet wszystkich liści $\mathcal{C}_{(2k-1)n+2j-2}$ jest $q_{k,j}$.

W szczególności, druga i trzecia współrzędna każdego wierzchołka zależy tylko od jego wysokości w drzewie.²¹ Nasza teza wynika z punktu 2, zastosowanego do stanów końcowych. Na Obrazku 18 widzimy etykiety wierzchołka na poziomie $(2k-1)n+2j-2$ i jego synów.

Jeśli $\varrho(\beta, q) = (\alpha, p)$ to ruch „globalny” zmienia etykiety postaci²² $(\cdot, \beta, q)$ na $(\cdot, \perp, p)$ i przekazuje do następnej fazy informację o nowej literze α poprzez regułę lokalną zawierającą pary postaci $(\cdot, \alpha, q') \rightarrow (\cdot, \perp, q')$. Następujący po nim ruch lokalny służy do ustalenia właściwego symbolu taśmy automatu (drugiej współrzędnej) którą pobiera się z poprzedniej fazy. W kolejnym ruchu globalnym w podobny sposób symulujemy działanie maszyny na następnej klatce taśmy. Dla ilustracji rozpatrzmy sytuację gdy słowem początkowym jest $\beta_1\beta_2\beta_3$, a automat wierszowy przyjmuje kolejno takie konfiguracje:

$$\begin{aligned}
 &(\varepsilon, q_0, \beta_1\beta_2\beta_3) \rightarrow (\gamma_1, q_1, \beta_2\beta_3) \rightarrow (\gamma_1\gamma_2, q_2, \beta_3) \rightarrow (\gamma_1\gamma_2\gamma_3, q_3, \varepsilon) \rightarrow \\
 &\rightarrow (\varepsilon, p_0, \gamma_1\gamma_2\gamma_3) \rightarrow (\delta_1, p_1, \gamma_2\gamma_3) \rightarrow (\delta_1\delta_2, p_2, \gamma_3) \rightarrow (\delta_1\delta_2\delta_3, p_3, \varepsilon) \rightarrow \\
 &\rightarrow (\varepsilon, r_0, \delta_1\delta_2\delta_3) \rightarrow \dots
 \end{aligned}$$

Odpowiada to rozgrywce, w której na kolejnych poziomach drzewa mamy takie etykiety (nad strzałkami zaznaczono litery przekazywane do następnej fazy):

$$\begin{aligned}
 &(\cdot, \perp, q_0) \xrightarrow{\beta_1} (\cdot, \perp, q_0) \xrightarrow{\beta_2} (\cdot, \perp, q_0) \xrightarrow{\beta_3} (\cdot, \perp, q_0) \xrightarrow{\varepsilon} \\
 &(\cdot, \beta_1, q_0) \xrightarrow{\gamma_1} (\cdot, \perp, q_1) \rightarrow (\cdot, \beta_2, q_1) \xrightarrow{\gamma_2} (\cdot, \perp, q_2) \rightarrow (\cdot, \beta_3, q_2) \xrightarrow{\gamma_3} (\cdot, \perp, q_3) \rightarrow (\cdot, \varepsilon, q_3) \xrightarrow{\varepsilon} (\cdot, \perp, p_0) \rightarrow \\
 &(\cdot, \gamma_1, p_0) \xrightarrow{\delta_1} (\cdot, \perp, p_1) \rightarrow (\cdot, \gamma_2, p_1) \xrightarrow{\delta_2} (\cdot, \perp, p_2) \rightarrow (\cdot, \gamma_3, p_2) \xrightarrow{\delta_3} (\cdot, \perp, p_3) \rightarrow (\cdot, \varepsilon, p_3) \xrightarrow{\varepsilon} (\cdot, \perp, r_0) \rightarrow \\
 &(\cdot, \delta_1, r_0) \rightarrow \dots
 \end{aligned}$$

Z powyższych rozważań ostatecznie wynika

Twierdzenie 21.5 *Problem inhabitacji dla typów iloczynowych jest nierozstrzygalny.*

²¹Drzewo jest nam potrzebne tylko do zapewnienia cyklicznego zachowania gry.

²²Ignorujemy pierwszą współrzędną.

22 Logika drugiego rzędu

Przez logikę drugiego rzędu zazwyczaj rozumie się system powstały przez rozszerzenie rachunku predykatów (logiki pierwszego rzędu) o kwantyfikatory przebiegające zbiory i relacje. Z konstruktywnego punktu widzenia, celowość takiego rozszerzenia jest wątpliwa. Interpretacja BHK nadaje bowiem taki sens kwantyfikatorowi ogólnemu:

Konstrukcją (dowodem) formuły $\forall R W(R)$ jest metoda przekształcenia dowolnego znaczenia $\mathbf{R}$ zmiennej R w konstrukcję stwierdzenia $W(\mathbf{R})$.

Aby jednak wykonać konstrukcję dla $W(\mathbf{R})$ trzeba na ogół *znać* sam zbiór $\mathbf{R}$, tj. umieć go zdefiniować, a to nie zawsze jest możliwe, bo każda nieskończona dziedzina ma nieprzeliczalnie wiele podzbiorów. Dlatego konstruktywne rozumienie kwantyfikatora drugiego rzędu jest inne: zmienna relacyjna przebiega predykaty definiowalne formułami. W intuicjonistycznej logice drugiego rzędu przyjmuje się więc (z grubsza²³) takie reguły:

$$\frac{\Gamma \vdash \forall X \varphi(X)}{\Gamma \vdash \varphi(\sigma)} \qquad \frac{\Gamma \vdash \varphi(X)}{\Gamma \vdash \forall X \varphi(X)} \quad (X \notin \text{FV}\Gamma)$$

Uwaga: Systemy klasycznej logiki drugiego rzędu, w których kwantyfikatory przebiegają dowolne zbiory, nie są efektywnie aksjomatyzowalne. Tak nie jest w przypadku intuicjonistycznej logiki drugiego rzędu, która jest aksjomatyzowalna i dlatego rekurencyjnie przeliczalna.

Jak powiedzieliśmy, w logice drugiego rzędu obok kwantyfikatorów relacyjnych występują także elementy zwykłego rachunku predykatów: wyrażenia i kwantyfikatory indywidualowe. Okazuje się jednak, że te ostatnie nie zawsze są istotne, a wiele własności logiki drugiego rzędu można badać w oderwaniu od jej aspektu indywidualowego. Jeżeli w formułach rachunku predykatów drugiego rzędu zaniedbać wszystko to, co odnosi się do indywidualów, to otrzymamy formuły zdaniowe z kwantyfikatorami. Na przykład rozpatrzmy formułę $N(x)$, która definiuje liczby naturalne:

$$\forall R(\forall z(R(z) \rightarrow R(sz)) \rightarrow R(0) \rightarrow R(x))$$

Usunięcie z niej informacji o indywidualach daje kwantyfikowaną formułę zdaniową:

$$\forall R((R \rightarrow R) \rightarrow R \rightarrow R),$$

w której jednoargumentowa zmienna relacyjna R została zastąpiona zmienną zdaniową. W podobny sposób można każdą formułę przekształcić w formułę zdaniową z kwantyfikatorami. Okazuje się, że ten zabieg, znacznie upraszczając składnię, nie zmienia wielu istotnych własności naszej logiki.

Zdaniowa logika drugiego rzędu

Na razie zatem będziemy się zajmować tylko zdaniową logiką drugiego rzędu, i to wyłącznie fragmentem implikacyjno-unwersalnym. Zaczynamy od składni. *Formułami* nazywamy zmienne zdaniowe $p, q, r, \dots$ i wyrażenia $(\sigma \rightarrow \tau)$, $\forall p \sigma$, gdzie σ i τ są formułami. Przez $\text{FVT}(\varphi)$

²³Pomijamy tu ścisłą definicję wyrażenia $\varphi(\sigma)$.

oznaczamy zbiór zmiennych zdaniowych występujących wolno w formule φ . Definiujemy go tak: $FVT(p) = \{p\}$, $FVT(\sigma \rightarrow \tau) = FVT(\sigma) \cup FVT(\tau)$, oraz $FVT(\forall p \sigma) = FVT(\sigma) - \{p\}$. Podstawienie formuły σ do formuły φ w miejsce wolnych wystąpień zmiennej p oznaczamy przez $\varphi[p := \sigma]$ (lub przez $\varphi[\sigma/p]$). Podstawienie jest definiowane przez indukcję:

- $p[p := \sigma] = \sigma$ oraz $q[p := \sigma] = q$;
- $(\psi \rightarrow \tau)[p := \sigma] = \psi[p := \sigma] \rightarrow \tau[p := \sigma]$;
- $(\forall p \psi)[p := \sigma] = \forall p \psi$;
- $(\forall q \psi)[p := \sigma] = \forall q \psi[p := \sigma]$, gdy $q \notin FVT(\sigma)$;

Utożsamiamy (alfa-konwersja) formuły różniące się tylko wyborem zmiennych związanych. Reguły naturalnej dedukcji dla naszej logiki są takie:

$$\begin{array}{c}
 (\text{Ax}) \Gamma, \varphi \vdash \varphi \\
 (\rightarrow \text{I}) \frac{\Gamma, \varphi \vdash \psi}{\Gamma \vdash \varphi \rightarrow \psi} \quad (\rightarrow \text{E}) \frac{\Gamma \vdash \varphi \rightarrow \psi \quad \Gamma \vdash \varphi}{\Gamma \vdash \psi} \\
 (\forall \text{I}) \frac{\Gamma \vdash \varphi}{\Gamma \vdash \forall p \varphi} \quad (p \notin FVT(\Gamma)) \quad (\forall \text{E}) \frac{\Gamma \vdash \forall p \varphi}{\Gamma \vdash \varphi[p := \vartheta]}
 \end{array}$$

Uwaga: Znaczeniem zmiennej zdaniowej może być dowolna formuła, zob. regułę $(\forall \text{E})$. Tę własność czasem nazywa się *full comprehension*. Znaczenie formuły $\varphi = \forall p \psi$ jest wyznaczone przez wszystkie możliwe podstawienia $\psi[p := \sigma]$. Ale ponieważ w roli σ może wystąpić samo φ , więc znaczenie formuły w sposób uwikłany zależy od niej samej. Inaczej mówiąc, nasza logika jest *impredykatywna*. Konsekwencją tej własności jest utrudnione definiowanie i dowodzenie własności przez zwykłą indukcję ze względu na „budowę formuły”.

Twierdzenie 22.1 (Löb, Gabbay/Sobolew) *Intuicjonistyczna logika zdaniowa drugiego rzędu jest nierozstrzygalna (nierozstrzygalne jest pytanie czy dana formuła jest twierdzeniem).*

Powyższy wynik kontrastuje ze znanym twierdzeniem z teorii złożoności: klasyczna logika kwantyfikowanych formuł boole’owskich jest PSPACE-zupełna. Zauważmy jednak, że w logice klasycznej, opartej na zerojedynkowej semantyce, każdy kwantyfikator można wyeliminować (kosztem wykładniczego wydłużenia formuły). Klasyczny rachunek zdań jest *funkcjonalnie zupełny*: każdą funkcję zerojedynkową można zdefiniować formułą. W przypadku intuicjonistycznym tak nie jest. Na przykład formuły $\neg \forall p(p \vee \neg p)$ i $\neg \neg p \rightarrow \exists q((p \rightarrow \neg q \vee q) \rightarrow p)$ nie są równoważne żadnej formule zdaniowej bez kwantyfikatorów.

23 System F

Zgodnie z interpretacją BHK, konstrukcją dla $\forall p \varphi$ jest metoda, która dla dowolnej formuły σ pozwala otrzymać konstrukcję dla $\varphi[p := \sigma]$. Taka konstrukcja jest więc funkcją określoną

na zbiorze formuł, która dla każdego σ przyjmuje wartość ze zbioru konstrukcji $\varphi[p := \sigma]$. Zbiór takich konstrukcji odpowiada (nieformalnie) produktowi $\prod_{p \in \mathbf{Prop}} \varphi[p := \sigma]$. Jeśli więc utożsamiać formuły i typy to formuła uniwersalna jest typem produktowym. Patrząc na to z innej strony, możemy powiedzieć, że term typu $\forall p \varphi$ reprezentuje procedurę polimorficzną z (formalnym) parametrem typowym p . Przekazanie do takiej procedury parametru aktualnego σ ustala typ bieżącej aktywacji jako $\varphi[p := \sigma]$.

System rachunku lambda, w którym typy są formułami zdaniowymi drugiego rzędu nazywamy *polimorficznym rachunkiem lambda*, *rachunkiem lambda drugiego rzędu*, *polimorficznym rachunkiem lambda drugiego rzędu* lub *systemem $\mathbf{F}$* . Poniższe reguły są zarówno regułami przypisania typów, jak też definicją składni. Wyrażenie $\Gamma \vdash M : \sigma$ czytamy „w otoczeniu Γ term M ma typ σ ”. Przez *otoczenie* rozumiemy zbiór deklaracji postaci $(x : \tau)$, gdzie x jest zmienną (przedmiotową) a τ jest typem. Żądamy przy tym, aby otoczenie zawsze było funkcją, tj. aby jedna zmienna nie była deklarowana dwa razy. Piszemy $\Gamma, x:\varphi$ zamiast $\Gamma \cup \{(x : \varphi)\}$ a przez $\text{FVT}(\Gamma)$ oznaczamy sumę $\bigcup \{\text{FVT}(\gamma) \mid (x : \gamma) \in \Gamma, \text{ dla pewnego } x\}$.

$$\begin{array}{c}
(\text{Var}) \quad \Gamma, x : \varphi \vdash x : \varphi \\
\\
(\text{Abs}) \quad \frac{\Gamma, x:\varphi \vdash M : \psi}{\Gamma \vdash (\lambda x:\varphi M) : \varphi \rightarrow \psi} \qquad (\text{App}) \quad \frac{\Gamma \vdash M : \varphi \rightarrow \psi \quad \Gamma \vdash N : \varphi}{\Gamma \vdash (MN) : \psi} \\
\\
(\text{Gen}) \quad \frac{\Gamma \vdash M : \varphi}{\Gamma \vdash (\Lambda p M) : \forall p \varphi} \quad (p \notin \text{FVT}(\Gamma)) \qquad (\text{Inst}) \quad \frac{\Gamma \vdash M : \forall p \varphi}{\Gamma \vdash M\sigma : \varphi[p := \sigma]}
\end{array}$$

Składnię systemu $\mathbf{F}$ zdefiniowaliśmy *w stylu Churcha*, tj. tak, że typy stanowią integralną część termów. Przy ustalonym otoczeniu, deklarującym typy zmiennych wolnych, każdy term ma dokładnie jeden typ, który można bez trudu ustalić.

W termach mogą występować zarówno zmienne przedmiotowe jak typowe (zdaniowe). Operatorami wiążącymi zmienne są obie lambdy i kwantyfikator występujący w typach. Symbol $\text{FVT}(\)$ odnosi się do zmiennych wolnych zdaniowych, a symbol FV do zmiennych wolnych przedmiotowych. Definicje są indukcyjne:

- $\text{FV}(x) = \{x\}$;
- $\text{FV}(\lambda x:\varphi M) = \text{FV}(M) - \{x\}$;
- $\text{FV}(MN) = \text{FV}(M) \cup \text{FV}(N)$;
- $\text{FV}(\Lambda p M) = \text{FV}(M)$;
- $\text{FV}(M\sigma) = \text{FV}(M)$;
- $\text{FVT}(x) = \emptyset$;
- $\text{FVT}(\lambda x:\varphi M) = \text{FVT}(\varphi) \cup \text{FVT}(M)$;
- $\text{FVT}(MN) = \text{FVT}(M) \cup \text{FVT}(N)$;
- $\text{FVT}(\Lambda p M) = \text{FVT}(M) - \{p\}$;

- $\text{FVT}(M\sigma) = \text{FVT}(M) \cup \text{FVT}(\sigma)$.

Operacja podstawienia występuje w dwóch wersjach: podstawienie typu (formuły) za zmienną typową (zdaniową) i podstawienie termu za zmienną przedmiotową. Definiujemy je, zakładając, że zmienne związane są tak dobrane, aby nie doszło do konfuzji, tj. stosujemy w razie potrzeby domyślną wymianę zmiennej związanej.

- $x[x := P] = P$ oraz $y[x := P] = y$;
- $(MN)[x := P] = M[x := P]N[x := P]$;
- $(\lambda x:\varphi M)[x := P] = \lambda x:\varphi M$;
- $(\lambda y:\varphi M)[x := P] = \lambda y:\varphi. M[x := P]$, gdzie $y \notin \text{FV}(P)$;
- $(M\tau)[x := P] = M[x := P]\tau$;
- $(\Lambda p M)[x := P] = \Lambda p M[x := P]$, gdzie $p \notin \text{FVT}(P)$;
- $x[p := \sigma] = x$;
- $(MN)[p := \sigma] = M[p := \sigma]N[p := \sigma]$;
- $(\lambda x:\varphi M)[p := \sigma] = \lambda x:\varphi[p := \sigma]. M[p := \sigma]$;
- $(M\tau)[p := \sigma] = M[p := \sigma]\tau[p := \sigma]$;
- $(\Lambda p M)[p := \sigma] = \Lambda p M$;
- $(\Lambda q M)[p := \sigma] = \Lambda q M[p := \sigma]$, gdzie $q \notin \text{FVT}(\sigma)$.

Przez $\Gamma[p := \sigma]$ oznaczmy otoczenie Γ , w którym każda deklaracja $(x : \tau)$ została zastąpiona przez $(x : \tau[p := \sigma])$.

Lemat 23.1

- (1) Jeśli $\Gamma \vdash M : \varphi$, to $\Gamma[p := \sigma] \vdash M[p := \sigma] : \varphi[p := \sigma]$.
- (2) Jeśli $\Gamma, x : \tau \vdash M : \varphi$ oraz $\Gamma \vdash P : \tau$ to $\Gamma \vdash M[x := P] : \varphi$.

Dowód: Indukcja ze względu na budowę M . Część (2) wymaga lematu podobnego do lematu 13.1. ■

Poniżej mamy kilka przykładów termów i ich typów (zamiast $\lambda x : \tau$ piszemy często λx^τ):

- $\vdash \Lambda p. (\lambda x^{\forall p(p \rightarrow p)}. x(q \rightarrow q)(xq) : \forall q(\forall p(p \rightarrow p) \rightarrow q \rightarrow q))$;
- $\vdash \Lambda p. \lambda f^{p \rightarrow p} \lambda x^p. f(fx) : \forall p((p \rightarrow p) \rightarrow (p \rightarrow p))$;
- $\vdash \lambda f^{\forall p(p \rightarrow q \rightarrow p)} \Lambda p \lambda x^p. f(q \rightarrow p)(fpx) : \forall p(p \rightarrow q \rightarrow p) \rightarrow \forall p(p \rightarrow q \rightarrow q \rightarrow p)$.

Twierdzenie 23.2 Dla dowolnego typu τ następujące warunki są równoważne:

- 1) Istnieje term zamknięty typu τ ;
- 2) Formuła τ jest twierdzeniem intuicjonistycznej logiki zdaniowej drugiego rzędu.

Dowód: Oczywiście. ■

Definicja 23.3 Przez beta redukcję w systemie **F** rozumiemy najmniejszą relację $\rightarrow_\beta$ w zbiorze termów zawierającą:

- $(\lambda x:\tau.M)N \rightarrow_\beta M[x := N]$;
- $(\Lambda\alpha.M)\tau \rightarrow_\beta M[\alpha := \tau]$,

i zamkniętą ze względu na konteksty, tj. $M \rightarrow_\beta N$ implikuje $MP \rightarrow_\beta NP$, $PM \rightarrow_\beta PN$, $\lambda x:\tau.M \rightarrow_\beta \lambda x:\tau.N$, $\Lambda p M \rightarrow_\beta \Lambda p N$ oraz $M\tau \rightarrow_\beta N\tau$, o ile odpowiednie wyrażenia są poprawnymi termami.

Jak zwykle w takich przypadkach, symbol $\rightarrow_\beta$ oznacza domknięcie przechodnio-zwrotne relacji redukcji $\rightarrow_\beta$. Indeks β często jest pomijany. Najważniejsze własności redukcji są takie:

Twierdzenie 23.4

- 1) Jeśli $\Gamma \vdash M : \tau$ oraz $M \rightarrow_\beta N$ to $\Gamma \vdash N : \tau$.
- 2) Własność Churcha-Rossera: Jeżeli $M \rightarrow_\beta N_1$ oraz $M \rightarrow_\beta N_2$ to $N_1 \rightarrow_\beta P$ i $N_2 \rightarrow_\beta P$, dla pewnego P .
- 3) Silna normalizacja: Nie istnieje nieskończony ciąg redukcji rozpoczynający się od jakiegokolwiek (poprawnego) termu.

Dowód: Część (1) wynika z lematu 23.1. Części (2) i (3) zostawiamy na razie bez dowodu. ■

Uwaga: Dla termów systemu **F** można też rozważać relację eta-redukcji, zadaną regułami:

- $\lambda x:\tau.Mx \rightarrow_\beta M$ (gdy $x \notin \text{FV}(M)$);
- $\Lambda\alpha.M\tau \rightarrow_\beta M$ (gdy $\alpha \notin \text{FVT}(M)$).

Liczby naturalne

Niech $\omega = \forall p((p \rightarrow p) \rightarrow (p \rightarrow p))$. Liczby naturalne są reprezentowane przez termy zamknięte typu ω .

$$\mathbf{n} = \Lambda p \lambda f : (p \rightarrow p) \lambda x : p.f(f(\cdots f(x)))$$

Na przykład $\mathbf{2} = \Lambda p.\lambda f^{p \rightarrow p} x^p.f(fx)$.

Mówimy, że funkcja $f : \mathbb{N}^k \rightarrow \mathbb{N}$ jest *definiowalna* w systemie **F**, gdy istnieje taki term $M : \omega^k \rightarrow \omega$, że dla dowolnych $n_1, \dots, n_k$ zachodzi $M\mathbf{n}_1 \dots \mathbf{n}_k \twoheadrightarrow_\beta \mathbf{f}(\mathbf{n}_1, \dots, \mathbf{n}_k)$.

Na przykład dodawanie, mnożenie i potęgowanie definiujemy tak:

- $\mathbf{Add} = \lambda mn. \Lambda p. \lambda fx. mpf(npfx)$;
- $\mathbf{Mult} = \lambda mn. \Lambda p. \lambda fx. mp(npfx)$;
- $\mathbf{Exp} = \lambda mn. \Lambda p. \lambda fx. m(p \rightarrow p)(np)$.

Spójniki zdaniowe

W logice zdaniowej drugiego rzędu fałsz można zdefiniować tak:

$$\perp = \forall p p$$

Nietrudno zauważyć, że z fałszu wszystko wynika. Jeśli $M : \perp$, to $M\tau : \tau$.

Można też zdefiniować pozostałe spójniki zdaniowe. Przypomnijmy odpowiednie reguły naturalnej dedukcji wraz z przypisaniem termów:

$$\mathbf{Koniunkcja:} \quad \frac{\Gamma \vdash M : \varphi \quad \Gamma \vdash N : \psi}{\Gamma \vdash \langle M, N \rangle : \varphi \wedge \psi} (\mathbf{I}\wedge) \quad \frac{\Gamma \vdash M : \varphi \wedge \psi}{\Gamma \vdash \pi_1(M) : \varphi} (\mathbf{E}\wedge) \quad \frac{\Gamma \vdash M : \varphi \wedge \psi}{\Gamma \vdash \pi_2(M) : \psi}$$

$$\mathbf{Alternatywa:} \quad \frac{\Gamma \vdash M : \varphi}{\Gamma \vdash \mathbf{inl}_{\varphi \vee \psi}(M) : \varphi \vee \psi} (\mathbf{I}\vee) \quad \frac{\Gamma \vdash M : \psi}{\Gamma \vdash \mathbf{inr}_{\varphi \vee \psi}(M) : \varphi \vee \psi}$$

$$\frac{\Gamma \vdash M : \varphi \vee \psi \quad \Gamma, x : \varphi \vdash P : \vartheta \quad \Gamma, y : \psi \vdash Q : \vartheta}{\Gamma \vdash \mathbf{case} M \mathbf{of} [x]P, [y]Q : \vartheta} (\mathbf{E}\vee)$$

Redukcje związane z koniunkcją i alternatywą są następujące:

$$\begin{aligned} \pi_1(\langle M, N \rangle) &\rightarrow M, \\ \pi_2(\langle M, N \rangle) &\rightarrow N; \\ \mathbf{case} \mathbf{inl}(P) \mathbf{of} [x]M, [y]N &\rightarrow M[x := P]; \\ \mathbf{case} \mathbf{inr}(Q) \mathbf{of} [x]M, [y]N &\rightarrow N[y := Q]. \end{aligned}$$

Definicję alternatywy w systemie **F** otrzymamy, jeśli spróbujemy wyrazić w logice drugiego rzędu pojęcie sumy jako najmniejszego zbioru zawierającego oba składniki:

$$x \in A \cup B \Leftrightarrow \forall P (A \subseteq P \rightarrow B \subseteq P \rightarrow x \in P),$$

lub inaczej:

$$(A \cup B)(x) \Leftrightarrow \forall P (\forall y (A(y) \rightarrow P(y)) \rightarrow \forall y (B(y) \rightarrow P(y)) \rightarrow P(x)).$$

Po usunięciu informacji o indywiduach otrzymujemy definicję alternatywy:

$$\sigma \vee \tau = \forall p((\sigma \rightarrow p) \rightarrow (\tau \rightarrow p) \rightarrow p),$$

gdzie zmienna p jest nowa, tj. $p \notin \text{FVT}(\sigma) \cup \text{FVT}(\tau)$. Pozostaje sprawdzić, że nasza definicja jest dobra. Wystarczy w tym celu zdefiniować **inl**, **inr** i **case** w taki sposób, żeby reguły (IV) i (EV) były poprawne. Robimy to tak:

- **inl** _{$\sigma \vee \tau$} $M = \Lambda p \lambda u^{\sigma \rightarrow p} \lambda v^{\tau \rightarrow p}.uM$;
- **inr** _{$\sigma \vee \tau$} $M = \Lambda p \lambda u^{\sigma \rightarrow p} \lambda v^{\tau \rightarrow p}.vM$;
- **case** M **of** $[x]P^\vartheta, [y]Q^\vartheta = M\vartheta(\lambda x^\sigma.P)(\lambda y^\tau.Q)$.

Nietrudno sprawdzić, że typowanie jest poprawne, co w szczególności oznacza, że nasza definicja spełnia reguły wnioskowania dla alternatywy. W istocie spełniony jest silniejszy warunek: także redukcje są zgodne z oczekiwaniami, tyle tylko, że może wymagać więcej niż jednego kroku:

$$\begin{aligned} \text{case } \mathbf{inl}(P) \text{ of } [x]M, [y]N &\rightarrow M[x := P]; \\ \text{case } \mathbf{inr}(Q) \text{ of } [x]M, [y]N &\rightarrow N[y := Q]. \end{aligned}$$

Koniunkcję definiujemy tak:

$$\sigma \wedge \tau = \forall p((\sigma \rightarrow \tau \rightarrow p) \rightarrow p),$$

gdzie zmienna p jest nowa (patrz wyżej), a parę i rzuty tak:

- $\langle M, N \rangle = \Lambda p \lambda y^{\sigma \rightarrow \tau \rightarrow p}.yMN$;
- $\pi_1(M) = M\sigma(\lambda x^\sigma \lambda y^\tau.x)$;
- $\pi_2(M) = M\tau(\lambda x^\sigma \lambda y^\tau.y)$.

Łatwo widzieć, że ta definicja spełnia zarówno reguły typowania jak i redukcji.

Kwantyfikator szczegółowy

Kwantyfikator szczegółowy drugiego rzędu ma takie reguły naturalnej dedukcji:

$$\begin{aligned} (\exists I) \quad & \frac{\Gamma \vdash \sigma[p := \tau]}{\Gamma \vdash \exists p. \sigma} \\ (\exists E) \quad & \frac{\Gamma \vdash \exists p. \sigma \quad \Gamma \vdash \rho \quad (p \notin \text{FV}(\Gamma, \rho))}{\Gamma, \sigma \vdash \rho} \end{aligned}$$

Interpretujemy go tak: typ $\exists p. \sigma$ jest typem abstrakcyjnym, w którym zmienna p reprezentuje typ prywatny. Na przykład, jeśli

$$\text{stos}(p) = p \times (p \rightarrow p) \times (p \rightarrow \omega) \times (\omega \times p \rightarrow p),$$

to typ $\exists p \text{stos}(p)$ jest abstrakcyjnym typem struktury stosowej, w której mamy stos pusty jako stałą oraz operacje *pop*, *top* i *push*.)

Operacja **pack** tworzy obiekt typu abstrakcyjnego ukrywając prywatną część typu:

$$(\text{pack}) \frac{\Gamma \vdash M : \sigma[p := \tau]}{\Gamma \vdash \mathbf{pack} \ M, \tau \ \mathbf{to} \ \exists p. \sigma : \exists p. \sigma}$$

Taki obiekt może być użyty tam, gdzie implementacja typu prywatnego nie ma znaczenia:

$$(\text{unpack}) \frac{\Gamma \vdash M : \exists p. \sigma \quad \Gamma, x : \sigma \vdash N : \rho \quad (p \notin FV(\Gamma, \rho))}{\Gamma \vdash \mathbf{let} \ M \ \mathbf{be} \ x, p \ \mathbf{in} \ N : \rho}$$

Następująca reguła redukcji mówi o tym, co się wtedy dzieje:

$$\mathbf{let} \ (\mathbf{pack} \ M, \tau \ \mathbf{to} \ \exists p. \sigma) \ \mathbf{be} \ x, p \ \mathbf{in} \ N \longrightarrow_{\beta} N[p := \tau][x := M].$$

Definicja kwantyfikatora szczegółowego w systemie **F** przypomina definicję alternatywy, i może być objaśniona w podobny sposób. Mamy tu „wytarcie” definicji sumy uogólnionej wyrażonej w języku drugiego rzędu. Taka suma to najmniejszy zbiór zawierający wszystkie składniki. Poniżej, zmienna q musi być nowa, tj. $q \notin FVT(\sigma) \cup \{p\}$.

$$\exists p \sigma = \forall q (\forall p (\sigma \rightarrow q) \rightarrow q).$$

Zauważmy, że powyższa definicja jest wzmocnieniem definicji klasycznej przez prawo de Morgana:

$$\neg \forall p \neg \sigma = \forall p (\sigma \rightarrow \perp) \rightarrow \perp$$

Teraz możemy zdefiniować **pack** i **let**:

- $\mathbf{pack} \ M, \tau \ \mathbf{to} \ \exists p. \sigma = \Lambda q \lambda z^{\forall p (\sigma \rightarrow q)}. z \tau M;$
- $\mathbf{let} \ M \ \mathbf{be} \ x, p \ \mathbf{in} \ N^{\rho} = M \rho (\Lambda p \lambda x^{\sigma}. N).$

i sprawdzić, że wszystko działa jak trzeba.

24 System F w stylu Curry’ego

Wariant systemu **F** w stylu Curry’ego określony jest przez następujące reguły przypisania typów dla termów „czystego” rachunku lambda. (Typy i otoczenia typowe określone są tak samo jak dla wersji Churcha.)

$$\begin{array}{c} \Gamma, x:\tau \vdash x : \tau \\ \\ (\text{APP}) \frac{\Gamma \vdash N : \tau \rightarrow \sigma \quad \Gamma \vdash M : \tau}{\Gamma \vdash NM : \sigma} \qquad \qquad \frac{\Gamma, x : \tau \vdash M : \sigma}{\Gamma \vdash \lambda x. M : \tau \rightarrow \sigma} \ (\text{ABS}) \end{array}$$

$$(\text{GEN}) \frac{\Gamma \vdash M : \sigma}{\Gamma \vdash M : \forall p \sigma} \quad (p \notin FV(\Gamma)) \qquad \frac{\Gamma \vdash M : \forall p \sigma}{\Gamma \vdash M : \sigma[p := \tau]} \quad (\text{INST})$$

Reguły dla abstrakcji i aplikacji są takie same jak dla rachunku z typami prostymi w stylu Curry’ego. Reguły wprowadzania i eliminacji kwantyfikatora ogólnego (zwane odpowiednio regułami *generalizacji* and *instancjacji*) reprezentują ideę *polimorfizmu implicite*: obiekt o typie uniwersalnym $\forall p \tau$ ma wszystkie typy postaci $\tau[p := \sigma]$.

Definicja 24.1 Operacja *wycierania typów* jest określona równaniami:

$$\begin{aligned} |x| &= x \\ |MN| &= |M||N| \\ |\lambda x:\sigma. M| &= \lambda x. |M| \\ |\Lambda p. M| &= |M| \\ |M\tau| &= |M| \end{aligned}$$

Fakt 24.2 Asercja $\Gamma \vdash M : \tau$ jest wyprowadzalna w stylu Curry’ego wtedy i tylko wtedy, gdy istnieje taki term M_0 w stylu Churcha, że $|M_0| = M$, oraz $\Gamma \vdash M_0 : \tau$.

Dowód: Łatwy. ■

O termie N w stylu Churcha można myśleć jak o wyprowadzeniu typu dla termu $|N|$. Następujące twierdzenie w istocie stwierdza, że każda redukcja w $|N|$ odpowiada pewnej redukcji w N .

Twierdzenie 24.3 Jeśli $\Gamma \vdash M : \tau$ w stylu Curry’ego, oraz $M \rightarrow_\beta M'$ to $\Gamma \vdash M' : \tau$.

Zanim zajmiemy się dowodem twierdzenia 24.3, zauważmy, że sprawa wcale nie jest taka oczywista.

Przykład 24.4 Rozpatrzmy następujące przypisanie typu:

$$x : p \rightarrow \forall q (q \rightarrow q) \vdash \lambda y. xy : p \rightarrow q \rightarrow q.$$

Chociaż mamy eta-redukcję $\lambda y. xy \rightarrow_\eta x$, to jednak:

$$x : p \rightarrow \forall q (q \rightarrow q) \not\vdash x : p \rightarrow q \rightarrow q.$$

Zauważmy, że termowi $\lambda y. xy$ odpowiada w naszym przykładzie term Churcha $\lambda y:p. xyq$, który nie jest eta-redeksem. A więc przyczyną tego, że system **F** w stylu Curry’ego nie jest zamknięty ze względu na η -redukcje jest to, że wycieranie typów odsłania „nowe” redeksy. John Mitchell pokazał, że określając odpowiednią relację $\sqsubseteq$ i rozszerzając system o regułę

$$\frac{\Gamma \vdash M : \tau, \quad \tau \sqsubseteq \sigma}{\Gamma \vdash M : \sigma},$$

otrzymuje się rachunek zamknięty ze względu na η -redukcje.

Dowód twierdzenia 24.3

Na rozgrzewkę mamy takie proste obserwacje:

Lemat 24.5 *Dla termów w stylu Churcha:*

1. $|M[p := \tau]| = |M|$ oraz $|M[x := N]| = |M|[x := |N|]$;
2. *Jeśli $M \rightarrow_\beta N$ to $|M| \rightarrow_\beta |N|$.*

Dowód twierdzenia 24.3 wymaga oczywiście odpowiedniego lematu o generowaniu. Aby taki lemat sformułować potrzebujemy pewnej pomocniczej relacji na typach. Piszemy $\sigma \preceq \tau$, gdy typ σ jest postaci $\forall \vec{p} \sigma'$, gdzie $\forall \vec{p}$ jest ciągiem kwantyfikatorów uniwersalnych, typ σ' nie zaczyna się od kwantyfikatora, oraz zachodzi równość $\tau = \forall \vec{q}. \sigma'[\vec{p} := \vec{\rho}]$, dla pewnych typów $\vec{\rho}$, i pewnych zmiennych $\vec{q}$, nie występujących wolno w σ . Wtedy $\Gamma \vdash M : \sigma$ implikuje $\Gamma \vdash M : \tau$ i wystarczy do tego reguły (GEN) i (INST). Mamy też nieco słabsze twierdzenie odwrotne:

Lemat 24.6 *Jeśli asercję $\Gamma \vdash M : \tau$ otrzymano z $\Gamma \vdash M : \sigma$ za pomocą ciągu zastosowań reguł (GEN) i (INST), to $\forall \vec{p} \sigma \preceq \forall \vec{q} \tau$, gdzie $\vec{p} = \text{FVT}(\sigma) - \text{FVT}(\Gamma)$ oraz $\vec{q} = \text{FVT}(\tau) - \text{FVT}(\Gamma)$.*

Dowód: Indukcja ze względu na liczbę kroków. ■

Lemat 24.7 *Jeśli $\Gamma \vdash M : \tau$ to $\Gamma[p := \rho] \vdash M : \tau[p := \rho]$.*

Dowód: Indukcja ze względu na rozmiar wyprowadzenia $\Gamma \vdash M : \tau$ ■

Lemat 24.8 (o generowaniu)

1. *Jeśli $\Gamma \vdash x : \tau$ to $\Gamma(x) \preceq \tau$.*
2. *Jeśli $\Gamma \vdash MN : \tau$ to $\Gamma \vdash M : \varrho \rightarrow \sigma$ oraz $\Gamma \vdash N : \varrho$ dla pewnych ϱ i σ , takich że $\forall \vec{p} \sigma \preceq \tau$, gdzie $\vec{p} = \text{FVT}(\sigma) - \text{FVT}(\Gamma)$.*
3. *Jeśli $\Gamma \vdash \lambda x M : \tau$, oraz $x \notin \text{Dom}(\Gamma)$, to $\tau = \forall \vec{p}(\varrho \rightarrow \sigma)$, gdzie $\Gamma, x:\varrho \vdash M : \sigma$ a zmienne $\vec{p}$ nie są wolne w Γ .*

Dowód: Wyprowadzenie typu dla zmiennej zaczyna się zawsze od aksjomatu $\Gamma \vdash x : \Gamma(x)$, po którym następuje ciąg zastosowań „stacjonarnych” reguł (GEN) i (INST). Teza wynika więc bezpośrednio z lematu 24.6. Podobnie jest dla aplikacji. W przypadku abstrakcji trzeba dodatkowo skorzystać z lematu 24.7. ■

Poniższe reguły tworzą system przypisania typów, który jest „syntaktycznie zorientowany” w takim sensie: ostatnia reguła użyta w wyprowadzeniu typu dla termu M jest wyznaczona przez postać termu M .

$$\begin{array}{ll}
(\text{VAR}') & \Gamma \vdash x : \tau \quad \text{jeśli } (x : \sigma) \text{ is in } \Gamma \text{ i } \sigma \preceq \tau \\
(\text{APP}') & \frac{\Gamma \vdash M : \tau \rightarrow \rho, \quad \Gamma \vdash N : \tau}{\Gamma \vdash (MN) : \rho} \quad \text{jeśli } \forall \vec{p} \rho \preceq \sigma, \text{ gdzie } \vec{p} = \text{FVT}(\rho) - \text{FVT}(\Gamma) \\
(\text{ABS}') & \frac{\Gamma, x:\tau \vdash M : \sigma}{\Gamma \vdash \lambda x.M : \forall \vec{p}(\tau \rightarrow \sigma)} \quad \text{jeśli } \vec{p} \cap \text{FVT}(\Gamma) = \emptyset
\end{array}$$

Konsekwencją lematu o generowaniu jest:

Lemat 24.9 *Asercja $\Gamma \vdash M : \sigma$ jest wyprowadzalna w $\mathbf{F}$ wtedy i tylko wtedy, gdy jest wyprowadzalna w systemie jak wyżej.*

Dowód twierdzenia: Przypuśćmy teraz, że $\Gamma \vdash (\lambda x.M)N : \sigma$. Na mocy lematu 24.9 mamy $\Gamma, x:\tau \vdash M : \rho$ oraz $\Gamma \vdash N : \tau$, dla pewnych τ i ρ . Na dodatek $\forall \vec{p} \rho \preceq \sigma$, gdzie $\vec{p} = \text{FVT}(\rho) - \text{FVT}(\Gamma)$. Stąd nietrudno wywnioskować $\Gamma \vdash M[x := N] : \rho$ a to z kolei implikuje $\Gamma \vdash M[x := N] : \sigma$.

Termy typowalne i nie

Używając polimorfizmu, można przypisywać typy rozmaitym termom, które nie są typowalne w typach prostych. Na przykład termy $\lambda x.xx$ i $\mathbf{2K}$ są teraz typowalne. Jak się zaraz okaże, termy typowalne mają własność SN, ale nie na odwrót: kontrprzykładem jest silnie normalizowalny term

$$(\lambda zy. y(z\mathbf{I})(z\mathbf{K}))(\lambda x.xx),$$

który nie jest typowalny w systemie $\mathbf{F}$. Przyczyną jest niemożność znalezienia takiego typu dla $(\lambda x.xx)$, który „pasowałby” zarówno do argumentu $\mathbf{I}$ jak i $\mathbf{K}$. Innym przykładem jest $\mathbf{22K}$. Zauważmy jednak, że term $\mathbf{2(2K)}$ jest typowalny.

Niestety, zachodzi następujące twierdzenie:

Twierdzenie 24.10 (Wells) *Problem typowalności²⁴ i problem sprawdzania typu²⁵ są dla systemu $\mathbf{F}$ nierozstrzygalne.*

25 Silna normalizacja

Naiwna próba uogólnienia metody Taita na typy polimorficzne polega na zdefiniowaniu:

$$\llbracket \forall p \tau \rrbracket = \bigcap \{ \llbracket \tau[p := \sigma] \rrbracket \mid \sigma \text{ — dowolny typ} \}.$$

Ale ta definicja nie jest dobrze ufundowana: nie da się zdefiniować $\llbracket \forall p \tau \rrbracket$ bez wcześniejszego zdefiniowania np. $\llbracket \tau[p := \forall p \tau] \rrbracket$, co z kolei wymaga definicji $\llbracket \forall p \tau \rrbracket$. Pomysł Girarda, który

²⁴Dany term M , czy istnieją takie Γ i τ , że $\Gamma \vdash M : \tau$?

²⁵Dane są Γ , M i τ . Czy $\Gamma \vdash M : \tau$?

wykorzystamy, zwany „metodą kandydatów”, polega na tym, aby nie definiować zbiorów $\llbracket \tau \rrbracket$ w sposób „absolutny”, w szczególności nie przyjmować z góry $\llbracket p \rrbracket := \text{SN}$ dla zmiennych. Zamiast tego należy określić warunki jakie zbiory $\llbracket \tau \rrbracket$ powinny spełniać i rozważać dowolne wybory zbiorów $\llbracket p \rrbracket$, jak gdyby to były „wartościowania” zmiennych typowych.

Powiemy więc, że zbiór termów X jest *kandydatem*²⁶ (albo, że jest *rodziną normalną*), jeżeli X ma trzy własności z lematów 13.5–13.7:

- 1) $X \subseteq \text{SN}$.
- 2) Jeśli $N_1, \dots, N_k \in \text{SN}$ to $xN_1 \dots N_k \in X$.
- 3) Jeśli $M[x := N_0]N_1 \dots N_k \in X$, oraz $N_0 \in \text{SN}$ to $(\lambda x.M)N_0N_1 \dots N_k \in X$.

Niech $\mathcal{G}$ będzie rodziną wszystkich kandydatów. Nietrudno zauważyć, że $\text{SN} \in \mathcal{G}$ a zatem $\mathcal{G} \neq \emptyset$. *Wartościowaniem* w $\mathcal{G}$ nazwiemy dowolną funkcję $\xi : U \rightarrow \mathcal{G}$. Teraz dla dowolnego wartościowania ξ i dowolnego typu τ definiujemy zbiór $\llbracket \tau \rrbracket_\xi$:

$$\begin{aligned} \llbracket p \rrbracket_\xi &= \xi(p); \\ \llbracket \sigma \rightarrow \tau \rrbracket_\xi &= \{M \mid \forall N (N \in \llbracket \sigma \rrbracket_\xi \Rightarrow MN \in \llbracket \tau \rrbracket_\xi)\}; \\ \llbracket \forall p.\sigma \rrbracket_\xi &= \bigcap_{X \in \mathcal{G}} \llbracket \sigma \rrbracket_{\xi(p \mapsto X)}. \end{aligned}$$

Lemat 25.1 *Dla dowolnych ξ i σ zachodzi $\llbracket \sigma \rrbracket_\xi \in \mathcal{G}$.*

Dowód: Dowód jest przez indukcję ze względu na τ . Jeśli τ jest zmienną to $\llbracket \tau \rrbracket_\xi = \text{SN} \in \mathcal{G}$ i dobrze. Przypadek $\tau = \forall p \sigma$ wynika z prostej obserwacji: iloczyn rodziny kandydatów jest kandydatem. Niech więc $\tau = \sigma \rightarrow \rho$ i niech $M \in \llbracket \sigma \rightarrow \rho \rrbracket_\xi$. Weźmy zmienną $x : \sigma$. Z założenia indukcyjnego (2) mamy $x \in \llbracket \sigma \rrbracket_\xi$, więc z definicji $Mx \in \llbracket \rho \rrbracket_\xi$. A więc $Mx \in \text{SN}$, z założenia indukcyjnego (1). Tym bardziej $Mx \in \text{SN}$. Pokazaliśmy (1).

W punkcie (2) mamy pokazać, że $xN_1 \dots N_k P \in \llbracket \rho \rrbracket_\xi$, dla każdego $P \in \llbracket \sigma \rrbracket_\xi$. Ale $P \in \text{SN}$ z założenia indukcyjnego (1), więc teza wynika z założenia indukcyjnego (2) dla ρ .

Zostaje warunek (3). Niech więc $M[x := N_0]N_1 \dots N_k \in \llbracket \sigma \rightarrow \rho \rrbracket_\xi$, oraz $N_0 \in \text{SN}$. Mamy pokazać, że dla $P \in \llbracket \sigma \rrbracket_\xi$ musi zachodzić $(\lambda x.M)N_0N_1 \dots N_k P \in \llbracket \rho \rrbracket_\xi$. Ale to wynika z założenia indukcyjnego (3) dla ρ . ■

Lemat 25.2

- $\llbracket \sigma[p := \tau] \rrbracket_\xi = \llbracket \sigma \rrbracket_{\xi(p \mapsto (\llbracket \tau \rrbracket_\xi))}$.
- $\llbracket \sigma \rrbracket_\xi = \llbracket \sigma \rrbracket_{\xi'}$, gdy $\xi(p) = \xi'(p)$ dla $p \in \text{FVT}(\sigma)$.

Dowód: Indukcja ze względu na σ . ■

²⁶Oryginalna nazwa: *candidat de reductibilité*

Lemat 25.3 *Niech $\Gamma \vdash M : \tau$ oraz $\text{FV}(M) = \{x_1, \dots, x_k\}$ i przy tym $N_i \in \llbracket \Gamma(x_i) \rrbracket_\xi$, dla $i = 1, \dots, k$. Wtedy $M[x_1 := N_1, \dots, x_k := N_k] \in \llbracket \tau \rrbracket_\xi$.*

Dowód: Dowód jest przez indukcję ze względu na wyprowadzenie $\Gamma \vdash M : \tau$, przez przypadki w zależności od ostatniej użytej reguły.

(VAR) Przypuśćmy, że $\Gamma \vdash x_i : \Gamma(x_i)$. Wtedy term $M[x_1 := N_1, \dots, x_k := N_k]$ to po prostu N_i , więc teza wynika wprost z założenia.

(ABS) Przypuśćmy, że $\Gamma \vdash \lambda x.M : \sigma \rightarrow \rho$ otrzymano z przesłanki $\Gamma, x : \sigma \vdash M : \rho$. Niech $P \in \llbracket \sigma \rrbracket_\xi$. Mamy pokazać, że $((\lambda x.M)[x_1 := N_1, \dots, x_k := N_k])P \in \llbracket \rho \rrbracket_\xi$. Ponieważ możemy założyć, że x nie jest wolne w termach N_i , więc mamy

$$(\lambda x.M)[x_1 := N_1, \dots, x_k := N_k] = \lambda x.M[x_1 := N_1, \dots, x_k := N_k].$$

Z założenia indukcyjnego

$$M[x_1 := N_1, \dots, x_k := N_k][x := P] = M[x_1 := N_1, \dots, x_k := N_k, x := P] \in \llbracket \rho \rrbracket_\xi,$$

więc teza wynika z warunku (3) definicji kandydata.

(APP) Przypuśćmy, że $\Gamma \vdash MN : \tau$ otrzymano z przesłanek $\Gamma \vdash M : \sigma \rightarrow \tau$ i $\Gamma \vdash N : \sigma$. Z założenia indukcyjnego otrzymujemy, że $M[x_1 := N_1, \dots, x_k := N_k] \in \llbracket \sigma \rightarrow \tau \rrbracket_\xi$ a także $N[x_1 := N_1, \dots, x_k := N_k] \in \llbracket \sigma \rrbracket_\xi$. Stąd $(MN)[x_1 := N_1, \dots, x_k := N_k] \in \llbracket \tau \rrbracket_\xi$.

(GEN) Przypuśćmy, że $\Gamma \vdash M : \forall p\tau$ otrzymano z przesłanki $\Gamma \vdash M : \tau$, i przy tym p nie należy do $\text{FVT}(\Gamma)$. Niech $M' = M[x_1 := N_1, \dots, x_k := N_k]$, gdzie $N_i \in \llbracket \Gamma(x_i) \rrbracket_\xi$, dla $i = 1, \dots, k$. Z założenia indukcyjnego $M' \in \llbracket \tau \rrbracket_\nu$ przy dowolnym ν , spełniającym warunek $N_i \in \llbracket \Gamma(x_i) \rrbracket_\nu$, dla $i = 1, \dots, k$. Skoro $p \notin \text{FVT}(\Gamma(x_i))$ to $\llbracket \Gamma(x_i) \rrbracket_\xi = \llbracket \Gamma(x_i) \rrbracket_{\xi(p \mapsto X)}$, dla dowolnych i, X . Zatem $M' \in \llbracket \tau \rrbracket_{\xi(p \mapsto X)}$ dla dowolnego $X \in \mathcal{G}$, a więc $M' \in \llbracket \forall p\tau \rrbracket_\xi$.

(INST) Przypuśćmy, że $\Gamma \vdash M : \sigma[p := \tau]$ otrzymano z $\Gamma \vdash M : \forall p\sigma$. Z założenia indukcyjnego $M[x_1 := N_1, \dots, x_k := N_k] \in \llbracket \forall p\sigma \rrbracket_\xi$, a więc także $M[x_1 := N_1, \dots, x_k := N_k] \in \llbracket \sigma \rrbracket_{\xi(p \mapsto \llbracket \tau \rrbracket_\xi)} = \llbracket \sigma[p := \tau] \rrbracket_\xi$ i już. ■

Twierdzenie 25.4 *System $\mathbf{F}$ w wersji Curry'ego ma własność silnej normalizacji.*

Dowód: Jeśli $\Gamma \vdash M : \tau$ dla pewnych Γ i τ to z poprzedniego lematu wynika, że dla dowolnego ξ mamy $M \in \llbracket \tau \rrbracket_\xi$, bo $x \in \llbracket \Gamma(x) \rrbracket_\xi$ dla dowolnej zmiennej x . Teza wynika więc stąd, że $\llbracket \tau \rrbracket_\xi \subseteq \text{SN}$. ■

Twierdzenie 25.5 *System $\mathbf{F}$ w wersji Churcha ma własność silnej normalizacji.*

Dowód: Przypuśćmy, że mamy nieskończony ciąg redukcji

$$M = M_0 \rightarrow M_1 \rightarrow M_2 \rightarrow \dots$$

Nietrudno zauważyć, że wtedy

$$|M| = |M_0| \succ |M_1| \succ |M_2| \succ \dots$$

gdzie $|M_i|$ są jak w definicji 24.1, a symbol $\succ$ oznacza zawsze $\rightarrow$ lub $=$. Ponieważ termy typowalne w systemie $\mathbf{F}$ w wersji Curry’ego mają własność SN (twierdzenie 25.4) więc poczynając od pewnego miejsca, termy $|M_n|$ muszą być identyczne. To znaczy, że prawie wszystkie redukcje w naszym ciągu są „typowe”, tj. postaci $(\lambda p.N)\tau \rightarrow_q N[p := \tau]$. Skoro każdorazowo znika jedno wystąpienie symbolu Λ , to nasz ciąg nie może być nieskończony. ■

Twierdzenie 25.6 *System $\mathbf{F}$ ma własność Churcha-Rossera.*²⁷

Dowód: Własność Churcha-Rossera wynika z lematu Newmana (silna normalizacja wraz ze słabą własnością Churcha-Rossera implikuje własność Churcha-Rossera). ■

26 System $\mathbf{F}_\omega$

O typie τ , który ma zmienną wolną p , można myśleć jak o operacji $\lambda p \tau$, która zaaplikowana do argumentu σ daje w wyniku pewien typ $\tau(\sigma)$. Ten sposób myślenia został zalegalizowany w systemie $\mathbf{F}_\omega$, zwanym polimorficznym rachunkiem lambda wyższego rzędu. Taką operację $\lambda p \tau$ nazywamy tutaj *konstruktorem typowym*. Raz dopuściwszy operacje na typach, nie widzimy nic zdrożnego we wprowadzeniu także operacji na konstruktorach, operacji na takich operacjach itd.

Formalizację powyższego zaczynamy od pojęcia *rodzaju* (ang. *kind*). Mamy stałą $*$ (rodzaj wszystkich typów) a ponadto jeśli ∇_1 i ∇_2 są rodzajami, to $\nabla_1 \Rightarrow \nabla_2$ jest rodzajem. Dla każdego rodzaju ∇ , definiujemy *konstruktory* tego rodzaju przez indukcję:

- *Zmienna konstruktorowa* rodzaju ∇ jest konstruktorem rodzaju ∇ .
- Jeśli φ jest konstruktorem rodzaju $\nabla_1 \Rightarrow \nabla_2$ i τ jest konstruktorem rodzaju ∇_1 , to $\varphi\tau$ jest konstruktorem rodzaju ∇_2 .
- Jeśli α jest zmienną konstruktorową rodzaju ∇_1 i τ jest konstruktorem rodzaju ∇_2 , to $\lambda\alpha \tau$ jest konstruktorem rodzaju $\nabla_1 \Rightarrow \nabla_2$.
- Jeśli α jest zmienną konstruktorową dowolnego rodzaju i τ jest konstruktorem rodzaju $*$, to $\forall\alpha \tau$ jest konstruktorem rodzaju $*$.
- Jeśli τ i σ są konstruktorami rodzaju $*$, to $\tau \rightarrow \sigma$ jest konstruktorem rodzaju $*$.

Piszemy $\tau : \nabla$ na oznaczenie tego, że τ jest konstruktorem rodzaju ∇ . Konstruktory rodzaju $*$ nazywamy *typami*. Jak dla $\mathbf{F}$ podstawienie (konstruktora na zmienną konstruktorową) definiujemy przez indukcję.

²⁷W obu wersjach.

- $\alpha[\alpha := \varphi] = \varphi$;
- $\beta[\alpha := \varphi] = \beta$;
- $(\vartheta\psi)[\alpha := \varphi] = \vartheta[\alpha := \varphi]\psi[\alpha := \varphi]$;
- $(\sigma \rightarrow \rho)[\alpha := \varphi] = \sigma[\alpha := \varphi] \rightarrow \rho[\alpha := \varphi]$;
- $(\forall\alpha\sigma)[\alpha := \varphi] = \forall\alpha\sigma$;
- $(\forall\beta\sigma)[\alpha := \varphi] = \forall\beta\sigma$, jeśli $\alpha \notin FV(\sigma)$;
- $(\forall\beta\sigma)[\alpha := \varphi] = \forall\beta\sigma[\alpha := \varphi]$, jeśli $\alpha \in FV(\sigma)$ i $\beta \notin FV(\varphi)$;
- $(\lambda\alpha\vartheta)[\alpha := \varphi] = \lambda\alpha\vartheta$;
- $(\lambda\beta\vartheta)[\alpha := \varphi] = \lambda\beta\vartheta$, jeśli $\alpha \notin FV(\vartheta)$;
- $(\lambda\beta\vartheta)[\alpha := \varphi] = \lambda\beta\vartheta[\alpha := \varphi]$, jeśli $\alpha \in FV(\vartheta)$ i $\beta \notin FV(\varphi)$;

Zauważmy, że mamy teraz dwie operacje wiążące zmienne konstruktorowe i typowe: lambdę i kwantyfikator. Alfa-konwersja jest określona warunkami

- $\forall\alpha\tau =_{\alpha} \forall\beta(\tau[\alpha := \beta])$, gdzie β nie jest wolne w τ ;
- $\lambda\alpha\tau =_{\alpha} \lambda\beta(\tau[\alpha := \beta])$, gdzie β nie jest wolne w τ ; β does not occur free in τ ;
- $\tau \rightarrow \rho =_{\alpha} \tau' \rightarrow \rho'$ i $\forall\alpha\tau =_{\alpha} \forall\alpha\tau'$ gdzie $\tau =_{\alpha} \tau'$ i $\rho =_{\alpha} \rho'$;
- $\tau\rho =_{\alpha} \tau'\rho'$ i $\lambda\alpha\tau =_{\alpha} \lambda\alpha\tau'$, gdzie $\tau =_{\alpha} \tau'$ i $\rho =_{\alpha} \rho'$;

Oczywiście utożsamiamy alfa-równoważne konstruktory.

Na konstruktorach określa się beta-redukcję:

$$(\beta) \quad (\lambda\alpha\sigma)\tau \Longrightarrow \sigma[\alpha := \tau].$$

Z silnej normalizacji dla typów skończonych łatwo wywnioskować, że każdy konstruktor jest silnie normalizowalny. Własność Churcha-Rossera też zachodzi. Przez $nf(\sigma)$ oznaczmy postać normalną typu σ .

Termy Churcha

System $\mathbf{F}_{\omega}$ w wersji Churcha różni się od systemu $\mathbf{F}$ tym, że zamiast aplikacji i abstrakcji typowej mamy aplikację i abstrakcję konstruktorową. Jeśli więc $M : \forall\alpha\sigma$, gdzie $\alpha : \nabla$, to dopuszczamy termy postaci $M\varphi : nf(\sigma[\alpha := \varphi])$, gdzie φ jest dowolnym konstruktorem rodzaju ∇ . A jeśli zmienna konstruktorowa $\alpha : \nabla$ nie występuje wolno w typie żadnej zmiennej wolnej termu M , to możemy utworzyć polimorficzną abstrakcję $\Lambda\alpha M$ typu $\forall\alpha M$. Tak określony rachunek ma własność Churcha-Rossera i własność silnej normalizacji.

Termy Curry'ego

System $\mathbf{F}_\omega$ w wersji Curry'ego jest zadany następującymi regułami przypisania typów.

VAR	$\Gamma \vdash x : \sigma$	gdy $\Gamma(x) = \sigma$
APP	$\frac{\Gamma \vdash M : \sigma \rightarrow \tau \quad \Gamma \vdash N : \sigma}{\Gamma \vdash (M \ N) : \tau}$	
ABS	$\frac{\Gamma \cup \{x : \sigma\} \vdash M : \tau}{\Gamma \vdash (\lambda x \ M) : \sigma \rightarrow \tau}$	
INST	$\frac{\Gamma \vdash M : \forall \alpha \ \sigma}{\Gamma \vdash M : nf(\sigma[\alpha := \tau])}$	
GEN	$\frac{\Gamma \vdash M : \sigma}{\Gamma \vdash M : \forall \alpha \ \sigma} \quad \alpha \notin \text{FV}(\Gamma)$	

Zamiast normalizacji typu w regule (INST) można przyjąć dodatkową regułę konwersji:

CNV	$\frac{\Gamma \vdash M : \sigma}{\Gamma \vdash M : \tau}$	gdy $\sigma =_\beta \tau$
-----	---	---------------------------

Przykład

Rozpatrzmy następujący term

$$M = (\lambda x \ xx\mathbf{K})2,$$

gdzie $\mathbf{K} = \lambda xy. x$ i $2 = \lambda fy. f(fy)$. Ten term jest typowalny w $\mathbf{F}_\omega$, a jedyne nietrywialne rodzaje jakich potrzebujemy to $* \Rightarrow * \text{ i } * \Rightarrow * \Rightarrow *$. Uwaga: ten przykład jest trudniejszy niż term $22\mathbf{K}$, bo mamy tylko jedną dwójkę.

Poniżej p i q są zmiennymi typowymi, δ , γ i β są zmiennymi konstruktorowymi rodzaju $* \Rightarrow *$, a ξ jest zmienną konstruktową rodzaju $* \Rightarrow * \Rightarrow *$.

Niech

$$\kappa = \forall p \delta (p \rightarrow q \rightarrow p),$$

i niech

$$\mathcal{L} = \forall p \gamma (\xi p (\forall p \delta (p \rightarrow \gamma p)) \rightarrow \xi(\beta p) (\forall p \delta (p \rightarrow \gamma(\gamma p)))).$$

Rozpatrzmy otoczenie złożone z deklaracji:

$$\{y : \xi p \kappa, f : \mathcal{L}\}$$

W tym otoczeniu wyprowadzimy

$$fy : \xi(\beta p) (\forall p \delta (p \rightarrow q \rightarrow q \rightarrow p)),$$

bo zmienna γ w typie f może być zastąpiona konstruktorem $\lambda p. q \rightarrow p$. Inną możliwość to podstawienie na γ konstruktora $\lambda p. q \rightarrow q \rightarrow p$, oraz podstawienie βp na p . To pozwala wyprowadzić

$$f(fy) : \xi(\beta(\beta p))\sigma,$$

gdzie

$$\sigma = \forall p \delta (p \rightarrow q \rightarrow q \rightarrow q \rightarrow q \rightarrow p).$$

Zauważmy, że p i γ nie są wolne w $\mathcal{L}$, więc w pustym otoczeniu mamy:

$$2 : \forall \xi p \beta (\mathcal{L} \rightarrow \forall p \gamma (\xi p \kappa \rightarrow \xi(\beta(\beta p))\sigma)).$$

Oznaczmy powyższy typ przez τ . Chcemy znaleźć typ dla termu $xx\mathbf{K}$ w otoczeniu, w którym zmienna x ma typ τ . Typem $\mathbf{K}$ będzie oczywiście κ . Typ drugiego wystąpienia x otrzymamy jako instancję τ : podstawimy na ξ rzutowanie typowe $\lambda p_1 p_2. p_1$. Otrzymamy typ $\forall p \beta (\forall p \gamma (p \rightarrow \beta p) \rightarrow \forall p \gamma (p \rightarrow \beta(\beta p)))$, alfa-równoważny typowi

$$\tau_2 = \forall p \gamma (\forall p \delta (p \rightarrow \gamma p) \rightarrow (\forall p \delta (p \rightarrow \gamma(\gamma p)))).$$

Pierwsze wystąpienie x otrzyma typ otrzymany z τ przez podstawienie drugiego rzutowania $\lambda p_1 p_2. p_2$, a mianowicie typ $\forall p \beta (\tau_2 \rightarrow \forall p \gamma (\kappa \rightarrow \forall p \delta (p \rightarrow q \rightarrow q \rightarrow q \rightarrow q \rightarrow p)))$. Pozostaje pozbyć się z tego typu początkowych kwantyfikatorów (za pomocą trywialnego podstawienia) i mamy taki typ dla x :

$$\tau_1 = \tau_2 \rightarrow \forall p \gamma (\kappa \rightarrow \sigma).$$

Poprawne typowanie dla $xx\mathbf{K}$ jest już łatwe.

Drugi przykład

Teraz (już bez dowodu) przykład termu, który ma własność silnej normalizacji, ale nie jest typowalny w $\mathbf{F}_\omega$. Oto on:

$$(\lambda x. z(x1)(x1'))(\lambda y. yyy)$$

Powyżej 1 oznacza $\lambda fu. fu$, a $1'$ oznacza $\lambda vg. gv$.

Nierozstrzygalność

Fakt 26.1 *Następujące problemy są nierozstrzygalne dla systemu $\mathbf{F}_\omega$:*

- *Sprawdzanie typu:* dane Γ, M, τ , pytamy czy $\Gamma \vdash M : \tau$?
- *Typowalność:* dane M , pytamy czy istnieją takie Γ, τ , że $\Gamma \vdash M : \tau$?

Szkic dowodu części pierwszej: Skorzystamy z nierozstrzygalności problemu unifikacji drugiego rzędu (twierdzenie 15.5). Wiadomo, że problem ten jest już nierozstrzygalny dla sygnatury złożonej z jednego dwuargumentowego symbolu funkcyjnego $\rightarrow$ (w notacji infiksowej) i dwóch stałych α i β . Co więcej, bez straty ogólności można ograniczyć się do pojedynczych równań postaci $\tau = \rho$, gdzie niewiadome (dla uproszczenia wszystkie niewiadome

są dwuargumentowymi symbolami funkcyjnymi) występujące w τ i ρ tworzą rozłączne zbiory. (Zauważmy bowiem, że dla dwuargumentowych niewiadomych ζ i ζ' , równania $\zeta\alpha\beta = \zeta'\alpha\beta$ i $\zeta\beta\alpha = \zeta'\beta\alpha$ „wymuszają” równość $\zeta = \zeta'$, tj. odpowiednie współrzędne każdego rozwiązania muszą być równe).

Założmy więc, że dane jest równanie $\tau = \rho$, gdzie niewiadome $\zeta_1, \dots, \zeta_k$ występują w τ , a niewiadome $\zeta_{k+1}, \dots, \zeta_n$ występują w ρ . Pytanie o rozwiązanie tego równania sprowadzamy do pytania o to, czy termowi $\mathbf{K}y(x(xy))$ można przypisać typ p w otoczeniu złożonym z deklaracji $(y : \forall pp)$ i $(x : \forall \zeta_1, \dots, \zeta_n(\tau \rightarrow \rho))$. Jest tak wtedy i tylko wtedy, gdy możliwa jest taka instancjacja zmiennych ζ_i , przy której możliwe jest składanie x ze sobą, tj. wtedy, gdy nasze równanie ma rozwiązanie. ■

Twierdzenie 26.2 *Problem niepustości typu w systemie $\mathbf{F}_\omega$ jest nierozstrzygalny.*

Dla dowodu twierdzenia 26.2 przypomnimy pojęcie *automatu dwulicznikowego*. Można go zdefiniować jako trójkę uporządkowaną $\mathcal{M} = \langle Q, q_0, q_f, \delta \rangle$, gdzie Q to oczywiście zbiór stanów, q_0 i q_f to stan początkowy i końcowy, a funkcja przejścia δ przypisuje każdemu stanowi (oprócz q_f) pewną *instrukcję*. Instrukcje mogą być takie (dla $i = 1, 2$ oraz $p, p' \in Q$):

1. $c_1 := c_1 + 1$; **goto** p ;
2. $c_2 := c_2 + 1$; **goto** p ;
3. $c_1 := c_1 - 1$; **goto** p ;
4. $c_2 := c_2 - 1$; **goto** p ;
5. **if** $c_1 = 0$ **then goto** p **else goto** p' ;
6. **if** $c_2 = 0$ **then goto** p **else goto** p' .

Konfiguracja automatu to trójka $\langle q, m, n \rangle$, gdzie $q \in Q$ oraz $m, n \in \mathbb{N}$. Konfiguracja *początkowa* ma postać $\langle q_0, 0, 0 \rangle$, a każda konfiguracja postaci $\langle q_f, m, n \rangle$ jest *końcowa*. Relacja przejścia $\rightarrow_{\mathcal{M}}$ pomiędzy konfiguracjami jest określona tak:

- $\langle q, m, n \rangle \rightarrow_{\mathcal{M}} \langle p, m + 1, n \rangle$, gdy $\delta(q)$ jest postaci (1);
- $\langle q, m, n \rangle \rightarrow_{\mathcal{M}} \langle p, m, n + 1 \rangle$, gdy $\delta(q)$ jest postaci (2);
- $\langle q, m + 1, n \rangle \rightarrow_{\mathcal{M}} \langle p, m, n \rangle$, gdy $\delta(q)$ jest postaci (3);
- $\langle q, m, n + 1 \rangle \rightarrow_{\mathcal{M}} \langle p, m, n \rangle$, gdy $\delta(q)$ jest postaci (4);
- $\langle q, 0, n \rangle \rightarrow_{\mathcal{M}} \langle p, 0, n \rangle$, gdy $\delta(q)$ jest postaci (5);
- $\langle q, m + 1, n \rangle \rightarrow_{\mathcal{M}} \langle p', m + 1, n \rangle$, gdy $\delta(q)$ jest postaci (5);
- $\langle q, m, 0 \rangle \rightarrow_{\mathcal{M}} \langle p, m, 0 \rangle$, gdy $\delta(q)$ jest postaci (6);
- $\langle q, m, n + 1 \rangle \rightarrow_{\mathcal{M}} \langle p', m, n + 1 \rangle$, gdy $\delta(q)$ jest postaci (6);

Jak widać, wykonanie instrukcji (1) i (2) jest niemożliwe w przypadku gdy odpowiedni licznik ma wartość zero. Jak zwykle $\rightarrow_{\mathcal{M}}$ oznacza domknięcie przechodnio-zwrotne relacji $\rightarrow_{\mathcal{M}}$. Mówimy, że maszyna $\mathcal{M}$ *zatrzymuje się*, gdy $\langle q_0, 0, 0 \rangle \rightarrow_{\mathcal{M}} \langle q_f, m, n \rangle$, dla pewnych m, n .

Twierdzenie 26.3 *Problem stopu dla automatów dwulicznikowych (czy dany automat się zatrzymuje?) jest nierozstrzygalny.*

Szkic dowodu: Naturalnym uogólnieniem pojęcia automatu dwulicznikowego jest automat z dowolną liczbą liczników. Nietrudno pokazać, że obliczenie dowolnej maszyny Turinga może być symulowane przez pewien automat licznikowy (przy pomocy odpowiedniego kodowania zawartości taśmy, położenia głowicy itd. jako liczb naturalnych). Pozostaje więc pokazać, jak zastąpić np. pięć liczników przez dwa.

Najpierw zauważmy, że używając drugiego licznika do celów pomocniczych możemy pomnożyć wartość licznika pierwszego przez ustaloną liczbę, na przykład 7. W tym celu najpierw zmniejszamy licznik drugi aż do zera, a potem powtarzamy taką czynność: odejmując 1 od licznika pierwszego, dodajemy 7 (tj. siedmiokrotnie dodajemy jedynkę) do licznika drugiego. Tak postępujemy aż do wyzerowania pierwszego licznika. Podobnie możemy dzielić przez 7, a także sprawdzić, czy wartość licznika jest podzielna przez 7.

Piątkę liczb naturalnych (zawartość pięciu liczników i, j, k, l, m) reprezentujemy za pomocą liczby $2^i 3^j 5^k 7^l 11^m$, którą przechowujemy jako wartość jednego z dwóch liczników. Drugi licznik służy do celów pomocniczych. Powiększeniu lub pomniejszeniu o 1 licznika l odpowiada teraz pomnożenie lub podzielenie naszej liczby przez 7. Możemy także sprawdzić, czy $l = 0$. W ten sposób obliczenie używające pięciu liczników zrealizujemy z pomocą dwóch. ■

Przejdźmy do rzeczy. Piszemy $\Gamma \vdash \tau$, gdy $\Gamma \vdash M : \tau$, dla pewnego M . W szczególności $\vdash \tau$ oznacza, że typ τ jest *niepusty*, tj. istnieje zamknięty term M typu τ . Interesuje nas *problem niepustości typu*, czyli pytanie „Czy istnieje term zamknięty danego typu?”. Łatwo widzieć, że ten problem jest równoważny zadaniu „Czy $\Gamma \vdash \tau$, dla danych Γ, τ ?”

Problem stopu dla automatów dwulicznikowych sprowadzimy do pytania: *Dane Γ i τ , czy $\Gamma \vdash \tau$?* Załóżmy, że dany jest automat $\mathcal{M} = \langle Q, q_0, q_N, \delta \rangle$, gdzie $Q = \{q_0, \dots, q_N\}$. Na początek ustalmy takie zmienne konstruktorowe:

$$Q_i : * \Rightarrow * \Rightarrow *, \text{ dla dowolnego } i = 1, \dots, N;$$

$$f : * \Rightarrow *; \quad 0 : *, \quad OK : *.$$

Teraz zbudujemy pewne otoczenie Γ , reprezentujące automat $\mathcal{M}$. Z każdą instrukcją $\delta(q)$ zwiążemy jedną lub dwie zmienne dekladowane w Γ . Typy tych zmiennych są następujące:

- $\forall xy (Q_i xy \rightarrow Q_j (fx)y)$, gdy $\delta(q_i) = „c_1 := c_1 + 1; \text{goto } q_j;”$
- $\forall xy (Q_i xy \rightarrow Q_j x(fy))$, gdy $\delta(q_i) = „c_2 := c_2 + 1; \text{goto } q_j;”$
- $\forall xy (Q_i (fx)y \rightarrow Q_j xy)$, gdy $\delta(q_i) = „c_1 := c_1 - 1; \text{goto } q_j;”$
- $\forall xy (Q_i x(fy) \rightarrow Q_j xy)$, gdy $\delta(q_i) = „c_2 := c_2 - 1; \text{goto } q_j;”$
- $\forall y (Q_i 0y \rightarrow Q_j 0y)$ oraz $\forall xy (Q_i (fx)y \rightarrow Q_k (fx)y)$,
gdy $\delta(q_i) = „\text{if } c_2 = 0 \text{ then goto } q_j \text{ else goto } q_k;”$

- $\forall x(\mathbb{Q}_i x 0 \rightarrow \mathbb{Q}_j x 0)$ oraz $\forall xy(\mathbb{Q}_i x(\mathbf{f}y) \rightarrow \mathbb{Q}_k x(\mathbf{f}y))$,
gdzie $\delta(q_i) = \text{„if } c_2 = 0 \text{ then goto } q_j \text{ else goto } q_k\text{.”}$

Na koniec do otoczenia Γ dokładamy jeszcze deklarację zmiennych takich typów:

$$\mathbb{Q}_0 00 \quad \text{oraz} \quad \forall xy(\mathbb{Q}_N xy \rightarrow 0K)$$

Oczywiście nazwy zadeklarowanych zmiennych tak naprawdę są bez znaczenia.

Na potrzeby tego dowodu wprowadźmy jeszcze oznaczenie $\underline{m} = \mathbf{f}^k(0)$. Oczywiście, $\underline{m} : *$.

Lemat 26.4 *Jeśli $\langle q_i, m, n \rangle \rightarrow_{\mathcal{M}} \langle q_N, m', n' \rangle$ to $\Gamma, \mathbb{Q}_i \underline{m} \underline{n} \vdash 0K$.*

Dowód: Indukcja ze względu na długość obliczenia rozpoczynającego się od konfiguracji $\langle q_i, m, n \rangle$. Jeśli to już jest konfiguracja końcowa, to należy wykorzystać zmienną typu $\forall xy(\mathbb{Q}_N xy \rightarrow 0K)$. W kroku indukcyjnym należy zauważyć, że jeśli $\langle q_i, m, n \rangle \rightarrow_{\mathcal{M}} \langle q_j, m', n' \rangle$ to $\Gamma, \mathbb{Q}_i \underline{m} \underline{n} \vdash \mathbb{Q}_j \underline{m'} \underline{n'}$. ■

Lemat 26.5 *Jeśli $\Gamma \vdash \mathbb{Q}_j \underline{m} \underline{n}$, to $\langle q_0, 0, 0 \rangle \rightarrow_{\mathcal{M}} \langle q_j, m, n \rangle$.*

Dowód: Istnieje taki term M , że $\Gamma \vdash M : \mathbb{Q}_j \underline{m} \underline{n}$. Można zakładać, że M jest w postaci normalnej. Dowód jest przez indukcję ze względu na rozmiar M . Z powodu swojego typu, M nie może być abstrakcją. Musi to być zmienna lub aplikacja. Jedyne typy w Γ , który zaczyna się od konstruktora reprezentującego stan, to zmienna $\mathbb{Q}_0 00$. Ale jeśli to ma być typ M , to znaczy, że $m = n = j = 0$ i teza zachodzi w sposób trywialny.

Jeśli M jest aplikacją, to ma postać $z\vec{N}$, gdzie z jest zmienną deklarowaną w Γ , a $\vec{N}$ jest ciągiem typów i termów. Typ zmiennej z ma postać $\forall x(\dots \rightarrow \mathbb{Q}_j uv)$ lub $\forall xy(\dots \rightarrow \mathbb{Q}_j uv)$ i odpowiada zastosowaniu pewnej instrukcji automatu $\mathcal{M}$. Przypuśćmy na przykład, że z jest zmienną typu $\forall xy(\mathbb{Q}_i x(\mathbf{f}y) \rightarrow \mathbb{Q}_j xy)$, odpowiadającego instrukcji $\delta(q_i) = c_2 := c_2 - 1; \text{goto } q_j$. Wtedy $M = z \underline{m} \underline{n} M'$ gdzie $\Gamma \vdash M' : \mathbb{Q}_i \underline{m}(\mathbf{f} \underline{n})$. Z założenia indukcyjnego wnioskujemy, że $\langle q_0, 0, 0 \rangle \rightarrow_{\mathcal{M}} \langle q_i, m, n + 1 \rangle$. Ponadto oczywiście $\langle q_i, m, n + 1 \rangle \rightarrow_{\mathcal{M}} \langle q_j, m, n \rangle$, więc ostatecznie $\langle q_0, 0, 0 \rangle \rightarrow_{\mathcal{M}} \langle q_j, m, n \rangle$, jak miało być. Podobnie postępujemy w przypadku pozostałych instrukcji. ■

Dowód twierdzenia 26.2: Z lematów 26.4 i 26.5 łatwo otrzymujemy równoważność:

$$\text{Automat } \mathcal{M} \text{ zatrzymuje się wtedy i tylko wtedy, gdy } \Gamma \vdash 0K.$$

To kończy dowód twierdzenia 26.2. ■

Uwaga: Zauważmy, że w dowodzie wykorzystaliśmy tylko zmienne konstruktorowe rodzaju $*$, $* \Rightarrow * \text{ i } * \Rightarrow * \Rightarrow *$, przy czym tylko zmienne rodzaju $*$ były wiązane kwantyfikatorami.

Podziękowania

Dziękuję Pani Małgorzacie Maciejewskiej, Panom Tomaszowi Domańskiemu, Stanisławowi Findeisenowi, Danielowi Hansowi, Szczepanowi Hummelowi, Wojciechowi Jaworskiemu, Michałowi Kotowskiemu, Michałowi Misiakowi, Filipowi Murlakowi, Łukaszowi Osipiukowi, Pawłowi Parysowi, Jakubowi Pochrybniakowi, Michałowi Rutkowskiemu, Adamowi Ślaskiemu, Mateuszowi Zakrzewskiemu i dr. Aleksemu Schubertowi za wykrycie rozmaitych błędów we wcześniejszych wersjach tych notatek.