

Andrzej Nowicki

O INKLUZYWNEJ TEORII KWANTYFIKACJI

W pierwszej części pracy omówione zostały założenia tkwiące u podstaw inkluzywnej teorii kwantyfikacji oraz pewne własności tej teorii. Dyskutowane są związki między inkluzywną i standardową teorią kwantyfikacji, przedstawione są również niektóre argumenty przemawiające za, bądź też przeciw stosowaniu teorii inkluzywnej w miejsce teorii standardowej. W części drugiej omówiono różne sformułowania inkluzywnej teorii kwantyfikacji: A. Mostowskiego¹, T. Hailperina², W. V. O. Quine'a³ oraz R. K. Meyera i K. Lamberta⁴. W trzeciej części pracy zaprezentowana została pewna koncepcja inkluzywnej teorii predykatów, oparta, podobnie jak teoria Meyera i Lamberta, na idei rozważania w języku pierwszego rzędu nazw nieoznaczających. Zdefiniowano semantycznie cztery relacje konsekwencji i pokazano, w jaki sposób założenia dotyczące realizowania się w dziedzinach nazw indywiduowych języka warunkują postać teorii kwantyfikacji. Omawiane relacje konsekwencji zostały również scharakteryzowane aksjomatycznie, przy czym podane aksjomatyki różnią się jedynie sposobem potraktowania formuł postaci $\forall x \alpha \rightarrow \alpha \left[\frac{x}{t} \right]$, gdzie t jest nazwą (termem) języka pierwszego rzędu.

¹ A. Mostowski, On the rules of proof in the pure functional calculus of first order, "Journal of Symbolic Logic" 1951, vol. 16, s. 107-111.

² T. Hailperin, Quantification theory and empty individual domain, "Journal of Symbolic Logic" 1953, vol. 18, s. 197-200.

³ W. V. O. Quine, Quantification and the empty domain, "Journal of Symbolic Logic" 1954, vol. 19, s. 177-179.

⁴ R. K. Meyer, K. Lambert, Universally free logic and standard quantification theory, "Journal of Symbolic Logic" 1968, vol. 33, s. 8-26.

I

Standardowe sformułowania teorii kwantyfikacji nie uwzględniają interpretacji języka pierwszego rzędu w pustych dziedzinach indywiduów. W konsekwencji formułą logicznie poprawną nazywa się każdą formułę, która prawdziwa jest przy dowolnej interpretacji języka w niepustej dziedzinie obiektów. Wśród tak zdefiniowanych formuł logicznie poprawnych znajdują się jednak formuły których nie można uznać za prawdziwe w dziedzinie pustej, np.:

$$\forall xP(x) \rightarrow \exists xP(x), \exists xP(x) \vee \exists x \neg P(x), \exists x(P(x) \rightarrow P(x))$$

(P jest symbolem predykatu 1-arg.). Jeśli więc interpretacje w pustej dziedzinie obiektów chcielibyśmy potraktować na równych prawach z innymi interpretacjami, to powinniśmy zmienić zakres pojęcia formuły logicznie poprawnej. Można zatem pokusić się o skonstruowanie teorii, której twierdzenia będą prawdziwe w każdej interpretacji, włączając dziedzinę pustą. Teorię tak pomyślaną nazywa się inkluzywną teorią kwantyfikacji.

Fakt, że standardowa teoria kwantyfikacji stosuje się wyłącznie do dziedzin niepustych jest powszechnie znany. Prowadzić on może do przekonania, że istnienie przynajmniej jednego obiektu jest postulatem ontologicznym rachunku predykatów. Zatrzymajmy się chwilę nad tym stwierdzeniem. Otóż wspomniany postulat ontologiczny głoszący istnienie przynajmniej jednego obiektu nie musi być i chyba nie jest organiczną częścią rachunku predykatów, jeśli tylko przez rachunek predykatów rozumieć będziemy nie standardową teorię kwantyfikacji, lecz teorię mającą na celu wyjaśnienie reguł rządzących posługiwaniem się kwantyfikatorami. Bez tego postulatu istnienia obywateli się inkluzywna teoria kwantyfikacji. Teoria standardowa, przyjmując założenie o istnieniu przynajmniej jednego obiektu w rozpatrywanych dziedzinach, nie czyni tego, jak się wydaje, dlatego, że kwantyfikacja w dziedzinie pustej jest nie do pomyślenia, lecz jedynie w celu uproszczenia semantyki (trudności związane z włączeniem dziedziny pustej będą omówione później). Należy zauważyć, że dyskutowany postulat istnienia nie powinien być utożsamiany z metafizycznym założeniem o istnieniu przedmiotów w ogóle, dotyczy on bowiem istnienia obiektów w dziedzinach dopuszczalnych interpretacji języka pierwszego rzędu. W ten właśnie sposób trzeba rozumieć termin "postulat ontologiczny".

Ponieważ inkluzywna teoria kwantyfikacji odrzuca postulat o istnieniu przynajmniej jednego obiektu, więc tym samym oparta jest ona na innych założeniach niż standardowa teoria kwantyfikacji. To przede wszystkim należy mieć na uwadze przy porównywaniu obu teorii. Zauważmy, że zdanie stwierdzające istnienie przynajmniej jednego obiektu daje się sformułować w języku pierwszego rzędu (języku przedmiotowym). Może mieć ono postać formuły $\exists x1$, gdzie 1 jest symbolem zdania prawdziwego (jest spójnikiem O-argumentowym); jeśli nie chcemy mieć w języku takiego symbolu, to zamiast 1 możemy użyć pewnego (dowolnego) zdania będącego tautologią. Zdanie $\exists x1$ nie jest tezą inkluzywnej teorii kwantyfikacji, jest natomiast tezą teorii standardowej. Ten ostatni fakt jest odzwierciedleniem dyskutowanego wcześniej (metajęzykowego) założenia o niepustości rozpatrywanych dziedzin.

Jednym z zadań inkluzywnej teorii predykatów jest ustalenie znaczenia kwantyfikacji w pustej dziedzinie obiektów. Zanotujmy następującą uwagę Quine'a⁵ dotyczącą tej kwestii: "Łatwy uzupełniający test umożliwia nam rozstrzygnięcie czy dana formuła jest prawdziwa w dziedzinie pustej. Powinniśmy jedynie wszystkie kwantyfikacje uniwersalne potraktować jako prawdziwe, natomiast egzystencjalne jako fałszywe i zastosować metodę tablic prawdziwościowych". Dodajmy tutaj, że test Quine'a dotyczy formuł domkniętych; w jego ujęciu teorii kwantyfikacji jedynie formuły domknięte mogą być twierdzeniami. Status zmiennych indywidualnych wolnych i związane z tym zagadnienie prawdziwości formuł otwartych nie są przez Quine'a rozważane. Fakt ten podkreślamy tutaj dlatego, że pewne interesujące założenia dotyczące zmiennych i stałych indywidualnych stanowiły punkt wyjścia dla twórców tzw. uniwersalnie wolnej logiki - pewnej wersji inkluzywnej teorii kwantyfikacji, którą później omówimy.

Sformułowanie testu sprawdzającego prawdziwość formuły w dziedzinie pustej nie wyczerpuje wszystkich kwestii związanych z inkluzywną teorią kwantyfikacji. Quine zauważa: "Dowód twierdzenia mógłby składać się po prostu z dowodów w teorii standardowej i następującego potem sprawdzenia podaną metodą. Możemy jednak być zainteresowani ujrzeć proste i autonomiczne sformułowanie, nie takie, które składa się, jak powyższe, ze standardowej teorii i reguły okrawania"⁶. Oczywiście mówiąc o inkluzywnej teorii kwan-

⁵ Quine, op. cit.

⁶ Ibidem.

tyfikacji, takiego właśnie sformułowania oczekujemy - niezależnego od teorii standardowej, jednorodnego w potraktowaniu wszystkich interpretacji, łącznie z dziedziną pustą i posiadającego odrębną aksjomatykę.

Obok argumentów przemawiających za traktowaniem dziedziny pustej na równych prawach z innymi dziedzinami indywiduów, istnieją również argumenty przeciwne. Spotkać można m. in. następujące stanowisko: "zastrzeżenie, że dziedziny puste wyklucza się z rozważań, nie jest w istocie żadnym ograniczeniem, ponieważ nie może być pustych dziedzin indywiduów"⁷. Ze stanowiskiem tym trudno jest polemizować inaczej, jak również stanowczo twierdzić, że puste dziedziny indywiduów istnieją.

Niekiedy sądzi się, że trywialność zbioru pustego jest wystarczającym argumentem dla wykluczenia go z rozważań. W wielu teoriach matematycznych badających struktury pewnego określonego rodzaju, u góry przyjmuje się, że dziedzina każdej z rozważanych struktur jest zbiorem niepustym. Jednak teoria tak podstawowa, jak teoria kwantyfikacji, teoria operująca terminem "poprawność logiczna" nie powinna, jak się wydaje, ograniczać swoich zastosowań wyłącznie do dziedzin niepustych, tym bardziej, że włączenie dziedziny pustej zmienia zakres pojęcia "formuła logicznie poprawna".

Inkluzywna teoria kwantyfikacji stosując się do większej ilości struktur, będąc teorią ogólniejszą niż standardowa teoria kwantyfikacji, posiada tym samym mniej tez. Dlatego też można by przeciwko niej wysunąć następujący zarzut: "Włączenie dziedziny pustej oznaczałoby będzie wyrzeczenie się pewnej ilości formuł które prawdziwe są wszędzie indziej, są więc ogólnie użyteczne"⁸. Zatrzymajmy się nieco dłużej nad tym stanowiskiem i rozpatrzmy pewne rozumowanie pozostające z nim w ścisłym związku. Przyjmijmy, że mamy do czynienia z czystym językiem predykatów (bez symboli stałych indywiduowych i symboli funkcji) i ponadto, że interesują nas jedynie zdania (formuły domknięte) tego języka. Oznaczmy symbolami T_{st} i T_{inkl} zbiory tez standardowej i inkluzywnej teorii kwantyfikacji. Symbolem F_0 oznaczmy zbiór wszystkich zdań fałszywych w dziedzinie pustej; zbiór F_0 może być wyznaczony

⁷ W. I. M. Kneale, The Development of Logic, Oxford 1962, s. 706-707, cyt. za: G. Hunter, Metalogic, MacMillan, London 1971, s. 206.

⁸ Quine, op. cit.

dzięki istnieniu testu sprawdzającego prawdziwość formuły w dziedzinie pustej (test ten cytowaliśmy wcześniej). Zachodzi wówczas równość $T_{inkl} = T_{st} \setminus F_0$. Zbiór tez teorii inkluzywnej może być więc otrzymany ze zbioru tez teorii standardowej. Teoria inkluzywna wydaje się zatem być teorią wtórną, zaś teoria standardowa mieć znaczenie podstawowe; dysponując nią bowiem możemy otrzymać teorię inkluzywną, a ponadto mamy w dalszym ciągu pod ręką zbiór T_{st} - zbiór formuł "ogólnie użytecznych". Podobne rozumowanie można jednak przeprowadzić dla wykazania, że być może jest przeciwnie, że to właśnie teoria inkluzywna ma znaczenie podstawowe. Należy mianowicie pokazać, że $T_{st} = \{\alpha : (\exists x_1 \rightarrow \alpha) \in T_{inkl}\}$ albo też, że $T_{st} = Cn_{inkl}(\exists x_1)$, gdzie Cn_{inkl} oznacza operację konsekwencji inkluzywnej teorii kwantyfikacji. Tak więc nie jest prawdą że w ramach teorii inkluzywnej utracimy zbiór T_{st} - formuł "ogólnie użytecznych". Zbiór ten można tu również zdefiniować semantycznie - przez ograniczenie się do interpretacji o niepustych dziedzinach. Na marginesie zauważmy, że porównując inkluzywną i standardową teorię kwantyfikacji, formułami ogólnie użytecznymi należałoby raczej nazwać zbiór tez teorii inkluzywnej a nie standardowej.

Wspomnieliśmy już o występującym w wielu teoriach matematycznych założeniu, że dziedzina każdej z rozważanych struktur jest zbiorem niepustym. Nie dzieje się tak bez powodu. Otóż w badaniu struktur pewnego ustalonego typu stosuje się niekiedy język pierwszego rzędu zadany przez typ rozpatrywanych struktur. Przykładowo, dla opisu algebr typu zadanego przez pewien zbiór symboli reprezentujących działania w tych algebrach, definiuje się język pierwszego rzędu, którego specyficznymi symbolami są wspomniane symbole funkcyjne i symbol identyczności. Wprowadzając taki język dla opisu struktur danego typu, wykorzystuje się przy tym standardową teorię kwantyfikacji. Jednak teoria ta stosuje się tylko do dziedzin niepustych; fakt ten jest przyczyną przyjmowania w definicji struktury, iż dziedzina każdej z rozważanych struktur jest zbiorem niepustym. Dlatego też chcąc opisywać struktury zarówno przy użyciu specyficznego dla nich języka pierwszego rzędu (języka przedmiotowego, języka "wewnętrznego"), jak i "z zewnątrz" (w metajęzyku), zmuszeni jesteśmy wykluczyć struktury o dziedzinie pustej albo też zastosować inkluzywną teorię kwantyfikacji w miejsce teorii standardowej. To drugie rozwiązanie wydaje się bardziej naturalne. Jest ono również korzystne z tech-

nicznego punktu widzenia, gdyż wykluczenie dziedziny pustej prowadzi niekiedy do pewnych niepotrzebnych komplikacji teorii. Pokażemy to na przykładzie zaczerpniętym z algebry uniwersalnej.

Przyjmuje się, że nośnik każdej algebry jest zbiorem niepustym. Dla ustalonej algebry $\mathcal{U}$ rozważmy rodzinę wszystkich jej podalgebr, niech $\text{Sub}(\mathcal{U})$ będzie zbiorem nośników tych podalgebr (każda podalgebra jest z definicji algebrą, a więc ma nośnik niepusty). Jeśli w sygnaturze algebry $\mathcal{U}$ nie występują symbole operacji 0-argumentowych, to położmy $\overline{\text{Sub}}(\mathcal{U}) = \text{Sub}(\mathcal{U}) \cup \{\emptyset\}$, w przeciwnym wypadku niech $\overline{\text{Sub}}(\mathcal{U}) = \text{Sub}(\mathcal{U})$. Struktura $(\overline{\text{Sub}}(\mathcal{U}), \subseteq)$ ma bardzo interesujące własności i pełni ważną rolę w algebrze uniwersalnej. Jest ona kratą algebraiczną, zaś $\overline{\text{Sub}}(\mathcal{U})$ jest (algebraicznym) systemem domknięć⁹, podczas gdy rodzina $\text{Sub}(\mathcal{U})$ nie musi być nawet domknięta na operację brania iloczynu dwóch zbiorów należących do $\text{Sub}(\mathcal{U})$. Gdyby jednak zmienić definicję algebry w ten sposób, żeby nie wykluczać przypadku, że nośnik algebry może być zbiorem pustym, to nie musielibyśmy później w sposób sztuczny wzbogacać rodziny $\text{Sub}(\mathcal{U})$ i definiować zbioru $\overline{\text{Sub}}(\mathcal{U})$. Rolę struktury $(\overline{\text{Sub}}(\mathcal{U}), \subseteq)$ pełniłaby wówczas struktura $(\text{Sub}(\mathcal{U}), \subseteq)$. Wyjaśnienia wymagać może jeszcze następująca kwestia. Otóż mimo tego iż w definicji algebry nie żądaliśmy, żeby nośnik był niepusty, to jednak w przypadku występowania w sygnaturze pewnej algebry symboli operacji 0-argumentowych (symbole stałych), nośnik tej algebry musiałby być zbiorem niepustym.

Podany przykład sugeruje, że niekiedy korzystnie byłoby przyjąć, że struktury mogą mieć pustą dziedzinę, jeśli tylko pozwalają na to inne warunki nakładane na rozpatrywane struktury. Jednak odrzucenie założenia o tym, że dziedzina struktury jest zbiorem niepustym musi iść w parze z przyjęciem inkluzywnej teorii kwantyfikacji w języku pierwszego rzędu stosowanego do opisu struktur rozważanego typu.

II

Dotychczas kilkakrotnie posługiwaliśmy się pojęciem interpretacji, obecnie uściślimy to pojęcie. Przyjmijmy, że $\mathcal{L}$ jest czystym (bezfunkcyjnym) językiem pierwszego rzędu. Strukturą (dla

⁹ G. Grätzer, Universal Algebra. Van Nostrand, Princeton, 1968, s. 47.

$\mathcal{L}$) nazywać będziemy dowolny zbiór M wraz z funkcją przyporządkowującą każdemu symbolowi n -argumentowego predykatu języka $\mathcal{L}$ pewną n -argumentową relację w zbiorze M . Nazwijmy waluacją w strukturze $\mathcal{A}$ dowolną funkcję częściową ze zbioru Var (zmiennych języka $\mathcal{L}$) w zbiór M (dziedzinę struktury $\mathcal{A}$). Interpretacja (dla $\mathcal{L}$) będzie parą złożoną ze struktury $\mathcal{A}$ i waluacji w tej strukturze. Spełnianie formuły definiowane jest w odniesieniu do interpretacji, zaś prawdziwość formuły w odniesieniu do struktury.

W większości sformułowań standardowej teorii kwantyfikacji przyjmuje się, że każda waluacja wartościuje wszystkie zmienne języka $\mathcal{L}$, czyli jest funkcją pełną ze zbioru Var w zbiór M (waluacje takie nazywać będziemy waluacjami pełnymi, zaś interpretacje przez nie wyznaczone - interpretacjami pełnymi). Zauważmy, że w dziedzinie pustej nie istnieją waluacje pełne, dlatego też prawdziwość formuły w strukturze o dziedzinie pustej nie może być określona na podstawie pojęcia spełniania dla interpretacji pełnych.

W inkluzywnej teorii kwantyfikacji musimy więc posługiwać się waluacjami częściowymi; w strukturze o dziedzinie pustej istnieje bowiem tylko jedna waluacja, mianowicie waluacja nigdzie nie określona. Jednak również w teorii standardowej zamiast waluacji pełnych rozważać można waluacje częściowe. Jest znanym faktem, że dla ustalonej formuły α języka $\mathcal{L}$, struktury $\mathcal{A}$ i waluacji pełnej s w tej strukturze, spełnianie formuły α w interpretacji pełnej $(\mathcal{A}, s)$ zależy jedynie od tego, jakie wartości przyjmuje waluacja s na zbiorze $\text{Var}(\alpha)$ (zmiennych wolnych formuły α); wartościowanie przez waluację s pozostałych zmiennych nie jest istotne. Dlatego też definicję prawdziwości formuły α w strukturze $\mathcal{A}$ oprócz można wyłącznie na takich interpretacjach $\mathcal{I} = (\mathcal{A}, s)$, w których s jest funkcją częściową, określoną na zbiorze $\text{Var}(\alpha)$. Dla formuł domkniętych wystarczy ograniczyć się do waluacji, która nie jest określona dla żadnej zmiennej. Znajduje to wyraz w następującej uwadze A. Tarskiego¹⁰, "Co się tyczy pojęcia prawdy to należy zauważyć, że - na gruncie powyższej koncepcji - zdanie, tj. funkcję bez zmiennych wolnych, spełniać może jeden tylko ciąg, mianowicie ciąg »pusty«, nie posiadający ani jednego wyrazu; prawdziwymi wypadnie więc nazwać takie zdania, które ciąg »pusty« istotnie spełnia".

¹⁰ A. T a r s k i, Pojęcie prawdy w językach nauk dedukcyjnych, Warszawa 1933.

Przedstawioną powyżej koncepcję waluacji częściowych zastosował Mostowski¹¹ w semantyce dla inkluzywnej teorii kwantyfikacji. Dla dowolnej struktury definiuje on prawdziwość formuły w tej strukturze, określa zbiór formuł logicznie poprawnych (tj. prawdziwych we wszystkich strukturach), podaje aksjomatykę zbioru tez teorii inkluzywnej i udowadnia twierdzenie o pełności, wykorzystując przy tym twierdzenie o pełności dla teorii standardowej

Kilka kwestii związanych z teorią Mostowskiego zasługuje na uwagę. Pierwsza dotyczy tzw. pustego wiązania kwantifikatorowego. Hailperin¹² zauważa, że w następstwie przyjętej przez Mostowskiego definicji prawdziwości, formułę $\forall x \exists y (P(y) \wedge \neg P(y))$ należy uznać za fałszywą w strukturze z dziedziną pustą (tak samo inne formuły $\forall x \alpha$, gdzie α jest fałszywa i x nie występuje w α). Pokazuje on, że teorię bardziej elegancką otrzymamy przyjmując, że każda formuła postaci $\forall x \alpha$ jest prawdziwa w dziedzinie pustej. Wymaga to innego niż uczynił to Mostowski, bardziej naturalnego, jak się wydaje, sformułowania warunku spełniania dla formuły $\forall x \alpha$. Quine¹³ podziela stanowisko Hailperina, ale jego argumentacja jest bardziej stanowcza; wykazuje, że puste wiązania kwantifikatorowe nie tylko możemy, lecz musimy traktować inaczej niż uczynił to Mostowski. W inkluzywnej teorii kwantyfikacji Hailperina i Quine'a formuły α i $\forall x \alpha$ (dla $x \notin \text{Var}(\alpha)$) nie muszą być równoważne (podobnie formuły α i $\exists x \alpha$, $x \notin \text{Var}(\alpha)$), natomiast w standardowej teorii kwantyfikacji i w teorii inkluzywnej Mostowskiego formuły te są równoważne: kwantifikator pusto wiążący może być pominięty w każdej formule.

W sformułowanej przez Mostowskiego inkluzywnej teorii kwantyfikacji prawdziwe i logicznie poprawne mogą być nie tylko formuły domknięte, ale również otwarte. W dziedzinie pustej każda formuła otwarta jest prawdziwa. Konsekwencją tego jest stwierdzenie, że reguła Modus Ponens nie zachowuje prawdziwości i poprawności logicznej. Żeby to pokazać weźmy pod uwagę dwie formuły α i β ; niech $\alpha = P(x) \vee \neg P(x)$, $\beta = \exists x (P(x) \vee \neg P(x))$. Formuły α i $\alpha \rightarrow \beta$ są prawdziwe w dziedzinie pustej (ponieważ są otwarte) i są one logicznie poprawne, natomiast formuła β jest w dziedzinie pustej fałszywa (dlatego też nie jest ona logicznie poprawna). Regułę α ,

¹¹ Mostowski, op. cit.

¹² Hailperin, op. cit.

¹³ Quine, op. cit.

$\alpha \rightarrow \beta / \beta$ można jednak zachować, nakładając pewne warunki na zmienne wolne formuł α i β ; żądając mianowicie, żeby $\text{Var}(\alpha) \subseteq \text{Var}(\beta)$. Tak ograniczoną regułę Modus Ponens stosuje Mostowski w aksjomatyce dla zbioru tez inkluzywnej teorii kwantyfikacji (tylko to ograniczenie różni aksjomatykę Mostowskiego od aksjomatyki Churcha dla teorii standardowej).

Omawiane problemy dotyczące prawdziwości formuł otwartych i ograniczenia stosowania reguły Modus Ponens nie występują w sformułowaniach inkluzywnej teorii kwantyfikacji przedstawionych w pracach Hailperina¹⁴ i Quine'a¹⁵. Sformułowania te są bowiem wzorowane na standardowej teorii kwantyfikacji w ujęciu Quine'a¹⁶, który przyjmuje, że twierdzeniami teorii mogą być jedynie formuły domknięte. Obie wspomniane prace koncentrują się wokół aksjomatyzacji zbioru tez inkluzywnej teorii kwantyfikacji; aksjomatyka Quine'a jest uproszczeniem aksjomatyki podanej przez Hailperina i jest nieznaczna tylko przeróbką aksjomatyki teorii standardowej, przedstawionej w drugim wydaniu książki Quine'a¹⁷. Układ aksjomatów dla inkluzywnej teorii kwantyfikacji, pokrewny aksjomatyce Hailperina i Quine'a, znaleźć można także w książce A. Grzegorzczaka¹⁸.

Inkluzywna teoria kwantyfikacji zgodna jest z koncepcją tzw. logik wolnych, które również odrzucają pewne założenia egzystencyjne tkwiące u podstaw standardowej teorii kwantyfikacji. "Lambert chce traktować wzrastające zainteresowanie się logikami wolnymi, jako oznakę takiego spojrzenia na logikę, które wyraża się w aforyzmie, że logika ma wstręt do istnienia" - czytamy w pracy Meyera i Lamberta¹⁹. Cechą charakterystyczną logik wolnych jest odrzucenie założeń dotyczących realizowania się w dziedzinach stałych indywiduowych języka. Nie żąda się, żeby każda stała indywiduowa była interpretowana jako nazwa obiektu istniejącego w dziedzinie realnej przedmiotów. Nie przesądza to oczywiście o

¹⁴ H a i l p e r i n, op. cit.

¹⁵ Q u i n e, op. cit.

¹⁶ W. V. O. Q u i n e, Mathematical logic. Harvard University Press, Cambridge (Mass.) 1940.

¹⁷ Q u i n e, Mathematical logic... oraz toż, wyd. II, Cambridge (Mass.) 1951.

¹⁸ A. G r z e g o r c z y k, Zarys logiki matematycznej, Warszawa 1981, s. 144-145.

¹⁹ M e y e r, L a m b e r t, op. cit.

kształcie teorii logicznej, gdyż możliwe są tu różne rozwiązania semantyczne. Jeśli dodatkowo zażądamy, żeby twierdzenia teorii były prawdziwe w każdej dziedzinie z włączeniem dziedziny pustej, to otrzymamy teorię będącą połączeniem logiki wolnej i inkluzywnej teorii kwantyfikacji. Teoria taka, przedstawiona została przez Meyera i Lamberta²⁰ i nazwana przez nich uniwersalnie wolną logiką. Oba postulaty, które legły u podstaw uniwersalnie wolnej logiki są niezależne, na co zwracają uwagę Meyer i Lambert jednak, jak piszą, naturalne jest żądanie uwzględnienia pustej dziedziny indywiduów w obecności postulatu dopuszczającego, żeby stałe indywiduowe nie realizowały się jako elementy dziedziny (dziedziny realnej - w terminologii Meyera i Lamberta).

Język pierwszego rzędu, w którym Meyer i Lambert prowadzą swoje rozważania zawiera symbol pewnego wyróżnionego predykatu: $E!$, " $E!x$ " czytamy jako " x istnieje". Każdy element dziedziny realnej spełnia relację odpowiadającą predykatowi $E!$. Dyskutując problem stałych indywiduowych, podając przykłady zdań języka naturalnego zawierających nazwy indywidualne, Meyer i Lambert nie rozpatrują jednak języka ze stałymi indywiduowymi; operują czystym językiem pierwszego rzędu, bez symboli stałych i symboli funkcji. Według autorów jest to tylko pozornie paradoks, ponieważ, jak piszą, rozważają oni czystą, a nie stosowaną logikę i wobec tego nie muszą odróżniać zmiennych od stałych indywiduowych, albo też mogą traktować zmienne wolne jako miejsce dla stałych indywiduowych.

Semantyka dla uniwersalnie wolnej logiki Meyera i Lamberta ma charakter wielostopniowy. Dla każdej interpretacji realnej określa się tzw. nominalne interpretacje, będące rozszerzeniem realnej interpretacji. Z kolei dla dowolnej nominalnej interpretacji definiuje się tzw. punkty logiczne - inne nominalne interpretacje o tej samej dziedzinie.

Interpretacja realna zadana jest przez: 1) dowolny zbiór M nazwany dziedziną realną (M może być zbiorem pustym), 2) waluację częściową w zbiór M , 3) przyporządkowanie symbolom predykatów relacji w zbiorze M . Ponieważ nie wszystkie zmienne muszą być interpretowane, więc nie wszystkie formuły atomowe mają określoną wartość (T lub F); spójniki i kwantyfikatory rozumiane są tak, jak w silnej trójwartościowej logice Kleene'go, w której oprócz

²⁰ Ibidem.

wartości T i F występuje wartość n - nieokreśloności bądź niezde-terminowania.

Interpretacja nominalna zadana jest przez: 1) dziedzinę realną M , 2) niepustą dziedzinę nominalną N taką, że $M \subseteq N$, 3) waluację pełną w zbiór N , 4) przyporządkowane każdemu symbolowi n -argumentowego predykatu n -argumentowej relacji częściowej (R^+, R^-) w zbiorze N ($R^+, R^- \subseteq N^n$, $R^+ \cap R^- = \emptyset$; R^+ i R^- nie muszą się dopełniać w N^n , żąda się jednak, żeby $M^n \subseteq R^+ \cup R^-$). Podobnie jak w realnej interpretacji, również w interpretacji nominalnej nie wszystkie formuły atomowe mają określoną wartość jednak powodem jest tu sposób w jaki interpretujemy predykaty, a nie zmienne indywiduowe. Działanie kwantyfikatorów ograniczone jest do zbioru M , zaś znaczenie spójników i kwantyfikatorów znów jest takie, jak w logice trójwartościowej Kleene'go.

Jeśli w punkcie 4) definicji interpretacji nominalnej zażądamy, żeby relacja częściowa (R^+, R^-) była relacją pełną (tzn. $R^+ \cup R^- = N^n$), to otrzymamy definicję punktu logicznego. W punkcie logicznym każda formuła ma określoną wartość (T lub F), można więc zastąpić logikę trójwartościową klasyczną logiką dwuwartościową.

Dla ustalonej interpretacji nominalnej $\mathcal{I}$ dowolną formułę nazywa się $\mathcal{I}$ -poprawną, gdy dla każdego punktu logicznego będącego uzupełnieniem interpretacji $\mathcal{I}$, formuła ta przyjmuje wartość T. Z kolei, dla ustalonej interpretacji realnej $\mathcal{I}$ formułę nazywa się $\mathcal{I}$ -poprawną, gdy jest ona $\mathcal{I}$ -poprawna dla każdej interpretacji nominalnej $\mathcal{I}$ będącej rozszerzeniem realnej interpretacji $\mathcal{I}$. Wreszcie, formułę nazywamy logicznie poprawną, gdy jest ona poprawna dla każdej realnej interpretacji. Pojęcia $\mathcal{I}$ -, $\mathcal{I}$ -poprawności oraz poprawności logicznej są niezupełne; mogą istnieć formuły takie, że ani one, ani ich negacje nie są ($\mathcal{I}$ -, $\mathcal{I}$ - bądź logicznie) poprawne.

Przyjęcie takiej właśnie semantyki dla uniwersalnie wolnej logiki jest wynikiem przeprowadzonej przez Meyera i Lamberta analizy prawdziwości różnych zdań języka naturalnego zawierających nazwy nieoznaczające (nazwy takie traktuje się jak stałe indywiduowe bądź zmienne, które nie realizują się w dziedzinie realnej). Rozpatrują one odmienne, niekiedy więc sprzeczne, stanowisko dotyczące prawdziwości takich zdań. Nie opowiadają się za żadnym z tych stanowisk; ich wielopoziomowa semantyka ma za zadanie stanowiska te pogodzić, wskazać, gdzie leży źródło nieporozumień. Staje się to możliwe dzięki temu, że w ich teorii funkcjonuje wiele pojęć prawdziwości i poprawności (prawdziwość formuły w interpretacji

realnej, prawdziwość w interpretacji nominalnej, prawdziwość w punkcie logicznym, poprawność w interpretacji realnej, poprawność w interpretacji nominalnej, logiczna poprawność).

Meyer i Lambert podają również aksjomatykę dla zbioru tez uniwersalnie wolnej logiki. Wyróżnioną i istotną rolę pełni w niej predykat $E!$. Zamiast spotykanego w standardowej teorii kwantyfikacji schematu aksjomatów: $\forall x\alpha \rightarrow \alpha[\frac{x}{y}]$ (o ile x jest wolne dla y w α), w uniwersalnie wolnej logice przyjęli oni schemat: $\forall x\alpha \rightarrow (E!y \rightarrow \alpha[\frac{x}{y}])$ (gdy x jest wolne dla y w α). Predykat $E!$ jest ponadto scharakteryzowany aksjomatem $\forall xE!(x)$. W przedstawionym przez Meyera i Lamberta dowodzie twierdzenia o pełności, wykorzystuje się pełność teorii standardowej oraz podobieństwo aksjomatyki. W uniwersalnie wolnej logice, inaczej niż w inkluzywnej teorii kwantyfikacji Hailperina i Quine'a, twierdzeniami teorii mogą być zarówno formuły domknięte, jak i otwarte; reguła Modus Ponens nie podlega jednak żadnym ograniczeniom, jak to było w teorii Mostowskiego. Pełna lista schematów aksjomatów i reguł uniwersalnie wolnej logiki przedstawia się następująco:

101. Jeśli A jest formułą tautologiczną, to A jest aksjomatem.
102. $\forall x(A \rightarrow B) \rightarrow (A \rightarrow \forall xB)$, gdzie $x \notin \text{Var}(A)$.
103. $\forall xA(x) \rightarrow (E!(y) \rightarrow A(y))$; $A(y)$ powstaje z $A(x)$ przez podstawienie zmiennej y w miejsce każdego wolnego występowania zmiennej x ; przyjmujemy, że x jest wolne dla y w $A(x)$.
104. $\forall xE!(x)$.
105. Z A i $A \rightarrow B$ wnioskuj B .
106. Z A wnioskuj $\forall xA$.
107. $\forall x(A \rightarrow B) \rightarrow (\forall xA \rightarrow \forall xB)$.

Przyjęto, że w języku występuje implikacja, negacja, kwantyfikator ogólny, zmienne i symbole predykatów. Reguła 106 może być ograniczona jedynie do aksjomatów, natomiast schemat 102 może być uproszczony do postaci $A \rightarrow \forall xA(x \notin \text{Var}(A))$.

Istotą zaproponowanego przez Meyera i Lamberta podejścia do teorii kwantyfikacji jest sposób potraktowania zmiennych. Zmienne wolne, na które można patrzeć jak na stałe indywidualne (nazwy indywidualne) nie muszą realizować się w dziedzinie interpretacji realnej, zaś zmienne związane kwantyfikatorami przebiegają tylko zbiór obiektów realnych - dziedzinę realną. Dlatego też formuły $\forall xP(x) \rightarrow P(x)$, $P(x) \rightarrow \exists xP(x)$ nie są tezami systemu. Mostowski, który również dopuszczał, żeby formuły otwarte były prawdziwe lub fałszywe, przyjął inny punkt widzenia. W każdej formule otwartej

interpretował on zmienne wolne jako nazwy obiektów dziedziny realnej; dlatego w jego teorii wymienione formuły są tezami.

III

Przedstawimy teraz pewną koncepcję inkluzywnej teorii kwantyfikacji, która ma wiele cech wspólnych z teorią Meyera i Lamberta, jest jednak mniej rozbudowana w warstwie semantycznej. U podstaw tej teorii leżą uboższe założenia ontologiczne, postuluje się bowiem istnienie jedynie dziedziny realnej obiektów, podczas gdy semantykę uniwersalnie wolnej logiki Meyera i Lamberta buduje się na podstawie pojęć dziedziny realnej i dziedziny nominalnej. Zamiast trzech różnych definicji interpretacji, które występują w teorii Meyera i Lamberta (interpretacji realnej, interpretacji nominalnej i punktu logicznego) będziemy posługiwać się interpretacjami jednego rodzaju, określonymi nad dziedziną realną. Nie utracimy jednak korzyści, które w uniwersalnie wolnej logice płyną w wielości typów rozważanych interpretacji i możliwości posługiwania się kilkoma pojęciami prawdziwości lub poprawności. Intuicje Meyera i Lamberta dotyczące różnych sposobów rozumienia prawdziwości (bądź poprawności) mogą być odtworzone w prezentowanej tu teorii jako pewne przypadki określenia prawdziwości w różnego typu strukturach. Nie będziemy szerzej dyskutować tego problemu, gdyż naszym celem jest jedynie sformułowanie teorii kwantyfikacji, a więc określenie logiki, przez którą rozumiemy operację konsekwencji, a tę można zdefiniować na podstawie pojęcia interpretacji.

Przyjmijmy ogólnie, że każda interpretacja (i struktura również) zadana jest nad pewną algebrą częściową i zdefiniujemy dwie inkluzywne operacje konsekwencji semantycznej: jedna określona będzie przez klasę wszystkich interpretacji zadanych nad algebrami częściowymi, druga natomiast przez interpretacje zadane nad algebrami. Obie operacje konsekwencji zostaną scharakteryzowane aksjomatycznie.

Nie będziemy zakładać, że w rozważanym przez nas języku istnieją wyróżnione symbole predykatów: istnienia bądź identyczności. Wymienione predykaty można scharakteryzować semantycznie i aksjomatycznie, jednak, inaczej niż w uniwersalnie wolnej logice przedstawiona przez nas aksjomatyka dla inkluzywnej teorii kwantyfikacji jest wolna od uwikłania wyróżnionych predykatów.

Przypuśćmy, że pewna interpretacja $\mathcal{I}$ jest zadana nad algebrą częściową $\mathcal{M}$. Weźmy pod uwagę zdanie $P(e)$, gdzie P jest symbolem predykatu 1-argumentowego, natomiast e jest stałą indywiduową rozważanego języka pierwszego rzędu (e może być traktowane jako symbol operacji 0-argumentowej). Zastanówmy się nad sposobem określenia prawdziwości zdania $P(e)$ przy interpretacji $\mathcal{I}$. W interpretacji $\mathcal{I}$ symbol P realizuje się jako pewna relacja 1-argumentowa na zbiorze M - nośniku algebry $\mathcal{M}$; oznaczmy tę relację $P^{\mathcal{I}}$. Ponieważ $\mathcal{M}$ jest algebrą częściową, każdy symbol operacji n -argumentowej realizuje się w $\mathcal{M}$ jako działanie częściowe n -argumentowe w zbiorze M ; działanie odpowiadające symbolowi f oznaczamy będziemy $f^{\mathcal{M}}$. W szczególności mamy do czynienia z działaniami częściowymi 0-argumentowymi, które albo są określone jako element dziedziny M , albo też nie są określone w ogóle. Ustalenie prawdziwości zdania $P(e)$ w przypadku, gdy $e^{\mathcal{M}}$ jest określone (e jest elementem zbioru M) polega na sprawdzeniu, czy $e^{\mathcal{M}} \in P^{\mathcal{I}}$. Jednak w przypadku gdy e nie realizuje się w $\mathcal{M}$ jako element zbioru M (tzn. $e^{\mathcal{M}}$ nie jest określone), prawdziwość bądź fałszywość zdania $P(e)$ (w interpretacji $\mathcal{I}$) musi być ustalona w inny sposób. Przyjmujemy, że wartość logiczna zdania $P(e)$ zadana jest wówczas przez pewną konwencję określoną w interpretacji $\mathcal{I}$.

Konwencje, o których mowa, powinny umożliwiać m. in. ustalenie wartości logicznej wszystkich zdań postaci $P(n)$, gdzie n jest dowolną nazwą nieoznaczającą rozważanego języka. Na konwencje nałożymy tylko jeden warunek, który w tym miejscu spróbujemy wyjaśnić na przykładzie. Weźmy pod uwagę dwie nazwy postaci $f(e_1)$ i $f(e_2)$ (f jest symbolem funkcji 1-argumentowej, natomiast e_1 i e_2 są symbolami funkcji 0-argumentowych, f możemy więc traktować jako funktor nazwotwórczy od jednego argumentu nazwowego, zaś e_1 , e_2 jako nazwy). Przypuśćmy, że w algebrze częściowej $\mathcal{M}$ działania $e_1^{\mathcal{M}}$ i $e_2^{\mathcal{M}}$ są określone (są to elementy zbioru M) i ponadto $e_1^{\mathcal{M}} = e_2^{\mathcal{M}}$, natomiast $f^{\mathcal{M}}$ (funkcja częściowa 1-argumentowa) nie jest określona na elemencie $e_1^{\mathcal{M}}$. Obie nazwy $f(e_1)$ i $f(e_2)$ są w $\mathcal{M}$ nazwami nieoznaczającymi, przyjmujemy jednak, że nie są one odróżnialne ze względu na jakąkolwiek konwencję obowiązującą w dowolnej interpretacji $\mathcal{I}$ nad algebrą częściową $\mathcal{M}$.

Przystąpmy obecnie do sformalizowania przedstawionych wyżej intuicji.

Dla dowolnego zbioru X , X^* oznaczać będzie zbiór wszystkich ciągów skończonych elementów zbioru X . Ciąg jednoelementowy, któ-

rego elementem jest a oznaczać będziemy również symbolem a , co nie powinno prowadzić do nieporozumień. Złożenie dwóch ciągów $\alpha, \beta \in X^*$ oznaczamy $\alpha\beta$. Typem funkcyjnym nazywać będziemy parę (F, o) , gdzie F jest dowolnym zbiorem, natomiast o jest funkcją ze zbioru F w zbiór liczb naturalnych. Położmy $F_n = \{f \in F : o(f) = n\}$; F_n nazywać będziemy zbiorem symboli funkcji n -argumentowych, F_0 nazywamy zbiorem stałych. Typ relacyjny definiujemy podobnie jak typ funkcyjny.

Algebrą częściową typu $\tau = (F, o)$ nazywamy zbiór M wraz z funkcją φ określoną na zbiorze F : jeśli $f \in F_n$, to $\varphi(f)$ jest działaniem częściowym n -argumentowym w zbiorze M . Jeśli $\mathcal{M} = (M, \varphi)$ jest algebrą częściową i $f \in F$, to zamiast $\varphi(f)$ będziemy pisać $f^{\mathcal{M}}$. Algebrę częściową nazywamy algebrą, jeśli dla każdego $f \in F$, $f^{\mathcal{M}}$ jest działaniem wszędzie określonym. Nie wykluczamy przypadku, że nośnik algebry lub algebry częściowej jest zbiorem pustym.

Językiem termów nazywamy parę (Var, τ) , gdzie $\tau = (F, o)$ jest typem, zaś Var - zbiorem takim że $\text{Var} \cap F = \emptyset$. Zbiór termów dla języka termów (Var, τ) oznaczamy $\mathcal{T}(\text{Var}, \tau)$, (termy określamy jako ciągi: $\mathcal{T}(\text{Var}, \tau) \subseteq (\text{Var} \cup F)^*$; przyjmujemy tu znaną definicję zbioru $\mathcal{T}(\text{Var}, \tau)$). Zbiór Var nazywamy zbiorem zmiennych dla języka termów (Var, τ) . Dla termu $t \in \mathcal{T}(\text{Var}, \tau)$, $\text{Var}(t)$ oznaczać będzie zbiór wszystkich zmiennych występujących w t (t jest ciągiem). Będziemy pisać $t = t(x_0, \dots, x_{n-1})$ dla wyrażenia faktu, że $\text{Var}(t) \subseteq \{x_0, \dots, x_{n-1}\}$ i $x_0, \dots, x_{n-1}$ są wzajemnie różne.

Niech (Var, τ) będzie językiem termów, natomiast $\mathcal{M} = (M, \varphi)$ algebrą częściową typu τ . W aluacjā w $\mathcal{M}$ nazywamy dowolną funkcję częściową ze zbioru Var w zbiór M . Zbiór wszystkich waluacji w $\mathcal{M}$ oznaczamy $\text{Val}_{\mathcal{M}}(\text{Var})$. Jeśli $s \in \text{Val}_{\mathcal{M}}(\text{Var})$, to $\mathcal{D}_m(s)$ oznaczać będzie dziedzinę waluacji s . Niech $s \in \text{Val}_{\mathcal{M}}(\text{Var})$, $x \in \text{Var}$, $a \in M$; wówczas $s[\frac{x}{a}]$ jest waluacją zdefiniowaną następująco: dla $y \in \text{Var}$: $s[\frac{x}{a}](y) = \begin{cases} s(x), & \text{gdy } y \neq x \\ a, & \text{gdy } y = x. \end{cases}$

Dla $t \in \mathcal{T}(\text{Var}, \tau)$, $s \in \text{Val}_{\mathcal{M}}(\text{Var})$ zdefiniować można wartość termu t przy waluacji s , oznaczamy ją $t^{\mathcal{M}}, s$; $t^{\mathcal{M}}, s \in M$ lub $t^{\mathcal{M}}, s$ nie jest określone. Pominiemy tutaj definicję $t^{\mathcal{M}}, s$, którą łatwo jest zbudować. Nieokreśloność $t^{\mathcal{M}}, s$ może być wynikiem tego, że $f^{\mathcal{M}}$ (dla $f \in F$) są działaniami częściowymi lub też być konsekwencją tego, że s jest funkcją częściową.

Niech (V, τ) będzie językiem termów, natomiast $\mathcal{M} = (M, \varphi)$ algebrą częściową typu τ . Zdefiniujemy teraz zbiór $\mathcal{T}_{\mathcal{M}}(V)$, elementami którego będą zarówno elementy zbioru M , jak również nazwy nieoznaczające. Zbiór $\mathcal{T}_{\mathcal{M}}(V)$ jest zasadniczym elementem prezentowanej konstrukcji i posłuży nam później do określenia interpretacji. Elementy zbioru $\mathcal{T}_{\mathcal{M}}(V)$ są ciągami elementów ze zbioru $M \cup V \cup F$ (tak więc $\mathcal{T}_{\mathcal{M}}(V) \subseteq (M \cup V \cup F)^*$ i będą one budowane podobnie jak termy. Przyjmiemy dla uproszczenia, że zbiory M i $V \cup F$ są rozłączne, w przeciwnym wypadku każdy element zbioru M należałoby zaopatrzyć w indeks pozwalający właściwie go zinterpretować. Każdy element zbioru V będziemy traktować jako pewną nazwę nieoznaczającą.

Mając już teraz na uwadze późniejszą definicję interpretacji, wyjaśnimy, że zbiór V będzie tam występował jako część większego zbioru Var (tzn. $V \subseteq \text{Var}$); zmienne ze zbioru $\text{Var} \setminus V$, w odróżnieniu od zmiennych ze zbioru V , będą interpretowane jako elementy dziedziny (zbioru M).

- 1) Jeśli $a \in M$, to $a \in \mathcal{T}_{\mathcal{M}}(V)$.
- 2) Jeśli $x \in V$, to $x \in \mathcal{T}_{\mathcal{M}}(V)$.
- 3) Niech $f \in F_n$, $x_0, \dots, x_{n-1} \in \mathcal{T}_{\mathcal{M}}(V)$. Wówczas:
 - a) jeśli $x_0, \dots, x_{n-1} \in M$ i $f^{\mathcal{M}}(x_0, \dots, x_{n-1})$ jest określone, to $f^{\mathcal{M}}(x_0, \dots, x_{n-1}) \in \mathcal{T}_{\mathcal{M}}(V)$,
 - b) jeśli $x_0, \dots, x_{n-1} \in M$ i $f^{\mathcal{M}}(x_0, \dots, x_{n-1})$ nie jest określone, to $f x_0, \dots, x_{n-1} \in \mathcal{T}_{\mathcal{M}}(V)$,
 - c) jeśli $x_k \notin M$ dla pewnego k ($0 < k < n$), to $f x_0 \dots x_{n-1} \in \mathcal{T}_{\mathcal{M}}(V)$.

Zauważmy, że w definicji tej punkt 3a jest zbędny ze względu na punkt 1. Na zbiorze $\mathcal{T}_{\mathcal{M}}(V)$ określić można strukturę algebry typu τ - algebrę tę oznaczać będziemy symbolem $\underline{\mathcal{T}_{\mathcal{M}}}(V)$. Przyjmujemy że działania w algebrze $\underline{\mathcal{T}_{\mathcal{M}}}(V)$ zdefiniowane są w sposób określony w punkcie 3 definicji zbioru $\mathcal{T}_{\mathcal{M}}(V)$. Algebra częściowa $\mathcal{M}$ może być utożsamiana z relatywną podalgebrą algebry $\underline{\mathcal{T}_{\mathcal{M}}}(V)$. Połóżmy $\underline{\mathcal{T}_{\mathcal{M}}} = \underline{\mathcal{T}_{\mathcal{M}}}(\emptyset)$. Algebra $\underline{\mathcal{T}_{\mathcal{M}}} (= \underline{\mathcal{T}_{\mathcal{M}}}(\emptyset))$ ma bardzo interesujące własności (izomorficzna z nią algebra, będąca rozszerzeniem algebry częściowej $\mathcal{M}$, zdefiniowana jest przez G. Gratzera²¹). Zauważmy, że jeśli $\mathcal{M}$ jest algebrą, to $\mathcal{M}$ można utożsamiać z algebrą $\underline{\mathcal{T}_{\mathcal{M}}}$.

Niech $t \in \mathcal{T}(\text{Var}, \tau)$, $s \in \text{Val}_{\mathcal{M}}(\text{Var})$ i niech $\text{Var} \setminus \mathcal{M}(s) \subseteq V$. Zdefiniujemy $t^{<\mathcal{M}, s>} \in \mathcal{T}_{\mathcal{M}}(V)$.

²¹ G r ä t z e r, op. cit., s. 84-90.

- 1) Niech $x \in \text{Var}$. Wówczas $x^{<\mathcal{M}, s>} = \begin{cases} s(x), & \text{gdy } x \in \mathcal{M}(s) \\ x, & \text{gdy } x \in \text{Var} \setminus \mathcal{M}(s) \subseteq V. \end{cases}$
 - 2) Niech $f \in F_n$, $t_0, \dots, t_{n-1} \in \mathcal{T}(\text{Var}, \tau)$. Wówczas $(ft_0 \dots t_{n-1})^{<\mathcal{M}, s>} = f_{\mathcal{T}(\mathcal{M})(V)}(t_0^{<\mathcal{M}, s>}, \dots, t_{n-1}^{<\mathcal{M}, s>})$.
- Zauważmy, że $t^{<\mathcal{M}, s>} \in M$ wtw $t^{<\mathcal{M}, s>}$ jest określone (wówczas $t^{<\mathcal{M}, s>} = t^{<\mathcal{M}, s>}$).

Przejdźmy teraz do opisanego języka i zdefiniowania interpretacji. Niech $\mathcal{L}$ będzie językiem pierwszego rzędu zadany przez:

- 1) typ funkcyjny $\tau = (F, o_\tau)$,
- 2) typ relacyjny $\rho = (R, o_\rho)$,
- 3) zbiór Var - zmiennych indywidualnych wolnych (przyjmujemy, że Var jest zbiorem nieskończonym),
- 4) zbiór Vb - zmiennych indywidualnych związanych (Vb jest również zbiorem nieskończonym),
- 5) zbiór trzejelementowy $\{C, O, \forall\}$ (C nazywamy symbolem implikacji, O - symbolem zdania fałszywego, $\forall$ - symbolem kwantyfikatora ogólnego).

Przyjmujemy, że zbiory $F, R, \text{Var}, \text{Vb}, \{C, O, \forall\}$ są rozłączone; każda formuła języka $\mathcal{L}$ jest ciągiem elementów wymienionych zbiorów. Definicji zbioru formuł - $\text{For}(\mathcal{L})$ - nie będziemy tu przytaczać; ponieważ jednak odróżniamy zmienne wolne od związanych podamy jedynie fragment tej definicji, który pozwala budować formuły przy użyciu kwantyfikatora:

- 1) jeśli $\alpha \in \text{For}(\mathcal{L})$, $x \in \text{Var}$, $\xi \in \text{Vb}$ i ξ nie występuje w α , to $\forall \xi \alpha[\frac{x}{\xi}] \in \text{For}(\mathcal{L})$,

- 2) $\alpha[\frac{x}{\xi}]$ oznacza tu ciąg powstały w ciągu α przez zastąpienie elementu x elementem ξ ; podobnie można zdefiniować $\alpha[\frac{x}{t}]$ dla $t \in \mathcal{T}(\text{Var}, \tau)$.

W dalszej części pracy operując formułami postaci $\forall \xi \alpha[\frac{x}{\xi}]$ nie będziemy jawnie przytaczać warunku, że ξ nie występuje w formule α , łatwo bowiem można warunek ten dopisać wszędzie tam, gdzie jest on konieczny. Zamiast $\alpha\beta$ będziemy pisać $(\alpha \rightarrow \beta)$. Nawiasy, podobnie jak znak " $\rightarrow$ ", będą pełniły rolę pomocniczą.

I n t e r p r e t a c j a $\mathcal{I}$ języka $\mathcal{L}$ zadana jest przez:

- 1) algebrę częściową $\mathcal{M}$ typu τ ,
- 2) waluację $s_{\mathcal{I}} \in \text{Val}_{\mathcal{M}}(\text{Var})$ (położmy $V_{\mathcal{I}} = \text{Var} \setminus \mathcal{M}(s_{\mathcal{I}})$),
- 3) funkcję Ψ określoną na zbiorze R : jeśli $P \in R_n$, to $\Psi(P)$

jest relacją (pełną) n -argumentową w zbiorze $\mathcal{T}_{\mathcal{M}}(V_{\mathcal{J}})$. Zamiast $\Psi(P)$ będziemy pisać $P^{\mathcal{J}}$; $P^{\mathcal{J}} = (P_+^{\mathcal{J}}, P_-^{\mathcal{J}})$.

Dla interpretacji $\mathcal{J}$ i dowolnej formuły $\alpha \in \text{For}(\mathcal{L})$ określimy $\alpha^{\mathcal{J}} \in \{0, 1\}$. Wcześniej jednak podamy definicję $\alpha^{\mathcal{J}, s}$ dla dowolnej waluacji $s \in \text{Val}_{\mathcal{M}}(\text{Var})$ takiej, że $\text{Var} \setminus \mathcal{M}(s) \subseteq V_{\mathcal{J}}$.

- 1) $(P t_0 \dots t_{n-1})^{\mathcal{J}, s} = 1$ wtw $(t_0^{\langle \mathcal{M}, s \rangle}, \dots, t_{n-1}^{\langle \mathcal{M}, s \rangle}) \in P_+^{\mathcal{J}}$.
- 2) $0^{\mathcal{J}, s} = 0$.
- 3) $(\alpha + \beta)^{\mathcal{J}, s} = 1$ wtw $\alpha^{\mathcal{J}, s} = 0$ lub $\beta^{\mathcal{J}, s} = 1$.
- 4) $(\forall x \alpha[\frac{x}{a}])^{\mathcal{J}, s} = 1$ wtw dla każdego $a \in M$; $\alpha^{\mathcal{J}, s[\frac{x}{a}]} = 1$.

Położmy $\alpha^{\mathcal{J}} = \alpha^{\mathcal{J}, s_{\mathcal{J}}}$; jeśli $\alpha^{\mathcal{J}} = 1$, to formułę α nazywamy spełnioną w interpretacji $\mathcal{J}$.

Przyjmijmy następującą definicję semantycznych relacji konsekwencji. Niech $\mathbb{K}$ będzie klasą interpretacji. Klasa $\mathbb{K}$ wyznacza następującą relację konsekwencji w języku $\mathcal{L}$:

dla $A \subseteq \text{For}(\mathcal{L})$, $\alpha \in \text{For}(\mathcal{L})$; $A \stackrel{\mathbb{K}}{=} \alpha$ wtw dla każdej interpretacji $\mathcal{J} \in \mathbb{K}$ jeśli $\{\beta^{\mathcal{J}} : \beta \in A\} \subseteq \{1\}$, to $\alpha^{\mathcal{J}} = 1$.

Interesować nas będą następujące cztery semantyczne relacje konsekwencji:

1) konsekwencja $=$ wyznaczona przez klasę wszystkich interpretacji (klasa $\mathcal{J}\text{nt}$),

2) konsekwencja $\overset{\times}{=}$ wyznaczona przez klasę wszystkich interpretacji $\mathcal{J} = (\mathcal{M}, s_{\mathcal{J}}, \Psi)$ takich, że $\mathcal{M}$ jest algebrą (klasa $\mathcal{J}\text{nt}^*$),

3) konsekwencja $\overset{\circ}{=}$ wyznaczona przez klasę wszystkich interpretacji $\mathcal{J} = (\mathcal{M}, s_{\mathcal{J}}, \Psi)$ takich, że $s_{\mathcal{J}}$ jest waluacją pełną, tzn. $\mathcal{M}(s_{\mathcal{J}}) = \text{Var}$ (klasa $\mathcal{J}\text{nt}^{\circ}$),

4) konsekwencja $\overset{\oplus}{=}$ wyznaczona przez klasę wszystkich interpretacji $\mathcal{J} = (\mathcal{M}, s_{\mathcal{J}}, \Psi)$ takich, że $\mathcal{M}$ jest algebrą, natomiast $s_{\mathcal{J}}$ jest waluacją pełną (klasa $\mathcal{J}\text{nt}^{\oplus}$).

Zauważmy, że warunek żądający, żeby $s_{\mathcal{J}}$ było waluacją pełną wyklucza interpretacje o dziedzinie pustej. Konsekwencje $=$ i $\overset{\times}{=}$ są inkluzywne; jeśli $\mathbb{F}_0 \neq \emptyset$, to warunek, żeby $\mathcal{M}$ było algebrą również wyklucza dziedzinę pustą, jest to jednak wykluczenie naturalne. Konsekwencja $\overset{\oplus}{=}$ jest standardową relacją konsekwencji.

Następujące warunki wyznaczają w języku syntaktyczną relację konsekwencji $\vdash$:

- A1) jeśli α jest formułą tautologiczną, to $\vdash \alpha$,

$$A2) \vdash \forall x (\alpha \rightarrow \beta) \left[\frac{x}{x} \right] \rightarrow (\forall x \alpha \left[\frac{x}{x} \right] \rightarrow \forall x \beta \left[\frac{x}{x} \right]),$$

$$A3) \vdash \alpha \rightarrow \forall x \alpha \left[\frac{x}{x} \right], \text{ o ile } x \notin \text{Var}(\alpha) \text{ (wówczas } \alpha \left[\frac{x}{x} \right] = \alpha),$$

$$A4) \vdash \forall \eta (\forall x \alpha \left[\frac{x}{x} \right]) \left[\frac{Y}{\eta} \right] \rightarrow \forall x (\forall \eta \alpha \left[\frac{Y}{\eta} \right]) \left[\frac{x}{x} \right],$$

$$A5) \vdash \forall \eta (\forall x \alpha \left[\frac{x}{x} \right] \rightarrow \alpha \left[\frac{Y}{Y} \right]) \left[\frac{Y}{\eta} \right],$$

$$A6) \text{ jeśli } \vdash \alpha, \text{ to } \vdash \forall x \alpha \left[\frac{x}{x} \right],$$

$$A7) \{ \alpha \rightarrow \beta, \alpha \} \vdash \beta.$$

Status reguły A6 jest inny niż reguły A7. Podczas gdy A7 jest regułą konsekwencji $\vdash$, reguła A6 nie wyprowadza jedynie poza zbiór tez (można zresztą ograniczyć stosowanie A6 tylko do aksjomatów). Niech teraz konsekwencje $\vdash^x$, $\vdash^0$ i $\vdash^*$ będą wyznaczone warunkami A1-A4, A6, A7 i odpowiednio warunkiem $A5^x$, $A5^0$ i $A5^*$:

$$A5)^x \vdash^x \forall \eta_{n-1} (\dots (\forall \eta_0 (\forall x \alpha \left[\frac{x}{x} \right] \rightarrow \alpha \left[\frac{Y_0}{\eta_0} \right]) \dots) \left[\frac{Y_{n-1}}{\eta_{n-1}} \right], \text{ dla } t = t(Y_0, \dots, Y_{n-1}),$$

$$A5)^0 \vdash^0 \forall x \alpha \left[\frac{x}{x} \right] \rightarrow \alpha \left[\frac{Y}{Y} \right],$$

$$A5)^* \vdash^* \forall x \alpha \left[\frac{x}{x} \right] \rightarrow \alpha \left[\frac{Y}{Y} \right].$$

Zauważmy, że w obecności warunków $A5^0$ lub $A5^*$ można pominąć A4.

Zdefiniowane aksjomatycznie relacje konsekwencji odpowiadają tym, które wcześniej określiliśmy semantycznie tzn.: $\vdash = \vdash$, $\vdash^x = \vdash^x$, $\vdash^0 = \vdash^0$, $\vdash^* = \vdash^*$.

Dowód wymienionych powyżej twierdzeń o pełności zostanie przedstawiony w następnej pracy poświęconej prezentowanej tu teorii. Będą tam również omówione pewne interesujące kwestie związane z dołączeniem do rozważanego tu języka, symboli wyróżnionych predykatów: istnienia i identyczności. Istnieje kilka możliwości semantycznego i aksjomatycznego scharakteryzowania tych predykatów; dlatego też, być może, właściwiej byłoby mówić o kilku rodzajach identyczności bądź istnienia.

Przystąpmy teraz do określenia struktury. **S t r u k t u r a** typu (τ, ρ) (τ jest typem funkcyjnym, ρ - typem relacyjnym) nazywać będziemy parę $(\mathcal{M}, \theta)$ taką, że $\mathcal{M}$ jest algebrą częściową typu τ , natomiast θ jest funkcją określoną na zbiorze $\mathbb{R}$: jeśli $P \in \mathbb{R}_n$, to $\theta(P)$ jest relacją częściową n -argumentową w zbiorze $\mathcal{T}_{\mathcal{M}}$.

Jeśli $\delta = (\mathcal{M}, \theta)$ jest strukturą i $P \in R$ to zamiast $\theta(P)$ będziemy pisać P^δ ; $P^\delta = (P_+^\delta, P_-^\delta)$.

Interpretację $\mathcal{I} = (\mathcal{M}, s_{\mathcal{I}}, \psi)$ języka $\mathcal{L}$ nazywać będziemy uzupełnieniem struktury $\delta = (\mathcal{M}, \theta)$, jeśli dla każdego $P \in R$: $P_+^\delta \subseteq P_+^{\mathcal{I}}$ i $P_-^\delta \subseteq P_-^{\mathcal{I}}$. Zbiór wszystkich interpretacji będących uzupełnieniem struktury oznaczamy $\mathcal{I}nt(\delta)$. Położmy $\mathcal{I}nt^0(\delta) = \mathcal{I}nt(\delta) \cap \mathcal{I}nt^0$.

Zbiór formuł prawdziwych w strukturze δ zdefiniować można opierając się na pojęciu spełniania w interpretacjach ze zbioru $\mathcal{I}nt(\delta)$ lub też ze zbioru $\mathcal{I}nt^0(\delta)$. Rozpatrywać można pewne wyróżnione klasy struktur, m. in.: 1) klasę struktur $\delta = (\mathcal{M}, \theta)$ takich, że $\mathcal{M}$ jest algebrą (wówczas $\mathcal{T}_{\mathcal{M}} = M$), 2) klasy struktur $\delta = (\mathcal{M}, \theta)$ takich, że dla każdego $P \in R$: a) P^δ jest relacją pełną na zbiorze $\mathcal{T}_{\mathcal{M}}$ (tzn. $P_+^\delta \cup P_-^\delta = \mathcal{T}_{\mathcal{M}}^n$) b) $M^n \subseteq P_+^\delta \cup P_-^\delta$, c) $P_+^\delta \cup P_-^\delta = M^n$. Zauważmy jeszcze, że prawdziwość formuły w strukturze δ definiować możemy również na podstawie silnej trójwartościowej logiki Kleene'go i jest to dość naturalna logika dla rozważanych przez nas struktur. Wspomniane powyżej wielorakie możliwości określenia prawdziwości w strukturach i możliwość ograniczenia się do pewnej klasy struktur powodują, że występujące w uniwersalnie wolnej logice Meyera i Lamberta różne pojęcia interpretacji i prawdziwości dają się zrekonstruować w przedstawionej tutaj teorii.

Warto być może zauważyć, że rozważanie nazw nieoznaczających w języku pierwszego rzędu przynieść może określone korzyści. Zdania o przedmiotach nieistniejących, w szczególności zdania stwierdzające istnienie bądź nieistnienie pewnych przedmiotów mogą być wówczas formułowane w języku przedmiotowym. Jako przykład weźmy język teorii mnogości z symbolem $\in$ i potraktujmy każde wyrażenie postaci $\{\xi : \alpha[\frac{X}{\xi}]\}$ jako nazwę. Przyjmijmy, że aksjomatem teorii jest każda generalizacja formuł postaci $\forall \eta((\eta \in \{\xi : \alpha[\frac{X}{\xi}]\}) \leftrightarrow \alpha[\frac{X}{\eta}])$. Położmy $a_R = \{\xi : \xi \notin \xi\}$ (zbiór Russella). Otrzymamy wówczas $\forall \eta(\eta \in a_R \leftrightarrow \eta \notin \eta)$. Wykorzystajmy teraz następujący schemat aksjomatów: $\forall \eta \beta[\frac{Y}{\eta}] \rightarrow (E!(m) \rightarrow \beta[\frac{Y}{m}])$, gdzie m jest dowolną nazwą rozważanego języka; położmy $\beta = (y \in a_R \leftrightarrow y \notin y)$, natomiast $m = a_R$. Dostaniemy stąd: $\neg E!(a_R)$. Powyższy przykład pokazuje, że w omawianej teorii kwantyfikacji, inaczej niż w teorii standardowej, nie musimy się obawiać, że wprowadzanie do języka pewnych wyrażeń nazwowych może spowodować sprzeczność teorii; wprowadzenie do języka nazwy m jest niezależne od przyjęcia założenia:

$E!(m)$. W naszkicowanej powyżej "teorii mnogości", można zdefiniować np. $\emptyset = \{x: 0\}$ nie otrzymamy stąd jednak tezy: $E!(\emptyset)$ istnienie zbioru pustego musimy zagwarantować oddzielnym aksjomatem.

Uniwersytet Łódzki
Katedra Logiki i Metodologii Nauk

Andrzej Nowicki

ON THE INCLUSIVE THEORY OF QUANTIFICATION

The first part of this paper presents the foundations and some properties of the inclusive theory of quantification and analyses the relationship between inclusive and standard quantification theory. Some arguments concerning application of the inclusive theory instead of the standard theory have been discussed too. The second part of the article presents some versions of the inclusive theory of quantification formulated by: Mostowski (1951), Hailperin (1953), Quine (1954), Meyer and Lambert (1968).

In the third part, a new conception of the inclusive theory of predicates has been presented. This conception is based, analogically to Meyer and Lambert theory, on the idea of undesignating names the first-order language. Four consequence relations have been semantically defined, and the way, how the foundations concerning realizing in the domain of individual names condition the form of the quantification theory has been presented. The discussed consequence relations have been axiomatically characterised. The only difference between the axiomatics is the way of dealing with the formulas of the following type: $\forall x\alpha \rightarrow \alpha[\frac{x}{t}]$, where t is the name (term) of the first-order language.