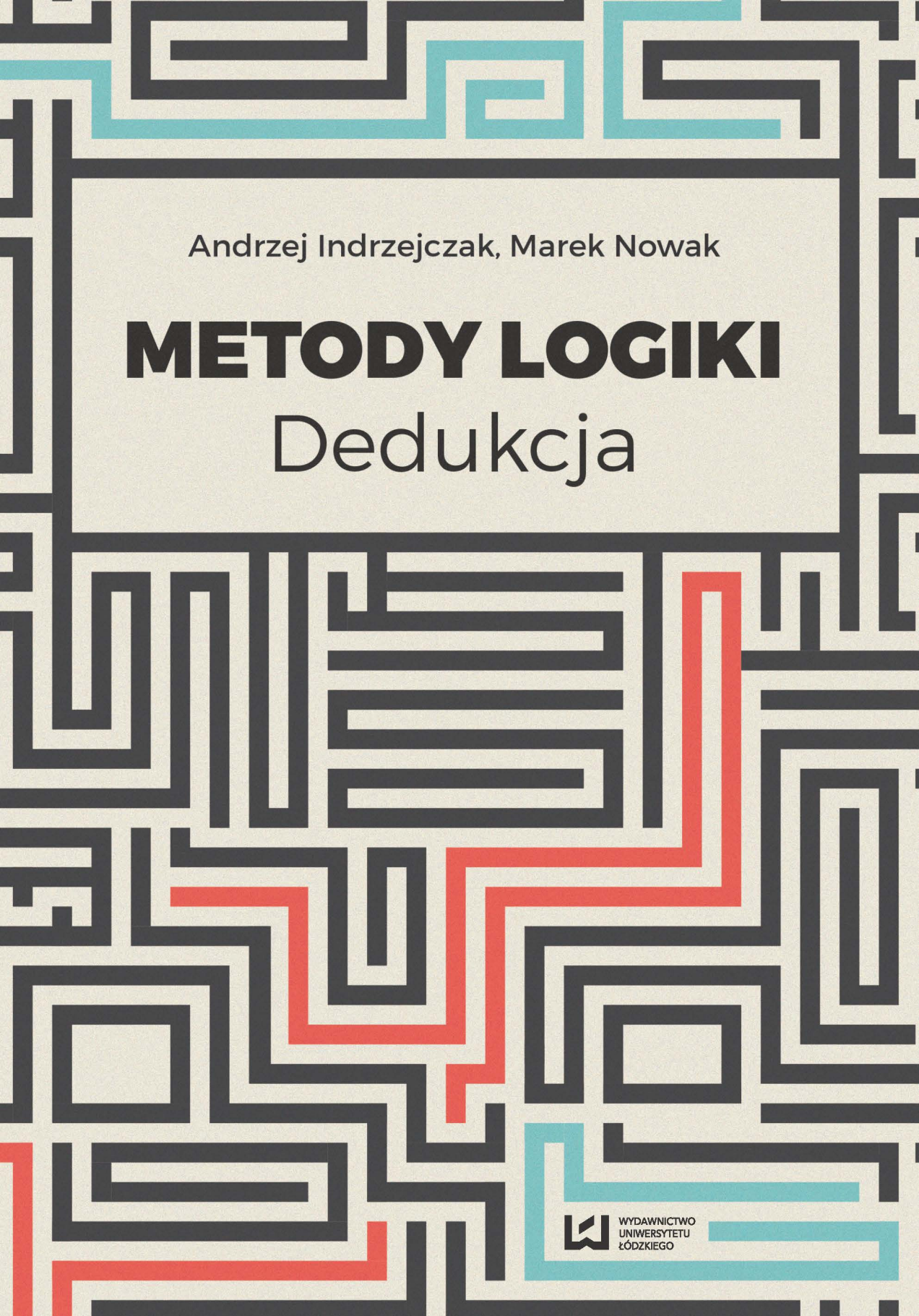




Andrzej Indrzejczak, Marek Nowak

METODY LOGIKI

Dedukcja



WYDAWNICTWO
UNIWERSYTETU
ŁÓDZKIEGO

METODY LOGIKI

Dedukcja



WYDAWNICTWO
UNIWERSYTETU
ŁÓDZKIEGO

Andrzej Indrzejczak, Marek Nowak

METODY LOGIKI

Dedukcja

Andrzej Indrzejczak, Marek Nowak – Uniwersytet Łódzki
Wydział Filozoficzno-Historyczny, Katedra Logiki i Metodologii Nauk
90-131 Łódź, ul. Lindleya 3/5

RECENZENT
Dariusz Surowik

REDAKTOR INICJUJĄCY
Damian Rusek

SKŁAD I ŁAMANIE
Andrzej Indrzejczak, Marek Nowak

PROJEKT OKŁADKI
Katarzyna Turkowska

Zdjęcie wykorzystane na okładce: © Depositphotos.com/silvercircle

Wydrukowano z gotowych materiałów dostarczonych do Wydawnictwa UŁ

© Copyright by Authors, Łódź 2016
© Copyright for this edition by Uniwersytet Łódzki, Łódź 2016
Publikacja jest udostępniona na licencji Creative Commons Uznanie autorstwa-Użycie
niekomercyjne-Bez utworów zależnych 4.0 (CC BY-NC-ND)

Wydane przez Wydawnictwo Uniwersytetu Łódzkiego
Wydanie I. W.07671.16.0.K

Ark. druk. 9,0

ISBN 978-83-8088-359-8
e-ISBN 978-83-8088-360-4
<https://doi.org/10.18778/8088-360-4>

Wydawnictwo Uniwersytetu Łódzkiego
90-131 Łódź, ul. Lindleya 8
www.wydawnictwo.uni.lodz.pl
e-mail: ksiegarnia@uni.lodz.pl
tel. (42) 665 58 63

Spis treści

Wstęp	9
1 Dowodzenie w logice klasycznej	13
1.1 Klasyczny rachunek zdań	13
1.1.1 Język KRZ	13
1.1.2 Aksjomatyzacja KRZ	16
1.1.3 Dowód	17
1.2 Dedukcja naturalna	20
1.2.1 Pierwotne reguły inferencji	20
1.2.2 Proste dedukcje	21
1.2.3 Dowody założeniowe wprost	23
1.2.4 Dowodzenie nie wprost	24
1.2.5 Dowody a dedukcje	26
1.2.6 Równoważności	28
1.3 Zaawansowana dedukcja	29
1.3.1 Stosowanie założeń dodatkowych	29
1.3.2 Poddowody warunkowe	30
1.3.3 Poddowody nie wprost	32
1.3.4 Poddowody wielokrotne i zagnieżdżone	33
1.4 Dodatkowe środki dowodowe	36
1.4.1 Reguły wtórne	36
1.4.2 Reguły obustronne	38
1.4.3 Dodatkowe reguły konstrukcji dowodu	42
1.4.4 Dodatkowe sposoby dowodzenia równoważności	45
1.5 Klasyczny rachunek kwantyfikatorów	47
1.5.1 Języki pierwszego rzędu	48
1.5.2 Zmienne wolne i związane	51
1.5.3 Podstawianie i zastępowanie	52

1.6	Dowodzenie w rachunku kwantyfikatorów	54
1.6.1	Reguły inferencji dla $\forall$ i $\exists$	54
1.6.2	Reguły konstrukcji dowodu dla kwantyfikatorów	58
1.6.3	Reguły wtórne	62
1.6.4	Reguły dla identyczności	64
1.7	Uwagi końcowe	68
1.7.1	Strategie dowodzenia	68
1.7.2	Dowody nieformalne	72
2	Dowodzenie w arytmetyce liczb naturalnych i teorii zbiorów	75
2.1	Arytmetyka elementarna	75
2.1.1	Aksjomaty	75
2.1.2	Dowody indukcyjne	76
2.2	Arytmetyka liczb naturalnych z dodawaniem	77
2.2.1	Aksjomaty i podstawowe własności dodawania	77
2.2.2	Relacja porządku	81
2.3	Arytmetyka z dodawaniem i mnożeniem	84
2.3.1	Aksjomaty i podstawowe własności mnożenia	84
2.4	Teoria mnogości	86
2.4.1	Naiwna teoria zbiorów	86
2.4.2	Paradoks Russella	88
2.5	Teoria zbiorów Zermelo-Fraenkla	89
2.5.1	Aksjomaty teorii mnogości ZF^- (bez aksjomatów ufundowania i wyboru)	89
2.5.2	Inkluzja zbiorów	93
2.5.3	Zbiór pusty	95
2.5.4	Zbiór potęgowy zbioru	97
2.5.5	Suma zbioru	98
2.5.6	Para zbiorów, zbiór jednoelementowy	99
2.5.7	Operacje boolowskie na zbiorach, zbiór n -elementowy	100
2.5.8	Przekrój zbioru niepustego	105
2.6	Algebra Boole'a zbiorów	107
2.6.1	Ciało zbiorów	107
2.6.2	Algebra Boole'a	110
2.7	Relacje i funkcje	112
2.7.1	Para uporządkowana. Produkt kartezjański dwóch zbiorów	112
2.7.2	Relacje binarne	115
2.7.3	Funkcje	119

2.8	Zbiory ufundowane	126
2.8.1	Teoria ZF^- z aksjomatem Ω	127
2.8.2	Aksjomat regularności (ufundowania)	136
2.9	Interpretacja arytmetyki elementarnej w teorii ZF	137
2.9.1	Operacja następnika	137
2.9.2	Indukcja	139
Bibliografia		143

Wstęp

Logika współczesna dostarcza formalnych narzędzi umożliwiających precyzyjną analizę poprawności rozumowań. Najważniejszym sposobem uzasadniania głoszonych twierdzeń jest bez wątpienia dowód. Na gruncie logiki wypracowano wiele typów systemów dowodzenia takich jak systemy aksjomatyczne, systemy rezolucji, systemy sekwentowe czy systemy dedukcji naturalnej. Te ostatnie zostały zaproponowane jako najbardziej zbliżone do faktycznej praktyki dowodzenia stosowanej przede wszystkim przez matematyków. W prezentowanym opracowaniu skupimy się na opisie struktury i zastosowań ostatniego z wymienionych sposobów kodyfikacji dowodzenia. W rozdziale pierwszym zaprezentujemy jedno z jego możliwych realizacji w odniesieniu do logiki klasycznej. Dowody i dedukcje będą tutaj prezentowane w czysto formalny sposób. Prezentację dedukcji naturalnej poprzedzimy krótkim omówieniem aksjomatycznej formalizacji klasycznego rachunku zdań. W drugim rozdziale zaprezentowane zostaną dwie teorie aksjomatyczne: arytmetyka liczb naturalnych w ujęciu Peano i obszerny fragment teorii mnogości w ujęciu Zermelo i Fraenkela. Dowody twierdzeń w tych teoriach zostaną przedstawione skrótowo, w sposób niesformalizowany ale na tyle dokładny, że można je „przepisać” jako w pełni formalne dowody w systemie dedukcji naturalnej z części pierwszej, wzbogaconym o aksjomaty rozważanych teorii. Takie rozwiązanie pokazuje nam w naturalny sposób związek dedukcji naturalnej z faktycznie prezentowanymi szkicami dowodów w literaturze matematycznej a ponadto jest ekonomiczne, gdyż formalne dowody ważnych twierdzeń są zwykle bardzo długie.

Zanim przejdziemy do charakterystyki dedukcji naturalnej wypada poświęcić nieco uwagi historii zagadnienia. Dowody, często całkiem zadowalające z punktu widzenia stosowanych obecnie kryteriów, konstruowali już matematycy greccy z kręgu szkoły pitagorejskiej działający w VI–V w. p.n.e. Wypracowali oni wiele reguł i technik, które należą do wypróbowanego repertuaru środków dowodzenia stosowanego współcześnie. Należą do nich np.

takie środki dowodzenia jak dowód warunkowy, dowód nie wprost, dowodzenie z alternatywy, czy reguły wnioskowania takie jak modus ponendo ponens, modus tollendo tollens, modus ponendo tollens¹. W rozdziale 1 objaśnimy działanie wszystkich wymienionych wyżej technik dowodzenia.

Intuicyjne stosowanie środków dowodzenia należy jednak odróżnić od świadomej refleksji nad ich prawomocnością. Ta pojawiła się nieco później w pismach logicznych Arystotelesa (384–322 p.n.e.), który wprowadził pojęcie dedukcji jako niezawodnego sposobu rozumowania gwarantującego, że ze zdań prawdziwych, przyjętych jako przesłanki rozumowania, wywnioskujemy jedynie prawdziwe wnioski. Dokonał on również pierwszej próby kodyfikacji logiki nazw, która systematyzowała pewną klasę sposobów rozumowania. Sylogistyka, czyli rachunek nazw Arystotelesa, stanowiła wprawdzie podstawę nauczania logiki aż do wieku XX, jednak jej zakres zastosowań był zbyt wąski, zwłaszcza z punktu widzenia potrzeb matematyki. Alternatywne próby kodyfikacji zasad rozumowania, prowadzące do konstrukcji logiki zdań, podejmowali filozofowie stoicy, jednak wiedza, którą dysponujemy na temat ich dokonań jest bardzo fragmentaryczna.

U Arystotelesa pojawia się też refleksja nad budową systemów aksjomatycznych jako sposobów systematyzowania wiedzy z danej dziedziny. W systemie aksjomatycznym ustala się pewną, zazwyczaj niewielką, ilość pojęć pierwotnych dla danej teorii, oraz zestaw aksjomatów, które są twierdzeniami przyjętymi bez dowodu, zazwyczaj uznawanymi za intuicyjnie oczywiste. System rozbudowuje się w sensie pojęciowym przez dołączanie definicji nowych pojęć oraz w sensie merytorycznym przez dołączanie nowych twierdzeń udowodnionych na podstawie aksjomatów i wcześniej udowodnionych twierdzeń. Tym samym systemy aksjomatyczne można uznać za najstarszą formę systemów dowodzenia, chociaż niektórzy badacze (np. Corcoran [6]) twierdzą, że stoicy prezentowali swoją logikę zdań w postaci systemu dedukcji naturalnej. Nie ulega wątpliwości, że pierwszy dojrzały system aksjomatyczny obejmujący całość ówczesnej matematyki został zaprezentowany w „Elementach” Euklidesa. Praca ta przez ponad dwa tysiąclecia stanowiła wzór stosowania systemów aksjomatycznych, inspirując wielu myślicieli (np. Kartezjusza (1596–1647)) do przenoszenia tych zasad poza matematykę.

Jednak starożytne rozumienie systemu aksjomatycznego, z punktu widzenia dzisiejszych wymagań, pozostawiało sporo do życzenia. Jeżeli chodzi o prezentację podstawy systemu, to była ona zazwyczaj podawana w sposób dość intuicyjny. Jeżeli chodzi o sposób rozbudowy systemu, to nie określono

¹Nazwy są późniejsze, łacińskie.

precyzyjnie ani zasad definiowania ani dostępnego zakresu sposobów dowodzenia i samej konstrukcji dowodu. Problemy te podjęte zostały dopiero pod koniec XIX wieku przez niemieckiego matematyka Gottloba Fregego, „ojca” współczesnej logiki matematycznej, a doprecyzowane ostatecznie w tzw. szkole Dawida Hilberta. Współcześnie systemy dowodzenia, w szczególności aksjomatyczne, prezentuje się w postaci całkowicie sformalizowanej, tzn. jako konstrukcje podane w języku sztucznym i spełniające określone warunki poprawności. W konstrukcji systemu aksjomatycznego przestano przywiązywać wagę do takich tradycyjnie formułowanych postulatów jak „oczywistość” aksjomatów; kłopoty z paradoksami w podstawach matematyki, uzyskanymi na podstawie – jak się wydawało – zupełnie oczywistych założeń, uczyniły takie wymagania bezużytecznymi. Bardziej podstawowym wymogiem stała się niesprzeczność zbioru aksjomatów, czyli niemożność wydedukowania z nich jakiegoś zdania i jego zaprzeczenia. Ryzyko takie wiąże się jednak nie tylko z kwestią doboru aksjomatów; do sprzeczności może też doprowadzić niefrasobliwy sposób definiowania lub stosowanie niepoprawnych reguł wnioskowania. Stąd we współczesnej teorii systemów formalnych do tych kwestii przywiązuje się bardzo dużą wagę.

Zasady budowy i stosowania sformalizowanego systemu aksjomatycznego zilustrujemy na przykładzie konstrukcji (jednego z wielu możliwych) systemów aksjomatycznych Klasycznego Rachunku Zdań (KRZ). Ukazuje on na bardzo prostym przykładzie ogólne własności charakterystyczne systemów aksjomatycznych – „bogatych” w aksjomaty a „ubogich” w reguły. Dwa dalsze przykłady teorii aksjomatycznych zostaną zaprezentowane w rozdziale drugim. Rozdział 1 został przygotowany przez Andrzeja Indrzejczaka a rozdział 2 przez Marka Nowaka.

Rozdział 1

Dowodzenie w logice klasycznej

1.1 Klasyczny rachunek zdań

Logika klasyczna często jest dzielona na dwie składowe: klasyczny rachunek zdań (KRZ) i klasyczny rachunek kwantyfikatorów (KRK). Prezentacja taka ma swe uzasadnienie zarówno dydaktyczne jak i merytoryczne. KRZ jest łatwiejszy do omówienia a ponadto posiada pewne pożyteczne własności (np. rozstrzygalność), których KRK nie posiada. KRZ jest systemem prostszym, gdyż ogranicza się do analizy takich związków logicznych, które zależą tylko od występowania pewnego typu spójników. W związku z tym zarówno język KRZ jak i jego semantyka (której tutaj nie omawiamy) są bardzo proste. Niezależnie od swojej prostoty KRZ okazuje się narzędziem wystarczającym nawet do realizacji bardzo wymagających zadań, a jego struktura jest wystarczająco bogata aby umożliwić prezentację rozmaitych technik dedukcji.

1.1.1 Język KRZ

Formalny język dowolnej logiki lub pozalogicznej teorii ustalamy przez scharakteryzowanie jego słownika i reguł składniowych. Jest to czysto syntaktyczne ujęcie, w którym nie odwołujemy się do żadnego pojęcia semantycznej interpretacji. W obrębie słownika ustala się jakie występują w nim kategorie wyrażen zmiennych i stałych. Zmienne nie mają ustalonego znaczenia a jedynie ustalony zakres możliwych podstawień, np. zmienne nazwowe to symbole, które mogą reprezentować dowolne wyrażenia nazwowe (można za nie podstawiać dowolne nazwy). Wyrażenia stałe to takie, których znaczenie jest precyzowane na gruncie danego systemu, np. w aksjomatach lub regułach dowodzenia. Stałe dzielimy na logiczne i pozalogiczne, te ostatnie występują

tylko w językach formalnych teorii (aksjomatycznych). Opisu języka dokonujemy w metajęzyku, który jest zazwyczaj mieszanką (pewnej części) języka naturalnego i dodatkowych symboli tzw. metazmiennych. W dalszym ciągu jako symboli metazmiennych będziemy używać liter greckich, a jako symboli zmiennych liter łacińskich. Odróżnienie (sztucznego) języka danej logiki czy teorii, zwanego też językiem przedmiotowym, od metajęzyka używanego do jego opisu jest gwarantem uniknięcia problemów takich jak np. paradoks kłamcy, wynikających z samozwrotności języków naturalnych. Użycie metajęzyka, a konkretnie metazmiennych, pozwala też uzyskać niezbędny poziom ogólności w omawianiu reguł i systemów dowodzenia, np. stwierdzać coś o dowolnym zdaniu, bez względu na jego formę.

Język KRZ jest bardzo prosty. Przyjmujemy w nim tylko jeden nieskończony, przeliczalny zbiór zmiennych zdaniowych ZZ. Będą one reprezentowane przez litery $p, q, r, s, \dots, p_1, q_1, \dots, p_2, \dots$. Zmienne tego typu mają za zadanie reprezentować dowolne zdania oznajmujące. Jedyne (pierwotne) stałe tego języka to spójniki oznaczane symbolami: $\neg$ (negacja), $\wedge$ (koniunkcja), $\vee$ (alternatywa), $\rightarrow$ (implikacja). Pierwszy z nich jest jednoargumentowy, tzn. łączy się z jednym zdaniem jako swym argumentem, powstałe w ten sposób zdanie nazywamy również negacją. Pozostałe są dwuargumentowe, tzn. łączą ze sobą dwa zdania jako swoje argumenty. Intuicyjnie, negacja odpowiada wyrażeniu „nieprawda, że”, koniunkcja – „i”, alternatywa – „lub”, a implikacja – „jeżeli ... to”. Należy jednak pamiętać, że chodzi tu jedynie o przybliżoną charakterystykę; podane zwroty w języku polskim są bowiem wieloznaczne¹, natomiast odpowiednie spójniki języka KRZ mają w nim precyzyjnie określone znaczenie. W metajęzyku będziemy używać małych liter greckich $\varphi, \psi, \chi, \dots$ na oznaczenie dowolnych zdań (formuł) języka KRZ; a ich zbiory oznaczymy za pomocą wielkich liter $\Gamma, \Delta, \Pi, \dots$. Możemy teraz podać kluczową definicję zdania albo formuły² w tym języku.

Definicja 1.1 (Zbiór formuł KRZ) *FOR jest najmniejszym zbiorem spełniającym warunki:*

1. $ZZ \subset FOR$
2. jeżeli $\varphi \in FOR$, to $\neg\varphi \in FOR$
3. jeżeli $\varphi, \psi \in FOR$, to $(\varphi \wedge \psi), (\varphi \vee \psi), (\varphi \rightarrow \psi) \in FOR$

¹Obszerną dyskusję na ten temat znaleźć można w Indrzejczak [12].

²W przypadku KRZ te dwa określenia mają takie samo znaczenie; rozróżnimy je w przypadku języka rachunku kwantyfikatorów.

W przypadku implikacji pierwszy (lewy) argument nazywamy poprzednikiem, a drugi następnikiem implikacji.

Podana definicja jest standardową definicją rekurencyjną, co umożliwia stosowanie w odniesieniu do danego systemu tzw. dowodów indukcyjnych przez tzw. indukcję strukturalną lub przy wykorzystaniu miar takich, jak długość (ilość wszystkich symboli) lub złożoność (ilość stałych logicznych) formuły.

W definicji użyliśmy dodatkowo nawiasów jako środków interpunkcyjnych. Są one często wymieniane w opisie słownika jako dodatkowa kategoria tzw. symboli pomocniczych. Można ich uniknąć stosując tzw. notację polską wprowadzoną przez Jana Łukasiewicza. Aby uprościć zapis zastosujemy konwencję pomijania zbędnych nawiasów, w szczególności zewnętrznych. Aby jeszcze bardziej zredukować ich liczbę stosujemy konwencję odnośnie siły wiązania spójników, wg. kolejności $\neg, \wedge, \vee, \rightarrow$. Dodatkowo omijając będziemy wewnętrzne nawiasy w przypadku wielokrotnego powtórzenia operacji łącznych, tj. koniunkcji i alternatywy. Konwencje te pozwalają formułę:

$$((p \vee (\neg q \wedge r)) \rightarrow ((p \wedge \neg s) \vee (q \wedge (s \wedge t))))$$

zapisać następująco:

$$p \vee \neg q \wedge r \rightarrow p \wedge \neg s \vee q \wedge s \wedge t$$

Celowo ograniczyliśmy zasób spójników będących pierwotnymi stałymi języka, pozwala to bowiem zredukować ilość aksjomatów. Można jednak dołączyć już teraz dodatkowe stałe za pomocą definicji. Będziemy używać symboli $\leftrightarrow$, $\perp$ i $\top$ (dwuargumentowy spójnik równoważności, stałe zdaniowe falsum i verum) jako definicyjnych skrótów:

- $(\varphi \leftrightarrow \psi) := ((\varphi \rightarrow \psi) \wedge (\psi \rightarrow \varphi))$
- $\perp := (p \wedge \neg p)$; $\top := \neg \perp$

W przypadkach użycia $\leftrightarrow$ przyjmujemy, że siła wiązania tego spójnika jest jeszcze słabsza niż pozostałych spójników dwuargumentowych, czyli tam gdzie nawiasy nie wskazują inaczej traktujemy $\leftrightarrow$ jako główny spójnik wyrażenia.

1.1.2 Aksjomatyzacja KRZ

Kolejnym krokiem jest podanie zbioru aksjomatów, czyli formuł (lub ich schematów) przyjętych bez dowodu, oraz zbioru reguł dowodzenia, które pozwalają z uznanych wcześniej formuł dedukować kolejne. Reguły takie będziemy przedstawiać za pomocą schematów $\varphi_1, \dots, \varphi_n / \psi$, gdzie każde φ_i , $i \leq n$, to przesłanka, a ψ to wydedukowany wniosek.

Istnieje wiele zestawów aksjomatów dla KRZ, poniżej podajemy wersję KRZ z regułą MP jako jedyną regułą inferencji oraz z następującą listą schematów aksjomatów podaną przez Hilberta.

1. $\varphi \rightarrow (\psi \rightarrow \varphi)$
2. $(\varphi \rightarrow (\psi \rightarrow \chi)) \rightarrow ((\varphi \rightarrow \psi) \rightarrow (\varphi \rightarrow \chi))$
3. $\varphi \wedge \psi \rightarrow \varphi$
4. $\varphi \wedge \psi \rightarrow \psi$
5. $\varphi \rightarrow (\psi \rightarrow \varphi \wedge \psi)$
6. $\varphi \rightarrow \varphi \vee \psi$
7. $\psi \rightarrow \varphi \vee \psi$
8. $(\varphi \rightarrow \chi) \rightarrow ((\psi \rightarrow \chi) \rightarrow (\varphi \vee \psi \rightarrow \chi))$
9. $(\neg\varphi \rightarrow \neg\psi) \rightarrow (\psi \rightarrow \varphi)$

Jedyną regułą jest reguła MP (modus ponens, reguła odrywania) postaci:
 $\varphi \rightarrow \psi, \varphi / \psi$.

Podany wyżej system aksjomatyczny jest przykładem formalizacji inwariantnej, tj. takiej, w której zamiast skończonego zbioru aksjomatów wyrażonych w języku przedmiotowym danej logiki używamy zbioru schematów wyrażonych w metajęzyku. Np. konkretnymi aksjomatami tego systemu są formuły: $p \rightarrow (q \rightarrow p)$, $p \wedge q \rightarrow (r \rightarrow p \wedge q)$, $r \rightarrow (r \rightarrow r)$ i nieskończenie wiele innych podstawień schematu 1. Należy podkreślić, że za metazmienne w schematach można podstawiać zarówno różne zmienne, jak i formuły złożone (np. $p \wedge q$ za φ w drugim przypadku). Zauważmy też, że za różne metazmienne można podstawiać tę samą formułę (r podstawione za obie metazmienne

w trzecim przypadku) ale oczywiście nie można podstawiać różnych formuł za tę samą metazmienną w jej różnych wystąpieniach w schemacie, np. $p \rightarrow (q \rightarrow r)$ nie jest podstawieniem aksjomatu 1.

W rezultacie zbiór aksjomatów jest nieskończony, ale nie musimy wprowadzać do zbioru reguł, jako pierwotnej, reguły podstawiania za zmienne zdaniowe. Dla krótkości w dalszym ciągu aksjomatami będziemy nazywać zarówno same schematy jak i ich konkretne instancje w języku przedmiotowym.

1.1.3 Dowód

Definicja dowodu w systemie aksjomatycznym jest bardzo prosta. Podamy ją tutaj w taki sposób, że równocześnie zdefiniujemy pojęcie relacji dowiedliwości lub dedukowalności (relację tę oznaczmy symbolem $\vdash$) zachodzącej między zbiorami formuł a pojedynczymi formułami, pojęcie tezy KRZ, oraz pojęcie (nie)sprzecznego zbioru formuł.

Definicja 1.2 • $\Gamma \vdash \varphi$ *wtw istnieje skończony ciąg formuł (dowód), w którym ostatnim elementem jest φ i gdzie każdy element jest:*

1. *podstawieniem aksjomatu lub*
2. *elementem zbioru Γ (przesłanką, założeniem) lub*
3. *wynikiem zastosowania MP do poprzedzających wyrazów ciągu*

- φ *jest tezą wtw $\vdash \varphi$ ($\Gamma = \emptyset$).*
- Γ *jest spreczny wtw $\Gamma \vdash \perp$, w przeciwnym wypadku jest niespreczny.*

W definicji tej (jak i w dalszej części pracy) użyliśmy skrótu „wtw” dla zwrotu „wtedy i tylko wtedy gdy”. Jak widać każdy dowód jest skończonym ciągiem formuł i jest albo dowodem tezy ($\vdash \varphi$) – gdy zbiór przesłanek jest pusty, albo dowodem uzasadniającym poprawność pewnego rozumowania ($\Gamma \vdash \varphi$) – gdy zbiór przesłanek jest niepusty. W drugim przypadku będziemy mówić o dedukcji (derywacji) φ z (przesłanek) Γ , zachowując zasadniczo określenie dowód tylko dla dowodów tez. Zauważmy też, że w dowodzie tezy w systemie aksjomatycznym, każdy wiersz poprzedzający konkluzję też jest tezą, chociaż zazwyczaj zainteresowani jesteśmy tylko formułą z wiersza ostatniego.

Definicja dowodu w systemie aksjomatycznym jest bardzo prosta, głównie dlatego, że ilość reguł jest ograniczona do minimum. Daje to liczne korzyści w badaniach metalogicznych, czyli teoretycznej refleksji nad własnościami takiego systemu. Jednak z tego samego powodu konstrukcja dowodu konkretnej tezy czy dedukcji wniosku z przesłanek w systemie aksjomatycznym nastrocza zazwyczaj wiele kłopotów i wymaga sporej wyobraźni, nawet gdy wynik wydaje się dość oczywisty. Dla większej przejrzystości dowody i dedukcje będziemy przedstawiać jako ciągi w trzech kolumnach. Kolumna pierwsza to numeracja kolejnych wierszy, kolumna druga to właściwy dowód (dedukcja) czyli ciąg formuł. Wreszcie ostatnia kolumna to ciąg uzasadnień poszczególnych wierszy; podajemy tam czy dana formuła jest przesłanką, założeniem, aksjomatem, czy formułą wydedukowaną – w tym przypadku, z jakich wierszy poprzedzających i za pomocą jakiej reguły. Dla ilustracji podajemy dowód tezy $p \rightarrow p$:

1	$p \rightarrow (p \rightarrow p)$	aksjomat 1
2	$p \rightarrow ((p \rightarrow p) \rightarrow p)$	aksjomat 1
3	$(p \rightarrow ((p \rightarrow p) \rightarrow p)) \rightarrow ((p \rightarrow (p \rightarrow p)) \rightarrow (p \rightarrow p))$	aksjomat 2
4	$(p \rightarrow (p \rightarrow p)) \rightarrow (p \rightarrow p)$	2, 3, <i>MP</i>
5	$p \rightarrow p$	1, 4, <i>MP</i>

Dowód nie jest wprawdzie zbyt długi ale wymaga pewnego wysiłku aby znaleźć odpowiednie podstawienia aksjomatów, które pozwolą na dedukcję poszukiwanego wyniku. Zazwyczaj prościej niż dowody tez można uzyskać dowody schematów rozumowań. Przykładowo, dowód rozumowania $p \vdash p$ zgodnie z definicją dowodu składa się tylko z jednego wiersza zawierającego formułę p jako zarazem założenie i konkluzję dowodu.

Krótko mówiąc systemy aksjomatyczne są wygodne dla uprawiania teorii ale niewygodne dla praktycznego dowodzenia twierdzeń. W argumentacji, uzasadnianiu, wnioskowaniu, stosujemy znacznie więcej reguł i technik wnioskowania niż reguła modus ponens. Nie wszystkie z nich są oczywiście poprawne z punktu widzenia logików, jednak wiele z nich, zwłaszcza te, które stosują w swej praktyce dowodowej matematycy, jest dedukcyjnych. Można oczywiście poprawić praktyczną wydolność systemu aksjomatycznego poprzez dołączanie do niego dodatkowych reguł (tzw. reguł wtórnych) W istocie, tak jak każdy wiersz dowodu dostarcza nam tezy, tak każda dedukcja $\Gamma \vdash \varphi$ ze skończonego zbioru formuł Γ , dostarcza uzasadnienia dla poprawności reguły o postaci Γ / φ . Reguły takie w dalszym ciągu nazywać

będziemy regułami inferencji. Przykładowo, poprawność reguły $\varphi, \psi / \varphi \wedge \psi$ możemy uzasadnić z pomocą następującej dedukcji:

1	p	założenie
2	q	założenie
3	$p \rightarrow (q \rightarrow p \wedge q)$	aksjomat 5
4	$q \rightarrow p \wedge q$	1, 3, MP
5	$p \wedge q$	2, 4, MP

W dedukcji tej użyliśmy wprowadzić konkretnych formuł języka KRZ, a nie zmiennych φ, ψ , łatwo jednak zauważyć, że bez względu na to jakie formuły znajdują się w powyższej dedukcji na miejscu zmiennych p, q , struktura dowodu nie ulegnie zmianie. Taki sposób uzasadniania wtórnych środków dowodzenia będziemy stosować również w dalszej części rozdziału, poświęconej dedukcji naturalnej.

Szczególną rolę w poszerzaniu asortymentu środków dowodowych systemu aksjomatycznego odgrywa tzw. twierdzenie o dedukcji (TD). Aczkolwiek odpowiada ono tradycyjnemu sposobowi dowodzenia warunkowego stosowanemu już w starożytności, to formalnie jego poprawność wykazał dopiero w latach 30-tych XX w. J. Herbrand. Nie jest to zwykła reguła wtórna postaci $\Gamma \vdash \varphi$, ale raczej formalny wyraz pewnej strategii dowodzenia, która pozwala uzyskać dowód tezy o postaci implikacji na podstawie dedukcji następnika tej implikacji z jej poprzednika jako założenia dowodu. O ekonomiczności takiego dowodzenia może świadczyć porównanie dowodu tezy $p \rightarrow p$ z dedukcją $p \vdash p$. Zastosowanie twierdzenia o dedukcji pozwala uzyskać $\vdash p \rightarrow p$ od razu na podstawie tej dedukcji. Ogólne sformułowanie twierdzenia o dedukcji (TD) wygląda następująco:

(TD) Jeżeli $\Gamma, \varphi \vdash \psi$, to $\Gamma \vdash \varphi \rightarrow \psi$.

To ogólniejsze sformułowanie dopuszcza, że w dedukcji ψ korzystamy też z założeń należących do Γ . Oczywiście w takim przypadku wydedukowana implikacja nie jest tezą. Zauważmy, że sukcesywne stosowanie TD do kolejnych założeń pozwala jednak zawsze otrzymać dowód tezy o postaci tzw. implikacji wstępującej (IW):

Jeżeli $\varphi_1, \dots, \varphi_n \vdash \psi$, to $\vdash \varphi_1 \rightarrow (\dots \rightarrow (\varphi_n \rightarrow \psi) \dots)$

Pokazuje to, że każda reguła inferencji może być przekształcona na tezę o odpowiednim kształcie.

1.2 Dedukcja naturalna

Zarysowane wyżej poszerzanie zbioru środków dowodzenia zamienia nam w istocie rzeczy wyjściowy system aksjomatyczny w inny system z innymi regułami i inną definicją dowodu, która powinna nowe reguły uwzględnić (por. np. formalizacje systemów aksjomatycznych w Porębska i Suchoń [20]). Można jednak do kwestii konstrukcji praktycznego systemu dowodzenia podejść w bardziej bezpośredni sposób niż przez stopniowe przerabianie systemu aksjomatycznego. Ze względu na niepraktyczność systemów aksjomatycznych logicy stosunkowo wcześniej zadbali o skonstruowanie wygodniejszych w praktyce systemów konstrukcji dowodu. Jednym z nich są systemy dedukcji naturalnej (DN), zaprojektowane niezależnie w latach 30-tych XX w. przez Gerharda Gentzena [9] i Stanisława Jaśkowskiego [14]. Należą one do najbardziej popularnych systemów dowodzenia ze względu na dużą łatwość w konstruowaniu dowodu oraz wykorzystywanie środków dowodzenia znanych od starożytności. W przeciwieństwie do systemów aksjomatycznych, systemy DN są „bogate” w reguły, a „ubogie” w aksjomaty, co nie tylko znacznie upraszcza praktykę budowania dowodu, ale stoi w zgodzie z przekonaniem, że logika ma być właśnie systemem dedukcji – aksjomaty powinny charakteryzować teorie pozalogiczne³. Taki sposób wykładu przyjmujemy w tej pracy; w rozdziale pierwszym szcharakteryzujemy logikę klasyczną w terminach dedukcji naturalnej, w rozdziale drugim zaprezentujemy dwie ważne teorie matematyczne w postaci aksjomatycznej, ale z użyciem wszystkich środków dowodowych dostarczonych na gruncie systemu DN.

Przedstawiony dalej system DN zasadniczo pochodzi od Słupeckiego i Borkowskiego [3] i oparty jest na oryginalnym systemie Jaśkowskiego. Żeby zwiększyć jeszcze bardziej jego walory dydaktyczne zmieniliśmy nieco jego konstrukcję (dobór reguł pierwotnych) i porządek prezentacji. Opis i zastosowania oryginalnego systemu Słupeckiego i Borkowskiego można znaleźć w wielu podręcznikach (np. [4, 5, 22]). Omówienie konkurencyjnych systemów DN zawiera Indrzejczak [13].

1.2.1 Pierwotne reguły inferencji

Zacznijmy od przedstawienia kilku podstawowych reguł inferencji, które – jak pamiętamy – pozwalają z pewnej ilości przesłanek wydedukować odpo-

³Wydaje się, że już Arystoteles stał na takim stanowisku, gdyż nie traktował logiki jako nauki ale „organon” czyli narzędzie uprawiania innych nauk.

wiedni wniosek. W systemie DN koniunkcja, alternatywa i implikacja scharakteryzowane są za pomocą następujących reguł:

$$\begin{array}{ll}
 (\wedge D) & \varphi, \psi / \varphi \wedge \psi \\
 (\wedge E) & \varphi \wedge \psi / \varphi \text{ oraz } \varphi \wedge \psi / \psi \\
 (\vee D) & \varphi / \varphi \vee \psi \text{ oraz } \psi / \varphi \vee \psi \\
 (\vee E) & \varphi \vee \psi, \neg \varphi / \psi \\
 (\rightarrow E) & \varphi \rightarrow \psi, \varphi / \psi
 \end{array}$$

Każdy spójnik (za wyjątkiem implikacji, do której wrócimy niżej) scharakteryzowany jest za pomocą pary reguł, z których jedno są regułami dołączania (formuły z takim spójnikiem do dowodu) a drugie regułami eliminacji. Stąd w nazwach litery *D* i *E*. Reguła eliminacji implikacji to po prostu MP (modus ponens). W przypadku reguł z dwiema (lub więcej) przesłankami kolejność ich występowania w dedukcji jest dowolna. W systemie dla KRZ nie występują żadne aksjomaty. Łatwo zauważyć, że obie reguły dla koniunkcji oddają dokładnie treść trzech aksjomatów charakteryzujących ten spójnik w systemie aksjomatycznym. Podobnie jest z regułą $(\vee D)$, która odpowiada aksjomatom 6 i 7; wszystkie na zasadzie odwrotności twierdzenia o dedukcji: jeżeli $\vdash \varphi \rightarrow \psi$, to $\varphi \vdash \psi$ i jeżeli $\varphi \rightarrow (\chi \rightarrow \psi)$, to $\varphi, \chi \vdash \psi$. Tylko reguła $(\vee E)$ nie odpowiada, w tym sensie, żadnemu aksjomatowi, jest ona natomiast tradycyjną regułą sformułowaną przez stoików i zwaną „psim sylogizmem”⁴ albo modus tollendo ponens (tryb stwierdzający przez zaprzeczenie).

1.2.2 Proste dedukcje

Zanim pokażemy, jak w systemie DN składającym się jedynie z reguł inferencji można budować dowody tez, skoncentrujemy się na konstrukcji dedukcji uzasadniających poprawność rozumowań. Poniższy przykład ilustruje zastosowanie podanych reguł w dedukcji uzasadniającej poprawność rozumowania: $p \rightarrow q \wedge r, r \rightarrow s \vee t, p \wedge \neg s \vdash t \wedge q$

⁴Podobno posłużył się nią w trakcie tropienia zająca pies Chryzypa, wybitnego stoickiego filozofa.

1	$p \rightarrow q \wedge r$	$p.$
2	$r \rightarrow s \vee t$	$p.$
3	$p \wedge \neg s$	$p.$
4	p	3, $\wedge E$
5	$q \wedge r$	1, 4, $\rightarrow E$
6	q	5, $\wedge E$
7	r	5, $\wedge E$
8	$s \vee t$	2, 7, $\rightarrow E$
9	$\neg s$	3, $\wedge E$
10	t	8, 9, $\vee E$
11	$t \wedge q$	6, 10, $\wedge D$

Podobnie jak w systemie aksjomatycznym w dedukcji umieszczamy jako osobne wiersze przesłanki (w kolumnie uzasadnień oznaczone literą $p.$), zazwyczaj na początku dedukcji, chociaż można je wprowadzać sukcesywnie później. Następnie stosując dostępne reguły staramy się rozbić złożone formuły za pomocą reguł eliminacji tak aby dostać te wyrażenia, z których za pomocą reguł dołączania wydedukujemy wniosek. Jest to bardzo ogólne zalecenie, które w praktyce może ulec wielu modyfikacjom omawianym dalej. W podanym przykładzie, aby wydedukować $t \wedge q$, musimy najpierw otrzymać obie zmienne; t z przesłanki 2 a q z 1. Jednak żadnej z tych przesłanek nie da się rozbić od razu w taki sposób aby otrzymać potrzebne nam zmienne; potrzebujemy p lub r aby móc do nich najpierw zastosować ($\rightarrow E$). p uzyskujemy z przesłanki 3 przez zastosowanie ($\wedge E$), co pozwala otrzymać następnie w analogiczny sposób q (jeden człon wniosku już mamy) i r . Ta ostatnia zmienna jest potrzebna do dedukcji następnika przesłanki 2, który zawiera t . Jednak tym razem jest to alternatywa i aby wyizolować t z pomocą ($\vee E$) potrzebujemy $\neg s$. Chwila refleksji pokazuje, że można tę formułę wydedukować z przesłanki 3.

Jak widać proces dedukcji nie jest automatyczny. Wymaga dokładnego sprawdzenia czego potrzebujemy (co jest celem dedukcji), co z czego możemy wydedukować i jakie reguły nam się do tego mogą przydać. Co więcej dedukcję dla tych samych rozumowań można zbudować na rozmaite sposoby, stosując różne reguły i w różnej kolejności. W miarę jak będziemy poszerzać asortyment środków dowodzenia wprowadzimy do tekstu szereg uwag o charakterze heurystycznym, dotyczących strategii poszukiwania rozwiązania. Na początek należy zapamiętać 3 ważne uwagi:

1. Kolejność przesłanek w dowodzie jest dowolna;
2. przesłanki i wniosek mogą być od siebie oddalone o dowolną ilość wierszy;
3. zarówno przesłanki jak i wniosek to formuły zajmujące całe wiersze dowodu a nie ich część; inaczej, reguły stosujemy do całych formuł a nie do ich podformuł.

Ostatnią uwagę zilustrujemy prostym przykładem: z $p \vee (q \rightarrow r)$ i q nie możemy wydedukować ani r ani $p \vee r$ za pomocą ($\rightarrow E$).

1.2.3 Dowody założeniowe wprost

Aby w systemie DN budować nie tylko dedukcje ale dowody też korzystamy z pojęcia implikacji wstępującej (por. paragraf 1.2.1.) i odpowiedniej wersji TD, którą można wzmocnić do postaci równoważności:

$$\varphi_1, \dots, \varphi_n \vdash \psi \quad \text{wtw} \quad \vdash \varphi_1 \rightarrow (\dots \rightarrow (\varphi_n \rightarrow \psi) \dots)$$

Z lewej do prawej równoważność ta stwierdza, że jeżeli zbudowaliśmy odpowiednią dedukcję, to mamy tym samym dowód odpowiadającej jej tezy o kształcie implikacji wstępującej. Z prawej do lewej równoważność ta podaje nam przepis na budowę dowodów też implikacyjnych: sukcesywnie wypisz wszystkie poprzedniki jako założenia i udowodnij na ich podstawie następnik, który jest ostatnią podformułą nie będącą implikacją (w schemacie reprezentowana przez ψ). Udowodnijmy w ten sposób tzw. sylogizm hipotetyczny (SH): $(p \rightarrow q) \rightarrow ((q \rightarrow r) \rightarrow (p \rightarrow r))$. Naszym φ_1 jest $p \rightarrow q$, φ_2 jest $q \rightarrow r$, $\varphi_3 = p$, a $\psi = r$. Dowód wygląda następująco:

1	$p \rightarrow q$	$z.$
2	$q \rightarrow r$	$z.$
3	p	$z.$
4	q	1, 3, $\rightarrow E$
5	r	2, 4, $\rightarrow E$

Dokonując wstępnego rozbioru tezy implikacyjnej na poprzedniki (są to założenia oznaczone w dowodzie przez $z.$) i dowodzony następnik trzeba uważać aby przez nieuwagę nie potraktować jako osobnego założenia implikacji, która nie jest samodzielnym poprzednikiem ale jego częścią. Przykładowo dowodząc tezy $((p \rightarrow q) \rightarrow r) \rightarrow (\neg p \rightarrow r)$ nie można jako założenia przyjąć $p \rightarrow q$; poprawnie przyjętymi założeniami dowodu wprost są $(p \rightarrow q) \rightarrow r$

oraz $\neg p$. Podobnie w dowodzie tezy $(p \rightarrow (q \rightarrow r)) \rightarrow ((p \rightarrow q) \rightarrow (p \rightarrow r))$ należy jako założenia przyjąć formuły $p \rightarrow (q \rightarrow r)$, $p \rightarrow q$ i p , natomiast błędem będzie rozbicie pierwszej formuły na p i $q \rightarrow r$. Krótko mówiąc, dowolne φ_i w schemacie też może być implikacją ale wypisujemy je w całości jako i -ty wiersz dowodu.

1.2.4 Dowodzenie nie wprost

Wszystkie dotąd zaprezentowane dedukcje i dowody były przeprowadzone wprost, tzn. punktem wyjścia były przesłanki (założenia) a punktem dojścia wniosek. Równie popularną formą dowodów (dedukcji) są dowody i dedukcje nie wprost, tradycyjnie określane jako dowody przeprowadzone przez „reductio ad absurdum” albo dowody apagogiczne. Istotą takiej strategii jest przyjęcie w punkcie wyjścia negacji wniosku jako dodatkowego założenia. Odtąd naszym celem staje się dedukcja sprzeczności czyli dowolnej pary zdań, z których jedno jest negacją drugiego. Oczywiście dedukcja wniosku (jak w dowodzeniu wprost) też spełnia ten warunek, gdyż wtedy parą zdań sprzecznych jest wniosek i jego negacja, którą założyliśmy wyżej. Generalnie chodzi jednak o to aby uzyskać jakąkolwiek parę wyrażeń sprzecznych.

Idea dowodzenia nie wprost opiera się na semantycznym założeniu o dwuwartościowości logicznej – każde zdanie jest bądź prawdziwe bądź fałszywe (wtedy jego negacja jest prawdziwa). Jeżeli wniosek istotnie wynika logicznie z przesłanek, to musi być prawdziwy gdy przesłanki są prawdziwe. Dołączenie jego negacji daje zatem zbiór, którego wszystkie elementy nie mogą być prawdziwe. Dedukcja pary zdań, które oba nie mogą być równocześnie prawdziwe, stanowi potwierdzenie tego faktu skoro wykorzystywane reguły gwarantują zachodzenie wynikania. W takiej sytuacji uznajemy, że zdanie, które – niejako na próbę – zanegowaliśmy zostało udowodnione na bazie przyjętych przesłanek.

Formalnie, dowodzenie nie wprost opiera się na twierdzeniu o dedukcji nie wprost (TDN) o postaci:

(TDN) Jeżeli $\Gamma, \neg\varphi \vdash \perp$, to $\Gamma \vdash \varphi$.

W celu łatwiejszego zapisu dowodów nie wprost dołączymy do naszego zestawu reguł dwie dodatkowe:

$(\perp D)$ $\varphi, \neg\varphi / \perp$
 $(\perp E)$ $\perp / \varphi$

Odtąd przyjmujemy, że jeżeli w dedukcji pojawi się $\perp$ to jest ona zakończona (z sukcesem). Do tego celu będziemy przede wszystkim używać reguły ($\perp D$). ($\perp E$) natomiast pokazuje nam pewną niebezpieczną własność sprzeczności, którą można scharakteryzować jako jej „eksplozywność”. Zauważmy, że φ nie jest zdaniem o określonej postaci; ze sprzeczności wynika dowolne zdanie!

Przeanalizujemy następujący przykład: $p \vee q, q \rightarrow r, r \rightarrow \neg q \vdash p$. Zauważmy, że po wypisaniu trzech przesłanek nasze dotąd stosowane reguły nie pozwalają na żaden ruch; każda z tych formuł potrzebuje dodatkowej przesłanki, aby można było zastosować odpowiednią regułę. Nie mając q ani r nie możemy zastosować ($\rightarrow E$) do drugiej lub trzeciej przesłanki, ale jeżeli dodamy $\neg p$ jako założenie nie wprost, to możemy zastosować ($\vee E$) do przesłanki pierwszej. Kolejne kroki wyglądają następująco:

1	$p \vee q$	p .
2	$q \rightarrow r$	p .
3	$r \rightarrow \neg q$	p .
4	$\neg p$	$zn.$
5	q	1, 4, $\vee E$
6	r	2, 5, $\rightarrow E$
7	$\neg q$	3, 6, $\rightarrow E$
8	$\perp$	5, 7, $\perp D$

Zauważmy, że dzięki dowodzeniu nie wprost możemy przeprowadzać dowody tez, które nie są implikacjami; stosowne przykłady podamy dalej. Zresztą nawet w przypadku tez implikacyjnych w wielu przypadkach dowód nie wprost wydaje się jedynym możliwym w naszym systemie DN. Przykładowe tezy tego typu to $\vdash p \rightarrow p$, $\vdash p \rightarrow (q \rightarrow p)$ (podstawienie aksjomatu 1) czy $\vdash p \rightarrow (\neg p \rightarrow q)$ (tzw. prawo Dunsza Szkota). Ich natychmiastowe dowody nie wprost uzyskujemy przez wypisanie założeń wraz z założeniem nie wprost, które natychmiast daje sprzeczność. W przypadku prawa Dunsza Szkota możliwy jest też dowód wprost, gdyż sprzeczność dają już założenia p i $\neg p$ z których dedukujemy $\perp$ a potem q z pomocą naszych reguł dla $\perp$.

Zauważmy, że wprowadzone w tym paragrafie dodatkowe reguły umożliwiają nie tylko inny sposób wykazywania poprawności rozumowań ale również wykazywanie, że pewne zbiory zdań są sprzeczne. Np. sprzeczny jest zbiór: $p \rightarrow \neg q, q \vee r, p \wedge \neg r$ co można wykazać w następujący sposób:

1	$p \rightarrow \neg q$	$p.$
2	$q \vee r$	$p.$
3	$p \wedge \neg r$	$p.$
4	p	3, $\wedge E$
5	$\neg r$	3, $\wedge E$
6	$\neg q$	1, 4, $\rightarrow E$
7	r	2, 6, $\vee E$
8	$\perp$	5, 7, $\perp D$

W przykładzie tym nie wprowadziliśmy żadnego założenia nie wprost; sprzeczność generują same formuły z analizowanego zbioru.

1.2.5 Dowody a dedukcje

Przepis na budowanie dowodów założeniowych z wykorzystaniem schematu implikacji wstępującej przydaje się również przy konstruowaniu dedukcji wniosków z przesłanek dla poprawnych rozumowań. Rozważmy przykład: $p \rightarrow (q \rightarrow r), p \wedge s \rightarrow t, r \wedge v \rightarrow s \vdash p \rightarrow (q \wedge v \rightarrow t)$. Łatwo zauważyć po wypisaniu trzech przesłanek, że nie można do nich zastosować żadnej reguły, pozwalającej pomyśleć o dedukcji wniosku. Możemy oczywiście stosować ($\wedge D$) lub ($\vee D$) ale te reguły nie dadzą nam poszukiwanego efektu. Wniosek jest jednak wstępującą implikacją i można potraktować go jak tezę do udowodnienia, tzn. wypisać jako założenia kolejne poprzedniki, a to pozwala nam na dedukcję następnika:

1	$p \rightarrow (q \rightarrow r)$	$p.$
2	$p \wedge s \rightarrow t$	$p.$
3	$r \wedge v \rightarrow s$	$p.$
4	p	$z.$
5	$q \wedge v$	$z.$
6	$q \rightarrow r$	1, 4, $\rightarrow E$
7	q	5, $\wedge E$
8	v	5, $\wedge E$
9	r	6, 7, $\rightarrow E$
10	$r \wedge v$	8, 9, $\wedge D$
11	s	3, 10, $\rightarrow E$
12	$p \wedge s$	4, 11, $\wedge D$
13	t	2, 12, $\rightarrow E$

W podobny sposób możemy konstruować dedukcje nie wprost. Oto przykład dedukcji dla $p \rightarrow (\neg q \rightarrow r \vee s)$, $\neg s \vdash p \rightarrow (\neg r \rightarrow q)$, w której również oprócz samych przesłanek użyjemy założeń zaczerpniętych z analizy wniosku.

1	$p \rightarrow (\neg q \rightarrow r \vee s)$	$p.$
2	$\neg s$	$p.$
3	p	$z.$
4	$\neg r$	$z.$
5	$\neg q$	$zn.$
6	$\neg q \rightarrow r \vee s$	$1, 3, \rightarrow E$
7	$r \vee s$	$5, 6, \rightarrow E$
8	s	$4, 7, \vee E$
9	$\perp$	$2, 8, \perp D$

W tym przypadku widać, że obie przesłanki nie wystarczają do dalszej dedukcji gdyż $\neg s$ nie jest poprzednikiem pierwszej przesłanki. Wypisanie założeń będących poprzednikami wniosku też nie wystarcza – gdybyśmy nie dodali negacji następnika, to nasza dedukcja zatrzymałaby się na wierszu 6.

Stosując powyższą strategię w konstrukcji dedukcji, zwłaszcza takich gdzie wnioskiem jest implikacja, musimy pamiętać, że tylko wniosek, czyli formuła po prawej stronie $\vdash$, pozwala na analizę w terminach implikacji wstępującej i uzyskanie w ten sposób założeń umożliwiających konstrukcję dedukcji. Nie wolno w ten sposób traktować przesłanek, czyli formuł występujących po lewej stronie $\vdash$. Np. pierwsza przesłanka z ostatniego przykładu też jest implikacją wstępującą ale niedopuszczalnym błędem byłoby potraktowanie jej w ten sposób i wypisanie w osobnych wierszach formuł p i $\neg q$. Jedynym uproszczeniem jakiego możemy się dopuścić przy wypisywaniu przesłanek dedukcji jest rozbitcie przesłanki o postaci koniunkcji na jej składniki. Jest to skrót, którego poprawność uzasadnia reguła $(\wedge E)$.

TD pokazuje, że to czy mówimy o dowodzie tezy czy dedukcji wniosku z przesłanek jest w pewnym sensie względne. Również podawane przykłady pokazują, że ten sam ciąg formuł, jeżeli wykasujemy w nim w kolumnie uzasadnień oznaczenia $p.$ i $z.$ może być potraktowany jako dowód lub jako dedukcję wniosku z przesłanek. Przykładowo, powyżej dowiedliśmy sylogizmu hipotetycznego (SH) jako tezy, ale równie dobrze dowód ten można odczytać jako dedukcję uzasadniającą dla dowolnego z poniższych schematów rozumowań:

$$\begin{aligned}
p \rightarrow q &\vdash (q \rightarrow r) \rightarrow (p \rightarrow r) \\
p \rightarrow q, q \rightarrow r &\vdash p \rightarrow r \\
p \rightarrow q, q \rightarrow r, p &\vdash r
\end{aligned}$$

wystarczy zamiast *z.* wstawić *p.* w kolumnie uzasadnień dla odpowiednich wierszy. Z tego względu w dalszym ciągu będziemy dla uproszczenia rozważań generalnie mówić o dowodach.

1.2.6 Równoważności

W systemach formalnych często występują twierdzenia o postaci równoważności. Ponieważ nie wprowadzaliśmy spójnika równoważności jako pierwotnego w naszym języku więc podobnie w systemie DN użyjemy dwóch reguł o charakterze definicyjnym, w tym sensie, że redukują one użycie spójnika równoważności do użycia innych spójników pierwotnych (w tym przypadku implikacji). Odpowiednie reguły to:

$$\begin{aligned}
(\leftrightarrow E) \quad & \varphi \leftrightarrow \psi / \varphi \rightarrow \psi \quad \text{oraz} \quad \varphi \leftrightarrow \psi / \psi \rightarrow \varphi \\
(\leftrightarrow D) \quad & \varphi \rightarrow \psi, \psi \rightarrow \varphi / \varphi \leftrightarrow \psi
\end{aligned}$$

Zauważmy, że obie reguły można wyprowadzić z definicji $\leftrightarrow$; pierwszą przez zastosowanie ($\wedge E$) a drugą przez zastosowanie ($\wedge D$) do obu przesłanek. W szczególności druga reguła dostarcza nam ogólnej strategii dowodzenia równoważności: należy skonstruować niezależne dowody dwóch implikacji. Przykładowo możemy w ten sposób udowodnić prawo eksportacji/importacji $p \rightarrow (q \rightarrow r) \leftrightarrow p \wedge q \rightarrow r$ rozbijając dowód na dwie części:

($\rightarrow$)

1	$p \rightarrow (q \rightarrow r)$	$z.$
2	$p \wedge q$	$z.$
3	p	2, $\wedge E$
4	q	2, $\wedge E$
5	$q \rightarrow r$	1, 3, $\rightarrow E$
6	r	4, 5, $\rightarrow E$

($\leftarrow$)

1	$p \wedge q \rightarrow r$	$z.$
2	p	$z.$
3	q	z
4	$p \wedge q$	$2, 3, \wedge D$
5	r	$1, 4, \rightarrow E$

Tezę otrzymujemy z obu dowodów przez ($\leftrightarrow D$). Nawiasem mówiąc ostatnia teza pozwala formalnie uzasadnić zasygnalizowany w poprzednim paragrafie skrót pozwalający na rozpisanie koniunkcji w poprzedniku tezy na jej składniki jako osobne założenia dowodu. Formalnie możemy to ująć w poszerzonym schemacie (TD):

$$\varphi_1, \dots, \varphi_n \vdash \psi \text{ wtw } \vdash \varphi_1 \rightarrow (\dots \rightarrow (\varphi_n \rightarrow \psi) \dots) \text{ wtw } \vdash \varphi_1 \wedge \dots \wedge \varphi_n \rightarrow \psi$$

1.3 Zaawansowana dedukcja

1.3.1 Stosowanie założeń dodatkowych

W dowodach i dedukcjach często oprócz przesłanek czy założeń pierwotnych (wprost i nie wprost) używa się założeń dodatkowych. Przykładowo, chcąc udowodnić $p \vee q \rightarrow r, (p \rightarrow r) \rightarrow s \vdash s \vee t$ po wypisaniu obu założeń jesteśmy w kłopotcie, gdyż oba są implikacjami a nie mamy żadnego poprzednika. Dodanie założenia nie wprost tj. $\neg(s \vee t)$ w niczym naszej sytuacji nie zmienia. Aby wydedukować $s \vee t$ potrzebujemy s (gdyż t wogóle w założeniach nie występuje), a tę formułę można dostać z drugiego założenia przez ($\rightarrow E$). Potrzebna jest jednak formuła $p \rightarrow r$; skąd ją uzyskać? Pomocne może być odwołanie się do TD, gdyż wiemy, że aby ze zbioru przesłanek wydedukować implikację wystarczy wydedukować jej następnik przy (dodatkowym) założeniu poprzednika. Dedukcja dla naszego przykładu wygląda wtedy tak:

1	$p \vee q \rightarrow r$	$p.$
2	$(p \rightarrow r) \rightarrow s$	$p.$
3	p	$z.$
4	$p \vee q$	$3, \vee D$
5	r	$1, 4, \rightarrow E$
6	$p \rightarrow r$	$1, 3 - 5, \text{TD}$
7	s	$2, 6, \rightarrow E$
8	$s \vee t$	$7, \vee D$

Uzasadnienie dla wiersza 6 należy rozumieć następująco: ponieważ $p \vee q \rightarrow r, p \vdash r$, zatem, przez TD, $p \vee q \rightarrow r \vdash p \rightarrow r$. Zauważmy, że w ten sposób, nieformalnie wprowadziliśmy do zestawu naszych reguł brakującą dotąd regułę dołączania implikacji, która odpowiada zastosowaniu TD. Należy jednak podkreślić, że użycie takich dodatkowych założeń wiąże się z pewnym ryzykiem. Przeanalizujemy poniższą „dedukcję” $(p \rightarrow q) \wedge p$ z przesłanki $p \vee r \rightarrow q$:

1	$p \vee r \rightarrow q$	$p.$
2	p	$z.$
3	$p \vee r$	2, $\vee D$
4	q	1, 3, $\rightarrow E$
5	$p \rightarrow q$	1, 2 – 4, TD
6	$(p \rightarrow q) \wedge p$	2, 5, $\wedge D$

Jednak $(p \rightarrow q) \wedge p$ nie wynika logicznie z $p \vee r \rightarrow q$, a przecież poprawny system dedukcyjny ma pozwalać tylko na dedukcję wniosków logicznie wynikających z przesłanek. Na szczęście podany wyżej ciąg formuł nie jest poprawną dedukcją; błąd popełniony jest w ostatnim wierszu przy zastosowaniu ($\wedge D$). Problem polega na tym, że w momencie wykorzystania TD (czy też reguły dołączania implikacji) dodatkowe założenie traci swą moc. Ani ono, ani żadna formuła, którą z tego założenia wydedukowaliśmy nie jest już dalej w dedukcji lub dowodzie dostępne. Należy zatem wprowadzić dodatkowe środki formalne, które uniemożliwią nam popełnianie tego typu błędów.

1.3.2 Poddowody warunkowe

Przyjmujemy, że założenie dodatkowe otwiera nam w obrębie naszego dotychczasowego dowodu, nowy dowód, który z chwilą zastosowania reguły ($\rightarrow D$) (tzn. TD) zostaje zakończony. Ten nowy dowód jest poddowodem dotychczasowego, w tym sensie, że w jego obrębie można wykorzystywać wszystkie formuły, które pojawiły się w dowodzie wcześniej, przed wprowadzeniem założenia dodatkowego (tj. w dowodzie nadrzędnym). Natomiast po jego zakończeniu (poprzez zastosowanie ($\rightarrow D$) lub innych reguł podobnie działających, które opiszemy niżej) nie możemy w dowodzie nadrzędnym użyć ani tego założenia ani żadnej formuły, którą z niego wydedukowaliśmy. Działanie założeń dodatkowych i opartych na nich poddowodów można porównać do działania pożyczek; wprowadzenie założenia dodatkowego jest „pożyczeniem” formuły, której dotąd nie mieliśmy. Nowe środki pozwalają

nam otrzymać więcej niż było dostępne z wykorzystaniem dotychczasowych ale trzeba się z nich w końcu rozliczyć – oddać pożyczkę. Zastosowanie reguły takiej jak $(\rightarrow D)$ likwiduje nasze „zadłużenie” zamykając odpowiedni poddowód.

Aby uniknąć problemów takich jak w powyższym przykładzie błędnej dedukcji musimy wprowadzić dodatkowe środki, które umożliwiają nam na bieżąco kontrolę poddowodu. W naszym przypadku będzie to oryginalne rozwiązanie Jaśkowskiego, użyte również przez Słupeckiego i Borkowskiego. Wraz z wprowadzeniem dodatkowego założenia „zamrażamy” numerację dowodu nadrzędnego i zaczynamy w jej obrębie osobną numerację dla wierszy poddowodu. Przykładowo, jeżeli ostatni wiersz dowodu poprzedzający wprowadzenie dodatkowego założenia ma numer n , to założenie otrzymuje numer $n.1$ a kolejne wiersze numery $n.2, n.3, \dots$. Stosując regułę $(\rightarrow D)$ kończymy poddowód, wprowadzona implikacja należy już do dowodu nadrzędnego zatem jej numer to $n+1$. Korzystając z tej konwencji notacyjnej podaną wyżej dedukcję zapiszemy następująco:

1	$p \vee q \rightarrow r$	$p.$
2	$(p \rightarrow r) \rightarrow s$	$p.$
2.1	p	$zd.$
2.2	$p \vee q$	2.1, $\vee D$
2.3	r	1, 2.2, $\rightarrow E$
3	$p \rightarrow r$	2.1 – 2.3, $\rightarrow D$
4	s	2, 3, $\rightarrow E$
5	$s \vee t$	4, $\vee D$

$zd.$ oznacza założenie dodatkowe, przy $(\rightarrow D)$ podajemy jako uzasadnienie numery całego poddowodu od założenia dodatkowego do wiersza zawierającego następnik wydedukowanej implikacji. Próba przeprowadzenia niepoprawnej „dedukcji” z poprzedniego paragrafu wygląda następująco:

1	$p \vee r \rightarrow q$	$p.$
1.1	p	$zd.$
1.2	$p \vee r$	1.1, $\vee D$
1.3	q	1, 1.2, $\rightarrow E$
2	$p \rightarrow q$	1.1 – 1.3, $\rightarrow D$
3	$(p \rightarrow q) \wedge p$	2, 1.1, $\wedge D$

w wierszu 3 użyliśmy niepoprawnie $(\wedge D)$, gdyż p nie jest już dostępne z chwilą gdy zastosowaliśmy $(\rightarrow D)$ i wróciliśmy do nadrzędnej dedukcji.

Generalnie, jeżeli stosujemy dowolną regułę inferencji do (przynajmniej jednej) przesłanki, która występuje w poddowodzie, to wniosek nie może się znajdować w dowodzie nadrzędnym.

Zauważmy, że $(\rightarrow D)$ nie jest regułą inferencji; nie dołącza nam nowej formuły na podstawie formuł wyżej występujących. Raczej dołącza nową formułę do dowodu nadrzędnego na podstawie całej dedukcji, którą skonstruowaliśmy w obrębie dowodu podrzędnego. Reguły takie określamy mianem reguł konstrukcji dowodu. Mówiąc najprościej przekształcają one jedne dowody (dedukcje) na inne, podczas gdy reguły inferencji przekształcają jedne formuły na inne.

1.3.3 Poddowody nie wprost

Inną regułą konstrukcji dowodu jest reguła $(\neg E)$, która pozwala wydedukować z dowolnego zestawu przesłanek (założeń) formułę φ , jeżeli z dodatkowego założenia $\neg\varphi$ wydedukowaliśmy $\perp$. Jak widać stanowi ona uogólnienie metody dowodu nie wprost. Jako przykład rozważmy dedukcję dla $p \rightarrow q \vee r, p \wedge (r \rightarrow \neg p) \vdash q \wedge p$:

1	$p \rightarrow q \vee r$	$p.$
2	$p \wedge (r \rightarrow \neg p)$	$p.$
3	p	$2, \wedge E$
4	$r \rightarrow \neg p$	$2, \wedge E$
5	$q \vee r$	$1, 3, \rightarrow E$
5.1	$\neg q$	$zdn.$
5.2	r	$5, 5.1, \vee E$
5.3	$\neg p$	$4, 5.2, \rightarrow E$
5.4	$\perp$	$3, 5.3, \perp D$
6	q	$5.1 - 5.4, \neg E$
7	$q \wedge p$	$3, 6, \wedge D$

W podanym przykładzie to, czego nam brakuje po wypisaniu przesłanek, to nie jakaś implikacja ale formuła q . Zatem reguła $(\rightarrow D)$ jest bezużyteczna, ale jako założenie dodatkowe możemy wprowadzić $\neg q$ i z jego pomocą wydedukować $\perp$, co dzięki zastosowaniu $(\neg E)$ daje nam w dowodzie nadrzędnym poszukiwane q . Założenie dodatkowe nie wprost oznaczamy przez *zdn.* a dla poddowodu stosujemy taką samą konwencję notacyjną jak przy $(\rightarrow D)$. Zauważmy, że $\perp$ pojawia się nie w dowodzie głównym ale w poddowodzie, i co za tym idzie zakończeniu ulega poddowód a nie dowód główny.

Reguła ta pozwala również zbudować dowody wielu tez nie będących implikacjami. Oto dowód prawa niesprzeczności w postaci $\neg(p \wedge \neg p)$:

1	$\neg\neg(p \wedge \neg p)$	<i>zn.</i>
1.1	$\neg(p \wedge \neg p)$	<i>zdn.</i>
1.2	$\perp$	1, 1.1, $\perp D$
2	$p \wedge \neg p$	1.1 – 1.2, $\neg E$
3	p	2, $\wedge E$
4	$\neg p$	2, $\wedge E$
5	$\perp$	3, 4, $\perp D$

oraz prawa wyłączonego środka $\neg p \vee p$:

1	$\neg(\neg p \vee p)$	<i>zn.</i>
1.1	$\neg p$	<i>zdn.</i>
1.2	$\neg p \vee p$	1.1, $\vee D$
1.3	$\perp$	1, 1.2, $\perp D$
2	p	1.1 – 1.3, $\neg E$
3	$\neg p \vee p$	2, $\vee D$
4	$\perp$	1, 3, $\perp D$

1.3.4 Poddowody wielokrotne i zagnieżdżone

W obrębie dowodu możemy otwierać dowolnie wiele poddowodów, zarówno niezależnych od siebie jak i wzajemnie podporządkowanych. Spójrzmy na dowód tezy $(p \vee q \rightarrow r \wedge s) \rightarrow (p \rightarrow r) \wedge (q \rightarrow s)$:

1	$p \vee q \rightarrow r \wedge s$	<i>z.</i>
1.1	p	<i>zd.</i>
1.2	$p \vee q$	1.1, $\vee D$
1.3	$r \wedge s$	1, 1.2, $\rightarrow E$
1.4	r	1.3, $\wedge E$
2	$p \rightarrow r$	1.1 – 1.4, $\rightarrow D$
2.1	q	<i>zd.</i>
2.2	$p \vee q$	2.1, $\vee D$
2.3	$r \wedge s$	1, 2.2, $\rightarrow E$
2.4	s	2.3, $\wedge E$
3	$q \rightarrow s$	2.1 – 2.4, $\rightarrow D$
4	$(p \rightarrow r) \wedge (q \rightarrow s)$	2, 3, $\wedge D$

W tym przypadku, następnik jest koniunkcją więc nie można z niego otrzymać dalszych założeń (rozbijanie koniunkcji na składniki można zastosować tylko do przesłanek – por. uwaga w podrozdziale 1.3.3.). Jednak każdy ze składników koniunkcji jest implikacją a to oznacza, że można każdy tych składników wydedukować osobno korzystając z $(\rightarrow D)$ a następnie zastosować do nich $(\wedge D)$. Każdy z poddowodów jest niezależny od drugiego, co oznacza, że nie możemy w obrębie drugiego wykorzystywać tych formuł, które pojawiły się w poprzednim. Np. byłoby błędem wpisanie w wierszu 2.2 formuły s jako wydedukowanej z wiersza 1.3. Pierwszy z poddowodów został już zakończony z chwilą zastosowania $(\rightarrow D)$ i nie mamy prawa wykorzystywać dalej żadnej formuły, która w jego obrębie się pojawiła. Choć wydaje się to zbędną pedanterią, w obrębie drugiego poddowodu musimy powtórzyć niemal identyczną dedukcję aby dowieść drugiej implikacji.

Należy podkreślić, że generalnie w takich przypadkach poddowody są zależne od dowodu nadrzędnego ale niezależne od siebie. Oznacza to, że w każdym z nich mamy prawo używać formuł z dowodu nadrzędnego, ale nie mamy prawa użyć w drugim z nich żadnej formuły, która jest elementem poprzedniego, zamkniętego już poddowodu. Nieco inna sytuacja ma miejsce gdy dany poddowód nie został jeszcze zamknięty, ale w jego obrębie decydujemy się wprowadzić kolejne założenie dodatkowe. Wtedy nowo otwarty poddowód stanowi część poprzedniego poddowodu jak i nadrzędnego dowodu. Jego (podwójną) zależność wizualnie oddaje użycie dodatkowej kolumny liczb w numeracji wiersza. Zilustrujmy to dowodem nie wprost tezy $\neg(p \rightarrow q) \rightarrow p \wedge \neg q$:

1	$\neg(p \rightarrow q)$	$z.$
2	$\neg(p \wedge \neg q)$	$zn.$
2.1	p	$zd.$
2.1.1	$\neg q$	$zdn.$
2.1.2	$p \wedge \neg q$	2.1, 2.1.1, $\wedge D$
2.1.3	$\perp$	2, 2.1.2, $\perp D$
2.2	q	2.1.1 – 2.1.3, $\neg E$
3	$p \rightarrow q$	2.1 – 2.2, $\rightarrow D$
4	$\perp$	1, 3, $\perp D$

W powyższym przykładzie decydujemy się na dowód nie wprost a sprzeczność chcemy uzyskać przez dedukcję $p \rightarrow q$. W tym celu zaczynamy poddowód z p jako założeniem dodatkowym, a celem tego poddowodu jest dedukcja q . Aby to zrobić zaczynamy kolejny poddowód (nie wprost), którego

założeniem dodatkowym jest $\neg q$, gdyż jeżeli uzyskamy w nim $\perp$, to da nam to poszukiwane q w poddowodzie nadrzędnym. Zaczynając kolejny poddowód w obrębie niezakończonego dotąd innego poddowodu stosujemy takie samo oznaczanie wierszy, tzn. „zamrażamy” ostatni numer wiersza dowodu nadrzędnego i zaczynamy kolejną kolumnę liczb.

Biorąc pod uwagę omówione wyżej sposoby konstruowania poddowodów możemy schematycznie obie reguły konstrukcji dowodu zapisać w następujący sposób:

$$\begin{array}{lll}
 \vdots & & \vdots \\
 \sigma.k.1 & \varphi & zd. \\
 \vdots & & \vdots \\
 \sigma.k.n & \psi & \\
 \sigma.k+1 & \varphi \rightarrow \psi & \sigma.k.1 - \sigma.k.n, \rightarrow D
 \end{array}
 \qquad
 \begin{array}{lll}
 \vdots & & \vdots \\
 \sigma.k.1 & \neg\varphi & zdn. \\
 \vdots & & \vdots \\
 \sigma.k.n & \perp & \\
 \sigma.k+1 & \varphi & \sigma.k.1 - \sigma.k.n, \neg E
 \end{array}$$

gdzie σ oznacza skończony (możliwe, że pusty) ciąg liczb naturalnych, a k, n dowolne liczby.

Generalnie, możemy w trakcie konstrukcji dowodu dołączać tyle dodatkowych założeń ile chcemy i o takim kształcie jaki potrzebujemy, należy jednak pamiętać, że każde z nich otwiera nowy poddowód a z poddowodów można na poziom dowodów nadrzędnych dostać się wyłącznie za pomocą wyznaczonych do tego reguł konstrukcji dowodu. Jeżeli celem dowodu nadrzędnego jest np. dedukcja formuły φ i taka formuła pojawi się w poddowodzie, to nasz cel nadal nie został osiągnięty; została ona wprowadzona, ale z pomocą dodatkowych założeń. To powinno narzucać pewne ograniczenia w „radosnym” dołączaniu dodatkowych założeń do dowodu – dołączać je powinniśmy tylko wówczas, gdy widzimy, że dana formuła pozwoli nam na dedukcję innych potrzebnych nam formuł i gdy widzimy, że jest jakaś szansa, żeby nowo otwarty poddowód kiedyś zakończyć.

Obie reguły wiążą się z następującymi strategiami dowodzenia:

1. Jeżeli w danym momencie konstruowania dedukcji, uświadamiamy sobie, że potrzebna jest nam implikacja $\varphi \rightarrow \psi$, to wprowadzamy φ jako założenie dodatkowe i staramy się w nowym poddowodzie wydedukować ψ .

2. Jeżeli w danym momencie konstruowania dedukcji, uświadamiamy sobie, że potrzebna jest nam formuła φ , to wprowadzamy $\neg\varphi$ jako założenie dodatkowe nie wprost i staramy się w nowym poddowodzie wydedukować $\perp$.

1.4 Dodatkowe środki dowodowe

Ważnym sposobem skracania i ułatwiania konstrukcji dowodów i dedukcji jest wykorzystywanie w nich udowodnionych uprzednio tez oraz rozmaitych reguł wtórnych. W przypadku wykorzystywania wcześniej dowiedzionych tez odwołujemy się dodatkowo do użycia tzw. reguły podstawiania. Wstawienie do danego dowodu tezy, którą wyżej udowodniliśmy działa na zasadzie użycia „macro” – pojedynczy wiersz zawierający tezę zawsze możemy rozwinąć wstawiając do aktualnego dowodu wcześniej przeprowadzony dowód użyciej tezy. Jednak do efektywnego użycia tez potrzebujemy reguły (a raczej metareguły) podstawiania, która w dowolnej tezie pozwala za wszystkie wystąpienia danej zmiennej podstawić inną formułę. Np. udowodniliśmy wyżej prawo niesprzeczności w postaci $\neg(p \wedge \neg p)$ jednak w innym dowodzie możemy potrzebować go w sformułowaniu $\neg(q \wedge \neg q)$ albo $\neg((p \rightarrow q) \wedge \neg(p \rightarrow q))$. W pierwszym przypadku dokonaliśmy podstawienia q za p , w drugim $p \rightarrow q$. Formalnie efekt podstawienia będziemy zapisywać następująco: $\varphi[p/\psi]$, co oznacza, że wszystkie wystąpienia zmiennej p w formule φ zostały zamienione na wystąpienia formuły ψ . Musimy pamiętać, że regułę podstawiania stosujemy tylko do tez i podstawiamy tylko za zmienne. Natomiast podstawiać możemy nie tylko zmienne ale formuły dowolnego kształtu i długości, pamiętając tylko, że podstawiamy za wszystkie wystąpienia danej zmiennej i za każde podstawiamy tę samą formułę. Uzyskane w ten sposób wyrażenie też jest tezą.

1.4.1 Reguły wtórne

Innym ważnym sposobem zwiększania ekonomii dowodzenia jest dołączanie reguł wtórnych. Zastosowanie każdej reguły wtórnej też przypomina użycie zdefiniowanego wcześniej „macro” i zawsze można je w dowodzie zastąpić dedukcją uzasadniającą poprawność tej reguły. Oczywiście nie należy przesadzać z ilością takich reguł. Pamiętajmy, że na bazie każdej dedukcji $\Gamma \vdash \varphi$ możemy dołączyć kolejną regułę wtórną Γ / φ ale przydatne są jedynie takie, które mają prosty schemat łatwy do zapamiętania i potencjalnie często występujący w dowodach. Oto kilka przydatnych reguł wtórnych o prostej postaci:

- $(SH) \quad \varphi \rightarrow \psi, \psi \rightarrow \chi / \varphi \rightarrow \chi$
 $(PNE) \quad \neg\neg\varphi / \varphi$
 $(MT) \quad \varphi \rightarrow \psi, \neg\psi / \neg\varphi$
 $(MTP) \quad \varphi \vee \psi, \neg\psi / \varphi$
 $(DI) \quad \varphi / \psi \rightarrow \varphi \text{ oraz } \neg\varphi / \varphi \rightarrow \psi$

Pierwsza z nich to reguła sylogizmu hipotetycznego, dla której uzasadnienie stanowi dedukcja tezy (SH) podana w paragrafie 1.3.3. Druga to reguła eliminacji podwójnej negacji, trzecia to tzw. modus tollendo tollens (tryb obalający przez obalenie) a czwarta to wariant ($\vee E$) (modus tollendo ponens), w którym negowany jest drugi argument alternatywy. (DI) to dwie dodatkowe reguły dołączania implikacji oparte na aksjomacie 1 i prawie Dunsza Szkota. Uzasadnienia poprawności tych reguł są proste; zademonstrujemy uzasadnienie dla (MT), w którym wykorzystamy (PNE):

1	$p \rightarrow q$	$p.$
2	$\neg q$	$p.$
3	$\neg\neg p$	$zn.$
4	p	3, PNE
5	q	1, 4, $\rightarrow E$
6	$\perp$	2, 5, $\perp D$

Dowód jest prosty, zwłaszcza, że dla jego skrócenia użyliśmy (PNE). Bez zastosowania tej reguły wtórnej musielibyśmy pomiędzy wierszem 3 i 4 wstawić poddowód zawierający dwa wiersze: 3.1 $\neg p \text{ } zn.$ i 3.2 $\perp$ uzyskane przez ($\perp D$) z wierszy 3 i 3.1. Oczywiście uzasadnieniem dla wiersza 4 byłaby wtedy reguła ($\neg E$). Zauważmy też, że uzasadnienia są dedukcjami w języku, natomiast schematy reguł wprowadzamy w metajęzyku. Jest tak dlatego, że w danej dedukcji, która uzasadnia poprawność reguły zmienne $p, q, \dots$ moglibyśmy zastąpić dowolnymi formułami i wynik byłby nadal poprawny, oczywiście pod warunkiem, że za każde wystąpienie danej zmiennej podstawimy tę samą formułę. Jest to rodzaj poszerzenia działania reguły podstawiania z tez na schematy poprawnych rozumowań.

Reguła (PNE) znajduje rozliczne zastosowania, zwłaszcza w dowodach nie wprost dla rozmaitych tez nie będących implikacjami. Łatwo zauważyć, że podany w paragrafie 1.4.3. dowód prawa niesprzeczności można też uprościć eliminując zastosowanie ($\neg E$). Korzystając z (PNE) można łatwo udowodnić też regułę odwrotną: $\varphi / \neg\neg\varphi$. W dalszym ciągu będziemy się po prostu odwoływać do obustronnej reguły podwójnej negacji o postaci:

$$(PN) \quad \neg\neg\varphi // \varphi$$

W podanym schemacie $//$ oznacza, że reguła jest poprawna zarówno z lewej do prawej jak i z prawej do lewej. Odwołując się do (PN) można uzasadnić wiele wariantów reguł, w których schematach występuje negacja przynajmniej jednej przesłanki, np:

$$(MT') \quad \varphi \rightarrow \neg\psi, \psi / \neg\varphi$$

$$(MTP') \quad \varphi \vee \neg\psi, \psi / \varphi$$

(czytelnik łatwo znajdzie inne warianty)

Co więcej można w dogodny sposób uogólnić zarówno regułę dowodu nie wprost jak i $(\neg E)$. W obu przypadkach, jeżeli formuła, którą chcemy (po zanegowaniu) przyjąć jako zn. (zdn.) jest postaci $\neg\varphi$, to zamiast wprowadzać jako założenie $\neg\neg\varphi$ możemy od razu wpisać φ . W ten sposób druga z reguł staje się formalnie rzecz biorąc regułą $(\neg D)$, gdyż po dedukcji $\perp$ w poddowodzie, do dowodu nadrzędnego dołączamy $\neg\varphi$.

1.4.2 Reguły obustronne

Wartość użytkowa reguł obustronnych, takich jak (PN) jest bardzo duża. W wielu systemach DN są one wprowadzane jako osobna grupa pierwotnych reguł ze względu na swą użyteczność, dlatego poświęcimy im nieco więcej uwagi. Oto lista najważniejszych reguł wtórnych tego typu:

$$(NA) \quad \neg(\varphi \vee \psi) // \neg\varphi \wedge \neg\psi$$

$$(NK) \quad \neg(\varphi \wedge \psi) // \neg\varphi \vee \neg\psi$$

$$(NI) \quad \neg(\varphi \rightarrow \psi) // \varphi \wedge \neg\psi$$

$$(IA) \quad \varphi \rightarrow \psi // \neg\varphi \vee \psi \text{ oraz } \neg\varphi \rightarrow \psi // \varphi \vee \psi$$

$$(PK) \quad \varphi \wedge \psi // \psi \wedge \varphi$$

$$(PA) \quad \varphi \vee \psi // \psi \vee \varphi$$

$$(LK) \quad \varphi \wedge (\psi \wedge \chi) // (\varphi \wedge \psi) \wedge \chi$$

$$(LA) \quad \varphi \vee (\psi \vee \chi) // (\varphi \vee \psi) \vee \chi$$

$$(DKA) \quad \varphi \wedge (\psi \vee \chi) // (\varphi \wedge \psi) \vee (\varphi \wedge \chi)$$

$$(DAK) \quad \varphi \vee (\psi \wedge \chi) // (\varphi \vee \psi) \wedge (\varphi \vee \chi)$$

$$(RED) \quad \varphi // \varphi \wedge \varphi // \varphi \vee \varphi // \varphi \wedge \top // \varphi \vee \perp$$

Szczególnie przydatne są reguły (NA) , (NK) (negacji koniunkcji i alternatywy, zwane często regułami DeMorgana) oraz (NI) (negacja implikacji), które pozwalają (w wersji od lewej do prawej) przekształcić zanegowaną formułę złożoną na formułę niezanegowaną. Para reguł zamiany implikacji na

alternatywę (IA) oraz pozostałe reguły o charakterze algebraicznym (przemienności, łączności, dystrybucji) mają duże znaczenie dla sprowadzania dowolnych formuł do postaci normalnych (por. Malinowski [17]). Ostatni wiersz zbiorczo podaje kilka reguł redukcji wyrażenia złożonego do prostego. W obu ostatnich regułach redukcji (od prawej strony) użyliśmy $\top$ i $\perp$ w nieco ogólniejszym znaczeniu niż w definicjach podanych w paragrafie 1.1.1; $\top$ oznacza tutaj dowolną tezę a $\perp$ dowolną formułę sprzeczną.

Uzasadnienie dla (NI) (w jedną stronę) podane było w paragrafie 1.4.4. Tutaj zaprezentujemy dowód uzasadniający dla (NA):

($\rightarrow$)

1	$\neg(p \vee q)$	$z.$
1.1	p	$zdn.$
1.2	$p \vee q$	1.1, $\vee D$
1.3	$\perp$	1, 1.2, $\perp D$
2	$\neg p$	1.1 – 1.3, $\neg D$
2.1	q	$zdn.$
2.2	$p \vee q$	2.1, $\vee D$
2.3	$\perp$	1, 2.2, $\perp D$
3	$\neg q$	2.1 – 2.3, $\neg D$
4	$\neg p \wedge \neg q$	2, 3, $\wedge D$

($\leftarrow$)

1	$\neg p \wedge \neg q$	$z.$
2	$p \vee q$	$zn.$
3	$\neg p$	1, $\wedge E$
4	$\neg q$	1, $\wedge E$
5	q	2, 3, $\vee E$
6	$\perp$	4, 5, $\perp D$

W obu dowodach użyliśmy wtórnych reguł dowodzenia nie wprost. Dowód uzasadniający dla (NK) jest prosty przy użyciu (NA).

Dowody uzasadniające reguły przemienności i łączności są łatwe dla koniunkcji ale nie dla alternatywy. Oto dowód dla (jednego kierunku) (LA):

1	$p \vee (q \vee r)$	$z.$
2	$\neg((p \vee q) \vee r)$	$zn.$
2.1	p	$zdn.$
2.2	$p \vee q$	2.1, $\vee D$
2.3	$(p \vee q) \vee r$	2.2, $\vee D$
2.4	$\perp$	2, 2.3, $\perp D$
3	$\neg p$	2.1 – 2.4, $\neg D$
4	$q \vee r$	1, 3, $\vee E$
4.1	q	$zdn.$
4.2	$p \vee q$	4.1, $\vee D$
4.3	$(p \vee q) \vee r$	4.2, $\vee D$
4.4	$\perp$	2, 4.3, $\perp D$
5	$\neg q$	4.1 – 4.4, $\neg D$
6	r	4, 5, $\vee E$
7	$(p \vee q) \vee r$	6, $\vee D$
8	$\perp$	2, 7, $\perp D$

Dowody reguł dystrybucji zaprezentujemy w następnym paragrafie, z wykorzystaniem wtórnej reguły konstrukcji dowodu, która ułatwia dowody dla tez i rozumowań zawierających alternatywy.

Dysponowanie regułami obustronnymi pozwala jeszcze na jedno usprawnienie konstrukcji dowodów. Reguły jednostronne można stosować tylko wtedy, gdy podstawienia przesłanek to całe wiersze dowodowe. Reguły obustronne można stosować również wówczas, gdy przesłanka jest podformułą formuły zajmującej wiersz dowodu. Np. (PN) można zastosować do formuły postaci $\neg p \vee \neg \neg q \rightarrow r$ dedukując z niej $\neg p \vee q \rightarrow r$ a tę ostatnią możemy przekształcić przy użyciu (IA) na $(p \rightarrow q) \rightarrow r$. W obu zastosowaniach dokonaliśmy zamiany podformuły danej formuły na jej równoważnik. Formalnie rzecz ujmując za tego typu zabiegami kryje się zastosowanie wtórnej reguły (a raczej metareguły – podobnie jak przy podstawianiu za tezy) zastępowania, zwanej też regułą ekstensjonalności, o postaci:

$$(RZ) \quad \varphi \leftrightarrow \psi, \chi / \chi[\varphi//\psi]$$

$\chi[\varphi//\psi]$ oznacza zastąpienie dowolnego (przynajmniej jednego) wystąpienia φ jako podformuły χ przez ψ . W przeciwieństwie do reguły podstawiania zastąpić możemy dowolną formułę a nie tylko zmienną i nie musimy zastępować w ten sposób wszystkich wystąpień φ . Natomiast ψ , czyli formuła zastępująca, nie może być jakąkolwiek formułą ale musi być formułą

równoważną φ (przynajmniej na gruncie przeprowadzanego dowodu), co oddaje pierwsza przesłanka schematu. W podanym wyżej przykładzie odwołaliśmy się wprawdzie do reguły obustronnej ale to nic nie zmienia gdyż każdej regule obustronnej $\varphi // \psi$ odpowiada teza postaci $\varphi \leftrightarrow \psi$, która może być użyta jako pierwsza przesłanka.

Poniżej przykład dowodu przeprowadzonego z użyciem reguł wtórnych dla konwersu (tj. implikacji odwrotnej) aksjomatu 2 czyli tezy $((p \rightarrow q) \rightarrow (p \rightarrow r)) \rightarrow (p \rightarrow (q \rightarrow r))$:

1	$(p \rightarrow q) \rightarrow (p \rightarrow r)$	$z.$
2	p	$z.$
3	q	$z.$
4	$\neg r$	$zn.$
5	$p \wedge \neg r$	$2, 4, \wedge D$
6	$\neg(p \rightarrow r)$	$5, NI$
7	$\neg(p \rightarrow q)$	$1, 6, MT$
8	$p \wedge \neg q$	$7, NI$
9	$\neg q$	$8, \wedge E$
10	$\perp$	$3, 9, \perp D$

Korzystając z innych reguł wtórnych można zbudować jeszcze krótszy dowód tej tezy. Np. wykorzystując (DI) do założenia z wiersza 3 można zbudować dowód wprost liczący tylko 6 wierszy. Czytelnik może to łatwo sprawdzić, jak również to, że bez użycia reguł wtórnych dowód ten wymagałby wprowadzenia poddowodów.

Chociaż (NI) jest regułą obustronną, w powyższym przykładzie została wykorzystana w sposób standardowy, na przesłance będącej formułą z całego wiersza dowodu. Wykorzystanie reguły zastępowania i konsekwentne użycie równoważników pozwala jeszcze bardziej skracać dowody tez o postaci równoważności. Ponieważ poprzednia teza daje się wzmocnić do postaci równoważności (jako konwers aksjomatu) więc zademonstrujemy tę strategię na tym samym przykładzie:

1	$(p \rightarrow q) \rightarrow (p \rightarrow r)$	<i>z.</i>
2	$\neg(p \rightarrow q) \vee (p \rightarrow r)$	<i>IA</i>
3	$\neg(p \rightarrow q) \vee (\neg p \vee r)$	<i>IA</i>
4	$p \wedge \neg q \vee (\neg p \vee r)$	<i>NI</i>
5	$(p \vee (\neg p \vee r)) \wedge (\neg q \vee (\neg p \vee r))$	<i>PA, DAK</i>
6	$((p \vee \neg p) \vee r) \wedge (\neg q \vee (\neg p \vee r))$	<i>LA</i>
7	$\neg q \vee (\neg p \vee r)$	<i>RED</i>
8	$(\neg q \vee \neg p) \vee r$	<i>LA</i>
9	$(\neg p \vee \neg q) \vee r$	<i>PA</i>
10	$\neg p \vee (\neg q \vee r)$	<i>LA</i>
11	$p \rightarrow (\neg q \vee r)$	<i>IA</i>
12	$p \rightarrow (q \rightarrow r)$	<i>IA</i>

W dowodzie nie zaznaczaliśmy w kolumnie uzasadnień użycia (*RZ*) (w wierszach 3, 4, 6, 9, 12) oraz przekształcanych wierszy, gdyż w każdym przypadku był to poprzedni wiersz. Zastosowanie (*RED*) (i (*RZ*)) do wiersza 6 opiera się na tym, że $(p \vee \neg p) \vee r$ jest tezą. Zauważmy, że dowodząc w ten sposób tez o postaci równoważności nie musimy robić dwóch dowodów, gdyż ostatni wiersz dowodu jest równoważny pierwszemu. Czasem, dla lepszego podkreślenia tego faktu, takie algebraiczne dowody zapisuje się w następujący sposób:

1	$(p \rightarrow q) \rightarrow (p \rightarrow r)$	$\leftrightarrow$	$\neg(p \rightarrow q) \vee (p \rightarrow r)$	<i>IA</i>
2		$\leftrightarrow$	$\neg(p \rightarrow q) \vee (\neg p \vee r)$	<i>IA</i>
3		$\leftrightarrow$	$p \wedge \neg q \vee (\neg p \vee r)$	<i>NI</i>
4		$\leftrightarrow$	$(p \vee (\neg p \vee r)) \wedge (\neg q \vee (\neg p \vee r))$	<i>PA, DAK</i>
5		$\leftrightarrow$	$((p \vee \neg p) \vee r) \wedge (\neg q \vee (\neg p \vee r))$	<i>LA</i>
6		$\leftrightarrow$	$\neg q \vee (\neg p \vee r)$	<i>RED</i>
7		$\leftrightarrow$	$(\neg q \vee \neg p) \vee r$	<i>LA</i>
8		$\leftrightarrow$	$(\neg p \vee \neg q) \vee r$	<i>PA</i>
9		$\leftrightarrow$	$\neg p \vee (\neg q \vee r)$	<i>LA</i>
10		$\leftrightarrow$	$p \rightarrow (\neg q \vee r)$	<i>IA</i>
11		$\leftrightarrow$	$p \rightarrow (q \rightarrow r)$	<i>IA</i>

1.4.3 Dodatkowe reguły konstrukcji dowodu

Oprócz wtórnych reguł inferencji warto zapoznać się z dodatkową regułą konstrukcji dowodu rozgałęzionego (*DR*). Odpowiada ona tradycyjnym re-

gułom wnioskowania z członów alternatywy czy też używania rozmaitych reguł tzw. dylematu konstrukcyjnego prostego, któremu odpowiada aksjomat 8. Stosujemy (*DR*) jeżeli mamy w dowodzie alternatywę, której nie możemy rozbić gdyż nie dysponujemy negacją żadnego z jej członów, która umożliwia zastosowanie ($\vee E$) lub (*MTP*). Jeżeli dodatkowo zauważamy, że z każdego z członów tej alternatywy da się wydedukować pewną formułę χ , która byłaby przydatna w dalszym dowodzeniu, to można to zrobić poprzez utworzenie dwóch, niezależnych od siebie, poddowodów, z których każdy rozpoczyna założenie będące innym członem naszej alternatywy. Każdy z tych poddowodów można zakończyć jeżeli uzyskamy w nim poszukiwaną formułę χ . Wówczas umieszczamy χ w dowodzie nadrzędnym jako wydedukowane z naszej alternatywy z pomocą *DR*. Zademonstrujemy tę technikę przy pomocy dowodu jednego z praw dystrybucji (na razie w jedną stronę):

1	$p \vee (q \wedge r)$	$z.$
2.1	p	$zdr.$
2.2	$p \vee q$	2.1, $\vee D$
2.3	$p \vee r$	2.1, $\vee D$
2.4	$(p \vee q) \wedge (p \vee r)$	2.2, 2.3, $\wedge D$
3.1	$q \wedge r$	$zdr.$
3.2	q	3.1, $\wedge E$
3.3	r	3.1, $\wedge E$
3.4	$p \vee q$	3.2, $\vee D$
3.5	$p \vee r$	3.3, $\vee D$
3.6	$(p \vee q) \wedge (p \vee r)$	3.4, 3.5, $\wedge D$
4	$(p \vee q) \wedge (p \vee r)$	1, 2.1 – 2.4, 3.1 – 3.6, <i>DR</i>

W przykładzie tym alternatywa jest jedyną formułą dostępną w dowodzie. Nawet gdybyśmy zdecydowali się na dowód nie wprost to i tak negacja $(p \vee q) \wedge (p \vee r)$ da nam co najwyżej (po użyciu (*NK*)) jeszcze jedną alternatywę, z którą nie bardzo wiadomo co zrobić. Dwa poddowody (od 2.1 do 2.4 i od 3.1 do 3.6) dają nam proste dedukcje poszukiwanej konkluzji wyprowadzone z obu członów alternatywy z wiersza 1. Są one oznaczone skrótem *zdr.* (założenie dowodu rozgałęzionego). Zauważmy, że oba poddowody są niezależne od siebie, co oddaje ich numeracja; pierwszy to jakby rozbudowany krok 2, a drugi, krok 3, dowodu nadrzędnego. Celowo użyliśmy już w pierwszym poddowodzie kolejnej liczby naturalnej zamiast zostawiać 1 jako numer pierwszej kolumny, gdyż i tak dla drugiego dowodu musimy wprowadzić kolejną liczbę aby odróżnić od siebie oba poddowody. Zamiast

umieszczać jeden poddowód bezpośrednio pod drugim można je umieścić obok siebie. W wierszu 4 mamy przeniesienie konkluzji obu poddowodów do dowodu nadrzędnego. Uzasadnienie odwołuje się do wiersza zawierającego alternatywę oraz do obu poddowodów.

Reguła (*DR*) może być użyta również w przypadku, gdy alternatywa zawiera więcej członów, ale w takim przypadku musimy zbudować tyle poddowodów, ile jest członów alternatywy i każdy z nich musi się zakończyć tą samą formułą. Należy odróżnić to od sytuacji, gdy w ramach tego samego dowodu ponownie stosujemy (*DR*) w obrębie jednego z poddowodów. W taki sposób można dowieść np. konwersu poprzedniej implikacji:

1	$(p \vee q) \wedge (p \vee r)$	$z.$
2	$p \vee q$	1, $\wedge E$
3	$p \vee r$	1, $\wedge E$
4.1	p	$zdr.$
4.2	$p \vee (q \wedge r)$	4.1, $\vee D$
5.1	q	$zdr.$
5.2.1	p	$zdr.$
5.2.2	$p \vee (q \wedge r)$	5.2.1, $\vee D$
5.3.1	r	$zdr.$
5.3.2	$q \wedge r$	5.1, 5.3.1, $\wedge D$
5.3.3	$p \vee (q \wedge r)$	5.3.2, $\vee D$
5.4	$p \vee (q \wedge r)$	3, 5.2.1 – 5.2.2, 5.3.1 – 5.3.3, <i>DR</i>
6	$p \vee (q \wedge r)$	2, 4.1 – 4.2, 5.1 – 5.4, <i>DR</i>

Dowód powyższy choć w pełni poprawny nie jest specjalnie przejrzysty ze względu na dużą liczbę poddowodów. Umieściliśmy go tu tylko, aby zdemonstrować, w jaki sposób można jeden dowód rozgałęziony przeprowadzić w obrębie drugiego. Powyższą tezę znacznie prościej można jednak dowieść nie wprost z wykorzystaniem (*NA*) (zostawiamy to czytelnikowi). Możliwe jest też nieco inne zastosowanie (*DR*), które nie wymaga dwukrotnej eksploatacji tej reguły. Zauważmy, że reguła ta może być też stosowana nawet gdy w dowodzie powyżej nie występuje żadna alternatywa. Oba poddowody zacząć można od dowolnej formuły φ i jej negacji. W tym przypadku implicite odwołujemy się do prawa wyłączonego środka o postaci $\varphi \vee \neg\varphi$. W naszym przypadku w dowodzie wprawdzie jest dostępna alternatywa, a nawet dwie, ale taka strategia pozwala szybciej osiągnąć sukces:

1	$(p \vee q) \wedge (p \vee r)$	$z.$
2	$p \vee q$	$1, \wedge E$
3	$p \vee r$	$1, \wedge E$
4.1	p	$zdr.$
4.2	$p \vee (q \wedge r)$	$4.1, \vee D$
5.1	$\neg p$	$zdr.$
5.2	q	$2, 5.1, \vee E$
5.3	r	$3, 5.1, \vee E$
5.4	$q \wedge r$	$5.2, 5.3, \wedge D$
5.5	$p \vee (q \wedge r)$	$5.4, \vee D$
6	$p \vee (q \wedge r)$	$4.1 - 4.2, 5.1 - 5.5, DR$

W uzasadnieniu wiersza 6, tym razem nie podajemy numeru wiersza z alternatywą – jest to teza $p \vee \neg p$, do której odwołaliśmy się implicite. Efekt jest, w tym wypadku, krótszy i bardziej przejrzysty.

Warto odnotować jeszcze jeden wariant (DR). Reguła ta może być użyta w wersji nie wprost. Otóż jeżeli w przynajmniej jednym poddowodzie wydedukujemy $\perp$, to poszukiwana przez nas formuła χ i tak daje się wydedukować przez ($\perp E$), zatem poddowód też można zamknąć i w dowodzie nadrzędnym dopisać χ .

Korzystanie z dowodu rozgałęzionego jest dogodną strategią, gdy dysponujemy w dowodzie alternatywą, dla której nie mamy negacji żadnego jej członu. Zauważmy, że można tę strategię zastosować też do implikacji, wykorzystując regułę (IA). Z drugiej strony dla dowodzenia alternatywy też można wykorzystać (IA) a następnie strategię opartą na zastosowaniu ($\rightarrow D$).

1.4.4 Dodatkowe sposoby dowodzenia równoważności

Biorąc pod uwagę częste użycie równoważności w dowodzeniu warto zapamiętać kilka reguł wtórnych ułatwiających przeprowadzanie dowodów:

(RMP)	$\varphi \leftrightarrow \psi, \varphi / \psi$ oraz $\varphi \leftrightarrow \psi, \psi / \varphi$
(RMT)	$\varphi \leftrightarrow \psi, \neg \varphi / \neg \psi$ oraz $\varphi \leftrightarrow \psi, \neg \psi / \neg \varphi$
(NR)	$\neg(\varphi \leftrightarrow \psi) // \neg \varphi \leftrightarrow \psi // \varphi \leftrightarrow \neg \psi$

Dwie pierwsze reguły to warianty modus ponens i modus tollens dla równoważności. Dowody uzasadniające ich poprawność są oczywiste, w przeciwieństwie do dowodu dla reguły negacji równoważności, który zamieszczamy

poniżej. Ograniczymy się do wykazania równoważności (obustronnej poprawności) pierwszej pary gdyż druga łatwo daje się dowieść w podobny sposób przy wykorzystaniu symetrii równoważności.

($\rightarrow$)

1	$\neg(p \leftrightarrow q)$	$p.$
2	$\neg((p \rightarrow q) \wedge (q \rightarrow p))$	1, RZ
3	$\neg(p \rightarrow q) \vee \neg(q \rightarrow p)$	2, NK
4.1	$\neg p$	$zd.$
4.2	$\neg p \vee q$	4.1, $\vee D$
4.3	$p \rightarrow q$	4.2, IA
4.4	$\neg(q \rightarrow p)$	3, 4.3, MTP
4.5	$q \wedge \neg p$	4.4, NI
4.6	q	4.5, $\wedge E$
5	$\neg p \rightarrow q$	4.1 – 4.6, $\rightarrow D$
5.1	q	$zd.$
5.1.1	p	$zdn.$
5.1.2	$\neg q \vee p$	5.1.1, $\vee D$
5.1.3	$q \rightarrow p$	5.1.2, IA
5.1.4	$\neg(p \rightarrow q)$	3, 5.1.3, MTP
5.1.5	$p \wedge \neg q$	5.1.4, NI
5.1.6	$\neg q$	5.1.5, $\wedge E$
5.1.7	$\perp$	5.1, 5.1.6, $\perp D$
5.2	$\neg p$	5.1.1 – 5.1.7, $\neg D$
6	$q \rightarrow \neg p$	5.1 – 5.2, $\rightarrow D$
7	$\neg p \leftrightarrow q$	5, 6, $\leftrightarrow D$

($\leftarrow$)

1	$\neg p \leftrightarrow q$	$z.$
2	$p \leftrightarrow q$	$zn.$
2.1	p	$zdr.$
2.2	q	2, 2.1, RMP
2.3	$\neg p$	1, 2.2, RMP
2.4	$\perp$	2.1, 2.3, $\perp D$
3.1	$\neg p$	$zdr.$
3.2	q	1, 3.1, RMP
3.3	p	2, 3.2, RMP
3.4	$\perp$	3.1, 3.3, $\perp D$
4	$\perp$	2.1 – 2.4, 3.1 – 3.4, DR

W pierwszym dowodzie wykorzystaliśmy na początku regułę zastępowania i definicję równoważności zamieniając ją na koniunkcję implikacji w obrębie negacji, a następnie skorzystaliśmy z reguły negacji koniunkcji. Dowód zawiera jeszcze inne zastosowania reguł wtórnych; bez nich byłby znacznie dłuższy. W drugim dowodzie wykorzystaliśmy wielokrotnie wtórną regułę (*RMP*) a także wariant dowodu rozgałęzionego, w którym w obu poddowodach dedukujemy $\perp$.

Własności równoważności oraz pewne wtórne środki dowodzenia omówione wyżej pozwalają na wprowadzenie pewnej strategii dowodzenia twierdzeń, która pozwala na znaczne zaoszczędzenie wysiłku. Otóż w praktyce matematycznej często formułuje się twierdzenia o postaci: $\varphi_1, \varphi_2, \dots, \varphi_n$ są równoważne. W pierwszym odruchu próbowalibyśmy udowodnić dla każdej pary formuł z osobna, że są one równoważne, można jednak skrócić wydatnie zadanie korzystając z pewnych własności przechodniości implikacji. Przykładowo, dla $n = 3$, zamiast dowodzić osobno $\varphi_1 \leftrightarrow \varphi_2, \varphi_2 \leftrightarrow \varphi_3, \varphi_1 \leftrightarrow \varphi_3$, co daje łącznie sześć dowodów implikacji, wystarczy dowieść trzy z nich, np. (1) $\varphi_1 \rightarrow \varphi_2$, (2) $\varphi_2 \rightarrow \varphi_3$, (3) $\varphi_3 \rightarrow \varphi_1$. Pozostałe w prosty sposób można udowodnić przy pomocy tych trzech: korzystając z (1) i (2) udowodnimy $\varphi_1 \rightarrow \varphi_3$, przy pomocy (2) i (3) udowodnimy $\varphi_2 \rightarrow \varphi_1$, wreszcie (3) i (1) daje dowód $\varphi_3 \rightarrow \varphi_2$. W każdym przypadku wprowadzamy obie implikacje do dowodu jako udowodnione wcześniej tezy i stosujemy (*SH*).

Strategię tę zilustrujemy w paragrafie 1.7.4.

1.5 Klasyczny rachunek kwantyfikatorów

Jak wiadomo KRZ jest logiką zbyt słabą by umożliwić analizę dowolnych rozumowań w języku naturalnym lub w językach teorii formalnych. Możemy wprawdzie znaleźć przykłady nawet bardzo wyrafinowanych rozumowań filozoficznych, dla których KRZ jest narzędziem wystarczającym. Jednak, jeżeli poprawność rozumowania zależy nie od występowania odpowiednich spójników lecz od wewnętrznej struktury zdań prostych, to KRZ zawodzi nawet w przypadkach banalnych. Przykładowo, rozumowanie: *Dumbo jest słoniem. Wszystkie słonie mają trąby. Zatem Dumbo ma trąbę.* raczej nie budzi wątpliwości odnośnie swojej poprawności. Jednak próba formalnego wykazania tego faktu przy użyciu KRZ jest niewykonalna. Ponieważ wszystkie zdania są od siebie różne i, z punktu widzenia KRZ, proste, więc formalizacja w KRZ daje nam schemat: $p, q \ / \ r$. Wystarczy przypisać prawdziwość obu przesłankom i fałsz wnioskowi aby pokazać, że jest to schemat niepoprawnego ro-

zumowania na gruncie KRZ. Ale to nie znaczy, że analizowane rozumowanie jest niepoprawne, lecz, że KRZ jest niewystarczającym narzędziem analizy. Potrzebna jest logika, które umożliwi analizę struktury wewnętrznej zdań prostych oraz uwzględnia dodatkowe stałe logiczne – kwantyfikatory: ogólny ($\forall$ – czytamy „dla każdego”) i szczegółowy ($\exists$ – czytamy „dla jakiegoś”).

Taką logiką jest Klasyczny Rachunek Kwantyfikatorów (KRK) zwany też często rachunkiem predykatów lub po prostu logiką klasyczną. KRK określana jest również jako logika pierwszego rzędu ze względu na to, że kwantyfikacja dopuszczona jest jedynie dla zmiennych indywiduowych (lub nazwowych), czyli rzędu pierwszego (denotujących obiekty), a nie rozważa się kwantyfikacji po zmiennych predykatywnych (drugiego rzędu – denotujących zbiory obiektów). Nakłada to pewne ograniczenia na możliwości ekspresji KRK ale, w zasadzie, stanowi ona wystarczającą bazę dla formalizacji wielu znanych teorii matematycznych.

1.5.1 Języki pierwszego rzędu

Język KRK jest naturalnym poszerzeniem języka KRZ, które powstaje przez wprowadzenie większej ilości kategorii syntaktycznych oraz dołączenie nowych logicznych stałych zwanych kwantyfikatorami. W istocie rzeczy powinniśmy mówić nie tyle o języku co o językach pierwszego rzędu, gdyż na gruncie logiki rozważa się tylko ogólny schemat języka, w którym oprócz zmiennych indywiduowych, używamy schematycznych symboli oznaczających nazwy (termy) lub własności i relacje (predykaty). Konkretnie języki uzyskujemy wprowadzając zamiast nieokreślonych symboli (w istocie zmiennych, choć nie poddawanych kwantyfikacji) konkretne stałe pozalogiczne.

Słownik dowolnego języka pierwszego rzędu jest to suma rozłącznych zbiorów:

- $VAR = \{x, y, z, \dots\}$ – przeliczalny zbiór zmiennych indywiduowych
- $CON = \{a, b, c, \dots\}$ – przeliczalny zbiór symboli nazw indywiduowych
- $FUN = \{f, g, h, \dots\}$ – przeliczalny zbiór symboli funkcyjnych
- $PRED = \{A, B, C, \dots\}$ – przeliczalny zbiór symboli predykatów
- $LOG = \{\neg, \wedge, \vee, \rightarrow, \leftrightarrow, \forall, \exists, =\}$

Zbiór VAR jest stały we wszystkich językach, jego elementy mają za zadanie reprezentować nieokreślone obiekty oraz występować w powiązaniu

z kwantyfikatorami. Elementy CON , w przeciwieństwie do zmiennych ze zbioru VAR , denotują konkretne obiekty, są ich nazwami. Elementy FUN są używane do budowy nazw złożonych takich jak np. „ojciec Adama”. Wyrażenie takie formalnie zapiszemy $f(a)$ lub fa , gdzie a symbolizuje imię własne „Adam”, a f funktor nazwotwórczy „ojciec...”. W szczególności, w teoriach matematycznych elementy FUN denotują operacje przeprowadzane na obiektach badanych w danej teorii, np. operację dodawania (dwóch) liczb lub operację sumowania zbiorów. Predykaty to funktory zdaniotwórcze służące do wyrażenia własności obiektów lub relacji zachodzących między nimi. Zarówno w przypadku symboli funkcyjnych jak i predykatów ich istotną cechą jest ilość argumentów, które muszą być do nich dodane aby utworzyć nazwę lub zdanie. Formalnie ujmujemy to tak, że elementy FUN i $PRED$ mają określoną argumentowość (lub arność), którą często zaznacza się w górnym indeksie, np. f^2, A^1 . Zazwyczaj będziemy te indeksy pomijać wtedy, gdy argumentowość symbolu jest jasna z kontekstu. Intuicyjnie argumentowość predykatu oznacza, że reprezentuje on relację zachodzącą między obiektami reprezentowanymi przez jego argumenty (predykaty dwu- i więcej argumentowe) lub własność obiektu (predykaty jednoargumentowe). Zmienne zdaniowe z KRZ traktuje się po prostu jako predykaty zero-argumentowe (podobnie, zbiór CON można zredukować do zbioru FUN traktując nazwy jako funktory zero-argumentowe).

Zauważmy, że symbole w zbiorach CON , FUN , $PRED$ mają charakter parametryczny, gdyż nie oznaczają konkretnych pozalogicznych stałych odpowiednich kategorii. Operowanie nimi jest przydatne w ogólnych rozważaniach nad logiką pierwszego rzędu i dowolnymi językami. W konkretnych językach elementy zbiorów CON , FUN , $PRED$ będą oznaczane specjalnymi symbolami zamiast liter a, b, f, P itp. Wówczas zbiory CON , FUN , $PRED$ są zazwyczaj skończone i mogą być puste – zawierają one stałe pozalogiczne (nazwy, operacje, relacje) specyficzne dla danej teorii.

Oto przykłady słowników konkretnych języków:

Teoria grup: $CON = \{1\}$, $FUN = \{\circ\}$ (dwuargumentowy)

Teoria krat: $CON = \{1, 0\}$, $FUN = \{\times, +\}$ (oba dwuargumentowe)

Teoria mnogości: $PRED = \{\in\}$ (dwuargumentowy)

Arytmetyka: $CON = \{0\}$, $FUN = \{s\}$ (jednoargumentowy)

Oczywiście, w praktyce tak oszczędne języki są niewygodne i wprowadza się definicyjnie kolejne stałe pozalogiczne. Np. w teorii mnogości stałą $\emptyset$, operacje $\cap, \cup, -$, P i predykaty $\subseteq, \subset$ itp., w arytmetyce stałe $1, 2, 3, \dots$

operacje $+$, $\cdot$, predykaty $\leq$, $<$ itd.

Kwantyfikatory zawsze występują w połączeniu z konkretnymi zmiennymi, np. $\forall x, \exists y$ (dla dowolnego x , dla jakiegoś y) oraz z formułami, w których takie zmienne winny się znajdować (choć nie jest to konieczne, dopuszczamy przypadek tzw. pustej kwantyfikacji). Toteż oba kwantyfikatory należą do grupy wyrażeń zwanych operatorami, których zadaniem jest „wiązanie” zmiennych w wyrażeniach.

W metajęzyku używać będziemy tych samych symboli oraz dodatkowych metazmiennych τ_i dla oznaczenia dowolnych termów (tak jak φ, ψ dla dowolnych formuł). Formalnie zbiór termów (wyrażeń nazwowych) *TERM* możemy zdefiniować następująco:

1. $VAR \cup CON \subseteq TERM$;
2. jeżeli $f^n \in FUN$, to $f^n(\tau_1, \dots, \tau_n) \in TERM$, gdzie $\tau_i \in TERM$ dla dowolnego $i \leq n$ (przy czym termy $\tau_1, \dots, \tau_n$ nie muszą być różne);
3. nic więcej nie należy do zbioru *TERM*.

W praktyce będziemy pomijali zarówno indeks górny charakteryzujący argumentowość f jak i nawiasy, jeżeli nie będzie to prowadzić do nieporozumień – ale nie zawsze da się tego uniknąć np. $fxgya$ jest dwuznaczne. Można jednoznacznie je zapisać jako $fxg(ya)$, gdy f^2 i g^2 , lub $fxg(y)a$ gdy f^3 , a g^1 . Zauważmy też, że zazwyczaj dla konkretnych języków stosuje się zapis infiksowy, a nie prefiksowy, tzn. zamiast $+xy$ napiszemy $x + y$. Przykłady termów w konkretnych językach:

$sx, ss0, s(x + sy) + s(ss0 + z)$ – w arytmetyce

$x \cap (y \cup -z), P(x \cap y)$ – w teorii mnogości

Definicja zbioru formuł *FOR* w KRK musi ulec odpowiedniej modyfikacji i zawiera teraz następujące warunki:

1. Jeżeli $P^n \in PRED$, to $P^n(\tau_1, \dots, \tau_n) \in FOR$, gdzie $\tau_i \in TERM$ dla dowolnego $i \leq n$ (przy czym termy $\tau_1, \dots, \tau_n$ nie muszą być różne);
2. jeżeli $\varphi \in FOR$, to $\neg\varphi \in FOR$;
3. jeżeli $\varphi, \psi \in FOR$, to $(\varphi \wedge \psi), (\varphi \vee \psi), (\varphi \rightarrow \psi) \in FOR$;
4. jeżeli $\varphi \in FOR$, to dla dowolnej $x \in VAR$, $\forall x\varphi, \exists x\varphi \in FOR$;

Formuła składająca się z predykatu i jego argumentów to *formuła atomowa*. Natomiast formuła φ występująca w formule $\forall x(\varphi)$ to *zasięg kwantyfikatora* $\forall x$ (analogicznie dla $\exists x$).

W praktyce będziemy pomijali zarówno indeks górny charakteryzujący argumentowość P jak i nawiasy, jeżeli nie będzie to prowadzić do nieporozumień, analogicznie jak w przypadku wyrażeń funkcyjnych (ale np. $Pfxy$ jest dwuznaczne, można je rozumieć jako $Pf(xy)$ lub $Pf(x)y$). Konwencje pomijania nawiasów w formułach złożonych są takie same jak w języku KRZ, w szczególności kwantyfikatory traktujemy analogicznie jak negację, czyli zachowujemy nawias jeżeli zasięg kwantyfikatora jest formułą złożoną, a pomijamy jeżeli jest formułą atomową. Będziemy też stosować konwencję pomijania symboli takich samych kwantyfikatorów, jeżeli występują w bezpośrednim sąsiedztwie, tzn. zamiast $\forall x\forall y\forall z\varphi$ napiszemy $\forall xyz\varphi$.

Zauważmy, że w praktyce, dla konkretnych języków w przypadku termów i predykatów stosuje się zapis infiksowy, tzn. zamiast $\leq xy$ napiszemy $x \leq y$.

Przykłady formuł:

$sx \leq ss0, x \leq sy \rightarrow y = x, \forall xy(x + sy \leq sx + sy)$ – w arytmetyce
 $x \in y, x \subseteq x \cap y, \exists x(x = y \cap \emptyset)$ – w teorii mnogości

Dysponując tak zdefiniowanym językiem KRK możemy z łatwością dokonać formalizacji prostego rozumowania podanego wyżej. Zapisać je możemy następująco: $Sa, \forall x(Sx \rightarrow Tx) / Ta$, gdzie a reprezentuje nazwę „Dumbo”, a S, T – predykaty jednoargumentowe „... jest słoniem”, „... ma trąbę”. Zasad formalizowania zdań języka naturalnego w KRK nie będziemy w tym miejscu objaśniać⁵, gdyż naszym celem jest prezentacja reguł, które umożliwią wykazanie poprawności podanego przykładu. Jednak zanim scharakteryzujemy reguły systemu dedukcji naturalnej dla KRK musimy omówić pewne kwestie terminologiczne.

1.5.2 Zmienne wolne i związane

Okoliczność, że te same zmienne mogą być w tym samym wyrażeniu wolne lub związane przez któryś z kwantyfikatorów zmusza do poczynienia dodatkowych ustaleń i rozróżnienia między zmienną a jej wystąpieniami. Wystąpienie zmiennej x w zasięgu $\forall x$ lub $\exists x$ jest *związane*, w przeciwnym wypadku jest *wolne*. Zmienna jest związana (wolna) w formule, jeżeli ma w niej co najmniej jedno związane (wolne) wystąpienie. Jak widać dana zmien-

⁵Por. np. Indrzejczak [12].

na może być w tej samej formule zarazem wolna i związana. Sporadycznie będziemy też używać zapisu $\varphi(x)$ dla oznaczenia formuły z wolną zmienną x (zawierającą co najmniej jedno wolne wystąpienie x). W podanej wyżej przesłance $\forall x(Sx \rightarrow Tx)$ pierwsze wystąpienie zmiennej x (przy kwantyfikatorze) jest traktowane tylko jako wskaźnik mówiący, która zmienna jest przez ten kwantyfikator związana. Formuła w nawiasie otwartym bezpośrednio za kwantyfikatorem, to zasięg tego kwantyfikatora. W przesłance oba wystąpienia zmiennej x w zasięgu kwantyfikatora są związane. Gdyby zastąpić ją przez formułę $\forall xSx \rightarrow Tx$, to wtedy zasięg kwantyfikatora zostałby ograniczony do Sx i tylko pierwsze wystąpienie x byłoby związane. W formule o postaci $\forall x(Rxy \rightarrow \exists yRxy)$ oba wystąpienia x są związane ale jedynie drugie wystąpienie y jest związane; pierwsze jest wolne.

Niech $\varphi := \forall x(\forall y(Ax \rightarrow Ryz) \rightarrow Bx \wedge \exists xSxy \vee Cx)$, wtedy zbiór wszystkich zmiennych występujących w φ to $\{x, y, z\}$, zbiór zmiennych wolnych to $\{y, z\}$ a związanych $\{x, y\}$. Zauważmy, że y jest zarówno wolne (w drugim wystąpieniu) jak i związane (w pierwszym). Ponadto różne wystąpienia x są związane przez różne kwantyfikatory (w Ax , Bx i Cx przez ogólny, a w Sxy przez szczegółowy).

1.5.3 Podstawianie i zastępowanie

Podobnie jak w KRZ tak i w KRK możemy wprowadzić podstawianie i zastępowanie, tym razem na wyrażeniach nazwowych. Ponadto, w przeciwieństwie do KRZ, operacje te na gruncie KRK są obligatoryjne. Dla poprawnego scharakteryzowania reguł dowolnego rachunku dla KRK niezbędne jest precyzyjne zdefiniowanie operacji podstawiania na zmiennych wolnych w formułach (termach). Podstawiamy dowolne termy ale tylko za zmienne wolne w danej formule czy termie. $\varphi[x/\tau]$ (ewentualnie $\tau_1[x/\tau_2]$) oznacza rezultat *podstawiania* termu τ za x w formule φ (termu τ_2 za x w termie τ_1).

Intuicyjnie dość łatwo jest sformułować warunki właściwego podstawiania, tj. takiego, które z formuły prawdziwej w pewnej interpretacji nie uczyni formuły fałszywej. Są one następujące:

1. Wszystkie wolne wystąpienia x są zastąpione przez τ ;
2. jeżeli τ zawiera zmienne, to pozostają one wolne w φ .

Oba warunki wiążą się z tym, że chcemy, aby operacja podstawiania zachowywała prawdziwość formuł w modelach, tzn. żeby wykluczona była sytuacja, że formuła prawdziwa zamieni się w fałszywą. Obostrzenie pierwsze

wynika nie tylko stąd, że dziedziną operacji podstawiania są tylko zmienne wolne, ale również stąd, że wybiórcze potraktowanie wystąpień podstawianej zmiennej może prowadzić od formuły prawdziwej do fałszywej. Np. gdy w $x = x$, które jest prawdziwe w dowolnej interpretacji podstawimy a za tylko jedno wystąpienie x , to otrzymamy formułę fałszywą w dowolnej interpretacji, w której x i a mają inne desygnaty.

Obostrzenie drugie również jest niezbędne dla uniknięcia podobnych kłopotów. Zilustrujmy to intuicyjnie na prostym przykładzie; $\forall x \exists y, x < y$ jest prawdziwe w arytmetyce liczb naturalnych, ale $\exists y, x < y[x/y] := \exists y, y < y$ jest fałszywe, choć $\exists y, x < y[x/\tau]$ będzie prawdziwe dla każdego τ , w którym y nie występuje.

Aby uniknąć problemów trzeba również ściślej określić kiedy term τ jest poprawnie podstawialny za zmienną x w φ . W tym celu wprowadzimy kolejne pojęcie:

Definicja 1.3 (τ jest termem wolnym za x w φ) *wtw żadne wolne wystąpienie x w φ nie występuje w obrębie podformuły postaci $Qy\psi$ (nie jest w zasięgu Qy), gdzie Q symbolizuje dowolny kwantyfikator a y występuje w τ .*

W dalszym ciągu będziemy używali zapisu $\varphi[x/\tau]$ tylko w przypadku poprawnych podstawień, tj. takich, że τ jest termem wolnym za x w φ .

W ostatnim warunku definicji podstawiania określona jest sytuacja kiedy (poprawne) podstawianie termu nie jest wykonalne, gdyż albo y nie jest wolne w formule, w której dokonujemy podstawienia albo występująca w podstawianym termie zmienna zostanie związana przez jakiś kwantyfikator. Sytuację taką określa się jako kolizję zmiennych. W podanym w paragrafie 1.6.2 przykładzie $(\forall x(\forall y(Ax \rightarrow Ryz) \rightarrow Bx \wedge \exists xSxy \vee Cx))$ ze względu na pierwsze obostrzenie nie jest możliwe $[x/\tau]$ dla żadnego τ , gdyż x nie jest wolne w φ . Ze względu na drugie obostrzenie nie jest możliwe $[z/x], [z/y]$ i $[y/x]$ ani podstawienie żadnego termu, który te zmienne zawiera.

Oprócz podstawiania wyróżnimy zastępowanie. W przeciwieństwie do podstawiania nie zastępujemy tylko zmiennych, ale dowolne termy i nie musimy zastępować wszystkich ich wystąpień. $\varphi[\tau_1//\tau_2]$ (ewentualnie $\tau[\tau_1//\tau_2]$) oznacza rezultat *zastąpienia* termu τ_1 przez τ_2 w φ (ewentualnie w termie τ). Zauważmy, że zastępowanie nie jest operacją w takim sensie jak podstawianie, tzn. nie przypisuje danej formule dokładnie jednej formuły będącej wynikiem podstawienia, ale zbiór formuł. Np. $Pxx[x//y]$ daje trzy możliwe rezultaty: Pyy, Pxy, Pyx .

1.6 Dowodzenie w rachunku kwantyfikatorów

W systemie DN dla KRK zachowane są wszystkie reguły pierwotne i wtórne, które wprowadziliśmy dla KRZ. Dodatkowo, jako pierwotne, wprowadzamy cztery reguły dołączania i eliminacji dla kwantyfikatorów. Dwie z nich są regułami inferencji a dwie regułami konstrukcji dowodu.

1.6.1 Reguły inferencji dla $\forall$ i $\exists$

Zacniemy od omówienia pierwotnych reguł inferencji dla KRK; oto ich schematy:

$$\begin{array}{l} (\forall E) \quad \forall x\varphi / \varphi[x/\tau] \\ (\exists D) \quad \varphi[x/\tau] / \exists x\varphi \end{array}$$

Reguły te są intuicyjnie oczywiste. Sens $(\forall E)$ oddaje twierdzenie, że cokolwiek można orzekać o wszystkich obiektach pewnego rodzaju, to można orzekać o każdym z nich, natomiast $(\exists D)$ odpowiada twierdzeniu, że jeżeli konkretny obiekt spełnia pewien warunek, to istnieje pewien obiekt spełniający ten warunek. Pomimo oczywistej poprawności i formalnej prostoty obie reguły często stosowane są w błędny sposób, toteż poświęcimy trochę uwagi omówieniu tego co jest dopuszczalne, a co nie jest, przy ich zastosowaniu. Problemy wiążą się z faktem, że – w przeciwieństwie do reguł z KRZ – oba schematy dopuszczają rozmaite konkretyzacje ze względu na dokonywane podstawienia termów za zmienną we wniosku $(\forall E)$ i przesłance $(\exists D)$. Oznacza to w szczególności, że obie reguły można wielokrotnie (z różnym skutkiem) stosować do tej samej przesłanki. Pamiętać też musimy, że podstawienia muszą spełniać opisane wyżej warunki poprawności. Przeanalizujemy następującą (ćwiczeniową) dedukcję:

1	$\forall xAx$	$p.$
2	Ax	1, $\forall E$
3	Aa	1, $\forall E$
4	Az	1, $\forall E$
5	Afy	1, $\forall E$
6	$Af(ag(xy))$	1, $\forall E$

W przykładzie mamy rozmaite zastosowania tej samej reguły do przesłanki z wiersza 1. Pierwsze (wiersz 2) jest trywialne w tym sensie, że term podstawiany za x jest identyczny z tą zmienną – po prostu odłączamy kwantyfikator uwalniając x . W pozostałych przypadkach dokonujemy faktycznego

podstawienia, bądź stałej nazwowej (wiersz 3) bądź innej zmiennej (wiersz 4), bądź termu złożonego, nawet o bardzo skomplikowanej strukturze (wiersz 6). W tym ostatnim przypadku zauważmy, że wyjściowa zmienna x może ponownie pojawić się jako argument wyrażenia funkcyjnego.

W formułach złożonych lub zawierających termy złożone ze zmiennymi związanymi przez $\forall$ postępujemy podobnie. Oto przykład:

1	$\forall x(Af(xx) \rightarrow Bx)$	$p.$
2	$Af(xx) \rightarrow Bx$	$1, \forall E$
3	$Af(aa) \rightarrow Ba$	$1, \forall E$
4	$Af(zz) \rightarrow Bz$	$1, \forall E$
5	$Af(g(x)g(x)) \rightarrow Bg(x)$	$1, \forall E$

Ogólnie rzecz biorąc stosowanie ($\forall E$) w przypadku gdy φ (tj. zasięg naszego kwantyfikatora) nie zawiera innych kwantyfikatorów nie sprawia raczej problemów. Musimy pamiętać tylko aby nie łamać ogólnego warunku poprawności dla reguł inferencji, tzn. stosować je tylko do głównej stałej logicznej wyrażenia (całej formuły w wierszu), oraz pamiętać, że podstawienia dokonujemy na wszystkich wystąpieniach zmiennej, która jest związana przez $\forall$. W związku z tym niepoprawne są następujące wnioskowania:

- (a) z $\forall x Ax \vee Bx$ do $Ax \vee Bx$ bądź do $Ay \vee By$
- (b) z $\forall x(Ax \vee Bx)$ do $Ay \vee Bx$ bądź do $Ay \vee Bz$

W (a) $\forall$ nie jest główną stałą wyrażenia. Wprawdzie w pierwszym przypadku sama inferencja jest poprawna ale wymagałaby bardziej złożonego dowodu (np. rozgałęzionego i z zastosowaniem ($\forall E$) do samej formuły $\forall x Ax$); druga jest całkowicie niepoprawna gdyż podstawiliśmy y za x w Bx choć nie było w zasięgu kwantyfikatora. W (b) błędne wnioskowania wynikają stąd, że nie podstawiliśmy nowego termu za wszystkie wystąpienia x lub podstawiliśmy różne termy (y i z) za różne wystąpienia.

Znacznie większe problemy mogą się pojawić gdy φ (tj. zasięg) zawiera inne kwantyfikatory. Bezpieczne zastosowania ($\forall E$) to zastosowania trywialne (tj. ta sama zmienna) lub takie gdzie podstawiamy stałą nazwową lub wyrażenie funkcyjne nie zawierające zmiennych. Oczywiście musimy pamiętać o tym, że podstawiamy wybrany term tylko za te wystąpienia zmiennej, które są związane przez eliminowany kwantyfikator. W związku z tym niepoprawne jest następujące wnioskowanie:

- (c) z $\forall x(Ax \vee \forall x(Bx \rightarrow Cx))$ do $Ay \vee \forall x(By \rightarrow Cy)$ bądź do $Ay \vee \forall y(By \rightarrow Cy)$

gdyż tylko pierwsze wystąpienie x (przy Ax) jest związane przez eliminowany kwantyfikator; pozostałe są związane przez inny kwantyfikator, zatem dopuszczalna inferencja przy podstawieniu y za x pozwala jedynie na wniosek $Ay \vee \forall x(Bx \rightarrow Cx)$.

Dodatkowe problemy mogą się pojawić gdy podstawiany term zawiera zmienne. Niepoprawne jest np. następujące wnioskowanie:

$$(d) \text{ z } \forall x(Ax \vee \forall y(By \rightarrow Cx)) \text{ do } Ay \vee \forall y(By \rightarrow Cy)$$

gdyż y w ostatnim wystąpieniu stało się związane przez inny kwantyfikator. W podanym przypadku powiemy, że zastosowanie $(\forall E)$ jest zabronione zarówno dla y jak i dla dowolnego termu złożonego, który tę zmienną zawiera (np. fy , $g(fa)y$ itp.), gdyż y nie jest wolny za x w tym wyrażeniu.

Podobne uwagi dotyczą zastosowań $(\exists D)$. Przyjrzyjmy się następującej dedukcji:

1	$Axx \wedge Bxa \vee Cfxz$	$p.$
2	$\exists x(Axx \wedge Bxa \vee Cfxz)$	1, $\exists D$
3	$\exists y(Ayy \wedge Bya \vee Cfxz)$	1, $\exists D$
4	$\exists y(Axx \wedge Bxy \vee Cfxz)$	1, $\exists D$
5	$\exists z(Axx \wedge Bxa \vee Cfxz)$	1, $\exists D$
6	$\exists y(Axx \wedge Bxa \vee Cfy)$	1, $\exists D$
7	$\exists y(Axx \wedge Bxa \vee Cy)$	1, $\exists D$
8	$\exists y(Ayx \wedge Bxa \vee Cfxz)$	1, $\exists D$
9	$\exists vy(Ayv \wedge Bxa \vee Cfxz)$	8, $\exists D$

Jak widać w wierszach 2 i 3 można przez zastosowanie $(\exists D)$ skwantyfikować wszystkie wolne wystąpienia danej zmiennej albo w takim kształcie w jakim występuje w przesłance (trywialne podstawienie x za x) albo przez użycie innej zmiennej (y zamiast x). Można skwantyfikować również stałą nazwową (wiersz 4). Wiersze 5 – 7 pokazują, że w przypadku termów złożonych możemy kwantyfikować nie tylko ich proste argumenty (wiersz 5 i 6) ale również całe te wyrażenia (wiersz 7) bez względu na stopień ich złożoności. Wiersz 8 i 9 pokazują, że stosując tę regułę nie jesteśmy zobligowani do kwantyfikowania wszystkich wolnych wystąpień danej zmiennej (czy generalnie termu) w przesłance – można dokonać wyboru wystąpień, które chcemy związać z użyciem kwantyfikatora egzystencjalnego, a nawet użyć różnych kwantyfikatorów. Te ostatnie zastosowania mogą w pierwszym odczuciu wydać się niepoprawne przez skojarzenie z żądaniem podstawiania za wszystkie wystąpienia danej zmiennej przy stosowaniu $(\forall E)$. Zauważmy jednak, że

w poprzednim przypadku kontrola poprawności była zgodna z zastosowaniem reguły tj. od przesłanki do wniosku, natomiast tutaj poprawność sprawdzamy w kierunku odwrotnym do zastosowania reguły, tzn. od wniosku do przesłanki. Łatwo sprawdzić, że formuła z wiersza 1 spełnia warunki poprawnego podstawienia względem zasięgu $\exists$ w każdym z wierszy 2 – 8, a formuła z wiersza 8 względem formuły z wiersza 9. Stosowanie takiej procedury sprawdzania poprawności zastosowania reguły pozwala wykazać, że następujące wnioskowanie jest niepoprawne:

(e) z Axy do $\exists y Ayy$, gdyż Axy nie jest możliwe do uzyskania z Ayy przez podstawienie x za y

Np. wnioskując w taki sposób moglibyśmy wyprowadzić ewidentnie fałszywą formułę $\exists y, y < y$ z formuły $x < y$, prawdziwej przy dowolnej interpretacji, która zmiennej x przyporządkuje liczbę mniejszą od zmiennej y .

Na koniec zaobserwujmy, że formuła $\varphi[x/\tau]$ uzyskana z zasięgu danego kwantyfikatora (czyli z formuły φ) przy stosowaniu dowolnej z omawianych reguł, może mieć dokładnie taką samą strukturę jak φ ale może ją zmieniać w tym sensie, że $\varphi[x/\tau]$ ma mniej różnych wyrażeń nazwowych niż φ . Przykładowo, w podanej niżej dedukcji:

1	$\forall x(Ax \rightarrow Ba)$	$p.$
2	$Ax \rightarrow Ba$	1, $\forall E$
3	$Ab \rightarrow Ba$	1, $\forall E$
4	$Aa \rightarrow Ba$	1, $\forall E$

Formuły z wiersza 2 i 3 mają taką samą strukturę jak φ (modulo użycie innego termu w wierszu 3), natomiast formuła z wiersza 4 ma mniej różnych argumentów nazwowych. Ponieważ $(\exists D)$ jest, w sensie stosowanej operacji podstawiania, odwróceniem $(\forall E)$ więc uzyskujemy w ten sposób pewną wskazówkę o charakterze strategicznym. $(\forall E)$ może służyć do ujednolicania termów w wyrażeniu, natomiast $(\exists D)$ do ich różnicowania. Poniższa dedukcja ilustruje oba zastosowania:

1	$\forall xy(Ax \vee Bya)$	$p.$
2	$\forall y(Aa \vee Bya)$	1, $\forall E$
3	$Aa \vee Baa$	2, $\forall E$
4	$\exists x(Aa \vee Bax)$	3, $\exists D$

$$\begin{array}{ll} 5 & \exists yx(Ay \vee Bax) \quad 4, \exists D \\ 6 & \exists zy(Ay \vee Bzx) \quad 5, \exists D \end{array}$$

1.6.2 Reguły konstrukcji dowodu dla kwantyfikatorów

Zarówno dołączanie $\forall$ jak i eliminacja $\exists$ wykorzystują reguły konstrukcji dowodu o postaci:

$$\begin{array}{ccccccc} \Gamma & z. & \sigma.i & \exists x\varphi & & & \\ \vdots & & & \vdots & & & \\ \sigma.k \quad \varphi & & \sigma.k.1 & \varphi[x/a] \quad ze. & & & \\ \vdots & & & \vdots & & & \\ \sigma.n \quad \forall x\varphi & \sigma.k, \forall D & \sigma.k.n & \psi & & & \\ & & \sigma.k+1 & \psi & \sigma.i, \sigma.k.1 - \sigma.k.n, \exists E & & \end{array}$$

Do obu reguł dołączone są warunki poprawności. W przypadku $(\forall D)$ zmienna x nie może być wolna w żadnym założeniu, które jest aktywne w chwili zastosowania reguły (albo wcale w założeniach nie występuje albo jest związana). Dla poprawności $(\exists E)$ wymagamy, by a było stałą nazwową nie występującą dotąd w dowodzie i nie występującą w formule ψ . Poddowód zainicjowany założeniem egzystencjalnym oznaczonym $ze.$ jest pierwszym wierszem w dowodzie, w którym występuje ta stała.

$(\forall D)$ nie wymaga wprowadzania dodatkowych założeń więc pozornie wygląda jak reguła inferencji i to prostsza niż $(\forall E)$, gdyż nie wymagająca podstawiania – przesłanka ma być identyczna z zasięgiem $\forall$ we wniosku. Jedyne co musimy sprawdzić to czy zmienna, którą chcemy skwantyfikować uniwersalnie nie jest wolna w założeniach. Bez tego ograniczenia moglibyśmy z Ax (jako założenia lub przesłanki) wywnioskować $\forall xAx$, co z pewnością jest zabiegiem niepoprawnym (np. niech x oznacza 2 a A predykat ... „jest liczbą parzystą”). Sens dodanego ograniczenia jest następujący: uogólnić możemy warunek dotyczący arbitralnej zmiennej, czyli takiej, o której nic się wyżej nie zakładało (nie ma jej w ogóle w założeniach dowodu lub jest tam zmienną związaną).⁶

⁶ W rozdziale 2 wprowadzona zostanie inna reguła pozwalająca na dowodzenie formuł postaci $\forall x\varphi$. Jest to reguła indukcji matematycznej, powszechnie stosowana w dowodach w wielu teoriach matematycznych). W przeciwieństwie do reguły $(\forall D)$ nie jest ona jednak regułą uniwersalną lecz stosuje się do zbiorów, których elementy można odpowiednio uporządkować.

Intuicyjny sens $(\exists E)$ jest następujący: jeżeli prawdą jest $\exists x\varphi$, to istnieje pewien, co najmniej jeden, obiekt, który spełnia ten warunek. Nazwijmy go umownie a i podstawmy to imię za x , przy czym nazwa a musi być nowa, gdyż nie możemy zakładać, że jest to jeden z obiektów, o których wyżej coś już ustaliliśmy. Nazwa ta przysługuje rozważanemu obiektowi jedynie czasowo, aby umożliwić nam przeprowadzenie stosownych inferencji. Formuła ψ jest z tego założenia wydedukowana ale nie mówi nic o obiekcie doraźnie nazwanym a , zatem jej prawdziwość nie zależy od przyjętego założenia $\varphi[x/a]$ i możemy dodatkowy dowód zakończyć a ψ potraktować jako wydedukowane z $\exists x\varphi$. Zauważmy, że reguła $(\exists E)$ przypomina regułę (DR) , w której również kształt dodatkowych założeń zależy od wyżej występującej formuły, a poddowód kończy się przepisaniem wydedukowanej formuły do dowodu nadrzędnego. Analogia ta ma dość ścisły charakter gdyż w każdym skończonym zbiorze obiektów formułę egzystencjalną można potraktować jako alternatywę odnośnych formuł zastosowanych do kolejnych obiektów ze zbioru. Np. w zbiorze 2-elementowym zawierającym obiekty o nazwach a i b , zamiast pisać $\exists x\varphi$ napiszemy $\varphi[x/a] \vee \varphi[x/b]$ a efekt uzyskany przez $(\exists E)$ możemy uzyskać przez zastosowanie (DR) .

Rozważmy przykład dowodu $\forall x(Ax \rightarrow Bx), \forall y(By \rightarrow Cy) \vdash \forall z(Az \rightarrow Cz)$

1	$\forall x(Ax \rightarrow Bx)$	$p.$
2	$\forall y(By \rightarrow Cy)$	$p.$
3	$Az \rightarrow Bz$	1, $\forall E$
4	$Bz \rightarrow Cz$	2, $\forall E$
4.1	Az	$zd.$
4.2	Bz	3, 4.1, $\rightarrow E$
4.3	Cz	4, 4.2, $\rightarrow E$
5	$Az \rightarrow Cz$	4.1 – 4.3, $\rightarrow D$
6	$\forall z(Az \rightarrow Cz)$	5, $\forall D$

Zauważmy, że dokonując eliminacji $\forall$ na obu przesłankach nie dokonujemy dowolnych podstawień (choć użycie dowolnych termów byłoby w tym przypadku poprawne) ale w obu przypadkach jest to zmienna z . Jest tak dlatego, że wniosek jest formułą, w której z jest skwantyfikowane przez $\forall$ a reguła $(\forall D)$ nie dopuszcza żadnych podstawień w przesłance – z musi tam już wcześniej się pojawić. Dla poprawnego zastosowania tej reguły musimy tylko sprawdzić, że z nie występuje jako zmienna wolna w aktywnych założeniach lub przesłankach. W przykładzie zmienna z jest wprowadzie wolna

w założeniu z wiersza 4.1 ale w momencie zastosowania ($\forall D$) to założenie nie jest już aktywne, gdyż powyżej zostało zamknięte (wraz z całym poddowodem) przez zastosowanie ($\rightarrow D$). z jest również wolne w wierszach 3 i 4 ale to nie są założenia. Zauważmy, że kształt wniosku wyznacza główną strategię (dedukujemy $Az \rightarrow Cz$ aby zastosować do niej ($\forall D$)) ale również i podstrategię. Ponieważ zasięg $\forall$ we wniosku jest implikacją więc zakładamy dodatkowo poprzednik i dedukujemy następnik aby zastosować ($\rightarrow D$). W powyższym przykładzie można było wprawdzie uniknąć wprowadzania poddowodu i zastosować (SH) do wierszy 3 i 4 uzyskując od razu wiersz 5. Jednak takie skróty nie zawsze są możliwe a zastosowanie obu sprzężonych ze sobą strategii jest powszechne gdyż często mamy do czynienia z dowodzonymi implikacjami poprzedzonymi kwantyfikatorami ogólnymi. W szczególności, strategia ta dobrze pracuje w przypadku tez o postaci implikacji skwantyfikowanych ogólnie. Przykładowo aby dowieść tezy $\forall x(\forall xAx \rightarrow Ax)$ możemy zacząć od poddowodu kończącego się zastosowaniem ($\rightarrow D$) a następnie zastosować ($\forall D$); wygląda to tak:

1.1	$\forall xAx$	$zd.$
1.2	Ax	1.1, $\forall E$
2	$\forall xAx \rightarrow Ax$	1.1 – 1.2, $\rightarrow D$
3	$\forall x(\forall xAx \rightarrow Ax)$	2, $\forall D$

Rozważmy dedukcję $\forall xy(Ax \wedge By), \forall x(Ax \wedge Cx \rightarrow Dx), \exists xCx \vdash \exists yx(Bx \wedge Dy)$

1	$\forall xy(Ax \wedge By)$	$p.$
2	$\forall x(Ax \wedge Cx \rightarrow Dx)$	$p.$
3	$\exists xCx$	$p.$
3.1	Ca	$ze.$
3.2	$\forall y(Aa \wedge By)$	1, $\forall E$
3.3	$Aa \wedge By$	3.2, $\forall E$
3.4	$Aa \wedge Ca \rightarrow Da$	2, $\forall E$
3.5	Aa	3.3, $\wedge E$
3.6	$Aa \wedge Ca$	3.5, 3.1, $\wedge D$
3.7	Da	3.4, 3.6, $\rightarrow E$
3.8	By	3.3, $\wedge E$
3.9	$By \wedge Da$	3.7, 3.8, $\wedge D$
3.10	$\exists x(Bx \wedge Da)$	3.9, $\exists D$
3.11	$\exists yx(Bx \wedge Dy)$	3.10, $\exists D$
4	$\exists yx(Bx \wedge Dy)$	3, 3.1 – 3.11, $\exists E$

W przykładzie mamy nie tylko ilustrację zastosowania ($\exists E$) ale również typowej strategii preferencji przy stosowaniu reguł eliminacji kwantyfikatorów. Jeżeli mamy wybór, to zawsze najpierw stosujemy ($\exists E$), tzn. wprowadzamy założenie egzystencjalne (*ze.*) z nową stałą a dopiero potem stosujemy ($\forall E$) aby podstawić tę stałą tam gdzie jest potrzebna. Gdybyśmy zaczęli dowód od serii zastosowań ($\forall E$) do przesłanek 1 i 2, to byłyby to oczywiście poprawne wnioskowania ale zupełnie bezużyteczne, gdyż po wykonaniu ($\exists E$) (tj. wprowadzeniu założenia egzystencjalnego z nową stałą) i tak trzeba by było ponownie stosować ($\forall E$) do przesłanek z wierszy 1 i 2 aby za x podstawić a . Oczywiście, nie zawsze struktura formuł w dowodzie pozwala najpierw stosować ($\exists E$) a dopiero potem ($\forall E$). Wtedy należy pamiętać, że do danej formuły z kwantyfikatorem ogólnym można zastosować ($\forall E$) ponownie uzyskując potrzebne podstawienie. Poniższy dowód tezy $\neg\forall x\exists y(Ax \wedge \neg Ay)$ pokazuje zastosowanie tej strategii:

1	$\forall x\exists y(Ax \wedge \neg Ay)$	<i>zn.</i>
2	$\exists y(Ax \wedge \neg Ay)$	1, $\forall E$
2.1	$Ax \wedge \neg Aa$	<i>ze.</i>
2.2	$\neg Aa$	2.1, $\wedge E$
2.3	$\exists y(Aa \wedge \neg Ay)$	1, $\forall E$
2.3.1	$Aa \wedge \neg Ab$	<i>ze.</i>
2.3.2	Aa	2.3.1, $\wedge E$
2.3.3	$\perp$	2.2, 2.3.2, $\perp D$
2.4	$\perp$	2.3, 2.3.1 – 2.3.3, $\exists E$
3	$\perp$	2, 2.1 – 2.4, $\exists E$

Niestety czasem zmuszeni jesteśmy wielokrotnie naprzemiennie stosować ($\exists E$) i ($\forall E$) zanim wreszcie uda nam się ujednolicić stałe nazwowe.

Poniżej dwa przykłady błędnych dedukcji, które pokazują do czego może prowadzić nieprzestrzeganie warunków ograniczających dla obu reguł konstrukcji dowodu.

1	$\exists xAx \wedge \exists xBx$	<i>z.</i>
2	$\exists xAx$	1, $\wedge E$
3	$\exists xBx$	1, $\wedge E$
3.1	Aa	<i>ze.</i>
3.1.1	Ba	<i>ze.</i>
3.1.2	$Aa \wedge Ba$	3.1, 3.1.1, $\wedge D$
3.1.3	$\exists x(Ax \wedge Bx)$	3.1.2, $\exists D$

$$\begin{array}{lll} 3.2 & \exists x(Ax \wedge Bx) & 3, 3.1.1 - 3.1.3, \exists E \\ 4 & \exists x(Ax \wedge Bx) & 2, 3.1 - 3.2, \exists E \end{array}$$

Łatwo znaleźć interpretację pokazującą, że powyższe wnioskowanie jest niepoprawne, np. niech A oznacza „... jest parzysta” a B - „... jest nieparzysta” w zbiorze liczb naturalnych, wtedy założenie jest prawdziwe a wniosek fałszywy. Drugie zastosowanie ($\exists E$) jest niepoprawne, gdyż założenie egzystencjalne inicjujące drugi poddowód zawiera stałą a , która już wyżej została użyta. Jeżeli w wierszu 3.1.1 wprowadzimy (poprawnie) jako założenie formułę Bb , to wtedy nie uda nam się udowodnić $\exists x(Ax \wedge Bx)$ (chyba, że w błędny sposób zastosujemy ($\exists D$) do $Aa \wedge Bb$ kwantyfikując różne stałe tym samym kwantyfikatorem).

Łamiąc zasady poprawności reguł możemy też „udowodnić” $\forall x \exists y Rxy \rightarrow \exists y \forall x Rxy$, dla której to formuły kontrprzykładem jest np. interpretacja, która w zbiorze liczb naturalnych przypisze predykatorowi R relację „ x jest mniejsze od y ”. Oto przykład:

$$\begin{array}{lll} 1 & \forall x \exists y Rxy & z. \\ 2 & \exists y Rxy & 1, \forall E \\ 2.1 & Rxa & ze. \\ 2.2 & \forall x Rxa & 2.1, \forall D \\ 2.3 & \exists y \forall x Rxy & 2.2, \exists D \\ 3 & \exists y \forall x Rxy & 2, 2.1 - 2.3, \exists E \end{array}$$

Tym razem błędnie zastosowano ($\forall D$) w wierszu 2.2, gdyż zmienna x jest wolna w założeniu tego poddowodu.

1.6.3 Reguły wtórne

W DN dla KRK poprawne są wszystkie reguły wtórne jakie wprowadziliśmy dla KRZ. W szczególności wtórna reguła dołączania (podstawień) też do dowodu. Zauważmy, że jeżeli wykorzystujemy podstawienia też zawierające wolne zmienne to można je poprawnie skwantyfikować ogólnie. Np.

$$\begin{array}{lll} 1 & Ax \vee \neg Ax & t. \\ 2 & \forall x(Ax \vee \neg Ax) & 1, \forall D \end{array}$$

oraz

1	$Axy \vee \neg Axy$	$t.$
2	$\exists z(Axz \vee \neg Axy)$	1, $\exists D$
3	$\forall y \exists z(Axz \vee \neg Axy)$	2, $\forall D$
4	$\forall xy \exists z(Axz \vee \neg Axy)$	3, $\forall D$

są poprawnymi dowodami skwantyfikowanych wersji prawa wyłączonego środka. W obu przypadkach uniwersalnie kwantyfikowana zmienna nie jest wolna w założeniach czy przesłankach ale w podstawieniu tezy.

Bardzo przydatne w dowodach są reguły De Morgana dla kwantyfikatorów:

$$\neg \forall x \varphi // \exists x \neg \varphi$$

$$\neg \exists x \varphi // \forall x \neg \varphi$$

Oto dowód pierwszej z nich:

($\rightarrow$)

1	$\neg \forall x Ax$	$z.$
2	$\neg \exists x \neg Ax$	$zn.$
2.1	$\neg Ax$	$zdn.$
2.2	$\exists x \neg Ax$	2.1, $\exists D$
2.3	$\perp$	2, 2.2, $\perp D$
3	Ax	2.1 – 2.3, $\neg E$
4	$\forall x Ax$	3, $\forall D$
5	$\perp$	1, 4, $\perp D$

($\leftarrow$)

1	$\exists x \neg Ax$	$z.$
2	$\forall x Ax$	$zn.$
2.1	$\neg Aa$	$ze.$
2.2	Aa	2, $\forall E$
2.3	$\perp$	2.1, 2.2, $\perp D$
3	$\perp$	2, 2.1 – 2.3, $\exists E$

Warto zapamiętać jeszcze następujące reguły wtórne; ich dowody pomijamy gdyż nie nastęrczają szczególnych problemów.

- (DOK) $\forall x(\varphi \wedge \psi) // \forall x\varphi \wedge \forall x\psi$
- (DOA) $\forall x\varphi \vee \forall x\psi / \forall x(\varphi \vee \psi)$
- (DSA) $\exists x(\varphi \vee \psi) // \exists x\varphi \vee \exists x\psi$
- (DSK) $\exists x(\varphi \wedge \psi) / \exists x\varphi \wedge \exists x\psi$
- (PO) $\forall xy\varphi // \forall yx\varphi$
- (PS) $\exists xy\varphi // \exists yx\varphi$
- (PSO) $\exists x\forall y\varphi / \forall y\exists x\varphi$

1.6.4 Reguły dla identyczności

KRK bez dołączenia identyczności jest logiką niewystarczającą dla formalizacji teorii matematycznych. Język KRK poszerzamy poprzez dodanie (jako nowej stałej logicznej) dwuargumentowego predykatu $=$, który zwyczajowo będziemy wpisywać pomiędzy jego argumenty. Najprościej można scharakteryzować identyczność przy pomocy aksjomatu i reguły:

- (ID) $\tau = \tau$
- (RL) $\tau_1 = \tau_2, \varphi / \varphi[\tau_1/\tau_2]$

Aksjomat identyczności stwierdza po prostu, że dowolny obiekt o nazwie τ jest identyczny sam ze sobą, innymi słowy wyraża on własność zwrotności dla relacji identyczności. Reguła Leibniza (RL), zwana też często regułą ekstensjonalności dla identyczności, pozwala zastąpić w dowolnym kontekście jedną nazwę przez drugą, o ile obie denotują ten sam obiekt. Reguła ta jest obciążona dwoma warunkami poprawności zastosowania:

- (a) jeżeli τ_1 zawiera zmienne to muszą być one wolne w tym wystąpieniu τ_1 , które ulega zastąpieniu;
- (b) jeżeli τ_2 zawiera zmienne to muszą one być wolne za τ_1 w tym jego wystąpieniu, które ulega zastąpieniu.

Relacja identyczności jest relacją równoważnościową, co oznacza, że oprócz własności zwrotności przysługują jej też własności symetrii i przechodności. Oto dowód tej ostatniej:

1.1	$x = y \wedge y = z$	$z.$
1.2	$x = y$	1.1, $\wedge E$
1.3	$y = z$	1.1, $\wedge E$
1.4	$x = z$	1.2, 1.3, RL
2	$x = y \wedge y = z \rightarrow x = z$	1.1 – 1.4, $\rightarrow D$
3	$\forall z(x = y \wedge y = z \rightarrow x = z)$	2, $\forall D$
4	$\forall yz(x = y \wedge y = z \rightarrow x = z)$	3, $\forall D$
5	$\forall xyz(x = y \wedge y = z \rightarrow x = z)$	4, $\forall D$

ilustrujący typową strategię dowodzenia ogólnie skwantyfikowanych implikacji. W wierszu 1.4 zastosowaliśmy (RL) w taki sposób, że $\tau_1 = \tau_2$ to $y = z$ a naszą φ jest $x = y$, w której zastąpiliśmy y przez z .

Aby lepiej zapoznać się z działaniem obu nowych reguł udowodnimy, że następujące formuły są równoważne:

- A $\exists xAx \wedge \forall xy(Ax \wedge Ay \rightarrow x = y)$
- B $\exists x(Ax \wedge \forall y(Ay \rightarrow x = y))$
- C $\exists x\forall y(Ay \leftrightarrow x = y)$

co oznacza, że każda z nich może być użyta w celu charakteryzacji istnienia dokładnie jednego obiektu spełniającego warunek opisany predykatem A . W dowodzeniu równoważności tych formuł skorzystamy ze strategii opisanej w paragrafie 1.5.4. tzn. ograniczymy się jedynie do wykazania, że $A \vdash B$, $B \vdash C$ i $C \vdash A$:

$A \vdash B$

1	$\exists xAx \wedge \forall xy(Ax \wedge Ay \rightarrow x = y)$	$z.$
2	$\exists xAx$	1, $\wedge E$
3	$\forall xy(Ax \wedge Ay \rightarrow x = y)$	1, $\wedge E$
3.1	Aa	$ze.$
3.1.1	Ay	$zd.$
3.1.2	$Aa \wedge Ay \rightarrow a = y$	3, $\forall E$
3.1.3	$Aa \wedge Ay$	3.1, 3.1.1, $\wedge D$
3.1.4	$a = y$	3.1.2, 3.1.3, $\rightarrow E$
3.2	$Ay \rightarrow a = y$	3.1.1 – 3.1.4, $\rightarrow D$
3.3	$\forall y(Ay \rightarrow a = y)$	3.2, $\forall D$
3.4	$Aa \wedge \forall y(Ay \rightarrow a = y)$	3.1, 3.3, $\wedge D$
3.5	$\exists x(Ax \wedge \forall y(Ay \rightarrow x = y))$	3.4, $\exists D$
4	$\exists x(Ax \wedge \forall y(Ay \rightarrow x = y))$	2, 3.1 – 3.5, $\exists E$

Zauważmy, że w wierszu 3.1.2 mamy dwukrotne zastosowanie ($\forall E$) do wiersza 3.

$B \vdash C$

1	$\exists x(Ax \wedge \forall y(Ay \rightarrow x = y))$	$z.$
1.1	$Aa \wedge \forall y(Ay \rightarrow a = y)$	$ze.$
1.2	Aa	1.1, $\wedge E$
1.3	$\forall y(Ay \rightarrow a = y)$	1.1, $\wedge E$
1.3.1	$a = y$	$zd.$
1.3.2	Ay	1.2, 1.3.1, RL
1.4	$a = y \rightarrow Ay$	1.3.1 – 1.3.2, $\rightarrow D$
1.5	$Ay \rightarrow a = y$	1.3, $\forall E$
1.6	$Ay \leftrightarrow a = y$	1.4, 1.5, $\leftrightarrow D$
1.7	$\forall y(Ay \leftrightarrow a = y)$	1.6, $\forall D$
1.8	$\exists x \forall y(Ay \leftrightarrow x = y)$	1.7, $\exists D$
2	$\exists x \forall y(Ay \leftrightarrow x = y)$	1, 1.1 – 1.8, $\exists E$

$C \vdash A$

1	$\exists x \forall y(Ay \leftrightarrow x = y)$	$z.$
1.1	$\forall y(Ay \leftrightarrow a = y)$	$ze.$
1.2	$Ay \leftrightarrow a = y$	1.1, $\forall E$
1.3	$Ax \leftrightarrow a = x$	1.1, $\forall E$
1.3.1	$Ax \wedge Ay$	$zd.$
1.3.2	Ax	1.3.1, $\wedge E$
1.3.3	Ay	1.3.1, $\wedge E$
1.3.4	$a = x$	1.3, 1.3.2, RMP
1.3.5	$a = y$	1.2, 1.3.3, RMP
1.3.6	$x = y$	1.3.4, 1.3.5, RL
1.4	$Ax \wedge Ay \rightarrow x = y$	1.3.1 – 1.3.6, $\rightarrow D$
1.5	$\forall xy(Ax \wedge Ay \rightarrow x = y)$	1.4, $\forall D$
1.6	$Aa \leftrightarrow a = a$	1.1, $\forall E$
1.7	$a = a$	ID
1.8	Aa	1.6, 1.7, RMP
1.9	$\exists x Ax$	1.8, $\exists D$
1.10	$\exists x Ax \wedge \forall xy(Ax \wedge Ay \rightarrow x = y)$	1.5, 1.9, $\wedge D$
2	$\exists x Ax \wedge \forall xy(Ax \wedge Ay \rightarrow x = y)$	1, 1.1 – 1.10, $\exists E$

W wierszu 1.5 zastosowaliśmy od razu dwukrotnie ($\forall D$).

Dowody, w których bazuje się na ciągu przekształceń identyczności wykorzystując (LR) często zapisuje się w podobny sposób jak dowody bazujące na przekształceniach równoważności (por. ostatni przykład z paragrafu 1.5.2.). Dla przykładu zademonstrujemy dowód prostego twierdzenia teorii grup, która opisuje własności pewnych elementarnych struktur algebraicznych znajdujących rozliczne zastosowania⁷. Język teorii grup zawiera jedną stałą „1” i jedną operację dwuargumentową „ $\circ$ ”, które charakteryzują następujące trzy aksjomaty⁸:

$$A1 \quad \forall xyz(x \circ (y \circ z) = (x \circ y) \circ z)$$

$$A2 \quad \forall x(x \circ 1 = x)$$

$$A3 \quad \forall y \exists x(y \circ x = 1)$$

Udowodnimy następujące twierdzenie teorii grup: $\forall xy(x \circ y = 1 \rightarrow y \circ x = 1)$:

1.1	$x \circ y$	$=$	1	$zd.$
1.2	$\exists x(y \circ x$	$=$	1)	$A3, \forall E$
1.2.1	$y \circ b$	$=$	1	$ze.$
1.2.2	$y \circ x$	$=$	$(y \circ x) \circ 1$	$A2, \forall E$
1.2.3		$=$	$(y \circ x) \circ (y \circ b)$	1.2.1, 1.2.2, RL
1.2.4		$=$	$((y \circ x) \circ y) \circ b$	$A1, \forall E, 1.2.3, RL$
1.2.5		$=$	$(y \circ (x \circ y)) \circ b$	$A1, \forall E, 1.2.4, RL$
1.2.6		$=$	$(y \circ 1) \circ b$	1.1, RL
1.2.7		$=$	$y \circ b$	$A2, \forall E, 1.2.6, RL$
1.2.8		$=$	1	1.2.1, RL
1.3	$y \circ x$	$=$	1	1.2, 1.2.1 – 1.2.8, $\exists E$
2	$x \circ y = 1$	$\rightarrow$	$y \circ x = 1$	1.1 – 1.3, $\rightarrow D$
3	$\forall y(2)$			2, $\forall D$
4	$\forall xy(2)$			3, $\forall D$

Zauważmy, że w wierszach 1.2.2 – 1.2.8 lewa strona identyczności jest stała, co widać po zastosowaniu ($\exists E$) i przeniesieniu ostatniej identyczności do wiersza 1.3. W 1.2.2 wykorzystaliśmy aksjomat 2, dokonując podstawienia $y \circ x$ za x ; podobnych podstawień w aksjomacie 1 użyliśmy w wierszach

⁷ Jedno z nich znaleźć można w paragrafie 2.7.3.

⁸ Z uwagi na trzeci aksjomat można wprowadzić dodatkową operację jednoargumentową – w taki sposób scharakteryzowane zostaną grupy w paragrafie 2.7.3.

1.2.4 i 1.2.5. W ostatnich dwóch wierszach zastosowaliśmy również skrót notacyjny wpisując zamiast formuły numer jej wiersza w zasięgu dołączanego kwantyfikatora ogólnego. Skrót taki jest możliwy gdyż dołączając $\forall$ niczego nie zmieniamy w formule kwantyfikowanej.

1.7 Uwagi końcowe

Na koniec zwrócimy uwagę na dwie kwestie: proces konstruowania dowodu i ekonomizacja jego zapisu.

1.7.1 Strategie dowodzenia

Dowodzenie nie jest czynnością mechaniczną, wymaga wyobraźni i wprawy. Zaczynamy dowód zawsze kierując się pewnym celem głównym (dedukcja wniosku z przesłanek lub dowód tezy) ale w trakcie dowodzenia często pojawiają się cele częściowe, czyli formuły, które są nam potrzebne na danym etapie konstrukcji dowodu. Koncentrując się na celu częściowym nie powinniśmy tracić „z oczu” celu nadrzędnego gdyż prowadzi to często do mnożenia poddowodów, co komplikuje cały dowód i często utrudnia jego zakończenie. Nie należy zapominać o dowodach nie wprost ale lepiej ich nie nadużywać (gdy widzimy możliwość uzyskania dowodu wprost), gdyż zazwyczaj dowody takie są dłuższe a wielu matematyków uważa je za mało „eleganckie”.

Zazwyczaj konstruując dowód łączymy ze sobą dwa tryby poszukiwania rozwiązania. Pierwszy ma charakter progresywny; analizujemy przesłanki (założenia) i próbujemy z nich wydedukować formuły, które zbliżą nas do rozwiązania zadania. Drugi ma charakter regresywny; analizujemy poszukiwany wniosek i próbujemy „od końca” znaleźć takie formuły, z których da się on wydedukować. Proces konstrukcji dowodu to ciągle przeplatanie tych trybów, kierowane postawionym celem, ostatecznym lub częściowym.

Szkicując dowód musimy zacząć od wypisania założeń oraz – poniżej albo obok – poszukiwanego wniosku. Analizując przesłanki należy cały czas porównywać je z formułą, którą chcemy wydedukować i starać się zastosować do nich te reguły, które mogą zbliżyć nas do wyniku. Analiza formuły, która jest celem dowodu dostarcza wskazówek odnośnie ogólnej strategii dowodzenia lub przynajmniej kroków poprzedzających ostatni wiersz dowodu. Przykładowo, jeżeli mamy dowieść koniunkcji, to zastanawiamy się skąd wydedukujemy oba jej człony – możemy je nawet prowizorycznie wpisać nad wnioskiem. Jeżeli mamy dowieść formuły skwantyfikowanej ogólnie, to wiemy, że przedostatni krok dowodu to zasięg tego kwantyfikatora i musimy

pomyśleć jak go wydedukować. Często jest to implikacja więc nasuwająca się strategia to rozpoczęcie poddowodu (z poprzednikiem tej implikacji jako założeniem dodatkowym) i próba dowiedzenia następnika. Strategię tę ilustrował np. dowód przechodniości identyczności.

Należy pamiętać, że implikacja i alternatywa są wzajemnie zastępowalne za pomocą reguły (*IA*). W związku z tym jeżeli celem (ogólnym lub częściowym) jest alternatywa to można posłużyć się ($\rightarrow D$) do jej udowodnienia (czyli wprowadzić negację jednego z jej członów jako założenie dodatkowe i dowodzić drugiego). Odwrotnie, jeżeli dysponujemy w dowodzie implikacją ale nie mamy poprzednika ani negacji następnika, to możemy potraktować ją jako alternatywę i zastosować dowód rozgałęziony.

W tekście liczne dowody zostały opatrzone komentarzami, które pomagają w zrozumieniu jak zostały skonstruowane. Poniżej dla ilustracji zastosowania omawianych strategii zademonstrujemy bardziej detalicznie proces konstrukcji dowodu dla rozumowania $\exists x(Ax \wedge \forall y(By \rightarrow Rxy)) \vdash \forall y(By \rightarrow \exists x(Ax \wedge Rxy))$. Zaczynamy od wypisania przesłanki a ponieważ jest ona skwantyfikowana egzystencjalnie więc jedyny narzucający się krok to wprowadzenie założenia egzystencjalnego co daje:

$$\begin{array}{ll} 1 & \exists x(Ax \wedge \forall y(By \rightarrow Rxy)) \quad p. \\ 1.1 & Aa \wedge \forall y(By \rightarrow Ray) \quad ze. \end{array}$$

Uzyskane założenie jest koniunkcją a spojrzenie na wniosek pokazuje, że elementy przesłanki występują w nim w innym porządku, zatem rozsądne jest zastosowanie ($\wedge E$) co daje:

$$\begin{array}{lll} 1 & \exists x(Ax \wedge \forall y(By \rightarrow Rxy)) & p. \\ 1.1 & Aa \wedge \forall y(By \rightarrow Ray) & ze. \\ 1.2 & Aa & 1.1, \wedge E \\ 1.3 & \forall y(By \rightarrow Ray) & 1.1, \wedge E \end{array}$$

W tym momencie, zanim np. zastosujemy ($\forall E$) do wiersza 1.3, lepiej jest zastanowić się nad wnioskiem. Jest tak dlatego, że stosowanie ($\forall E$) dopuszcza rozmaite podstawienia za y więc nie należy tego robić na ślepo. Wniosek jest skwantyfikowany ogólnie zatem możemy wyobrazić sobie, że ostatni krok to zastosowanie ($\forall D$), co daje nam taki schemat dowodu:

1	$\exists x(Ax \wedge \forall y(By \rightarrow Rxy))$	$p.$
1.1	$Aa \wedge \forall y(By \rightarrow Ray)$	$ze.$
1.2	Aa	$1.1, \wedge E$
1.3	$\forall y(By \rightarrow Ray)$	$1.1, \wedge E$
	$\vdots$	
$n - 1$	$By \rightarrow \exists x(Ax \wedge Rxy)$	
n	$\forall y(By \rightarrow \exists x(Ax \wedge Rxy))$	$n - 1, \forall D$

Naszym celem jest teraz dedukcja implikacji z wiersza $n - 1$, zatem dobrą strategią jest konstrukcja poddowodu, który zacznie się od założenia dodatkowego By a po uzyskaniu $\exists x(Ax \wedge Rxy)$ pozwoli na zastosowanie $(\rightarrow D)$. Myślimy więc o dowodzie mającym postać:

1	$\exists x(Ax \wedge \forall y(By \rightarrow Rxy))$	$p.$
1.1	$Aa \wedge \forall y(By \rightarrow Ray)$	$ze.$
1.2	Aa	$1.1, \wedge E$
1.3	$\forall y(By \rightarrow Ray)$	$1.1, \wedge E$
	$\vdots$	
$n - 2.1$	By	$zd.$
	$\vdots$	
$n - 2.k$	$\exists x(Ax \wedge Rxy)$	$n - 2.i, \exists D$
$n - 1$	$By \rightarrow \exists x(Ax \wedge Rxy)$	$n - 2.1 - n - 2.k, \rightarrow D$
n	$\forall y(By \rightarrow \exists x(Ax \wedge Rxy))$	$n - 1, \forall D$

Pojawia się tu jednak pewna trudność: planowany poddowód, który doprowadzi do zastosowania $(\rightarrow D)$ chyba nie może być przeprowadzony bezpośrednio nad wierszem $n - 1$, gdyż wcześniej zaczęliśmy poddowód z założeniem egzystencjalnym, który też musi ulec zamknięciu przez zastosowanie $(\exists E)$. Teoretycznie jest możliwe, że mielibyśmy po prostu dwa niezależne poddowody; ten, który zaczęliśmy wcześniej (wiersz 1.1) zamknąłby się wcześniej a poniżej znalazłby się niezależny poddowód prowadzący do dedukcji implikacji z wiersza $n - 1$. Jednak wydaje się to złym tropem biorąc pod uwagę, że w następniku tej implikacji znajduje się formuła Ax a tę można pozyskać z założenia egzystencjalnego Aa (zamiana a na x dokona się przez zastosowanie $(\exists D)$). Zatem oba poddowody nie będą niezależne lecz jeden będzie zawarty w drugim. Rzut oka na implikację z wiersza $n - 1$ pokazuje, że stała nazwowa a w nim nie występuje, zatem można tę formułę prze-

nieść z poddowodu poprzez zastosowanie ($\exists E$). Prowadzi to do następującej modyfikacji naszego szkicu dowodu:

1	$\exists x(Ax \wedge \forall y(By \rightarrow Rxy))$	$p.$
1.1	$Aa \wedge \forall y(By \rightarrow Ray)$	$ze.$
1.2	Aa	1.1, $\wedge E$
1.3	$\forall y(By \rightarrow Ray)$	1.1, $\wedge E$
	$\vdots$	
1.k.1	By	$zd.$
	$\vdots$	
1.k.j	$\exists x(Ax \wedge Rxy)$	1.k.i, $\exists D$
1.k + 1.	$By \rightarrow \exists x(Ax \wedge Rxy)$	1.k.1 – 1.k.i, $\rightarrow D$
2	$By \rightarrow \exists x(Ax \wedge Rxy)$	1, 1.1 – 1.k + 1, $\exists E$
3	$\forall y(By \rightarrow \exists x(Ax \wedge Rxy))$	2, $\forall D$

Przy okazji widzimy, że jedyne sensowne zastosowanie ($\forall E$) do formuły z wiersza 1.3. to takie, które pozostawi zmienną y , gdyż taka występuje w założeniu z wiersza 1.k.1. Zatem otrzymujemy:

1	$\exists x(Ax \wedge \forall y(By \rightarrow Rxy))$	$p.$
1.1	$Aa \wedge \forall y(By \rightarrow Ray)$	$ze.$
1.2	Aa	1.1, $\wedge E$
1.3	$\forall y(By \rightarrow Ray)$	1.1, $\wedge E$
1.4	$By \rightarrow Ray$	1.3, $\forall E$
1.4.1	By	$zd.$
	$\vdots$	
1.4.j	$\exists x(Ax \wedge Rxy)$	1.4.i, $\exists D$
1.5.	$By \rightarrow \exists x(Ax \wedge Rxy)$	1.k.1 – 1.k.i, $\rightarrow D$
2	$By \rightarrow \exists x(Ax \wedge Rxy)$	1, 1.1 – 1.k + 1, $\exists E$
3	$\forall y(By \rightarrow \exists x(Ax \wedge Rxy))$	2, $\forall D$

Pozostaje wypełnić konkretnymi inferencjami ostatni wykropkowany odcinek co jest zadaniem łatwym i daje w rezultacie następujący dowód:

1	$\exists x(Ax \wedge \forall y(By \rightarrow Rxy))$	$p.$
1.1	$Aa \wedge \forall y(By \rightarrow Ray)$	$ze.$
1.2	Aa	1.1, $\wedge E$
1.3	$\forall y(By \rightarrow Ray)$	1.1, $\wedge E$
1.4	$By \rightarrow Ray$	1.3, $\forall E$
1.4.1	By	$zd.$
1.4.2	Ray	1.4, 1.4.1, $\rightarrow E$
1.4.3	$Aa \wedge Ray$	1.2, 1.4.2, $\wedge D$
1.4.4	$\exists x(Ax \wedge Rxy)$	1.4.3, $\exists D$
1.5	$By \rightarrow \exists x(Ax \wedge Rxy)$	1.4.1 – 1.4.4, $\rightarrow D$
2	$By \rightarrow \exists x(Ax \wedge Rxy)$	1, 1.1 – 1.5, $\exists E$
3	$\forall y(By \rightarrow \exists x(Ax \wedge Rxy))$	2, $\forall D$

1.7.2 Dowody nieformalne

Dowody formalne mają tę zaletę, że można sprawdzić ich poprawność w sposób prawie mechaniczny. Ich wadą jest nadmierna długość. Wprawdzie użycie rozmaitych środków w postaci reguł wtórnych pozwala skracać nieraz w znaczący sposób długość dowodów ale wraz ze wzrostem biegłości w dowodzeniu twierdzeń pojawia się potrzeba szukania oszczędności w inny sposób, poprzez pomijanie oczywistych kroków dowodowych. Związane jest z tym pewne ryzyko, gdyż określenie pewnych kroków w dowodzie jako oczywistych ma charakter subiektywny. W podręcznikowych ujęciach rozmaitych działów matematyki przedstawia się zazwyczaj nieformalne i skrótowe szkice dowodów zawierające tylko kluczowe i najtrudniejsze fragmenty wnioskowań. Przykładowo, dowody też będących ogólnymi kwantyfikacjami formuł implikacyjnych zazwyczaj są prezentowane jako dowody odpowiednich implikacji. I tak podany w poprzednim paragrafie dowód przechodniości identyczności podawać można w takim skrócie:

1.1	$x = y$	$z.$
1.2	$y = z$	$z.$
1.3	$x = z$	1.1, 1.2, RL
2	$x = y \wedge y = z \rightarrow x = z$	1.1 – 1.4, $\rightarrow D$
3	$\forall xyz(2)$	2, $\forall D$

gdzie dodatkowo rozpisaliśmy od razu koniunkcję jako dwa założenia i skróciliśmy zapis trzykrotnego zastosowania ($\forall D$). Biorąc pod uwagę schematyczność zastosowań tej ostatniej reguły często po prostu pomija się jej za-

stosowania traktując je jako domyślne. Wtedy nie ma potrzeby stosować też dodatkowej numeracji dla poddowodu; stosujemy po prostu schemat implikacji wstępującej do formuły, która jest (domyślnie) ogólnie skwantyfikowana. W tym konkretnym przypadku dowodzimy $x = y \wedge y = z \rightarrow x = z$, np. w taki sposób:

1	$x = y$	$z.$
2	$y = z$	$z.$
3	$x = z$	1.1, 1.2, RL
4	$x = y \wedge y = z \rightarrow x = z$	1 – 3, $\rightarrow D$

lub nawet tak:

DOWÓD: Załóżmy, że $x = y$ i $y = z$. Przez (RL) mamy $x = z$. $\square$

Ostatni symbol oznacza koniec dowodu. Inne popularne skróty i symbole często używane przy nieformalnym zapisie dla zaznaczenia końca dowodu to $\dashv$ lub skrót „Q.E.D” („quod erat demonstrandum” co oznacza „co było do udowodnienia”) lub po polsku „c.b.d.u”.

W skróconym lub nieformalnym zapisie dowodu często nie zaznacza się też explicite, że odłączenie $\exists$ generuje poddowód a przekształcenia identyczności przedstawia jako jeden ciąg, w którym ostateczny wynik to identyczność łącząca skrajne termny. Np. podany wyżej dowód w teorii grup można zapisać następująco:

DOWÓD: Załóżmy, że $x \circ y = 1$. Ponieważ z aksjomatu 3 mamy $y \circ b = 1$ dla pewnego b , więc, korzystając z aksjomatu 2 i 1 mamy:

$$y \circ x = (y \circ x) \circ 1 = (y \circ x) \circ (y \circ b) = ((y \circ x) \circ y) \circ b = (y \circ (x \circ y)) \circ b = (y \circ 1) \circ b = y \circ b = 1$$

Zatem $x \circ y = 1$ implikuje $y \circ x = 1$ dla dowolnego x, y . $\square$.

Powyższe przykłady ilustrują jedynie niektóre sposoby stosowane w celu efektywnego skracania zapisu dowodów. Warto jednak pamiętać, że nieformalne szkice niestety często pomagają ukryć błędy. Toteż zanim dokonamy nieformalnego zapisu warto rozpisać – przynajmniej dla siebie – dowód w sposób formalny i w miarę pełny aby mieć pewność, że nie popełniliśmy żadnego błędu.

W następnym rozdziale zaprezentujemy dwie ważne teorie: arytmetykę liczb naturalnych w ujęciu Giuseppe Peano i obszerny fragment teorii mnogości w ujęciu Zermelo i Fraenkela. Dowody twierdzeń obu teorii będą zaprezentowane w sposób nieformalny ze względów, które wyżej podaliśmy. Oczy-

wiście zapis każdego dowodu w czysto formalnej postaci jest możliwy; takie dowody może czytelnik znaleźć np. w podręcznikach Borkowskiego [4, 5], a w wersji dostosowanej do komputerowego środowiska MIZAR w podręczniku Nowaka [19]. Zachęcamy do „przepisania” kilku wybranych dowodów podanych dalej jako pełnych formalnych dowodów w systemie dedukcji naturalnej.

Rozdział 2

Dowodzenie w arytmetyce liczb naturalnych i teorii zbiorów

2.1 Arytmetyka elementarna

Arytmetyka elementarna jest najprostszą z teorii liczb naturalnych. Ujmuje ona liczby naturalne bez uwzględnienia działań dodawania i mnożenia.

Arytmetyka elementarna jest teorią pierwszego rzędu zdefiniowaną aksjomatycznie w języku z jedną pierwotną stałą indywidualną 0 reprezentującą liczbę zero oraz jednym pierwotnym 1-argumentowym symbolem funkcyjnym s , reprezentującym tzw. operację następnika (której wartością na danej liczbie naturalnej jest liczba następna po niej tzn. o jeden większa, zwana właśnie następnikiem danej liczby).

2.1.1 Aksjomaty

(AN1) $\forall x(s(x) \neq 0)$,

Liczba zero nie jest następnikiem żadnej liczby naturalnej.

(AN2) $\forall x \forall y(s(x) = s(y) \rightarrow x = y)$,

Dowolne dwie liczby naturalne są identyczne, jeśli ich następniki są identyczne.

(Aksjomat indukcji) $[\psi(0) \wedge \forall x(\psi(x) \rightarrow \psi(s(x)))] \rightarrow \forall x(\psi(x))$,

gdzie $\psi(x)$ jest formułą języka, w której x jest przynajmniej jedną

zmienną wolną, zaś $\psi(0)$ oraz $\psi(s(x))$ ¹ są formułami uzyskanymi z $\psi(x)$ przez zastąpienie każdego wolnego występowania zmiennej x w $\psi(x)$ odpowiednio termami: 0 oraz $s(x)$.

Każda liczba naturalna x spełnia formułę $\psi(x)$, o ile liczba 0 ją spełnia oraz jest tak, że następnik dowolnej liczby spełnia $\psi(x)$ ilekroć ta liczba ją spełnia.

Aksjomat indukcji, podobnie jak niektóre aksjomaty teorii ZF czy też pewnik Cantora, jest schematem dla aksjomatów, z którego uzyskujemy właściwy aksjomat (formułę języka) wstawiając w miejsce $\psi(x)$ konkretną formułę języka teorii.

Praktycznie wyprowadzamy przy użyciu aksjomatu indukcji wszystkie twierdzenia arytmetyki elementarnej o postaci: $\forall x(\psi(x))$, a więc stwierdzające przysługiwanie pewnych własności wszystkim liczbom naturalnym. Mówimy wówczas, że dowody takich twierdzeń są indukcyjne.

2.1.2 Dowody indukcyjne

Aby dowieść indukcyjnie twierdzenia postaci $\forall x(\psi(x))$, wykazujemy prawdziwość poprzednika w aksjomacie indukcji dla tej konkretnej postaci formuły $\psi(x)$. Twierdzenie $\forall x(\psi(x))$ otrzymujemy jako następnik w tym aksjomacie. Ponieważ ów poprzednik jest koniunkcją, wykazujemy więc każdy z jej członów:

(1) $\psi(0)$ (tzw. zerowy krok indukcyjny),

(2) $\forall x(\psi(x) \rightarrow \psi(s(x)))$.

Aby dowieść (2) wybieramy dowolne x oraz zakładamy, że $\psi(x)$ (tzw. założenie indukcyjne). Następnie dążymy do wykazania: $\psi(s(x))$ (tzw. teza indukcyjna).

Wyprowadzimy, przy użyciu aksjomatu indukcji, dwa twierdzenia. Pierwsze mówi, że następnik dowolnej liczby naturalnej jest od niej różny. Drugie, że jakkolwiek liczba naturalna różna od zera jest następnikiem jakiejś liczby.

TWIERDZENIE 1. $\forall x(s(x) \neq x)$.

¹W notacji z 1.5.3 zapis byłby mniej skrótowy: $\psi[x/0]$, $\psi[x/s(x)]$.

DOWÓD: Jako $\psi(x)$ kładziemy w aksjomacie indukcji formułę: $s(x) \neq x$. Zastosowanie tego aksjomatu w tym dowodzie będzie możliwe po wykazaniu prawdziwości dwóch wyrażeń, których koniunkcja jest poprzednikiem aksjomatu indukcji zapisanego dla owej konkretnej formuły $\psi(x)$. Dowodzimy zatem:

(1) $s(0) \neq 0$ (zerowy krok indukcyjny) oraz

(2) $\forall x(s(x) \neq x \rightarrow s(s(x)) \neq s(x))$.

Naturalnie (1) jest bezpośrednim wnioskiem z (AN1). Aby wykazać (2), założymy, że $s(x) \neq x$ (założenie indukcyjne) dla dowolnie wybranego x . Musimy wykazać tezę indukcyjną postaci $s(s(x)) \neq s(x)$. Założymy nie wprost, że $s(s(x)) = s(x)$. Wówczas na mocy (AN2), $s(x) = x$. Sprzeczność z założeniem indukcyjnym. $\square$

TWIERDZENIE 2. $\forall x(x = 0 \vee \exists y(x = s(y)))$.

DOWÓD: Krok zerowy: $0 = 0 \vee \exists y(0 = s(y))$, jest oczywiście spełniony.

Założenie indukcyjne: $x = 0 \vee \exists y(x = s(y))$, dla dowolnie wybranego, ustalonego x . Dowodzimy tezę indukcyjnej postaci: $s(x) = 0 \vee \exists y(s(x) = s(y))$.

Niech założenie indukcyjne będzie prawdziwe w ten sposób, że $x = 0$. Wówczas $s(x) = s(0)$, więc $\exists y(s(x) = s(y))$. Zatem $s(x) = 0 \vee \exists y(s(x) = s(y))$. Niech założenie indukcyjne będzie prawdziwe w ten sposób, że $\exists y(x = s(y))$. Wówczas dla pewnego a , $x = s(a)$ i konsekwentnie $s(x) = s(s(a))$, stąd $\exists y(s(x) = s(y))$, a więc $s(x) = 0 \vee \exists y(s(x) = s(y))$. $\square$

2.2 Arytmetyka liczb naturalnych z dodawaniem

Arytmetyka liczb naturalnych z dodawaniem jest teorią pierwszego rzędu określoną aksjomatycznie na języku arytmetyki elementarnej rozszerzonym o 2-argumentowy symbol funkcyjny $+$, interpretowany w modelu standardowym arytmetyki elementarnej jako operacja „zwykłego” dodawania.

2.2.1 Aksjomaty i podstawowe własności dodawania

Oprócz aksjomatów arytmetyki elementarnej (AN1), (AN2), (Aksjomat indukcji) występują tu dwa aksjomaty, nadające znaczenie symbolowi $+$:

$$(AN3) \quad \forall y(0 + y = y),$$

$$(AN4) \quad \forall x \forall y(s(x) + y = s(x + y)).$$

Aksjomaty te określają wartość operacji dodawania na dowolnych liczbach naturalnych ($1 \stackrel{def}{=} s(0)$, $2 \stackrel{def}{=} s(s(0))$, ..., $n \stackrel{def}{=} s^n(0)$, ...) w sposób następujący: w myśl twierdzenia 2, każda liczba naturalna jest bądź liczbą 0, bądź następnikiem jakiejś liczby naturalnej; (AN3) ustala wartość operacji dodawania dla ciągu $(0, y)$, gdzie y jest dowolną liczbą naturalną, zaś (AN4) określa wartość tej operacji w pozostałych przypadkach, a więc wówczas, gdy pierwszym argumentem operacji dodawania jest jakakolwiek liczba naturalna różna od 0 (liczba $s(x)$), a drugim dowolna liczba naturalna y . Aby w myśl (AN4) obliczyć wartość operacji dodawania na ciągu $(s(x), y)$, trzeba wcześniej znać wartość tej operacji na ciągu (x, y) . Aby ją z kolei obliczyć, bierzemy pod uwagę (AN3), gdy $x = 0$, co zakończy obliczenie, albo znowu (AN4), gdy $x \neq 0$, a więc $x = s(z)$ dla pewnej liczby naturalnej z . Wówczas jednak należy znać wartość operacji dodawania na ciągu (z, y) , a więc znowu korzystamy z (AN3), gdy $z = 0$ (obliczenie skończone), bądź z (AN4), gdy $z \neq 0$ itd. aż do chwili, gdy pierwszy argument operacji będzie liczbą 0, zatem do obliczeń zastosujemy (AN3). Przykładowo:

$$\begin{aligned} 2 + 3 &=_{wg \text{ def}} s(s(0)) + s(s(s(0))) =_{wg (AN4)} s(s(0) + s(s(s(0)))) =_{wg (AN4)} \\ &s(s(0 + s(s(s(0))))) =_{wg (AN3)} s(s(s(s(0)))) =_{wg \text{ def}} 5. \end{aligned}$$

Naturalnie wszystkie twierdzenia arytmetyki elementarnej są twierdzeniami arytmetyki z dodawaniem, choć nie na odwrót. Arytmetyka z dodawaniem jest bogatszą teorią, zaś dowodzenie jej twierdzeń opisujących operację dodawania jest trudniejsze niż dowodzenie twierdzeń w arytmetyce elementarnej. Przykładowo wykażemy, że dodawanie jest operacją przemenną.

TWIERDZENIE 3. $\forall x \forall y(x + y = y + x)$.

DOWÓD: Połóżmy $\phi(x) := \forall y(x + y = y + x)$. Wówczas dowodzone twierdzenie jest następnikiem aksjomatu indukcji, którego poprzednik ma postać:

$$(1) \quad \forall y(0 + y = y + 0) \wedge \forall x[\forall y(x + y = y + x) \rightarrow \forall y(s(x) + y = y + s(x))].$$

Aby wykazać (1) najpierw indukcyjnie dowodzimy

$$(2) \quad \forall y(0 + y = y + 0), \text{ tzn. } \phi(0).$$

Położmy więc $\psi(y) := 0 + y = y + 0$. Wówczas (2) jest następnikiem aksjomatu indukcji, którego poprzednikiem jest

$$(3) \quad 0 + 0 = 0 + 0 \wedge \forall y(0 + y = y + 0 \rightarrow 0 + s(y) = s(y) + 0).$$

Naturalnie pierwszy człon koniunkcji (3) (formuła $\psi(0)$) jest spełniony. Aby wykazać prawdziwość drugiego członu, tzn. formuły: $\forall y(\psi(y) \rightarrow \psi(s(y)))$, weźmy dowolną liczbę y i załóżmy, że $0 + y = y + 0$. Na mocy (AN3) mamy wówczas: $y = y + 0$, zatem $s(y) = s(y + 0)$. Stąd na mocy (AN4), $s(y) = s(y) + 0$. Lecz z drugiej strony, na podstawie (AN3), $0 + s(y) = s(y)$, zatem $0 + s(y) = s(y) + 0$. W ten sposób wykazaliśmy (3), a więc w konsekwencji również (2), czyli wykonaliśmy zerowy krok indukcyjny w dowodzie twierdzenia.

Aby dowieść drugiego członu koniunkcji (1), tzn. formuły: $\forall x(\phi(x) \rightarrow \phi(s(x)))$, rozważmy dowolną ustaloną liczbę x oraz załóżmy

$$(4) \quad \forall y(x + y = y + x).$$

Aby wykazać

$$(5) \quad \forall y(s(x) + y = y + s(x)),$$

połóżmy $\chi(y) := s(x) + y = y + s(x)$. Dowód (5) będzie więc indukcyjny. Musimy wykazać

$$(6) \quad s(x) + 0 = 0 + s(x) \text{ (tzn. } \chi(0)) \text{ oraz}$$

$$(7) \quad \forall y(s(x) + y = y + s(x) \rightarrow s(x) + s(y) = s(y) + s(x)) \text{ (tzn. } \forall y(\chi(y) \rightarrow \chi(s(y))))).$$

Formuła (6) jest oczywistym wnioskiem z (2). Dowodzimy więc formuły (7). Weźmy dowolną liczbę y i załóżmy, że

$$(8) \quad s(x) + y = y + s(x).$$

Wówczas $s(x) + s(y) =_{na \text{ mocy (AN4)}} s(x + s(y)) =_{na \text{ mocy (4)}} s(s(y) + x) =_{na \text{ mocy (AN4)}} s(s(y + x)) =_{na \text{ mocy (4)}} s(s(x + y)) =_{na \text{ mocy (AN4)}} s(s(x) + y) =_{na \text{ mocy (8)}} s(y + s(x)) =_{na \text{ mocy (AN4)}} s(y) + s(x)$.

W ten sposób wykazano prawdziwość formuły (5), wobec czego również formuły (1), zatem na mocy aksjomatu indukcji dowód twierdzenia jest zakończony. $\square$

Wykazujemy, nie stosując jednak wielokrotnie aksjomatu indukcji, lecz standardowy sposób dowodzenia zdania ogólnego, że działanie dodawania jest łączne.

TWIERDZENIE 4. $\forall z \forall y \forall x (x + (y + z) = (x + y) + z)$.

DOWÓD: Weźmy pod uwagę dowolne dwie liczby naturalne y, z . Wykazujemy indukcyjnie:

$$\forall x (x + (y + z) = (x + y) + z).$$

Stosując dwukrotnie aksjomat (AN3) czynimy krok zerowy indukcji: $0 + (y + z) = y + z = (0 + y) + z$. Przyjmujemy założenie indukcyjne:

$$x + (y + z) = (x + y) + z$$

i wykazujemy tezę indukcyjną:

$$s(x) + (y + z) = (s(x) + y) + z,$$

stosując parokrotnie aksjomat (AN4): $s(x) + (y + z) = s(x + (y + z)) = s((x + y) + z) = s(x + y) + z = (s(x) + y) + z$. $\square$

Nie w każdym dowodzie wykorzystujemy aksjomat indukcji. Często wystarczy wykorzystać inne twierdzenia (których dowody są indukcyjne). W sformułowaniu kolejnego twierdzenia, zgodnie z popularną konwencją, pomijamy duże kwantyfikatory (szczególnie w twierdzeniach nie dowodzonych bezpośrednio indukcyjnie).

Twierdzenie 5. $x + y = 0 \leftrightarrow (x = 0 \wedge y = 0)$.

Dowód: ($\rightarrow$): Niech $x + y = 0$ oraz nie wprost, $y \neq 0$. Wówczas na mocy twierdzenia 2, $y = s(z)$, dla pewnego z . Zatem według twierdzenia 3 oraz (AN3), $x + y = x + s(z) = s(z) + x = s(z + x)$, a więc z założenia, $s(z + x) = 0$, co jest sprzeczne z (AN1). Analogicznie, gdy $x \neq 0$.

($\leftarrow$): Na podstawie (AN3). $\square$

Obecnie poświęćmy uwagę twierdzeniom mającym zastosowanie w rozwiązywaniu równań.

Twierdzenie 6. $\forall y \forall x (x + y = x \rightarrow y = 0)$.

Dowód: Rozważmy dowolną liczbę naturalną y . Wykazujemy indukcyjnie, że

$$\forall x (x + y = x \rightarrow y = 0).$$

Warunek $0 + y = 0 \rightarrow y = 0$, zachodzi na mocy (AN3). Załóżmy, że dla jakiejś liczby naturalnej x , $x + y = x \rightarrow y = 0$. Aby wykazać, że $s(x) + y = s(x) \rightarrow y = 0$, przypuśćmy, że $s(x) + y = s(x)$. Wtedy według (AN4): $s(x + y) = sx$, zatem z (AN2): $x + y = x$. Wobec założenia indukcyjnego, $y = 0$. $\square$

Twierdzenie 7. $\forall x \forall y \forall z (z + x = z + y \rightarrow x = y)$.

Dowód: Weźmy dowolne liczby naturalne x, y . Wykazujemy indukcyjnie, że $\forall z (z + x = z + y \rightarrow x = y)$. Warunek $0 + x = 0 + y \rightarrow x = y$, zachodzi na mocy (AN3). Załóżmy, że dla jakiejś liczby naturalnej z zachodzi: $z + x =$

$z + y \rightarrow x = y$. Aby wykazać, że $s(z) + x = s(z) + y \rightarrow x = y$, przypuśćmy, że $s(z) + x = s(z) + y$. Wówczas według (AN4), (AN2) mamy $z + x = z + y$, co na mocy założenia indukcyjnego implikuje: $x = y$. $\square$

2.2.2 Relacja porządku

Wzbogacamy język arytmetyki o dwuargumentowy predykat $\leq$.

DEFINICJA. $\forall x \forall y (x \leq y \leftrightarrow \exists z (z + x = y))$.

Najpierw dowodzimy, że relacja $\leq$ jest częściowym porządkiem, tzn. jest zwrotna (twierdzenie 8), antysymetryczna (twierdzenie 9) i przechodnia (twierdzenie 10).

TWIERDZENIE 8. $\forall x (x \leq x)$.

DOWÓD: Oczywisty na mocy (AN3). $\square$

TWIERDZENIE 9. $\forall x \forall y (x \leq y \wedge y \leq x \rightarrow x = y)$.

DOWÓD: Załóżmy, że $x \leq y$ oraz $y \leq x$. Zatem z definicji relacji $\leq$, dla pewnych liczb naturalnych a, b mamy $a + x = y$ oraz $b + y = x$. Wobec tego $b + (a + x) = x$, co na mocy twierdzenia 4 oraz twierdzenia 3 daje: $x + (b + a) = x$ i w konsekwencji $b + a = 0$, zgodnie z twierdzeniem 6. Dlatego, na mocy twierdzenia 5, $b = 0$ oraz $a = 0$. Ostatecznie według (AN3), $x = y$. $\square$

TWIERDZENIE 10. $\forall x \forall y \forall z (x \leq y \wedge y \leq z \rightarrow x \leq z)$.

DOWÓD: Załóżmy, że $x \leq y$ oraz $y \leq z$. Wtedy dla pewnych liczb naturalnych a, b mamy $a + x = y$ oraz $b + y = z$. Stąd $b + (a + x) = z$. Stosując twierdzenie 4 i definicję relacji $\leq$ otrzymujemy $x \leq z$. $\square$

Liczba 0 jest najmniejszą liczbą naturalną:

TWIERDZENIE 11. $\forall x (0 \leq x)$.

DOWÓD: Oczywisty na mocy (AN3) i twierdzenia 3. $\square$

TWIERDZENIE 12. $\forall x (x \leq s(x))$.

DOWÓD: Ponieważ na mocy (AN4), (AN3): $s(0) + x = s(0 + x) = s(x)$, więc $x \leq s(x)$. $\square$

TWIERDZENIE 13. $\forall x(\neg s(x) \leq x)$.

DOWÓD: Załóżmy nie wprost, że $s(a) \leq a$ dla pewnej liczby naturalnej a . Wówczas na mocy twierdzenia 12 oraz twierdzenia 9, $s(a) = a$, co jest sprzeczne z twierdzeniem 1. $\square$

Twierdzenia 1, 12 oraz 14 łącznie stwierdzają, że następnik danej liczby naturalnej jest najmniejszą ze wszystkich liczb większych od tej danej liczby.

TWIERDZENIE 14. $x \leq y \wedge x \neq y \rightarrow s(x) \leq y$.

DOWÓD: Załóżmy, że $x \leq y$ oraz $x \neq y$. Wówczas dla pewnej liczby z , $z + x = y$ oraz, dzięki (AN3), $z \neq 0$. Stąd, na mocy twierdzenia 2, $z = s(a)$ dla pewnej liczby a . Zatem $s(a) + x = y$, a więc z (AN4), $s(a + x) = y$. Stosując twierdzenie 3 oraz aksjomat (AN4) otrzymujemy: $a + s(x) = y$, co daje $s(x) \leq y$. $\square$

Częściowy porządek $\leq$ jest liniowy, tzn. relacja $\leq$ jest spójna.

TWIERDZENIE 15. $\forall x \forall y (x \leq y \vee y \leq x)$.

DOWÓD: (indukcyjny) Krok zerowy: $\forall y (0 \leq y \vee y \leq 0)$, jest spełniony na mocy twierdzenia 11. Przyjmujemy założenie indukcyjne $\forall y (x \leq y \vee y \leq x)$ i dowodzimy tezy indukcyjnej: $\forall y (s(x) \leq y \vee y \leq s(x))$. Rozważmy zatem dowolną liczbę naturalną y i na podstawie założenia indukcyjnego przypuśćmy, że pierwszy z dysjunktów: $x \leq y$ zachodzi. Mamy oczywiście alternatywę: $x = y$ lub $x \neq y$. Gdy $x = y$ to $s(x) = s(y)$, dlatego, $y \leq s(x)$ na mocy twierdzenia 12. Gdy zaś $x \neq y$, to według twierdzenia 14, $s(x) \leq y$. Przypuśćmy teraz, że drugi z dysjunktów z założenia indukcyjnego: $y \leq x$, zachodzi. Wówczas, na podstawie twierdzenia 12 otrzymujemy $y \leq s(x)$. $\square$

TWIERDZENIE 16. $x \leq y \leftrightarrow s(x) \leq s(y)$.

DOWÓD: ($\rightarrow$): Załóżmy, że $x \leq y$. Zachodzi $x = y$ lub $x \neq y$. Gdy $x = y$, to $s(x) = s(y)$, zatem $s(x) \leq s(y)$, na mocy twierdzenia 8. Gdy zaś $x \neq y$, to $s(x) \leq y$, na mocy twierdzenia 14, zatem $s(x) \leq s(y)$ na podstawie twierdzeń 12, 10.

($\leftarrow$): Załóżmy, że $s(x) \leq s(y)$. Gdyby zachodziło $y \leq x \wedge y \neq x$, to wówczas według twierdzenia 14 byłoby $s(y) \leq x$, zatem wobec założenia i twierdzenia 10, $s(x) \leq x$, co jest niemożliwe na mocy twierdzenia 13. Wnioskujemy więc, iż $\neg y \leq x \vee y = x$, co według twierdzeń 15, 8 implikuje $x \leq y$. $\square$

Nie istnieje największa liczba naturalna:

Twierdzenie 17. $\neg \exists x \forall y (y \leq x)$.

Dowód: Załóżmy nie wprost, że dla pewnej liczby naturalnej a , $\forall y (y \leq a)$. Wówczas $s(a) \leq a$, co jest niemożliwe wobec twierdzenia 13. $\square$

Aksjomat indukcji zazwyczaj wykorzystywany jest dla dowodów indukcyjnych w bardziej popularnej postaci, możliwej do napisania dopiero w języku arytmetyki z dodawaniem:

$$(\psi(0) \wedge \forall x (\psi(x) \rightarrow \psi(x+1))) \rightarrow \forall x (\psi(x)).$$

Traktujemy powyższe wyrażenie jako twierdzenie, będące oczywistym wnioskiem z formuły (Aksjomat indukcji) oraz twierdzenia:

Twierdzenie 18. $\forall x (s(x) = x + s(0))$.

Dowód: Rozważmy dowolną liczbę x . Wówczas na mocy twierdzenia 3 oraz (AN4), (AN3) mamy: $x + s(0) = s(0) + x = s(0 + x) = s(x)$. $\square$

Rozważmy uogólnienie aksjomatu indukcji w następującej postaci.

Twierdzenie 19. *Dla każdej liczby naturalnej y oraz dowolnej formuły $\phi(x)$ języka arytmetyki, z co najmniej jedną zmienną wolną x ,*

$$[\phi(y) \wedge \forall x (y \leq x \rightarrow (\phi(x) \rightarrow \phi(x+1)))] \rightarrow \forall x (y \leq x \rightarrow \phi(x)).$$

Dowód: Załóżmy

- (1) $\phi(y)$ oraz
- (2) $\forall x (y \leq x \rightarrow (\phi(x) \rightarrow \phi(x+1)))$.

Aby wykazać następnik dowodzonej implikacji korzystamy z aksjomatu indukcji dla formuły $\psi(x)$ postaci $y \leq x \rightarrow \phi(x)$. Zatem najpierw dowodzimy:

$$y \leq 0 \rightarrow \phi(0).$$

Przypuśćmy więc, że $y \leq 0$. Wtedy $y = 0$, na mocy twierdzeń 11 i 9. Stąd i z (1), $\phi(0)$. Teraz wykazujemy:

$$\forall x [(y \leq x \rightarrow \phi(x)) \rightarrow (y \leq x + 1 \rightarrow \phi(x + 1))].$$

Weźmy pod uwagę dowolną liczbę naturalną x i założmy, że

$$(3) \quad y \leq x \rightarrow \phi(x) \text{ oraz}$$

$$(4) \quad y \leq x + 1.$$

Oczywiście zachodzi $y = x + 1$ lub $y \neq x + 1$. Gdy $y = x + 1$, to $\phi(x + 1)$, na mocy (1). Jeśli zaś $y \neq x + 1$, to według (4) oraz twierdzeń 14, 18 mamy $s(y) \leq s(x)$, co implikuje $y \leq x$, na mocy twierdzenia 16. Stąd i z (3) otrzymujemy $\phi(x)$ i ostatecznie z (2), $\phi(x + 1)$. $\square$

2.3 Arytmetyka z dodawaniem i mnożeniem

Arytmetyka liczb naturalnych z dodawaniem i mnożeniem jest teorią pierwszego rzędu określoną aksjomatycznie na języku arytmetyki z dodawaniem rozszerzonym o 2-argumentowy symbol funkcyjny $\cdot$, interpretowany jako operacja mnożenia.

2.3.1 Aksjomaty i podstawowe własności mnożenia

Aksjomatyka tej teorii złożona jest z aksjomatów arytmetyki z dodawaniem oraz dwóch kolejnych aksjomatów opisujących operację mnożenia, analogicznie jak aksjomaty (AN3), (AN4) opisują dodawanie:

$$(AN5) \quad \forall y (0 \cdot y = 0),$$

$$(AN6) \quad \forall x \forall y (s(x) \cdot y = (x \cdot y) + y).$$

Wykazujemy, że działanie mnożenia jest przemienne.

Twierdzenie 20. $\forall x \forall y (x \cdot y = y \cdot x)$.

Dowód: Indukcyjny. Aby wykonać krok zerowy wykażemy najpierw indukcyjnie, że

$$(1) \quad \forall y (y \cdot 0 = 0).$$

Oczywiście $0 \cdot 0 = 0$ na mocy (AN5). Przyjmujemy założenie indukcyjne: $y \cdot 0 = 0$ i dowodzimy tezę: $s(y) \cdot 0 = 0$. Mamy więc dzięki (AN6), założeniu indukcyjnemu i (AN3): $s(y) \cdot 0 = (y \cdot 0) + 0 = 0 + 0 = 0$.

Teraz zaczynamy dowód indukcyjny twierdzenia 20. Weźmy dowolną liczbę y . Wówczas z (AN5) i (1) otrzymujemy: $0 \cdot y = 0 = y \cdot 0$. W ten sposób dowiedliśmy, iż $\forall y (0 \cdot y = y \cdot 0)$.

Pozostaje wykazać dla wybranej dowolnie liczby naturalnej x tezę indukcyjną

$$\forall y (s(x) \cdot y = y \cdot s(x)),$$

przy założeniu indukcyjnym

$$(2) \quad \forall y (x \cdot y = y \cdot x).$$

W tym celu wykażemy indukcyjnie najpierw, dla owej ustalonej liczby x , że

$$(3) \quad \forall y ((y \cdot x) + y = y \cdot s(x)).$$

Krok zerowy: $(0 \cdot x) + 0 = 0 + 0 = 0 = 0 \cdot s(x)$, zachodzi na mocy (AN5) i (AN3). Załóżmy, że

$$(y \cdot x) + y = y \cdot s(x) \text{ i wykażmy, że}$$

$$(s(y) \cdot x) + s(y) = s(y) \cdot s(x).$$

Mamy zatem, najpierw z (AN6), następnie z założenia indukcyjnego, twierdzeń 3, 4, (AN4) i w końcu (AN6): $s(y) \cdot s(x) = (y \cdot s(x)) + s(x) = (y \cdot x) + y + s(x) = (y \cdot x) + s(x) + y = (y \cdot x) + s(x + y) = (y \cdot x) + s(y + x) = (y \cdot x) + s(y) + x = ((y \cdot x) + x) + s(y) = (s(y) \cdot x) + s(y)$.

Zakładamy (2) i bierzemy pod uwagę dowolną liczbę y . Wówczas dzięki (AN6), (2) (3) mamy $s(x) \cdot y = (x \cdot y) + y = (y \cdot x) + y = y \cdot s(x)$. $\square$

Łączności mnożenia będziemy dowodzić tak jak w Grzegorzczuk[10], zatem posilkując się rozdzielnością mnożenia względem dodawania.

TWIERDZENIE 21. $\forall x \forall y \forall z (x \cdot (y + z) = (x \cdot y) + (x \cdot z))$.

DOWÓD: Indukcyjny. Formuła $\forall y \forall z (0 \cdot (y + z) = (0 \cdot y) + (0 \cdot z))$ jest dowodliwa dzięki (AN5) i (AN3). Wykazujemy, iż $\forall y \forall z (s(x) \cdot (y + z) = (s(x) \cdot y) + (s(x) \cdot z))$ przy założeniu: $\forall y \forall z (x \cdot (y + z) = (x \cdot y) + (x \cdot z))$. Weźmy więc pod uwagę dowolne y, z . Wówczas $s(x) \cdot (y + z) = x \cdot (y + z) + (y + z) = (x \cdot y) + (x \cdot z) + (y + z) = ((x \cdot y) + y) + ((x \cdot z) + z) = (s(x) \cdot y) + (s(x) \cdot z)$, dzięki (AN6), założeniu indukcyjnemu, twierdzeniom 3 i 4 oraz ponownemu zastosowaniu (AN6). $\square$

TWIERDZENIE 22. $\forall x \forall y \forall z ((y + z) \cdot x = (y \cdot x) + (z \cdot x))$.

DOWÓD: Weźmy dowolne x, y, z . Z twierdzeń 20, 21 mamy $(y + z) \cdot x = x \cdot (y + z) = (x \cdot y) + (x \cdot z) = (y \cdot x) + (z \cdot x)$. $\square$

TWIERDZENIE 23. $\forall x \forall y \forall z (x \cdot (y \cdot z) = (x \cdot y) \cdot z)$.

DOWÓD: Indukcyjny. Krok zerowy: $\forall y \forall z (0 \cdot (y \cdot z) = (0 \cdot y) \cdot z)$ jest spełniony na mocy (AN5). Wykazujemy tezę indukcyjną $\forall y \forall z (s(x) \cdot (y \cdot z) = (s(x) \cdot y) \cdot z)$ przy założeniu $\forall y \forall z (x \cdot (y \cdot z) = (x \cdot y) \cdot z)$. Zatem, $s(x) \cdot (y \cdot z) = (x \cdot (y \cdot z)) + (y \cdot z) = ((x \cdot y) \cdot z) + (y \cdot z) = ((x \cdot y) + y) \cdot z = (s(x) \cdot y) \cdot z$, dzięki (AN6), założeniu indukcyjnemu, twierdzeniu 22 oraz ponownemu zastosowaniu aksjomatu (AN6). $\square$

2.4 Teoria mnogości

Kolejnym naszym celem jest prezentacja najbardziej popularnej teorii mnogości, tzw. teorii ZF , czyli teorii Zermelo-Fraenkla. Jej język wyposażony jest w jedyny pierwotny predykat 2-argumentowy $\in$ (poza naturalnie predykatem identyczności: $=$, traktowanym jako stała logiczna) i nie zawiera żadnych pierwotnych stałych indywidualnych ani pierwotnych symboli funkcyjnych. Predykat „ $\in$ ” czytamy: „należy do” lub „jest elementem” (używamy zapisu: $x \notin y$ jako równoznacznego z negacją formuły atomowej: $x \in y$).

2.4.1 Naiwna teoria zbiorów

Teoria ZF może być więc traktowana jako jedna z teorii ustalających i precyzujących znaczenie terminu: „jest elementem mnogości (zbioru)” czy też „jest częścią mnogości”. Historycznie pierwszą z takich teorii formalnych jest tzw. naiwna teoria mnogości, sformułowana nieaksjomatycznie w drugiej połowie XIX stulecia przez Georga Cantora. Miała ona fundamentalny mankament – była sprzeczna. Sformułowanie teorii niesprzecznej, lecz zachowującej podstawowe intuicje znaczeniowe Cantora było celem badań na początku XX w. Jednym z owoców tych badań jest właśnie teoria ZF . Aby przybliżyć owe intuicje Cantora rozważmy przez chwilę tzw. aksjomatyczną teorię naiwną. Określają ją następujące aksjomaty:

$$(Ax =) \forall x \forall y (\forall z (z \in x \leftrightarrow z \in y) \rightarrow x = y),$$

$$(AxCan) \exists y \forall x (x \in y \leftrightarrow \phi(x)),$$

gdzie $\phi(x)$ jest dowolną formułą języka teorii, w której x jest przynajmniej jedną zmienną wolną oraz w której y nie jest zmienną wolną.

$(Ax =)$ zwany jest aksjomatem identyczności lub ekstensjonalności. $(AxCan)$ zwany jest aksjomatem lub pewnikiem Cantora. Nie jest to właściwie aksjomat (formuła języka) lecz schemat aksjomatu: w zależności od konkretnej postaci formuły $\phi(x)$ uzyskujemy z $(AxCan)$ konkretny aksjomat teorii.

Oba aksjomaty mają na celu formalnie ujmować następującą intuicję: dla dowolnie pomyślanej własności istnieje zbiór (mnogość) tych i tylko tych obiektów, którym ta własność przysługuje. Owa własność reprezentowana jest formalnie formułą $\phi(x)$ w $(AxCan)$. Właśnie $(AxCan)$ stwierdza istnienie zbioru y tych i tylko tych obiektów x , którym „własność” $\phi(x)$ przysługuje. Jednakże samo wyrażenie $(AxCan)$ nie wystarcza dla oddania owej intuicji. Bowiem niejawnie jest w niej mowa o dokładnie jednym zbiorze tych i tylko tych obiektów, którym dana własność przysługuje. Tymczasem $(AxCan)$ nie gwarantuje wcale istnienia dokładnie jednego zbioru y złożonego z obiektów x , dla których prawdą jest, że $\phi(x)$. Przykładowo rozważmy $\phi(x)$ postaci

$$\forall z(x \neq z \rightarrow x \in z),$$

odpowiadającą własności: jest częścią czegośkolwiek różnego od siebie. Wówczas z $(AxCan)$ otrzymujemy zdanie:

$$\exists y \forall x(x \in y \leftrightarrow \forall z(x \neq z \rightarrow x \in z)).$$

Zinterpretujmy je w strukturze relacyjnej $(\{a, b, c\}, \in)$, gdzie a, b, c są różnymi od siebie obiektami oraz relacja $\in$ (nie odróżniana tu symbolicznie od predykatu $\in$) jest taka, że $c \in a$ i $c \in b$. Wówczas łatwo sprawdzić, że prawdziwe są formuły:

$$\forall x(x \in a \leftrightarrow \forall z(x \neq z \rightarrow x \in z)) \text{ oraz}$$

$$\forall x(x \in b \leftrightarrow \forall z(x \neq z \rightarrow x \in z))$$

(nie odróżniamy tu obiektów a, b od stałych indywidualnych będących ich nazwami w tej interpretacji), tzn. istnieją dwa różne „zbiory” y dla których prawdą jest

$$\forall x(x \in y \leftrightarrow \forall z(x \neq z \rightarrow x \in z)).$$

Dodanie aksjomatu $(Ax =)$ powoduje, że ów zbiór y , którego istnienie stwierdza $(AxCan)$ jest dokładnie jeden. $(Ax =)$ stwierdza utożsamienie

zbiorów mających te same elementy. W podanej tu przykładowo interpretacji oczywiście $(Ax =)$ jest fałszywy, bowiem obiekty a, b mają jedyny element c (jeden i ten sam), nie są zaś identyczne.

Bardziej ogólnie, załóżmy, że w obecności $(Ax=)$, dla ustalonej (konkretnej) formuły $\phi(x)$ prawdziwe są dwie formuły uzyskane z $(AxCan)$:

$$(1) \quad \forall x(x \in a \leftrightarrow \phi(x)) \text{ oraz}$$

$$(2) \quad \forall x(x \in b \leftrightarrow \phi(x)).$$

Wówczas mamy:

$$(3) \quad \forall x(x \in a \leftrightarrow x \in b).$$

Weźmy bowiem dowolny x i w celu pokazania równoważności: $x \in a \leftrightarrow x \in b$, załóżmy, że $x \in a$. Wówczas z (1) mamy: $\phi(x)$, zatem z (2): $x \in b$. Tym samym mamy implikację: $x \in a \rightarrow x \in b$. Dowód implikacji odwrotnej: $x \in b \rightarrow x \in a$, jest analogiczny.

(3) w połączeniu z $(Ax =)$ w postaci: $\forall x(x \in a \leftrightarrow x \in b) \rightarrow a = b$, prowadzi do: $a = b$. Zatem zbiór, którego istnienie stwierdza $(AxCan)$ (dla ustalonej $\phi(x)$) jest dokładnie jeden.

2.4.2 Paradoks Russella

Niestety, jak wiadomo, intuicja, której formalnym ujęciem są aksjomaty $(Ax =)$, $(AxCan)$ jest naiwna, bowiem jest nieuzasadniona by nie rzec fałszywa. Teoria oparta na tych aksjomatach jest bowiem sprzeczna (tzn. jest zbiorem wszystkich zdań swojego języka). Rozważmy mianowicie formułę $\phi(x)$ postaci: $x \notin x$, uzyskując z $(AxCan)$ aksjomat:

$$\exists y \forall x(x \in y \leftrightarrow x \notin x).$$

Oznaczmy ów jedyny zbiór, którego istnienie ta formuła stwierdza symbolem a . Wówczas

$$\forall x(x \in a \leftrightarrow x \notin x), \text{ a stąd}$$

$$a \in a \leftrightarrow a \notin a,$$

Zdanie to, będące przecież postaci $A \leftrightarrow \neg A$, w żadnej interpretacji nie jest prawdziwe. Sprzeczność ta nosi nazwę antynomii Russella, przyczyniła się ona do rozwoju badań teoriomnogościowych owocujących innymi niż naiwna teoriami mnogości.

2.5 Teoria zbiorów Zermelo-Fraenkla

Teoria ZF nie posiada przedstawionego mankamentu teorii naiwnej. Oparta jest ona na innym zestawie aksjomatów (zachowany jest aksjomat identyczności), jednakże formuły postaci $(AxCan)$ są w niej obecne, choć nie dla wszystkich formuł $\phi(x)$. Z powodu obecności $(Ax=)$ w ZF , zbiór y , którego istnienie stwierdza w teorii ZF formuła $\exists y \forall x (x \in y \leftrightarrow \phi(x))$, jest jedyny. W ogólności, jest on oznaczany w znany sposób, jako: $\{x : \phi(x)\}$ (zbiór wszystkich takich x , że $\phi(x)$).

Zatem w teorii ZF również stwierdza się istnienie zbiorów, których elementami są te i tylko te zbiory, którym przysługuje „własność” $\phi(x)$. Jednakże nie dla każdej „własności” taki zbiór istnieje. Przykładowo, można wykazać w ZF , że, zgodnie z krytyką Russella teorii naiwnej, nie istnieje zbiór wszystkich zbiorów nie będących swoimi elementami ($\phi(x)$ postaci $x \notin x$), co (w obecności aksjomatu ufundowania) jest równoważne nieistnieniu zbioru wszystkich zbiorów; innym przykładem jest nieistnienie zbioru wszystkich liczb porządkowych ($\phi(x)$ postaci: x jest liczbą porządkową).

2.5.1 Aksjomaty teorii mnogości ZF^- (bez aksjomatów ufundowania i wyboru)

W literaturze przedmiotu spotyka się różne zestawy aksjomatów teorii ZF^- (por. np. Beth[2], Fraenkel[7], Fraenkel, Bar-Hillel, Levy[8], Guzik, Zbierski[11], Jech[15], Kuratowski, Mostowski[16], Morse[18], Schoenfield[21], Takeuti, Zaring[23]). Wybieramy zestaw sześciu aksjomatów z encyklopedycznej wersji teorii ZF^- Schoenfield[21]. Aksjomaty są formułami, w których poza stałymi logicznymi, tzn. spójnikami $\neg, \wedge, \vee, \rightarrow, \leftrightarrow$, kwantyfikatorami $\forall, \exists$ oraz predykatem identyczności $=$, występują zmienne indywidualne x, y, z, u, v, w oraz jedyny predykat: $\in$.

Adekwatne objaśnienie znaczenia niektórych aksjomatów jest możliwe dopiero po wprowadzeniu do teorii odpowiednich pojęć. Wówczas bowiem jest możliwy inny, równoważny zapis owych aksjomatów, w znacznie krótszej postaci, w której obok predykatu $\in$ pojawiają się symbole tych pojęć.

Aksjomat identyczności (lub ekstensjonalności):

$(Ax =) \forall x \forall y (\forall z (z \in x \leftrightarrow z \in y) \rightarrow x = y)$.

Zbiór x jest tożsamy ze zbiorem y , o ile zachodzi $\forall z (z \in x \leftrightarrow z \in y)$ czyli gdy x, y mają te same elementy.

Aksjomat identyczności umożliwia definiowanie danego zbioru (o ile on istnieje) przez podanie jakie zbiory są jego elementami. Wiedząc bowiem jakie zbiory stanowią wszystkie elementy danego zbioru, mamy, dzięki $(Ax =)$, jednoznacznie ten zbiór wyznaczony.

Nie oznacza to, że dany zbiór można utożsamić z „sekwencją” czy „ekspozycją” jego elementów. Status ontyczny zbioru „eksponowanych” elementów jest taki sam jak status każdego z jego elementów (są to obiekty „świata” zbiorów), lecz różny od statusu „ekspozycji”.

Aksjomat Zermelo (lub podzbiorów albo selekcji):

$$(AxZ)_\phi \quad \forall x(\phi(x) \rightarrow x \in z) \rightarrow \exists y \forall x(x \in y \leftrightarrow \phi(x)),$$

gdzie $\phi(x)$ jest dowolną formułą języka teorii, w której x jest przynajmniej jedną zmienną wolną oraz w której y nie jest zmienną wolną.

Podobnie jak rozważane wcześniej aksjomat indukcji oraz aksjomat Cantora, (AxZ) nie jest pojedynczym (konkretnym) aksjomatem teorii, lecz klasą aksjomatów, z której „wyjmuje” się konkretny aksjomat $(AxZ)_\phi$, dla konkretnej formuły $\phi(x)$.

Ponadto, jak widać, zmienna z jest w $(AxZ)_\phi$ zmienną wolną. Jednakże wszystkie formuły teorii ZF (dotyczy to jakiegokolwiek teorii pierwszego rzędu), zatem również jej aksjomaty, winny być domknięte, tzn. nie występują w nich zmienne wolne. Fakt, że w $(AxZ)_\phi$ pojawia się co najmniej jedna zmienna wolna z – być może w konkretnie wziętej formule $\phi(x)$ występują jeszcze inne zmienne wolne niż zmienna x , które to zmienne są wolne w całej formule $(AxZ)_\phi$ – jest oparty na powszechnie przyjętej konwencji notacyjnej, według której formułę $\psi(x_1, \dots, x_n)$, w której $x_1, \dots, x_n$ są wszystkimi jej zmiennymi wolnymi postrzega się jako formułę uniwersalnie domkniętą: $\forall x_1 \dots \forall x_n(\psi(x_1, \dots, x_n))$, tzn. przed ową formułą ze zmiennymi wolnymi „widzi” się kwantyfikatory uniwersalne wiążące wszystkie te zmienne wolne. W przypadku zapisu formuły $(AxZ)_\phi$ zastosowanie owej konwencji notacyjnej jest bardzo wygodne, bowiem to, jakie zmienne mają być wiązane kwantyfikatorami uniwersalnymi umieszczonymi przed $(AxZ)_\phi$ zależy przecież od konkretnej postaci formuły $\phi(x)$.

Jest widoczne, że następnik implikacji $(AxZ)_\phi$ jest identyczny z formułą $(AxCan)$. Funkcja jaką spełnia $(AxZ)_\phi$ w teorii ZF jest więc podobna do tej jaką aksjomat Cantora pełni w naiwnej teorii mnogości. Mianowicie, dla konkretnej formuły $\phi(x)$, z aksjomatu $(AxZ)_\phi$ wnioskujemy również istnie-

nie zbioru y tych i tylko tych zbiorów x , dla których zachodzi $\phi(x)$. Jednakże wnioskowanie to jest uprawnione jedynie wówczas, gdy $\phi(x)$ jest taką formułą, że spełniony jest poprzednik $(AxZ)_\phi : \forall x(\phi(x) \rightarrow x \in z)$, gdzie z jest jakimś dowolnie wybranym, ustalonym zbiorem.

Stwierdzenie prawdziwości owego poprzednika zabezpiecza klasę aksjomatów (AxZ) przed sprzecznością. Bowiem to, że jest on prawdziwy, oznacza, że elementy zbioru y , którego istnienie stwierdza następnik w $(AxZ)_\phi$ należą do wcześniej danego zbioru z ; innymi słowy, $(AxZ)_\phi$ umożliwia utworzenie zbioru y z elementów danego zbioru z , czyli utworzenie *podzbioru* zbioru z . Intuicyjnie rzecz ujmując, nie widać nic absurdalnego czy prowadzącego do sprzeczności w możliwości „łączenia” niektórych elementów danego zbioru (tutaj oznaczonego symbolem „ z ”) w całość zwaną zbiorem (tutaj oznaczonym symbolem „ y ”).

To zabezpieczenie przed sprzecznością powoduje jednakże, że „siła dedukcyjna” klasy formuł (AxZ) jest słabsza niż klasy formuł $(AxCan)$. Stąd teoria ZF^- jest wyposażona w inne jeszcze aksjomaty niż teoria naiwna, tak, aby intuicyjnie pojmowany „świat” zbiorów, który miał być ujęty teorią naiwną, był opisany w teorii ZF^- .

Z punktu widzenia historii rozwoju aksjomatyki teorii ZF należy podkreślić, że aksjomat podany przez Zermelo oraz występujący w literaturze pod nazwą aksjomatu Zermelo ma właściwie postać następującą:

$$(AxZer)_\phi \exists y \forall x (x \in y \leftrightarrow (x \in z \wedge \phi(x))).$$

Łatwo wykazać, że $(AxZ)_\phi$ jest implikowany przez $(AxZer)_\phi$: załóżmy bowiem $(AxZer)_\phi$ oraz poprzednik

$$(1) \forall x (\phi(x) \rightarrow x \in z)$$

implikacji $(AxZ)_\phi$. Niech a będzie zbiorem, którego istnienie $(AxZer)_\phi$ stwierdza. Wówczas mamy:

$$(2) \forall x (x \in a \leftrightarrow (x \in z \wedge \phi(x))).$$

Wykazujemy:

$$(3) \forall x (x \in a \leftrightarrow \phi(x)).$$

$(\rightarrow)$: Niech $x \in a$. Wówczas z (2): $x \in z \wedge \phi(x)$, zatem $\phi(x)$.

$(\leftarrow)$: Załóżmy, że zachodzi $\phi(x)$. Wówczas z (1): $x \in z$. Zatem $x \in z \wedge \phi(x)$. Ostatecznie $x \in a$ na mocy (2).

Na podstawie (3) otrzymujemy bezpośrednio następnik implikacji $(AxZ)_\phi : \exists y \forall x (x \in y \leftrightarrow \phi(x))$.

Równie łatwo jest wykazać, że klasa formuł (AxZ) implikuje $(AxZer)_\phi$ dla dowolnej formuły ϕ : rozważmy bowiem dowolną $\phi(x)$ oraz załóżmy $(AxZ)_\psi$, gdzie $\psi(x)$ jest postaci: $x \in z \wedge \phi(x)$, tzn. załóżmy, że spełniona jest formuła:

$$(4) \quad \forall x ((x \in z \wedge \phi(x)) \rightarrow x \in z) \rightarrow \exists y \forall x (x \in y \leftrightarrow (x \in z \wedge \phi(x))).$$

Ponieważ poprzednik implikacji (4) jest prawdziwy, mamy więc następnik czyli $(AxZer)_\phi$.

Aksjomat zbioru potęgowego:

$$(AxP) \quad \forall u \exists y \forall x (x \in y \leftrightarrow \forall z (z \in x \rightarrow z \in u)).$$

Dla dowolnego zbioru u istnieje *zbiór potęgowy* zbioru u , tzn. zbiór, którego elementami są wszystkie podzbiory zbioru u i tylko one (formuła $\forall z (z \in x \rightarrow z \in u)$ oznacza, że x jest podzbiorem zbioru u).

Aksjomat sumy:

$$(Ax \cup) \quad \forall u \exists y \forall x (x \in y \leftrightarrow \exists z (z \in u \wedge x \in z)).$$

Dla dowolnego zbioru u istnieje *suma* zbioru u , tzn. zbiór, którego elementami są te i tylko te zbiory, które są elementami jakiegoś elementu zbioru u .

Aksjomat podstawiania (lub zastępowania):

$$(AxSUB)_\psi \quad \forall y \exists z (\psi(y, z) \wedge \forall v (\psi(y, v) \rightarrow v = z)) \rightarrow \\ \exists u \forall z (z \in u \leftrightarrow \exists y (y \in x \wedge \psi(y, z))),$$

gdzie $\psi(y, z)$ jest dowolną formułą języka teorii, w której y, z są zmiennymi wolnymi, zaś $\psi(y, v)$ jest formułą otrzymaną z $\psi(y, z)$ przez zastąpienie każdego wolnego występowania zmiennej z w $\psi(y, z)$ zmienną v .

Pod pewnymi względami $(AxSUB)_\psi$ zachowuje się analogicznie jak $(AxZ)_\phi$. Jest, jak widać, klasą różnych aksjomatów w zależności od różnych

formułę $\psi(y, z)$. Ponadto wnioskujemy z konkretnego aksjomatu $(AxSUB)_\psi$ istnienie pewnego zbioru u , gdy poprzednik implikacji $(AxSUB)_\psi$ jest spełniony.

Ów poprzednik jest spełniony wówczas, gdy formuła $\psi(y, z)$ jest taka, że dla każdego zbioru y istnieje dokładnie jeden zbiór z taki, że zachodzi $\psi(y, z)$. Krótko mówiąc, wówczas, gdy dysponujemy formułą $\psi(y, z)$ ustalającą jakieś przyporządkowanie każdemu zbiorowi y dokładnie jednego zbioru z , mamy gwarancję istnienia zbioru u , o którym mowa w następniku implikacji $(AxSUB)_\psi$. Ów zbiór u jest zależny od jeszcze jednego parametru jakim jest zbiór oznaczony w $(AxSUB)_\psi$ zmienną (wolną) x . Mając więc dowolnie wybrany, ustalony zbiór x , na mocy $(AxSUB)_\psi$ stwierdzamy istnienie zbioru u tych i tylko tych zbiorów z , które są jednoznacznie przyporządkowane elementom zbioru x według formuły ψ .

Aksjomat nieskończoności:

$$(Ax\infty) \exists x[\exists y(y \in x \wedge \forall z(z \notin y)) \wedge \forall y(y \in x \rightarrow \exists z(z \in x \wedge \forall u(u \in z \leftrightarrow (u \in y \vee u = y))))].$$

Jest to jedyny z aksjomatów teorii ZF , który bezwarunkowo i niezależnie od istnienia innych zbiorów stwierdza istnienie pewnego zbioru, oznaczonego tu zmienną x , o własnościach opisanych formułą w nawiasach kwadratowych, i ze względu na te własności nazywanego zbiorem *indukcyjnym*. Własności te powodują, że zbiór indukcyjny ma nieskończenie wiele elementów. Aksjomat ten stwierdza więc istnienie w „świecie” zbiorów zbioru nieskończonego.

2.5.2 Inkluzja zbiorów

Mimo ubóstwa pojęć pierwotnych, teoria ZF^- dysponuje niejakim bogactwem pojęć definiowanych. Tak jak w każdej teorii pierwszego rzędu, pojęcia te dzielimy na dwa rodzaje: relacje i operacje (w tym tzw. 0-argumentowe operacje tożsame z wyróżnionymi obiektami). Po stronie językowej odpowiadają tym pojęciom, odpowiednio predykaty i symbole funkcyjne. W związku z tym mamy dwa rodzaje definicji nominalnych (tzn. definiujących wyrażenia): definicje predykatów oraz definicje symboli funkcyjnych.

Definicja predykatu wprowadza do języka teorii nowy predykat, w ten sposób, że formuła atomowa, zbudowana w oparciu o ten predykat, jest skrótem dla pewnej, zazwyczaj dość złożonej formuły, w której ów nowy predykat naturalnie nie występuje. Relację czy stosunek wyrażony definiowanym predykatem można więc również określić bez jego wprowadzania – przy pomo-

cy owej formuły złożonej. W konsekwencji, kosztem długości zapisu, każdą z relacji rozważanych w teorii (poza pierwotnymi) można by wyrażać bez jej nazwy – predykatu, a wyłącznie przez odwoływanie się do formuł zawierających tylko terminy pierwotne teorii. Czego się naturalnie ze względów praktycznych nie czyni.

W przeciwieństwie do definicji predykatu, definicja symbolu funkcyjnego na ogół nie określa żadnego skrótu dla innej formuły, lecz dostarcza nazwy dla operacji, której istnienie jest niejawnie opisane w aksjomatach lub twierdzeniach teorii. Bez takiej definicji nie można operacji z aksjomatów czy twierdzeń „wydobyć”. Jeżeli operacja jest 0-argumentowa, czyli gdy jest ona wyróżnionym obiektem, jego istnienie i jedyność muszą być zagwarantowane w owych aksjomatach czy twierdzeniach. Dopiero po pokazaniu tych „gwarancji” można definicyjnie wprowadzić stałą indywidualną nazywającą ów obiekt. Podobnie, aksjomaty czy twierdzenia teorii muszą gwarantować istnienie n -argumentowego ($n \geq 1$) przyporządkowania (operacji) dowolnej sekwencji n obiektów dokładnie jednego obiektu, aby można było definicyjnie wprowadzić symbol funkcyjny wyrażający to przyporządkowanie.

Wprowadzanie pojęć niepierwotnych teorii ZF zaczniemy od 2-argumentowego stosunku *inkluzji*.

Uwaga: Na oznaczanie obiektów opisywanych w teorii, a więc zbiorów, będziemy w dalszym ciągu używać zmiennych indywidualnych w postaci zarówno małych jak i wielkich liter alfabetu łacińskiego, traktując zmienne postaci x, X jako różne. Przypominamy, że stosujemy skrót „wtw” dla wyrażenia „wtedy i tylko wtedy gdy”.

DEFINICJA. Definiujemy 2-argumentowy predykat *inkluzji* $\subseteq$ (zawierania się zbiorów): $X \subseteq Y \leftrightarrow \forall x(x \in X \rightarrow x \in Y)$.

(Dla dowolnych zbiorów X, Y zbiór X zawiera się w zbiorze Y lub X jest podzbiorem zbioru Y wtw każdy element zbioru X jest elementem zbioru Y).

TWIERDZENIE 1. $X \subseteq X$.

DOWÓD: Rozważmy dowolny zbiór X . Wówczas prawdziwa jest formuła $\forall x(x \in X \rightarrow x \in X)$ (biorąc bowiem pod uwagę dowolny zbiór x , prawdziwa jest implikacja $x \in X \rightarrow x \in X$). Z definicji inkluzji, $X \subseteq X$. $\square$

TWIERDZENIE 2. $(X \subseteq Y \wedge Y \subseteq X) \rightarrow X = Y$.

DOWÓD: Załóżmy, że

- (1) $X \subseteq Y$,
- (2) $Y \subseteq X$.

Wykażemy najpierw, że

- (3) $\forall x(x \in X \leftrightarrow x \in Y)$.

Weźmy więc dowolny x . Dowodzimy równoważności $x \in X \leftrightarrow x \in Y$.

($\rightarrow$): Niech $x \in X$. Ponieważ na mocy (1) z definicji inkluzji mamy: $\forall x(x \in X \rightarrow x \in Y)$, zatem również: $x \in X \rightarrow x \in Y$, więc $x \in Y$.

Analogicznie dowodzimy odwrotnej implikacji korzystając z (2). Skorzystajmy teraz z aksjomatu ($Ax =$) uzyskując wyrażenie $\forall x(x \in X \leftrightarrow x \in Y) \rightarrow X = Y$. Zatem, na mocy (3), $X = Y$. $\square$

Uwaga. Twierdzenie 2 jest bardzo szeroko stosowane dla dowodzenia równości zbiorów. Aby wykazać, że $X = Y$ dowodzi się inkluzji ($\subseteq$) tzn., że $X \subseteq Y$; następnie inkluzji ($\supseteq$) tzn., że $Y \subseteq X$ (zob. np. dowód twierdzenia 6).

TWIERDZENIE 3. $(X \subseteq Y \wedge Y \subseteq Z) \rightarrow X \subseteq Z$.

DOWÓD: Załóżmy, że

- (1) $X \subseteq Y$ oraz
- (2) $Y \subseteq Z$.

Aby wykazać, że $X \subseteq Z$ musimy zgodnie z definicją inkluzji dowieść, że $\forall x(x \in X \rightarrow x \in Z)$. Rozważmy więc jakiś zbiór x . Załóżmy, że $x \in X$. Z (1) mamy: $\forall x(x \in X \rightarrow x \in Y)$. Zatem $x \in Y$. Ale z (2): $\forall x(x \in Y \rightarrow x \in Z)$, zatem $x \in Z$. $\square$

2.5.3 Zbiór pusty

Weźmy na chwilę pod uwagę aksjomat nieskończoności ($Ax\infty$); oznaczmy zbiór, którego istnienie on stwierdza symbolem a . Rozważmy następnie aksjomat podzbiorów (AxZ) $_{\phi}$, taki, że formuła $\phi(x)$ jest postaci $x \neq x$:

$\forall z[\forall x(x \neq x \rightarrow x \in z) \rightarrow \exists y\forall x(x \in y \leftrightarrow x \neq x)]$ (po dokonaniu uniwersalnego domknięcia).

Stąd otrzymujemy:

$$\forall x(x \neq x \rightarrow x \in a) \rightarrow \exists y \forall x(x \in y \leftrightarrow x \neq x).$$

Lecz poprzednik tej implikacji jest spełniony (dla dowolnego x , poprzednik implikacji: $x \neq x \rightarrow x \in a$ jest fałszywy, zatem jest ona prawdziwa). Stąd prawdziwy w dziedzinie zbiorów jest następnik $\exists y \forall x(x \in y \leftrightarrow x \neq x)$. Zbiór (dokładnie jeden), którego istnienie on stwierdza oznaczamy symbolem $\emptyset$ i nazywamy zbiorem *pustym*:

DEFINICJA. $\forall x(x \in \emptyset \leftrightarrow x \neq x)$, lub $\emptyset = \{x : x \neq x\}$.

Uwaga. Symbol $\emptyset$ jest stałą indywidualną wprowadzoną definicyjnie do języka teorii.

TWIERDZENIE 4. $\forall x(x \notin \emptyset)$ (*zbiór pusty nie ma elementów*).

DOWÓD: Załóżmy nie wprost, że $\neg \forall x(x \notin \emptyset)$. Wówczas $\exists x(x \in \emptyset)$, zatem dla jakiegoś zbioru a , $a \in \emptyset$. Wówczas z definicji zbioru $\emptyset$, $a \neq a$ co jest niemożliwe. $\square$

TWIERDZENIE 5. $\emptyset \subseteq X$.

DOWÓD: Rozważmy dowolny zbiór X . Wówczas formuła $\forall x(x \in \emptyset \rightarrow x \in X)$ jest prawdziwa (biorąc bowiem pod uwagę dowolny zbiór x , na mocy twierdzenia 4 fałszywy jest poprzednik implikacji $x \in \emptyset \rightarrow x \in X$, zatem jest ona prawdziwa). Wobec tego z definicji inkluzji, $\emptyset \subseteq X$. $\square$

TWIERDZENIE 6. $\forall x(x \notin X) \rightarrow X = \emptyset$ (*zbiór $\emptyset$ jest jedynym zbiorem nie mającym elementów*).

DOWÓD: Załóżmy, że $\forall x(x \notin X)$. Dowód równości $X = \emptyset$ opieramy na twierdzeniu 2.

($\subseteq$): Pokazujemy, że $X \subseteq \emptyset$, tzn., że $\forall x(x \in X \rightarrow x \in \emptyset)$. Weźmy więc pod uwagę zbiór x . Na mocy założenia: $x \notin X$, zatem implikacja $x \in X \rightarrow x \in \emptyset$ jest prawdziwa (fałszywy poprzednik).

($\supseteq$): Należy pokazać, że $\emptyset \subseteq X$, ale to jest oczywistym wnioskiem z twierdzenia 5.

Ostatecznie na mocy twierdzenia 2 (w postaci: $(X \subseteq \emptyset \wedge \emptyset \subseteq X) \rightarrow X = \emptyset$) wnosimy, że $X = \emptyset$. $\square$

WNIOSEK. Dla dowolnego zbioru X , $\forall x(x \notin X) \leftrightarrow X = \emptyset$ (inaczej: $X \neq \emptyset \leftrightarrow \exists x(x \in X)$).

DOWÓD: $(\rightarrow)$: Z twierdzenia 6.

$(\leftarrow)$: Z twierdzenia 4. $\square$

2.5.4 Zbiór potęgowy zbioru

Biorąc pod uwagę definicję inkluzji, możemy aksjomat zbioru potęgowego (AxP) zapisać w postaci:

$$\forall u \exists y \forall x (x \in y \leftrightarrow x \subseteq u).$$

Rozważmy więc dowolny zbiór U . Wówczas ten jedyny zbiór y , którego istnienie stwierdza formuła: $\exists y \forall x (x \in y \leftrightarrow x \subseteq U)$ oznaczamy symbolem $P(U)$ i nazywamy *zbiorem potęgowym zbioru U* (jest to zbiór wszystkich podzbiorów zbioru U):

DEFINICJA. Dla dowolnego zbioru U ,

$$\forall x (x \in P(U) \leftrightarrow x \subseteq U) \text{ lub } P(U) = \{x : x \subseteq U\}.$$

Uwaga. Symbol P jest jednoargumentowym symbolem funkcyjnym nazywającym jednoargumentową operację przyporządkowującą każdemu zbiorowi jego zbiór potęgowy.

TWIERDZENIE 7. Dla dowolnego zbioru $U : \emptyset, U \in P(U)$.

DOWÓD: Ponieważ $\emptyset \subseteq U$ oraz $U \subseteq U$ (twierdzenia 5 i 1), więc z definicji zbioru potęgowego, $\emptyset \in P(U)$ oraz $U \in P(U)$. $\square$

WNIOSEK. Dla dowolnego zbioru U , $P(U) \neq \emptyset$.

DOWÓD: Na podstawie twierdzenia 7 i wniosku z twierdzenia 6. $\square$

TWIERDZENIE 8. $X \subseteq Y \leftrightarrow P(X) \subseteq P(Y)$.

DOWÓD: $(\rightarrow)$: Załóżmy, że $X \subseteq Y$. Niech $A \in P(X)$. Wówczas z definicji zbioru potęgowego, $A \subseteq X$. Zatem z założenia, na mocy twierdzenia 3, $A \subseteq Y$, czyli $A \in P(Y)$. Wobec dowolności wyboru zbioru A stwierdzamy, że $P(X) \subseteq P(Y)$.

($\leftarrow$): Załóżmy, że $P(X) \subseteq P(Y)$. Ponieważ $X \in P(X)$ (twierdzenie 7), więc z założenia otrzymujemy: $X \in P(Y)$. Stąd $X \subseteq Y$. $\square$

Twierdzenie 9. $\forall x(x \in P(\emptyset) \leftrightarrow x = \emptyset)$ (zbiór potęgowy zbioru pustego ma dokładnie jeden element: zbiór pusty).

Dowód: ($\rightarrow$): Załóżmy, że $x \in P(\emptyset)$. Zatem $x \subseteq \emptyset$. Wówczas, na mocy twierdzeń 5 oraz 2, $x = \emptyset$.

($\leftarrow$): Na mocy twierdzenia 7. $\square$

Twierdzenie 10. $\forall x[x \in P(P(\emptyset)) \leftrightarrow (x = \emptyset \vee x = P(\emptyset))]$ (zbiór potęgowy zbioru potęgowego zbioru pustego ma dokładnie dwa elementy: zbiór pusty oraz zbiór potęgowy zbioru pustego).

Dowód: ($\rightarrow$): Załóżmy, że $x \in P(P(\emptyset))$. Wówczas

(1) $x \subseteq P(\emptyset)$.

W celu wykazania alternatywy $x = \emptyset \vee x = P(\emptyset)$, załóżmy, że $x \neq \emptyset$.

Wówczas, na mocy wniosku z twierdzenia 6, niech

(2) $a \in x$.

Wtedy z (1): $a \in P(\emptyset)$, a więc na mocy twierdzenia 9 mamy:

(3) $a = \emptyset$.

Wykazujemy teraz, że

(4) $P(\emptyset) \subseteq x$.

Niech więc $y \in P(\emptyset)$. Wówczas na podstawie twierdzenia 9, $y = \emptyset$. Zatem z (3): $y = a$ i konsekwentnie na mocy (2): $y \in x$.

Z (1) i (4) wnioskujemy, na mocy twierdzenia 2, iż $x = P(\emptyset)$.

($\leftarrow$): Na mocy twierdzenia 7. $\square$

2.5.5 Suma zbioru

Weźmy pod uwagę aksjomat sumy ($Ax \cup$) oraz dowolny zbiór u . Ten jedyny zbiór y , którego istnienie stwierdza formuła:

$$\exists y \forall x(x \in y \leftrightarrow \exists z(z \in u \wedge x \in z))$$

nazywamy *sumą (uogólnioną)* $\bigcup u$ zbioru u .

DEFINICJA. Dla dowolnego zbioru u :

$$\forall x(x \in \bigcup u \leftrightarrow \exists z(z \in u \wedge x \in z)) \quad \text{lub} \quad \bigcup u = \{x : \exists z(z \in u \wedge x \in z)\}.$$

Uwaga. Symbol $\bigcup$ jest jednoargumentowym symbolem funkcyjnym nazywającym operację przyporządkowującą każdemu zbiorowi u jego sumę $\bigcup u$ czyli zbiór tych i tylko tych elementów, które należą do przynajmniej jednego elementu zbioru u .

TWIERDZENIE 11. *Dla dowolnego zbioru u :*

- (1) $\forall x(x \in u \rightarrow x \subseteq \bigcup u)$, czyli $u \subseteq P(\bigcup u)$,
- (2) dla dowolnego y , $\forall x(x \in u \rightarrow x \subseteq y) \rightarrow \bigcup u \subseteq y$,
- (3) $\bigcup P(u) = u$,
- (4) $\bigcup \emptyset = \emptyset$.

DOWÓD: dla (1): Załóżmy, że $x \in u$. Aby wykazać inkluzję $x \subseteq \bigcup u$, niech $y \in x$. Zatem $\exists z(z \in u \wedge y \in z)$. Wobec tego z definicji sumy, $y \in \bigcup u$.

dla (2): Załóżmy, że y jest takim zbiorem, że $\forall x(x \in u \rightarrow x \subseteq y)$. Niech $v \in \bigcup u$. Wówczas $\exists z(z \in u \wedge v \in z)$. Zatem dla pewnego zbioru a , $a \in u$ oraz $v \in a$. Wobec tego z założenia mamy $a \subseteq y$ (bo $a \in u$) skoro więc $v \in a$, to $v \in y$, co dowodzi inkluzji $\bigcup u \subseteq y$.

dla (3): ($\supseteq$): Na mocy (1) zapisanego dla zbioru $P(u)$ w miejsce u mamy: $u \in P(u) \rightarrow u \subseteq \bigcup P(u)$, zatem według twierdzenia 7, $u \subseteq \bigcup P(u)$. Aby dowieść inkluzji przeciwnej zauważmy, że $\forall x(x \in P(u) \leftrightarrow x \subseteq u)$ i zastosujmy (2), gdzie w miejscu zbioru u wystąpi $P(u)$, zaś w miejscu zbioru y zbiór u . Z (2) mamy wówczas: $\forall x(x \in P(u) \rightarrow x \subseteq u) \rightarrow \bigcup P(u) \subseteq u$. Zatem, wobec prawdziwości poprzednika, otrzymujemy: $\bigcup P(u) \subseteq u$.

dla (4): Załóżmy nie wprost, że $\bigcup \emptyset \neq \emptyset$. Wtedy dla pewnego a , $a \in \bigcup \emptyset$. Zatem $\exists z(z \in \emptyset \wedge a \in z)$ skąd dla pewnego b , $b \in \emptyset$, co jest niemożliwe. $\square$

2.5.6 Para zbiorów, zbiór jednoelementowy

Rozważmy aksjomat podstawiania $(AxSUB)_\psi$ dla formuły $\psi(y, z)$ postaci: $(y = \emptyset \wedge z = X) \vee (y \neq \emptyset \wedge z = Y)$ oraz dla zbioru oznaczonego zmienną x postaci $P(P(\emptyset))$:

$$\forall X \forall Y [\forall y \exists z (\psi(y, z) \wedge \forall v (\psi(y, v) \rightarrow v = z)) \rightarrow \exists u \forall z (z \in u \leftrightarrow \exists y (y \in P(P(\emptyset)) \wedge \psi(y, z)))]$$

Weźmy pod uwagę dowolne zbiory X, Y . Wówczas spełniony jest poprzednik implikacji w $(AxSUB)_\psi$, tzn. formuła $\forall y \exists z (\psi(y, z) \wedge \forall v (\psi(y, v) \rightarrow v = z))$. Weźmy bowiem pod uwagę dowolny zbiór y . Naturalnie $y = \emptyset$

lub $y \neq \emptyset$. Gdy $y = \emptyset$, mamy: $\psi(y, X) \wedge \forall v(\psi(y, v) \rightarrow v = X)$, gdy zaś $y \neq \emptyset$ wówczas prawdą jest: $\psi(y, Y) \wedge \forall v(\psi(y, v) \rightarrow v = Y)$.

Formuła $\psi(y, z)$ ustala więc jednoznaczne przyporządkowanie każdemu zbiorowi y dokładnie jednego zbioru z , mianowicie, zbiorowi pustemu przyporządkowany jest zbiór X , zaś każdemu zbiorowi niepustemu przyporządkowany jest zbiór Y .

Prawdziwy jest więc następnik implikacji w $(AxSUB)_\psi$ czyli formuła

$$(1) \quad \exists u \forall z (z \in u \leftrightarrow \exists y (y \in P(P(\emptyset)) \wedge \psi(y, z))).$$

Wyrażenie $\exists y (y \in P(P(\emptyset)) \wedge \psi(y, z))$ jest równoważne na mocy twierdzenia 10 formule

$$\exists y [(y = \emptyset \vee y = P(\emptyset)) \wedge ((y = \emptyset \wedge z = X) \vee (y \neq \emptyset \wedge z = Y))],$$

która z kolei (bierzemy tu pod uwagę wniosek z twierdzenia 7 w postaci: $P(\emptyset) \neq \emptyset$) jest równoważna wyrażeniu $z = X \vee z = Y$. Ostatecznie z (1) otrzymujemy:

$$\exists u \forall z (z \in u \leftrightarrow (z = X \vee z = Y)).$$

Ów jedyny zbiór u , którego istnienie stwierdza ostatnia formuła oznaczamy w postaci: $\{X, Y\}$ i nazywamy *parą zbiorów* lub, gdy $X \neq Y$ – *zbiorem 2-elementowym*:

DEFINICJA. Dla dowolnych zbiorów X, Y ,

$$\forall z (z \in \{X, Y\} \leftrightarrow (z = X \vee z = Y)) \text{ lub } \{X, Y\} = \{z : z = X \vee z = Y\}.$$

Ponadto dla dowolnego zbioru X , $\{X\} = \{X, X\}$, tzn. $\{X\} = \{z : z = X\}$ bądź $\forall z (z \in \{X\} \leftrightarrow z = X)$. Zbiór $\{X\}$ nazywamy *singletonem* lub zbiorem *1-elementowym*, którego jedynym elementem jest X .

PRZYKŁAD. Na podstawie twierdzeń 9 i 10 oraz $(Ax=)$, $P(\emptyset) = \{\emptyset\}$ oraz $P(P(\emptyset)) = \{\emptyset, P(\emptyset)\}$.

2.5.7 Operacje boolowskie na zbiorach, zbiór n -elementowy

DEFINICJA. Dla dowolnych zbiorów x, y zbiór $\bigcup\{x, y\}$, czyli sumę parę zbiorów $\{x, y\}$ nazywamy *sumą zbiorów* x, y i oznaczamy w postaci: $x \cup y$.

Uwaga. Symbol $\cup$ jest 2-argumentowym symbolem funkcyjnym nazywającym operację przyporządkowującą zbiorom x, y ich sumę $x \cup y$.

Twierdzenie poniżej konstryuuje pojęcie sumy dwóch zbiorów w znany sposób:

Twierdzenie 12. *Dla dowolnych zbiorów x, y ,
 $\forall z(z \in x \cup y \leftrightarrow (z \in x \vee z \in y))$ (suma dwóch zbiorów jest zbiorem tych i tylko tych zbiorów, które należą do co najmniej jednego z tych dwóch zbiorów).*

Dowód: $(\rightarrow)$: Załóżmy, że $z \in x \cup y$, tzn. $z \in \bigcup\{x, y\}$. Wówczas dla pewnego $a \in \{x, y\}$, $z \in a$. Z definicji pary zbiorów, $a = x$ lub $a = y$. Niech $a = x$. Wówczas $z \in x$, stąd mamy: $z \in x \vee z \in y$. Analogicznie, gdy $a = y$.

$(\leftarrow)$: Załóżmy, że $z \in x$ lub $z \in y$. Niech $z \in x$. Ponieważ $x = x$, więc z definicji pary zbiorów, $x \in \{x, y\}$. Zatem $\exists u(u \in \{x, y\} \wedge z \in u)$. Ostatecznie z definicji sumy, $z \in \bigcup\{x, y\}$, tzn. $z \in x \cup y$. Analogicznie, gdy $z \in y$. $\square$

Twierdzenie 13. *Dla dowolnych zbiorów A, B, X :*

- (1) $A \subseteq A \cup B$, $B \subseteq A \cup B$,
- (2) $(A \subseteq X \wedge B \subseteq X) \rightarrow A \cup B \subseteq X$,
- (3) $A \subseteq B \leftrightarrow A \cup B = B$.

Dowód: dla (1): Niech $x \in A$. Wówczas $x \in A \vee x \in B$, zatem według twierdzenia 12, $x \in A \cup B$ czyli $A \subseteq A \cup B$. Analogicznie dla drugiej inkluzji.

dla (2): Załóżmy, że $A \subseteq X, B \subseteq X$. Niech $x \in A \cup B$. Wówczas $x \in A \vee x \in B$. Załóżmy, że $x \in A$. Wtedy $x \in X$ (skoro $A \subseteq X$). Załóżmy, że $x \in B$. Wówczas również $x \in X$ (bo $B \subseteq X$). Ostatecznie $x \in X$, czyli $A \cup B \subseteq X$.

dla (3): $(\rightarrow)$: Załóżmy, że $A \subseteq B$. Wykazujemy, że $A \cup B = B$.

$(\subseteq)$: Niech $x \in A \cup B$. Wówczas $x \in A \vee x \in B$. Musimy wykazać, że $x \in B$. Wystarczy więc rozważyć przypadek, gdy alternatywa ta jest prawdziwa w ten sposób, że jej lewy człon jest prawdziwy, tzn. $x \in A$. Wówczas z założenia ($A \subseteq B$) mamy: $x \in B$.

$(\supseteq)$: Inkluzja $B \subseteq A \cup B$ zachodzi na mocy (1).

$(\leftarrow)$: Załóżmy, że $A \cup B = B$. Niech $x \in A$. Wówczas $x \in A \vee x \in B$. Stąd $x \in A \cup B$. Zatem z założenia, $x \in B$. Ostatecznie $A \subseteq B$. $\square$

Definicja. Dla dowolnego $n = 2, 3, \dots$, dla dowolnych zbiorów $x_1, \dots, x_n, x_{n+1}$, $\{x_1, \dots, x_n, x_{n+1}\} = \{x_1, \dots, x_n\} \cup \{x_{n+1}\}$.

Gdy wszystkie zbiory $x_1, \dots, x_n$ ($n \geq 1$) są różne od siebie, zbiór $\{x_1, \dots, x_n\}$ nazywamy *n-elementowym*.

Zbiór X nazywamy *skończonym*, gdy $X = \emptyset$ lub istnieje liczba naturalna $n \geq 1$ oraz zbiory $x_1, \dots, x_n$ takie, że $X = \{x_1, \dots, x_n\}$. Zbiór nazywamy *nieskończonym*, gdy nie jest on skończony.

Uwaga. Symbol $\{\dots\}$ traktujemy jako *n*-argumentowy symbol funkcyjny nazywający *n*-argumentową operację przyporządkowującą sekwencji *n* zbiorów zbiór, którego elementami są wszystkie zbiory z tej sekwencji i tylko one. Intuicyjnie rzecz biorąc, ilość elementów zbioru *n*-elementowego wynosi *n*.

Twierdzenie 14. Dla dowolnego $n \geq 1$, dla dowolnych zbiorów $x_1, \dots, x_n$, $(\{\}) \forall x(x \in \{x_1, \dots, x_n\} \leftrightarrow (x = x_1 \vee \dots \vee x = x_n))$.

Dowód: (indukcyjny) Dla $n = 1$ oraz $n = 2$ warunek $(\{\})$ jest spełniony na mocy definicji singletonu oraz definicji pary zbiorów.

Założmy, że dla pewnego $n \geq 1$ warunek $(\{\})$ zachodzi dla dowolnych zbiorów $x_1, \dots, x_n$. Mamy wykazać, że dla dowolnych zbiorów $x_1, \dots, x_n, x_{n+1}$, $\forall x(x \in \{x_1, \dots, x_n, x_{n+1}\} \leftrightarrow (x = x_1 \vee \dots \vee x = x_n \vee x = x_{n+1}))$.

Weźmy więc pod uwagę jakieś zbiory $x_1, \dots, x_n, x_{n+1}$. Wówczas z definicji zbioru $\{x_1, \dots, x_n, x_{n+1}\}$, Tw.12, założenia indukcyjnego oraz definicji singletonu mamy: $x \in \{x_1, \dots, x_n, x_{n+1}\}$ wtw $x \in \{x_1, \dots, x_n\}$ lub $x \in \{x_{n+1}\}$ wtw $(x = x_1 \vee \dots \vee x = x_n)$ lub $x = x_{n+1}$ wtw $x = x_1 \vee \dots \vee x = x_n \vee x = x_{n+1}$. $\square$

Przykład. $P(\{x, y\}) = \{\emptyset, \{x\}, \{y\}, \{x, y\}\}$ dla dowolnych zbiorów x, y .

Rozważmy aksjomat podzbiorów $(AxZ)_\phi$, gdzie formuła $\phi(x)$ jest postaci $x \in u \wedge x \in v$, oraz w którym zbiór oznaczony zmienną z jest zbiorem u :

$$\forall u \forall v [\forall x ((x \in u \wedge x \in v) \rightarrow x \in u) \rightarrow \exists y \forall x (x \in y \leftrightarrow (x \in u \wedge x \in v))].$$

Weźmy pod uwagę dowolne zbiory u, v . Wówczas poprzednik implikacji w $(AxZ)_\phi$ jest prawdziwy, mamy zatem następnik

$$\exists y \forall x (x \in y \leftrightarrow (x \in u \wedge x \in v)),$$

co prowadzi do definicji *iloczynu* $u \cap v$ zbiorów u, v :

DEFINICJA. Dla dowolnych zbiorów u, v ,

$\forall x(x \in u \cap v \leftrightarrow (x \in u \wedge x \in v))$ lub $u \cap v = \{x : x \in u \wedge x \in v\}$ (iloczyn dwóch zbiorów jest zbiorem tych i tylko tych zbiorów, które należą do każdego z tych dwóch zbiorów).

Zbiory u, v nazywamy *rozłącznymi*, gdy $u \cap v = \emptyset$ (tzn. gdy nie mają one wspólnego elementu).

TWIERDZENIE 15. Dla dowolnych zbiorów A, B, X :

- (1) $A \cap B \subseteq A, A \cap B \subseteq B$,
- (2) $(X \subseteq A \wedge X \subseteq B) \rightarrow X \subseteq A \cap B$,
- (3) $A \subseteq B \leftrightarrow A \cap B = A$.

DOWÓD: dla (1): Niech $x \in A \cap B$. Wówczas $x \in A \wedge x \in B$. Stąd $x \in A$, ostatecznie $A \cap B \subseteq A$. Analogicznie dla drugiej inkluzji.

dla (2): Załóżmy, że $X \subseteq A, X \subseteq B$. Niech $x \in X$. Wówczas z założenia, $x \in A, x \in B$, czyli $x \in A \cap B$. Ostatecznie, $X \subseteq A \cap B$.

dla (3): $(\rightarrow)$: Załóżmy, że $A \subseteq B$. Wykazujemy, że $A \cap B = A$. Inkluzja $(\subseteq)$ zachodzi na mocy (1).

$(\supseteq)$: Niech $x \in A$. Wówczas z założenia, $x \in B$, zatem $x \in A \wedge x \in B$, czyli $x \in A \cap B$.

$(\leftarrow)$: Załóżmy, że $A \cap B = A$. Aby dowieść, że $A \subseteq B$ przypuśćmy, że $x \in A$. Wówczas z założenia, $x \in A \cap B$. Stąd $x \in A \wedge x \in B$, a więc $x \in B$.

□

Postępując analogicznie z aksjomatem podzbiorów jak powyżej, lecz dla formuły $\phi(x)$ postaci: $x \in u \wedge x \notin v$, uzyskujemy definicję *różnicy* $u - v$ zbiorów u, v :

DEFINICJA. Dla dowolnych zbiorów u, v ,

$\forall x(x \in u - v \leftrightarrow (x \in u \wedge x \notin v))$ lub $u - v = \{x : x \in u \wedge x \notin v\}$.

TWIERDZENIE 16. Dla dowolnych zbiorów A, B, X :

- (1) $A - B \subseteq A$,
- (2) $(A - B) \cap B = \emptyset$,
- (3) $(X \subseteq A \wedge X \cap B = \emptyset) \rightarrow X \subseteq A - B$,
- (4) $A \subseteq B \leftrightarrow A - B = \emptyset$.

DOWÓD: dla (1): Niech $x \in A - B$. Wówczas $x \in A \wedge x \notin B$. Stąd $x \in A$. Zatem $A - B \subseteq A$.

dla (2): Załóżmy, że $(A - B) \cap B \neq \emptyset$. Wówczas dla jakiegoś zbioru x , $x \in (A - B) \cap B$. Wtedy $x \in A - B$ oraz $x \in B$. Z definicji różnicy zbiorów, $x \in A \wedge x \notin B$. Stąd $x \notin B$. Sprzeczność.

dla (3): Załóżmy, że $X \subseteq A$ oraz $X \cap B = \emptyset$. Aby wykazać, że $X \subseteq A - B$, weźmy $x \in X$. Wówczas z założenia, że $X \subseteq A$ mamy: $x \in A$. Ponadto $x \notin B$ (gdyby $x \in B$, to $x \in X \cap B$, zatem $X \cap B \neq \emptyset$, co jest niemożliwe na mocy założenia). Zatem $x \in A - B$, czyli wobec dowolności wyboru elementu x , $X \subseteq A - B$.

dla (4): $(\rightarrow)$: Załóżmy, że $A \subseteq B$ oraz że $A - B \neq \emptyset$. Wówczas istnieje taki x , że $x \in A - B$, tzn. $x \in A$ oraz $x \notin B$. Lecz skoro $A \subseteq B$ (założenie) oraz $x \in A$, więc $x \in B$. Sprzeczność.

$(\leftarrow)$: Załóżmy, że $A - B = \emptyset$. W celu wykazania, że $A \subseteq B$ załóżmy, że $x \in A$. Naturalnie $x \in B \vee x \notin B$. Gdyby $x \notin B$, to $x \in A - B$ i konsekwentnie $A - B \neq \emptyset$, co jest niemożliwe na mocy założenia. Zatem $x \in B$, co dowodzi inkluzji $A \subseteq B$. $\square$

DEFINICJA. Niech U będzie dowolnym zbiorem. Dla dowolnego $A \in P(U)$ niech $-A = U - A$. Przy ustalonym zbiorze U , $-$ jest jednoargumentową operacją zwaną operacją *dopełnienia*, tzn. $-A$ nazywamy dopełnieniem zbioru A (do zbioru U).

TWIERDZENIE 17. Dla dowolnych $A, B \in P(U)$:

- (1) $-(A \cup B) = -A \cap -B$,
- (2) $-(A \cap B) = -A \cup -B$,
- (3) $--A = A$,
- (4) $A \subseteq B \leftrightarrow -B \subseteq -A$,
- (5) $-A \cup A = U$,
- (6) $-A \cap A = \emptyset$,
- (7) $-U = \emptyset$,
- (8) $-\emptyset = U$.

DOWÓD: dla (1): $(\subseteq)$: Niech $x \in -(A \cup B)$. Wówczas $x \in U$ oraz $x \notin A \cup B$. Stąd $x \notin A$ oraz $x \notin B$. Zatem $x \in -A$ oraz $x \in -B$, czyli $x \in -A \cap -B$.

$(\supseteq)$: Niech $x \in -A \cap -B$. Zatem $x \in -A$ oraz $x \in -B$, czyli $x \in U$ oraz $x \notin A$ i $x \notin B$. Stąd $x \notin A \cup B$, czyli $x \in -(A \cup B)$.

dla (2): $(\subseteq)$: Niech $x \in -(A \cap B)$. Wówczas $x \in U$ oraz $x \notin A \cap B$. Stąd $x \notin A \vee x \notin B$. Załóżmy, że $x \notin A$. Wtedy $x \in -A$, zaś $-A \subseteq -A \cup -B$, zatem $x \in -A \cup -B$. Załóżmy, że $x \notin B$. Wówczas $x \in -B$ lecz $-B \subseteq -A \cup -B$, więc $x \in -A \cup -B$.

($\supseteq$): Niech $x \in -A \cup -B$. Wówczas $x \in -A \vee x \in -B$. Załóżmy, że $x \in -A$. Wtedy $x \in U$ oraz $x \notin A$. Stąd $x \notin A \vee x \notin B$, czyli $x \notin A \cap B$. Zatem $x \in -(A \cap B)$. Załóżmy, że $x \in -B$. Wówczas analogicznie jak przed chwilą otrzymujemy: $x \in -(A \cap B)$.

dla (3): ($\subseteq$): Niech $x \in --A$. Zatem $x \in U$ oraz $x \notin -A$. Ostatnie wyrażenie jest równoważne formule $x \notin U \vee x \in A$. Skoro więc $x \in U$, to $x \in A$.

($\supseteq$): Niech $x \in A$. Ponieważ $A \subseteq U$, więc $x \in U$. Mamy: $x \notin -A$ (gdyby $x \in -A$, to $x \notin A$ wbrew założeniu). Ostatecznie $x \in U - (-A)$, tzn. $x \in --A$.

dla (4): ($\rightarrow$): Załóżmy, że $A \subseteq B$. Wówczas na mocy twierdzenia 13(3), $A \cup B = B$. Zatem $-(A \cup B) = -B$, co na mocy (1) implikuje: $-A \cap -B = -B$, więc również $-B \cap -A = -B$ i konsekwentnie na mocy twierdzenia 15(3), $-B \subseteq -A$.

($\leftarrow$): Załóżmy, że $-B \subseteq -A$. Wówczas na mocy przed chwilą dowiedzionej implikacji otrzymujemy: $--A \subseteq --B$, zatem z (3): $A \subseteq B$.

dla (5): Ponieważ $-A \subseteq U$ oraz $A \subseteq U$, więc na mocy twierdzenia 13(2), $-A \cup A \subseteq U$. Aby dowieść inkluzji ($\supseteq$) załóżmy, że $x \in U$. Naturalnie $x \in A \vee x \notin A$. Gdy $x \in A$, to oczywiście $x \in -A \cup A$ (bo $A \subseteq -A \cup A$), gdy zaś $x \notin A$, to $x \in -A$, a ponieważ $-A \subseteq -A \cup A$, więc $x \in -A \cup A$.

dla (6): Załóżmy, że $-A \cap A \neq \emptyset$. Niech $x \in -A \cap A$. Wówczas $x \in -A$ oraz $x \in A$. Jednakże skoro $x \in -A$, to $x \notin A$. Sprzeczność.

dla (7): Na podstawie (5), (1), (3) i (6) mamy: $-U = -(-A \cup A) = --A \cap -A = A \cap -A = \emptyset$ (bo naturalnie $A \cap -A = -A \cap A$).

dla (8): Z (7): $-U = \emptyset$, zatem $--U = -\emptyset$, więc z (3), $U = -\emptyset$. $\square$

2.5.8 Przekrój zbioru niepustego

Zastosujmy aksjomat podzbiorów w postaci:

$$(AxZ)_\phi \forall u[\forall x(\forall z(z \in u \rightarrow x \in z) \rightarrow x \in \bigcup u) \rightarrow \exists y \forall x(x \in y \leftrightarrow \forall z(z \in u \rightarrow x \in z))].$$

Rozważmy dowolny zbiór u . Naturalnie $u = \emptyset$ lub $u \neq \emptyset$. Gdy $u = \emptyset$, wyrażenie $\forall z(z \in u \rightarrow x \in z)$ (dla ustalonego dowolnego zbioru x) jest prawdziwe, bowiem poprzednik implikacji: $z \in u \rightarrow x \in z$ dla dowolnego zbioru z jest fałszywy. Tymczasem wyrażenie $x \in \bigcup u$ jest na mocy twierdzenia 11(4) fałszywe. Ostatecznie, poprzednik implikacji w nawiasach kwadratowych w $(AxZ)_\phi$ jest fałszywy. Nie możemy więc oderwać następnika.

Rozważmy jednakże przypadek: $u \neq \emptyset$. Wówczas ów poprzednik jest prawdziwy. Weźmy bowiem pod uwagę dowolny zbiór x oraz załóżmy, że $\forall z(z \in u \rightarrow x \in z)$. Skoro $u \neq \emptyset$, więc dla pewnego zbioru a , $a \in u$. Wówczas z założenia, $x \in a$, zatem $\exists z(z \in u \wedge x \in z)$, czyli $x \in \bigcup u$.

Ostatecznie, dla dowolnego niepustego zbioru u odrywamy następnik w $(AxZ)_\phi$, co prowadzi do definicji *iloczynu (uogólnionego)* lub *przekroju* $\bigcap u$ niepustego zbioru u , jako zbioru tych wszystkich zbiorów x , które należą do każdego elementu zbioru u :

DEFINICJA. Dla dowolnego niepustego zbioru u ,

$$\forall x(x \in \bigcap u \leftrightarrow \forall z(z \in u \rightarrow x \in z)) \text{ lub } \bigcap u = \{x : \forall z(z \in u \rightarrow x \in z)\}.$$

Przekrój niepustego zbioru ma analogiczne własności do podanych w twierdzeniu 11(1),(2) dla sumy zbioru:

TWIERDZENIE 18. Dla dowolnego niepustego zbioru u :

- (1) $\forall x(x \in u \rightarrow \bigcap u \subseteq x)$,
- (2) dla dowolnego y , $\forall x(x \in u \rightarrow y \subseteq x) \rightarrow y \subseteq \bigcap u$.

DOWÓD: dla (1): Załóżmy, że $x \in u$. Niech $a \in \bigcap u$. Wówczas z definicji przekroju, $\forall z(z \in u \rightarrow a \in z)$. Stąd i z założenia mamy: $a \in x$, co dowodzi inkluzji $\bigcap u \subseteq x$.

dla (2): Załóżmy, że y jest takim zbiorem, że $\forall x(x \in u \rightarrow y \subseteq x)$. Niech $a \in y$. Aby wykazać, że $a \in \bigcap u$ (co dowiedzie inkluzji $y \subseteq \bigcap u$) musimy dowieść, że $\forall z(z \in u \rightarrow a \in z)$. Niech więc $z \in u$. Wówczas z założenia otrzymujemy: $y \subseteq z$ i dalej, skoro $a \in y$, to $a \in z$, co dowodzi wyrażenia $\forall z(z \in u \rightarrow a \in z)$. $\square$

Przekrój pary zbiorów jest iloczynem jej elementów:

TWIERDZENIE 19. Dla dowolnych zbiorów x, y : $\bigcap \{x, y\} = x \cap y$.

DOWÓD: Rozważmy dowolny zbiór a . Mamy: $a \in \bigcap \{x, y\}$ wtw $\forall z(z \in \{x, y\} \rightarrow a \in z)$ wtw $\forall z((z = x \vee z = y) \rightarrow a \in z)$ wtw $(a \in x \wedge a \in y)$ wtw $a \in x \cap y$, co dowodzi, że $\forall v(v \in \bigcap \{x, y\} \leftrightarrow v \in x \cap y)$. Zatem z aksjomatu równości $(Ax =)$ w postaci $\forall v(v \in \bigcap \{x, y\} \leftrightarrow v \in x \cap y) \rightarrow \bigcap \{x, y\} = x \cap y$ uzyskujemy: $\bigcap \{x, y\} = x \cap y$. $\square$

2.6 Algebra Boole'a zbiorów

2.6.1 Ciało zbiorów

DEFINICJA. Zbiór $\mathcal{C} \subseteq P(U)$ nazywamy *ciałem zbiorów* (dokładniej *ciałem podzbiorów* zbioru U), gdy

- (1) $\mathcal{C} \neq \emptyset$,
- (2) $\forall X (X \in \mathcal{C} \rightarrow -X \in \mathcal{C})$,
- (3) $\forall X \forall Y ((X \in \mathcal{C} \wedge Y \in \mathcal{C}) \rightarrow X \cap Y \in \mathcal{C})$.

PRZYKŁAD. Zbiór potęgowy $P(U)$ jest ciałem zbiorów (ciałem wszystkich podzbiorów zbioru U), bowiem $P(U) \neq \emptyset$ i gdy $X \in P(U)$, to $-X \in P(U)$ oraz gdy $X, Y \in P(U)$, czyli $X, Y \subseteq U$, to $X \cap Y \subseteq U$, zatem $X \cap Y \in P(U)$.

Innym przykładem ciała podzbiorów zbioru U jest zbiór $\{\emptyset, U\}$. Oczywiście warunek (1) definicji ciała zbiorów jest dla $\mathcal{C} = \{\emptyset, U\}$ spełniony. Warunek (2) jest spełniony na mocy twierdzenia 17(7),(8), natomiast warunek (3) na mocy oczywistych faktów: $U \cap U = U$, $\emptyset \cap \emptyset = \emptyset$ oraz $\emptyset \cap U = U \cap \emptyset = \emptyset$.

TWIERDZENIE 20. Niech $\mathcal{C}$ będzie dowolnym ciałem podzbiorów zbioru U . Wówczas

- (1) $\forall X \forall Y ((X \in \mathcal{C} \wedge Y \in \mathcal{C}) \rightarrow X \cup Y \in \mathcal{C})$,
- (2) $\emptyset \in \mathcal{C}$ oraz $U \in \mathcal{C}$.

DOWÓD: Niech $\mathcal{C} \subseteq P(U)$ będzie ciałem podzbiorów zbioru U .

dla (1): Załóżmy, że $X, Y \in \mathcal{C}$. Wówczas na mocy warunku (2) definicji ciała zbiorów, $-X, -Y \in \mathcal{C}$. Stąd, na mocy warunku (3): $-X \cap -Y \in \mathcal{C}$. Wobec tego, na podstawie twierdzenia 17(1), $-(X \cup Y) \in \mathcal{C}$. Zatem znowu z warunku (2) definicji, $-(X \cup Y) \in \mathcal{C}$, co na mocy twierdzenia 17(3) daje: $X \cup Y \in \mathcal{C}$.

dla (2): Na mocy warunku (1) definicji ciała zbiorów, niech $X \in \mathcal{C}$. Wówczas na podstawie warunku (2) tej definicji, $-X \in \mathcal{C}$. Zatem na podstawie warunku (3): $-X \cap X \in \mathcal{C}$, czyli na mocy twierdzenia 17(6), $\emptyset \in \mathcal{C}$. Konsekwentnie, na mocy (2) definicji ciała zbiorów, otrzymujemy: $-\emptyset \in \mathcal{C}$. Zatem według twierdzenia 17(8), $U \in \mathcal{C}$. $\square$

WNIOSEK. Ciało $\{\emptyset, U\}$ jest najmniejszym (w sensie inkluzji) ciałem podzbiorów zbioru U .

DOWÓD: Na mocy twierdzenia 20(2) zbiór $\{\emptyset, U\}$ zawiera się w każdym ciecie $\mathcal{C}$ podzbiorów zbioru U . $\square$

PRZYKŁAD. Wszystkie ciała podzbiorów zbioru $U = \{a, b, c\}$:
 $\{\emptyset, U\}, \{\emptyset, \{a\}, \{b, c\}, U\}, \{\emptyset, \{b\}, \{a, c\}, U\}, \{\emptyset, \{c\}, \{a, b\}, U\}, P(U)$.

Ciało zbiorów jest przykładem zbioru *zamkniętego na operacje*: $\cap, \cup, -$. Mówimy bowiem, że zbiór jest *zamknięty* na daną n -argumentową operację, gdy jej wartość na dowolnej sekwencji (tzn. obiekt przyporządkowany sekwencji) n elementów tego zbioru jest również elementem tego zbioru. Gdy zbiór jest zamknięty na daną operację, to mówimy, iż jest to *operacja na tym zbiorze*.

Rozważmy obecnie ilość elementów największego (w sensie inkluzji) ciała podzbiorów danego n -elementowego zbioru U , czyli po prostu ilość wszystkich podzbiorów zbioru U (ilość elementów zbioru potęgowego n -elementowego zbioru). W tym celu sformułujmy pomocniczy fakt, wykorzystywany dalej w dowodzie twierdzenia ustalającego ową ilość. Pojawiające się tu zbiory $\mathcal{A}, \mathcal{B}$ istnieją na mocy aksjomatu podzbiorów.

LEMAT. Niech U będzie dowolnym niepustym zbiorem oraz $a \in U$. Niech ponadto $\mathcal{A} = \{X : X \subseteq U \wedge a \notin X\}, \mathcal{B} = \{X : X \subseteq U \wedge a \in X\}$. Wówczas $\mathcal{A} \cup \mathcal{B} = P(U)$ oraz $\mathcal{A} \cap \mathcal{B} = \emptyset$. Ponadto zbiory $\mathcal{A}, \mathcal{B}$ mają tyle samo elementów.

DOWÓD: Dowodzimy pierwszej równości. ($\subseteq$): Niech $X \in \mathcal{A} \cup \mathcal{B}$. Wówczas $X \in \mathcal{A} \vee X \in \mathcal{B}$. Gdy $X \in \mathcal{A}$, to $X \subseteq U$ oraz gdy $X \in \mathcal{B}$, to również $X \subseteq U$ z definicji zbiorów $\mathcal{A}, \mathcal{B}$, zatem ostatecznie $X \in P(U)$, czyli $\mathcal{A} \cup \mathcal{B} \subseteq P(U)$.

($\supseteq$): Załóżmy, że $X \in P(U)$. Wówczas $X \subseteq U$. Naturalnie $a \in X \vee a \notin X$. Gdy $a \in X$, to wówczas $X \in \mathcal{B}$, zatem $X \in \mathcal{A} \vee X \in \mathcal{B}$, czyli $X \in \mathcal{A} \cup \mathcal{B}$. Gdy zaś $a \notin X$, to $X \in \mathcal{A}$, czyli znowu $X \in \mathcal{A} \vee X \in \mathcal{B}$, skąd $X \in \mathcal{A} \cup \mathcal{B}$. Ostatecznie, $P(U) \subseteq \mathcal{A} \cup \mathcal{B}$.

Dowodzimy drugiej równości. Załóżmy, że $\mathcal{A} \cap \mathcal{B} \neq \emptyset$. Niech więc $X \in \mathcal{A} \cap \mathcal{B}$. Wówczas $X \in \mathcal{A}$ oraz $X \in \mathcal{B}$. Zatem $a \notin X$ oraz $a \in X$ z definicji zbiorów $\mathcal{A}, \mathcal{B}$. Sprzeczność.

Aby wykazać, że zbiory $\mathcal{A}, \mathcal{B}$ mają tę samą ilość elementów, weźmy pod uwagę dowolny $X \in \mathcal{A}$. Wówczas $X \subseteq U$ oraz $a \notin X$. Zatem $X \cup \{a\} \in \mathcal{B}$ (bo $X \cup \{a\} \subseteq U$ oraz $a \in X \cup \{a\}$).

Rozważmy przyporządkowanie każdemu elementowi X ze zbioru $\mathcal{A}$ elementu $X \cup \{a\}$ ze zbioru $\mathcal{B}$. Przyporządkowanie to ma dwie następujące własności:

(w1) dla dowolnych $X_1, X_2 \in \mathcal{A}$, $(X_1 \neq X_2 \rightarrow X_1 \cup \{a\} \neq X_2 \cup \{a\})$,

(w2) dla dowolnego $Y \in \mathcal{B}$ istnieje $X \in \mathcal{A}$ taki, że $Y = X \cup \{a\}$.

Aby dowieść (w1) założymy: $X_1 \cup \{a\} = X_2 \cup \{a\}$ dla jakichś $X_1, X_2 \in \mathcal{A}$ i wykażemy równość $X_1 = X_2$.

($\subseteq$): Niech $x \in X_1$. Wówczas $x \in X_1 \cup \{a\}$. Zatem z założenia, $x \in X_2 \cup \{a\}$ czyli $x \in X_2 \vee x = a$. Lecz $x \neq a$ (gdyby bowiem $x = a$, to $x \notin X_1$, bo $a \notin X_1$ skoro $X_1 \in \mathcal{A}$). Zatem $x \in X_2$. Analogicznie dowodzi się inkluzji ($\supseteq$).

Aby dowieść (w2) rozważmy dowolny $Y \in \mathcal{B}$. Wówczas $Y \subseteq U$ oraz $a \in Y$. Ponieważ $Y - \{a\} \subseteq Y$, więc $Y - \{a\} \subseteq U$. Oczywiście $a \notin Y - \{a\}$. Zatem $Y - \{a\} \in \mathcal{A}$. Ponadto $(Y - \{a\}) \cup \{a\} = Y$. Mamy bowiem: $Y - \{a\} \subseteq Y$ oraz $\{a\} \subseteq Y$ (bo $a \in Y$). Stąd na mocy Tw.13(2), $(Y - \{a\}) \cup \{a\} \subseteq Y$. Aby dowieść odwrotnej inkluzji założymy, że $x \in Y$. Naturalnie $x = a \vee x \neq a$. Jeśli $x = a$, to $x \in \{a\}$, zaś $\{a\} \subseteq (Y - \{a\}) \cup \{a\}$, czyli $x \in (Y - \{a\}) \cup \{a\}$. Jeśli $x \neq a$, to $x \notin \{a\}$, zatem $x \in Y - \{a\}$, a ponieważ $Y - \{a\} \subseteq (Y - \{a\}) \cup \{a\}$, więc $x \in (Y - \{a\}) \cup \{a\}$.

Ostatecznie dla wybranego dowolnie $Y \in \mathcal{B}$ znaleźliśmy $X \in \mathcal{A}$ taki, że $Y = X \cup \{a\}$ (ów $X = Y - \{a\}$).

Jak widać, fakt, że zbiory $\mathcal{A}, \mathcal{B}$ mają tyle samo elementów, jest bezpośrednim wnioskiem z (w1) i (w2). Własność (w1) bowiem mówi, że w zbiorze $\mathcal{B}$ nie ma mniej elementów niż w zbiorze $\mathcal{A}$, zaś (w2) mówi, że w $\mathcal{B}$ nie ma więcej elementów niż w $\mathcal{A}$. $\square$

TWIERDZENIE 21. *Dla dowolnej liczby naturalnej n , dla dowolnego n -elementowego zbioru U , zbiór potęgowy $P(U)$ jest 2^n -elementowy.*

DOWÓD: (indukcyjny) Dla $n = 0$, $U = \emptyset$. Wówczas $P(U) = \{\emptyset\}$ (twierdzenie 9) czyli $P(U)$ jest 2^0 -elementowy. Ustalmy dowolne $n \geq 0$.

Założenie indukcyjne: dla dowolnego n -elementowego zbioru U , $P(U)$ jest 2^n -elementowy.

Teza indukcyjna (mamy ją dowieść): dla dowolnego $(n+1)$ -elementowego zbioru U , zbiór $P(U)$ jest 2^{n+1} -elementowy.

Rozważmy $(n+1)$ -elementowy zbiór U . Wówczas $U \neq \emptyset$. Niech $a \in U$ będzie dowolnie wybranym elementem. Oznaczmy:

$$\mathcal{A} = \{X : X \subseteq U \wedge a \notin X\}, \quad \mathcal{B} = \{X : X \subseteq U \wedge a \in X\}.$$

Na mocy lematu mamy natychmiast:

(1) $\mathcal{A} \cup \mathcal{B} = P(U)$ oraz

(2) $\mathcal{A} \cap \mathcal{B} = \emptyset$.

Wykażemy:

$$(3) \mathcal{A} = P(U - \{a\}).$$

Dowód (3) opieramy na następującej równoważności:

$$(4) (X \subseteq U \wedge a \notin X) \text{ wtw } \forall x(x \in X \rightarrow x \in U - \{a\}), \text{ dla dowolnego } X.$$

Aby dowieść (4) $(\rightarrow)$ załóżmy, że $X \subseteq U$ oraz $a \notin X$. Weźmy dowolny x . Niech $x \in X$. Wówczas $x \in U$. Ponadto $x \neq a$ (gdyby $x = a$, to $a \in X$ wbrew założeniu) czyli $x \notin \{a\}$. Ostatecznie $x \in U - \{a\}$.

(4) $(\leftarrow)$: Na odwrót, załóżmy, że $\forall x(x \in X \rightarrow x \in U - \{a\})$. Aby wykazać inkluzję $X \subseteq U$, weźmy $x \in X$. Wówczas z założenia, $x \in U - \{a\}$, czyli $x \in U \wedge x \notin \{a\}$, skąd otrzymujemy, że $x \in U$. Aby wykazać, iż $a \notin X$, przypuśćmy, że $a \in X$. Wówczas z założenia: $a \in U - \{a\}$, co oznacza, że $a \in U \wedge a \notin \{a\}$. W szczególności więc $a \notin \{a\}$, czyli $a \neq a$, co jest niemożliwe.

Wykazujemy teraz (3):

$(\subseteq)$: Niech $X \in \mathcal{A}$. Wówczas $X \subseteq U \wedge a \notin X$. Zatem na mocy (4), $X \subseteq U - \{a\}$, czyli $X \in P(U - \{a\})$.

$(\supseteq)$: Na odwrót, załóżmy, że $X \in P(U - \{a\})$, tzn. $X \subseteq U - \{a\}$. Wówczas z (4) mamy natychmiast: $X \subseteq U \wedge a \notin X$, a zatem $X \in \mathcal{A}$.

Bezpośrednio na mocy lematu mamy:

(5) Zbiory $\mathcal{A}, \mathcal{B}$ mają tyle samo elementów.

Ponieważ zbiór $U - \{a\}$ jest n -elementowy, więc stosując do tego zbioru założenie indukcyjne, otrzymujemy: zbiór $P(U - \{a\})$ jest 2^n -elementowy. Zatem, na mocy (3), zbiór $\mathcal{A}$ jest 2^n -elementowy. Wobec tego, zgodnie z (5), zbiór $\mathcal{B}$ jest również 2^n -elementowy. Ostatecznie, w oparciu o (1) i (2) stwierdzamy, że zbiór $P(U)$ ma $2^n + 2^n$ elementów (bo ilość elementów sumy skończonych rozłącznych zbiorów jest równa sumie ilości elementów tych zbiorów), tzn. zbiór $P(U)$ jest 2^{n+1} -elementowy. $\square$

2.6.2 Algebra Boole'a

DEFINICJA. Dowolny zbiór A wraz z następującymi operacjami na nim: 2-argumentowymi $\wedge, \vee$, 1-argumentową $-$, oraz wyróżnionymi elementami 1, 0, spełniającymi równości:

$$\begin{array}{ll} (1) & x \wedge x = x, & x \vee x = x, \\ (2) & x \wedge y = y \wedge x, & x \vee y = y \vee x, \\ (3) & x \wedge (y \wedge z) = (x \wedge y) \wedge z, & x \vee (y \vee z) = (x \vee y) \vee z, \\ (4) & x \wedge (x \vee y) = x, & x \vee (x \wedge y) = x, \\ (5) & x \wedge (y \vee z) = (x \wedge y) \vee (x \wedge z), & x \vee (y \wedge z) = (x \vee y) \wedge (x \vee z), \end{array}$$

$$(6) \quad x \wedge 1 = x,$$

$$x \vee 0 = x,$$

$$(7) \quad -x \wedge x = 0,$$

$$-x \vee x = 1,$$

nazywamy *algebrą Boole'a*.

PRZYKŁAD. Układ $(\{1, 0\}, \wedge, \vee, -, 1, 0)$, gdzie $1, 0$ są wartościami logicznymi odpowiednio prawdy i fałszu oraz $\wedge, \vee, -$ są operacjami na tych wartościach logicznych odpowiadającymi warunkom prawdziwości w logice klasycznej dla spójników odpowiednio koniunkcji, alternatywy i negacji, tzn. dla dowolnych $x, y \in \{1, 0\}$:

$x \wedge y = 1$, gdy $x = y = 1$, w przeciwnym wypadku $x \wedge y = 0$,

$x \vee y = 0$, gdy $x = y = 0$, w przeciwnym wypadku $x \vee y = 1$,

$-1 = 0, -0 = 1$,

jest 2-elementową algebrą Boole'a.

TWIERDZENIE 22. *Niech $\mathcal{C} \subseteq P(U)$ będzie ciałem podzbiorów zbioru U . Wówczas $(\mathcal{C}, \cap, \cup, -, U, \emptyset)$, gdzie $\cap, \cup$ są odpowiednio operacjami iloczynu i sumy dwóch zbiorów oraz $-$ jest operacją dopełnienia (do zbioru U) jest algebrą Boole'a.*

DOWÓD: Rzeczywiście, z definicji ciała zbiorów, $\cap$ oraz $-$ są odpowiednio 2- i 1-argumentowymi operacjami na zbiorze $\mathcal{C}$. Na mocy twierdzenia 20(1) również $\cup$ jest operacją na zbiorze $\mathcal{C}$, zaś zgodnie z twierdzeniem 20(2), U oraz $\emptyset$ są elementami zbioru $\mathcal{C}$. Aby więc dowieść, że układ $(\mathcal{C}, \cap, \cup, -, U, \emptyset)$ jest algebrą Boole'a należy wykazać, że równości (1) - (7) z definicji algebry Boole'a, zapisane dla operacji $\cap, \cup, -$ oraz elementów wyróżnionych $U, \emptyset$, są spełnione.

Równości (1), (2) oraz (3) są w oczywisty sposób spełnione.

Weźmy dowolne $X, Y, Z \in \mathcal{C}$. Dla równości (4): $X \cap (X \cup Y) \subseteq X$ na mocy twierdzenia 15(1). Ponieważ $X \subseteq X$ oraz $X \subseteq X \cup Y$ (twierdzenia 1 oraz 13(1)), więc według twierdzenia 15(2), $X \subseteq X \cap (X \cup Y)$. Ostatecznie, $X \cap (X \cup Y) = X$. Analogicznie dowodzimy drugą z równości (4).

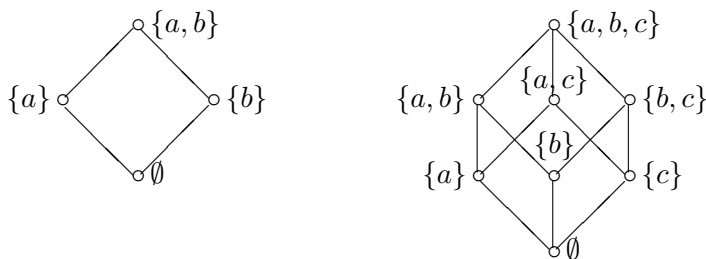
Dla (5): $X \cap Y \subseteq Y$ (twierdzenie 15(1)) oraz $Y \subseteq Y \cup Z$, (twierdzenie 13(1)) zatem na mocy twierdzenia 3, $X \cap Y \subseteq Y \cup Z$. Ponadto $X \cap Y \subseteq X$. Zatem, według twierdzenia 15(2), $X \cap Y \subseteq X \cap (Y \cup Z)$. Analogicznie wykazuje się, iż $X \cap Z \subseteq X \cap (Y \cup Z)$. Z dwóch ostatnich wyrażeń wnosimy, na mocy twierdzenia 13(2), iż $(X \cap Y) \cup (X \cap Z) \subseteq X \cap (Y \cup Z)$. Pozostaje

dowieść odwrotnej inkluzji. Niech $x \in X \cap (Y \cup Z)$. Wówczas $x \in X$ oraz $x \in Y \cup Z$. Stąd $x \in Y$ lub $x \in Z$. Niech zatem $x \in Y$. Wówczas $x \in X \cap Y$, stąd $x \in (X \cap Y) \cup (X \cap Z)$. Gdy zaś $x \in Z$, to $x \in X \cap Z$, zatem również $x \in (X \cap Y) \cup (X \cap Z)$. Ostatecznie $X \cap (Y \cup Z) \subseteq (X \cap Y) \cup (X \cap Z)$. Analogicznie dowodzimy drugą z równości (5).

Dla (6): Równości: $X \cap U = X$ oraz $X \cup \emptyset = X$ wynikają odpowiednio z twierdzenia 15(3) (bo oczywiście $X \subseteq U$), oraz równości (2) (przemienność sumy), twierdzenia 13(3) i twierdzenia 5.

Równości (7) to twierdzenie 17(6),(5). $\square$

PRZYKŁAD. Zbiór potęgowy dowolnego zbioru wraz z odpowiednimi operacjami jest algebrą Boole'a. Ilustracje algebr Boole'a wszystkich podzbiorów zbioru $\{a, b\}$ oraz zbioru $\{a, b, c\}$ są następujące:



Linia łamana prowadząca ku górze (na każdym swoim odcinku) od jednego do drugiego zbioru, oznacza, że zbiór położony niżej jest podzbiorem zbioru będącego wyżej.

O algebrach Boole'a dowodzi się twierdzenia w pewnym sensie odwrotnego do twierdzenia 22: każda algebra Boole'a ma „strukturę algebraiczną” pewnego ciała zbiorów (zob. np. Traczyk[24]).

2.7 Relacje i funkcje

2.7.1 Para uporządkowana. Produkt kartezjański dwóch zbiorów

Dla pary zbiorów $\{x, y\}$ zachodzi, jak łatwo sprawdzić, równość $\{x, y\} = \{y, x\}$. To znaczy, kolejność wymienienia elementów pary zbiorów jest dla

jej określenia nieistotna. Wprowadzimy pojęcie *uporządkowanej pary zbiorów*: $\langle x, y \rangle$ takiej, że $\langle x, y \rangle \neq \langle y, x \rangle$ o ile $x \neq y$.

DEFINICJA. Dla dowolnych zbiorów x, y , $\langle x, y \rangle = \{\{x\}, \{x, y\}\}$.

TWIERDZENIE 23. $\langle x, y \rangle = \langle u, v \rangle \rightarrow (x = u \wedge y = v)$.

DOWÓD: Załóżmy, że $\langle x, y \rangle = \langle u, v \rangle$. Wówczas z definicji pary uporządkowanej mamy:

$$(1) \quad \{\{x\}, \{x, y\}\} = \{\{u\}, \{u, v\}\}.$$

Oczywiście $x = y$ lub $x \neq y$. Gdy $x = y$, to $\{\{x\}, \{x, y\}\} = \{\{x\}\}$. Zatem z (1) mamy: $\{\{x\}\} = \{\{u\}, \{u, v\}\}$, co oznacza, że

$$(2) \quad \{x\} = \{u\} \text{ oraz}$$

$$(3) \quad \{x\} = \{u, v\}.$$

Z (2) otrzymujemy: $x = u$, zaś z (3): $x = u = v$, zatem skoro $y = x$, więc $y = v$.

Założmy, że $x \neq y$. Wówczas

$$(4) \quad \{x\} \neq \{x, y\}.$$

Wtedy na podstawie (1):

$$(\{x\} = \{u\} \text{ lub } \{x\} = \{u, v\}) \text{ oraz } (\{x, y\} = \{u\} \text{ lub } \{x, y\} = \{u, v\}).$$

Gdy $\{x\} = \{u\}$, to $\{x, y\} \neq \{u\}$ na mocy (4), zatem $\{x, y\} = \{u, v\}$. Gdy zaś $\{x\} = \{u, v\}$, to $\{x, y\} \neq \{u, v\}$, czyli $\{x, y\} = \{u\}$. Mamy zatem następujące dwa przypadki:

$$(5) \quad \{x\} = \{u\} \text{ i } \{x, y\} = \{u, v\},$$

$$(6) \quad \{x\} = \{u, v\} \text{ i } \{x, y\} = \{u\}.$$

Jednakże przypadek (6) nie może zachodzić, bowiem gdyby zachodził, to $x = u$ oraz $y = u$, czyli byłoby $x = y$. Zatem zachodzi przypadek (5), co oznacza, że $x = u$, a ponadto skoro $y \neq x$ czyli $y \neq u$, więc $y = v$. $\square$

WNIOSEK. $\langle x, y \rangle = \langle y, x \rangle \rightarrow x = y$.

DOWÓD: Oczywisty na podstawie twierdzenia 23. $\square$

Wykorzystamy aksjomat podzbiorów, aby dowieść istnienia zbioru wszystkich par uporządkowanych, z których pierwszy element należy do jednego danego zbioru, zaś drugi element pary do drugiego danego zbioru. W tym celu wykazujemy fakt pomocniczy:

LEMAT. $(u \in a \wedge v \in b) \rightarrow \langle u, v \rangle \subseteq P(a \cup b)$.

DOWÓD: Załóżmy, że $u \in a$ oraz $v \in b$. Wówczas $\{u\} \subseteq a \subseteq a \cup b$. Ponadto, $\{u, v\} \subseteq a \cup b$. Zatem $\{u\}, \{u, v\} \in P(a \cup b)$, czyli $\{\{u\}, \{u, v\}\} \subseteq P(a \cup b)$, a więc z definicji pary uporządkowanej: $\langle u, v \rangle \subseteq P(a \cup b)$. $\square$

Rozważmy w aksjomacie podzbiorów $(AxZ)_\phi$ formułę $\phi(x)$ postaci $\exists u \exists v (u \in a \wedge v \in b \wedge x = \langle u, v \rangle)$ oraz zbiór z postaci $P(P(a \cup b))$, otrzymując wyrażenie

$$\forall a \forall b [\forall x (\exists u \exists v (u \in a \wedge v \in b \wedge x = \langle u, v \rangle) \rightarrow x \in P(P(a \cup b))) \rightarrow \exists y \forall x (x \in y \leftrightarrow \exists u \exists v (u \in a \wedge v \in b \wedge x = \langle u, v \rangle))],$$

które, wobec prawdziwości poprzednika głównej implikacji (na mocy lematu) prowadzi do definicji *produktu kartezjańskiego* $a \times b$ dla dowolnych zbiorów a, b :

DEFINICJA. Dla dowolnych zbiorów a, b ,

$\forall x (x \in a \times b \leftrightarrow \exists u \exists v (u \in a \wedge v \in b \wedge x = \langle u, v \rangle))$ lub

$a \times b = \{x : \exists u \exists v (u \in a \wedge v \in b \wedge x = \langle u, v \rangle)\} = \{\langle u, v \rangle : u \in a \wedge v \in b\}$.

Produkt kartezjański $a \times b$ zbiorów a, b jest zbiorem wszystkich par uporządkowanych, których pierwszy element należy do zbioru a , zaś drugi do zbioru b .

Produkt $a \times a$ bywa oznaczany jako a^2 .

Uwaga. Ostatnia z równości w definicji produktu została uzyskana na podstawie konwencji notacyjnej, w myśl której sekwencja symboli $\{F(x_1, \dots, x_n) : \psi(x_1, \dots, x_n)\}$, gdzie F jest symbolem n -argumentowej operacji, zaś $x_1, \dots, x_n$ zmiennymi wolnymi w formule $\psi(x_1, \dots, x_n)$, oznacza zbiór wszystkich wartości $F(x_1, \dots, x_n)$ operacji F na tej sekwencji zbiorów $x_1, \dots, x_n$, dla której prawdziwe jest $\psi(x_1, \dots, x_n)$.

Używać również będziemy konwencji notacyjnej polegającej na pisaniu $\{x \in A : \psi(x)\}$ na oznaczenie zbioru $\{x : x \in A \wedge \psi(x)\}$.

Ponadto używać będziemy kwantyfikatorów o ograniczonym zakresie: $\forall \alpha(x)(\beta) =_{\text{def}} \forall x (\alpha(x) \rightarrow \beta)$ oraz $\exists \alpha(x)(\beta) =_{\text{def}} \exists x (\alpha(x) \wedge \beta)$, gdzie $\alpha(x)$ jest formułą, w której x jest zmienną wolną oraz β jest formułą. Na przykład będziemy pisać $\forall x \in A (x \in B)$ lub po prostu $\forall x \in A, x \in B$ zamiast pisać $\forall x (x \in A \rightarrow x \in B)$. Podobnie będziemy pisać $\exists x \in A (x \in B)$ lub $\exists x \in A, x \in B$ zamiast $\exists x (x \in A \wedge x \in B)$.

Każdą niepustą relację binarną określoną na zbiorze A można teraz uwi-
docznąć przez podkreślenie par do niej należących. Nazwijmy ekspozycję
podkreślonych par uporządkowanych – diagramem relacji R .

Jest jasne, że relacja R jest zwrotna, gdy przekątna macierzy o końcach
 $\langle a_1, a_1 \rangle, \langle a_n, a_n \rangle$ zawiera się w diagramie R , natomiast R jest prze-
ciwzwrotna, gdy owa przekątna nie ma wspólnych par z diagramem relacji
 R .

Ponadto, R jest symetryczna, gdy, ilekroć jakaś para z macierzy jest
w jej diagramie, tylekroć „zwierciadlane odbicie” tej pary względem przekąt-
nej (tzn. wartość symetrii osiowej względem przekątnej na tej parze) również
należy do diagramu R .

R jest przeciwsymetryczna, gdy, ilekroć jakaś para z macierzy jest
w diagramie R , tylekroć „zwierciadlane odbicie” tej pary względem prze-
kątnej nie jest w tym diagramie. W szczególności więc, żadna para z prze-
kątnej nie jest w diagramie takiej relacji (tzn. relacja przeciwsymetryczna
jest przeciwwzrotna), skoro „zwierciadlanym odbiciem” pary z przekątnej
jest ona sama.

Relacja R jest antysymetryczna, gdy dla par spoza przekątnej zacho-
wuje się ona jak relacja przeciwsymetryczna, tzn., ilekroć jakaś para spoza
przekątnej jest w jej diagramie, tylekroć „zwierciadlane odbicie” tej pary
nie znajduje się w diagramie; jednakże, w przeciwieństwie do relacji prze-
ciwsymetrycznej, jakiejkolwiek pary z przekątnej mogą należeć dla relacji
antysymetrycznej.

Relacja R jest spójna, gdy dla dowolnej pary uporządkowanej poza prze-
kątą, para ta lub jej „zwierciadlane odbicie” względem przekątnej są
w diagramie relacji R .

Jest widoczne, że niektóre z wymienionych wyżej własności relacji nie
są niezależne od siebie. Przykładowo, związki między przeciwwzrotnością
a przeciwsymetrią są postaci:

TWIERDZENIE 24. *Niech R będzie relacją binarną określoną na zbiorze A .*

- (1) *Jeżeli R jest przeciwsymetryczna, to R jest przeciwwzrotna.*
- (2) *Jeżeli R jest przeciwwzrotna i przechodnia, to R jest przeciwsymetryczna.*

DOWÓD: dla (1): Załóżmy, że $\forall x, y \in A (\langle x, y \rangle \in R \rightarrow \langle y, x \rangle \notin R)$ oraz nie
wprost, że R nie jest przeciwwzrotna, tzn. dla pewnego $x \in A$, $\langle x, x \rangle \in R$.
Wówczas z warunku przeciwsymetrii mamy: $\langle x, x \rangle \in R \rightarrow \langle x, x \rangle \notin R$.
Stąd $\langle x, x \rangle \notin R$. Sprzeczność.

dla (2): Załóżmy, że $\forall x \in A, \langle x, x \rangle \notin R$ oraz $\forall x, y, z \in A ((\langle x, y \rangle \in R \wedge \langle y, z \rangle \in R) \rightarrow \langle x, z \rangle \in R)$ i nie wprost, że R nie jest przeciwsymetryczna, tzn. dla pewnych $x, y \in A, \langle x, y \rangle \in R$ oraz $\langle y, x \rangle \notin R$. Z przechodniości mamy: $(\langle x, y \rangle \in R \wedge \langle y, x \rangle \in R) \rightarrow \langle x, x \rangle \in R$. Stąd $\langle x, x \rangle \in R$. Jednakże z przeciwwrotności, $\langle x, x \rangle \notin R$. Sprzeczność. $\square$

DEFINICJA. Niech $R \subseteq A \times B$. Przez *konwers relacji R* (lub *relację odwrotną do R*) rozumiemy relację $R^\sim \subseteq B \times A$ określoną następująco: $R^\sim = \{\langle x, y \rangle \in B \times A : \langle y, x \rangle \in R\}$, tzn. $\forall x \in B \forall y \in A, \langle x, y \rangle \in R^\sim$ wtw $\langle y, x \rangle \in R$.

Niech ponadto $S \subseteq B \times C$. Przez *złożenie* (lub *superpozycję*) relacji R, S rozumiemy relację $R \circ S \subseteq A \times C$ określoną następująco:

$R \circ S = \{\langle x, y \rangle \in A \times C : \exists z \in B (\langle x, z \rangle \in R \wedge \langle z, y \rangle \in S)\}$, tzn. $\forall x \in A \forall y \in C, \langle x, y \rangle \in R \circ S$ wtw $\exists z \in B (\langle x, z \rangle \in R \wedge \langle z, y \rangle \in S)$.

Zbiór $\text{id}_A = \{\langle x, x \rangle : x \in A\} = \{\langle x, y \rangle \in A^2 : x = y\}$ nazywamy relacją *identycznościową* lub *tożsamościową na zbiorze A* . Inaczej, $\forall x, y \in A, \langle x, y \rangle \in \text{id}_A$ wtw $x = y$.

Uwaga. Nietrudno zastosować aksjomat podzbiorów dla dowodów istnienia konwersu oraz złożenia relacji jak również relacji identycznościowej.

TWIERDZENIE 25. Niech $R \subseteq A \times B, S \subseteq B \times C, T \subseteq C \times D$. Wówczas

- (1) $R \circ (S \circ T) = (R \circ S) \circ T$,
- (2) $R \circ \text{id}_B = \text{id}_A \circ R = R$,
- (3) $(R \circ S)^\sim = S^\sim \circ R^\sim$.

DOWÓD: dla (1): Naturalnie $S \circ T \subseteq B \times D$, zaś $R \circ S \subseteq A \times C$. Zatem $R \circ (S \circ T) \subseteq A \times D$ oraz $(R \circ S) \circ T \subseteq A \times D$.

($\subseteq$): Niech dla $a \in A$ i $d \in D, \langle a, d \rangle \in R \circ (S \circ T)$. Wówczas dla pewnego $b \in B$:

- (4) $\langle a, b \rangle \in R$ oraz
- (5) $\langle b, d \rangle \in S \circ T$.

Z (5) dla pewnego $c \in C$:

- (6) $\langle b, c \rangle \in S$ oraz
- (7) $\langle c, d \rangle \in T$.

Zatem z (4) i (6): $\langle a, c \rangle \in R \circ S$, stąd i z (7): $\langle a, d \rangle \in (R \circ S) \circ T$.

Dla inkluzji ($\supseteq$) dowód jest analogiczny.

dla (2): Wykazujemy, że $R \circ \text{id}_B = R$. Niech $a \in A$ oraz $b \in B$.

($\subseteq$): Załóżmy, że $\langle a, b \rangle \in R \circ \text{id}_B$. Wówczas dla pewnego $x \in B, \langle a, x \rangle \in R$ oraz $\langle x, b \rangle \in \text{id}_B$. Stąd $x = b$, czyli $\langle a, b \rangle \in R$.

($\supseteq$): Niech $\langle a, b \rangle \in R$. Ponieważ $\langle b, b \rangle \in \text{id}_B$, więc z definicji złożenia relacji mamy: $\langle a, b \rangle \in R \circ \text{id}_B$.

Analogicznie wykazujemy, że $\text{id}_A \circ R = R$.

dla (3): Naturalnie $(R \circ S)^\sim \subseteq C \times A$ (bo $R \circ S \subseteq A \times C$) oraz $S^\sim \circ R^\sim \subseteq C \times A$ (bo $S^\sim \subseteq C \times B$ oraz $R^\sim \subseteq B \times A$).

($\subseteq$): Niech $c \in C$ oraz $a \in A$. Załóżmy, że $\langle c, a \rangle \in (R \circ S)^\sim$. Wówczas $\langle a, c \rangle \in R \circ S$, zatem dla pewnego $b \in B$, $\langle a, b \rangle \in R$ i $\langle b, c \rangle \in S$. Czyli $\langle c, b \rangle \in S^\sim$ oraz $\langle b, a \rangle \in R^\sim$. Ostatecznie, $\langle c, a \rangle \in S^\sim \circ R^\sim$.

($\supseteq$): Rozumowanie odwrotne do powyższego. $\square$

TWIERDZENIE 26. Niech $R_1, R_2 \subseteq A \times B$ oraz $S \subseteq B \times C$. Wówczas

- (1) $R_1 \subseteq R_2 \rightarrow R_1 \circ S \subseteq R_2 \circ S$,
- (2) $(R_1 \cap R_2)^\sim = R_1^\sim \cap R_2^\sim$,
- (3) $(R_1 \cup R_2)^\sim = R_1^\sim \cup R_2^\sim$.

DOWÓD: dla (1): Załóżmy, że $R_1 \subseteq R_2$. Niech dla $a \in A$ oraz $c \in C$, $\langle a, c \rangle \in R_1 \circ S$. Wówczas dla pewnego $b \in B$, $\langle a, b \rangle \in R_1$ oraz $\langle b, c \rangle \in S$. Z założenia zatem otrzymujemy: $\langle a, b \rangle \in R_2$ i ostatecznie, $\langle a, c \rangle \in R_2 \circ S$.

dla (2): Mamy: $\langle b, a \rangle \in (R_1 \cap R_2)^\sim$ wtw $\langle a, b \rangle \in R_1 \cap R_2$ wtw $\langle a, b \rangle \in R_1$ oraz $\langle a, b \rangle \in R_2$ wtw $\langle b, a \rangle \in R_1^\sim$ oraz $\langle b, a \rangle \in R_2^\sim$ wtw $\langle b, a \rangle \in R_1^\sim \cap R_2^\sim$, co na mocy aksjomatu identyczności implikuje równość $(R_1 \cap R_2)^\sim = R_1^\sim \cap R_2^\sim$.

dla (3): Analogicznie jak dla (2). $\square$

Nietrudno zauważyć, że wartości operacji złożenia i konwersu na relacjach binarnych określonych na danym zbiorze A są relacjami na tym zbiorze. Innymi słowy, zbiór wszystkich relacji określonych na zbiorze A , czyli zbiór $P(A \times A)$ jest zamknięty na te dwie operacje. Jako ciało zbiorów, jest on również zamknięty na operacje: $\cap, \cup, -$. Zbiór potęgowy $P(A \times A)$ wyposażony w operacje: $\cap, \cup, -, \circ, \sim$ z wyróżnionymi elementami: $\emptyset, A \times A, \text{id}_A$ nazywamy *algebrą relacji* lub *algebrą relacyjną*.

Operacje na relacjach mogą posłużyć do wyrażenia własności formalnych relacji określonych na danym zbiorze. Przykładowo:

TWIERDZENIE 27. Dla dowolnej relacji $R \subseteq A^2$:

- (1) R jest zwrotna na A wtw $\text{id}_A \subseteq R$,
- (2) R jest symetryczna na A wtw $R^\sim = R$
- (3) R jest przechodnia na A wtw $R \circ R \subseteq R$.

DOWÓD: dla (1): Mamy oczywiste równoważności:

R jest zwrotna wtw $\forall x \in A, \langle x, x \rangle \in R$ wtw $\text{id}_A \subseteq R$.

dla (2): $(\rightarrow)$: Załóżmy, że $\forall x, y \in A (\langle x, y \rangle \in R \rightarrow \langle y, x \rangle \in R)$. Aby wykazać równość $R = R^\sim$ dowodzimy inkluzji $(\subseteq)$: niech $\langle x, y \rangle \in R$. Wówczas na mocy symetrii: $\langle y, x \rangle \in R$, zatem z definicji konwersu, $\langle x, y \rangle \in R^\sim$. $(\supseteq)$: Niech $\langle x, y \rangle \in R^\sim$. Wówczas $\langle y, x \rangle \in R$, zatem z warunku symetrii, $\langle x, y \rangle \in R$.

$(\Leftarrow)$: Załóżmy, że $R = R^\sim$. Aby dowieść warunek symetrii weźmy $x, y \in A$ i załóżmy, że $\langle x, y \rangle \in R$. Wówczas z założenia, $\langle x, y \rangle \in R^\sim$ skąd $\langle y, x \rangle \in R$, zatem R jest symetryczna na A .

dla (3): $(\rightarrow)$: Załóżmy, że $\forall x, y, z \in A ((\langle x, y \rangle \in R \wedge \langle y, z \rangle \in R) \rightarrow \langle x, z \rangle \in R)$. Aby wykazać inkluzję $R \circ R \subseteq R$ załóżmy, że dla jakichś $x, y \in A, \langle x, y \rangle \in R \circ R$. Wówczas z definicji złożenia, $\langle x, a \rangle \in R$ oraz $\langle a, y \rangle \in R$ dla pewnego $a \in A$. Zatem z przechodniości otrzymujemy: $(\langle x, a \rangle \in R \wedge \langle a, y \rangle \in R) \rightarrow \langle x, y \rangle \in R$, skąd mamy natychmiast: $\langle x, y \rangle \in R$, co kończy dowód inkluzji $R \circ R \subseteq R$.

$(\Leftarrow)$: Załóżmy, że $R \circ R \subseteq R$. Aby wykazać warunek przechodniości rozważmy $x, y, z \in A$ i załóżmy, że $\langle x, y \rangle \in R$ oraz $\langle y, z \rangle \in R$. Wówczas $\langle x, z \rangle \in R \circ R$, zatem z założenia, $\langle x, z \rangle \in R$, co dowodzi przechodniości relacji R na zbiorze A . $\square$

Uwaga. Dla pozostałych formalnych własności relacji binarnych określonych na zbiorze A tzn. dla przeciwzwrotności, antysymetrii, przeciwsymetrii oraz spójności można równie łatwo podać ich odpowiedniki w terminach operacji na relacjach.

2.7.3 Funkcje

DEFINICJA. Relację binarną f na zbiorach A, B nazywamy *funkcją przekształcającą zbiór A w B* (co zapisujemy $f : A \rightarrow B$) gdy

- (1) $\forall x \in A \exists y \in B (\langle x, y \rangle \in f)$,
- (2) $\forall x \in A \forall y, z \in B ((\langle x, y \rangle \in f \wedge \langle x, z \rangle \in f) \rightarrow y = z)$.

Gdy relacja $f \subseteq A \times B$ jest funkcją przekształcającą zbiór A w B , to zapis $\langle x, y \rangle \in f$ interpretujemy następująco: element y ze zbioru B *jest przyporządkowany według funkcji f elementowi x ze zbioru A* . Wówczas warunek (1) definicji funkcji, równoważny skądinąd wyrażeniu $D(f) = A$, czytamy następująco: każdemu elementowi ze zbioru A jest przyporządko-

wany według f jakiś element ze zbioru B . Natomiast warunek (2) mówi, iż co najwyżej jeden element ze zbioru B może być przyporządkowany danemu elementowi ze zbioru A . Ostatecznie, biorąc pod uwagę interpretację obu warunków, postrzegamy funkcję f przekształcającą A w B jako *przyporządkowanie* każdemu elementowi zbioru A dokładnie jednego elementu ze zbioru B . Ten jedyny element y taki, że $\langle x, y \rangle \in f$, tzn. przyporządkowany elementowi x według funkcji f , oznacza się jako $f(x)$ i nazywa *wartością funkcji f dla argumentu x* . Mamy zatem:

$$\forall x \in A \forall y \in B (\langle x, y \rangle \in f \leftrightarrow y = f(x)).$$

Zbiór wszystkich funkcji przekształcających zbiór A w B oznaczamy: B^A , tzn.

$$B^A = \{f \subseteq A \times B : f \text{ jest funkcją przekształcającą } A \text{ w } B\}.$$

Uwaga. Dla dowolnego zbioru B , $B^\emptyset = \{\emptyset\}$ oraz $\emptyset^A = \emptyset$, gdy $A \neq \emptyset$.

Naturalnie zbiór oznaczany powyżej jako B^A istnieje, na podstawie aksjomatu podzbiorów, w którym formuła $\phi(x)$ jest postaci: x jest funkcją przekształcającą A w B , oraz zbiór z jest postaci: $P(A \times B)$.

Pierwsze z twierdzeń podaje oczywisty warunek konieczny i wystarczający na to, aby dwie funkcje przekształcające jeden zbiór w drugi były jednym i tym samym zbiorem; warunkiem tym jest identyczność wartości tych funkcji na wszystkich argumentach:

TWIERDZENIE 28. *Dla dowolnych $f, g \in B^A$, $f = g$ wtw $\forall x \in A, f(x) = g(x)$.*

DOWÓD: Niech $f, g \in B^A$.

($\rightarrow$): Załóżmy, że $f = g$. Weźmy $x \in A$. Wówczas $\langle x, f(x) \rangle \in f$, zatem $\langle x, f(x) \rangle \in g$. Lecz również $\langle x, g(x) \rangle \in g$. Wobec warunku (2) definicji funkcji, $f(x) = g(x)$.

($\leftarrow$): Załóżmy, że $\forall x \in A, f(x) = g(x)$. ($\subseteq$): niech $\langle x, y \rangle \in f$. Wówczas $y = f(x)$, zatem $y = g(x)$; lecz $\langle x, g(x) \rangle \in g$, więc $\langle x, y \rangle \in g$. Analogicznie dla inkluzji ($\supseteq$). $\square$

TWIERDZENIE 29. *Niech $f : A \rightarrow B$ oraz $g : B \rightarrow C$. Wówczas $f \circ g \in C^A$. (Złożenie dwóch funkcji jest funkcją.)*

DOWÓD: Jest oczywiste, że $f \circ g \subseteq A \times C$.

Wykazujemy warunek (1) definicji funkcji: niech $a \in A$. Wówczas $\langle a, f(a) \rangle \in f$ oraz $\langle f(a), g(f(a)) \rangle \in g$, zatem $\langle a, g(f(a)) \rangle \in f \circ g$ oraz $g(f(a)) \in C$.

Wykazujemy warunek (2) definicji funkcji: niech $\langle a, x \rangle, \langle a, y \rangle \in f \circ g$, gdzie $a \in A, x, y \in C$. Zatem dla pewnych $b, c \in B$, $\langle a, b \rangle \in f$ i $\langle b, x \rangle \in g$ oraz $\langle a, c \rangle \in f$ i $\langle c, y \rangle \in g$. Stąd, ponieważ f jest funkcją, na podstawie warunku (2) mamy: $b = c$; zatem $\langle b, x \rangle, \langle b, y \rangle \in g$. Lecz g jest funkcją, więc $x = y$. $\square$

Ustalmy *explicite* wartość złożenia dwóch funkcji na dowolnym argumencie:

TWIERDZENIE 30. Niech $f : A \rightarrow B$ oraz $g : B \rightarrow C$. Wówczas $\forall a \in A, (f \circ g)(a) = g(f(a))$.

DOWÓD: Niech $a \in A$. Wówczas dla dowolnego $c \in C$, $c = (f \circ g)(a)$ wtw $\langle a, c \rangle \in f \circ g$ wtw dla pewnego $b \in B$, $\langle a, b \rangle \in f$ i $\langle b, c \rangle \in g$ wtw dla pewnego $b \in B$, $b = f(a)$ i $c = g(b)$ wtw $c = g(f(a))$. Stąd $(f \circ g)(a) = g(f(a))$. $\square$

W ogólności nie jest tak, że relacja odwrotna do funkcji $f : A \rightarrow B$ jest funkcją przekształcającą zbiór B w A . Na przykład relacja odwrotna do funkcji $f : P(P(\emptyset)) \rightarrow P(P(\emptyset))$ określonej następująco: $f(\emptyset) = f(P(\emptyset)) = \emptyset$, nie spełnia żadnego z warunków (1), (2) definiujących funkcję. Obecnie ograniczymy zbiór B^A wszystkich funkcji przekształcających A w B do zbioru tylko takich funkcji, których konwers należy do zbioru A^B .

DEFINICJA. Funkcja $f : A \rightarrow B$ jest 1-1 (albo *jedno-jednoznaczna* lub *różnowartościowa*), gdy $\forall x, y \in A (f(x) = f(y) \rightarrow x = y)$ (tzn. gdy $\forall x, y \in A \forall z \in B (\langle x, z \rangle \in f \wedge \langle y, z \rangle \in f \rightarrow x = y)$).

$f : A \rightarrow B$ jest „na” (dokładniej: *przekształtka zbioru A na B*), gdy $\forall y \in B \exists x \in A (y = f(x))$ (inaczej, gdy $D^{-1}(f) = B$).

$f : A \rightarrow B$ jest *bijekcją*, gdy f jest 1-1 i „na”.

TWIERDZENIE 31. Relacja identycznościowa id_A jest bijekcją przekształcającą A na A .

DOWÓD: Oczywiście. $\square$

TWIERDZENIE 32. *Jeżeli $f : A \longrightarrow B$ jest bijekcją, to $f^\sim \in A^B$.*

DOWÓD: Niech $f : A \longrightarrow B$ będzie bijekcją. Wykazujemy, że $f^\sim$ jest funkcją przekształcającą B w A .

Warunek (1) definicji funkcji: Niech $b \in B$. Ponieważ f przekształca A na B , więc dla pewnego $a \in A$, $b = f(a)$, tzn. $\langle a, b \rangle \in f$. Stąd $\langle b, a \rangle \in f^\sim$.

Warunek (2) definicji funkcji: Niech $b \in B, x, y \in A$ oraz $\langle b, x \rangle, \langle b, y \rangle \in f^\sim$. Zatem $\langle x, b \rangle, \langle y, b \rangle \in f$. Ponieważ f jest 1-1, więc $x = y$. $\square$

Gdy relacja odwrotna do funkcji $f \in B^A$ jest funkcją przekształcającą zbiór B w A , wówczas związki między tymi funkcjami są następującej postaci:

TWIERDZENIE 33. *Niech $f \in B^A$. Jeżeli $f^\sim$ jest funkcją przekształcającą B w A , to:*

- (1) *dla dowolnych $b \in B, a \in A$, $f^\sim(b) = a$ wtw $b = f(a)$,*
- (2) *$\forall b \in B, f(f^\sim(b)) = b$ oraz $\forall a \in A, f^\sim(f(a)) = a$,*
- (3) *$f^\sim \circ f = id_B$ oraz $f \circ f^\sim = id_A$.*

DOWÓD: Niech $f \in B^A$ oraz $f^\sim \in A^B$.

dla (1): $f^\sim(b) = a$ wtw $\langle b, a \rangle \in f^\sim$ wtw $\langle a, b \rangle \in f$ wtw $b = f(a)$.

dla (2): Na mocy (1) ponieważ $f^\sim(b) \in A$ mamy: $f^\sim(b) = f^\sim(b)$ wtw $b = f(f^\sim(b))$, skąd $b = f(f^\sim(b))$ oraz, ponieważ $f(a) \in B$, więc $f^\sim(f(a)) = a$ wtw $f(a) = f(a)$, skąd $f^\sim(f(a)) = a$.

dla (3): (3) jest wnioskiem z (2), twierdzenia 30, definicji relacji (funkcji) identycznościowej oraz twierdzenia 28. $\square$

Aby wykazać, że ograniczenie zbioru B^A do zbioru bijekcji przekształcających zbiór A na B jest również niezbędne, a nie tylko wystarczające (twierdzenie 32) na to, by relacje odwrotne do nich były funkcjami przekształcającymi zbiór B w A , wykorzystamy warunek konieczny i wystarczający dla bycia bijekcją, sformułowany w następującym twierdzeniu:

TWIERDZENIE 34. *Dla dowolnej funkcji $f : A \longrightarrow B$, f jest bijekcją wtw istnieje funkcja $g : B \longrightarrow A$ taka, że $f \circ g = id_A$ oraz $g \circ f = id_B$.*

DOWÓD: $(\rightarrow)$: na mocy twierdzeń 32 oraz 33(3).

$(\leftarrow)$: Niech $f \in B^A$ oraz $g \in A^B$ będą takie, że $f \circ g = id_A$ oraz $g \circ f = id_B$. Załóżmy, że $f(x) = f(y)$ dla $x, y \in A$. Wówczas $g(f(x)) = g(f(y))$.

Lecz z założenia, $g(f(x)) = x$ oraz $g(f(y)) = y$. Zatem $x = y$ czyli f jest 1-1. Niech $b \in B$. Na mocy założenia, $f(g(b)) = b$ oraz oczywiście $g(b) \in A$. Zatem f przekształca A na B . $\square$

Twierdzenie 35. *Niech $f \in B^A$. f jest bijekcją wtw $f^\sim \in A^B$.*

Dowód: $(\rightarrow)$: Twierdzenie 32.

$(\leftarrow)$: Niech $f^\sim$ będzie funkcją przekształcającą zbiór B w A . Wówczas na mocy twierdzeń 33(3) oraz 34, f jest bijekcją. $\square$

Twierdzenie 36. *Niech $f \in B^A$. Jeżeli $f^\sim \in A^B$, to $f^\sim$ jest bijekcją.*

Dowód: Niech $f^\sim : B \rightarrow A$. Zastosujmy twierdzenie 34 kładąc w miejsce funkcji f funkcję $f^\sim$ (zamieniając zbiór A na B oraz B na A). Wówczas w oparciu o twierdzenie 33(3), wnosimy, iż $f^\sim$ jest bijekcją. $\square$

Podamy algebraiczny opis zbioru $\text{Bij}(A)$ wszystkich bijekcji przekształcających zbiór A na siebie (naturalnie istnienie takiego zbioru jest gwarantowane przez aksjomat podzbiorów). W tym celu sformułujemy najpierw twierdzenie mówiące, że operacja złożenia funkcji zachowuje własność bycia bijekcją:

Twierdzenie 37. *Złożenie bijekcji jest bijekcją.*

Dowód: Niech $f : A \rightarrow B$ oraz $g : B \rightarrow C$ będą bijekcjami. Niech $x, y \in A$ oraz załóżmy, że $(f \circ g)(x) = (f \circ g)(y)$. Wówczas na mocy twierdzenia 30, $g(f(x)) = g(f(y))$. Ponieważ g jest 1-1, $f(x) = f(y)$, a stąd $x = y$, bo f jest 1-1. Ostatecznie, $f \circ g$ jest różnowartościowa.

Niech $c \in C$. Wówczas $c = g(b)$ dla jakiegoś $b \in B$, bo g przekształca B na C . Lecz $b = f(a)$ dla jakiegoś $a \in A$, bo f przekształca A na B . Zatem $c = g(f(a)) = (f \circ g)(a)$, czyli $f \circ g$ przekształca A na C . $\square$

Definicja. Dowolny zbiór A z operacjami 2-argumentową $\circ$, 1-argumentową $-$, oraz wyróżnionym elementem 1 spełniającymi równości, dla dowolnych $x, y, z \in A$:

$$(x \circ y) \circ z = x \circ (y \circ z),$$

$$x \circ 1 = 1 \circ x = x,$$

$$x \circ -x = -x \circ x = 1,$$

nazywamy *grupą*.

TWIERDZENIE 38. $(\text{Bij}(A), \circ, \sim, \text{id}_A)$, gdzie $\text{Bij}(A)$ jest zbiorem wszystkich bijekcji przekształcających zbiór A na A , jest grupą.

DOWÓD: Na podstawie twierdzenia 37, operacja złożenia relacji $\circ$ jest operacją na zbiorze $\text{Bij}(A)$. Według twierdzeń 32 i 36 operacja konwersu relacji $\sim$ jest operacją na zbiorze $\text{Bij}(A)$. Wreszcie, na mocy twierdzenia 31, $\text{id}_A \in \text{Bij}(A)$.

Równość $(f \circ g) \circ h = f \circ (g \circ h)$ jest prawdziwa już dla dowolnych relacji binarnych f, g, h określonych na zbiorze A (twierdzenie 25(1)). Podobnie równości $f \circ \text{id}_A = f$ oraz $\text{id}_A \circ f = f$ zachodzą już dla dowolnej relacji binarnej f na zbiorze A (twierdzenie 25(2)). Równości $f \circ f^\sim = \text{id}_A$ oraz $f^\sim \circ f = \text{id}_A$ dla dowolnej $f \in \text{Bij}(A)$, są konsekwencją twierdzenia 33(3). $\square$

DEFINICJA. Niech $f : A \longrightarrow B$. Dla dowolnego $X \subseteq A$, zbiór $\vec{f}(X) = \{b \in B : \exists a \in X, b = f(a)\} = \{f(a) : a \in X\}$ nazywamy *obrazem zbioru X według funkcji f* .

Dla dowolnego $Y \subseteq B$, zbiór $\overleftarrow{f}(Y) = \{a \in A : f(a) \in Y\}$ nazywamy *przeciwbrazem zbioru Y według funkcji f* .

Uwaga. Naturalnie istnienie obrazu i przeciwbrazu zbioru według danej funkcji jest gwarantowane przez aksjomat podzbiorów.

TWIERDZENIE 39. Dla dowolnych $f : A \longrightarrow B$, $X, Y \subseteq A$, $Z \subseteq B$:

- (1) $X \subseteq Y \rightarrow \vec{f}(X) \subseteq \vec{f}(Y)$,
- (2) $\vec{f}(X \cup Y) = \vec{f}(X) \cup \vec{f}(Y)$,
- (3) $\vec{f}(X \cap Y) \subseteq \vec{f}(X) \cap \vec{f}(Y)$,
- (4) $f \text{ jest } 1-1 \rightarrow \vec{f}(X \cap Y) = \vec{f}(X) \cap \vec{f}(Y)$,
- (5) $X \subseteq \overleftarrow{f}(\vec{f}(X))$,
- (6) $\vec{f}(\overleftarrow{f}(Z)) \subseteq Z$,
- (7) $f \text{ jest bijekcją} \rightarrow (X = \overleftarrow{f}(\vec{f}(X)) \wedge \vec{f}(\overleftarrow{f}(Z)) = Z)$.

DOWÓD: dla (1): Załóżmy, że $X \subseteq Y$. Weźmy $b \in \vec{f}(X)$. Wówczas $b = f(a)$ dla pewnego $a \in X$. Zatem z założenia, $a \in Y$, czyli $f(a) \in \vec{f}(Y)$, tzn. $b \in \vec{f}(Y)$.

dla (2): Ponieważ $X \subseteq X \cup Y$ oraz $Y \subseteq X \cup Y$, więc na mocy (1) mamy: $\vec{f}(X) \subseteq \vec{f}(X \cup Y)$ oraz $\vec{f}(Y) \subseteq \vec{f}(X \cup Y)$, zatem $\vec{f}(X) \cup \vec{f}(Y) \subseteq \vec{f}(X \cup Y)$. Aby dowieść odwrotnej inkluzji założymy, że $b \in \vec{f}(X \cup Y)$. Zatem dla pewnego $a \in X \cup Y$, $b = f(a)$. Mamy więc: $a \in X$ lub $a \in Y$. Gdy $a \in X$, to oczywiście $f(a) \in \vec{f}(X) \subseteq \vec{f}(X) \cup \vec{f}(Y)$, czyli $b \in \vec{f}(X) \cup \vec{f}(Y)$. Analogicznie gdy $a \in Y$.

dla (3): Ponieważ $X \cap Y \subseteq X$ oraz $X \cap Y \subseteq Y$, więc na mocy (1), $\vec{f}(X \cap Y) \subseteq \vec{f}(X)$ oraz $\vec{f}(X \cap Y) \subseteq \vec{f}(Y)$, zatem $\vec{f}(X \cap Y) \subseteq \vec{f}(X) \cap \vec{f}(Y)$.

dla (4): Załóżmy, że f jest różnowartościowa. Na mocy (3) wystarczy wykazać inkluzję $\vec{f}(X) \cap \vec{f}(Y) \subseteq \vec{f}(X \cap Y)$. Niech więc $b \in \vec{f}(X) \cap \vec{f}(Y)$. Wtedy $b \in \vec{f}(X)$ oraz $b \in \vec{f}(Y)$. Zatem dla pewnego $a_1 \in X$, $b = f(a_1)$ oraz dla pewnego $a_2 \in Y$, $b = f(a_2)$. Wówczas, skoro $f(a_1) = f(a_2)$, więc z założenia, $a_1 = a_2$. Zatem $a_1 \in Y$ (bo $a_2 \in Y$). Ostatecznie, $a_1 \in X \cap Y$, skąd $f(a_1) \in \vec{f}(X \cap Y)$, tzn. $b \in \vec{f}(X \cap Y)$.

dla (5): Niech $a \in X$. Wówczas $f(a) \in \vec{f}(X)$, zatem z definicji przeciwoobrazu: $a \in \overleftarrow{f}(\vec{f}(X))$.

dla (6): Niech $b \in \overleftarrow{f}(\vec{f}(Z))$. Wówczas dla pewnego $a \in \vec{f}(Z)$, $b = f(a)$. Lecz wtedy, z definicji przeciwoobrazu, $f(a) \in Z$, zatem $b \in Z$.

dla (7): Załóżmy, że f jest 1-1 oraz „na”. Aby wykazać pierwszą z równości wystarczy, na mocy (5), wykazać inkluzję $\overleftarrow{f}(\vec{f}(X)) \subseteq X$. Niech więc $a \in \overleftarrow{f}(\vec{f}(X))$. Wówczas $f(a) \in \vec{f}(X)$, zatem dla pewnego $a_1 \in X$, $f(a) = f(a_1)$. Ponieważ f jest 1-1, więc $a = a_1$, stąd $a \in X$.

Aby wykazać drugą równość wystarczy, na mocy (6) wykazać inkluzję $Z \subseteq \overleftarrow{f}(\vec{f}(Z))$. Załóżmy więc, że $b \in Z$. Ponieważ funkcja f przekształca zbiór A na B , zaś $Z \subseteq B$, więc dla pewnego $a \in A$, $b = f(a)$. Stąd $f(a) \in Z$, zatem $a \in \overleftarrow{f}(Z)$, co z kolei oznacza, że $f(a) \in \overleftarrow{f}(\vec{f}(Z))$. Ostatecznie, $b \in \overleftarrow{f}(\vec{f}(Z))$. $\square$

Twierdzenie 39 nie jest naturalnie wyczerpującym przeglądem własności obrazu i przeciwoobrazu zbioru. Można go rozszerzyć podając na przykład odpowiedniki warunków (1), (2), (3), (4) dla przeciwoobrazu:

$$U \subseteq V \rightarrow \overleftarrow{f}(U) \subseteq \overleftarrow{f}(V),$$

$$\begin{aligned}\overleftarrow{f}(U \cup V) &= \overleftarrow{f}(U) \cup \overleftarrow{f}(V), \\ \overleftarrow{f}(U \cap V) &= \overleftarrow{f}(U) \cap \overleftarrow{f}(V), \text{ dla dowolnych } U, V \subseteq B.\end{aligned}$$

Na koniec zdefiniujemy jeszcze użyteczne pojęcia obcięcia funkcji oraz ciągu.

DEFINICJA. Niech $f \in B^A$ oraz $X \subseteq A$. Naturalnie zbiór par uporządkowanych $\{ \langle x, f(x) \rangle : x \in X \}$ jest funkcją przekształcającą zbiór X w zbiór B . Nazywamy ją *obcięciem funkcji f do zbioru X* i oznaczamy: $f \upharpoonright X$.

DEFINICJA. Niech A będzie niepustym zbiorem. Dowolną funkcję $a : \{1, \dots, n\} \rightarrow A$ nazywamy *n -wyrazowym ciągiem elementów ze zbioru A* i oznaczamy $(a(1), \dots, a(n))$.

Dowolną funkcję $a : \{1, 2, \dots\} \rightarrow A$ nazywamy *ciągiem (nieskończonym) elementów ze zbioru A* i oznaczamy $(a(1), a(2), \dots)$.

Wartość ciągu a na argumentie $i \in \{1, \dots, n, \dots\}$, czyli $a(i)$, nazywamy *i -tym wyrazem ciągu a* .

2.8 Zbiory ufundowane

Obecnie skupimy uwagę na roli aksjomatu ufundowania (regularności) w teorii Zermelo-Fraenkla ZF . W tym celu rozważamy najpierw teorię bez tego aksjomatu, zastępując go innym, nie należącym ani do ZF^- , ani do ZF (formuła Ω). W ten sposób zapewniamy istnienie zbiorów *nie mających elementu minimalnego* oraz tzw. zbiorów *nieufundowanych*. Aby oddać intuicje związane z pojęciami takich zbiorów, wprowadzamy nieformalne pojęcie *nieskończonego zejścia zbiorów*. Twierdzenia 41, 42, 44, a więc twierdzenia, w których to pojęcie występuje, należy zatem traktować nieformalnie – nie należą one do ZF . Ponadto twierdzenie 40 również nie jest twierdzeniem tej teorii, lecz z innego powodu: jego sformułowanie i dowód wymagają aksjomatu (Ω). Natomiast twierdzenia 43, 45, 46, 47, 48, 49 należą do teorii ZF^- , bowiem w ich sformułowaniu i dowodach nie wykorzystuje się formuły (Ω).

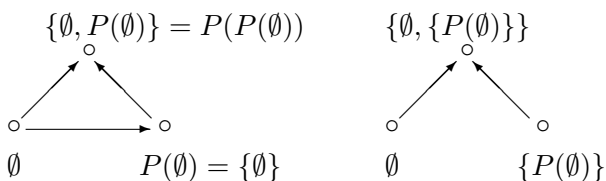
DEFINICJA. Dla dowolnych zbiorów x, y , y jest *elementem minimalnym* zbioru x , gdy $y \in x$ oraz $y \cap x = \emptyset$.

PRZYKŁAD. Zbiór $\emptyset$ jest (jedynym) elementem minimalnym zbioru $\{\emptyset, P(\emptyset)\}$. Zbiory $\emptyset$, $\{P(\emptyset)\}$ są elementami minimalnymi zbioru $\{\emptyset, \{P(\emptyset)\}\}$.

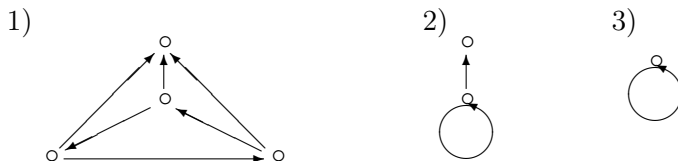
W celu określania elementów minimalnych w danym zbiorze, wygodnie jest posłużyć się ilustracją tzw. grafu zbioru.

Fakt, że $y \in x$ ilustrujemy w postaci: $y \circ \longrightarrow \circ x$, tzn. różne zbiory x, y oznaczamy różnymi punktami, a zachodzący między nimi stosunek należenia zaznaczamy odpowiednio strzałką. Przez ilustrację grafu zbioru x będziemy rozumieć rysunek, na którym znajdują się punkty odpowiadające zbiorowi x oraz wszystkim jego elementom, oraz strzałki pokazujące wszystkie stosunki należenia jakie zachodzą między zbiorami, którym odpowiadają owe punkty. Element minimalny zbioru x będzie wówczas oznaczony punktem, do którego nie „wchodzi” żadna strzałka „wychodząca” z jakiegoś elementu tego zbioru.

PRZYKŁAD. Grafy dwóch zbiorów z przykładu powyżej:



Poniżej podajemy przykłady grafów zbiorów, które nie posiadają elementu minimalnego (gdyby takie zbiory istniały). Każdy z tych zbiorów ilustrowany jest punktem najwyżej na rysunku położonym. Pierwszy z tych zbiorów jest 3-elementowy, pozostałe dwa zbiory są jednoelementowe:



2.8.1 Teoria ZF^- z aksjomatem Ω

W dalszym ciągu do zestawu aksjomatów teorii ZF^- dodajemy formułę (por. np. Aczel[1]):

$$(\Omega) \exists y \forall x (x \in y \leftrightarrow x = y).$$

Oznaczmy zbiór, którego istnienie ona stwierdza symbolem Ω . Wówczas mamy:

$$(1) \quad \forall x(x \in \Omega \leftrightarrow x = \Omega).$$

Z (1) otrzymujemy natychmiast: $\Omega \in \Omega$ oraz Ω jest jedynym takim zbiorem x , że $x \in \Omega$. Krótko mówiąc Ω jest singletonem, którego jedynym elementem jest właśnie Ω . Precyzyjniej, można to wykazać następująco. Z definicji zbioru 1-elementowego mamy:

$$(2) \quad \forall x(x \in \{\Omega\} \leftrightarrow x = \Omega).$$

Zatem z (1) oraz (2): $\forall x(x \in \Omega \leftrightarrow x \in \{\Omega\})$, skąd przy użyciu aksjomatu identyczności uzyskujemy:

$$(3) \quad \Omega = \{\Omega\}.$$

Zauważmy, że aksjomat identyczności nie gwarantuje wcale jedyności zbioru, którego istnienie stwierdza formuła (Ω) , czyli różnych zbiorów Ω o własności (3) może istnieć więcej niż jeden. Mówiąc dalej o zbiorze Ω mamy na myśli jakikolwiek ze zbiorów o własności (3). Jak widać, graf takiego zbioru jest przedstawiony powyżej na trzecim diagramie.

TWIERDZENIE 40. *Zbiór Ω nie ma elementu minimalnego.*

DOWÓD: Oczywiście jedynym elementem zbioru Ω jest Ω , lecz Ω nie jest elementem minimalnym zbioru Ω , bowiem $\Omega \cap \Omega \neq \emptyset$, skoro $\Omega \cap \Omega = \Omega = \{\Omega\} \neq \emptyset$. $\square$

DEFINICJA. Sekwencję zbiorów (niekoniecznie różnych) $x_1, x_2, \dots, x_n, \dots$ nazwiemy *nieskończonym wejściem* (zejściem), gdy $x_1 \in x_2, x_2 \in x_3, \dots, x_{n-1} \in x_n, x_n \in x_{n+1}, \dots$ ($\dots, x_{n+1} \in x_n, x_n \in x_{n-1}, \dots, x_3 \in x_2, x_2 \in x_1$).

PRZYKŁAD. Sekwencja zbiorów $\emptyset, \{\emptyset\}, \{\{\emptyset\}\}, \dots, \{\dots\{\emptyset\}\dots\}, \dots$ jest nieskończonym wejściem. Natomiast sekwencja $\Omega, \{\Omega\}, \{\{\Omega\}\}, \dots, \{\dots\{\Omega\}\dots\}, \dots$ jest jednocześnie nieskończonym wejściem i zejściem. Z definicji singletonu jest bowiem oczywiste, że $\Omega \in \{\Omega\} \in \{\{\Omega\}\} \in \dots \in \{\dots\{\Omega\}\dots\} \in \dots$. Jednakże również mamy: $\dots \in \{\dots\{\Omega\}\dots\} \in \dots \in \{\{\Omega\}\} \in \{\Omega\} \in \Omega$, skoro na mocy (3), dowolny singleton $\{\dots\{\Omega\}\dots\}$ jest identyczny ze zbio-

rem Ω , zatem sam jest jedynym swoim elementem (powyższa sekwencja ma oczywiście postać: $\Omega, \Omega, \dots, \Omega, \dots$).

Uwaga. Pojęcie nieskończonego zejścia (wejścia), tak jak zostało tu wprowadzone, formalnie nie jest pojęciem teorii ZF^- , bowiem bazuje ono na pojęciu sekwencji zbiorów, które z kolei jest tutaj pojęciem wyłącznie intuicyjnym, pierwotnie zakładanym, a więc nie należącym formalnie do teorii ZF^- . Powodem, dla którego wprowadzamy pojęcie nieskończonego zejścia jest objaśnienie intuicji związanych z pojęciem *zbioru nieufundowanego*.

W dalszym ciągu twierdzenia, w których sformułowaniu bądź dowodzie pojawia się pojęcie nieskończonego zejścia będą oznaczane symbolem $*$.

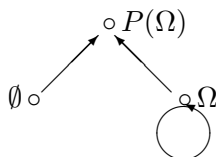
DEFINICJA. Powiemy, że zbiór x ma *nieskończone zejście*, gdy istnieje nieskończone zejście postaci: $x, x_1, x_2, \dots, x_n, \dots$. Wówczas będziemy mówić również, że takie nieskończone zejście jest *nieskończonym zejściem zbioru x* .

PRZYKŁAD. Naturalnie zbiór Ω ma nieskończone zejście: $\Omega, \Omega, \dots, \Omega, \dots$

*TWIERDZENIE 41: *Każdy niepusty zbiór nie mający elementu minimalnego ma nieskończone zejście.*

DOWÓD: Niech x będzie niepustym zbiorem nie mającym elementu minimalnego. Wówczas dla pewnego zbioru y_1 mamy: $y_1 \in x$ (bo $x \neq \emptyset$). Z założenia, y_1 nie jest elementem minimalnym zbioru x , zatem $y_1 \cap x \neq \emptyset$. Istnieje więc zbiór y_2 taki, że $y_2 \in y_1$ oraz $y_2 \in x$. Choć y_2 jest elementem zbioru x , to jednak nie jest jego elementem minimalnym. Zatem istnieje zbiór y_3 taki, że $y_3 \in y_2$ oraz $y_3 \in x$. I tak dalej. Jest widoczne, że sekwencja $x, y_1, y_2, y_3, \dots$ jest nieskończonym zejściem zbioru x . $\square$

Twierdzenie odwrotne do twierdzenia 41 nie jest prawdziwe. Istnieją zbiory (oczywiście w ramach ZF^- z aksjomatem (Ω)) z nieskończonym zejściem mające element minimalny. Np. zbiór $P(\Omega)$ ma tę własność: zbiór $\emptyset$ jest jego elementem minimalnym, chociaż sekwencja $\{\emptyset, \Omega\}, \Omega, \Omega, \dots, \Omega, \dots$ jest jego nieskończonym zejściem, jak na to wskazuje graf zbioru $P(\Omega)$:

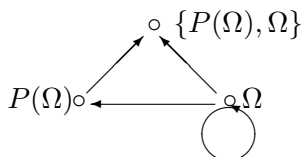


DEFINICJA. Mówimy, że zbiór x jest *ufundowany*, gdy dla dowolnego zbioru y , jeżeli $x \in y$, to y ma element minimalny.

Gdy zbiór nie jest ufundowany, mówimy, że jest on *nieufundowany*. Zatem x jest nieufundowany, gdy jest on elementem pewnego zbioru nie mającego elementu minimalnego.

PRZYKŁAD. Zbiór $\emptyset$ jest ufundowany. Jeśli bowiem $\emptyset \in y$, to $\emptyset$ jest elementem minimalnym zbioru y , bo $\emptyset \cap y = \emptyset$.

Zbiór $P(\Omega)$ nie jest ufundowany. Rozważmy bowiem zbiór $\{P(\Omega), \Omega\}$. Naturalnie $P(\Omega) \in \{P(\Omega), \Omega\}$, lecz zbiór $\{P(\Omega), \Omega\}$ nie ma elementu minimalnego:



*TWIERDZENIE 42. *Każdy zbiór nieufundowany ma nieskończone zejście.*

DOWÓD: Niech x będzie zbiorem nieufundowanym. Wówczas dla pewnego zbioru y , $x \in y$ oraz y nie ma elementu minimalnego. Zatem x nie jest elementem minimalnym zbioru y , czyli $x \cap y \neq \emptyset$. Niech więc $y_1 \in x$ oraz $y_1 \in y$. y_1 nie jest elementem minimalnym zbioru y , zatem $y_1 \cap y \neq \emptyset$. Niech więc $y_2 \in y_1$ oraz $y_2 \in y$. I tak dalej. Naturalnie sekwencja $x, y_1, y_2, \dots$ jest nieskończonym zejściem zbioru x . $\square$

Jak widać na podstawie twierdzeń 41 i 42, niepuste zbiory nie mające elementu minimalnego oraz zbiory nieufundowane zachowują się podobnie ze względu na posiadanie nieskończonego zejścia. Spróbujmy bliżej scharakteryzować związki między tymi klasami zbiorów, w ramach teorii ZF^- (twierdzenia 43, 45, 46, 47) oraz ZF^- z aksjomatem (Ω) (twierdzenie 44).

TWIERDZENIE 43. *Każdy niepusty zbiór nie mający elementu minimalnego jest zbiorem nieufundowanym. (Inaczej, każdy niepusty zbiór ufundowany ma element minimalny.)*

DOWÓD: Niech $x \neq \emptyset$ będzie zbiorem nie mającym elementu minimalnego. Rozważmy zbiór $x \cup \{x\}$. Wówczas naturalnie $x \in x \cup \{x\}$. Wykażemy,

że $x \cup \{x\}$ nie ma elementu minimalnego, co skończy dowód. Załóżmy nie wprost, że y jest elementem minimalnym zbioru $x \cup \{x\}$. Wówczas

- (1) $y \in x \cup \{x\}$ oraz
- (2) $y \cap (x \cup \{x\}) = \emptyset$.

Konsekwencją (2) jest oczywiście

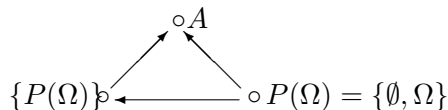
- (3) $y \cap x = \emptyset$,

gdyby bowiem dla pewnego z , $z \in y \cap x$, to $z \in y$ oraz $z \in x$, skąd $z \in x \cup \{x\}$, zatem $z \in y \cap (x \cup \{x\})$, co jest niemożliwe wobec (2).

Z (1) mamy natychmiast: $y \in x$ lub $y = x$. Gdy $y \in x$, to wobec (3), y jest elementem minimalnym zbioru x , wbrew założeniu. Gdy zaś $y = x$, to według (3), $x = \emptyset$, co również jest sprzeczne z założeniem. $\square$

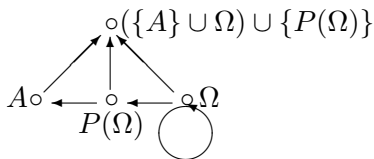
TWIERDZENIE 44. *Twierdzenie odwrotne do twierdzenia 43 jest fałszywe, tzn. istnieją zbiory nieufundowane posiadające element minimalny.*

DOWÓD: Wykażmy, że zbiór $A = \{\{P(\Omega)\}, P(\Omega)\}$ nie jest ufundowany, lecz posiada element minimalny. Graf zbioru A jest postaci:



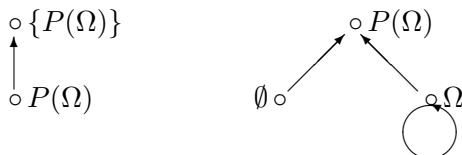
Zatem $P(\Omega)$ jest elementem minimalnym zbioru A .

Oczywiście $A \in (\{A\} \cup \Omega) \cup \{P(\Omega)\}$. Jednakże zbiór $(\{A\} \cup \Omega) \cup \{P(\Omega)\}$ nie ma elementów minimalnych:



Zatem A nie jest ufundowany. $\square$

Nie jest również tak, że zbiór nieufundowany musi mieć element, który nie ma elementów minimalnych. Zbiór A z dowodu twierdzenia 44 jest nieufundowany, lecz każdy jego element posiada element minimalny. Grafy elementów zbioru A są postaci:



Okazuje się jednak, że zbiór nieufundowany musi mieć element, który nie jest zbiorem ufundowanym:

Twierdzenie 45. *Dla dowolnego zbioru x , jeżeli x nie jest ufundowany, to istnieje zbiór y taki, że $y \in x$ oraz y nie jest ufundowany. (Inaczej, jeżeli każdy element zbioru x jest ufundowany, to x jest ufundowany.)*

Dowód: Dowodzimy drugiego sformułowania twierdzenia. Załóżmy, że $\forall y(y \in x \rightarrow y \text{ jest ufundowany})$. Aby wykazać, że x jest ufundowany, załóżmy, że $x \in z$. Wykażemy, że z ma element minimalny. Naturalnie x jest elementem minimalnym zbioru z lub x nie jest elementem minimalnym zbioru z . Załóżmy więc, że x nie jest elementem minimalnym zbioru z . Wówczas $x \cap z \neq \emptyset$. Niech $a \in x$ oraz $a \in z$. Ponieważ z założenia każdy element zbioru x jest ufundowany, więc a jest ufundowany. Zatem, skoro $a \in z$, to z ma element minimalny. $\square$

Twierdzenie odwrotne do twierdzenia 45 jest również prawdziwe.

Twierdzenie 46. *Dowolny zbiór mający nieufundowany element jest nieufundowany. (Inaczej, każdy element dowolnego zbioru ufundowanego jest zbiorem ufundowanym.)*

Dowód: Załóżmy, że x, y są zbiorami takimi, że

- (1) $y \in x$ oraz
- (2) y jest nieufundowany.

Na mocy (2) niech z będzie takim zbiorem, że

- (3) $y \in z$ oraz
- (4) z nie ma elementu minimalnego.

Rozważmy zbiór $z \cup \{x\}$. Naturalnie $x \in z \cup \{x\}$. Pokażemy że zbiór $z \cup \{x\}$ nie ma elementu minimalnego. Załóżmy nie wprost, że a jest elementem minimalnym tego zbioru, tzn.

- (5) $a \in z \cup \{x\}$ oraz
- (6) $a \cap (z \cup \{x\}) = \emptyset$.

Wówczas z (5): $a \in z$ lub $a = x$. Załóżmy, że $a \in z$. Na mocy (4), a nie jest elementem minimalnym zbioru z , zatem $a \cap z \neq \emptyset$. Stąd również $a \cap (z \cup \{x\}) \neq \emptyset$; sprzeczność z (6). Załóżmy, że $a = x$. Wówczas z (1) mamy: $y \in a$. Ponadto z (3), $y \in z \cup \{x\}$. Czyli $y \in a \cap (z \cup \{x\})$. Zatem znowu $a \cap (z \cup \{x\}) \neq \emptyset$.

Ostatecznie, x nie jest ufundowany. $\square$

Różnicę między pojęciami zbioru nieufundowanego oraz niepustego zbioru nie posiadającego elementu minimalnego można intuicyjnie wyjaśnić w oparciu o pojęcie nieskończonego zejścia. Choć zarówno zbiór nieufundowany jak i niepusty zbiór nie mający elementu minimalnego mają nieskończone zejścia (Twierdzenia 42 i 41), to jednak zbiór posiadający nieskończone zejście może mieć element minimalny, tymczasem jest on zbiorem nieufundowanym. Bowiem:

**zbiór nieufundowany to taki zbiór, który ma nieskończone zejście.*

Aby to stwierdzenie uzasadnić weźmy pod uwagę jakikolwiek zbiór x mający nieskończone zejście: $x, y_1, \dots, y_n, \dots$ i rozważmy klasę wszystkich zbiorów w tym zejściu (sekwencji) występujących. Zbiór x jest elementem tej klasy, lecz klasa ta nie ma elementu minimalnego, skoro każdy zbiór z tej sekwencji ma element występujący (bezpośrednio za nim) w tej sekwencji. Zatem zbiór x jest nieufundowany.

Uzasadnienie to miałoby nieformalną wartość, gdyby wyrażenie „klasa” można było zastąpić wyrażeniem „zbiór” lub też zmodyfikować definicję zbioru nieufundowanego w sposób następujący: zbiór x jest *nieufundowany*, gdy istnieje klasa (mnogość, ekspozycja) zbiorów, w której x się znajduje oraz każdy zbiór tej klasy ma element w niej występujący.

Pierwsze z tych rozwiązań odrzucamy – nie mamy bowiem w ogólności żadnych gwarancji istnienia zbioru wszystkich zbiorów występujących w jakimś nieskończonym zejściu. Drugie rozwiązanie łatwo akceptujemy, lecz wyłącznie w ramach rozważań nieformalnych (tzn. poza teorią ZF^-), które przecież od zakończenia dowodu twierdzenia 46 prowadzimy, skoro posługujemy się pojęciem nieskończonego zejścia.

Łatwo zauważyć, że przy tak zmienionym pojęciu zbioru nieufundowanego, można bez trudności zmodyfikować dowód twierdzenia 42, aby uzasadnić twierdzenie: **każdy zbiór nieufundowany ma nieskończone zejście.*

Naturalnie w ramach teorii ZF^- czy ZF funkcjonuje pojęcie zbioru nieufundowanego (czy ufundowanego) wcześniej zdefiniowane bez odwoływania się do pojęcia nieskończonego zejścia. Nie można w ZF^- utożsamiać (formalnie) zbioru nieufundowanego ze zbiorem mającym nieskończone zejście, bowiem w ramach tej teorii nie mamy przecież pojęcia nieskończonego zejścia.

Patrzając jednak nieformalnie na zbiory nieufundowane jako na te i tylko te, które posiadają nieskończone zejście, jaśniejsze stają się twierdzenia, które funkcjonują bez pojęcia nieskończonego zejścia:

twierdzenie 43 – na mocy twierdzenia 41, **niepusty zbiór nie mający elementu minimalnego ma nieskończone zejście, zatem jest nieufundowany,*

twierdzenie 44 – **istnieją zbiory mające nieskończone zejście oraz element minimalny,*

twierdzenie 45 – **jeżeli zbiór x ma nieskończone zejście, powiedzmy $x, y_1, y_2, \dots, y_n, \dots$, to $y_1 \in x$ oraz y_1 ma nieskończone zejście $y_1, y_2, \dots, y_n, \dots$,*

twierdzenie 46 – **jeżeli $y \in x$ oraz y ma nieskończone zejście $y, y_1, y_2, \dots, y_n, \dots$, to zbiór x ma nieskończone zejście $x, y, y_1, y_2, \dots, y_n, \dots$*

W teorii ZF^- klasa wszystkich zbiorów niepustych nie mających elementu minimalnego zawiera się w klasie wszystkich zbiorów nieufundowanych (twierdzenie 43). Zatem, gdyby w ZF^- nie było zbiorów nieufundowanych, to nie byłoby tam również niepustych zbiorów nie mających elementu minimalnego. Jednakże łatwo zauważyć, że również odwrotnie, gdyby nie było tam niepustych zbiorów bez elementu minimalnego, to nie byłoby również zbiorów nieufundowanych:

Twierdzenie 47. *Każdy niepusty zbiór ma element minimalny wtw każdy zbiór jest ufundowany.*

Dowód: ($\rightarrow$): Załóżmy, że każdy niepusty zbiór posiada element minimalny oraz nie wprost, niech a będzie zbiorem, który nie jest ufundowany. Wówczas dla pewnego zbioru y nie mającego elementu minimalnego, $a \in y$. Lecz wówczas $y \neq \emptyset$, zatem z założenia musi mieć element minimalny. Sprzeczność.

($\leftarrow$): Na mocy twierdzenia 43. $\square$

Odnótujmy dwie istotne własności zbiorów ufundowanych:

TWIERDZENIE 48: $\forall x(x \text{ jest ufundowany} \rightarrow x \notin x)$.

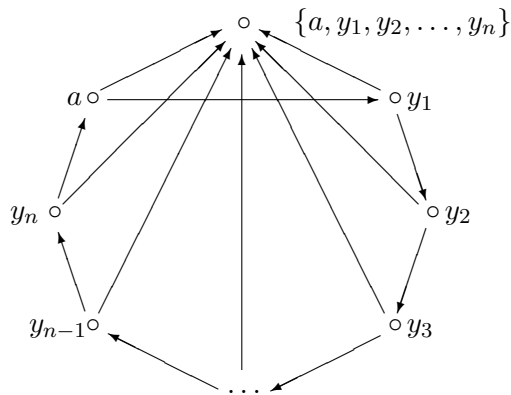
DOWÓD: Załóżmy nie wprost, że a jest ufundowany oraz $a \in a$. Ponieważ $a \in \{a\}$, więc zbiór $\{a\}$ ma element minimalny (z definicji zbioru ufundowanego). Niech b będzie elementem minimalnym zbioru $\{a\}$. Wówczas $b \in \{a\}$ oraz $b \cap \{a\} = \emptyset$. Zatem $b = a$ i konsekwentnie $a \cap \{a\} = \emptyset$. Tymczasem z założenia, $a \in a$. Ponieważ $a \in \{a\}$, więc $a \in a \cap \{a\}$, czyli $a \cap \{a\} \neq \emptyset$. Sprzeczność. $\square$

TWIERDZENIE 49. Dla dowolnego zbioru x , jeżeli x jest ufundowany, to nie istnieje liczba naturalna n oraz zbiory $y_1, y_2, \dots, y_n$ takie, że $x \in y_1 \wedge y_1 \in y_2 \wedge \dots \wedge y_{n-1} \in y_n \wedge y_n \in x$.

DOWÓD: Załóżmy nie wprost, że a jest zbiorem ufundowanym oraz dla pewnych $y_1, y_2, \dots, y_n$ mamy:

$$(1) \quad a \in y_1 \in y_2 \in \dots \in y_{n-1} \in y_n \in a.$$

Ponieważ a jest ufundowany oraz $a \in \{a, y_1, y_2, \dots, y_n\}$, więc zbiór $\{a, y_1, y_2, \dots, y_n\}$ ma element minimalny. Tymczasem na mocy (1), jak wskazuje ilustracja grafu tego zbioru, zbiór ten nie ma elementu minimalnego:



$\square$

2.8.2 Aksjomat regularności (ufundowania)

W języku pierwotnym teorii mnogości aksjomat regularności ma postać:

$$(AxR) \forall x[\exists y(y \in x) \rightarrow \exists y(y \in x \wedge \forall z(z \in y \rightarrow z \notin x))].$$

Każdy *niepusty* zbiór (tzn. taki zbiór x , że $\exists y(y \in x)$) posiada *element minimalny* (tzn. taki element y , który ze zbiorem x nie ma wspólnych elementów: $\forall z(z \in y \rightarrow z \notin x)$).

Można sformułować ów aksjomat następująco:

$$(AxR)' \forall x(x \neq \emptyset \rightarrow \exists y(y \text{ jest elementem minimalnym } x)),$$

lub równoważnie na mocy twierdzenia 47:

$$(AxR)'' \forall x (x \text{ jest ufundowany}).$$

W teorii ZF (z aksjomatem regularności) nie istnieją zbiory takie jak Ω (negacja formuły (Ω) jest twierdzeniem ZF) czyli zbiory nieufundowane. Według nieformalnych rozważań, takie zbiory nie mają „pełnego fundamentu”, co się przejawia w posiadaniu nieskończonego zejścia. Jednakże pojęcie nieskończonego zejścia nie mieści się w ramach, często nieświadomie akceptowanego, atomizmu metodologicznego B. Russella, według którego jakiegokolwiek konstrukty formalne winny być na początku budowane z fundamentalnych cegiełek (atomów) – tzw. *building blocks*. Zbiór, jako konstrukt formalny, winien być „rozkładalny” w skończonej ilości kroków na ostateczne elementarne części. Tymczasem zbiór posiadający nieskończone zejście, „rozkładalny” jest w nieskończoność.

W ZF , na podstawie aksjomatu ufundowania, otrzymujemy jako wnioski z twierdzeń 48 oraz 49 odpowiednio następujące zdania:

WNIOSEK Z TWIERDZENIA 48. $\forall x(x \notin x)$,

WNIOSEK Z TWIERDZENIA 49. *Dla dowolnego zbioru x nie istnieje liczba naturalna n oraz zbiory $y_1, y_2, \dots, y_n$ takie, że $x \in y_1 \wedge y_1 \in y_2 \wedge \dots \wedge y_{n-1} \in y_n \wedge y_n \in x$.*

Konsekwencją wniosku z twierdzenia 48 jest

Twierdzenie 50. $\neg\exists y\forall x(x \in y)$ (nie istnieje zbiór wszystkich zbiorów).

Dowód: Załóżmy nie wprost, że $\exists y\forall x(x \in y)$. Wówczas dla pewnego zbioru V mamy: $\forall x(x \in V)$. Zatem $V \in V$, co jest niemożliwe na mocy wniosku z twierdzenia 48. $\square$

2.9 Interpretacja arytmetyki elementarnej w teorii ZF

Interpretacja arytmetyki elementarnej w teorii ZF polega na wykazaniu, że każde twierdzenie arytmetyki elementarnej jest twierdzeniem teorii mnogości ZF , o ile wcześniej, po pierwsze, ustalono do jakich zbiorów (opisywanych w ZF) odnoszą się kwantyfikatory w języku arytmetyki, czyli jakie zbiory mają być postrzegane jako liczby naturalne, oraz po drugie, zinterpretowano w ZF pierwotne pojęcia arytmetyki elementarnej: liczbę zero oraz następnik, tzn. określono jaki wyróżniony zbiór (wśród zbiorów określonych jako liczby naturalne) odpowiada stałej 0 oraz jaka operacja na zbiorach (określonych jako liczby naturalne) odpowiada symbolowi funkcijnemu s .

Naturalnie, aby wykazać, że każde twierdzenie arytmetyki elementarnej jest twierdzeniem teorii ZF , wystarczy wykazać, że każdy aksjomat arytmetyki (po zinterpretowaniu) jest twierdzeniem ZF .

2.9.1 Operacja następnika

Na początek zdefiniujemy jednoargumentową operację S zwaną operacją następnika, określoną dla wszystkich zbiorów:

DEFINICJA. Dla dowolnego zbioru x , $S(x) = x \cup \{x\}$. To znaczy, $\forall y(y \in S(x) \leftrightarrow (y \in x \vee y = x))$.

Następnik $S(x)$ zbioru x jest więc takim zbiorem, że jednocześnie $x \in S(x)$ oraz $x \subseteq S(x)$.

Twierdzenie 51. Dla dowolnego zbioru x , $S(x)$ jest najmniejszym zbiorem (względem inkluzji) wśród wszystkich zbiorów y takich, że $x \in y$ oraz $x \subseteq y$.

DOWÓD: Załóżmy, że $x \in y$ oraz $x \subseteq y$. Wówczas $\{x\} \subseteq y$, zatem $x \cup \{x\} \subseteq y$, czyli $S(x) \subseteq y$. $\square$

Zauważmy ponadto, że dla dowolnego zbioru x nie istnieje taki zbiór y , że $x \in y$ oraz $y \in S(x)$. Gdyby bowiem taki y istniał, to byłoby: $y \in x$ lub $y = x$, lecz oba człony tej alternatywy są zabronione, na mocy wniosków z odpowiednio twierdzenia 49 oraz 48.

Kolejne dwa twierdzenia charakteryzujące operację S , wskazują na możliwość jej wykorzystania jako interpretacji symbolu funkcyjnego „ s ” z języka arytmetyki:

TWIERDZENIE 52. $\forall x(S(x) \neq \emptyset)$.

DOWÓD: Oczywisty na mocy definicji operacji następnika. $\square$

TWIERDZENIE 53. $\forall x \forall y(S(x) = S(y) \rightarrow x = y)$.

DOWÓD: Załóżmy, że $S(x) = S(y)$. Ponieważ według wniosku z twierdzenia 49 wykluczona jest koniunkcja $x \in y \wedge y \in x$, więc prawdziwa jest alternatywa $x \notin y \vee y \notin x$. Załóżmy, że $x \notin y$. Ponieważ $x \in S(x)$, więc z założenia dowodu uzyskujemy: $x \in S(y)$, zatem $x \in y \vee x = y$, skąd $x = y$. Analogicznie postępujemy przy założeniu, że $y \notin x$. $\square$

Jak widać, twierdzenie 52 jest interpretacją aksjomatu (AN1), o ile zbiór pusty uznamy za interpretację stałej 0. Naturalnie twierdzenie 53 jest interpretacją aksjomatu (AN2) arytmetyki elementarnej. Na uwagę zasługuje fakt, że kwantyfikatory pojawiające się w tych dwóch twierdzeniach odnoszą się do wszelkich zbiorów, zatem również do tych, które później określimy jako liczby naturalne.

Ograniczenie kwantyfikacji przy interpretacji aksjomatów arytmetyki w ZF do specjalnej klasy zbiorów, okazuje się konieczne. Nie oczekujemy bowiem, że interpretacja aksjomatu indukcji, w której kwantyfikatory odnoszą się do wszelkich zbiorów (gdy interpretacjami 0 i s są odpowiednio $\emptyset, S$) – nawet przy jakimś rozsądnym ograniczeniu klasy formuł $\psi(x)$ (teraz będących formułami języka ZF) – okaże się twierdzeniem teorii ZF . Gdyby tak było, to każde twierdzenie arytmetyki elementarnej, według tejże interpretacji (której dziedziną jest klasa wszystkich zbiorów), byłoby twierdzeniem ZF . Tak jednakże nie jest. Rozważmy na przykład twierdzenie 2 arytmetyki elementarnej, po zinterpretowaniu w postaci: $\forall x(x = \emptyset \vee \exists y(x = S(y)))$. By-

najmniej nie jest to twierdzenie teorii mnogości. Weźmy bowiem pod uwagę zbiór $\{\{\emptyset\}\}$. Wówczas prawdą jest, że $\{\{\emptyset\}\} \neq \emptyset \wedge \neg \exists y(\{\{\emptyset\}\} = S(y))$. Gdyby bowiem dla pewnego zbioru $y : \{\{\emptyset\}\} = S(y)$, to mielibyśmy: $y \in \{\{\emptyset\}\}$, zatem $y = \{\emptyset\}$; lecz wówczas $S(y) = y \cup \{y\} = \{\emptyset\} \cup \{\{\emptyset\}\} = \{\emptyset, \{\emptyset\}\}$ i ostatecznie $\{\{\emptyset\}\} = \{\emptyset, \{\emptyset\}\}$, co oczywiście nie zachodzi.

To, że kluczową rolę w interpretacji arytmetyki elementarnej odgrywa aksjomat indukcji jest widoczne na podstawie faktu, iż same twierdzenia 52 oraz 53 nie są wystarczające do uznania, że jedyną możliwą interpretacją symbolu „ s ” z języka arytmetyki jest operacja następnika S . Gdyby na przykład interpretować ów symbol jako operację zbioru potęgowego P (oraz 0 jako zbiór $\emptyset$), z aksjomatów (AN1), (AN2) otrzymujemy twierdzenia teorii ZF :

$$\forall x(P(x) \neq \emptyset) \text{ (wniosek z twierdzenia 7),}$$

$$\forall x \forall y(P(x) = P(y) \rightarrow x = y) \text{ (na mocy twierdzeń 1, 8, 2).}$$

2.9.2 Indukcja

Przedstawimy interpretację aksjomatu indukcji, w której symbol „ s ” reprezentuje operację S , uzyskując z tego aksjomatu twierdzenie teorii ZF .

Stosując operację następnika można zapisać aksjomat nieskończoności w krótszej postaci:

$$(Ax \infty)' \quad \exists x(\emptyset \in x \wedge \forall y(y \in x \rightarrow S(y) \in x)).$$

Zbiór, którego istnienie ów aksjomat stwierdza, nosi nazwę zbioru indukcyjnego:

DEFINICJA. Dla dowolnego zbioru x , x jest *indukcyjny*, gdy $\emptyset \in x$ oraz $\forall y(y \in x \rightarrow S(y) \in x)$.

Zbiór jest więc indukcyjny, gdy jego elementem jest zbiór $\emptyset$ oraz gdy jest on zamknięty na operację następnika.

W ten sposób uzyskujemy najprostsza postać aksjomatu nieskończoności:

$$(Ax \infty)'' \quad \exists x(x \text{ jest indukcyjny}).$$

Pojęcie zbioru indukcyjnego umożliwia wyodrębnienie spośród zbiorów tych, których mnogość stanowić będzie dziedzinę modelu dla aksjomatów arytmetyki elementarnej. Zbiory te nazwiemy więc liczbami naturalnymi. Właśnie do tych zbiorów ograniczymy kwantyfikację interpretując aksjomaty arytmetyki elementarnej.

DEFINICJA. Dla dowolnego zbioru x , x jest *liczbą naturalną*, gdy $\forall y(y \text{ jest indukcyjny} \rightarrow x \in y)$ (zbiór jest liczbą naturalną, gdy jest elementem każdego zbioru indukcyjnego).

Bezpośrednio z definicji liczby naturalnej oraz zbioru indukcyjnego otrzymujemy:

TWIERDZENIE 54. *Zbiór $\emptyset$ jest liczbą naturalną.*

Mnogość liczb naturalnych jest zamknięta na operację następnika:

TWIERDZENIE 55. $\forall x(x \text{ jest liczbą naturalną} \rightarrow S(x) \text{ jest liczbą naturalną})$.

DOWÓD: Załóżmy, że zbiór x jest liczbą naturalną, tzn. $\forall y(y \text{ jest indukcyjny} \rightarrow x \in y)$. Aby wykazać, że $S(x)$ jest liczbą naturalną rozważmy dowolny zbiór indukcyjny y . Wówczas z założenia, $x \in y$. Zatem z definicji zbioru indukcyjnego mamy: $S(x) \in y$, co wobec dowolności wyboru zbioru y dowodzi, że $S(x)$ jest liczbą naturalną. $\square$

Wykażemy teraz, że owa „mnogość” liczb naturalnych jest po prostu zbiorem (tzn. obiektem, którego istnienie jest dowodliwe w ZF). W tym celu rozważmy w aksjomacie podzbiorów formułę $\phi(x)$ postaci: x jest liczbą naturalną. Oznaczmy symbolem A jakikolwiek ze zbiorów, których istnienie stwierdza aksjomat nieskończoności, uzyskując prawdziwe zdanie: A jest indukcyjny. Wówczas natychmiast z definicji liczby naturalnej otrzymujemy prawdziwy poprzednik implikacji $(AxZ)_\phi$,

$$\forall x(x \text{ jest liczbą naturalną} \rightarrow x \in A),$$

zatem również prawdziwy następnik w aksjomacie podzbiorów,

$$\exists y \forall x(x \in y \leftrightarrow x \text{ jest liczbą naturalną}),$$

co w konsekwencji prowadzi do definicji zbioru liczb naturalnych N :

DEFINICJA. $\forall x(x \in N \leftrightarrow x \text{ jest liczbą naturalną})$ lub $N = \{x : x \text{ jest liczbą naturalną}\}$.

Jako bezpośredni wniosek z twierdzeń 54 oraz 55 uzyskujemy

TWIERDZENIE 56. N jest indukcyjny.

Dzięki twierdzeniu 56, symbole funkcyjne $0, s$ występujące w aksjomatach arytmetyki elementarnej możemy interpretować jako odpowiednio, zbiór $\emptyset$ oraz operacja następnika S (skoro $\emptyset \in N$ oraz zbiór N jest zamknięty na tę operację). Jest oczywiste, że na podstawie twierdzeń 52 oraz 53 mamy

$\forall x \in N (S(x) \neq \emptyset)$, oraz

$\forall x, y \in N (S(x) = S(y) \rightarrow x = y)$,

czyli dwa pierwsze aksjomaty arytmetyki elementarnej są prawdziwe w dziedzinie N (są one przecież prawdziwe w dziedzinie wszystkich zbiorów). Aby wykazać, że aksjomat indukcji jest również prawdziwy w dziedzinie N dowiedzmy najpierw, że zbiór liczb naturalnych jest najmniejszym (względem inkluzji) zbiorem indukcyjnym.

TWIERDZENIE 57: $\forall x(x \text{ jest indukcyjny} \rightarrow N \subseteq x)$.

DOWÓD: Załóżmy, że x jest indukcyjny. Niech $y \in N$. Skoro więc y jest liczbą naturalną, to $y \in x$, co dowodzi inkluzji: $N \subseteq x$. $\square$

Interpretacja aksjomatu indukcji ma postać:

TWIERDZENIE 58. Dla dowolnej formuły $\psi(x)$ języka teorii ZF ,
 $[\psi(\emptyset) \wedge \forall x \in N(\psi(x) \rightarrow \psi(S(x)))] \rightarrow \forall x \in N(\psi(x))$.

DOWÓD: Załóżmy, że

(1) $\psi(\emptyset)$ oraz

(2) $\forall x \in N(\psi(x) \rightarrow \psi(S(x)))$.

Wstawmy w aksjomacie podzbiorów $(AxZ)_\phi$ jako $\phi(x)$ formułę $x \in N \wedge \psi(x)$ uzyskując poprzednik tego aksjomatu w postaci:

$\forall x((x \in N \wedge \psi(x)) \rightarrow x \in N)$.

Oderwijmy więc następnik:

$\exists y \forall x(x \in y \leftrightarrow (x \in N \wedge \psi(x)))$.

Stąd dla pewnego zbioru B ,

(3) $\forall x(x \in B \leftrightarrow (x \in N \wedge \psi(x)))$.

Wykażmy, że B jest zbiorem indukcyjnym. Ponieważ $\emptyset \in N$, więc na mocy (1) i (3), $\emptyset \in B$. Załóżmy, że $y \in B$. Wówczas z (3), $y \in N$ oraz $\psi(y)$. Zatem z założenia (2), $\psi(S(y))$. Ponadto, skoro $y \in N$, więc $S(y) \in N$ (twierdzenie 55). Ostatecznie z (3), $S(y) \in B$, czyli B jest indukcyjny.

Zatem na podstawie twierdzenia 57 otrzymujemy

(4) $N \subseteq B$.

Aby więc dowieść, że $\forall x \in N(\psi(x))$ rozważmy dowolny $x \in N$. Wówczas $x \in B$ na mocy (4) czyli, biorąc pod uwagę (3), zachodzi $\psi(x)$. $\square$

Warto podkreślić, że w interpretacji aksjomatu indukcji jaką jest twierdzenie 58, nie ma żadnych ograniczeń dla formuł $\psi(x)$. Zastosujmy twierdzenie 58 w przypadku, gdy formuła $\psi(x)$ nie ma swojego odpowiednika w języku arytmetyki elementarnej:

Twierdzenie 59. *Każdy element dowolnej liczby naturalnej jest liczbą naturalną, tzn. $\forall x \in N \forall y(y \in x \rightarrow y \in N)$.*

Dowód: Połóżmy $\psi(x) := \forall y(y \in x \rightarrow y \in N)$. Wówczas dowodzone twierdzenie jest następnikiem implikacji twierdzenia 58 (aksjomatu indukcji). Udowodnijmy więc poprzednik. Formuła $\psi(\emptyset)$, a więc formuła postaci: $\forall y(y \in \emptyset \rightarrow y \in N)$ jest naturalnie prawdziwa.

Weźmy dowolne $x \in N$ oraz przyjmijmy założenie indukcyjne, że $\psi(x)$, tzn. $\forall y(y \in x \rightarrow y \in N)$. W celu wykazania, że prawdą jest $\psi(S(x))$, tzn. formuła postaci: $\forall y(y \in S(x) \rightarrow y \in N)$, weźmy dowolny y i załóżmy, że $y \in S(x)$. Wówczas $y \in x$ lub $y = x$. Gdy $y \in x$, to z założenia indukcyjnego mamy: $y \in N$, gdy zaś $y = x$, to oczywiście $y \in N$ (bo $x \in N$). $\square$

Bibliografia

- [1] ACZEL P., *Non-Well-Founded Sets*, CSLI, Lecture Notes 14, Stanford 1988.
- [2] BETH E. W., *The Foundations of Mathematics*, North Holland, Amsterdam 1959
- [3] BORKOWSKI L., J. SŁUPECKI, 'A Logical System based on rules and its applications in teaching Mathematical Logic', *Studia Logica*, 7: 71–113, 1958.
- [4] BORKOWSKI L., *Logika formalna*, PWN, Warszawa 1970.
- [5] BORKOWSKI L., *Wprowadzenie do logiki i teorii mnogości*, Wyd. KUL, Lublin 1991.
- [6] CORCORAN, J. 'Aristotle's Natural Deduction System', w: J. Corcoran (red.), *Ancient Logic and its Modern Interpretations*, Reidel, Dordrecht 1972.
- [7] FRAENKEL A. A., *Abstract Set Theory*, North Holland, Amsterdam 1976
- [8] FRAENKEL A. A., J. BAR-HILLEL, A. LEVY, *Foundations of set theory*, North Holland, Amsterdam 1973
- [9] GENTZEN, G., 'Untersuchungen über das Logische Schliessen', *Mathematische Zeitschrift* 39:176–210 and 39:405–431, 1934.
- [10] GRZEGORCZYK A., *Zarys arytmetyki teoretycznej*, PWN, Warszawa 1971
- [11] GUZICKI W., P. ZBIERSKI, *Podstawy teorii mnogości*, PWN Warszawa 1978

- [12] INDRZEJCZAK A., *Elementy logiki*, wyd. WSHE, Łódź 2004.
- [13] INDRZEJCZAK A., *Natural Deduction, Hybrid Systems and Modal Logics*, Springer 2010.
- [14] JAŚKOWSKI, S., 'On the Rules of Suppositions in Formal Logic' *Studia Logica* 1:5–32, 1934.
- [15] JECH T. J., *Lectures in set theory*, Lectures Notes in Mathematics 217, Springer-Verlag 1971
- [16] KURATOWSKI K., A. MOSTOWSKI, *Teoria mnogości*, PWN, Warszawa 1978
- [17] MALINOWSKI, G., *Logika ogólna*, PWN, Warszawa 2010.
- [18] MORSE A. P., *A Theory Of Sets*, Academic Press, Orlando 1986
- [19] NOWAK M., *Elementy teorii mnogości w Mizarze*, Łódź 1990.
- [20] PORĘBSKA M., W. SUCHOŃ, *Elementarne wprowadzenie w logikę formalną*, PWN, Warszawa 1991.
- [21] SCHOENFIELD J. R., 'Axioms of set theory', w: J. Barwise (red.) *Handbook of Mathematical Logic*, North Holland, Amsterdam 1977
- [22] SŁUPECKI J., L. BORKOWSKI, *Elementy logiki matematycznej i teorii mnogości*, PWN, Warszawa 1984.
- [23] TAKEUTI G., W. M. ZARING, *Introduction to Axiomatic Set Theory*, Springer-Verlag, New York 1971
- [24] TRACZYK T., *Wstęp do teorii algebr Boole'a*, PWN, Warszawa 1970