

6. WYKŁAD 6: RACHUNEK λ : PROGRAMOWANIE W RACHUNKU λ , KOMBINATOR PUNKTU STAŁEGO, KODY DE BRUIJN'A.

6.1. Programowanie w rachunku λ . Chcąc rozwinąć rachunek λ do pełnowartościowego języka funkcyjnego, musimy najpierw pokazać, w jaki sposób można zakodować często używane datatypy takie jak zmienne boole'owskie, liczby naturalne, czy listy w rachunku λ . Jako że wszystkie wartości w rachunku λ są λ -abstrakcjami, wszystkie wzmiankowane datatypy również są zakodowane z λ -abstrakcjami. Gdy raz zobaczymy w jaki sposób można zakodować specyficzne datatypy, będziemy mogli używać ich jak wbudowanych datatypów.

6.1.1. Zmienne boole'owskie Churcha. Główną cechą charakteryzującą zmienne boole'owskie jest możliwość wyboru jednej z dwóch różnych opcji. Na przykład, boole'owska prawda dokonuje wyboru pierwszej z dwóch możliwych opcji, tak jak w wyrażeniu SML-a `if true then  $e_1$  else  $e_2$` . Tym samym wartości boole'owskie w rachunku λ , zwane tu wartościami boole'owskimi Churcha, zapisujemy następująco:

$$\begin{aligned} \text{tt} &= \lambda t. \lambda f. t \\ \text{ff} &= \lambda t. \lambda f. f \end{aligned}$$

Dalej, konstrukcja warunkowa `if  $e$  then  $e_1$  else  $e_2$` jest definiowana następująco:

$$\text{if } e \text{ then } e_1 \text{ else } e_2 = ee_1e_2$$

Oto proste przykłady redukcji konstrukcji warunkowych przy użyciu strategii wywołania według nazwy:

$$\begin{aligned} \text{if tt then } e_1 \text{ else } e_2 &= \text{tte}_1e_2 = (\lambda t. \lambda f. t)e_1e_2 \mapsto (\lambda f. e_1)e_2 \mapsto e_1 \\ \text{if ff then } e_1 \text{ else } e_2 &= \text{ffe}_1e_2 = (\lambda t. \lambda f. f)e_1e_2 \mapsto (\lambda f. f)e_2 \mapsto e_2 \end{aligned}$$

Operatory logiczne na zmiennych boole'owskich są zdefiniowane następująco:

$$\begin{aligned} \text{and} &= \lambda x. \lambda y. xy\text{ff} \\ \text{or} &= \lambda x. \lambda y. x\text{tty} \\ \text{not} &= \lambda x. x\text{fftt} \end{aligned}$$

Jako przykład rozważmy następujące ciągi redukcji e_1e_2 gdy $e_1 \mapsto^* \text{tt}$ oraz $e_1 \mapsto^* \text{ff}$, odpowiednio, przy użyciu strategii wywołania według nazwy:

$$\begin{array}{ll} \text{and}_{e_1e_2} & \text{and}_{e_1e_2} \\ \mapsto^* e_1e_2\text{ff} & \mapsto^* e_1e_2\text{ff} \\ \mapsto^* \text{tte}_2\text{ff} & \mapsto^* \text{ffe}_2\text{ff} \\ \mapsto^* e_2 & \mapsto^* \text{ff} \end{array}$$

Ciąg po lewej stronie pokazuje, że jeśli zachodzi $e_1 \mapsto^* \text{tt}$, to $\text{and}_{e_1e_2}$ oznacza taką samą wartość prawdy, co e_2 . Ciąg po prawej stronie pokazuje, że jeśli zachodzi $e_1 \mapsto^* \text{ff}$, to $\text{and}_{e_1e_2}$ wartościuje się do ff niezależnie od e_2 .

6.1.2. Pary uporządkowane. Główną cechą charakteryzującą pary jest możliwość przechowania dwóch niepowiązanych wartości i odczytania każdej z nich wtedy, gdy jest to konieczne. Tym samym chcąc reprezentować parę (e_1, e_2) , musimy zbudować λ -abstrakcję, która zwraca e_1 lub e_2 , gdy zastosujemy ją do tt i ff , odpowiednio. Operatory rzutowania traktują parę jak λ -abstrakcję i stosują ją do tt albo ff .

$$\begin{aligned} \text{pair} &= \lambda x. \lambda y. \lambda b. bxy\text{ff} \\ \text{fst} &= \lambda p. p\text{tty} \\ \text{snd} &= \lambda p. p\text{ff} \end{aligned}$$

Jako przykład zredukujmy $\text{fst}(\text{paire}_1 e_2)$ z użyciem strategii wywołania według nazwy. Zauważmy, że para $e_1 e_2$ wartościuje się do $\lambda b. b e_1 e_2$, które z kolei oczekuje wartości boole'owskiej b celem wyboru e_1 lub e_2 . Jeżeli podstawimy tt w miejsce b , to wówczas $e_1 e_2$ zredukuje się do e_1 .

$$\begin{aligned} \text{fst}(\text{paire}_1 e_2) &\mapsto (\text{paire}_1 e_2) \text{tt} \\ &\mapsto^* (\lambda b. b e_1 e_2) \text{tt} \\ &\mapsto \text{tte}_1 e_2 \\ &\mapsto^* e_1 \end{aligned}$$

6.1.3. *Liczebniki Churcha.* Główną cechą charakteryzującą liczbę naturalną n jest możliwość przekazania informacji, że daną czynność należy powtórzyć n razy. W przypadku rachunku λ liczba n jest zakodowana jako λ -abstrakcja $\hat{n}$ zwana **liczebnikiem Churcha**, która bierze daną funkcję f i zwraca $f^n = \underbrace{f \circ f \circ \dots \circ f}_{n \text{ razy}}$. Zauważmy, że f^0 jest funkcją identycznościową $\lambda x. x$, ponieważ f zastosowane 0 razy

do swego argumentu x niczego nie zmienia, tak więc $\hat{1}$ jest sama funkcją identycznościową $\lambda f. f$.

$$\begin{aligned} \hat{0} &= \lambda f. f^0 = \lambda f. \lambda x. x \\ \hat{1} &= \lambda f. f^1 = \lambda f. \lambda x. f x \\ \hat{2} &= \lambda f. f^2 = \lambda f. \lambda x. f(f x) \\ \hat{3} &= \lambda f. f^3 = \lambda f. \lambda x. f(f(f x)) \\ &\vdots \\ \hat{n} &= \lambda f. f^n = \lambda f. \lambda x. f(f(f \dots (f x) \dots)) \end{aligned}$$

Jeśli będziemy odczytywać f jak **S**, to wówczas $\hat{n} f x$ zwróci nam reprezentację liczby naturalnej n znaną z wcześniejszych wykładów.

Zdefiniujmy teraz operacje arytmetyczne na liczbach naturalnych. Operacja dodawania $\text{add} \hat{m} \hat{n}$ zwraca $\widehat{m+n}$, czyli λ -abstrakcję biorącą funkcję f i zwracającą f^{m+n} . Ponieważ f^{m+n} może być zapisane jako $\lambda x. f^{m+n} x$, skonstruujemy add następująco (celem odróżnienia liczb naturalnych, na przykład liczby n , od odpowiadających im postaci zakodowanych, na przykład $\hat{n}$, będziemy używać $\hat{m}$ i $\hat{n}$ jako zmiennych):

$$\begin{aligned} \text{add} &= \lambda \hat{m}. \lambda \hat{n}. \lambda f. f^{m+n} \\ &= \lambda \hat{m}. \lambda \hat{n}. \lambda f. \lambda x. f^{m+n} x \\ &= \lambda \hat{m}. \lambda \hat{n}. \lambda f. \lambda x. f^m (f^n x) \\ &= \lambda \hat{m}. \lambda \hat{n}. \lambda f. \lambda x. \hat{m} f (\hat{n} f x) \end{aligned}$$

Zauważmy, że f^m otrzymujemy jako $\hat{m} f$ (i analogicznie dla f^n).

Operacja mnożenia może być zdefiniowana na dwa sposoby. Prostą sztuczką jest wykorzystanie równości $m \cdot n = \underbrace{m + m + \dots + m}_{n \text{ razy}}$. Innymi słowy $m \cdot n$ otrzymujemy przez dodanie m do zera dokładnie n

razy. Jako że $\text{add} \hat{m}$ jest w założeniu funkcją dodającą m do swego argumentu, zastosujemy $\text{add} \hat{m}$ do $\hat{0}$ dokładnie n razy celem otrzymania $\widehat{m \cdot n}$ – lub, równoważnie, zastosujemy $(\text{add} \hat{m})^n$ do $\hat{0}$:

$$\begin{aligned} \text{mult} &= \lambda \hat{m}. \lambda \hat{n}. (\text{add} \hat{m})^n \hat{0} \\ &= \lambda \hat{m}. \lambda \hat{n}. \hat{n} (\text{add} \hat{m}) \hat{0} \end{aligned}$$

Drugi pomysł na zdefiniowanie tej samej operacji (który w istocie może być prostszy do wymyślenia od podanego przed chwilą rozwiązania) polega na wykorzystaniu równania $f^{m \cdot n} = (f^m)^n = (\hat{m} f)^n = \hat{n} (\hat{m} f)$:

$$\text{mult} = \lambda \hat{m}. \lambda \hat{n}. \lambda f. \hat{n} (\hat{m} f)$$

Operacja odejmowania jest trudniejsza do zdefiniowania niż poprzednie dwie. Przypuśćmy bowiem, że mamy zdefiniowaną funkcję poprzednika **pred** obliczającą poprzednik danej liczby naturalnej: $\widehat{\text{pred } \hat{n}}$ zwraca $\widehat{n - 1}$ jeśli $n > 0$ lub $\hat{0}$ w przeciwnym wypadku. Chcąc zdefiniować operację odejmowania $\widehat{\text{sub } \hat{m} \hat{n}}$, która zwraca $\widehat{m - n}$ jeśli $m > n$ lub $\hat{0}$ w przeciwnym wypadku, musimy zastosować **pred** do $\hat{m}$ dokładnie n razy:

$$\begin{aligned} \text{sub} &= \lambda \hat{m}. \lambda \hat{n}. \text{pred}^n \hat{m} \\ &= \lambda \hat{m}. \lambda \hat{n}. (\hat{n} \text{pred}) \hat{m} \end{aligned}$$

Pozostaje, oczywiście, problem zdefiniowania funkcji **pred**. Zrobimy to przy użyciu pomocniczej funkcji **next**, która dla danej pary $\widehat{\text{pair } \hat{k} \hat{m}}$ ignoruje $\hat{k}$ i zwraca parę $\widehat{\text{pair } \hat{m} \hat{m} + 1}$:

$$\text{next} = \lambda p. \text{pair}(\text{snd } p)(\text{add}(\text{snd } p)\hat{1})$$

Można pokazać, że stosując **next** do pary $\widehat{\text{pair } \hat{0} \hat{0}}$ dokładnie n razy otrzymamy parę $\widehat{\text{pair } \hat{n} - 1 \hat{n}}$ jeśli $n > 0$ (przy użyciu pewnej strategii redukcyjnej):

$$\begin{aligned} \text{next}^0(\widehat{\text{pair } \hat{0} \hat{0}}) &\mapsto^* \widehat{\text{pair } \hat{0} \hat{0}} \\ \text{next}^1(\widehat{\text{pair } \hat{0} \hat{0}}) &\mapsto^* \widehat{\text{pair } \hat{0} \hat{1}} \\ \text{next}^2(\widehat{\text{pair } \hat{0} \hat{0}}) &\mapsto^* \widehat{\text{pair } \hat{1} \hat{2}} \\ &\vdots \\ \text{next}^n(\widehat{\text{pair } \hat{0} \hat{0}}) &\mapsto^* \widehat{\text{pair } \hat{n} - 1 \hat{n}} \end{aligned}$$

Jako że poprzednikiem zera jest tak czy inaczej zero, pierwszy składnik $\text{next}^n(\widehat{\text{pair } \hat{0} \hat{0}})$ koduje poprzednik n . Tym samym operacja **pred** jest zdefiniowana następująco:

$$\begin{aligned} \text{pred} &= \lambda \hat{n}. \text{fst}(\text{next}^n(\widehat{\text{pair } \hat{0} \hat{0}})) \\ &= \lambda \hat{n}. \text{fst}(\widehat{\text{next } \hat{0} \hat{0}}) \end{aligned}$$

6.2. Kombinator punktu stałego. Jako że siła wyrazu rachunku λ jest taka sama jak maszyn Turinga, każda maszyna Turinga może być symulowana za pomocą pewnego wyrażenia w rachunku λ . W szczególności istnieją wyrażenia w rachunku λ , które odpowiadają maszynom Turinga, które się nie zatrzymują, jak również istnieją maszyny Turinga obliczające **funkcje rekursywne**.

Jest stosunkowo łatwo znaleźć wyrażenie, którego redukcja się nie zatrzymuje. Przypuśćmy bowiem, że chcemy znaleźć wyrażenie **omega** takie, że $\text{omega} \mapsto \text{omega}$. Jako że redukuje się do samego siebie, jego redukcja nigdy się nie skończy. Przepisemy **omega** jako $(\lambda x. e)e'$ tak, aby β -redukcja mogła być zastosowana do całego wyrażenie ω . Otrzymujemy wówczas:

$$\text{omega} = (\lambda x. e)e' \mapsto [e'/x]e = \text{omega}.$$

Równość $\text{omega} = [e'/x]e = (\lambda x. e)e'$ sugeruje, że $e = e''$ dla pewnego wyrażenia e'' takiego, że $[e'/x]e'' = \lambda x. e$ (oraz $[e'/x]x = e'$):

$$\begin{aligned} \text{omega} &= [e'/x]e \\ &= [e'/x](e''x) && \text{wobec } e = e''x \\ &= [e'/x]e''[e'/x]x \\ &= [e'/x]e''e' \\ &= (\lambda x. e)e' && \text{wobec } [e'/x]e'' = \lambda x. e \end{aligned}$$

Wobec $e = e''x$ oraz $[e'/x]e'' = \lambda x.e$ otrzymujemy $[e'/x]e'' = \lambda x.e''x$. Podstawiając $e'' = x$ w $[e'/x]e'' = \lambda x.e''x$ otrzymujemy $e' = \lambda x.xx$. Tym samym **omega** może być zdefiniowana jak następuje:

$$\begin{aligned} \text{omega} &= (\lambda x.e)e' \\ &= (\lambda x.e''x)e' && \text{wobec } e = e''x \\ &= (\lambda x.xx)e' && \text{wobec } e'' = x \\ &= (\lambda x.xx)(\lambda x.xx) && \text{wobec } e' = \lambda x.xx \end{aligned}$$

Teraz możemy pokazać, że redukcja **omega** zdefiniowana jak wyżej nigdy się nie zatrzyma.

Zobaczmy teraz w jaki sposób możemy definiować funkcje rekursywne w rachunku λ . Zaczniemy od założenia, że mamy dany **konstruktor funkcji rekursywnej** $\text{fun}fx.e$, który definiuje funkcję rekursywną f o argumentach x i zawartości e . Zauważmy, że zawartość e może zawierać odnośniki do f . Naszym celem będzie pokazanie, że $\text{fun}fx.e$ jest w istocie **słodzikim syntaktycznym** w takim sensie, że może zostać zapisany jako już istniejące wyrażenie w rachunku λ i tym samym jego dodanie nie zmienia siły wyrazu rachunku λ .

Jako roboczy przykład rozważmy funkcję silnia **fac**:

$$\text{fac} = \text{fun } f \text{ } n. \text{if eq } n \hat{0} \text{ then } \hat{1} \text{ else mult } n(f(\text{pred } n))$$

Semantycznie f w zawartości powyższej funkcji oznacza definiowaną właśnie funkcję **fac**. Na początek mechanicznie tworzymy λ -abstrakcję $\text{FAC} = \lambda f. \lambda n. e$ z $\text{fac} = \text{fun}fx.e$:

$$\text{FAC} = \lambda f. \lambda n. \text{if eq } n \hat{0} \text{ then } \hat{1} \text{ else mult } n(f(\text{pred } n))$$

Zauważmy, że **FAC** ma zupełnie inną charakterystykę niż **fac**: podczas gdy **fac** bierze liczbę naturalną n i zwraca inną liczbę naturalną, **FAC** bierze funkcję f i zwraca inną funkcję (gdyby **fac** oraz **FAC** mogły mieć typy, **fac** miałby typ $\text{nat} \rightarrow \text{nat}$, natomiast **FAC** miałby typ $(\text{nat} \rightarrow \text{nat}) \rightarrow (\text{nat} \rightarrow \text{nat})$).

Główny pomysł, na którym opiera się konstrukcja funkcji **FAC** polega na tym, że przy danej **częściowej** implementacji f funkcji silnia, $\text{FAC}f$ zwraca **poprawioną** implementację tej samej funkcji. Przypuśćmy bowiem, że f poprawnie oblicza silnię każdej liczby naturalnej nie większej od n . Wówczas $\text{FAC}f$ poprawnie obliczy silnię z każdej liczby naturalnej nie większej od $n+1$ – jest to więc istotna poprawa implementacji. Zauważmy też, że $\text{FAC}f$ poprawnie oblicza silnię z zera, niezależnie od f . W szczególności nawet wtedy, gdy weźmiemy przekazującą możliwie najmniej informacji funkcja $f = \lambda n. \text{omega}$ (która niczego nie robi, ponieważ nigdy niczego nie zwróci), $\text{FAC}f$ poprawnie obliczy silnię z zera. Tym samym możemy sobie wyobrazić nieskończony ciąg funkcji $\text{fac}_0, \text{fac}_1, \dots, \text{fac}_i, \dots$ zaczynający się od $\text{fac}_0 = \text{FAC}\lambda n. \text{omega}$ i kolejno wykorzystuje równość $\text{fac}_{i+1} = \text{FAC}\text{fac}_i$:

$$\begin{aligned} \text{fac}_0 &= \text{FAC}\lambda n. \text{omega} \\ \text{fac}_1 &= \text{FAC}\text{fac}_0 = \text{FAC}^2\lambda n. \text{omega} \\ \text{fac}_2 &= \text{FAC}\text{fac}_1 = \text{FAC}^3\lambda n. \text{omega} \\ &\vdots \\ \text{fac}_i &= \text{FAC}\text{fac}_{i-1} = \text{FAC}^{i+1}\lambda n. \text{omega} \\ &\vdots \end{aligned}$$

Zauważmy, że fac_i poprawnie oblicza silnię z każdej liczby naturalnej nie większej niż i . Wobec tego, jeśli ω oznacza nieskończoną liczbę naturalną (o której, w uproszczeniu, możemy myśleć jak o liczbie naturalnej większej od dowolnej liczby naturalnej), możemy wziąć fac_ω jako poprawną implementację funkcji silnia **fac**, a zatem $\text{fac} = \text{fac}_\omega$.

Kolejna ważna obserwacja, jaką poczynimy, jest następująca: mając daną poprawną implementację funkcji silnia **fac**, FACfac zwraca nową poprawną implementację funkcji silnia. Innymi słowy, jeśli **fac** jest

poprawną implementacją funkcji silnia, to wówczas

$$\lambda n. \text{if eq } n \hat{0} \text{ then } \hat{1} \text{ else mult } n(\text{fac}(\text{pred } n))$$

jest również poprawną implementacją funkcji silnia. Jako że obydwie funkcje są *de facto* identyczne w takim sensie, że zwracają taki sam wynik po zastosowaniu do dowolnego argumentu, możemy zapisać $\text{fac} = \text{FACfac}$. Jeśli podstawimy fac_ω za fac w równaniu, otrzymamy $\text{fac}_\omega = \text{fac}_{\omega+1}$, co również ma sens, jako że $\omega \leq \omega + 1$ wobec definicji dodawania oraz $\omega + 1 \leq \omega$ wobec definicji liczby ω .

Wydaje się zatem, że **FAC** zawiera wszelką informację potrzebną do obliczenia $\text{fac} = \text{fac}_\omega = \text{FACfac}$. Istotnie, okazuje się, że fac otrzymujemy przez zastosowanie **kombinatora punktu stałego** fix do **FAC**, czyli że $\text{fac} = \text{fixFAC}$, gdzie fix jest zdefiniowany następująco:

$$\text{fix} = \lambda F. (\lambda f. F(\lambda x. f f x)) (\lambda f. F(\lambda x. f f x))$$

Zakładamy tutaj użycie stratego wywołania według wartości. Dla strategii wywołania według nazwy upraszczamy $\lambda x. f f x$ do $f f$ i stosujemy następujący kombinator punktu stałego fix_{CBN} :

$$\text{fix}_{CBN} = \lambda F. (\lambda f. F(f f)) (\lambda f. F(f f))$$

Aby zrozumieć jak działa kombinator punktu stałego fix musimy najpierw zrozumieć pojęcie **punktu stałego**. Punktem stałym funkcji f będziemy nazywali wartość v taką, że $v = f(v)$. Na przykład punktem stałym funkcji $f(x) = 2 - x$ jest 1, ponieważ $f(1) = 1$. Tak jak wskazuje sama nazwa, fix bierze funkcję F (która sama zamienia funkcję f na nową funkcję f') i zwraca jej punkt stały. Innymi słowy, $\text{fix}F$ jest punktem stałym F :

$$\text{fix}F = F(\text{fix}F)$$

Nieformalnie wyrażenie po lewej stronie zamienia się w wyrażenie po prawej stronie poprzez ciąg następujących kroków, w których używamy symbolu $\approx$ do podkreślenia braku rygoru formalnego, który nie jest dostatecznie uzasadniony samą tylko β -redukcją:

$$\begin{aligned} \text{fix}F &\mapsto g \ g && \text{gdzie } g = \lambda f. F(\lambda x. f f x) \\ &= (\lambda f. F(\lambda x. f f x))g \\ &\mapsto F(\lambda x. g g x) \\ &\approx F(g \ g) && \text{ponieważ } \lambda x. g g x \approx g g \\ &\approx F(\text{fix}F) && \text{ponieważ } \text{fix}F \mapsto g g \end{aligned}$$

Możemy teraz wyjaśnić, dlaczego fixFAC zwraca implementację funkcji silnia. Zgodnie z naturą kombinatora punktu stałego fix mamy

$$\text{fixFAC} = \text{FAC}(\text{fixFAC}).$$

Innymi słowy, fixFAC zwraca funkcję f spełniającą warunek $f = \text{FAC}f$, a więc dokładnie ten warunek, jaki musi spełniać fac . Tym samym możemy traktować fixFAC równoważnie z fac .

Inny sposób wyjaśnienia zachowania fixFAC jest następujący: przypuśćmy, że chcemy obliczyć $\text{fac}n$ dla dowolnie dużej liczby naturalnej n . Ponieważ fixFAC jest punktem stałym **FAC**, otrzymujemy następujące równanie:

$$\begin{aligned} \text{fixFAC} &= \text{FAC}(\text{fixFAC}) \\ &= \text{FAC}(\text{FAC}(\text{fixFAC})) = \text{FAC}^2(\text{fixFAC}) \\ &= \text{FAC}^2(\text{FAC}(\text{fixFAC})) = \text{FAC}^3(\text{fixFAC}) \\ &\vdots \\ &= \text{FAC}^n(\text{FAC}(\text{fixFAC})) = \text{FAC}^{n+1}(\text{fixFAC}) \end{aligned}$$

Kluczowa jest tu następująca obserwacja: $\text{FAC}^{n+1}(\text{fixFAC})$ poprawnie oblicza silnię dowolnej liczby naturalnej nie większej od n **niezależnie** od tego, co robi fixFAC . Jako że $\text{fixFAC} = \text{FAC}^{n+1}(\text{fixFAC})$, widzimy

że `fixFAC` poprawnie oblicza silnie **dowolnej** liczby naturalnej. Innymi słowy, `fixFAC` robi dokładnie to samo, co `fac`.

Reasumując, chcąc zakodować funkcję rekursywną `funfx.e` w rachunku λ , zaczynamy od λ -abstrakcji $F = \lambda f.\lambda x.e$, a następnie `fixF` **automagicznie** zwraca funkcję zachowującą się tak samo, jak `funfx.e`.

6.3. Kody de Bruijna. Jako że λ -abstrakcje mają w założeniu oznaczać funkcje matematyczne z formalnymi argumentami, nazwy zmiennych powinny być integralną częścią syntaksu rachunku λ . Na przykład, wydaje się nieuniknione, aby w użyć formalny argument, powiedzmy x , przy definiowaniu funkcji identycznościowej. Z drugiej strony wybór konkretnej nazwy dla zmiennej nie ma wpływu na znaczenie λ -abstrakcji. Na przykład $\lambda x.x$ oraz $\lambda y.y$ oznaczają taką samą funkcję identycznościową, aczkolwiek używają różnych nazw argumentów. W ogólności α -konwersja pozwala nam na przepisanie dowolnej λ -abstrakcji do postaci nowej λ – *abstrakcji* z inną nazwą zmiennej związanej. Obserwacja ta pozwala przypuszczać, że istnieje sposób reprezentowania zmiennych w rachunku λ bez użycia konkretnych nazw zmiennych. Przykładem takiej beznazwowej interpretacji zmiennych są **kody de Bruijna**.

Główny pomysł wykorzystywany przez kody de Bruijna polega na przypisaniu każdej zmiennej liczby całkowitej, zwanej kodem de Bruijna, zamiast nazwy. Kody de Bruijna mogą być w szczególności ujemne, ale zajmujemy się tu tylko dodatnimi kodami. Z grubsza biorąc kod de Bruijna zlicza **λ -wiązania** takie jak λx , λy oraz λz leżące pomiędzy daną zmienną a odpowiadającym jej jednoznacznie wyznaczonym λ -wiązaniem. Na przykład, x w zawartości $\lambda x.x$ ma przypisany kod de Bruijna 0, ponieważ pomiędzy x oraz λx nie występuje żadne λ -wiązanie. Z drugiej strony x w zawartości $\lambda x.\lambda y.xy$ ma kod de Bruijna 1, ponieważ pomiędzy x oraz λx występuje jeszcze λ -wiązanie λy . Tym samym kod de Bruijna zmiennej określa jej względne położenie wobec odpowiadającego jej λ -wiązania. W konsekwencji widać, że taka sama zmienna może mieć przypisane różne wartości kodów de Bruijna w zależności od swego położenia.

Jako że wszystkie zmienne są teraz reprezentowane przez liczby całkowite, nie ma potrzeby wprowadzania *explicite* zmiennych w λ -abstrakcjach. W istocie jest to niemożliwe, ponieważ tej samej zmiennej mogą zostać przyporządkowane różne kody de Bruijna. Tym samym wyrażenia z kodami de Bruijna, albo – jak je będziemy nazywać – **wyrażenia de Bruijna**, są zdefiniowane indukcyjnie jak następuje:

$$\begin{aligned} \text{wyrażenie de Bruijna } M &::= n|\lambda.M|M \quad M \\ \text{kod de Bruijna } n &::= 0|1|2|\dots \end{aligned}$$

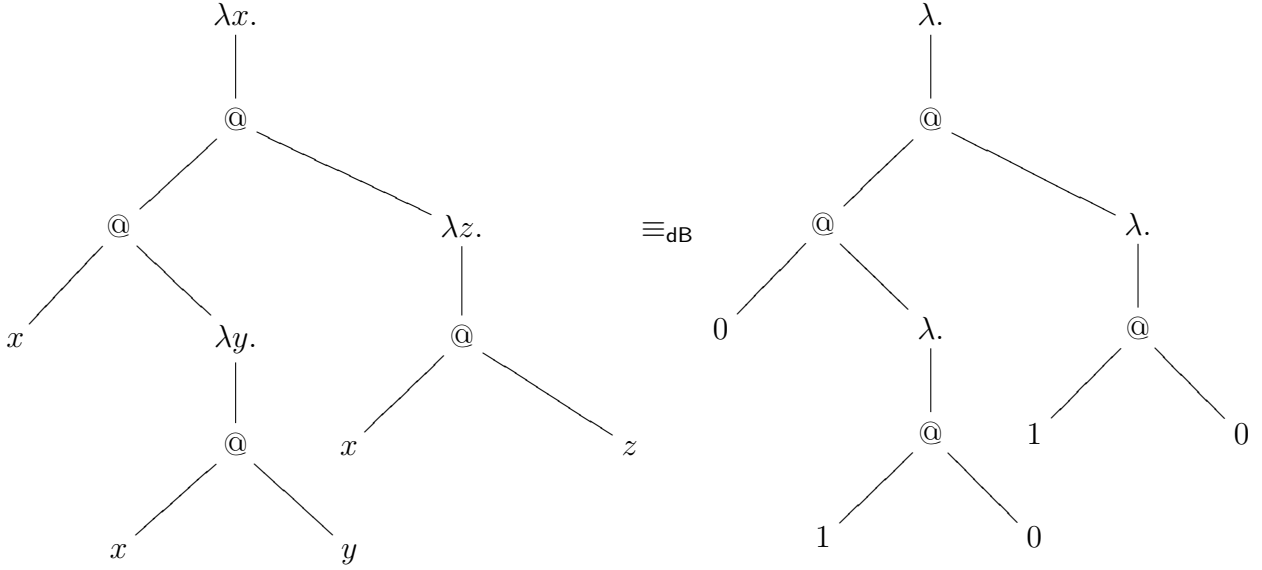
Dla wyrażeń de Bruijna będziemy używali metazmiennych M i N , zaś dla kodów de Bruijna metazmiennych m , n oraz i .

Będziemy pisali $e \equiv_{\text{dB}} M$ dla oznaczenia, że zwykłe wyrażenie e jest konwertowane do wyrażenia de Bruijna M . Czasami wystarczy dosłownie policzyć liczbę λ -wiązań pomiędzy poszczególnymi zmiennymi i odpowiadającymi im λ -wiązaniami, tak jak w następujących przykładach:

$$\begin{aligned} \lambda x.x &\equiv_{\text{dB}} \lambda.0 \\ \lambda x.\lambda y.xy &\equiv_{\text{dB}} \lambda.\lambda.1 \ 0 \\ \lambda x.x(\lambda y.xy) &\equiv_{\text{dB}} \lambda.0(\lambda.1 \ 0) \end{aligned}$$

W ogólności, wszelako, konwertowanie zwykłego wyrażenia e do wyrażenia de Bruijna wymaga jednak zinterpretowania e jako obiektu o strukturze pewnego drzewka, nie zaś obiektu o strukturze liniowej. Jako przykład rozważmy wyrażenie $\lambda x.(x(\lambda y.xy))(\lambda z.xz)$. Licząc bezpośrednio λ -wiązania otrzymujemy wyrażenie de Bruijna $\lambda.(0(\lambda.1 \ 0))(\lambda.2 \ 0)$, w którym ostatnie wystąpienie zmiennej x ma przypisany błędny kod de Bruijna 2, wynikający ze zliczenia λy oraz λz . Intuicyjnie czujemy jednak, że ostatnie wystąpienie zmiennej x musi mieć przypisany kod de Bruijna 1, ponieważ odpowiadające jej λ -wiązanie może być zlokalizowane z pominięciem λ -wiązania λy . Tym samym prawidłowa metoda konwertowania wyrażenia e do wyrażenia de Bruijna polega na zliczeniu λ -wiązań występujących pomiędzy daną

zmienną a odpowiadającym jej λ -wiązanemu występującym w reprezentacji e jako drzewa. Na przykład $\lambda x.(x(\lambda y.xy))(\lambda z.xz) \equiv_{\text{dB}} \lambda.(0(\lambda.1\ 0))(\lambda.1\ 0)$, zgodnie z poniższą ilustracją:



6.4. Podstawianie. Celem wykorzystania kodów de Bruijna do implementowania semantyki operacyjnej rachunku λ , będziemy potrzebować definicję podstawiania dla wyrażeń de Bruijna, z której, z kolei, będziemy w stanie wywieść definicję β -redukcji. Chcemy mianowicie zdefiniować $\sigma_0(M, N)$ tak, aby zachodził następujący związek:

$$\begin{aligned} (\lambda x.e)e' &\mapsto [e'/x]e \\ \equiv_{\text{dB}} &\quad \equiv_{\text{dB}} \\ (\lambda.M)N &\mapsto \sigma_0(M, N) \end{aligned}$$

Innymi słowy, stosując $\lambda.M$ do N , lub podstawiając N za zmienną związaną w $\lambda.M$, otrzymamy $\sigma_0(M, N)$ (znaczenie indeksu 0 w $\sigma_0(M, N)$ zostanie wyjaśnione później).

Zamiast zaczynać od kompletnej definicji $\sigma_0(M, N)$, zaczniemy od kilku prostych przykładów, które wykorzystamy do poprawienia pierwotnej naiwnej wersji definicji. Rozważmy następujący przykład, w którym redeks jest podkreślony:

$$\begin{aligned} \lambda x.\lambda y.\underline{(\lambda z.xyz)(\lambda w.w)} &\mapsto \lambda x.\lambda y.xy(\lambda w.w) \\ \equiv_{\text{dB}} &\quad \equiv_{\text{dB}} \\ \lambda.\lambda.\underline{(\lambda.2\ 1\ 0)}(\lambda.0) &\mapsto \lambda.\lambda.10(\lambda.0) \end{aligned}$$

Zauważmy, że 0, które odpowiada zmiennej związanej z w λ -abstrakcji $\lambda z.x\ y\ z$, jest zamieniona przez argument $\lambda.0$. Pozostałe kody 1 oraz 2 są zmniejszone o 1, ponieważ λ -wiązanie λz znika. Te dwie obserwacje prowadzą nas do następującej częściowej definicji $\sigma_0(M, N)$:

$$\begin{aligned} \sigma_0(M_1 M_2, N) &= \sigma_0(M_1, N) \sigma_0(M_2, N) \\ \sigma_0(0, N) &= N \\ \sigma_0(m, N) &= m - 1 && \text{jeśli } m > 0. \end{aligned}$$

Aby zobaczyć jak definiujemy pozostały przypadek $\sigma_0(\lambda.M, N)$, rozważmy następujący przykład, w którym redeks jest podkreślony:

$$\begin{array}{ccc} \lambda x.\lambda y.\underline{(\lambda z.(\lambda u.x\ y\ z\ u))(\lambda w.w)} & \mapsto & \lambda x.\lambda y.(\lambda u.x\ y(\lambda w.w)u) \\ \equiv_{dB} & & \equiv_{dB} \\ \lambda.\lambda.(\lambda.(\lambda.3\ 2\ 1\ 0))(\lambda.0) & \mapsto & \lambda.\lambda.(\lambda.2\ 1(\lambda.0)0) \end{array}$$

Zauważmy, że – inaczej niż w pierwszym przykładzie – 0 pozostaje nietknięte, ponieważ odpowiada do zmiennej związanej u w $\lambda u.x\ y\ z\ u$, podczas gdy 1 odpowiada z i tym samym jest zamienione przez $\lambda.0$. Powodem dla którego 1 jest teraz zastąpione przez $\lambda.0$ jest to, że kod de Bruijna m poza $\lambda.M$ wskazuje na taką samą zmienną jak $m + 1$ wewnątrz $\lambda.M$, to znaczy wewnątrz M . Ta obserwacja prowadzi nas do równania $\sigma_0(\lambda.M, N) = \lambda.\sigma_1(M, N)$, gdzie $\sigma_1(M, N)$ jest zdefiniowane jak następuje:

$$\begin{array}{ll} \sigma_1(M_1 M_2, N) & = \sigma_1(M_1, N) \sigma_1(M_2, N) \\ \sigma_1(0, N) & = 0 \\ \sigma_1(1, N) & = N \\ \sigma_1(m, N) & = m - 1 \quad \text{jeśli } m > 1. \end{array}$$

W powyższych dwóch przykładach zobaczyliśmy, że indeks n w $\sigma_n(M, N)$ pełni rolę wskaźnika "granicznego": m pozostaje bez zmian, jeśli $m < n$, jest zastępowane przez N , jeśli $m = n$ i jest zmniejszane o 1, jeśli $m > n$. Alternatywnie możemy spojrzeć na n w $\sigma_n(M, N)$ jako na liczbę λ -wiązań otaczających M tak, jak ilustruje poniższy wzór:

$$\sigma_0(\underbrace{\lambda.\lambda.\dots\lambda}_n.M, N) = \underbrace{\lambda.\lambda.\dots\lambda}_n \sigma_n(M, N)$$

Poniższa definicja $\sigma_n(M, N)$ wykorzystuje n jako wskaźnik graniczny i uogólnia związek pomiędzy $\sigma_0(\lambda.M, N)$ oraz $\lambda.\sigma_1(M, N)$:

$$\begin{array}{ll} \sigma_n(M_1 M_2, N) & = \sigma_n(M_1, N) \sigma_n(M_2, N) \\ \sigma_n(\lambda.M, N) & = \lambda.\sigma_{n+1}(M, N) \\ \sigma_n(m, N) & = m \quad \text{jeśli } m < n, \\ \sigma_n(n, N) & = N \\ \sigma_n(m, N) & = m - 1 \quad \text{jeśli } m > n. \end{array}$$

Poniższy przykład łączy dwa podane powyżej:

$$\begin{array}{ccc} \lambda x.\lambda y.\underline{(\lambda z.(\lambda u.x\ y\ z\ u)(x\ y\ z))(\lambda w.w)} & \mapsto & \lambda x.\lambda y.(\lambda u.x\ y(\lambda w.w)u)(xy(\lambda w.w)) \\ \equiv_{dB} & & \equiv_{dB} \\ \lambda.\lambda.(\lambda.(\lambda.3\ 2\ 1\ 0)(2\ 1\ 0))(\lambda.0) & \mapsto & \lambda.\lambda.(\lambda.2\ 1(\lambda.0)0)(1\ 0(\lambda.0)) \end{array}$$

Dzięki użyciu kodów de Bruijna możemy wyeliminować użycie α -konwersji, ponieważ nigdy nie wystąpi konflikt pomiędzy zmiennymi związanymi. Mówiąc krótko, nie ma potrzeby zmieniania nazw zmiennych związanych, albowiem zmienne związane tak czy inaczej nie mają już nazw.

6.5. Przesunięcia. Jakkolwiek podana przed chwilą definicja $\sigma_n(M, N)$ będzie poprawnie funkcjonowała, gdy N jest zamknięta, może jednak nie działać gdy N będzie reprezentowało wyrażenie ze zmiennymi wolnymi. Mówiąc dokładniej, równość $\sigma_n(n, N) = N$ przestaje być prawdziwa gdy $n > 0$ i N reprezentuje wyrażenie ze zmiennymi wolnymi. Rozważmy bowiem następujący przykład, w którym redeks został

podkreślony:

$$\begin{aligned} \lambda x. \lambda y. (\lambda z. (\lambda u. z) z) (\lambda w. x \ y \ w) &\mapsto \lambda x. \lambda y. (\lambda u. \lambda w. x \ y \ w) (\lambda w. x \ y \ w) \\ &\equiv_{dB} \lambda. \lambda. (\lambda. (\lambda. 1) 0) (\lambda. 2 \ 1 \ 0) &\mapsto \lambda. \lambda. (\lambda. \lambda. 3 \ 2 \ 0) (\lambda. 2 \ 1 \ 0) \end{aligned}$$

Poprzednia definicja $\sigma_n(M, N)$ daje niewłaściwy rezultat, ponieważ $\sigma_1(1, \lambda. 2 \ 1 \ 0)$ daje $\lambda. 2 \ 1 \ 0$ zamiast $\lambda. 3 \ 2 \ 0$:

$$\begin{aligned} (\lambda. (\lambda. 1) 0) (\lambda. 2 \ 1 \ 0) &\mapsto \sigma_0((\lambda. 1) 0, \lambda. 2 \ 1 \ 0) \\ &= \sigma_0(\lambda. 1, \lambda. 2 \ 1 \ 0) \sigma_0(0, \lambda. 2 \ 1 \ 0) \\ &= (\lambda. \sigma_1(1, \lambda. 2 \ 1 \ 0)) \sigma_0(0, \lambda. 2 \ 1 \ 0) \\ &= (\lambda. \lambda. 2 \ 1 \ 0) (\lambda. 2 \ 1 \ 0) \\ &\neq (\lambda. \lambda. 3 \ 2 \ 0) (\lambda. 2 \ 1 \ 0) \end{aligned}$$

Aby zobaczyć, dlaczego w ogólności nie jest prawdziwy wzór $\sigma_n(n, N) = N$, przypomnijmy, że indeks n w $\sigma_n(n, N)$ oznacza liczbę λ -wiązań otaczających kod de Bruijna n :

$$\sigma_0(\underbrace{\lambda. \lambda. \dots \lambda.}_n, N) = \underbrace{\lambda. \lambda. \dots \lambda.}_n \sigma_n(n, N)$$

Wobec tego wszystkie kody de Bruijna w N odpowiadające zmiennym wolnym muszą być przesunięte o n po podstawieniu tak, aby poprawnie pomijały n λ -wiązań otaczających kod de Bruijna n . Na przykład:

$$\sigma_0(\underbrace{\lambda. \lambda. \dots \lambda.}_n, m) = \underbrace{\lambda. \lambda. \dots \lambda.}_n \sigma_n(n, m) = \underbrace{\lambda. \lambda. \dots \lambda.}_n m + n$$

(tutaj przez $m + n$ rozumiemy pojedynczy kod de Bruijna dodający m do n , nie zaś złożone wyrażenie de Bruijna składające się z m , n oraz $+$).

Oznaczmy przez $\tau^n(N)$ przesunięcie o n wszystkich kodów de Bruijna w N odpowiadających zmiennym wolnym. Możemy teraz zastosować wzór $\sigma_n(n, N) = \tau^n(N)$ w miejsce $\sigma_n(n, N) = N$. Częściowa definicja $\tau^n(N)$ dana jest przez warunki:

$$\begin{aligned} \tau^n(N_1 N_2) &= \tau^n(N_1) \tau^n(N_2) \\ \tau^n(m) &= m + n \end{aligned}$$

Pozostały przypadek $\tau^n(\lambda. N)$ nie może być, niestety, zdefiniowany indukcyjnie w oparciu o $\tau^n(N)$ na przykład jako $\tau^n(\lambda. N) = \lambda. \tau^n(N)$. Powodem tego jest fakt, że wewnątrz N nie każdy kod de Bruijna odpowiada zmiennej wolnej: 0 odpowiada λ -wiązanie w $\lambda. N$ i tym samym musi pozostać niezmienione.

Powyższa obserwacja sugeruje konieczność wprowadzenia jeszcze jednego indeksu "granicznego" pozwalającego rozpoznawać, czy dany kod de Bruijna odpowiada zmiennej wolnej, czy nie. Na przykład, jeśli wskaźnik graniczny dla $\lambda. N$ równy jest 0, powinien zwiększyć się o 1 wewnątrz N . Tym samym powinniśmy wprowadzić ogólniejszą notację $\tau_i^n(N)$ dla przesunięcia o n wszystkich kodów de Bruijna wewnątrz N odpowiadających zmiennym wolnym, gdzie kod de Bruijna m w N dla $m < i$ nie liczy się jako zmienna wolna. Formalnie definicja $\tau_i^n(N)$ będzie następująca:

$\tau_i^n(N_1 N_2)$	$=$	$\tau_i^n(N_1) \tau_i^n(N_2)$	
$\tau_i^n(\lambda. N)$	$=$	$\lambda. \tau_{i+1}^n(N)$	
$\tau_i^n(m)$	$=$	$m + n$	jeśli $m \geq i$
$\tau_i^n(m)$	$=$	m	jeśli $m < i$.

Tym samym kompletna definicja $\sigma_n(M, N)$ będzie następująca:

$$\begin{aligned}
\sigma_n(M_1 M_2, N) &= \sigma_n(M_1, N) \sigma_n(M_2, N) \\
\sigma_n(\lambda.M, N) &= \lambda.\sigma_{n+1}(M, N) \\
\sigma_n(m, N) &= m && \text{jeśli } m < n, \\
\sigma_n(n, N) &= \tau_0^n(N) \\
\sigma_n(m, N) &= m - 1 && \text{jeśli } m > n.
\end{aligned}$$

Wracając do omawianego wcześniej przykładu widzimy, że teraz $(\lambda.(\lambda.1)0)(\lambda.2 \ 1 \ 0)$ zredukuje się poprawnie:

$$\begin{aligned}
\underline{(\lambda.(\lambda.1)0)(\lambda.2 \ 1 \ 0)} &\mapsto \sigma_0((\lambda.1)0, \lambda.2 \ 1 \ 0) \\
&= \sigma_0(\lambda.1, \lambda.2 \ 1 \ 0) \sigma_0(0, \lambda.2 \ 1 \ 0) \\
&= (\lambda.\sigma_1(1, \lambda.2 \ 1 \ 0)) \sigma_0(0, \lambda.2 \ 1 \ 0) \\
&= (\lambda.\tau_0^1(1, \lambda.2 \ 1 \ 0)) \tau_0^0(0, \lambda.2 \ 1 \ 0) \\
&= (\lambda.\lambda.\tau_1^1(2 \ 1 \ 0)) \lambda.\tau_1^0(2 \ 1 \ 0) \\
&= (\lambda.\lambda.(2 + 1)(1 + 1)0)(\lambda.(2 + 0)(1 + 0)0) \\
&= (\lambda.\lambda.3 \ 2 \ 0)(\lambda.2 \ 1 \ 0)
\end{aligned}$$

Za każdym razem, gdy konwertujemy zwykłe wyrażenie e ze zmiennymi wolnymi $x_0, x_1, \dots, x_n$ do wyrażenia de Bruijna, możemy zamiast tego skonwertować $\lambda x_0. \lambda x_1. \dots \lambda x_n. e$, co efektywnie przypisze kody de Bruijna $0, 1, \dots, n$ zmiennym $x_0, x_1, \dots, x_n$, odpowiednio. Wówczas możemy myśleć o redukowaniu e jak o redukowaniu $\lambda x_0. \lambda x_1. \dots \lambda x_n. e$, gdzie n λ -wiązań jest ignorowanych. Tym sposobem możemy wykorzystywać kody de Bruijna do redukowania wyrażeń ze zmiennymi wolnymi.