

Krzysztof A. Wieczorek

Logika dla opornych

**Wszystko co powinniście wiedzieć o logice,
ale nie uważaliście na zajęciach**

Ilustracje: Barbara Wieczorek

copyright © 2002 by Krzysztof A. Wieczorek

WSTĘP

Celem tego podręcznika nie jest systematyczny wykład logiki. Książek takich jest już wystarczająco dużo, więc osoba głębiej zainteresowana tym przedmiotem na pewno nie będzie miała kłopotu ze znalezieniem czegoś odpowiedniego dla siebie. Niniejsza pozycja przeznaczona jest przede wszystkim dla tych, którzy pobieżnie zetknąwszy się z logiką, na przykład jako z przedmiotem wykładanym podczas krótkiego kursu na wyższej uczelni, z przerażeniem stwierdzili, że nic z tego nie rozumieją. Przyświeca mi cel pokazania takim osobom, że wbrew pozorom logika wcale nie jest taka trudna, jak by się to mogło początkowo wydawać, a jej nauka nie musi przypominać drogi przez mękę.

Większość tradycyjnych podręczników logiki najeżona jest technicznymi terminami, sucho brzmiącymi definicjami i twierdzeniami oraz skomplikowanymi wzorami. Brakuje im natomiast przykładów ilustrujących zawarty materiał teoretyczny i wyjaśniających bardziej złożone zagadnienia w sposób zrozumiały dla osób uważających się za „humanistów”, a nie „ścisłowców”. Sytuacja ta sprawia, że po zapoznaniu się z treścią takiego podręcznika lub po wysłuchaniu wykładu opracowanego na jego podstawie, adept logiki ma trudności z rozwiązaniem nawet bardzo prostych zdań umieszczanych na końcach rozdziałów lub w specjalnych zbiorach ćwiczeń z logiki. Taki stan rzeczy przyprawia o mdłości i ból głowy zarówno wielu wykładowców logiki zrozpaczonych rzekomą całkowitą niezdolnością do poprawnego myślenia okazywaną przez ich studentów, jak i tych ostatnich, zmuszonych do zaliczenia przedmiotu, z którego niemal nic nie rozumieją.

Doświadczenie zdobyte przeze mnie podczas lat nauczania logiki na różnych kierunkach uniwersyteckich wskazuje jednakże, iż najczęściej nieumiejętność rozwiązywania zadań z logiki nie jest wynikiem jakichkolwiek braków umysłowych studentów ani nawet ich lenistwa, ale po prostu przerażeniem wywoływanym przez gąszcz niezrozumiałych dla nich wzorów, twierdzeń i definicji. Panika ta widoczna jest szczególnie u osób obdarzonych bardziej humanistycznym typem umysłowości, alergicznie reagujących na wszystko, co kojarzy im się z matematyką.

Można oczywiście ubolewać nad tym, że tak wielu młodych ludzi nie chce pokonać w sobie uprzedzeń do logiki i zmuszać ich „dla ich dobra” do przyswajania tej wiedzy w tradycyjnej formie. Czy ma to jednak większy sens? Da się oczywiście sprawić, że uczeń poświęci tydzień czasu przed egzaminem (często wspomagając się przy tym różnego rodzaju chemicznymi „środkami dopingującymi”) na pamięciowe wykucie kilkudziesięciu twierdzeń i

praw, a następnie nauczy się ich mechanicznego stosowania. Nie zmieni to jednak faktu, iż student taki w dalszym ciągu nie będzie rozumiał istoty tego, co robi, ani jaki jest właściwie cel wykonywanych przez niego operacji.

Żyjemy obecnie w czasach, w których liczy się przede wszystkim szybkość i skuteczność działania. Większość ludzi nie ma czasu na zgłębianie teoretycznych podstaw jakiejś dziedziny – interesują ich przede wszystkim praktyczne umiejętności, sposób w jaki teoria przejawia się w praktyce. Przykładowo użytkownik komputera nie musi znać zasad jego budowy ani języków pisania programów. Wystarczy mu, że potrafi kopiować pliki na dyskietkę, włączyć kilka ulubionych programów, wie, co zrobić, gdy komputer się zawiesi, a w razie większych komplikacji ma telefon do kogoś, kto zna się na tym lepiej. Również ucząc się obsługi potrzebnych programów, przeciętny człowiek nie musi korzystać ze specjalistycznych książek dla informatyków wyjaśniających wszelkie możliwe szczegóły techniczne. Wystarczy, że sięgnie on do popularnego podręcznika z serii „dla opornych”. Książki takie wiele spraw znacznie upraszczają, wiele trudnych problemów pomijają, ograniczając się do tego, co najważniejsze. Jeżeli jednak coś można ułatwić, przedstawić w sposób zrozumiały, nawet kosztem pewnej trywializacji, to dlaczego tego nie zrobić? Nie wszystko co ważne, musi być od razu trudne i opisane technicznym językiem.

Z podobnym nastawieniem pisana jest niniejsza książka. Wiele spraw jest w niej uproszczonych. Starłem się posługiwać zrozumiałym językiem, unikając gdzie tylko się da technicznego żargonu. Może to sprawić, że przedstawiona w ten sposób logika wyda się komuś nadmiernie spłycona. Być może jest tak faktycznie, jednak, podkreślam to raz jeszcze, celem tego podręcznika nie jest systematyczny wykład logiki, ale przede wszystkim pomoc w opanowaniu tego przedmiotu dla tych, którym wydaje się on niemal całkowicie niezrozumiały. Gdy stwierdzą oni, że logika nie jest wcale tak trudna, jak im się to początkowo wydawało, sięgną oni być może po podręcznik głębiej traktujący temat.

Jednocześnie książka ta może stać się zachętą do zainteresowania się logiką przez osoby, które nigdy się z tym przedmiotem nie zetknęły. Korzystając z zawartych tu przykładów, czytając odpowiedzi na pytania zwykle zadawane przez początkujących, widząc często popełniane błędy, mogą one przyswoić sobie podstawy logiki samodzielnie, bez pomocy nauczyciela.

Semestralny kurs logiki na wielu uniwersyteckich kierunkach trwa zwykle 60 godzin lekcyjnych. Jednakże zdarzają się kursy ograniczone do 30, 15, a nawet 10 godzin. W takim czasie doprawdy trudno jest nauczyć kogoś logiki. Można co najwyżej pokazać zarys tego przedmiotu. Studentom uczestniczącym w takich, z różnych względów skróconych, kursach,

niniejsza książka powinna przynieść szczególne korzyści. Może ona im pomóc w zrozumieniu tego, na wyjaśnienie czego nie starczyło czasu na wykładach lub ćwiczeniach, a jednocześnie pokazać, jak należy rozwiązywać zadania spotykane często na egzaminach i kolokwiach.

Jak korzystać z książki?

Celem tego podręcznika jest przede wszystkim wyrobienie u Ciebie, drogi Czytelniku, umiejętności rozwiązywania zadań spotykanych w standardowych podręcznikach do logiki. Najczęściej jednak rozwiązania przykładów wymagają pewnej podstawy teoretycznej. Potrzebna teoria, w formie bardzo okrojonej i uproszczonej, wprowadzana jest zwykle w początkowych partiach każdego rozdziału. Ponieważ, z uwagi na tę skróconość, nie wszystko w części teoretycznej może wydać Ci się od razu zrozumiałe, proponuję przeczytanie tych paragrafów dwa razy: na początku dla zapoznania się z podstawowymi pojęciami, a następnie po przerobieniu części praktycznej, w celu dokładniejszego zrozumienia i utrwalenia sobie przerobionego materiału. Jestem przekonany, że po takim powtórnym przeczytaniu fragmentów teorii w pełni jasne staną się sprawy, które początkowo wydawały się nie do końca klarowne.

W części teoretycznej przedstawiane są tylko konieczne podstawy – tyle, aby można było przystąpić do rozwiązywania pierwszych zadań. Wiele dalszych problemów omawianych jest później – gdy pojawiają się przy okazji praktycznych zadań. Rozwiązując te zadania, zapoznasz się, niejako mimochodem, z kolejnymi elementami teorii. Niektóre wiadomości teoretyczne zawarte są również w sekcjach „Uwaga na błędy” oraz „Często zadawane pytania”. Zawarte w książce przykłady uszeregowane są w kolejności od najprostszych do coraz trudniejszych. Umiejętności nabyte przy rozwiązywaniu jednych wykorzystywane są często w kolejnych zadaniach. Dobrane są one również w taki sposób, aby każdy z nich wskazywał jakiś inny problem techniczny lub teoretyczny.

Jeśli chcesz nauczyć się samodzielnego rozwiązywania zadań, nie powinieneś ograniczać się do śledzenia rozwiązań podanych przeze mnie krok po kroku. Doświadczenie wskazuje, że w takim momencie wydają się one banalnie proste; problemy pojawiają się jednak, gdy podobne rozwiązanie trzeba przedstawić samodzielnie. Dlatego po przerobieniu każdego działu spróbuj przepisać treść przykładów na osobną kartkę, rozwiąż je samodzielnie i dopiero wtedy porównaj wynik z podręcznikiem. W wielu wypadkach zobaczysz wtedy, iż nawet w pozornie prostych przykładach bardzo łatwo popełnić błędy. Nie powinno to jednak

powodować u nikogo większego niepokoju. Nic bowiem tak nie uczy, jak zrobienie błędu, dostrzeżenie go i następnie poprawienie. Tak więc – w dłuższej perspektywie – popełnianie błędów w początkowej fazie nauki jest nawet korzystne.

Z uwagi na to, że książka ta składa się przede wszystkim z przykładów, może ona posłużyć jako swojego rodzaju zbiór zdań z logiki. Osoby lepiej znające ten przedmiot nie muszą czytać drobiazgowych omówień poszczególnych ćwiczeń, i mogą od razu przystąpić do ich samodzielnego rozwiązania. Objaśnienia mogą się im przydać w sytuacjach, gdyby okazało się, że otrzymały w którymś miejscu nieprawidłowy wynik.

W niektórych miejscach tekstu tłustym drukiem wyróżnione zostały pojęcia szczególnie istotne w nauce logiki. Znaczenie tych pojęć powinieneś sobie przyswoić i dobrze zapamiętać. Definicje tych wyrażen i czasem dotyczące ich wyjaśnienia zawarte są również w znajdującym się na końcu książki słowniczku.



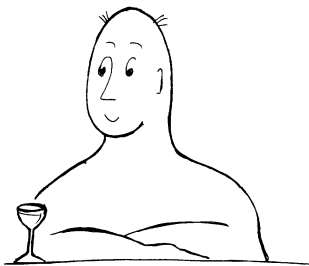
Rozdział I

KLASYCZNY RACHUNEK ZDAŃ.

Klasyczny rachunek zdań (w skrócie KRZ) jest jednym z najprostszych systemów logiki formalnej. W praktyce może on służyć do sprawdzania poprawności **wnioskowań**, czyli takich procesów myślowych, podczas których na podstawie uznania za prawdziwe jednych zdań (przesłanek) dochodzimy do uznania kolejnego zdania (wniosku). Dzięki znajomości KRZ każdy może się łatwo przekonać, że na przykład z takich przesłanek jak: *Jeśli na imprezie był Zdzisiek i Wacek, to impreza się nie udała* oraz *Impreza udała się* można wywnioskować iż: *Na imprezie nie było Zdziśka lub Wacka*. Posługując się metodami KRZ można również stwierdzić, iż nie rozumie poprawnie ten, kto z przesłanek: *Jeśli Wacek dostał wypłatę to jest w barze lub u Zdziśka* oraz *Wacek jest w barze* dochodzi do konkluzji: *Wacek dostał wypłatę*.

1.1. SCHEMATY ZDAŃ.

1.1.1. ŁYK TEORII.



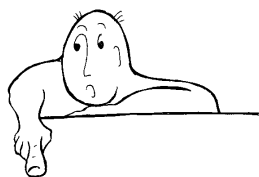
ŁYK TEORII

Pierwszą czynnością, jaką należy przećwiczyć rozpoczynając naukę klasycznego rachunku zdań, jest budowanie logicznych schematów zdań. Budowanie takich schematów przyrównać można do przekładu wyrażen „normalnego” języka, jakim ludzie posługują się na co dzień, na język logiki, w którym logicy sprawdzają poprawność danego rozumowania.

Termin „**zdanie**” oznacza w logice tylko i wyłącznie zdanie oznajmujące i schematy tylko takich zdań będziemy budować. Schematy pokazują nam położenie w zdaniach języka naturalnego zwrotów szczególnie istotnych z punktu widzenia logiki – niektórych z tak zwanych **stałych logicznych**: *nieprawda że, i, lub, jeśli... to, wtedy i tylko wtedy*. Zwroty te noszą w logice nazwy **negacji, koniunkcji, alternatywy, implikacji oraz równoważności** i będą w schematach zastępowane odpowiednimi symbolami: $\sim$ (negacja), $\wedge$ (koniunkcja), $\vee$ (alternatywa), $\rightarrow$ (implikacja), $\equiv$ (równoważność). Wymienione zwroty są (przynajmniej w takich znaczeniach, w jakich przyjmuje je logika)

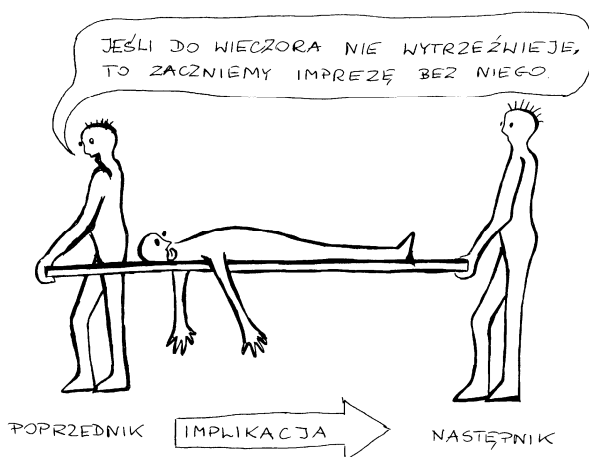
spójnikami łączącymi zdania, dlatego nazywamy je **spójnikami logicznymi**. Zdania proste, łączone przez spójniki logiczne zastępować będziemy w schematach literami: p, q, r, s, t... itd. Litery p, q, r... nazywamy **zmiennymi zdaniowymi** (ponieważ zastępują zdania języka naturalnego). Do budowy schematów będziemy też często używali nawiasów, które pełnią rolę podobną do znaków przestankowych w piśmie – pokazują jak schemat należy odczytać, które jego części wiążą się ze sobą ściślej, a które luźniej. Rola nawiasów stanie się jaśniejsza po przerobieniu kilku zadań praktycznych. Przykładowe schematy logiczne zdań mogą wyglądać następująco: $p \rightarrow q$, $\sim (p \wedge q)$, $p \vee (r \rightarrow \sim s)$, $[p \equiv (q \rightarrow r)] \wedge (s \rightarrow z)$.

Zdania wiązane przez spójniki logiczne nazywamy **członami** tych spójników. Człony równoważności niektórzy nazywają **stronami równoważności**, natomiast zdania wiązane przez implikację określamy najczęściej mianem **poprzednika** i **następnika** implikacji. Jak łatwo się domyśleć, poprzednik to zdanie znajdujące się przez „strzałką” implikacji, a następnik – zdanie po niej.



Uwaga na błędy!

Częstym błędem popełnianym przez studentów jest nazywanie poprzednikiem i następnikiem zdań łączonych przez spójniki inne niż implikacja. Powtórzmy więc jeszcze raz: **poprzednik i następnik występują wyłącznie przy implikacji.**

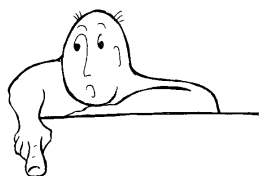


Mianem negacji, koniunkcji, alternatywy, implikacji oraz równoważności określa się w logice nie tylko spójniki, ale również całe zdania przy ich pomocy tworzone. Na przykład wyrażenie **Jeśli** Agnieszka zobaczy Ryszarda w tym stanie, **to** będzie rozczarowana nazywamy

zdaniami implikacyjnymi lub po prostu implikacją; zdanie *Ryszard wykazał się dużym sprytem* **lub** *po prostu dopisało mu szczęście* nazywamy alternatywą, itd.

Większość spójników (poza negacją) to tak zwane spójniki dwuargumentowe, co oznacza, że łączą one dwa zdania. Niekoniecznie muszą być to jednak zdania proste, równie dobrze mogą być to ujęte w nawiasy złożone wyrażenia. Na przykład w schemacie $p \vee q$ członami alternatywy są zdania proste oznaczane przez p i q . Jednakże członami koniunkcji w wyrażeniu $(p \rightarrow q) \wedge (r \vee s)$ są już wzięte w nawiasy zdania złożone: $(p \rightarrow q)$ oraz $(r \vee s)$. Stronami równoważności w kolejnym schemacie są jeszcze dłuższe zdania (ujęte w nawiasy klamrowy i kwadratowy) $\{[p \vee (q \rightarrow \sim r)] \wedge s\} \equiv [t \rightarrow (w \wedge z)]$

Wyrażenia łączone przez spójniki dwuargumentowe występują zawsze po obu stronach spójnika. Tak więc prawidłowe są zapisy: $p \rightarrow q$, $p \wedge (q \vee r)$, natomiast nieprawidłowe: $\rightarrow p q$, $p (q \vee r) \wedge$.



Uwaga na błędy!

W prawidłowo zapisanych schematach nie może nigdy zdarzyć się tak, aby występowały obok siebie dwie zmienne zdaniowe nie oddzielone spójnikiem (np. $p \rightarrow q r$), lub dwa spójniki dwuargumentowe (czyli wszystkie oprócz negacji) nie oddzielone zmienną (np. $p \vee \wedge q$)

Negacja jest tak zwanym spójnikiem jednoargumentowym, co oznacza, że nie łączy ona dwóch zdań, lecz wiąże się tylko z jednym. Podobnie jak w przypadku innych spójników nie musi być to zdanie proste, ale może być ujęte w nawias większa całość. W schemacie $\sim p$ negacja odnosi się do prostego zdania p , jednakże w $\sim [(p \rightarrow q) \wedge r]$, neguje ona całe wyrażenie ujęte w nawias kwadratowy.

Spójnik negacji zapisujemy zawsze **przed** wyrażeniem, do którego negacja się odnosi. Prawidłowy jest zatem zapis $\sim p$, natomiast błędny $p \sim$.



DO ZAPAMIĘTANIA:

Poniższa tabelka pokazuje podstawowe znaczenia spójników logicznych oraz prawidłowy sposób, w jaki występują one w schematach.

Nazwa spójnika	Symbol	Podstawowy odpowiednik w języku naturalnym	Przykładowe zastosowanie	
Negacja	$\sim$	nieprawda, że	$\sim p$	$\sim (p \vee q)$
Koniunkcja	$\wedge$	i	$p \wedge q$	$p \wedge (\sim q \equiv r)$
Alternatywa	$\vee$	lub	$p \vee q$	$(p \rightarrow q) \vee (r \wedge \sim s)$
Implikacja	$\rightarrow$	jeśli... to	$p \rightarrow q$	$(p \vee q) \rightarrow \sim r$
Równoważność	$\equiv$	wtedy i tylko wtedy	$p \equiv q$	$(p \wedge \sim q) \equiv (\sim r \rightarrow \sim s)$

1.1.2. PRAKTYKA: BUDOWANIE SCHEMATÓW ZDAŃ JĘZYKA NATURALNEGO.

Jak już wiemy z teorii, schemat ma za zadanie pokazać położenie w zdaniu spójników logicznych. Dlatego pisanie schematu dobrze jest rozpocząć od wytropienia w zdaniu zwrotów odpowiadających poszczególnym spójnikom – *nieprawda że*, *i*, *lub*, *jeśli... to*, *wtedy i tylko wtedy*. Dla ułatwienia sobie dalszej pracy symbole spójników można wtedy zapisać nad tymi zwrotami. Całą resztę badanego wyrażenia stanowić będą łączone przez spójniki zdania proste, które będziemy zastępowali przez zmienne zdaniowe. Symbole tych zmiennych również możemy dla ułatwienia zapisać nad ich odpowiednikami.

Przykład:

$$\begin{array}{ccc} p & \wedge & q \\ \underbrace{\hspace{10em}} & & \underbrace{\hspace{10em}} \\ \text{Zygryd czyści rewolwer} & \text{i} & \text{obmyśla plan zemsty.} \end{array}$$

W zdaniu tym znajdujemy jedno wyrażenie odpowiadające spójnikowi logicznemu – *i*, oraz dwa zdania proste – *Zygryd czyści rewolwer* oraz *(Zygryd) obmyśla plan zemsty*. W tym momencie z łatwością możemy już zapisać właściwy schemat całego zdania: $p \wedge q$.

Niektórzy wykładowcy mogą wymagać, aby po napisaniu schematu objaśnić również, co oznaczają poszczególne zmienne zdaniowe.

W takim wypadku piszemy:

$p \wedge q$,

p – Zygfryd czyści rewolwer, q – Zygfryd obmyśla plan zemsty.

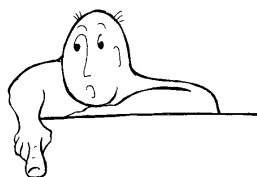
Przykład:

$\underbrace{\hspace{10em}}_p \quad \rightarrow \quad \underbrace{\hspace{10em}}_q$
Jeśli Marian zostanie prezesem, to Leszek straci pracę.

W przypadku implikacji, której składniki „*jeśli*” oraz „*to*” znajdują się w różnych miejscach zdania, strzałkę piszemy zawsze nad *to*. Schemat powyższego zdania to oczywiście

$p \rightarrow q$

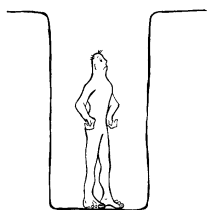
p – Marian zostanie prezesem, q – Leszek straci.



Uwaga na błędy!

Pisząc, co oznaczają poszczególne zmienne zdaniowe nie piszemy już wyrażień, które zastąpiliśmy spójnikami. Często spotykanym błędem, w zadaniach takich jak powyżej, jest napisanie, że p oznacza zdanie *jeśli Marian zostanie prezesem*. Jednakże *jeśli* zostało już przecież zastąpione symbolem „ $\rightarrow$ ”.

Po nabraniu pewnej wprawy można zrezygnować z pisania symboli spójników i zmiennych zdaniowych nad wyrażeniem, którego schemat budujemy. Jednakże trzeba wtedy zachować szczególną ostrożność w przypadku dłuższych zdań – łatwo jest bowiem „zgubić” jakiś spójnik lub zmienną.



1.1.3. UTRUDNIENIA I PUŁAPKI.

Czy to jest zdanie?

Często zdania łączone przez spójniki występują w „skrótowej” postaci.

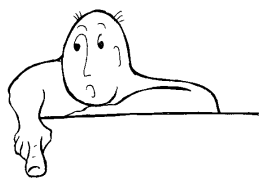
Przykład:

Wiesław zostanie ministrem kultury lub przemysłu ciężkiego.

W zdaniu tym wyrażenie „przemysłu ciężkiego”, to oczywiście skrót zdania „*Wiesław zostanie ministrem przemysłu ciężkiego*” i w taki sposób należy je traktować. Tak więc poprawny schemat zdania wygląda:

$$p \vee q$$

p – *Wiesław zostanie ministrem kultury*, q – *Wiesław zostanie ministrem przemysłu ciężkiego*.



Uwaga na błędy!

Napisanie, że q oznacza „przemysłu ciężkiego”, albo „przemysł ciężki” to duży błąd! Pamiętamy, że q to zmienna zdaniowa, a więc zastępuje ona zdanie. Wyrażania „przemysł ciężki” lub „przemysłu ciężkiego” zdaniami oczywiście nie są.

Czy to jest spójnik logiczny?

Wyrażenia odpowiadające spójnikom logicznym mogą występować w różnej postaci. Przykładowo spójnik alternatywy standardowo uznawany za odpowiadający słowu *lub* może się pojawić np. jako *albo*, czy też *bądź* . Jeszcze gorzej jest z koniunkcją – może się ona pojawić w postaci m.in.: *i*, *oraz*, *a także*, *a*, *lecz*, itd. Implikacji odpowiadają zwroty *jeśli... to*, *o ile... to*, *gdyby...*, *to*. Negacja to *nieprawda że*, *nie jest tak*, *że*, lub często po prostu samo *nie*. Najmniejszy kłopot jest z równoważnością – *wtedy i tylko wtedy*, ewentualnie *zawsze i tylko*

wtedy. Zwroty te są jednak rzadko spotykane — nie używa ich raczej nikt inny poza matematykami i logikami.

Przykład:

Zygmunt jest filozofem a Grzegorz biznesmenem.

$p \wedge q$

p – Zygmunt jest filozofem, q – Grzegorz jest biznesmenem.

Przykład:

Józef nie przyszedł na zebranie.

$\sim p$

p – Józef przyszedł na zebranie.

Przykład:

Albo Antoni jest ślepy, albo zakochany.

$p \vee q$

p – Antoni jest ślepy, q – Antoni jest zakochany.

Zauważmy, że pomimo dwukrotnego pojawienia się słowa „albo” mamy tu do czynienia tylko z jedną alternatywą. Zapis $\vee p \vee q$ nie mógłby się pojawić – nie jest on poprawnym wyrażeniem rachunku zdań.



DO ZAPAMIĘTANIA.

Poniższa tabelka pomoże utrwalić sobie znaczenia i symbole poszczególnych spójników logicznych.

Nazwa spójnika	Symbol	Podstawowy odpowiednik	Inne odpowiedniki
Negacja	$\sim$	nieprawda, że	nie jest tak, że; nie
Koniunkcja	$\wedge$	i	oraz; a także; lecz; a; ale
Alternatywa	$\vee$	lub	albo... albo; bądź

Implikacja	→	jeśli... to....	gdyby.... to...; o ile... to...
Równoważność	≡	wtedy i tylko wtedy	zawsze i tylko wtedy

To nie jest spójnik!

Bywa, że w zdaniu pojawi się wyrażenie pozornie odpowiadające któremuś ze spójników logicznych, ale użyte w innym znaczeniu (nie jako spójnik zdaniowy). W takim wypadku oczywiście nie wolno go zastępować symbolem spójnika.

Przykład:

Stefan i Krystyna są małżeństwem.

W zdaniu tym występuje wyrażenie *i*, ale nie łączy ono zdań. „Stefan” w tym wypadku nie jest zdaniem, ani też jego skrótem. Gdyby ktoś potraktował „Stefan” jako skrót zdania, otrzymałby bezsensowne wyrażenie: *Stefan jest małżeństwem*. Tak więc *Stefan i Krystyna są małżeństwem* to zdanie proste i jego schemat to tylko samo p.

Więcej spójników.

Często w zdaniu występuje więcej niż jeden spójnik. W takim wypadku należy na ogół skorzystać z nawiasów. Nawiasy wskazują, które zdania w sposób naturalny łączą się ze sobą bliżej, tworząc swego rodzaju całość. Jednocześnie nawiasy pokazują, który ze spójników pełni rolę tak zwanego **spójnika głównego**, czyli tego, który niejako spina całe zdanie, łączy ostatecznie wszystkie jego części. W każdym zdaniu złożonym musi być taki spójnik.

Przykład:

Jeżeli przeczytam podręcznik lub będę chodził na wykłady, to bez trudu zdam egzamin.

Prawidłowy schemat tego zdania to:

$$(p \vee q) \rightarrow r$$

Nawiasy pokazują, że zdania oznaczone zmiennymi p oraz q tworzą pewną całość i dopiero wzięte razem stanowią poprzednik implikacji. Implikacja pełni w tym schemacie rolę spójnika głównego – łączy ona wyrażenie w nawiasie oraz zmienną r.

Gdyby ktoś postawił nawiasy w złym miejscu i głównym spójnikiem uczynił alternatywę, czyli schemat wyglądałby: $p \vee (q \rightarrow r)$, to byłby to schemat następującego zdania: *Przeczytam podręcznik lub jeśli będę chodził na wykłady, to bez trudu zdam egzamin*, a więc innego, niż to, którego schemat mieliśmy napisać.

Przykład:

Nieprawda, że jeśli dopadnę drania, to od razu się z nim policzę.

Prawidłowy schemat to: $\sim (p \rightarrow q)$

Nawiasy są konieczne, aby pokazać, iż negacja jest tu spójnikiem głównym i odnosi się do całej implikacji *jeśli dopadnę drania, to od razu się z nim policzę*. Pozostawienie schematu bez nawiasów: $\sim p \rightarrow q$, wskazywało by, że negacja odnosi się tylko do prostego zdania p (głównym spójnikiem stałaby się wtedy implikacja), a więc byłby to schemat zdania *jeśli nie dopadnę drania, to od razu się z nim policzę*.

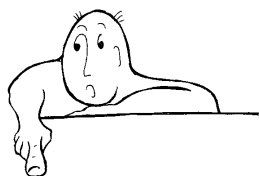
Przykład:

Jeżeli skończę studia to albo wyjadę za granicę, albo zostanę bezrobotnym.

Schemat tego zdania to: $p \rightarrow (q \vee r)$

Treść tego zdania wyraźnie wskazuje, że głównym spójnikiem jest w nim implikacja. Alternatywa została oddana przy pomocy zwrotu „albo...albo”.

Zauważmy, że gdyby zostało użyte słowo „lub”, mogłyby powstać wątpliwości, jaki spójnik pełni rolę głównego; wypowiadając zdanie *Jeżeli skończę studia to wyjadę za granicę lub zostanę bezrobotnym* ktoś mógł mieć bowiem na myśli alternatywę: istnieją dwie możliwości (1) wyjazdu za granicę w przypadku ukończenia studiów lub (2) zostania bezrobotnym (w domyśle – w przypadku nie ukończenia studiów). Wtedy schemat wyglądałby $(p \rightarrow q) \vee r$.



Uwaga na błędy!

Schemat w którym nawiasy nie wskazują jednoznacznie głównego spójnika, jest wieloznaczny (dopuszcza różne możliwości interpretacji). Takie wieloznaczne wyrażenia (np. $p \rightarrow q \vee r$ lub $p \wedge q \rightarrow r$) noszą nazwę **amfibolii**. Napisanie schematu będącego amfibolią traktowane jest jako błąd.

UWAGA!

Autorzy niektórych podręczników wprowadzają różne konwencje pozwalające pomijać nawiasy. Zasady te stwierdzają na przykład, że zasięg implikacji jest większy od zasięgu koniunkcji, a więc schemat $p \rightarrow q \wedge r$ należy domyślnie potraktować, tak jakby wyglądał on $p \rightarrow (q \wedge r)$. Ponieważ jednak nie wszyscy takie konwencje stosują, nie będziemy ich tu wprowadzać. Jedynym wyjątkiem jest stosowana dotąd bez wyjaśnienia, jednakże intuicyjnie oczywista zasada dotycząca negacji, mówiąca że jeśli nie ma nawiasów, to negacja odnosi się tylko do zmiennej, przed którą się znajduje. Na przykład w wyrażeniu $\sim p \vee q$ zanegowane jest tylko zdanie p ; nie ma zatem potrzeby zapisywania schematu w formie: $\sim (p) \vee q$, choć nie byłoby to błędem.

Gdzie dać ten nawias?

Czasami mogą powstać wątpliwości, gdzie należy postawić nawias, nawet gdy zdanie, którego schemat piszemy, na pewno nie jest amfibolią.

Przykład:

Jeżeli spotkam Wojtkę, to o ile nie będzie zbyt późno, to skoczmy na małe piwo.

W powyższym zdaniu mamy dwie implikacje (oddane przez „jeżeli” oraz „o ile”), łączące trzy zdania (w tym jedno zanegowane): $p \rightarrow \sim q \rightarrow r$. W schemacie takim musimy jednak przy pomocy nawiasów określić, która z implikacji stanowi główny spójnik zdania – czy schemat ma wyglądać: $(p \rightarrow \sim q) \rightarrow r$, czy też $p \rightarrow (\sim q \rightarrow r)$. Aby ten problem rozwiązać przyjrzyjmy się bliżej naszemu zdaniu – mówi ono, co się wydarzy, jeśli „spotkam Wojtkę”, a więc poprzednikiem głównej implikacji jest zdanie proste. Natomiast następnikiem sformułowanego w tym zdaniu warunku jest pewna implikacja „o ile nie będzie zbyt późno, skoczmy na małe piwo”. Tak więc mamy do czynienia z implikacją prowadzącą od zdania prostego do kolejnej implikacji, czyli prawidłowy jest schemat:

$$p \rightarrow (\sim q \rightarrow r)$$

To, że ten właśnie schemat jest właściwy, nie dla wszystkich może od razu być jasne. Jeśli ktoś nie jest o tym przekonany, niech spróbuje wypowiedzieć zdanie oparte na schemacie $(p \rightarrow \sim q) \rightarrow r$, wstawiając odpowiednie zdania proste za zmienne. Wyszłoby wtedy coś w rodzaju: „jeżeli jeśli spotkam Wojtkę to nie będzie zbyt późno, to skoczmy na małe piwo”.

Więcej nawiasów.

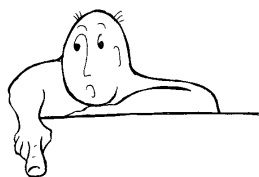
Czasem w zdaniu musi występować większa ilość nawiasów. Wskazują one niejako hierarchię wyrażen.

Przykład:

Nie jest prawdą, że jeśli skończę studia i prestiżowy kurs językowy to znajdę dobrze płatną pracę.

Poprawny schemat tego zdania to: $\sim [(p \wedge q) \rightarrow r]$

Nawias kwadratowy wskazuje, że negacja odnosi się do całego zdania złożonego i pełni rolę spójnika głównego. Natomiast nawias okrągły pokazuje, iż zdania p oraz q dopiero wzięte razem stanowią poprzednik implikacji.



Uwaga na błędy!

Pominięcie w powyższym przykładzie nawiasu kwadratowego: $\sim (p \wedge q) \rightarrow r$ sprawiłoby, że negacja odnosiłaby się jedynie do wyrażenia $(p \wedge q)$; zdanie, z implikacją jako głównym spójnikiem, musiałoby brzmieć wtedy: *Jeżeli nie ukończę studiów i prestiżowego kursu językowego, to znajdę dobrze płatną pracę.* Natomiast pominięcie nawiasu okrągłego: $\sim [p \wedge q \rightarrow r]$ sprawiłoby, że wyrażenie w nawiasie kwadratowym stałoby się amfibolią.

Przykład:

Jeżeli wybory wygra lewica to znów wzrosną podatki i spadnie tempo rozwoju gospodarczego, ale jeśli wygra prawica lub tak zwana centroprawica, to powstanie bardzo słaby rząd i albo będziemy przez cztery lata świadkami gorszących skandali, albo za rok będą nowe wybory.

Schemat tego zdania to: $[p \rightarrow (q \wedge r)] \wedge \{(s \vee t) \rightarrow [u \wedge (w \vee z)]\}$

Głównym spójnikiem zdania jest koniunkcja oddana przy pomocy słowa „ale”. Napisanie schematu pierwszego członu koniunkcji nie powinno sprawić nikomu większych

trudności. Większej uwagi wymaga schemat wyrażenia ujętego w nawias klamrowy. Głównym spójnikiem tej części jest implikacja – zdanie to mówi bowiem, co się wydarzy **jeśli** nastąpi warunek ujęty symbolicznie jako $s \vee t$. Gdy się to stanie, to po pierwsze będziemy mieli do czynienia z sytuacją opisaną przez zdanie u , a po drugie z alternatywą $w \vee z$. Zarówno u , jak i $(w \vee z)$ są więc, wzięte razem, następniem głównej implikacji.

Gdyby ktoś, błędnie, napisał schemat części w nawiasie klamrowym w sposób: $\{[(s \vee t) \rightarrow u] \wedge (w \vee z)\}$, wskazywało by to, że następniem implikacji jest tylko zdanie u , natomiast alternatywa $w \vee z$, stanowi osobną całość, niezależną od warunku $s \vee t$. Analizowane zdanie stwierdza jednak coś innego.

To samo zdanie – ta sama zmienna.

Czasem pewne zdanie proste pojawia się w kilkakrotnie w różnych miejscach zdania złożonego. W takich wypadkach należy wszędzie to zdanie zastąpić tę samą zmienną.

Przykład:

Jeśli Tadeusz zdąży na autobus, to przyjdzie, lub gdyby nie zdążył na autobus, to przełożymy nasze spotkanie.

$$(p \rightarrow q) \vee (\sim p \rightarrow r)$$

p – Tadeusz zdąży na autobus, q – Tadeusz przyjdzie, r – przełożymy nasze spotkanie.

Następnik przed poprzednikiem?

Czasami, na przykład ze względów stylistycznych, w zdaniu języka naturalnego mającego postać implikacji następnik występuje przed poprzednikiem implikacji. Przy pisaniu schematu należy tę kolejność odwrócić.

Przykład:

Populski przegra wybory, jeśli będzie uczciwy wobec konkurentów i nie będzie obiecywał gruszek na wierzbie.

Wprowadzie w zdaniu tym *Populski przegra wybory* pojawia się na samym początku, jest to jednak ewidentnie następnik implikacji. Prawidłowy schemat zatem wygląda następująco:

$$(p \wedge \sim q) \rightarrow r$$

p – Populski będzie uczciwy wobec konkurentów, q – Populski będzie obiecywał gruszki na wierzbie, r – Populski przegra wybory.

Ponieważ w implikacji w powyższym przykładzie nie występuje słowo „to”, dodatkową trudność może zrodzić kwestia postawienia strzałki w odpowiednim miejscu nad zdaniem – jeśli ktoś koniecznie chce to zrobić. W takim wypadku najlepiej postawić ją po zakończeniu całego zdania lub przed jego rozpoczęciem. Można też, przed napisaniem schematu, przeformułować zdanie, tak aby poprzednik i następnik znalazły się na właściwych miejscach: *Jeżeli Populski będzie uczciwy wobec konkurentów i nie będzie obiecywał gruszek na wierzbie, to przegra wybory.*



Warto zapamiętać!

Wątpliwości, co w danym przypadku jest poprzednikiem a co następnikiem, rozwiązać może użyteczna wskazówka, że poprzednikiem jest każdorazowo to, co znajduje się bezpośrednio po słowie „jeżeli” (jeżeli, o ile, gdy itp.). Następnik natomiast może znajdować się albo po poprzedniku oddzielony słowem „to”, albo na samym początku zdania, gdy „to” nie jest obecne.



1.1.4. CZĘSTO ZADAWANE PYTANIA.

Czy pojedynczy symbol zmiennej zdaniowej, na przykład samo p, to już jest schemat zdania?

Tak, schemat nie musi koniecznie zawierać spójników logicznych. Jeżeli w zdaniu nie ma wyrażeń odpowiadających spójnikom, to schemat takiego zdania składa się tylko z jednej zmiennej.

Czy zmienne w schemacie zdania muszą występować w kolejności p, q, r, s, t... itd.?

Nie, nie jest to konieczne. Wprawdzie przyjęło się jako pierwszą zmienną obierać p, a potem q, ale nie jest błędem rozpoczęcie schematu na przykład od r. Jest to co najwyżej mniej eleganckie rozwiązanie.

Czy w każdym schemacie musi być spójnik główny?

Tak, jeśli oczywiście schemat nie składa się jedynie z pojedynczej zmiennej. Schemat w którym nawiasy nie pokazują, który ze spójników jest główny, jest nieprawidłowy, ponieważ nie wiadomo, jak go należy odczytać. Przykładowo $p \wedge q \rightarrow r$ można by odczytać *p i jeśli q to*

r (gdyby głównym spójnikiem była koniunkcja) albo też *jeśli p i q to r* (gdyby głównym spójnikiem miała być implikacja).

Co więcej, jeśli mamy do czynienia ze formułą o znacznym stopniu złożoności, swoje spójniki główne muszą posiadać wszystkie ujęte w nawiasy zdania składowe. Na przykład w schemacie $\{[p \rightarrow (q \wedge r)] \vee s\} \equiv \sim [(s \vee t) \wedge z]$ głównym spójnikiem jest równoważność; Kolejne miejsce w hierarchii spójników zajmują alternatywa (główny spójnik lewej strony równoważności) oraz negacja (główny spójnik prawej strony równoważności). Następnie głównym spójnikiem wyrażenia w kwadratowym nawiasie z lewej strony jest implikacja, a w zanegowanym wyrażeniu w kwadratowym nawiasie z prawej strony – koniunkcja. Pominięcie któregośkolwiek z nawiasów uniemożliwiłoby określenie tych spójników.

Czy da się napisać schemat każdego zdania?

Tak, jeśli oczywiście jest to zdanie oznajmujące (bo tylko takie interesują nas w logice). Należy jednak pamiętać, że jeśli w zdaniu nie ma wyrażen odpowiadających spójnikom logicznym, to schematem tego zdanie będzie tylko „ p ”, choćby zdanie było bardzo długie.

Czy błędem jest „uproszczenie” sobie schematu poprzez pominięcie jakiegoś spójnika? Na przykład zapisanie schematu zdania „Jeśli spotkam Wojtkę lub Mateusza, to pójdziemy na piwo”, jako $p \rightarrow q$, gdzie p zostanie potraktowane jako „spotkam Wojtkę lub Mateusza”, zamiast $(p \vee q) \rightarrow r$?

Nie jest to błąd w ścisłym tego słowa znaczeniu. Czasem faktycznie, z różnych względów, pisze się takie uproszczone schematy. Tym niemniej na ogół, gdy w zadaniu należy napisać schemat zdania, rozumiany jest pod tym pojęciem tak zwany **schemat główny**, czyli zawierający wszystkie spójniki możliwe do wyróżnienia w zdaniu. Tak więc zapisanie schematu uproszczonego może zostać potraktowane jako błąd.

1.2. TABELKI ZERO-JEDYNKOWE I ICH ZASTOSOWANIE.

1.2.1. ŁYK TEORII.



ŁYK TEORII

Tak zwane tabelki zero-jedynkowe służą do określania prawdziwości lub fałszywości zdań zawierających spójniki logiczne. Prawdę lub fałsz nazywamy **wartością logiczną** zdania. W notacji logicznej symbol 0 oznacza zdanie fałszywe, natomiast 1 zdanie prawdziwe. Wartość logiczną zdania prostego zapisujemy zwykle pod (lub nad) odpowiadającą mu zmienną, wartość logiczną zdania złożonego zapisujemy pod głównym spójnikiem tego zdania.

Negacja

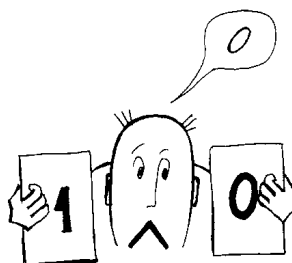
$\sim$	p
1	0
0	1

Tabela dla negacji ukazuje dość oczywistą prawidłowość, że negacja zmienia wartość logiczną zdania.

Gdy weźmiemy dowolne zdanie fałszywe (oznaczone – 0) i następnie zanegujemy je, to otrzymamy zdanie prawdziwe (oznaczone 1). Na przykład: Gdańsk jest stolicą Polski – fałsz, *Gdańsk nie jest stolicą Polski* – prawda. Natomiast poprzedzenie negacją zdania prawdziwego czyni z niego zdanie fałszywe. Na przykład: *Kraków leży nad Wisłą* – prawda, *Kraków nie leży nad Wisłą* – fałsz.

Koniunkcja

p	$\wedge$	q
0	0	0
0	0	1
1	0	0
1	1	1



KONIUNKCJA

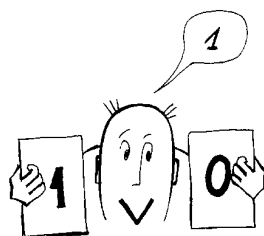
Jestem aż w trzech przypadkach fałszywa.
Dopiero, gdy oba człony są prawdziwe,
staje się prawdziwa.

Tabela dla koniunkcji pokazuje, że gdy przynajmniej jeden z członów tworzących koniunkcję jest fałszywy, to całe zdanie złożone też jest fałszywe. Aby zdanie było prawdziwe, prawdziwe muszą być oba człony koniunkcji.

Przykładowo, gdy ktoś stwierdza: *W tym roku byłem w Afryce i Australii*, a my skądinąd wiemy, że nie był on ani w Afryce, ani w Australii (oba człony koniunkcji fałszywe – pierwszy rząd w tabeli), to oczywiście całą wypowiedź należy uznać za fałszywą. Podobnie, gdyby okazało się, że wypowiadający zdanie był tylko w jednym z wymienionych miejsc (drugi i trzeci rząd w tabeli – jeden człon koniunkcji prawdziwy, a drugi fałszywy), to cała wypowiedź w dalszym ciągu pozostaje fałszywa. Dopiero w przypadku prawdziwości obu członów koniunkcji (ostatni wiersz tabeli) całe zdanie złożone należy uznać za prawdziwe.

Alternatywa

p	∨	q
0	0	0
0	1	1
1	1	0
1	1	1



ALTERNATYWA

Jestem prawdziwa aż w trzech przypadkach!
Dopiero, gdy oba człony są fałszywe, staję się fałszywa.

Tabela dla alternatywy pokazuje, iż jest ona zdaniem fałszywym tylko w jednym przypadku – gdy oba jej człony są fałszywe. Gdy przynajmniej jeden człon jest zdaniem prawdziwym – prawdziwa jest również cała alternatywa.

Gdy w prognozie pogody słyszymy, że *będzie padał deszcz lub śnieg*, tymczasem następnego dnia nie będzie ani deszczu, ani śniegu (czyli oba człony alternatywy okażą się zdaniami fałszywymi), to całą prognozę należy uznać za fałszywą. Gdy jednak spadnie sam deszcz (pierwszy człon prawdziwy), sam śnieg (drugi człon prawdziwy), lub też i śnieg i deszcz (oba człony alternatywy prawdziwe), zdanie mówiące że *będzie padał deszcz lub śnieg* okazuje się prawdziwe.

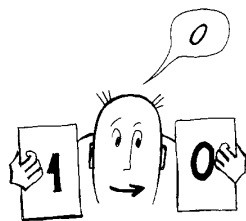
Uwaga na marginesie.

Jeżeli ktoś ma wątpliwości co do ostatniego wiersza tabelki dla alternatywy, to są to wątpliwości całkowicie uzasadnione. Tabela ta ilustruje bowiem tylko jedno ze znaczeń, w jakim alternatywa jest używana. Znaczenie to można opisać zwrotem *przynajmniej jedno z dwojga*; czy też *jedno lub drugie lub oba naraz* – jest to tak zwana alternatywa nierozłączna. W języku potocznym alternatywy używamy też często w znaczeniu *dokładnie jedno z dwojga*; *albo tylko jedno, albo tylko drugie* (alternatywa rozłączna). W takim rozumieniu

alternatywy w ostatnim wierszu tabelki powinno pojawić się zero. W niektórych systemach logicznych oba znaczenia alternatywy są starannie rozróżniane (jest to szczególnie istotne dla prawników) i oddawane przy pomocy różnych symboli (najczęściej $\perp$ – dla alternatywy rozłącznej).

Implikacja

p	$\rightarrow$	q
0	1	0
0	1	1
1	0	0
1	1	1



IMPLIKACJA

Jestem prawdziwa aż w trzech przypadkach!
Staje się fałszywa tylko wtedy, gdy poprzednik jest prawdziwy,
a następnik fałszywy.

Z tabelki dla implikacji możemy dowiedzieć się, że zdanie, którego głównym spójnikiem jest *jeśli... to* może być fałszywe tylko w jednym wypadku, mianowicie, gdy jego poprzednik jest prawdziwy, natomiast następnik fałszywy.

Jako przykładem ilustrującym tabelkę dla implikacji posłużymy się zdaniem wypowiedzianym przez ojca do dziecka: *Jeśli zdasz egzamin, to dostaniesz komputer*. Gdy następnie dziecko nie zdaje egzaminu i komputera nie dostaje (pierwszy wiersz tabeli – poprzednik i następnik implikacji fałszywe) lub gdy zdaje egzamin i dostaje komputer (ostatni wiersz tabeli – poprzednik i następnik implikacji prawdziwe), to nie powinno być wątpliwości, że obietnica ojca okazała się prawdziwa. Gdy natomiast dziecko zdaje egzamin, a jednak komputera nie dostaje (trzeci wiersz tabeli – poprzednik implikacji prawdziwy, a następnik fałszywy), należy wówczas uznać, że ojciec skłamał składając swoją obietnicę.

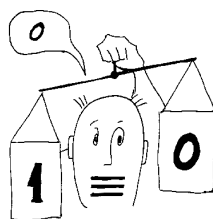
Pewne kontrowersje może budzić uznanie za prawdziwego zdania w przypadku, gdy poprzednik implikacji jest fałszywy, natomiast następnik prawdziwy (drugi wiersz tabeli), czyli w naszym przykładzie, gdy dziecko wprowadzi nie zdało egzaminu, a mimo to dostało komputer. Zauważmy jednak, że wbrew pozorom ojciec nie łamie wcale w takim przypadku obietnicy dania komputera po zdaniu egzaminie – nie powiedział on bowiem, że jest to jedyny przypadek, gdy dziecko może otrzymać komputer. Powiedzenie, że *jeśli zdasz egzamin, to dostaniesz komputer*, nie wyklucza wcale, że dziecko może również dostać komputer z innej okazji, na przykład na urodziny.

Powyższe wytłumaczenie drugiego wiersza tabelki dla implikacji może się wydawać nieco naciągane, a jest tak dlatego, że w języku potocznym często wypowiadamy zdania typu *jeśli... to* rozumiejąc przez nie *wtedy i tylko wtedy* (którego to zwrotu nikt raczej nie używa).

Jak za chwilę zobaczymy, tabelka dla równoważności różni się od tabelki implikacji tylko tym jednym kontrowersyjnym przypadkiem.

Równoważność

p	$\equiv$	q
0	1	0
0	0	1
1	0	0
1	1	1



RÓWNOWAŻNOŚĆ
Jestem fałszywa, gdy moje strony są różne.
Gdy strony się zgadzają, jestem prawdziwa.

Z uwagi na rzadkie występowanie w języku potocznym spójnika *wtedy i tylko wtedy* trudno jest wskazać przykłady obrazujące prawomocność powyższej tabelki.

Najłatwiejszym sposobem na zapamiętanie tabelki dla równoważności wydaje się skojarzenie, że aby równoważność była prawdziwa, obie jej strony muszą być „równoważne” sobie, to znaczy albo obie fałszywe (pierwszy wiersz tabeli), albo oba prawdziwe (ostatni wiersz). Gdy natomiast strony równoważności posiadają różne wartości logiczne (drugi i trzeci wiersz tabeli), cała równoważność jest fałszywa.



DO ZAPAMIĘTANIA:

Obecnie, dla utrwalenia, tabelki dla wszystkich spójników dwuargumentowych przedstawimy w formie skróconej „ściagi”:

p	q	$\wedge$	$\vee$	$\rightarrow$	$\equiv$
0	0	0	0	1	1
0	1	0	1	1	0
1	0	0	1	0	0
1	1	1	1	1	1

Znajomość powyższej tabelki jest konieczna do rozwiązywania zadań z zakresu rachunku zdań. Najlepiej więc od razu nauczyć się jej na pamięć. Wymaga to niestety pewnego wysiłku i czasu, ale bez tego rozwiązywanie dalszych przykładów będzie niemożliwe.

1.2.2. PRAKTYKA: ZASTOSOWANIE TABELEK.

Dzięki poznanym tabelkom możemy zawsze stwierdzić czy prawdziwe, czy też fałszywe jest zdanie złożone (niezależnie od jego długości), gdy tylko znamy wartości logiczne wchodzących w jego skład zdań prostych.

Przypomnijmy, że wartość logiczna całego zdania złożonego będzie zawsze zobrazowana symbolem 0 lub 1 znajdującym się pod głównym spójnikiem zdania (czyli spójnikiem ostatecznie wiążącym wszystkie elementy zdania).

Przykład:

Obliczymy wartość logiczną zdania $p \rightarrow (q \wedge r)$ przy założeniu, że zmienne p i q reprezentują zdanie prawdziwe, natomiast zmienna r – zdanie fałszywe, a więc zachodzi sytuacja:

$$\begin{array}{ccccc} p & \rightarrow & (q & \wedge & r) \\ 1 & & 1 & & 0 \end{array}$$

Wartość logiczną całego zdania reprezentować będzie symbol umieszczony pod głównym spójnikiem schematu, a więc pod implikacją. Aby określić wartość implikacji musimy znać wartość jej poprzednika i następnika. Poprzednikiem implikacji jest tu zdanie proste p i jego wartość mamy już podaną. Natomiast następnikiem jest tu całe ujęte w nawias wyrażenie $(p \wedge q)$, którego wartość musimy dopiero obliczyć. Robimy to korzystając z tabelki dla koniunkcji, a dokładniej jej wiersza mówiącego, że gdy pierwszy człon koniunkcji jest prawdziwy, a drugi fałszywy, to cała koniunkcja jest fałszywa. Mamy zatem sytuację:

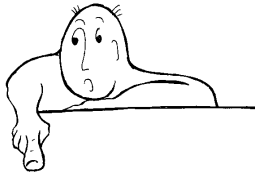
$$\begin{array}{ccccc} p & \rightarrow & (q & \wedge & r) \\ 1 & & \underline{1} & & \underline{0} \end{array}$$

(symbole podkreślone pokazują wartości, z których skorzystaliśmy do obliczeń)

W tym momencie możemy już określić wartość logiczną całego zdania, sprawdzając w tabelce jaką wartość przyjmuje implikacja, której poprzednik jest prawdziwy, a następnik fałszywy.

$$\begin{array}{ccccc} p & \rightarrow & (q & \wedge & r) \\ \underline{1} & & \underline{0} & & \underline{0} \end{array}$$

Ostatecznie widzimy, że całe zdanie jest fałszywe, ponieważ pod głównym spójnikiem otrzymaliśmy wartość 0.



Uwaga na błędy!

Częstym błędem popełnianym przez początkujących jest niedostrzeżenie, że zdanie wiązane przez spójnik jest złożone (np. następnik implikacji w powyższym przykładzie). Osoba popełniająca taki błąd może myśleć, że ostateczny wynik należy obliczyć biorąc pod uwagę p jako poprzednik implikacji, a samo q jako jej następnik, a więc:

$$p \rightarrow (q \wedge r)$$

$$\underline{1} \quad 1 \quad \underline{1} \quad 0 \quad 0$$

ŹLE!!!

Nie wolno tak jednak postępować w żadnym wypadku, ponieważ następnikiem implikacji jest całe wyrażenie ujęte w nawiasie, którego wartość znajduje się pod jego głównym spójnikiem, a więc koniunkcją.

Przykład:

Obliczymy teraz wartość logiczną zdania $(p \rightarrow q) \vee \sim r$, przy założeniach: $p = 1$, $q = 0$, $r = 0$, a więc:

$$\begin{array}{ccc} (p \rightarrow q) \vee \sim r \\ 1 \quad 0 \quad 0 \end{array}$$

W tym przypadku głównym spójnikiem jest alternatywa. Oba jej członów stanowią zdania złożone ($p \rightarrow q$ oraz $\sim r$), których wartości należy obliczyć najpierw. Korzystamy do tego z tabelki dla implikacji oraz dla negacji.

$$\begin{array}{ccc} (p \rightarrow q) \vee \sim r \\ \underline{1} \quad 0 \quad \underline{0} \quad 0 \end{array}$$

$$\begin{array}{ccc} (p \rightarrow q) \vee \sim r \\ 1 \quad 0 \quad 0 \quad 1 \quad \underline{0} \end{array}$$

Gdy znamy wartości logiczne obu członów alternatywy, możemy obliczyć ostateczny wynik. Czynimy to korzystając z tabelki dla alternatywy i biorąc pod uwagę wartości otrzymane pod implikacją oraz negacją, czyli głównymi spójnikami obu członów alternatywy.

$$\begin{array}{ccc} (p \rightarrow q) \vee \sim r \\ 1 \quad \underline{0} \quad 0 \quad 1 \quad \underline{1} \quad 0 \end{array}$$

Przykład:

Obliczymy wartość logiczną zdania: $\sim(p \wedge q) \equiv (\sim r \rightarrow \sim s)$ przy założeniach: $p = 1$, $q = 0$, $r = 1$, $s = 0$, a więc:

$$\sim(p \wedge q) \equiv (\sim r \rightarrow \sim s)$$

$$\begin{array}{cccc} 1 & 0 & 1 & 0 \end{array}$$

Głównym spójnikiem jest tu oczywiście równoważność. Obliczanie wartości jej stron rozpocząć musimy od obliczenia wartości koniunkcji w pierwszym nawiasie oraz negacji zdań prostych w drugim.

$$\sim(p \wedge q) \equiv (\sim r \rightarrow \sim s)$$

$$\begin{array}{cccc} \underline{1} & 0 & \underline{0} & \quad 1 \quad 0 \end{array}$$

$$\sim(p \wedge q) \equiv (\sim r \rightarrow \sim s)$$

$$\begin{array}{cccc} 1 & 0 & 0 & \quad 0 \underline{1} \quad 1 \underline{0} \end{array}$$

Następnie możemy określić wartość implikacji w drugim nawiasie, biorąc pod uwagę wartości otrzymane pod negacją r oraz negacją s (ponieważ poprzednikiem i następnikiem implikacji są zdania złożone $\sim r$ i $\sim s$):

$$\sim(p \wedge q) \equiv (\sim r \rightarrow \sim s)$$

$$\begin{array}{cccc} 1 & 0 & 0 & \quad \underline{0} \ 1 \ 1 \ \underline{1} \ 0 \end{array}$$

W tym momencie nie możemy jeszcze przystąpić do określenia wartości logicznej równoważności, ponieważ nie została obliczona do końca wartość jej lewej strony. Pierwszy człon równoważności to bowiem nie sama koniunkcja $(p \wedge q)$, ale dopiero negacja tej koniunkcji. Negacja jest tu głównym spójnikiem (dopiero ona spina koniunkcję w całość), musimy więc najpierw obliczyć wartość negacji:

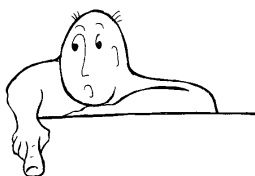
$$\sim(p \wedge q) \equiv (\sim r \rightarrow \sim s)$$

$$\begin{array}{cccc} 1 & 1 & \underline{0} & \quad 0 \underline{1} \ 1 \ 1 \ 0 \end{array}$$

Dopiero teraz możemy określić wartość całego zdania:

$$\sim(p \wedge q) \equiv (\sim r \rightarrow \sim s)$$

$$\begin{array}{cccc} \underline{1} & 1 & \underline{0} & \quad 0 \ 1 \ 1 \ \underline{1} \ 1 \ 0 \end{array}$$



Uwaga na błędy!

Jeśli negacja znajduje się przed nawiasem (jak w lewej stronie równoważności w przykładzie powyżej), to odnosi się ona do całego zdania w nawiasie, a nie tylko do jego pierwszego członu. Aby poznać wartość tej negacji (a zarazem całego zdania,

ponieważ negacja jest jego głównym spójnikiem) bierzemy pod uwagę główny spójnik wyrażenia w nawiasie, a więc:

$$\sim (p \wedge q) \\ \mathbf{1} \ \mathbf{1} \ \underline{\mathbf{0}} \ \mathbf{0} \qquad \text{DOBRZE}$$

a nie:

$$\sim (p \wedge q) \\ \mathbf{0} \ \underline{\mathbf{1}} \ \mathbf{0} \ \mathbf{0} \qquad \text{ŹLE!!!}$$

Przykład:

Obliczymy wartość formuły $[(p \equiv \sim q) \vee \sim r] \wedge \sim (\sim s \rightarrow z)$ przy założeniu, że zdania reprezentowane przez wszystkie zmienne są prawdziwe, a zatem:

$$[(p \equiv \sim q) \vee \sim r] \wedge \sim (\sim s \rightarrow z) \\ \mathbf{1} \quad \mathbf{1} \quad \mathbf{1} \quad \mathbf{1} \quad \mathbf{1}$$

W schemacie powyższym głównym spójnikiem jest koniunkcja łącząca zdanie w nawiasie kwadratowym z zanegowanym zdaniem w nawiasie okrągłym. W pierwszym kroku musimy obliczyć wartość negacji zdań prostych:

$$[(p \equiv \sim q) \vee \sim r] \wedge \sim (\sim s \rightarrow z) \\ \mathbf{1} \quad \mathbf{0} \ \underline{\mathbf{1}} \quad \mathbf{0} \ \underline{\mathbf{1}} \quad \mathbf{0} \ \underline{\mathbf{1}} \quad \mathbf{1}$$

Teraz możemy obliczyć wartość logiczną równoważności i implikacji w okrągłych nawiasach:

$$[(p \equiv \sim q) \vee \sim r] \wedge \sim (\sim s \rightarrow z) \\ \underline{\mathbf{1}} \ \mathbf{0} \ \underline{\mathbf{0}} \ \mathbf{1} \quad \mathbf{0} \ \mathbf{1} \quad \underline{\mathbf{0}} \ \mathbf{1} \ \mathbf{1} \ \underline{\mathbf{1}}$$

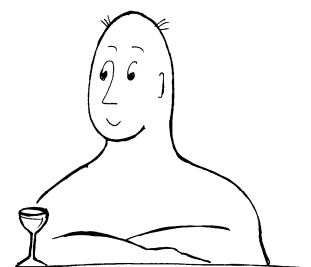
W kolejnym kroku obliczamy wartości logiczne alternatywy oraz negacji formuły w drugim okrągłym nawiasie:

$$[(p \equiv \sim q) \vee \sim r] \wedge \sim (\sim s \rightarrow z) \\ \mathbf{1} \ \underline{\mathbf{0}} \ \mathbf{0} \ \mathbf{1} \ \mathbf{0} \ \mathbf{0} \ \mathbf{1} \quad \mathbf{0} \ \mathbf{0} \ \mathbf{1} \ \underline{\mathbf{1}} \ \mathbf{1}$$

Ponieważ znamy już wartości członów głównej koniunkcji, możemy określić wartość logiczną całego zdania:

$$[(p \equiv \sim q) \vee \sim r] \wedge \sim (\sim s \rightarrow z) \\ \mathbf{1} \ \mathbf{0} \ \mathbf{0} \ \mathbf{1} \ \underline{\mathbf{0}} \ \mathbf{0} \ \mathbf{1} \ \mathbf{0} \ \underline{\mathbf{0}} \ \mathbf{0} \ \mathbf{1} \ \mathbf{1} \ \mathbf{1}$$

1.3. TAUTOLOGIE I KONTRTAUTOLOGIE.



ŁYK TEORII

1.3.1. ŁYK TEORII.

Jak łatwo zauważyć, formuły mogą okazywać się ostatecznie schematami zdań prawdziwych lub fałszywych w zależności od tego, jaką wartość przyjmują zdania proste wchodzące w ich skład. Przykładowo, gdy w schemacie $p \rightarrow \sim q$ za obie zmienne podstawimy zdania prawdziwe, cała implikacja okaże się fałszywa, gdy natomiast podstawimy za p i q zdania fałszywe, implikacja będzie prawdziwa.

Wśród formuł istnieją jednak też takie, które dają zawsze taki sam wynik, bez względu na wartość logiczną składających się na nie zdań prostych. Schematy, które w każdym przypadku dają ostatecznie zdanie prawdziwe nazywamy **tautologiami**; schematy, które generują zawsze zdania fałszywe – **kontrtautologiami**.

1.3.2. PRAKTYKA: SPRAWDZANIE STATUSU FORMUŁ.

Przykład:

Obliczymy wartości logiczne formuły $(p \rightarrow q) \rightarrow (\sim p \vee q)$ przy wszystkich możliwych podstawieniach zdań prawdziwych i fałszywych za zmienne zdaniowe. Ponieważ mamy dwie zmienne, mogą zajść cztery sytuacje:

$(p \rightarrow q) \rightarrow (\sim p \vee q)$			
0	0	0	0
0	1	0	1
1	0	1	0
1	1	1	1

Po obliczeniu wartości wyrażeń w nawiasach, będących poprzednikiem i następnikiem głównej implikacji otrzymamy:

$(p \rightarrow q) \rightarrow (\sim p \vee q)$			
<u>0</u>	<u>1</u>	<u>0</u>	<u>1</u>
<u>0</u>	<u>1</u>	<u>1</u>	<u>0</u>
<u>1</u>	<u>0</u>	<u>0</u>	<u>1</u>
<u>1</u>	<u>1</u>	<u>1</u>	<u>1</u>

Ostateczny wynik w każdym przypadku obliczamy następująco:

$$(p \rightarrow q) \rightarrow (\sim p \vee q)$$

0	<u>1</u>	0	1	1	0	<u>1</u>	0
0	<u>1</u>	1	1	1	0	<u>1</u>	1
1	<u>0</u>	0	1	0	1	<u>0</u>	0
1	<u>1</u>	1	1	0	1	<u>1</u>	1



TAUTOLOGIA

Ponieważ niezależnie od tego jak dobieraliśmy wartości logiczne zmiennych zdaniowych, otrzymaliśmy zawsze zdanie prawdziwe, badany schemat jest tautologią.

Przykład:

Sprawdzimy wartości logiczne formuły $(p \wedge \sim q) \wedge (p \rightarrow q)$ przy wszystkich możliwych podstawieniach zdań prawdziwych i fałszywych za zmienne zdaniowe. Ponieważ jest to dość prosty przykład i jego rozwiązanie zapewne nie sprawi nikomu kłopotu, nie będziemy jego analizy przeprowadzać krok po kroku.

$$(p \wedge \sim q) \wedge (p \rightarrow q)$$

0	<u>0</u>	1	0	0	0	<u>1</u>	0
0	<u>0</u>	0	1	0	0	<u>1</u>	1
1	<u>1</u>	1	0	0	1	<u>0</u>	0
1	<u>0</u>	0	1	0	1	<u>1</u>	1



KONTRTAUTOLOGIA

Badana formuła daje nam wyłącznie zdania fałszywe, niezależnie jakie zdania podstawimy w miejsce zmiennych. Jest to więc kontrtautologia.

Przykład:

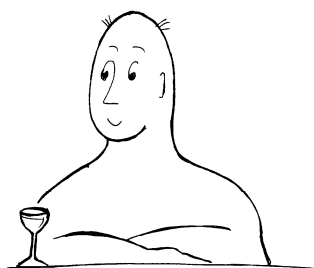
Zbadamy obecnie w podobny sposób formułę:

$$(\sim p \rightarrow \sim q) \vee (p \wedge \sim q)$$

1	0	<u>1</u>	1	0	1	0	<u>0</u>	1	0
1	0	<u>0</u>	0	1	0	0	<u>0</u>	0	1
0	1	<u>1</u>	1	0	1	1	<u>1</u>	1	0
0	1	<u>1</u>	0	1	1	1	<u>0</u>	0	1

W badanej formule w zależności od tego, jakie zdania podstawialiśmy za zmienne otrzymujemy ostatecznie czasem zdanie prawdziwe, a czasem fałszywe. Formuła nie jest więc ani tautologią ani kontrtautologią.

1.4. SKRÓCONA METODA ZEROJEDYNKOWA.



ŁYK TEORII

1.4.1. ŁYK TEORII.

Przedstawiona powyżej metoda badania statusu logicznego formuły (tego, czy jest ona tautologią, kontrtautologią, czy też ani tym, ani tym) nie jest ani najlepsza, ani jedyna. Pokazane przykłady miały za zadanie przede wszystkim usprawnienie umiejętności posługiwania się tabelkami zero-jedynkowymi i wyrobienie sobie ogólnej intuicji czym jest tautologia i kontrtautologia.

Poznana metoda badania formuł, polegająca na sprawdzaniu wszystkich możliwych podstawień zer i jedynek, jest jeszcze możliwa do zaakceptowania w przypadku formuł z dwiema lub ewentualnie trzema zmiennymi zdaniowymi. W przypadku formuł dłuższych staje się na całkowicie niewydolna – na przykład sprawdzenie statusu logicznego formuły mającej cztery zmienne wymagałoby zbadania szesnastu możliwości. Można sobie wyobrazić ile czasu by to zajęło i jak łatwo można by się było w trakcie tych obliczeń pomylić.

Dlatego też do badania formuł wykorzystuje się zwykle tak zwaną skróconą metodę zero-jedynkową (nazywaną też metodą nie wprost), która pozwala na udzielenie odpowiedzi, czy dana formuła jest tautologią lub kontrtautologią często już po rozpatrzeniu jednego przypadku.

Skróconej metodzie badania statusu logicznego formuł poświęcimy znaczną ilość czasu, ponieważ omówimy przy tej okazji różnego rodzaju problemy, jakie mogą się pojawić przy zastosowaniu tabel zero-jedynkowych również przy innych okazjach, na przykład przy sprawdzaniu poprawności wnioskowań.

Ogólna idea metody skróconej.

Wyobraźmy sobie, że chcemy się dowiedzieć, czy formuła jest tautologią, na razie jeszcze przy pomocy „zwykłej” metody polegającej na badaniu wszystkim możliwych podstawień zer i jedynek. Co by można było powiedzieć, gdyby już w pierwszym przypadku pod głównym spójnikiem badanego schematu pojawiło się zero? Oczywiście wiedzielibyśmy, że formuła na pewno już nie jest tautologią, bo przecież tautologia musi za każdym razem wygenerować zdanie prawdziwe. Wiedzę tę uzyskalibyśmy już po rozpatrzeniu jednego

przypadku, więc nie było by potrzeby rozważania kolejnych. Moglibyśmy udzielić w 100% pewnej odpowiedzi – badana formuła nie jest tautologią.

Na powyższej obserwacji opiera się właśnie skrócona metoda zero-jedynkowa. Polega ona bowiem na poszukiwaniu już w pierwszym podejściu takich podstawień zer i jedynek dla zmiennych zdaniowych, aby wykluczyć możliwość, że formuła jest tautologią. Dokładniejszy opis metody skróconej najlepiej przedstawić jest na przykładzie.

1.4.2. PRAKTYKA: WYKORZYSTANIE METODY SKRÓCONEJ.

Przykład:

Zbadamy przy pomocy metody skróconej, czy tautologią jest formuła $(p \rightarrow q) \rightarrow (p \vee q)$.

Gdybyśmy chcieli już w pierwszej linijce stwierdzić, że formuła nie jest tautologią, musielibyśmy znaleźć takie podstawienia zmiennych, aby pod głównym spójnikiem pojawiło się zero. Od tego więc zaczniemy:

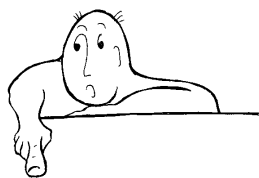
$$(p \rightarrow q) \rightarrow (p \vee q)$$

0

Wiemy zatem, że w poszukiwanym przez nas przypadku 0 musiałoby pojawić się pod spójnikiem implikacji. Gdy spojrzymy teraz do tabelki dla implikacji, zobaczymy, że może być ona fałszywa tylko w jednym przypadku – mianowicie jej poprzednik musi być prawdziwy, a następnik fałszywy. Aby więc w naszym przykładzie 0 mogło się pojawić tam, gdzie je postawiliśmy, prawdziwa musiałaby okazać się implikacja w pierwszym nawiasie, a fałszywa alternatywa w drugim. Otrzymujemy więc:

$$(p \rightarrow q) \rightarrow (p \vee q)$$

1 0 0



Uwaga na błędy!

Niektórzy początkujący adepci logiki widząc w tabelce, że aby implikacja była fałszywa, „p” musi być 1, a „q” – 0, wpisują jedynki pod wszelkimi możliwymi zmiennymi „p” w formule, a zera pod wszystkimi „q”, np.:

$$(p \rightarrow q) \rightarrow (p \vee q)$$

1 0 0 1 0

ŹLE!!!

Jest to oczywiście błąd. Zmienne „p” i „q” z tabelki należy rozumieć umownie, jako dowolny poprzednik i następnik implikacji. W naszym konkretnym przypadku poprzednikiem nie jest pojedyncze zdanie p, ale cała implikacja $p \rightarrow q$ (i to właśnie cała ta implikacja powinna posiadać wartość 1), zaś następnikiem nie proste zdanie q, ale alternatywa $p \vee q$ (i to ona musi być fałszywa), a więc:

$$\begin{array}{ccc} (p \rightarrow q) \rightarrow (p \vee q) \\ 1 \quad \underline{0} \quad 0 \end{array} \quad \text{DOBRZE}$$

W pierwszym nawiasie otrzymaliśmy jedynekę przy implikacji. W tabelce dla tego spójnika widzimy, że jedynka może się przy nim pojawić w trzech różnych sytuacjach. Ponieważ nie wiemy, który wariant wybrać, zostawiamy na razie tę implikację i przechodzimy do drugiego nawiasu. Mamy tu fałszywą alternatywę. W tabelce dla alternatywy widzimy, że jest ona fałszywa tylko w jednym przypadku – gdy oba jej człony są fałszywe. Tu zatem nie mamy żadnego wyboru. Musimy wpisać zera pod obydwoma zmiennymi zdaniowymi:

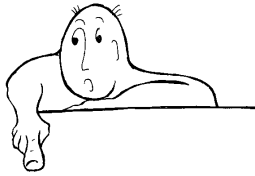
$$\begin{array}{ccccc} (p \rightarrow q) \rightarrow (p \vee q) \\ 1 \quad 0 \quad 0 \quad \underline{0} \quad 0 \end{array}$$

W tym momencie dowiedzieliśmy się, jakie powinny być wartości logiczne zmiennych p i q. Jako że wartości te muszą być oczywiście takie same w całym wyrażeniu (nie może być tak, aby jedno zdanie było w jednym miejscu prawdziwe, a w drugim fałszywe), przepisujemy je we wszystkie miejsca, gdzie zmienne p i q występują:

$$\begin{array}{ccccccc} (p \rightarrow q) \rightarrow (p \vee q) \\ 0 \quad 1 \quad 0 \quad 0 \quad \underline{0} \quad 0 \quad \underline{0} \end{array}$$

Widzimy, że wpisaliliśmy wartości logiczne we wszystkie możliwe miejsca. Pozostaje nam jeszcze sprawdzić, czy wszystko się zgadza. Jeżeli gdzieś mogła wkraść się jakaś nieprawidłowość, to jedynie w ostatnim kroku – tam gdzie przepisaliśmy wartości zmiennych p i q. Sprawdzamy zatem w tabelce, czy implikacja może być prawdziwa (tak wyszło w naszym przykładzie), gdy jej poprzednik i następnik są fałszywe (te wartości zmiennych przepisaliśmy z drugiego nawiasu). Wszystko się zgadza, implikacja taka jest prawdziwa. W innych miejscach formuły też wszystko musi się zgadzać, ponieważ wcześniej wszędzie wpisywaliśmy wartości logiczne wprost z tabelki.

Tak więc już w pierwszej linijce pokazaliśmy, że badana formuła może okazać się schematem zdania fałszywego, a zatem nie jest ona na pewno tautologią.



Uwaga na błędy!

W powyższym przykładzie wykazaliśmy jedynie, że formuła nie jest tautologią. Nie znaczy to jednak, iż jest ona kontrtautologią. Aby stwierdzić, że schemat jest kontrtautologią, musielibyśmy mieć pewność, że generuje on tylko i wyłącznie zdania fałszywe. My natomiast pokazaliśmy jedynie, że daje on takie zdanie w przynajmniej jednym przypadku. Sprawdzenie, czy formuła jest kontrtautologią wymagałoby obecnie posłużenia się metodą skróconą w inny sposób lub zastosowania metody zwykłej. Na razie wiemy tylko i wyłącznie, że nie jest ona tautologią.

Przykład:

Sprawdzimy przy pomocy skróconej metody, czy tautologią jest formuła:

$$(p \wedge q) \rightarrow (p \rightarrow q)$$

Jak zawsze w metodzie skróconej zaczynamy od sprawdzenia, czy formuła może stać się schematem zdania fałszywego, a zatem, czy pod głównym spójnikiem może pojawić się 0.

$$\begin{array}{c} (p \wedge q) \rightarrow (p \rightarrow q) \\ 0 \end{array}$$

Podobnie jak w poprzednim przykładzie mamy zero przy implikacji. Z tabelki dla tego spójnika wiemy, że w takim przypadku prawdziwy musi być poprzednik implikacji (a więc koniunkcja w pierwszym nawiasie), a fałszywy następnik (implikacja w drugim nawiasie):

$$\begin{array}{c} (p \wedge q) \rightarrow (p \rightarrow q) \\ 1 \quad \underline{0} \quad 0 \end{array}$$

W pierwszym nawiasie mamy prawdziwą koniunkcję. Z tabelki widzimy, że taka sytuacja możliwa jest tylko w jednym przypadku – oba człony koniunkcji muszą być prawdziwe:

$$\begin{array}{c} (p \wedge q) \rightarrow (p \rightarrow q) \\ 1 \quad \underline{1} \quad 0 \quad 0 \end{array}$$

Skoro znamy już wartości zmiennych p i q przepisujemy je wszędzie, gdzie te zmienne występują:

$$\begin{array}{c} (p \wedge q) \rightarrow (p \rightarrow q) \\ \underline{1} \quad \underline{1} \quad \underline{1} \quad 0 \quad 1 \quad 0 \quad 1 \end{array}$$

Podobnie jak poprzednio, musimy teraz jeszcze sprawdzić, czy wartości, które przepisaliśmy w ostatnim kroku zgadzają się z tymi, które wpisaliśmy wcześniej. W tym

momencie natykamy się na coś dziwnego. Okazuje się otrzymaliśmy fałszywą implikację, której zarówno poprzednik, jak i następnik są zdaniami prawdziwymi. Ale przecież sytuacja taka jest całkowicie niezgodna z tabelkami! Otrzymaliśmy ewidentną sprzeczność – coś, co nie ma prawa wystąpić:

$$(p \wedge q) \rightarrow (p \rightarrow q)$$

$$1 \ 1 \ 1 \quad 0 \quad \underline{1 \ 0 \ 1}$$

O czym może świadczyć pojawienie się sprzeczności? Aby to zrozumieć, dobrze jest prześledzić cały tok rozumowania od samego początku. Założyliśmy na początku 0 pod głównym spójnikiem całej formuły. Następnie wyciągaliśmy z tego konsekwencje, wpisując wartości, które musiałyby się pojawić, aby założone 0 faktycznie mogło wystąpić. Postępując w ten sposób doszliśmy do sprzeczności. Wynika z tego, że nasze założenie nie daje się utrzymać. Zero pod głównym spójnikiem nie może się pojawić, ponieważ prowadziłoby to do sprzeczności. A skoro pod głównym spójnikiem nie może być nigdy 0, to znaczy że zawsze jest tam 1, a to z kolei świadczy, że badana formuła jest tautologią.

Tautologiczność formuły wykazana została w jednej linijce. Po prostu zamiast pokazywać, że badany schemat zawsze daje zawsze zdania prawdziwe, udowodniliśmy, że nie może wygenerować on zdania fałszywego.

UWAGA!

Sposób, w jaki rozwiązany został powyższy przykład, nie jest jedynym możliwym. Zobaczmy, jak można to było zrobić inaczej.

Rozpoczynamy tak samo, wpisując 0 pod główną implikacją, a następnie 1 przy jej poprzedniku i 0 przy następniku:

$$(p \wedge q) \rightarrow (p \rightarrow q)$$

$$1 \quad \underline{0} \quad 0$$

Zauważmy teraz, że wcale nie musimy zaczynać od prawdziwej koniunkcji w pierwszym nawiasie. Również w drugim nawiasie mamy bowiem tylko jedną możliwość wpisania kombinacji zer i jedynek. Aby umieszczona tam implikacja była fałszywa, prawdziwy musi być jej poprzednik, a fałszywy następnik:

$$(p \wedge q) \rightarrow (p \rightarrow q)$$

$$1 \quad 0 \quad 1 \quad \underline{0} \quad 0$$

Gdy przepiszemy teraz otrzymane wartości zmiennych do pierwszego nawiasu otrzymamy:

$$(p \wedge q) \rightarrow (p \rightarrow q)$$

$$1 \ 1 \ 0 \ 0 \ \underline{1} \ 0 \ \underline{0}$$

Okazuje się, że tym razem również otrzymujemy sprzeczność, tyle że w innym miejscu:

$$(p \wedge q) \rightarrow (p \rightarrow q)$$

$$\underline{1} \ \underline{1} \ \underline{0} \ 0 \ 1 \ 0 \ 0$$

Użyteczna wskazówka:

Gdy sprawdzamy, czy formuła jest tautologią przy pomocy metody skróconej, nie jest istotne, gdzie pojawi się sprzeczność. Często może ona wystąpić w różnych miejscach, w zależności od tego, w jakiej kolejności wpisywaliśmy symbole 0 i 1 do formuły.

Wracając do omawianego przykładu, zobaczmy jeszcze inny sposób, w jaki sprzeczność mogła się ujawnić. Zaczynamy tak jak poprzednio:

$$(p \wedge q) \rightarrow (p \rightarrow q)$$

$$1 \quad \underline{0} \quad 0$$

Teraz zauważamy, że obu nawiasach mamy tylko jedną możliwość wpisania kombinacji 0 i 1 jedynek, więc je od razu jednocześnie wpisujemy:

$$(p \wedge q) \rightarrow (p \rightarrow q)$$

$$1 \ 1 \ 1 \ 0 \ 1 \ 0 \ 0$$

Tym razem również sprzeczność wystąpiła, choć może nie jest to widoczne na pierwszy rzut oka. Zmienna q okazuje się w jednym miejscu reprezentować zdanie prawdziwe, a jednocześnie w innym fałszywe. Taka sytuacja oczywiście nie jest możliwa.

$$(p \wedge q) \rightarrow (p \rightarrow q)$$

$$1 \ 1 \ \underline{1} \ 0 \ 1 \ 0 \ \underline{0}$$

Ponieważ dla właściwego posługiwania się skróconą metodą zero-jedynkową ważne jest zrozumienie całego toku rozumowania z nią związanego, przedstawimy go jeszcze raz.

Gdy chcemy dowiedzieć się, czy schemat jest tautologią, zaczynamy od postawienia symbolu 0 pod głównym spójnikiem, aby sprawdzić, czy formuła może choć w jednym przypadku wygenerować zdanie fałszywe.

Następnie wpisujemy zgodnie z tabelkami dla odpowiednich spójników symbole 0 i 1, w taki sposób w jaki musiałyby one występować, aby zero pod głównym spójnikiem mogło się pojawić. Czyniąc to wpisujemy tylko to, co wiemy na pewno. Gdy w jakimś miejscu mamy

dwie lub trzy możliwości wpisania symboli, nie wpisujemy tam chwilowo nic i przechodzimy dalej, szukając miejsca, gdzie jest tylko jedna możliwość.

Gdy symbol 0 lub 1 pojawi się pod jaką zmienną zdaniową, przepisujemy go wszędzie tam, gdzie dana zmienna występuje w formule.

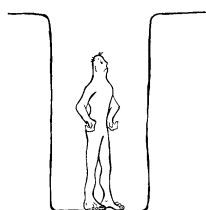
Na końcu sprawdzamy, czy w naszej formule nie pojawiła się przypadkiem sprzeczność (czy wszystko jest zgodne z tabelkami, czy też nie). Jeżeli sprzeczność (niezgodność z tabelkami) ma się gdzieś pojawić, to dzieje się to na ogół tam, gdzie w ostatnim kroku przepisaliliśmy wartości zmiennych. Jeżeli sprzeczności nigdzie nie ma, to znaczy, że formuła może okazać się schematem zdania fałszywego (takie założenie na początku przyjęliśmy wpisując 0 pod głównym spójnikiem), a więc nie jest ona tautologią. Gdy natomiast w formule pojawi się sprzeczność, oznacza to, że nie może ona wygenerować zdania fałszywego (przyjęte na początku założenie nie daje się utrzymać), a zatem jest ona tautologią.



DO ZAPAMIĘTANIA.

Jeszcze raz cała procedura w telegraficznym skrócie:

1. Zakładamy 0 pod głównym spójnikiem.
2. Wyciągamy z przyjętego założenia wszelkie konsekwencje, wpisując 0 i 1, tam gdzie istnieje tylko jedna możliwość ich wystąpienia.
3. Sprawdzamy, czy wszystko się zgadza z tabelkami (czy nie ma sprzeczności).
4. Ogłaszamy wynik według recepty: jest sprzeczność – formuła jest tautologią, nie ma sprzeczności – formuła nie jest tautologią.



1.4.3. UTRUDNIENIA I PUŁAPKI.

Uwaga na negacje.

Badane przez logików formuły są na ogół bardziej skomplikowane od omówionych w powyższych przykładach. Pierwsze utrudnienie mogą spowodować obecne w nich negacje.

Przykład:

$$(p \rightarrow q) \rightarrow (\sim q \rightarrow \sim p)$$

Rozpoczynamy od postawienia 0 pod głównym spójnikiem i wyciągamy z tego pierwszą konsekwencję:

$$\begin{array}{ccccccc} (p \rightarrow q) \rightarrow (\sim q \rightarrow \sim p) \\ 1 \quad \quad \underline{0} \quad \quad 0 \end{array}$$

Jedną możliwość wpisania kombinacji 0 i 1 mamy w drugim nawiasie. Aby implikacja była fałszywa, jej poprzednik musi być prawdziwy, a następnik fałszywy. Ważne jest tu jednak poprawne określenie co jest poprzednikiem i następnikiem badanej implikacji. Poprzednikiem jest zdanie złożone $\sim q$, a więc jedynek wskazującą na jego prawdziwość wpisujemy nad jego głównym spójnikiem – negacją; podobnie następnikiem jest złożone zdanie $\sim p$ i tu również wskazując jego fałszywość 0 wpisujemy pod negacją:

$$\begin{array}{ccccccc} (p \rightarrow q) \rightarrow (\sim q \rightarrow \sim p) \\ 1 \quad \quad 0 \quad 1 \quad \underline{0} \quad 0 \end{array}$$

Dopiero w tym momencie, korzystając z tabelki dla negacji, możemy wpisać wartości zdań p i q:

$$\begin{array}{ccccccc} (p \rightarrow q) \rightarrow (\sim q \rightarrow \sim p) \\ 1 \quad \quad 0 \quad \underline{1} \quad 0 \quad 0 \quad \underline{0} \quad 1 \end{array}$$

Po przepisaniu otrzymanych wartości do pierwszego nawiasu otrzymujemy sprzeczność: implikacja o prawdziwym poprzedniku i fałszywym następniku nie może być prawdziwa:

$$\begin{array}{ccccccc} (p \rightarrow q) \rightarrow (\sim q \rightarrow \sim p) \\ \underline{1 \quad 1 \quad 0} \quad 0 \quad 1 \quad 0 \quad 0 \quad 1 \end{array}$$

Badana formuła jest zatem tautologią.

Przykład:

Zbadamy, czy tautologią jest formuła $(p \rightarrow \sim q) \vee (\sim p \wedge q)$

Główny spójnik stanowi tu alternatywa, która jest fałszywa tylko w jednym przypadku – gdy oba jej człony są fałszywe:

$$\begin{array}{ccccccc} (p \rightarrow \sim q) \vee (\sim p \wedge q) \\ 0 \quad \quad \underline{0} \quad \quad 0 \end{array}$$

W pierwszym nawiasie mamy tylko jedną możliwość: aby implikacja była fałszywa jej poprzednik – p, musi być prawdziwy, a jej następnik – $\sim q$, fałszywy. Z tego ostatniego możemy od razu wpisać, że prawdziwe musi być q:

$$\begin{array}{ccccccc} (p \rightarrow \sim q) \vee (\sim p \wedge q) \\ 1 \quad \underline{0} \quad \underline{0} \quad 1 \quad 0 \quad \quad 0 \end{array}$$

Przepisujemy otrzymane wartości p i q do drugiego nawiasu:

$$(p \rightarrow \sim q) \vee (\sim p \wedge q)$$

$$\underline{1} \ 0 \ 0 \ \underline{1} \ 0 \quad 1 \ 0 \ 1$$

To jeszcze nie koniec zadania, ponieważ nie mamy wpisanej wartości negacji p. Skoro jednak samo p jest prawdziwe, to jego negacja musi być fałszywa:

$$(p \rightarrow \sim q) \vee (\sim p \wedge q)$$

$$1 \ 0 \ 0 \ 1 \ 0 \ 0 \ \underline{1} \ 0 \ 1$$

W powyższej formule nie występuje nigdzie sprzeczność. Członami koniunkcji w drugim nawiasie są: $\sim p$ oraz q. Negacja p jest fałszywa, a q prawdziwe – koniunkcja takich zdań (0 i 1) zgodnie z tabelkami musi być fałszywa.

Badana formuła nie jest tautologią.

Formuły z większą ilością nawiasów.

W dłuższych formułach pewne utrudnienia sprawić może wielość nawiasów wskazujących hierarchię spójników. W takich dłuższych formułach trzeba szczególną uwagę zwracać na wpisywanie symboli wartości logicznych we właściwe miejsca oraz na dokładne badanie, czy ostatecznie wystąpiła sprzeczność.

Przykład:

$$[(p \rightarrow q) \vee (r \rightarrow \sim p)] \rightarrow [p \rightarrow (q \vee \sim r)]$$

Głównym spójnikiem badanej formuły jest implikacja wiążąca wyrażenia w kwadratowych nawiasach. Aby implikacja była fałszywa, to jej poprzednik musi być prawdziwy, a następnik fałszywy – symbole jedynki i zera wpisujemy więc pod głównymi spójnikami każdego z wyrażen w kwadratowych nawiasach:

$$[(p \rightarrow q) \vee (r \rightarrow \sim p)] \rightarrow [p \rightarrow (q \vee \sim r)]$$

$$\quad \quad \quad 1 \quad \quad \quad \underline{0} \quad \quad \quad 0$$

W przypadku prawdziwej alternatywy w pierwszym nawiasie mamy trzy możliwości, więc na razie pomijamy to miejsce. W przypadku fałszywej implikacji w drugim nawiasie kwadratowym możemy wpisać, że prawdziwy jest jej poprzednik – czyli p, a fałszywy następnik – czyli alternatywa w nawiasie. Z tego ostatniego faktu wnioskujemy o fałszywości obu członów alternatywy – q oraz $\sim r$. W takim razie prawdziwe musi być oczywiście r:

$$[(p \rightarrow q) \vee (r \rightarrow \sim p)] \rightarrow [p \rightarrow (q \vee \sim r)]$$

$$\quad \quad \quad 1 \quad \quad \quad 0 \ 1 \ \underline{0} \ 0 \ 0 \ 0 \ 1$$

Otrzymane wartości zmiennych zdaniowych przepisujemy do wyrażenia w pierwszym kwadratowym nawiasie. Na ich podstawie obliczamy wartość $\sim p$, a następnie wartości implikacji w nawiasach okrągłych:

$$[(p \rightarrow q) \vee (r \rightarrow \sim p)] \rightarrow [p \rightarrow (q \vee \sim r)]$$

$$\underline{1} \ 0 \ \underline{0} \ 1 \ \underline{1} \ 0 \ 0 \ \underline{1} \quad 0 \ 1 \ 0 \ 0 \ 0 \ 0 \ 1$$

Teraz musimy sprawdzić, czy wszystko się zgadza. Ostatnie wartości jakie wpisaliśmy, to zera przy implikacjach w okrągłych nawiasach. Wartości te zgadzają się wprawdzie z wartościami zdań tworzących te implikacje (nie może być inaczej – przecież na podstawie tych zdań obliczyliśmy wartość implikacji zgodnie z tabelkami), kolidują natomiast z wartością alternatywy, której są członami. W tym właśnie miejscu tkwi sprzeczność – być może nie całkiem widoczna na pierwszy rzut oka:

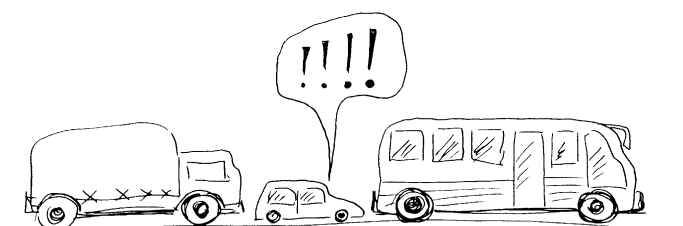
$$[(p \rightarrow q) \vee (r \rightarrow \sim p)] \rightarrow [p \rightarrow (q \vee \sim r)]$$

$$1 \ \underline{0} \ 0 \ \underline{1} \ 1 \ \underline{0} \ 0 \ 1 \quad 0 \ 1 \ 0 \ 0 \ 0 \ 0 \ 1$$

Badana formuła jest zatem tautologią.

Gdy pozornie utkniemy.

Czasami może się wydawać, że w badanej formule nie ma takiego miejsca, gdzie byłaby tylko jedna możliwość wpisania zer i jedynek. Często jednak okazuje się, że jest to tylko złudzenie i po bliższej analizie znajdujemy odpowiednie wyjście.



GDY POZORNIE UTKNIEMY ...

Przykład:

Sprawdzimy, czy tautologią jest formuła:

$$[(p \rightarrow q) \wedge (p \rightarrow r)] \rightarrow [p \rightarrow (q \wedge r)]$$

Po postawieniu zera przy głównej implikacji otrzymujemy jedynkę przy koniunkcji w pierwszym kwadratowym nawiasie oraz zero przy implikacji w drugim nawiasie kwadratowym. Z prawdziwości koniunkcji wyciągamy wniosek o prawdziwości obu jej

członów, a z fałszywości implikacji o prawdziwości p oraz fałszywości koniunkcji $q \wedge r$. Wartość p możemy przepisać w miejsca, gdzie zmienna ta jeszcze występuje:

$$\begin{array}{ccccccc} [(p \rightarrow q) \wedge (p \rightarrow r)] \rightarrow [p \rightarrow (q \wedge r)] \\ 1 & 1 & & 1 & 1 & 1 & \underline{0} & 1 & 0 & 0 \end{array}$$

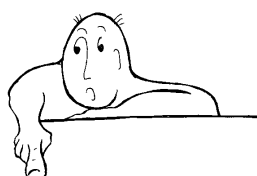
W tym momencie mogłoby się wydawać, że w każdym miejscu mamy po kilka możliwości wstawiania zer i jedynek. Jest to jednak tylko pozór. W dwóch pierwszych nawiasach okrągłych mamy prawdziwe implikacje. Ogólnie rzecz biorąc implikacja jest prawdziwa w trzech różnych przypadkach; zauważmy jednak, że my znamy obecnie również wartości poprzedników tych implikacji – są one prawdziwe. Gdy spojrzymy do tabelki dla implikacji, zobaczymy, że wśród trzech przypadków, gdy jest ona prawdziwa, jest tylko jeden taki, kiedy prawdziwy jest jej poprzednik – w przypadku tym prawdziwy musi być również następnik implikacji. Tak więc w rzeczywistości mamy tylko jedną możliwość określenia wartości zmiennych q i r w badanych implikacjach – muszą być one prawdziwe:

$$\begin{array}{ccccccccccc} [(p \rightarrow q) \wedge (p \rightarrow r)] \rightarrow [p \rightarrow (q \wedge r)] \\ \underline{1} & \underline{1} & 1 & 1 & \underline{1} & \underline{1} & 1 & 0 & 1 & 0 & 0 \end{array}$$

Po przepisaniu wartości q i r w inne miejsca, gdzie zmienne te występują, otrzymujemy ewidentną sprzeczność w koniunkcji q i r :

$$\begin{array}{ccccccccccc} [(p \rightarrow q) \wedge (p \rightarrow r)] \rightarrow [p \rightarrow (q \wedge r)] \\ 1 & 1 & 1 & 1 & 1 & 1 & 1 & \underline{0} & 1 & 0 & \underline{1} & \underline{0} & \underline{1} \end{array}$$

Badana formuła jest więc tautologią.



Uwaga na błędy!

Należy koniecznie zauważyć różnicę pomiędzy prawdziwą implikacją z prawdziwym poprzednikiem a prawdziwą implikacją z prawdziwym następnikiem. W pierwszym przypadku istnieje tylko jedna możliwość co do wartości drugiego członu (musi być 1), natomiast w drugim są dwie możliwości (0 lub 1):

$$\begin{array}{cc} p \rightarrow q & p \rightarrow q \\ \underline{1} & \underline{1} & 1 & ? & \underline{1} & \underline{1} \end{array}$$

Podobna różnica zachodzi pomiędzy prawdziwymi implikacjami z fałszywym następnikiem i poprzednikiem:

$$\begin{array}{cc} p \rightarrow q & p \rightarrow q \\ 0 & \underline{1} & \underline{0} & \underline{0} & \underline{1} & ? \end{array}$$

Zależności te powinny stać się jasne po dokładnym przeanalizowaniu tabelki dla implikacji.

Przykład:

Zbadamy, czy tautologią jest formuła: $\sim (p \rightarrow q) \vee [\sim (p \vee q) \vee (p \vee r)]$

Zaczynając od postawienia zera przy głównym spójniku, którym jest tu alternatywa, otrzymujemy fałszywe obydwie człony alternatywy, czyli negację formuły $p \rightarrow q$ (bo to stojąca przed nawiasem negacja jest tu głównym spójnikiem) oraz alternatywę w nawiasie kwadratowym:

$$\begin{array}{ccccccc} \sim (p \rightarrow q) & \vee & [\sim (p \vee q) & \vee & (p \vee r)] \\ 0 & & 0 & & 0 \end{array}$$

Skoro fałszywa jest negacja, to prawdziwa musi być formuła, do której negacja się odnosi. Natomiast z fałszywości alternatywy w nawiasie kwadratowym, wnioskujemy o fałszywości obu jej członów:

$$\begin{array}{ccccccc} \sim (p \rightarrow q) & \vee & [\sim (p \vee q) & \vee & (p \vee r)] \\ 0 & 1 & 0 & 0 & 0 & 0 \end{array}$$

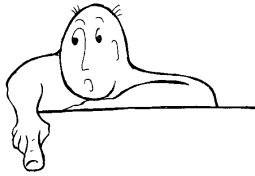
Znowu mamy fałszywą negację, a więc prawdziwa jest negowana przez nią formuła w nawiasie. Skoro natomiast fałszywa jest alternatywa $p \vee r$, to fałszywe są oba jej człony. Wartość zmiennej p przepisujemy tam, gdzie zmienna ta jeszcze występuje:

$$\begin{array}{ccccccc} \sim (p \rightarrow q) & \vee & [\sim (p \vee q) & \vee & (p \vee r)] \\ 0 & 0 & 1 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \end{array}$$

W pierwszym nawiasie mamy do czynienia z prawdziwą implikacją o fałszywym poprzedniku. W takim wypadku nic jeszcze nie wiemy o następniku – zgodnie z tabelkami może być on albo fałszywy albo prawdziwy. Natomiast w przypadku prawdziwej alternatywy z fałszywym pierwszym członem mamy tylko jedną możliwość – drugi człon musi być prawdziwy. Wpisujemy więc 1 pod q i przepisujemy ją tam, gdzie zmienna ta jeszcze występuje:

$$\begin{array}{ccccccc} \sim (p \rightarrow q) & \vee & [\sim (p \vee q) & \vee & (p \vee r)] \\ 0 & 0 & 1 & 1 & 0 & 0 & 0 & 1 & 1 & 0 & 0 & 0 & 0 \end{array}$$

W powyższej formule nie występuje nigdzie sprzeczność, a zatem nie jest ona tautologią.



Uwaga na błędy!

W przypadku prawdziwej alternatywy również nie w każdym przypadku możemy obliczyć wartość drugiego członu na podstawie znajomości wartości jednego członu oraz całej formuły. Możemy to uczynić jedynie wtedy, gdy alternatywa jest prawdziwa, a jeden z jej członów fałszywy – wtedy, zgodnie z tabelkami drugi musi być prawdziwy:

$$\begin{array}{cc} p \vee q \\ \underline{0} \ \underline{1} \ 1 \end{array}$$

$$\begin{array}{cc} p \vee q \\ 1 \ \underline{1} \ \underline{0} \end{array}$$

$$\begin{array}{cc} p \vee q \\ \underline{1} \ \underline{1} \ ? \end{array}$$

$$\begin{array}{cc} p \vee q \\ ? \ \underline{1} \ \underline{1} \end{array}$$

Podobnie w przypadku fałszywej koniunkcji możemy obliczyć wartość drugiego członu, tylko wtedy, gdy pierwszy jest prawdziwy:

$$\begin{array}{cc} p \wedge q \\ \underline{1} \ \underline{0} \ 0 \end{array}$$

$$\begin{array}{cc} p \wedge q \\ 0 \ \underline{0} \ \underline{1} \end{array}$$

$$\begin{array}{cc} p \wedge q \\ \underline{0} \ \underline{0} \ ? \end{array}$$

$$\begin{array}{cc} p \wedge q \\ ? \ \underline{0} \ \underline{0} \end{array}$$

Gdy utkniemy poważniej...

Przykład:

Sprawdźmy, czy tautologią jest formuła: $\{[p \rightarrow (q \wedge r)] \wedge (p \vee r)\} \rightarrow q$

Po założeniu fałszywości całej formuły, otrzymujemy 1 przy koniunkcji w nawiasie klamrowym i 0 przy q . Wartość q oczywiście przepisujemy, tam gdzie jeszcze q się pojawia. Z prawdziwości koniunkcji wnioskujemy o prawdziwości obu jej członów:

$$\begin{array}{cccccc} \{[p \rightarrow (q \wedge r)] \wedge (p \vee r)\} \rightarrow q \\ 1 \quad 0 \quad 1 \quad 1 \quad \underline{0} \quad 0 \end{array}$$

W tym momencie mogłoby się wydawać, że zupełnie nie wiadomo, co robić dalej. Jednakże przyjrzyjmy się bliżej koniunkcji $q \wedge r$. Jeden z członów tej koniunkcji jest fałszywy – a zatem, zgodnie z tabelkami – cała koniunkcja musi być fałszywa.

$$\begin{array}{cccccc} \{[p \rightarrow (q \wedge r)] \wedge (p \vee r)\} \rightarrow q \\ 1 \quad \underline{0} \ 0 \quad 1 \quad 1 \quad 0 \ 0 \end{array}$$

W tym momencie, na podstawie faktu, że prawdziwa implikacja z fałszywym następnikiem musi mieć fałszywy poprzednik, obliczamy wartość zmiennej p – 0, i przepisujemy ją, tam gdzie p występuje w alternatywie $p \vee q$.

$$\begin{array}{cccccc} \{[p \rightarrow (q \wedge r)] \wedge (p \vee r)\} \rightarrow q \\ 0 \ 1 \ \underline{0} \ \underline{0} \quad 1 \ 0 \ 1 \quad 0 \ 0 \end{array}$$

Ponieważ prawdziwa alternatywa z fałszywym pierwszym członem musi mieć prawdziwy drugi człon, wpisujemy 1 pod zmienną r w formule $p \vee r$ i przepisujemy tę wartość do koniunkcji $q \wedge r$.

$$\{[p \rightarrow (q \wedge r)] \wedge (p \vee r)\} \rightarrow q$$

$$0 \ 1 \ 0 \ 0 \ 1 \ 1 \ 0 \ \underline{1} \ 0 \ 0$$

Ponieważ przy takich podstawieniach w powyższej formule nie występuje nigdzie sprzeczność, nie jest ona tautologią.



WARTO ZAPAMIĘTAĆ.

Oto przypadki, gdzie można obliczyć wartość zdania złożonego na podstawie tylko jednego z jego członów:

$p \wedge q$	$p \wedge q$
$\underline{0} \ 0$	$0 \ \underline{0}$
$p \vee q$	$p \vee q$
$\underline{1} \ 1$	$1 \ \underline{1}$
$p \rightarrow q$	$p \rightarrow q$
$0 \ 1$	$\underline{1} \ 1$

Ogólnie – obliczenie wartości całego zdania złożonego jest możliwe na podstawie: fałszywości jednego z członów koniunkcji, prawdziwości jednego z członów alternatywy, fałszywości poprzednika implikacji oraz prawdziwości następnika implikacji.

Przykład:

Sprawdźmy, czy tautologią jest formuła: $\{[\sim (p \wedge q) \rightarrow r] \wedge (r \rightarrow p)\} \rightarrow (p \wedge q)$

Pierwsze kroki są oczywiste i wyglądają następująco:

$$\{[\sim (p \wedge q) \rightarrow r] \wedge (r \rightarrow p)\} \rightarrow (p \wedge q)$$

$$1 \quad 1 \quad 1 \quad \underline{0} \quad 0$$

W tym miejscu mogłoby się wydawać, że wszędzie mamy po kilka możliwości wpisania zer i jedynek. Zauważmy jednak, że znamy wartość koniunkcji $p \wedge q$ w ostatnim nawiasie, która to koniunkcja występuje też w jeszcze jednym miejscu. Możemy więc przepisać wartość tej koniunkcji, podobnie jak przepisujemy wartości zmiennych:

$$\{[\sim (p \wedge q) \rightarrow r] \wedge (r \rightarrow p)\} \rightarrow (p \wedge q)$$

$$0 \quad 1 \quad 1 \quad 1 \quad 0 \quad \underline{0}$$

Skoro koniunkcja $p \wedge q$ jest fałszywa, to jej negacja musi być prawdziwa. Na podstawie prawdziwości implikacji w nawiasie kwadratowym oraz prawdziwości jej poprzednika możemy obliczyć wartość $r = 1$, i przepisać ją:

$$\{[\sim(p \wedge q) \rightarrow r] \wedge (r \rightarrow p)\} \rightarrow (p \wedge q)$$

$$1 \quad \underline{0} \quad 1 \quad 1 \quad 1 \quad 1 \quad 1 \quad 0 \quad 0$$

Teraz możemy z łatwością obliczyć wartość p w implikacji $r \rightarrow p$ (1) i przepisać ją do obu koniunkcji $p \wedge q$. Mamy wtedy fałszywą koniunkcję z prawdziwym jednym członem – a zatem fałszywy musi być jej człon drugi – q .

$$\{[\sim(p \wedge q) \rightarrow r] \wedge (r \rightarrow p)\} \rightarrow (p \wedge q)$$

$$1 \quad 1 \quad 0 \quad 0 \quad 1 \quad 1 \quad 1 \quad \underline{1} \quad \underline{1} \quad 1 \quad 0 \quad 1 \quad 0 \quad 0$$

Przy takich podstawieniach nie ma żadnej sprzeczności, a zatem badana formuła nie jest tautologią.

PRAKTYCZNA RADA:

Co zrobić, gdy „utknę” i wydaje się, że nigdzie nie ma jednej możliwości wpisania zer i jedynek? Należy wówczas sprawdzić następujące rzeczy:

- czy przepisałem wszystkie wartości zmiennych w inne miejsca, gdzie zmienne występują,
- czy wpisałem wartości zmiennych, gdy obliczone są wartości ich negacji lub wartości negacji, gdy obliczone są wartości zmiennych (przy negacji jest zawsze tylko jedna możliwość),
- czy wpisałem wartości przy spójnikach dwuargumentowych, gdy znane są wartości obu ich członów,
- czy możliwe jest obliczenia wartości członu jakiegoś spójnika na podstawie znajomości wartości drugiego członu oraz całego zdania,
- czy możliwe jest gdzieś wpisanie wartości przy spójniku na podstawie znajomości wartości logicznej jednego z jego członów,
- czy można gdzieś przepisać wartość całego zdania złożonego.

Dwie możliwości od samego początku.

Czasem już na początku mamy dwie możliwości wpisania kombinacji zer i jedynek, na przykład gdy głównym spójnikiem jest równoważność.

Przykład:

$$[p \rightarrow (q \rightarrow r)] \equiv [(q \wedge \sim r) \rightarrow \sim p]$$

Sprawdzenie, czy powyższa formuła może być schematem zdania fałszywego wymaga rozpatrzenia dwóch możliwości:

$$\begin{array}{ccc} 1 & \underline{0} & 0 \\ [p \rightarrow (q \rightarrow r)] \equiv [(q \wedge \sim r) \rightarrow \sim p] \\ 0 & \underline{0} & 1 \end{array}$$

W przypadku „górnym” zacząć należy od prawej strony. Z fałszywości implikacji wiemy, że prawdziwy musi być jej poprzednik, czyli koniunkcja $q \wedge \sim r$, natomiast fałszywy następnik – $\sim p$. Z prawdziwości koniunkcji wyciągamy wniosek o prawdziwości jej członów. Wartość logiczna zdań r i p jest oczywiście odwrotna do wartości ich negacji:

$$\begin{array}{ccccccc} 1 & 0 & 1 & 1 & 1 & 0 & \underline{0} & 0 & 1 \\ [p \rightarrow (q \rightarrow r)] \equiv [(q \wedge \sim r) \rightarrow \sim p] \\ 0 & 0 & & & & & & 1 \end{array}$$

Po przepisaniu wartości zmiennych do lewej strony równoważności otrzymujemy:

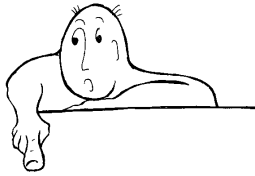
$$\begin{array}{ccccccc} 1 & 1 & 1 & 0 & 0 & 1 & 1 & 1 & 0 & 0 & 0 & 1 \\ [p \rightarrow (q \rightarrow r)] \equiv [(q \wedge \sim r) \rightarrow \sim p] \\ 0 & 0 & & & & & & 1 \end{array}$$

Pozostaje nam jeszcze obliczenie wartości implikacji $q \rightarrow r$. Ponieważ jej poprzednik jest prawdziwy, a następnik fałszywy, implikacja ta powinna być fałszywa:

$$\begin{array}{ccccccc} 1 & 1 & \underline{1} & 0 & \underline{0} & 0 & 1 & 1 & 1 & 0 & 0 & 0 & 1 \\ [p \rightarrow (q \rightarrow r)] \equiv [(q \wedge \sim r) \rightarrow \sim p] \\ 0 & 0 & & & & & & 1 \end{array}$$

Teraz musimy sprawdzić, czy to, co wpisaliśmy na końcu, nie stoi w sprzeczności z wartościami obliczonymi wcześniej. Fałszywa implikacja $q \rightarrow r$ jest jednocześnie następnikiem implikacji w nawiasie kwadratowym o poprzedniku p . Otrzymujemy tu sprzeczność, ponieważ cała implikacja w kwadratowym nawiasie wyszła nam prawdziwa, co jest niemożliwe przy prawdziwym poprzedniku i fałszywym następniku:

$$\begin{array}{ccccccc} \underline{1} & \underline{1} & 1 & \underline{0} & 0 & 0 & 1 & 1 & 1 & 0 & 0 & 0 & 1 \\ [p \rightarrow (q \rightarrow r)] \equiv [(q \wedge \sim r) \rightarrow \sim p] \\ 0 & 0 & & & & & & 1 \end{array}$$



Uwaga na błędy!

Otrzymanie sprzeczności w jednym z rozpatrywanych przypadków nie stanowi jeszcze dowodu, iż badana formuła jest tautologią. Należy pamiętać, że sprawdzanie tautologiczności formuły przy pomocy metody skróconej polega na stwierdzeniu niemożliwości wygenerowania przez dany schemat zdania fałszywego. Ponieważ w badanym przykładzie już na samym początku stwierdziliśmy istnienie dwóch przypadków w których formuła mogłaby okazać się schematem zdania fałszywego, wyeliminowanie jednego z nich (co dotąd zrobiliśmy), niczego jeszcze nie przesądza.

Musimy teraz zbadać drugi, „dolny” przypadek. Tu oczywiście rozpoczynamy od lewej strony, a otrzymane wartości zmiennych przepisujemy do strony prawej.

$$\begin{array}{cccccccccc} \underline{1} & \underline{1} & 1 & \underline{0} & 0 & 0 & 1 & 1 & 1 & 0 & 0 & 0 & 1 \\ [p \rightarrow (q \rightarrow r)] & \equiv & [(q \wedge \sim r) \rightarrow \sim p] \\ 1 & 0 & 1 & 0 & 0 & 0 & 1 & & 0 & 1 & & 1 \end{array}$$

Po obliczeniu wartości negacji zdań r oraz p , a następnie koniunkcji $q \wedge \sim r$, otrzymujemy sprzeczność z prawej strony równoważności:

$$\begin{array}{cccccccccc} \underline{1} & \underline{1} & 1 & \underline{0} & 0 & 0 & 1 & 1 & 1 & 0 & 0 & 0 & 1 \\ [p \rightarrow (q \rightarrow r)] & \equiv & [(q \wedge \sim r) \rightarrow \sim p] \\ 1 & 0 & 1 & 0 & 0 & 0 & 1 & \underline{1} & 1 & 0 & \underline{1} & \underline{0} & 1 \end{array}$$

Dopiero teraz, gdy okazało się, że niemożliwe jest wygenerowanie przez badaną formułę zdania fałszywego na żaden z dwóch teoretycznie możliwych sposobów, możemy stwierdzić, że schemat ten jest tautologią.

Przykład:

Zbadamy teraz, czy tautologią jest następująca formuła:

$$[p \rightarrow (\sim r \rightarrow q)] \equiv [(p \wedge \sim q) \vee (p \rightarrow r)]$$

Tu również głównym spójnikiem jest równoważność, która może dać zdanie fałszywe w dwóch przypadkach:

$$\begin{array}{ccc} 0 & \underline{0} & 1 \\ [p \rightarrow (\sim r \rightarrow q)] & \equiv & [(p \wedge \sim q) \vee (p \rightarrow r)] \\ 1 & \underline{0} & 0 \end{array}$$

W „górnym” przypadku należy rozpocząć od lewej strony. Po obliczeniu wartości zmiennych i przepisaniu ich na stronę prawą otrzymamy:

$$\begin{array}{cccccccccccc} 1 & 0 & 1 & 0 & 0 & 0 & 0 & 1 & & 0 & 1 & 1 & 0 \\ [p \rightarrow (\sim r \rightarrow q)] & \equiv & [(p \wedge \sim q) \vee (p \rightarrow r)] \\ 1 & & & & 0 & & & & & 0 & & & \end{array}$$

Teraz możemy obliczyć wartość negacji q , a następnie koniunkcji $p \wedge \sim q$ oraz implikacji $p \rightarrow r$ na podstawie wartości logicznej ich członów:

$$\begin{array}{cccccccccccc} 1 & 0 & 1 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 0 & 1 & 1 & 0 & 0 \\ [p \rightarrow (\sim r \rightarrow q)] & \equiv & [(p \wedge \sim q) \vee (p \rightarrow r)] \\ 1 & & & & 0 & & & & & 0 & & & & & \end{array}$$

Okazuje się, że przy takim podstawieniu zer i jedynek w badanej formule nie występuje żadna sprzeczność. Pokazaliśmy zatem, że formuła ta może być schematem zdania fałszywego, a więc na pewno nie jest tautologią. Badanie drugiej, „dolnej” możliwości nic tu zmieni, więc możemy go zaniechać.

Czasem nie trzeba wiedzieć wszystkiego.

Bywa, że nie musimy znać wartości wszystkich zmiennych, aby stwierdzić, że formuła jest tautologią – sprzeczność może pojawić się już wcześniej.

Przykład:

Zbadamy, czy tautologią jest formuła: $\{[r \rightarrow (q \wedge s)] \wedge [(p \vee s) \rightarrow r]\} \rightarrow (\sim q \rightarrow \sim p)$

Po standardowo rozpoczętym sprawdzaniu formuły otrzymujemy:

$$\begin{array}{cccccccc} \{[r \rightarrow (q \wedge s)] \wedge [(p \vee s) \rightarrow r]\} & \rightarrow & (\sim q \rightarrow \sim p) \\ 1 & 0 & & 1 & 1 & & 1 & \underline{0} & 1 & 0 & 0 & 0 & 1 \end{array}$$

Teraz możemy obliczyć wartość koniunkcji $q \wedge s$ na podstawie fałszywości jednego z jej członów oraz alternatywy $p \vee s$ na podstawie prawdziwości p :

$$\begin{array}{cccccccc} \{[r \rightarrow (q \wedge s)] \wedge [(p \vee s) \rightarrow r]\} & \rightarrow & (\sim q \rightarrow \sim p) \\ 1 & \underline{0} & 0 & & 1 & \underline{1} & 1 & & 1 & & 0 & 1 & 0 & 0 & 0 & 1 \end{array}$$

W pierwszym kwadratowym nawiasie mamy obecnie prawdziwą implikację z fałszywym następnikiem – a zatem fałszywy musi być również jej poprzednik, czyli r . Po przepisaniu wartości r do drugiego nawiasu otrzymujemy w nim sprzeczność, świadczącą o tym, że badana formuła jest tautologią:

$$\begin{array}{cccccccc} \{[r \rightarrow (q \wedge s)] \wedge [(p \vee s) \rightarrow r]\} & \rightarrow & (\sim q \rightarrow \sim p) \\ 0 & 1 & 0 & 0 & & 1 & 1 & \underline{1} & & \underline{1} & \underline{0} & & 0 & 1 & 0 & 0 & 0 & 1 \end{array}$$

Zauważmy, że sprzeczność pojawiła się, pomimo że nie poznaliśmy wartości zmiennej s ; sprzeczność ta jest od s niezależna – wystąpiłaby zarówno gdyby zdanie oznaczane przez s było prawdziwe, jak i wtedy, gdyby było ono fałszywe.

Może też zdarzyć się odwrotna sytuacja: sprzeczność nie pojawi się, niezależnie jakie zdanie podstawilibyśmy za jakąś zmienną.

Przykład:

Zbadamy, czy tautologią jest formuła: $[(p \vee q) \wedge r] \rightarrow \sim p$.

Po założeniu 0 pod głównym spójnikiem, niemal natychmiast otrzymujemy:

$$\begin{array}{ccccccc} [(p \vee q) \wedge r] \rightarrow \sim p \\ 1 & 1 & & 1 & 1 & \underline{0} & 0 & 1 \end{array}$$

W obecnej sytuacji nie mamy żadnych informacji pozwalających określić wartość zdania oznaczanego przez q . Zauważmy jednak, że jakiegokolwiek q by nie było, na pewno w badanej formule nie powstanie sprzeczność. W związku z tym możemy pod q wpisać dowolną wartość – cokolwiek bowiem tam wpisujemy, wykażemy, że formuła może być schematem zdania fałszywego (nie ma w tym żadnej sprzeczności), a więc nie jest ona tautologią:

$$\begin{array}{ccccccc} [(p \vee q) \wedge r] \rightarrow \sim p & \text{lub} & [(p \vee q) \wedge r] \rightarrow \sim p \\ 1 & 1 & 1 & 1 & 1 & \underline{0} & 0 & 0 & 1 & 1 & 1 & 0 & 1 & 1 & 0 & 0 & 1 \end{array}$$

Gdy nic już nie wiadomo...

Czasami może się zdarzyć i tak, że w jakimś momencie w badanej formule wszędzie są pod dwie lub nawet trzy możliwości wpisania kombinacji zer i jedynek.

Przykład:

Zbadamy, czy tautologią jest bardzo krótka formuła $(p \vee q) \rightarrow (r \wedge s)$.

$$\begin{array}{ccccccc} (p \vee q) \rightarrow (r \wedge s) \\ 1 & & \underline{0} & & 0 \end{array}$$

W takiej sytuacji wszędzie mamy po trzy możliwości. Nie powinno to jednak nikogo szczególnie przestraszyć, choć na początku może wyglądać groźnie. W istocie jest to sytuacja taka sama, jaka pojawiła się w ostatnim przykładzie, tyle że obecnie wystąpiła już na początku badania formuły i z niejako „większym natężeniem”.

Przypomnijmy sobie jednak istotę skróconej metody zero-jedynkowej. Polega ona na poszukiwaniu takich podstawień zer i jedynek, aby formuła dała zdanie fałszywe. Tutaj już na pierwszy rzut oka mamy takich możliwości sporo – wystarczy zatem wybrać dowolną z nich i wpisać, na przykład:

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

$$1 \ 1 \ 0 \ 0 \ 0 \ 0$$

W ten sposób pokazujemy, że formuła nie jest tautologią, ponieważ stała się schematem zdania fałszywego.

Równie dobrym rozwiązaniem byłoby też na przykład takie:

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

$$0 \ 1 \ 1 \ 0 \ 0 \ 1$$

1.4.4. KONTRTAUTOLOGIE.

Jak dotąd stosowaliśmy metodę skróconą do badania, czy formuła jest tautologią. Gdy przy jej pomocy odkrywaliśmy, że formuła tautologią nie jest, nie wiedzieliśmy jeszcze, czy jest ona kontrtautologią, czy też może być schematem zarówno zdań prawdziwych, jak i fałszywych. Teraz zobaczymy, jak sprawdzić przy pomocy metody skróconej, czy formuła jest kontrtautologią.

Procedura sprawdzania, czy formuła jest kontrtautologią różni się od sprawdzania tautologiczności jedynie wstępnym założeniem. Jak wiemy, kontrtautologia, to schemat dający wyłącznie zdania fałszywe. Aby zbadać przy pomocy metody skróconej, czy formuła jest kontrtautologią, musimy więc sprawdzić, czy może ona przynajmniej raz wygenerować zdanie prawdziwe. W praktyce wygląda to tak, że stawiamy 1 przy głównym spójniku zdania i znanymi już sposobami wyciągamy z tego wszelkie konsekwencje. Jeśli okaże się na końcu, że otrzymaliśmy sprzeczność, będzie to świadczyło, że formuła nie może być schematem zdania prawdziwego, a zatem jest kontrtautologią. Brak sprzeczności pokaże, że formuła przynajmniej raz może wygenerować zdanie prawdziwe, a więc nie jest kontrtautologią.

Przykład:

Zbadamy, czy kontrtautologią jest formuła: $\sim [(\sim p \vee q) \vee (q \rightarrow p)]$.

Ponieważ głównym spójnikiem badanego schematu jest negacja, musimy sprawdzić, czy istnieje możliwość, aby przy negacji tej pojawiła się wartość 1.

$$\sim [(\sim p \vee q) \vee (q \rightarrow p)]$$

$$\underline{1}$$

W kolejnych krokach wyciągamy wszelkie konsekwencje z przyjętego założenia. Jeżeli negacja ma być prawdziwa, to całe zdanie, do którego się ona odnosi (czyli alternatywa w kwadratowym nawiasie) musi być fałszywe. Jeśli fałszywa jest alternatywa, to fałszywe muszą być oba jej człony (zdania w nawiasach okrągłych). Otrzymujemy więc:

$$\sim [(\sim p \vee q) \vee (q \rightarrow p)]$$

$$\underline{1} \quad 0 \quad 0 \quad 0$$

W tym momencie mamy dwa miejsca, w których istnieje tylko jedna możliwość kombinacji zer i jedynek; nie jest istotne, od którego z nich zaczniemy. Gdy obliczymy najpierw wartość członów alternatywy w pierwszym nawiasie otrzymamy:

$$\sim [(\sim p \vee q) \vee (q \rightarrow p)]$$

$$1 \quad 0 \quad 1 \quad \underline{0} \quad 0 \quad 0 \quad 0$$

Po przepisaniu wartości zmiennych p i q do drugiego nawiasu otrzymujemy w nim ewidentną sprzeczność: implikacja z fałszywym poprzednikiem i prawdziwym następnikiem nie może być fałszywa.

$$\sim [(\sim p \vee q) \vee (q \rightarrow p)]$$

$$1 \quad 0 \quad 1 \quad 0 \quad 0 \quad 0 \quad \underline{0 \quad 0 \quad 1}$$

Widzimy zatem, że nie jest możliwa sytuacja, aby badana formuła okazała się schematem zdania prawdziwego; jest więc ona na pewno kontrtautologią.

Zauważmy na marginesie, że gdybyśmy najpierw obliczyli wartość członów implikacji w drugim nawiasie (gdzie też była tylko jedna możliwość), to otrzymalibyśmy sprzeczność przy alternatywie $\sim p \vee q$.

Przykład:

Zbadamy czy kontrtautologią jest formuła $\{(p \rightarrow q) \wedge \sim [(p \vee r) \rightarrow q]\} \wedge (q \rightarrow r)$.

Zaczynamy od postawienia symbolu 1 przy głównym spójniku, którym jest tu koniunkcja pomiędzy nawiasem klamrowym a okrągłym. Z prawdziwości tej koniunkcji wnosimy o prawdziwości obu jej członów, czyli koniunkcji w nawiasie klamrowym i implikacji w okrągłym:

$$\{(p \rightarrow q) \wedge \sim [(p \vee r) \rightarrow q]\} \wedge (q \rightarrow r)$$

$$\quad 1 \quad \quad \quad \underline{1} \quad \quad 1$$

Ponieważ prawdziwa jest koniunkcja w nawiasie klamrowym, prawdziwe muszą być oba jej człony: implikacja $p \rightarrow q$ oraz negacja wyrażenia w nawiasie kwadratowym. Jeżeli

prawdziwa jest negacja, to oczywiście fałszywe musi być zdanie, do którego się ona odnosi, czyli implikacja $(p \vee r) \rightarrow q$. Z kolei, jeśli fałszywa jest implikacja, to prawdziwy musi być jej poprzednik, a fałszywy następnik:

$$\{(p \rightarrow q) \wedge \sim [(p \vee r) \rightarrow q]\} \wedge (q \rightarrow r)$$

$$1 \quad \underline{1} \quad 1 \quad 1 \quad 0 \quad 0 \quad 1 \quad 1$$

Obliczoną wartość zmiennej q przepisujemy we wszystkie miejsca, gdzie zmienna ta występuje:

$$\{(p \rightarrow q) \wedge \sim [(p \vee r) \rightarrow q]\} \wedge (q \rightarrow r)$$

$$1 \quad 0 \quad 1 \quad 1 \quad 1 \quad 0 \quad \underline{0} \quad 1 \quad 0 \quad 1$$

Jedyne miejsce, w którym możemy coś wpisać ze stuprocentową pewnością, to pierwszy nawias okrągły. Jeżeli implikacja jest prawdziwa i jednocześnie ma fałszywy następnik, to fałszywy musi być również jej poprzednik. Oznaczamy więc p jako zdanie prawdziwe i przepisujemy tę wartość tam, gdzie jeszcze zdanie to występuje:

$$\{(p \rightarrow q) \wedge \sim [(p \vee r) \rightarrow q]\} \wedge (q \rightarrow r)$$

$$0 \quad \underline{1} \quad \underline{0} \quad 1 \quad 1 \quad 0 \quad 1 \quad 0 \quad 0 \quad 1 \quad 0 \quad 1$$

Obecnie możemy obliczyć wartość r w alternatywie $p \vee r$. Jeżeli alternatywa jest prawdziwa, a jeden jej człon jest fałszywy, to prawdziwy musi być człon drugi. Wpisujemy więc 1 przy zmiennej r i przepisujemy tę wartość pod r w implikacji w ostatnim nawiasie:

$$\{(p \rightarrow q) \wedge \sim [(p \vee r) \rightarrow q]\} \wedge (q \rightarrow r)$$

$$0 \quad 1 \quad 0 \quad 1 \quad 1 \quad 0 \quad 1 \quad 1 \quad 0 \quad 0 \quad 1 \quad 0 \quad 1 \quad 1$$

Ponieważ nigdzie nie występuje tu sprzeczność, pokazaliśmy, że badana formuła może być schematem zdania prawdziwego, a więc nie jest kontrtautologią.



1.4.5. CZĘSTO ZADAWANE PYTANIA.

Czy przy pomocy metody skróconej można od razu, „w jednej linijce” stwierdzić status logiczny formuły – zbadać czy jest ona tautologią, kontrtautologią czy też żadną z nich?

To zależy jak na to spojrzeć. Badanie czy formuła jest tautologią wymaga innego założenia, niż badanie czy jest kontrtautologią, więc w zasadzie należy zbadać przynajmniej dwie możliwości. Jednakże, gdy otrzymamy wynik „pozytywny” (to znaczy, że formuła jest tautologią lub jest kontrtautologią), to wiemy od razu, że nie jest ona niczym innym. Gdy natomiast otrzymamy wynik „negatywny”, to wiemy jedynie, że formuła czymś nie jest, dalej nie znając jej dokładnego statusu logicznego.

Czy formuła może „nie dać” się sprawdzić, czy jest tautologią lub kontrtautologią przy pomocy metody skróconej?

Sprawdzić przy pomocy metody skróconej da się zawsze. Jednakże czasami już na początku może pojawić się kilka możliwości do zbadania (na przykład gdyby ktoś chciał sprawdzić, czy tautologią jest formuła z koniunkcją jako głównym spójnikiem). W takich wypadkach metoda skrócona może stać się nieefektywna i wcale nie mniej pracochłonna od metody „zwykłej”.

1.5. PRAWDA LOGICZNA I ZDANIA WEWNĘTRZNIE SPRZECZNE.



ŁYK TEORII

1.5.1. ŁYK TEORII.

Jeśli schemat jakiegoś zdania języka naturalnego jest tautologią, to zdanie takie nazywamy **prawdą logiczną**. Zdanie będące prawdą logiczną jest prawdziwe ze względu na znaczenie tylko i wyłącznie użytych w nim spójników logicznych.

Zdania, których schematy są kontrtautologiami nazywamy **falszami logicznymi** lub **zdaniami wewnętrźnie sprzecznymi**. Zdania takie są fałszywe na mocy samych spójników logicznych, niezależnie od treści zdań składowych.

1.5.2. PRAKTYKA: SPRAWDZANIE, CZY ZDANIE JEST PRAWDĄ LOGICZNĄ LUB FAŁSZEM LOGICZNYM.

Sprawdzenie, czy zdanie jest prawdą logiczną jest bardzo proste i wymaga połączenia dwóch umiejętności: zapisywania schematu zdania oraz sprawdzania, czy schemat jest tautologią. Jeżeli schemat badanego zdania okaże się tautologią, stwierdzamy, że zdanie to jest prawdą logiczną, jeśli schemat tautologią nie jest, zdanie nie jest również prawdą logiczną.

Przykład:

Zbadamy bardzo proste zdanie: *Jutro będzie padać lub nie będzie padać.*

Schemat tego zdania, to oczywiście $p \vee \sim p$. Formuła $p \vee \sim p$ jest tautologią – gdybyśmy chcieli postawić 0 pod jej głównym spójnikiem, okazało by się, że zdanie p musi być jednocześnie prawdziwe i fałszywe, a więc otrzymalibyśmy sprzeczność.

$p \vee \sim p$
0 0 0 1

Ponieważ schemat zdania okazał się tautologią, to o zdaniu *Jutro będzie padać lub nie będzie padać* możemy powiedzieć, że jest ono prawdą logiczną. Łatwo zauważyć, że

faktycznie zdanie to nie może okazać się fałszywe – cokolwiek stanie się jutro, niezależnie jaka będzie pogoda, zdanie stwierdza coś, co na pewno się wydarzy.

Zauważmy, że takie bezwzględnie prawdziwe wyrażenia otrzymamy podstawiając dowolne zdanie za zmienną p w schemacie $p \vee \sim p$, na przykład *Zdam egzamin lub nie zdam egzaminu*, *Nasz prezes jest mądrym człowiekiem lub nie jest on mądrym człowiekiem* itp.

Przykład:

Sprawdźmy, czy prawdą logiczną jest zdanie: *O ile jest tak, że jeśli Jan jest zakochany, to jest zazdrosny, to jeśli Jan nie jest zazdrosny, to nie jest zakochany.*

Piszemy schemat zdania pamiętając o zastępowaniu tych samych zdań prostych tymi samymi zmiennymi:

$$(p \rightarrow q) \rightarrow (\sim q \rightarrow \sim p)$$

p – Jan jest zakochany, q – Jan jest zazdrosny.

Następnie sprawdzamy, czy powyższa formuła jest tautologią:

$$(p \rightarrow q) \rightarrow (\sim q \rightarrow \sim p)$$

<u>1</u>	<u>1</u>	<u>0</u>	0	1	0	0	1
----------	----------	----------	---	---	---	---	---

Okazuje się, że formuła nie może stać się schematem zdania fałszywego, a zatem jest tautologią. W związku z tym badane zdanie jest prawdą logiczną.

Przykład:

Sprawdźmy, czy prawdą logiczną jest zdanie: *Jeśli ten kamień jest diamentem, to przecina szkło lub jeśli nie jest diamentem, to nie przecina szkła.*

$$(p \rightarrow q) \vee (\sim p \rightarrow \sim q)$$

1	0	0	0	<u>0</u>	1	<u>0</u>	<u>1</u>	0
---	---	---	---	----------	---	----------	----------	---

Ponieważ schemat okazał się tautologią, badane zdanie jest prawdą logiczną.

Sprawdzenie, czy dane zdanie jest wewnętrznie sprzeczne jest równie proste. Jak łatwo się domyślić polega ono na napisaniu schematu zdania, a następnie zbadaniu, czy jest on kontrtautologią.

Przykład:

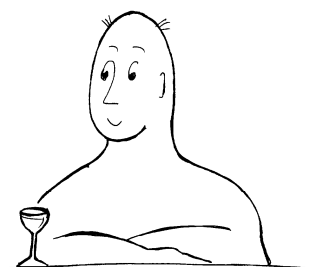
Zbadamy czy zdanie *Jeżeli jestem za, to nie jestem przeciw, ale ja jestem za i jestem przeciw* jest wewnętrznie sprzeczne.

$$(p \rightarrow \sim q) \wedge (p \wedge q)$$

$$\underline{1 \ 1 \ 0} \ 1 \ 1 \ 1 \ 1 \ 1$$

Ponieważ schemat badanego zdania jest kontrtautologią, samo zdanie jest wewnętrznie sprzeczne (jest fałszem logicznym).

1.6. WYNIKANIE LOGICZNE.



ŁYK TEORII

1.6.1. ŁYK TEORII.

Posługując się schematami zdań oraz tabelkami zero-jedynkowymi można sprawdzać poprawność logiczną prostych wnioskowań. W tym celu musimy najpierw zapoznać się z pojęciem wynikania logicznego.

Mówimy, że z pewnego zdania A wynika (w szerokim znaczeniu tego słowa) zdanie B, gdy nie jest możliwa sytuacja, aby zdanie A było prawdziwe, a jednocześnie B fałszywe. Czyli, ujmując rzecz inaczej, w przypadku gdy ze zdania A wynika zdanie B, to gdy tylko A jest prawdziwe, również prawdziwe musi być B.

I tak na przykład, ze zdania *Jan jest starszy od Piotra* wynika zdanie *Piotr jest młodszy od Jana*, bo nie jest możliwe, aby pierwsze było prawdziwe, a drugie fałszywe (lub, jak kto woli, gdy prawdziwe jest pierwsze zdanie, to i prawdziwe musi być drugie).

W logice pojęciem wynikania posługujemy się w bardzo ścisłym sensie, mówiąc o tak zwanym **wynikaniu logicznym**. W przykładzie powyżej mieliśmy do czynienia z wynikaniem w szerokim sensie, ale nie z wynikaniem logicznym. Stosunek wynikania uzależniony był tam od znaczenia słów „starszy” i „młodszy”; w przypadku wynikania logicznego to, że nie jest możliwa sytuacja, aby zdanie A było prawdziwe, a B fałszywe, uzależnione jest tylko i wyłącznie od obecnych w nich stałych logicznych (a więc, w przypadku rachunku zdań, od spójników logicznych).

To czy z jednego zdania wynika logicznie drugie możemy łatwo sprawdzić przy pomocy metody zero-jedynkowej, podobnie jak sprawdzamy, czy formuła jest tautologią lub kontrtautologią. Aby tego dokonać, musimy najpierw napisać schematy obu zdań. Schematy te piszemy na ogół w specjalnej formie – schemat pierwszego nad kreską, a pod kreską schemat drugiego:

schemat zdania A

schemat zdania B

Następnie sprawdzamy, czy jest możliwa sytuacja, aby zdanie A było prawdziwe, a B fałszywe. Wpisujemy symbol 1 przy głównym spójniku zdania A, a 0 przy głównym spójniku zdania B i wyciągamy z takich założeń wszelkie konsekwencje – podobnie jak to czyniliśmy

przy badaniu tautologii i kontrtautologii. Gdy okaże się, że ostatecznie nigdzie nie wystąpi sprzeczność, będzie to oznaczać, że sytuacja gdzie zdanie A jest prawdziwe, a B fałszywe może zaistnieć, a więc, zgodnie z definicją wynikania, że zdania A nie wynika logicznie zdanie B. Gdy natomiast wyciągając konsekwencje z przyjętego założenia dojdziemy do sprzeczności, będzie to wskazywać, że nie jest możliwe aby A było prawdziwe a B fałszywe, a zatem, że ze zdania A wynika logicznie zdanie B.



DO ZAPAMIĘTANIA:

W skrócie metoda badania czy z jednego zdania wynika zdanie drugie wygląda następująco:

- piszemy schematy zdań;
- zakładamy, że pierwsze zdanie jest prawdziwe, a drugie fałszywe;
- wyciągając z założonej sytuacji konsekwencje, sprawdzamy, czy może ona wystąpić;
- jeżeli otrzymamy sprzeczność, świadczy to, że ze zdania A wynika logicznie zdanie B; jeśli sprzeczności nie ma, ze zdania A nie wynika B.

1.6.2. PRAKTYKA: SPRAWDZANIE, CZY Z JEDNEGO ZDANIA WYNIKA DRUGIE.

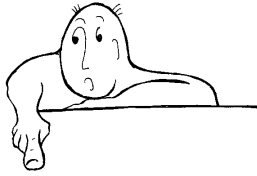
Przykład:

Sprawdźmy, czy ze zdania *Gospodarka rozwija się dobrze wtedy i tylko wtedy, gdy podatki nie są wysokie*, wynika logicznie zdanie *Jeżeli podatki są wysokie, to gospodarka nie rozwija się dobrze*.

Schematy powyższych zdań wyglądają następująco:

$$\frac{p \equiv \sim q}{q \rightarrow \sim p}$$

p – gospodarka rozwija się dobrze, q – podatki są wysokie.



Uwaga na błędy!

Należy bezwzględnie pamiętać o zastępowaniu tych samych zdań prostych występujących w różnych miejscach przez te same zmienne.

Sprawdzamy teraz, czy może zajść sytuacja, aby pierwsze zdanie było prawdziwe, a drugie fałszywe.

$$\begin{array}{r} 1 \\ p \equiv \sim q \\ \hline q \rightarrow \sim p \\ 0 \end{array}$$

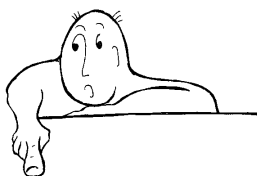
Z fałszywości implikacji możemy określić wartości logiczne zmiennych p oraz q i przenieść je do pierwszego zdania:

$$\begin{array}{r} 1 \quad 1 \quad 1 \\ p \equiv \sim q \\ \hline q \rightarrow \sim p \\ 1 \quad 0 \quad 0 \quad 1 \end{array}$$

Gdy na podstawie prawdziwości q obliczymy wartość prawej strony równoważności otrzymamy ewidentną sprzeczność – prawdziwą równoważność z jednym członem prawdziwym, a drugim fałszywym.

$$\begin{array}{r} \underline{1 \quad 1 \quad 0} \quad 1 \\ p \equiv \sim q \\ \hline q \rightarrow \sim p \\ 1 \quad 0 \quad 0 \quad 1 \end{array}$$

Widzimy zatem, że sytuacja aby pierwsze zdanie było prawdziwe, a drugie fałszywe nie jest możliwa. Możemy zatem powiedzieć, że ze zdania *Gospodarka rozwija się dobrze wtedy i tylko wtedy, gdy podatki nie są wysokie* wynika logicznie zdanie *Jeżeli podatki są wysokie, to gospodarka nie rozwija się dobrze*.



Uwaga na błędy!

W opisany wyżej sposób sprawdzamy zawsze, czy z pierwszego zdania wynika zdanie drugie, a nie na odwrót. Zdarza się, iż niektórzy nie zwracają uwagi na tę istotną różnicę i na zasadzie „coś z czegoś wynika” beztrząsco dają odpowiedź: *zdanie pierwsze wynika z drugiego*. Jest to bardzo duży błąd.

Przykład:

Sprawdzimy, czy ze zdania *Jeśli na imprezie był Zdzisiek i Wacek, to impreza się nie udała*, wynika logicznie zdanie *Jeśli impreza się nie udała, to był na niej Zdzisiek lub Wacek*.

Schematy powyższych zdań wyglądają następująco:

$$(p \wedge q) \rightarrow \sim r$$

$$\sim r \rightarrow (p \vee q)$$

Sprawdzamy teraz, czy możliwa jest sytuacja, aby pierwsze zdanie było prawdziwe, a drugie fałszywe.

$$\begin{array}{c} 1 \\ (p \wedge q) \rightarrow \sim r \end{array}$$

$$\sim r \rightarrow (p \vee q)$$

0

Z fałszywości implikacji na dole łatwo obliczamy wartości $\sim r$, oraz $p \vee q$, a następnie samych zmiennych p , q i r . Wartości tych zmiennych przenosimy do pierwszego zdania:

$$\begin{array}{cccc} 0 & 0 & 1 & 0 \\ (p \wedge q) \rightarrow \sim r \end{array}$$

$$\sim r \rightarrow (p \vee q)$$

$$1 \ 0 \ \underline{0} \ 0 \ 0 \ 0$$

Po obliczeniu wartości koniunkcji p i q oraz negacji r , okazuje się, że w badanych schematach wszystko się zgadza – nie ma żadnej sprzeczności:

$$0 \ 0 \ 0 \ 1 \ 1 \ 0$$

$$(p \wedge q) \rightarrow \sim r$$

$$\sim r \rightarrow (p \vee q)$$

$$1 \ 0 \ \underline{0} \ 0 \ 0 \ 0$$

Brak sprzeczności świadczy, że jak najbardziej możliwa jest sytuacja, aby pierwsze zdanie było prawdziwe, a drugie fałszywe. Stwierdzamy zatem, że w tym wypadku zdanie drugie nie wynika logicznie ze zdania pierwszego.

1.6.3. WYKORZYSTANIE POJĘCIA TAUTOLOGII.

Do sprawdzania, czy z jednego zdania wynika logicznie drugie zdanie, wykorzystać można również pojęcie tautologii. Jedno z ważniejszych twierdzeń logicznych, tak zwane twierdzenie o dedukcji, głosi bowiem co następuje: ze zdania A wynika logicznie zdanie B wtedy i tylko wtedy, gdy formuła $A \rightarrow B$ jest tautologią.

Aby, posługując się twierdzeniem o dedukcji, sprawdzić czy z jednego zdania wynika drugie, musimy napisać schematy tych zdań, następnie połączyć je spójnikiem implikacji, po czym sprawdzić, czy tak zbudowana formuła jest tautologią. Jeśli formuła jest tautologią, to oznacza to, iż ze zdania pierwszego wynika logicznie zdanie drugie; jeśli formuła tautologią nie jest, wynikanie nie zachodzi.

Przykład:

Sprawdzimy, tym razem przy pomocy twierdzenia o dedukcji, rozpatrywany już przykład – czy ze zdania *Gospodarka rozwija się dobrze wtedy i tylko wtedy, gdy podatki nie są wysokie*, wynika logicznie zdanie *Jeżeli podatki są wysokie, to gospodarka nie rozwija się dobrze*.

Formuła powstała z połączenia implikacją schematów zdań wygląda następująco:

$$(p \equiv \sim q) \rightarrow (q \rightarrow \sim p)$$

Sprawdzenie, czy jest ona tautologią jest bardzo proste:

$$(p \equiv \sim q) \rightarrow (q \rightarrow \sim p)$$

1	1	0	1	0	1	0	0	1
----------	----------	----------	---	---	---	---	---	---

Otrzymana sprzeczność świadczy, że formuła jest tautologią, a więc, zgodnie z twierdzeniem o dedukcji, ze zdania pierwszego wynika logicznie zdanie drugie.

Przykład:

Sprawdzimy przy pomocy twierdzenia o dedukcji czy ze zdania *Jeśli na imprezie był Zdzisiek i Wacek, to impreza się nie udała*, wynika logicznie zdanie *Jeśli nie było Zdziska i nie było Wacka, to impreza udała się*.

Po połączeniu implikacją schematów powyższych zdań otrzymujemy formułę:

$$[(p \wedge q) \rightarrow \sim p] \rightarrow [(\sim p \wedge \sim q) \rightarrow r]$$

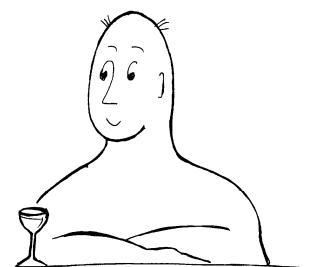
Po założeniu 0 pod głównym spójnikiem otrzymujemy ostatecznie:

$$[(p \wedge q) \rightarrow \sim r] \rightarrow [(\sim p \wedge \sim q) \rightarrow r]$$

$$000 \ 110 \ 0 \quad 101 \ 100 \ 0$$

Brak sprzeczności świadczy, że formuła nie jest tautologią. A zatem ze zdania pierwszego nie wynika logicznie zdanie drugie.

1.7. WNIOSKOWANIA.



ŁYK TEORII

1.7.1. ŁYK TEORII.

Wnioskowanie jest to proces myślowy, podczas którego na podstawie uznania za prawdziwe pewnych zdań (przesłanek) dochodzimy do uznania kolejnego zdania (konkluzji). Gdy ktoś na podstawie wiary, iż *jeśli jaskółki rano nisko latają, to po południu będzie deszcz*, oraz faktu, iż *dziś rano jaskółki nisko latają*, dochodzi do wniosku, że *dziś po południu będzie padać*, to jest to właśnie wnioskowanie.

Badanie logicznej poprawności wnioskowania wiąże się ściśle z pojęciem wynikania logicznego. Mówimy bowiem, iż wnioskowanie jest poprawne, jeśli wniosek wynika logicznie z przesłanek. Gdy badaliśmy, czy z jednego zdania wynika logicznie drugie zdanie, sprawdzaliśmy jednocześnie, jeszcze o tym nie wiedząc, poprawność bardzo prostego wnioskowania, w którym pierwsze zdanie pełni rolę jedynej przesłanki, a drugie wniosku. Obecnie zajmujemy się wnioskowaniami z większą ilością przesłanek.

Sprawdzenie poprawności wnioskowania rozpoczynamy od napisania schematów wszystkich zdań wchodzących w jego skład. Schematy przesłanek piszemy nad kreską, schemat wniosku pod kreską. Taki, znany już z poprzedniego rozdziału, układ schematów nazywamy **regułą wnioskowania** (lub regułą inferencji, albo po prostu regułą).

Nazwa „reguła” mogłaby sugerować, że jest to coś zawsze poprawnego – tak jednak nie jest; wśród reguł wyróżniamy bowiem **reguły dedukcyjne** (inaczej mówiąc **niezawodne**) i reguły **niededukcyjne** (**zawodne**). Reguła dedukcyjna (niezawodna), to taka, w której wniosek wynika logicznie z przesłanek, natomiast w przypadku reguły niededukcyjnej (zawodnej) wniosek nie wynika logicznie z przesłanek.

Badanie dedukcyjności reguły przeprowadzamy sprawdzając, czy możliwa jest sytuacja, aby wszystkie przesłanki były prawdziwe, a jednocześnie wniosek fałszywy. Jeśli sytuacja taka może wystąpić (nigdzie nie pojawia się sprzeczność) to znaczy to, że dana reguła jest niededukcyjna (zawodna), a to z kolei świadczy o tym, że oparte na tej regule wnioskowanie jest z logicznego punktu widzenia niepoprawne. Gdy natomiast założenie prawdziwości przesłanek i fałszywości wniosku doprowadzi do sprzeczności, świadczy to, że mamy do czynienia z regułą dedukcyjną (niezawodną), a zatem oparte na niej wnioskowanie jest poprawne.



DO ZAPAMIĘTANIA:

W skrócie sprawdzenie poprawności wnioskowania wygląda następująco:

- piszemy schematy zdań w postaci reguły;
- zakładamy, że wszystkie przesłanki są prawdziwe, a wniosek fałszywy;
- wyciągając z założonej sytuacji konsekwencje, sprawdzamy, czy może ona faktycznie wystąpić;
- jeżeli otrzymamy sprzeczność, świadczy to, że reguła jest dedukcyjna (niezawodna): wniosek wynika logicznie z przesłanek, a zatem badane wnioskowanie jest poprawne; jeśli sprzeczności nie ma, to znak, że reguła jest niededukcyjna (zawodna): wniosek nie wynika z przesłanek, a więc wnioskowanie jest logicznie niepoprawne.

1.7.2. PRAKTYKA: SPRAWDZANIE POPRAWNOŚCI WNIOSKOWAŃ.

Przykład:

Sprawdzimy poprawność wnioskowania: *Jeśli Wacek dostał wypłatę to jest w barze lub u Zenka. Wacka nie ma w barze. Zatem Wacek nie dostał wypłaty.*

We wnioskowaniu tym widzimy dwa zdania stanowiące przesłanki oraz oczywiście zdanie będące wnioskiem. Wniosek poznajemy zwykle po zwrotach typu „zatem”, „a więc” itp. Schematy zdań ułożone w formie reguły, na której opiera się powyższe wnioskowanie, wyglądają następująco:

$$p \rightarrow (q \vee r), \sim q$$

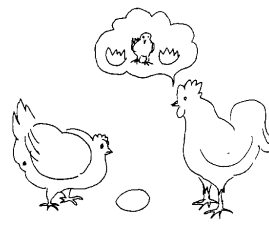
$$\sim p$$

Badając, czy reguła jest niezawodna, a więc, czy wniosek wynika z przesłanek, sprawdzamy, czy możliwa jest sytuacja aby wszystkie przesłanki były prawdziwe, a jednocześnie wniosek fałszywy:

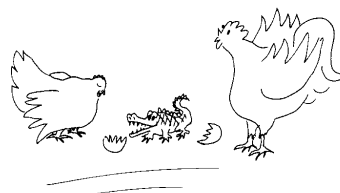
$$\begin{array}{cc} 1 & 1 \\ p \rightarrow (q \vee r), \sim q \end{array}$$

$$\begin{array}{c} \sim p \\ 0 \end{array}$$

PRZESŁANKI PRAWDZIWE



A WNIOSEK FAŁSZYWY



REGUŁA ZAWODNA

Dalsze kroki, które musimy wykonać przedstawiają się następująco: obliczamy wartości zdań p oraz q na podstawie znajomości wartości ich negacji; następnie przepisujemy te wartości i wiedząc, iż prawdziwa implikacja z prawdziwym poprzednikiem musi mieć prawdziwy następnik, wpisujemy wartość 1 nad spójnikiem alternatywy; znając wartość alternatywy oraz jednego z jej członów – q , obliczamy wartość $r - 1$:

$$\begin{array}{cccccc} 1 & 1 & 0 & 1 & 1 & 1 & 0 \\ p \rightarrow (q \vee r), \sim q & & & & & & \\ \hline & & \sim p & & & & \\ & & 0 & 1 & & & \end{array}$$

Ponieważ przy takich podstawieniach nie pojawia się nigdzie sprzeczność, wykazaliśmy że możliwa jest sytuacja, aby przesłanki były prawdziwe, a wniosek fałszywy. Powyższa reguła jest zatem zawodna, czyli jej wniosek nie wynika z przesłanek. Na podstawie tych faktów możemy dać ostateczną odpowiedź, iż badane wnioskowanie nie jest poprawne.

Przykład:

Zbadamy teraz poprawność wnioskowania będącego modyfikacją rozumowania z poprzedniego przykładu. *Jeśli Wacek dostał wypłatę to jest w barze lub u Zenka. Wacka nie ma w barze. Zatem Wacek nie dostał wypłaty lub jest u Zenka.*

Badając regułę, na której oparte jest wnioskowanie zaczynamy następująco:

$$\begin{array}{cccccc} 1 & & & & 1 & \\ p \rightarrow (q \vee r), \sim q & & & & & \\ \hline & & \sim p \vee r & & & \\ & & 0 & & & \end{array}$$

Następnie obliczamy wartości członów alternatywy we wniosku oraz wartość q . Wartości te przepisujemy do pierwszej przesłanki i stwierdzamy, że fałszywa musi być alternatywa $(q \vee r)$, ponieważ fałszywe są oba jej człony. Po bliższym przyjrzeniu się implikacji odkrywamy w niej sprzeczność:

$$\begin{array}{cccccc} \underline{1} & \underline{1} & 0 & \underline{0} & 0 & 1 & 0 \\ p \rightarrow (q \vee r), \sim q & & & & & & \\ \hline & & \sim p \vee r & & & & \\ & & 0 & 1 & 0 & 0 & \end{array}$$

Pokazaliśmy, że tym razem nie jest możliwa sytuacja, aby przesłanki były prawdziwe, a wniosek fałszywy. Powyższa reguła jest zatem niezawodna, a badane wnioskowanie poprawne.

UWAGA!

Badając dedukcyjność reguł, podobnie jak przy sprawdzaniu czy formuła jest tautologią lub kontrtautologią, sprzeczności mogą pojawić się w różnych miejscach. Na przykład w powyższym przykładzie ostateczny wynik mógł wyglądać następująco:

$$\begin{array}{cccc}
 1 & 1 & \underline{0 \ 1 \ 0} & 1 \ 0 \\
 p \rightarrow (q \vee r), \sim q & & & \\
 \hline
 & \sim p \vee r & & \\
 & 0 \ 1 \ 0 \ 0 & &
 \end{array}$$

Oczywiście jest to równie dobre rozwiązanie.

Przykład:

Sprawdzimy poprawność następującego wnioskowania: *Jeśli „Lolek” jest agentem, to agentem jest też „Bolek”, zaś nie jest nim „Tola”. Jeśli „Bolek” jest agentem, to jest nim też „Lolek” lub „Tola”. Jeśli jednak „Tola” nie jest agentem, to jest nim „Lolek” a nie jest „Bolek”. Tak więc to „Tola” jest agentem.*

Reguła na której oparte jest powyższe wnioskowanie wygląda następująco:

$$\begin{array}{c}
 p \rightarrow (q \wedge \sim r), \ q \rightarrow (p \vee r), \ \sim r \rightarrow (p \wedge \sim q) \\
 \hline
 r
 \end{array}$$

Po założeniu prawdziwości przesłanek oraz fałszywości wniosku, a następnie przepisaniu wszędzie wartości r otrzymujemy:

$$\begin{array}{cccccc}
 1 & & 0 & & 1 & & 0 & & 0 \ 1 \\
 p \rightarrow (q \wedge \sim r), \ q \rightarrow (p \vee r), \ \sim r \rightarrow (p \wedge \sim q) & & & & & & & & \\
 \hline
 & & & & r & & & & \\
 & & & & 0 & & & &
 \end{array}$$

Teraz możemy obliczyć wartość negacji r. W trzeciej przesłance mając prawdziwą implikację z prawdziwym poprzednikiem stwierdzamy, że prawdziwy musi być jej następnik – koniunkcja $p \wedge \sim q$. Teraz łatwo obliczamy wartości p oraz q i przepisujemy je. Po

obliczeniu wartości koniunkcji w pierwszej przesłance oraz alternatywy w drugiej otrzymujemy:

$$\begin{array}{cccccccccccc} \underline{1} & \underline{1} & 0 & 0 & 1 & 0 & 0 & 1 & 1 & 1 & 0 & 1 & 0 & 1 & 1 & 1 & 0 \\ p \rightarrow (q \wedge \sim r), & q \rightarrow (p \vee r), & \sim r \rightarrow (p \wedge \sim q) \\ \hline & & & & & & & & r & & & & & & & & \\ & & & & & & & & 0 & & & & & & & & \end{array}$$

Sprzeczność w pierwszej przesłance pokazuje, iż nie jest możliwa sytuacja, aby przesłanki były prawdziwe, a wniosek fałszywy. Wnioskowanie jest więc poprawne.

1.7.3. WYKORZYSTANIE POJĘCIA TAUTOLOGII.

Do sprawdzenia poprawności wnioskowania można również wykorzystać pojęcie tautologii, w podobny sposób, jak to czyniliśmy przy okazji sprawdzania, czy z jednego zdania wynika logicznie drugie zdanie. Twierdzenie o dedukcji mówi bowiem, że reguła jest niezawodna (a zatem oparte na niej wnioskowanie poprawne) gdy tautologią jest implikacja, której poprzednik stanowią połączone spójnikami koniunkcji przesłanki, a następnik – wniosek.

Przykład:

Zbadamy przy pomocy twierdzenia o dedukcji następujące wnioskowanie:

Jeżeli to nie Ted zastrzelił Billa, to zrobił to John. Jeśli zaś John nie zastrzelił Billa, to zrobił to Ted lub Mike. Ale Mike nie zastrzelił Billa. Zatem to Ted zastrzelił Billa.

Reguła na której opiera się wnioskowanie wygląda następująco:

$$\begin{array}{c} \sim p \rightarrow q, \sim q \rightarrow (p \vee r), \sim r \\ \hline p \end{array}$$

Aby móc skorzystać z twierdzenia o dedukcji musimy zbudować implikację, której poprzednik będą stanowić połączone spójnikami koniunkcji przesłanki, a następnik – wniosek. Praktycznie czynimy to tak, że bierzemy w nawias pierwszą przesłankę, łączymy ją koniunkcją z wziętą w nawias drugą przesłanką, bierzemy powstałe wyrażenie w nawias i łączymy koniunkcją z wziętą w nawias trzecią przesłanką, następnie bierzemy wszystkie przesłanki w jeden największy nawias i łączymy to wyrażenie z wnioskiem przy pomocy symbolu implikacji:

$$\{(\sim p \rightarrow q) \wedge [\sim q \rightarrow (p \vee r)]\} \wedge \sim r \rightarrow p$$

Następnie sprawdzamy, czy formuła ta jest tautologią. Ponieważ w powyższym schemacie mamy bardzo dużo nawiasów, trzeba to robić bardzo uważnie. Ważne jest, aby dobrze zlokalizować główny spójnik poprzednika implikacji:

$$\langle \{(\sim p \rightarrow q) \wedge [\sim q \rightarrow (p \vee r)]\} \wedge \sim r \rangle \rightarrow p$$

1 0 0

Ponieważ mamy prawdziwą koniunkcję, to prawdziwe muszą być oba jej człony – koniunkcja w nawiasie klamrowym oraz $\sim r$. Znowu mamy prawdziwą koniunkcję, z czego wnioskujemy o prawdziwości implikacji $\sim p \rightarrow q$ oraz $\sim q \rightarrow (p \vee r)$. Wartości p i r możemy przepisać tam, gdzie zmienne te jeszcze występują:

$$\langle \{(\sim p \rightarrow q) \wedge [\sim q \rightarrow (p \vee r)]\} \wedge \sim r \rangle \rightarrow p$$

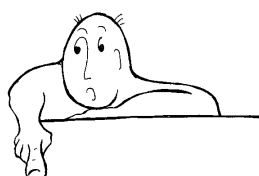
1 0 1 1 1 0 0 1 1 0 0 0

W pierwszym nawiasie mając prawdziwą implikację z prawdziwym poprzednikiem możemy obliczyć wartość q – 1. Po przepisaniu jej oraz obliczeniu wartości $\sim q$ i alternatywy $p \vee r$ otrzymujemy:

$$\langle \{(\sim p \rightarrow q) \wedge [\sim q \rightarrow (p \vee r)]\} \wedge \sim r \rangle \rightarrow p$$

1 0 1 1 1 0 1 1 0 0 0 1 1 0 0 0

Przy takim podstawieniu symboli 0 i 1 w badanej formule nie występuje nigdzie sprzeczność. Formuła nie jest więc tautologią, z czego wnioskujemy, że reguła na której opiera się wnioskowanie jest zawodna, a samo wnioskowanie niepoprawne.



Uwaga na błędy!

W powyższym przykładzie badaliśmy niezawodność (dedukcyjność) reguły korzystając z pojęcia tautologii. Nie wolno jednak mylić pojęć i mówić na przykład, że reguła jest (bądź nie jest) tautologią, albo że formuła jest (lub nie jest) dedukcyjna. Podkreślmy więc:

Tautologią może być (lub nie być) pojedyncza formuła.

Dedukcyjna (niezawodna) może być (lub nie być) reguła, czyli ciąg formuł.

Można badać dedukcyjność reguły korzystając z pojęcia tautologii, ale wtedy musimy najpierw zbudować odpowiednią formułę.



1.7.4. CZĘSTO ZADAWANE PYTANIA.

Czym wnioskowanie różni się od wynikania?

Wnioskowanie to pewien proces myślowy zachodzący w głowie rozumującej osoby, lub przykładowo zapisany na papierze. Wynikanie natomiast to związek mogący zachodzić pomiędzy przesłankami i wnioskiem. Wnioskowanie może być logicznie poprawne – wtedy gdy między przesłankami a wnioskiem zachodzi stosunek wynikania, lub logicznie niepoprawne, gdy stosunek taki nie zachodzi.

Czym różni się sprawdzenie poprawności wnioskowania, od sprawdzenia, czy z jednego zdania wynika logicznie drugie zdanie?

Praktycznie niczym się nie różni. Wnioskowania mogą mieć różną ilość przesłanek: jedną, dwie, trzy,... dziesięć,... sześćdziesiąt itd. Sprawdzając czy wnioskowanie jest poprawne, sprawdzamy czy wniosek wynika logicznie z przesłanek. Gdy mamy wnioskowanie z tylko jedną przesłanką, po prostu sprawdzamy, czy wniosek z niej wynika, a więc czy z jednego zdania wynika drugie zdanie. Mówiąc jeszcze inaczej: sprawdzenie, czy z jednego zdania wynika drugie zdanie jest po prostu sprawdzeniem poprawności wnioskowania mającego tylko jedną przesłankę.

SŁOWNICZEK.

Amfibolia – wyrażenie wieloznaczne, dopuszczające kilka możliwości interpretacji. Na gruncie rachunku zdań amfiboliami są wyrażenia, w których nie jest jednoznacznie określony spójnik główny. Np. $p \vee q \rightarrow r$ może być rozumiane jako implikacja $(p \vee q) \rightarrow r$, bądź też jako alternatywa $p \vee (q \rightarrow r)$. W języku naturalnym amfibolią jest na przykład zdanie: *Oskarżony zakopał łup wraz z teściową.*

Falsz logiczny – (zdanie wewnętrznie sprzeczne) – zdanie, którego schematem jest kontrtautologia.

Formuła – według ścisłej definicji formuła jest to wyrażenie zawierające zmienne. Możemy również powiedzieć, iż formułą danego rachunku logicznego nazywamy każde poprawnie zbudowane wyrażenie tego rachunku. Formułami klasycznego rachunku zdań są np.: p , $\sim q$, $(p \wedge q) \equiv \sim r$, $p \vee \sim (r \rightarrow s)$, natomiast nie są formułami tego rachunku wyrażenia: $p \sim q$, $\rightarrow (p \wedge q)$, $p \equiv \vee q$.

Kontrtautologia – formuła będąca schematem wyłącznie zdań fałszywych.

Prawda logiczna – zdanie, którego schematem jest tautologia.

Reguła – (reguła wnioskowania, reguła inferencji) ciąg formuł wśród których wyróżnione są przesłanki i wniosek. Można powiedzieć, że reguła jest schematem całego wnioskowania, tak jak formuła jest schematem pojedynczego zdania.

Reguła dedukcyjna – (reguła niezawodna) – reguła w której niemożliwe jest, aby przesłanki stały się schematami zdań prawdziwych, natomiast wniosek schematem zdania fałszywego. Oparte na takiej regule wnioskowanie jest logicznie poprawne (dedukcyjne).

Schemat główny zdania – jest to schemat zawierający wszystkie spójniki logiczne dające się wyodrębnić w zdaniu (najdłuższy możliwy schemat danego zdania). Np. w przypadku zdania *Jeżeli nie zarobię wystarczająco dużo lub obleję sesję na uczelni to nie pojadę na wakacje*, formuła $p \rightarrow q$ (p – *nie zarobię wystarczająco dużo lub obleję sesję na uczelni*, q – *nie pojadę na wakacje*) nie jest jego schematem głównym. Schemat główny tego zdania wygląda następująco: $(\sim p \vee q) \rightarrow \sim r$. (p – *zarobię wystarczająco dużo*, q – *obleję sesję na uczelni*, r – *pojadę na wakacje*). Mówiąc „schemat zdania” rozumiemy przez to na ogół domyślnie schemat główny.

Spójnik główny – spójnik niejako wiążący w całość całą formułę. W każdej formule musi być taki spójnik i może być on tylko jeden. W formule $(p \vee q) \rightarrow r$ spójnikiem głównym jest implikacja, w formule $p \vee (q \rightarrow r)$ – alternatywa, natomiast w $\sim [(p \vee q) \rightarrow r]$ negacja.

Spójnik logiczny – spójnikami logicznymi są wyrażenia *nieprawda, że; lub; i; jeśli...,to...;wtedy i tylko wtedy* w znaczeniu ściśle zdefiniowanym w tabelkach zero-jedynkowych.

Stała logiczna – stałe logiczne wraz ze zmiennymi i znakami interpunkcyjnymi (nawiasami) składają się na język danego rachunku logicznego. Do stałych logicznych KRZ zaliczamy spójniki logiczne.

Tautologia – formuła będąca schematem wyłącznie prawdziwych zdań. Innymi słowy, tautologia jest to formuła, która nie jest w stanie stać się schematem zdania fałszywego, niezależnie od tego, jakie zdania podstawialibyśmy za obecne w niej zmienne.

Wartość logiczna zdania – prawdziwość lub fałszywość zdania.

Wnioskowanie – proces myślowy, podczas którego na podstawie uznania za prawdziwe pewnych zdań (przesłanek) dochodzimy do uznania kolejnego zdania (konkluzji).

Zdanie – mówiąc „zdanie” rozumiemy przez to w logice „zdanie w sensie logicznym”. Zdaniem w sensie logicznym są tylko zdania oznajmujące.

Zdanie proste – zdanie w którym nie występuje żaden spójnik logiczny.

Zmienna zdaniowa – symbol, za który można podstawić zdanie. W klasycznym rachunku zdań zmienne zdaniowe symbolizowane są na ogół przez litery p, q, r, s , itd.

Rozdział II

SYLOGISTYKA.

WSTĘP.

Opisany w poprzednim rozdziale klasyczny rachunek zdań nie jest niestety narzędziem nadającym się do analizy wszelkich rozumowań. Aby się o tym przekonać, rozważmy następujące rozumowanie: *Każdy jamnik jest psem. Każdy pies jest ssakiem. Zatem każdy jamnik jest ssakiem.* Nawet dla osoby nie znającej logiki powinno być oczywiste, że jest to rozumowanie poprawne. Ci, którzy choć w zarysach przypominają sobie pojęcie wynikania logicznego łatwo zauważą, że nie jest możliwe, aby przesłanki były prawdziwe, a wniosek fałszywy, a więc wniosek, jak się wydaje, wynika z przesłanek. Spróbujmy jednak zbadać powyższe rozumowanie na gruncie rachunku zdań. Ponieważ ani przesłanki, ani wniosek nie zawierają w sobie spójników logicznych, ich schematami będą reprezentujące zdania proste pojedyncze zmienne zdaniowe. Reguła, na której wnioskowanie to jest oparte, wygląda zatem następująco:

$$\frac{p, q}{r}$$

Reguła ta nie jest oczywiście dedukcyjna, gdyż nic nie stoi na przeszkodzie, aby zaszła sytuacja:

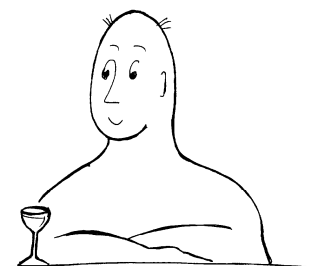
$$\frac{1 \quad 1}{p, q} \\ \frac{r}{0}$$

Jaki morał wynika z powyższego przykładu? Ktoś mógłby powiedzieć, że logika jest sprzeczna ze zdrowym rozsądkiem – rozumowanie w sposób oczywisty poprawne okazało się na gruncie logiki błędnym. Nie jest to jednak dobry wniosek. Prawda jest taka, że do analizy powyższego przykładu użyliśmy niewłaściwego narzędzia. Zamiast rachunku zdań należało tu bowiem wykorzystać system nazywany **sylogistyką** (teorią sylogizmów) lub czasem **rachunkiem nazw**.

Na marginesie dodajmy, że sylogistyka jest najstarszym systemem logicznym – opracowana została w IV w p.n.e przez greckiego filozofa Arystotelesa.

2.1. SCHEMATY ZDAŃ.

2.1.1. ŁYK TEORII.



ŁYK TEORII

Podobnie jak to było w przypadku rachunku zdań, poznanie teorii sylogizmów rozpoczniemy od nauki zapisywania schematów zdań. Na gruncie sylogistyki rolę stałych logicznych pełnią nie spójniki zdaniowe, ale cztery następujące zwroty: **każde... jest...**, **żaden... nie jest...**, **niektóre... są...**, **niektóre... nie są...**. Sporządzanie schematów zdań polegać będzie na wyszukiwaniu tych zwrotów i zastępowaniu ich odpowiednimi symbolami.

Przyjęło się, że zwrot **każde... jest...** oznaczany jest symbolem litery „a”, **żaden... nie jest...** – litery „e”, **niektóre... są...** – „i”, **niektóre... nie są...** – „o”. Łatwo zauważyć, że aby przy użyciu takich zwrotów powstały sensowne wyrażenia, w miejscach wykropkowanych znajdować się powinny nazwy, na przykład **każdy pies jest ssakiem**, **żaden student nie jest analfabetą**, **niektórzy politycy nie są złodziejami** itp. Z tego właśnie powodu, że elementami łączonymi przez stałe logiczne są tu nazwy, sylogistyka nazywana jest rachunkiem nazw.

każde ... jest...

żadne ... nie jest...

niektóre ... są ...

niektóre ... nie są ...

PRZY ZDANIACH TWIERDZĄCYCH ZAPAMIĘTUJEMY
DRUGĄ, LITERĘ, PIERWSZEGO WYRAZU, PRZY
ZDANIACH PRZECZĄCYCH – SZÓSTĄ, I PIĄTĄ.

W tym miejscu konieczne jest małe wyjaśnienie odnośnie nazw. Nikt nie ma wątpliwości, że nazwami są takie wyrażenia jak *pies*, *ssak*, *student*, czy *złodziej*. Trzeba jednak koniecznie zaznaczyć, że nazwa wcale nie musi składać się tylko z jednego rzeczownika – nazwami są również na przykład takie wyrażenia jak *duży pies*, *pilny student uniwersytetu*, czy też *złodziej poszukiwany listem gończym w całym kraju*. Nazwy nie muszą też odnosić się jedynie do obiektów fizycznych – mogą one wskazywać również „byty” bardziej abstrakcyjne – na przykład uczucia, własności czy też procesy dziejące się w czasie.

Nazwami są więc wyrażenia takie jak *wielka miłość, żelazne zdrowie, egzamin z logiki, strach przed sprawdzianem, wyprawa w kosmos lub zapalenie wyrostka robaczkowego*.

Obiekty wskazywane przez nazwy określamy mianem **desygnatów** danej nazwy. Tak więc na przykład każdy z nas jest desygnatem nazwy *człowiek*. Zbiór wszystkich desygnatów nazwy to **zakres** (lub inaczej: **denotacja**) nazwy.

Problematyka nazw dokładniej zostanie omówiona w rozdziale IV.

Zmienne odpowiadające nazwom w schematach sylogistycznych przyjęło się oznaczać przy pomocy dużych liter S oraz P – symbole te pochodzą one od łacińskich nazw *subiectum* – podmiot, oraz *praedicatum* – orzecznik.

Ponieważ w sylogistyce mamy tylko cztery stałe logiczne, a każda z nich może łączyć tylko dwie nazwy, w systemie tym istnieje możliwość napisania jedynie czterech rodzajów schematów: S a P – oznaczający zdanie *każde S jest P*, S e P – *żadne S nie jest P*, S i P – *niektóre S są P* (lub: *istnieją S będące P*), oraz S o P – *niektóre S nie są P* (lub: *istnieją S nie będące P*). Zdania tych czterech typów nazywamy **zdaniami kategorycznymi**.

Zdania kategoryczne typu *każde S jest P* oraz *żadne S nie jest P* nazywamy **zdaniami ogólnymi** – ponieważ stwierdzają one pewien fakt dotyczących wszystkich obiektów objętych nazwą S; zdania typu *niektóre S są P* oraz *niektóre S nie są P* nazywamy **zdaniami szczegółowymi** – bo mówią one tylko o niektórych S.

Dodatkowo zdania *każde S jest P* i *niektóre S są P* określamy jako **zдания twierdzące**, natomiast *żadne S nie jest P* oraz *niektóre S nie są P* **zdaniami przeczącymi**.

Oto tabelka systematyzująca powyższe wiadomości.

Zdania kategoryczne:

schemat	Zdanie	nazwa zdania
S a P	każde S jest P	zdanie ogólno-twierdzące
S e P	żadne S nie jest P	zdanie ogólno-przeczące
S i P	niektóre S są P (istnieją S będące P)	zdanie szczegółowo-twierdzące
S o P	niektóre S nie są P (istnieją S nie będące P)	zdanie szczegółowo-przeczące

Należy zwrócić uwagę na specjalne, nieco inne od potocznego, znaczenie zdań szczegółowych, jakie przyjmują one w sylogistyce. Zwroty *niektóre* oznaczają tu bowiem *przynajmniej niektóre*, a nie *tylko niektóre*.

Zdanie *niektóre S są P* stwierdza tu tylko tyle, że istnieją obiekty S będące jednocześnie P, nie mówiąc jednakże równocześnie (wbrew temu, co się potocznie przyjmuje), iż istnieją

też obiekty S nie będące P. Zdania *niektóre S są P* **nie należy** więc rozumieć, że **tylko niektóre** S są P, ale że istnieją pewne S (być może nawet wszystkie) będące P.

Tak więc na przykład na gruncie sylogistyki za prawdziwe uznać należy zdanie S i P, gdy za S podstawimy nazwę pies, a za P – ssak. Stwierdza ono bowiem *niektóre psy są ssakami* w znaczeniu, że istnieją psy będące jednocześnie ssakami, a nie że wśród wszystkich istniejących psów tylko część z nich jest ssakami.

Podobna sytuacja zachodzi w przypadku zdania szczegółowo-przeczącego. Stwierdza ono że *niektóre S nie są P*, w znaczeniu że istnieją obiekty S nie będące jednocześnie P, nie przesądzając jednak, czy są również obiekty S będące P. W związku z tym za prawdziwe należy uznać zdanie *niektórzy ludzie nie są ptakami* jako stwierdzające, iż istnieją ludzie nie będący ptakami.

2.1.2. PRAKTYKA: ZAPISYWANIE SCHEMATÓW ZDAŃ.

Ponieważ w sylogistyce mamy do czynienia jedynie z czterema możliwymi typami zdań, pisanie schematów wydaje się niezwykle proste. Jest tak faktycznie, choć, jak się za chwilę okaże, tu również kryć się mogą pewne utrudnienia.

Przykład:

Napiszemy schemat zdania: *Każdy szpak jest ptakiem.*

Schemat tego zdania to oczywiście:

S a P,

gdzie poszczególne zmienne oznaczają nazwy: S – *szpak*, P – *ptak*.

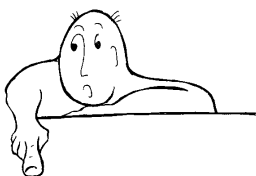
Przykład:

Napiszemy schemat zdania: *Niektórzy politycy nie są złodziejami.*

Schemat tego zdania to:

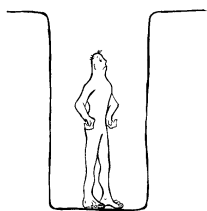
S o P

S – *polityk*, P – *złodziej*.



Uwaga na błędy!

Pisząc co oznaczają poszczególne zmienne nazwowe, podajemy nazwy w liczbie pojedynczej, a więc np. S oznacza nazwę *polityk*, a nie *politycy*, natomiast P *złodziej*, a nie *złodzieje*.



2.1.3. UTRUDNIENIA I PUŁAPKI.

Większość problemów mogących pojawić się przy pisaniu schematów zdań na gruncie sylogistyki wynika z faktu, iż w języku potocznym mało zdań ma formę dokładnie odpowiadającą któremuś ze schematów zdań kategoriycznych, a więc np. *każde [nazwa] jest [nazwa]* czy też *niektóre [nazwa] nie są [nazwa]* itd. Ze względów stylistycznych, brzmią one na ogół trochę (lub nawet całkiem) inaczej – a to, że są to w istocie zdania kategoriyczne odkrywamy dopiero po pewnym namyśle i odpowiedniej zmianie ich formy (choć oczywiście nie treści).

Czy to jest nazwa?

Często problemem może być ustalenie nazwy odpowiadającej zmiennej S lub P.

Przykład:

Napiszemy schemat zdania: *Niektórzy studenci są pilni.*

Wydaje się oczywiste, że mamy do czynienia ze zdaniem szczegółowo-twierdzącym, a więc jego schemat powinien wyglądać S i P. Problem może pojawić się jednak, gdy trzeba będzie określić, co oznacza zmienna P. Teoria mówi, że P musi odpowiadać jakaś nazwa – czy jednak wyrażenie *pilni*, (lub w liczbie pojedynczej *pilny*) jest nazwą? Otóż sam przymiotnik *pilny* nazwą jeszcze nie jest, jednakże w kontekście rozważanego zdania pełni on rolę skrótu wyrażenia *człowiek pilny* lub *osoba pilna* – i tak właśnie należy go potraktować. Tak więc ostateczne rozwiązanie zadania to:

S i P,

S – *student*, P – *człowiek pilny*.

Przykład:

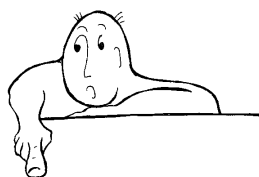
Napišemy schemat zdania: *Żaden uczony nie przeczytał wszystkich książek.*

Mamy tu oczywiście do czynienia ze zdaniem ogólnoprzeczącym, a więc jego schemat powinien wyglądać S e P. Podobnie jednak jak w poprzednim przykładzie trudność może tu sprawić określenie nazwy odpowiadającej zmiennej P – jak łatwo bowiem zauważyć, wyrażenie *przeczytał wszystkie książki* nazwą na pewno nie jest. Pierwszą narzucającą się możliwością jest uznanie za termin P wyrażenia *przeczytanie wszystkich książek* – jako nazwy pewnego procesu. W takim jednak wypadku po podstawieniu tej nazwy do schematu S e P otrzymalibyśmy wyrażenie *żaden uczony nie jest przeczytaniem wszystkich książek* – co nie jest oczywiście zdaniem, którego schemat mieliśmy napisać. Inną przychodzącą na myśl, choć również błędną, możliwością jest uznanie za P nazwy *książka* lub *każda książka*. Wtedy jednak również otrzymalibyśmy po podstawieniu nazw do schematu dość absurdalnie brzmiące wyrażenie – *żaden uczony nie jest każdą książką* lub coś podobnego. Prawidłowa odpowiedź jest taka, że zmienna P oznacza w przypadku badanego zdania nazwę – *człowiek, który przeczytał wszystkie książki* lub ewentualnie *ktoś, kto przeczytał wszystkie książki*. Po podstawieniu tego terminu do schematu S e P otrzymamy bowiem zdanie *żaden uczony nie jest człowiekiem, który przeczytał wszystkie książki* – a więc wyrażenie dokładnie odpowiadające treści zdania z przykładu, tylko nieco inaczej sformułowane.

Tak więc ostateczne rozwiązanie to:

S e P

S – *uczony*, P – *człowiek, który przeczytał wszystkie książki*.



Uwaga na błędy!

W powyższym przykładzie można łatwo popełnić pomyłkę uznając za P zdanie przeczące: *człowiek, który **nie** przeczytał wszystkich książek*. Jest to błąd, ponieważ przeczenie już zostało oddane przy pomocy stałej „e” oznaczającej *żaden **nie** jest*.

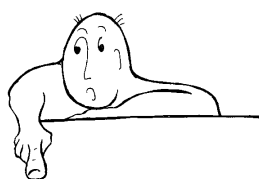
Przykład:

Napišemy schemat zdania: *Każdy, kto choć trochę poznał Józefa, wiedział, że nie można mu ufać.*

Oczywiste jest, iż mamy do czynienia ze zdaniem ogólno-twierdzącym, a więc jego schemat będzie wyglądał: S a P. Co jednak będą oznaczały zmienne S i P? Doświadczenie z poprzednich przykładów podpowiada, że P oznacza termin *ktoś, kto wiedział, że nie można ufać Józefowi*. Problem może tu jednak również sprawić określenie znaczenia zmiennej S. Na pewno nie jest to *Józef* – co łatwo sprawdzić, próbując podstawić tę nazwę do schematu *każde S jest P*. S w powyższym przykładzie oznacza nazwę – *ktoś, kto choć trochę poznał Józefa*. Tak więc mamy ostateczne rozwiązanie:

S a P

S – *ktoś, choć trochę poznał Józefa*, P – *ktoś, kto wiedział, że nie można ufać Józefowi*.



Uwaga na błędy!

W powyższym przykładzie błędem byłoby napisanie, że S oznacza **każdy**, *kto choć trochę poznał Józefa*. Słowo **każdy** zostało już bowiem oddane w symbolu „a”.

Przykład:

Napiszemy schemat zdania: *Niektórzy nie lubią zwierząt*.

Jest to oczywiście zdanie szczegółowo-przeczące, a więc o schemacie S o P. Zmiennej P odpowiada nazwa – *ktoś kto lubi zwierzęta* (pamiętamy, że *nie* zostało już oddane przy pomocy stałej „o”). Co jest jednak odpowiednikiem S? W badanym zdaniu nie widać żadnego wyrażenia, które można by za S podstawić – poza zwrotem o lubieniu zwierząt oraz wyrażeniem *niektórzy*, które zostaje oddane przez stałą „o” w zdaniu niczego więcej już nie ma. Jednakże treść zdania jasno wskazuje, że owi *niektórzy*, o których ono mówi, choć nie stwierdza tego wprost, to ludzie. Tak więc nazwa S to po prostu *człowiek*. Ostateczne rozwiązanie:

S o P

S – *człowiek*, P – *ktoś, kto lubi zwierzęta*.

Czy to jest stała logiczna?

Nie tylko odpowiadające zmiennym S oraz P nazwy mogą przybierać różnorodne formy; również stałe logiczne występują czasem pod zmienioną postacią.

Przykład:

Napiszemy schemat zdania: *Ktokolwiek twierdzi, że widział UFO, myli się lub kłamie.*

Wprawdzie w zdaniu tym nie występuje wprost żadne z wyrażeń odpowiadających stałym a, e, i, o, jednakże oczywiste jest, że *ktokolwiek* to odpowiednik zwrotu *wszyscy*, czy też *każdy*, a więc mamy do czynienia ze zdaniem ogólno-twierdzącym:

S a P

S – *ktoś, kto twierdzi, że widział UFO*, P – *ktoś, kto myli się lub kłamie.*

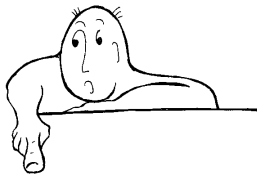
Przykład:

Napiszemy schemat zdania: *Nikt nie lubi gdy inni go krytykują.*

W tym wypadku *nikt*, to odpowiednik zwrotu *żaden*:

S e P

S – *człowiek*, P – *ktoś, kto lubi, gdy inni go krytykują.*



Uwaga na błędy!

Niektórzy mogą początkowo błędnie sądzić, że zmiennej S odpowiada nazwa *nikt* lub *ktoś, kto czegoś nie lubi*. Że nie są to dobre odpowiedzi łatwo się przekonać wstawiając te terminy za S w schemacie S e P.

Czy jest tam jakaś stała logiczna?

Czasem wyrażenie odpowiadające którejś ze stałych logicznych może być w ogóle nieobecne (nie ma go nawet w innej formie), jednakże można się go domyślić z treści zdania.

Przykład:

Napiszemy schemat zdania: *Kto rano wstaje, temu Pan Bóg daje.*

Wprawdzie w powyższym zdaniu nie ma wyrażenia *każdy*, *żaden*, ani *niektóry* (nawet w innej formie), jednakże zapewne każdy znający to powiedzenie uzna, że mamy do czynienia ze zdaniem ogólnym, odnoszącym się domyślnie do *wszystkich* ludzi. Tak więc schemat zdania wygląda następująco:

S a P

S – ktoś, kto rano wstaje, P – ktoś, komu Pan Bóg daje.

Co zrobić z negacją?

Zdarza się czasem, że mamy do czynienia z wyrażeniem, które stanowi negację któregoś ze zdań kategorycznych. Szczególnie często negacja występuje przy zdaniach ogólnotwierdzących.

Przykład:

Napiszemy schemat zdania: *Nie każdy polityk wierzy w to, co mówi.*

Na pierwszy rzut oka widać, że powyższe wyrażenie stanowi negację zdania S a P. Teoretycznie więc jego schemat można by zapisać $\sim (S a P)$ – i faktycznie czasami się tak robi. Jednakże w tradycyjnie ujętej sylogistyce negacje nie występują. Nie są one zresztą konieczne, ponieważ negację każdego ze zdań kategorycznych można oddać przy pomocy równoważnego mu innego zdania, już bez negacji. Po chwili zastanowienia każdy przyzna, że zdanie *nieprawda, że każde S jest P* mówi dokładnie to samo co *niektóre S nie są P*. Przy użyciu symboliki logicznej można by to zapisać $\sim (S a P) \equiv S o P$.

Wracając do naszego przykładu możemy zatem powiedzieć, że zdanie *nie każdy polityk wierzy w to, co mówi* równoważne jest zdaniu *niektórzy politycy nie wierzą w to, co mówią*. Tak więc jego schemat zapisać można:

S o P

S – polityk, P – osoba, która wierzy w to, co mówi.

DO ZAPAMIĘTANIA:



Oto jak można oddać negacje wszystkich zdań kategorycznych:

$\sim (S a P) \equiv S o P$

$\sim (S e P) \equiv S i P$

$\sim (S i P) \equiv S e P$

$\sim (S o P) \equiv S a P$

Przykład:

Napiszemy schemat zdania: *Nie jest prawdą, że niektórzy uczeni są nieomylni.*

Zdanie to stanowi negację zdania szczegółowo-twierdzącego (czyli $\sim (S i P)$), można więc je oddać przy pomocy schematu:

S e P

S – uczony, P – osoba nieomylna.

Gdzie S, a gdzie P?

Czasem trudność przy pisaniu schematu sprawić może określenie, która nazwa odpowiada zmiennej S, a która P.

Przykład:

Napiszemy schemat zdania: *Zły to ptak, co własne gniazdo kala.*

Podobnie jak w przypadku zdania *kto rano wstaje, temu Pan Bóg daje* można się domyślać, że powiedzenie to ma charakter zdania ogólnego o schemacie S a P. Czy jednak możemy uznać, że S odpowiada nazwie *zły ptak*, a P – *ptak kalający własne gniazdo*, jak by się to mogło wydawać na pierwszy rzut oka? W takim wypadku otrzymalibyśmy stwierdzenie, że *każdy zły ptak kala własne gniazdo*. Tymczasem w znanym powiedzeniu chodzi raczej o coś przeciwnego – że to *każdy ptak kalający własne gniazdo, jest zły*. Tak więc faktycznie mamy do czynienia ze zdaniem o schemacie S a P, jednakże nazwa odpowiadająca zmiennej S została w nim umieszczona na końcu, a odpowiadająca P – na początku. Tak więc ostateczne rozwiązanie to:

S a P

S – *ptak kalający własne gniazdo*, P – *zły ptak*.

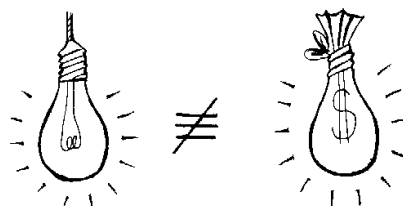
Przykład:

Napiszemy schemat zdania: *Nie wszystko złoto, co się świeci.*

Oczywiste wydaje się, że powyższe powiedzenie stanowi negację zdania o schemacie S a P, a więc ma ono formę S o P. Co jednak jest tu terminem S, a co P? Gdybyśmy określili S jak

złoto, a P jako *coś, co się świeci* i podstawili je do schematu S o P (lub $\sim (S \text{ a } P)$), otrzymalibyśmy zdanie stwierdzające, że niektóre rodzaje złota nie świecą się, lub też że nie jest prawdą, iż każde złoto się świeci. Jak widać nie jest to raczej to, o co chodzi w rozważanym przysłowiu.

Aby sprawę wyjaśnić zostawmy na chwilę negację i przyjrzyjmy się ogólnie zdaniom o formie *wszystko A, co B* – nie mówią one bynajmniej, że *każde A jest B*, ale odwrotnie, że to *każde B jest A*. Przykładowo *wszystko okazało się słuszne, co w życiu uczyniłem*, stwierdza, że każda rzecz, jaką w życiu zrobiłem, okazała się słuszną, a nie, że wszystkie rzeczy, jakie są słuszne, uczyniłem w swoim życiu.



NIE WSZYSTKO ZŁOTO, CO SIĘ ŚWIECI

Tak więc zdanie *nie wszystko złoto, co się świeci* stwierdza coś w rodzaju *nie jest prawdą, że każda rzecz świecąca się jest złotem*, czyli *niektóre rzeczy świecące się, nie są złotem*. Ostateczna odpowiedź to:

S o P

S – coś, co się świeci, P – złoto.

Co znaczy „tylko”?

Jako zdania kategoryczne można potraktować również wyrażenia ze zwrotem *tylko... są...*, choć na pierwszy rzut oka zwrot ten nie odpowiada żadnej z poznanych stałych logicznych.

Przykład:

Napiszemy schemat zdania: *Tylko kobiety są matkami*.

Intuicja podpowiada, że w powyższym przypadku mamy do czynienia ze zdaniem twierdzącym (nie ma w nim przeczenia) oraz ogólnym (stwierdza coś o wszystkich obiektach pewnego typu, a nie tylko o niektórych). Tak więc nasuwa się schemat S a P. Jest to faktycznie właściwy schemat – ważne jest jednak, abyśmy prawidłowo określili nazwy przyporządkowane zmiennym S oraz P. Gdyby za S podstawić nazwę *kobieta*, a za P – *matka* otrzymalibyśmy zdanie *każda kobieta jest matką*. Nie jest to na pewno zdanie równoważne stwierdzeniu *tylko kobiety są matkami* – widać to już na pierwszy rzut oka chociażby dlatego, że pierwsze z nich jest fałszywe, a drugie prawdziwe. Wyrażenie równoważne zdaniu z naszego przykładu, to *każda matka jest kobietą*.

Aby to dobrze zrozumieć, należy sobie wyobrazić, co to oznacza, że *tylko kobiety są matkami*. Znaczy to po prostu, iż wśród matek mamy tylko i wyłącznie kobiety, a więc ni mniej ni więcej, tylko właśnie *każda matka jest kobietą*. Tak więc ostateczne rozwiązanie to:

S a P

S – matka, P – kobieta.



DO ZAPAMIĘTANIA:

Zdania typu *tylko A są B* zawsze możemy przedstawić przy pomocy schematu S a P, gdzie $S = B$, $P = A$.

Przykład:

Zapiszemy schemat zdania: *Nie tylko artyści są zarozumiali*.

Schemat tego zdania to:

S o P

S – osoba zarozumiała, P – artysta.

Do powyższego rozwiązania dojść można na dwa sposoby. Jeden polega na wyobrażeniu sobie, co oznacza zdanie mówiące że *nie tylko artyści są zarozumiali*. Po chwili zastanowienia każdy powinien zobaczyć, że opisuje ono fakt, iż wśród osób zarozumiałych są też inni ludzie oprócz artystów, a więc inaczej mówiąc – *niektóre osoby zarozumiałe nie są artystami*.

Drugi sposób na otrzymanie prawidłowego schematu rozważanego zdania polega na zbudowaniu najpierw schematu zdania *tylko artyści są zarozumiali*, a następnie zanegowaniu go zgodnie z zasadami opisanymi wyżej w punkcie *co zrobić z negacją?*. Schemat zdania *tylko artyści są zarozumiali* to S a P, gdzie S – osoba zarozumiała, a P – artysta. Ponieważ ostatecznie musimy napisać schemat negacji tego stwierdzenia, znajdujemy zdanie równoważne negacji S a P, którym jest S o P.



2.1.4. CZĘSTO ZADAWANIE PYTANIA.

Czy na gruncie sylogistyki da się napisać schemat każdego zdania?

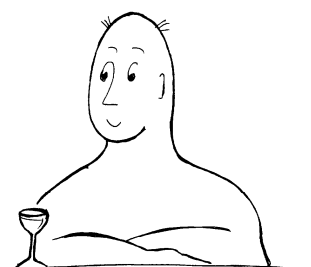
Nie. Na gruncie sylogistyki można pisać tylko schematy zdań kategoriowych, a więc zawierających zwroty: *każdy jest*, *żaden nie jest*, *niektóre są* i *niektóre nie są* (lub zwroty im równoważne). Gdy zdanie nie zawiera takiego zwrotu, napisanie jego schematu jest niemożliwe.

Czy nazwy koniecznie musimy oznaczać zmiennymi S oraz P?

Nie jest to konieczne, choć takie rozwiązanie jest bardzo mocno ugruntowane w tradycji. Dlatego też oznaczenie nazw innymi symbolami choć nie jest błędem, sprawia wrażenie mało eleganckiego. Jeżeli zachodzi potrzeba wykorzystania kolejnego symbolu na oznaczenie nowej nazwy (patrz niżej), używana jest zwykle litera M.

2.2. SPRAWDZANIE POPRAWNOŚCI SYLOGIZMÓW METODĄ DIAGRAMÓW VENNA.

2.2.1. ŁYK TEORII.



ŁYK TEORII

Co to jest sylogizm?

Sylogizm, to pewien ściśle określony rodzaj wnioskowania. Sylogizm zawsze musi składać się z trzech zdań kategoriycznych: dwóch przesłanek i wniosku. Dodatkowym warunkiem, jaki musi spełniać każdy sylogizm jest ilość nazw obecnych w owych trzech zdaniach – zawsze są to trzy nazwy. Tak więc oprócz zmiennych S oraz P w schematach zdań składających się na

sylogizm wykorzystać trzeba jeszcze trzeci symbol – zwykle jest to M.

Przykładowy sylogizm może wyglądać następująco: *Każdy człowiek szczęśliwy jest tolerancyjny. Niektórzy wychowawcy nie są tolerancyjni. Zatem niektórzy wychowawcy nie są szczęśliwi.*

Schematy powyższych zdań, zapisane w znanej z rachunku zdań formie reguły, przyjmują następującą postać:

P a M

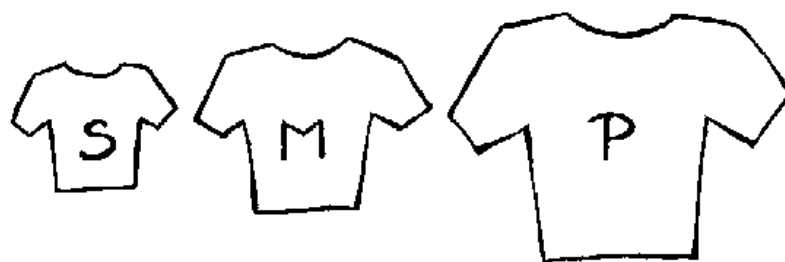
S o M

—

S o P

W sylogizmie ważne jest, które nazwy oznaczmy jaką zmienną. Przyjęte jest, aby symbole S oraz P zarezerwować dla nazw obecnych w konkluzji wnioskowania. Natomiast trzecia nazwa – ta, której nie ma w konkluzji, a która jest za to zawsze w obu przesłankach – oznaczana jest symbolem M. Tradycyjnie nazwę oznaczoną przez S nazywamy **terminem mniejszym** sylogizmu, nazwę oznaczoną P – **terminem większym**, natomiast nazwę M –

terminem średnim. Znajomość powyższej terminologii nie jest może najważniejsza dla rozwiązywania zadań z zakresów sylogizmów, ponieważ jednak jest to nazewnictwo stosowane w wielu podręcznikach logiki, dobrze jest je znać. Zapamiętanie określeń poszczególnych terminów nie powinno zresztą sprawić trudności nikomu, kto skojarzy je z popularnymi i ogólnie znanymi oznaczeniami odzieży, zgodnie z którymi S oznacza rozmiar mały, natomiast M – średni.



TERMINY : MNIEJSZY, ŚREDNI I WIĘKSZY

Kończąc rozważania na temat tradycyjnej terminologii dodajmy, że przesłanka, która obok nazwy oznaczanej M zawiera również termin P, nazywana jest **przesłanką większą** sylogizmu, natomiast ta, w której obok M występuje S, nazywana jest **przesłanką mniejszą**.

W przykładzie z początku tego paragrafu nazwa *wychowawca* stanowi zatem termin mniejszy, nazwa *człowiek szczęśliwy* termin większy, natomiast *człowiek tolerancyjny* termin średni. Przesłanka *każdy człowiek szczęśliwy jest tolerancyjny* jest przesłanką większą, natomiast *niektórzy wychowawcy nie są tolerancyjni* przesłanką mniejszą.

Sprawdzanie poprawności sylogizmu.

Sylogizm to rodzaj wnioskowania. Sprawdzenie poprawności sylogizmu, to zatem nic innego jak sprawdzenie poprawności wnioskowania. Jak pamiętamy z rachunku zdań wnioskowanie jest poprawne, gdy wniosek wynika logicznie z przesłanek, a to z kolei ma miejsce, gdy niezawodna jest reguła (czyli schemat całego wnioskowania), na której wnioskowanie jest oparte. Reguła jest niezawodna, gdy na mocy znaczenia stałych logicznych nie jest możliwa sytuacja, aby przesłanki były prawdziwe, natomiast wniosek fałszywy; lub, ujmując to samo innymi słowy, w przypadku niezawodnej reguły, jeśli przesłanki są prawdziwe, to prawdziwy musi być również i wniosek.

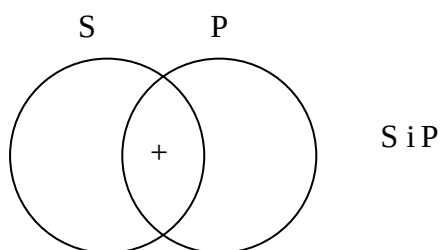
Na gruncie rachunku zdań niezawodność reguł badaliśmy przy pomocy tabelki zero-jedynkowych oddających znaczenie spójników logicznych. Ponieważ w teorii sylogizmów mamy stałe logiczne inne niż spójniki zdaniowe, konieczna jest tu odmienna metoda.

Przedstawimy obecnie najpopularniejszy sposób sprawdzania poprawności sylogizmów: metodę diagramów Venna.

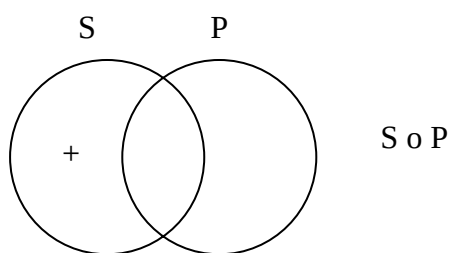
Diagramy Venna.

W diagramach Venna (nazywanych tak od nazwiska ich pomysłodawcy Johna Venna) koła symbolizują zbiory obiektów określanych przez poszczególne nazwy, a więc zakresy tych nazw. Znaki „+” oraz „-” w częściach tych kół informują, że w danym obszarze na pewno coś się znajduje lub też, że na pewno niczego tam nie ma.

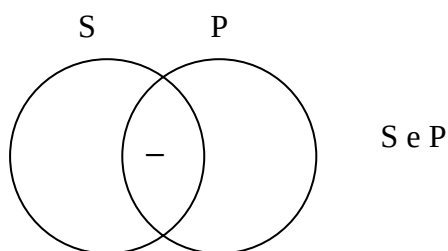
Oto, jak na diagramach Venna przedstawić można poszczególne zdania kategoryczne:



Zdanie mówiące, że *niektóre S są P* stwierdza, iż muszą istnieć jakieś obiekty w części wspólnej S oraz P. Symbolizuje to znak „+” w tej części rysunku. Na temat pozostałych obszarów diagramu zdanie S i P niczego nie mówi, dlatego nic do nich nie wpisujemy.



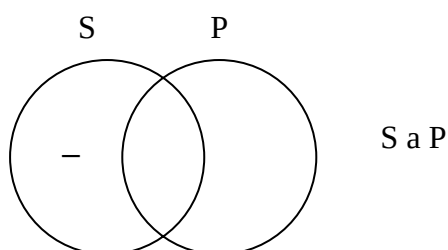
Zdanie *niektóre S nie są P* informuje, iż na pewno istnieją obiekty należące do zbioru S, a jednocześnie nie należące do P. Stąd znak „+” w części S znajdującej się poza zbiorem P. Odnosnie pozostałych obszarów diagramu zdanie S o P nie niesie żadnych informacji.



Zdanie *żadne S nie są P* stwierdza, że nie istnieją żadne obiekty należące jednocześnie do zbiorów S i P. Fakt ten uwidocznił jest przez znak „-” w części wspólnej tych zbiorów. Zauważmy, że zdanie typu $S e P$ nie informuje o istnieniu jakichkolwiek obiektów będących desygnatami nazw S lub P (może ono mówić na przykład *żaden krasnoludek nie jest jednorożcem*) – dlatego też niczego nie wpisujemy w pozostałe obszary diagramu.

Uwaga na marginesie.

W praktyce, przy rozwiązywaniu zadań związanych z sylogizmami, będziemy czasem korzystali z założenia, że obiekty będące desygnatami danej nazwy na pewno istnieją. Obecnie jednak, aby zbytnio nie zaciemniać obrazu, będziemy wpisywali do diagramu tylko to, co dane zdanie wprost stwierdza, pomijając informacje, jakie mogą z niego dodatkowo wynikać przy pewnych założeniach.



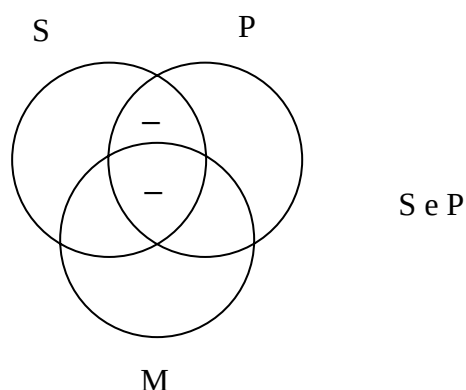
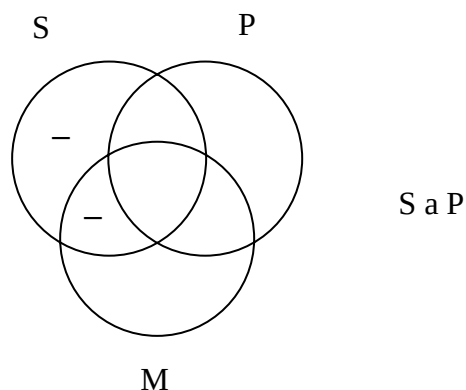
Zdanie *każde S jest P* informuje, że cokolwiek możemy określić nazwą S, podpada również pod nazwę P. Nie ma w związku z tym żadnych obiektów S nie będących jednocześnie P – stąd minus w lewej części diagramu. Zdanie to nie niesie jednak żadnej „pozytywnej” informacji, że jakiegokolwiek S faktycznie istnieje – stwierdza jedynie, że **jeżeli** coś jest S (o ile w ogóle istnieje) to jest również P. Dlatego też nie stawiamy znaku „+” w części środkowej.

Diagramy dla trzech nazw.

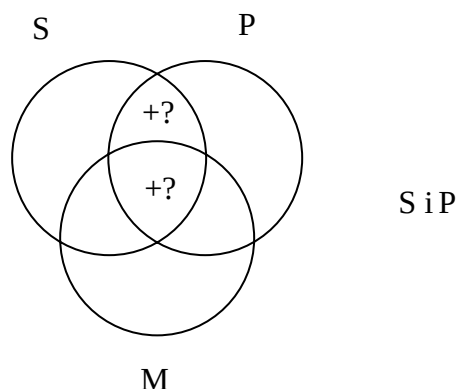
Powyżej przedstawione zostały diagramy Venna dla dwóch terminów. Jednakże w każdym sylogizmie występują trzy nazwy. Dlatego też do sprawdzania poprawności

sylogizmów potrzebna jest umiejętność zaznaczania poszczególnych zdań kategorycznych na diagramach złożonych z trzech kół.

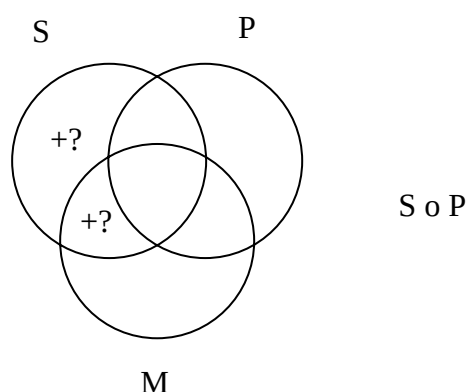
Tutaj prostsza jest sprawa dla zdań ogólnych – ich rysunki stanowią zwykłe rozszerzenie diagramów sporządzanych dla dwóch nazw. Gdy mamy do czynienia ze zdaniem $S \text{ a } P$ to pusty musi być cały obszar zbioru S leżący poza P , natomiast w przypadku zdania $S \text{ e } P$ pusty musi pozostać obszar wspólny tych zbiorów. Ponieważ teraz obszary te składają się z dwóch części, musimy postawić znaki „–” w obu tych kawałkach:



Nieco inaczej przedstawia się sytuacja w przypadku zdań szczegółowych. Rozpatrzmy najpierw zdanie $S \text{ i } P$. Stwierdza ono, że istnieją pewne obiekty w części wspólnej zbiorów S oraz P . Na rysunku obrazującym zależności między trzema nazwami obszar ten składa się z dwóch części. Zdanie $S \text{ i } P$ nie informuje jednak, w której z tych części coś się znajduje – może w jednej, może w drugiej, a może w obydwu. Zależy to od terminu M , o którym na razie nic nie wiemy. W związku z tym, wpisując symbole „+” w odpowiednich częściach, należy opatrzyć je znakami zapytania. Pytajniki te informują, że w danym obszarze na pewno jakieś elementy się znajdują, ale nie wiadomo w której jego części.



Z podobną sytuacją spotykamy się w przypadku zdania S o P. Informuje nas ono, że na pewno istnieją jakieś elementy w części zbioru S znajdującej się poza zbiorem P, ale nie określa, w którym fragmencie tego obszaru – w jednym, drugim, czy może obydwu.



Znajomość przedstawionych wyżej sposobów zaznaczania zdań kategoriycznych na diagramach konieczna jest do sprawdzania poprawności sylogizmów w takim samym stopniu, jak znajomość tabelk zero-jedynkowych była nieodzowna do badania prawidłowości wnioskowań na gruncie KRZ.



DO ZAPAMIĘTANIA:

Z powyższych rysunków warto zapamiętać następujące fakty.

- Zdania ogólne (S a P oraz S e P) dają nam zawsze minusy na diagramach, natomiast zdania szczegółowe (S i P oraz S o P) – plusy.
- Minusy są zawsze „pewne” (bez znaków zapytania) – wynika to z tego, że gdy jakiś obszar ma być pusty, to pusta musi być każdy jego część.
- Plusy są „niepewne” – gdy wiemy, że w danym obszarze, coś się znajduje, to nie oznacza to jeszcze, że wiemy w której jego części.

„Pewność” minusów i „niepewność” plusów na diagramach zilustrować można następującą analogią: gdy wiemy, że w jakimś mieszkaniu nikogo nie ma, to wiemy na pewno, że nikogo nie ma ani w kuchni, ani w pokoju („pewne” minusy w każdej części); gdy natomiast wiemy, że danym mieszkaniu ktoś jest, to nie znaczy to jeszcze, że wiemy, w którym jego pomieszczeniu.

Uwaga na marginesie.

W praktyce, gdy będziemy rozwiązywać zadania związane z sylogizmami, informacje zawarte w jednym zdaniu będą nam często jednoznacznie wskazywać, w którym miejscu należy wpisać znak „+” wynikający z drugiego zdania. W takich wypadkach plus ten będzie „pewny”.

2.2.2. PRAKTYKA: ZASTOSOWANIE DIAGRAMÓW VENNA.

Obecnie możemy przystąpić do sprawdzania poprawności sylogizmów. Oprócz umiejętności zaznaczania na diagramie poszczególnych typów zdań, przy badaniu sylogizmów musimy mieć w pamięci pojęcie wynikania logicznego. Sylogizm (jak każde wnioskowanie) jest bowiem wtedy poprawny, gdy jego wniosek wynika logicznie z przesłanek.

Badanie poprawności sylogizmów przy pomocy diagramów Venna składa się z dwóch kroków. W pierwszym z nich wpisujemy do diagramu wszystkie informacje, jakie niosą ze sobą przesłanki. W drugim kroku sprawdzamy, czy tak wypełniony diagram gwarantuje nam prawdziwość wniosku. Zdania będącego wnioskiem sylogizmu nie wpisujemy już jednak do diagramu. Musimy jedynie wyobrazić sobie, co by w diagramie musiało się znajdować, aby był on prawdziwy, a następnie sprawdzić, czy nasz diagram spełnia te warunki.

Jeśli okaże się, że prawdziwość konkluzji jest na wykonanym rysunku zagwarantowana, będzie to znak, że nie jest możliwa sytuacja, aby przesłanki były prawdziwe, a wniosek fałszywy, a więc że wniosek wynika z przesłanek, czyli sylogizm jest poprawny. Jeśli natomiast wypełnienie diagramu według przesłanek nie da nam pewności co do prawdziwości wniosku, będzie to oznaczało, że wniosek nie wynika z przesłanek (bo może być on fałszywy, pomimo prawdziwości przesłanek), a więc sylogizm nie jest logicznie poprawny. W takim przypadku zawsze możliwe jest stworzenie tak zwanego kontrprzykładu – diagramu ilustrującego sytuację, w której przesłanki są prawdziwe, a wniosek fałszywy.



DO ZAPAMIĘTANIA:

W skrócie procedura sprawdzania poprawności sylogizmów będzie wyglądała następująco:

- Piszemy schematy zdań wchodzących w skład sylogizmu.
- Rysujemy diagram składający się z trzech kół symbolizujących trzy nazwy występujące w sylogizmie.
- Wpisujemy do diagramu plusy i minusy, o których informują przesłanki sylogizmu.
- Patrzymy na rysunek i sprawdzamy, czy wypełniony na podstawie przesłanek diagram gwarantuje nam, że prawdziwe będzie zdanie stanowiące wniosek sylogizmu.
- Jeżeli rysunek gwarantuje prawdziwość konkluzji, oznacza to, że sylogizm jest poprawny; jeśli nie mamy pewności co do prawdziwości wniosku, oznacza to, że sylogizm jest niepoprawny.

Przykład:

Sprawdzimy poprawność sylogizmu przedstawionego we wstępie do tego rozdziału:
Każdy jamnik jest psem. Każdy pies jest ssakiem. Zatem każdy jamnik jest ssakiem.

Napisanie schematów przesłanek i wniosku nie powinno sprawić nikomu najmniejszej trudności. Pamiętać musimy jedynie, że jeśli chcemy być w zgodzie z tradycją, to wniosek naszego sylogizmu powinien mieć postać S P. Tak więc zacząć możemy od określenia, który termin należy oznaczyć jaką zmienną:

S – *jamnik*, P – *ssak*, M – *pies*.

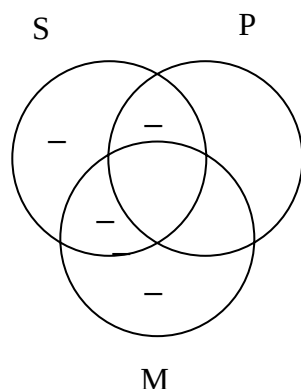
Reguła, na której opiera się badany sylogizm, jest następująca:

S a M

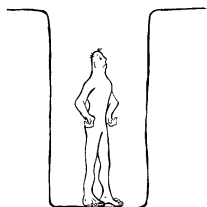
M a P

S a P

Teraz możemy narysować diagram i wpisać do niego to, co mówią przesłanki. Pierwsza przesłanka stwierdza, że pusty musi być obszar zbioru S leżący poza M, natomiast druga, że pusty musi być obszar zbioru M leżący poza P. Po wpisaniu w odpowiednie miejsca minusów otrzymujemy następujący diagram:



Do diagramu tego nie wpisujemy tego, co mówi wniosek sylogizmu, a jedynie patrzymy, czy wykonany na podstawie przesłanek rysunek, gwarantuje nam jego prawdziwość. Konkluzja naszego sylogizmu ma postać S a P, a więc aby była ona prawdziwa, pusty musi być obszar zbioru S leżący poza zbiorem P. Na wypełnionym diagramie w obu częściach tego obszaru znajdują się minusy, a więc mamy stuprocentową gwarancję, że jest on faktycznie pusty. Jest to znak, że wniosek wynika z przesłanek (musi być prawdziwy, jeśli tylko prawdziwe są przesłanki), a zatem **badany sylogizm jest poprawny**.



2.2.3. UTRUDNIENIA I PUŁAPKI.

Plus ze znakiem zapytania nie daje pewności!

Czasami może zdarzyć się sytuacja, że wniosek sylogizmu stwierdza, iż w danym obszarze coś się musi znajdować, natomiast na diagramie w miejscu tym będzie znak „+?”. Poniższy przykład ilustruje tę sytuację:

Przykład:

Zbadamy poprawność sylogizmu: *Każdy milioner jest bogaty. Niektórzy bogaci ludzie nie są szczęśliwi. Zatem niektórzy milionerzy nie są szczęśliwi.*

Schematy, na których opiera się powyższy sylogizm to:

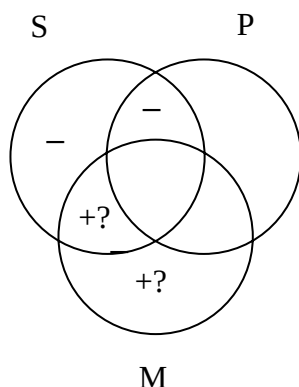
S a M

M o P

S o P

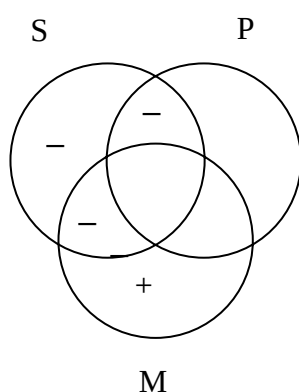
S – milioner, P – człowiek szczęśliwy, M – człowiek bogaty.

Po wpisaniu do diagramu informacji, jakie niosą ze sobą przesłanki, otrzymujemy następującą sytuację:



Teraz pozostaje nam sprawdzenie, czy tak wypełniony diagram gwarantuje nam prawdziwość konkluzji. Wniosek sylogizmu ma postać S o P, a więc stwierdza, że coś powinno znajdować się w obszarze zbioru S leżącym poza zbiorem P. Jak widać na rysunku w jednej części tego obszaru mamy znak „-” (na pewno więc nic tam nie ma), natomiast w drugiej „+?”. Czy taki plus ze znakiem zapytania daje nam gwarancję, że coś się w badanym obszarze znajduje? Oczywiście, że nie. Symbol ten wskazuje, że jakieś elementy mogą tam być, ale nie jest to pewne. Natomiast do tego, aby sylogizm uznać za poprawny, potrzebujemy stuprocentowej gwarancji prawdziwości konkluzji. Ponieważ w badanym przykładzie pewności takiej nie mamy, świadczy to o tym, że **sylogizm jest niepoprawny**.

O niepoprawności powyższego sylogizmu przekonuje diagram wypełniony w następujący sposób.



Rysunek ten stanowi graficzny kontrprzykład do badanej reguły. Widać na nim, że bez popadania w jakąkolwiek sprzeczność można wpisać do diagramu plusy i minusy w taki sposób, aby przesłanki były prawdziwe natomiast wniosek fałszywy. W przypadku reguły niezawodnej takie wypełnienie diagramu nie było by możliwe.

Kontrprzykład ukazujący zawodność reguły można też zbudować podstawiając do niej za zmienne S, P oraz M nazwy w taki sposób, że nie pozostawi to żadnych wątpliwości, iż przesłanki są prawdziwe, a wniosek fałszywy. W powyższym przykładzie może być to np.: S – *jamnik*, P – *pies*, M – *ssak*. Przesłanki powiedzą wtedy, że *każdy jamnik jest ssakiem* oraz *niektóre ssaki nie są psami* (prawda), natomiast wniosek: *niektóre jamniki nie są psami* (fałsz).

Uwaga na marginesie:

Do każdej zawodnej reguły na gruncie sylogistyki można zbudować kontrprzykład korzystając jedynie z nazw *kot*, *pies*, *jamnik*, *ssak*. W takim przypadku trzeba jednak wiedzieć, iż czasem zajdzie potrzeba oznaczenia dwóch zmiennych tą samą nazwą (np. S – *kot*, P – *kot*).

Można oczywiście też budować kontrprzykłady z innymi nazwami.

Kiedy znak „+” może być pewny?

Zdania szczegółowe każą nam wpisywać do pewnego obszaru diagramu znaki „+”, nie precyzując jednak dokładnie, w którą jego część. W praktyce często sprawa sama się wyjaśnia i miejsce wpisania symbolu „+” staje się oczywiste i jednoznaczne.

Przykład:

Zbadamy poprawność sylogizmu: *Żaden mędrzec nie jest fanatykiem jednej idei. Niektórzy uczeni są fanatykami jednej idei. Zatem niektórzy uczeni nie są mędrcami.*

Reguła na której oparty jest powyższy sylogizm jest następująca:

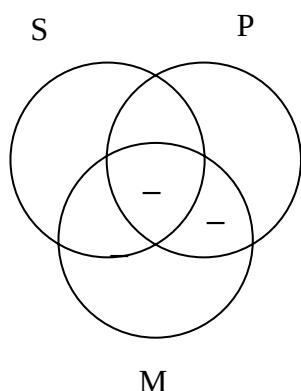
P e M

S i M

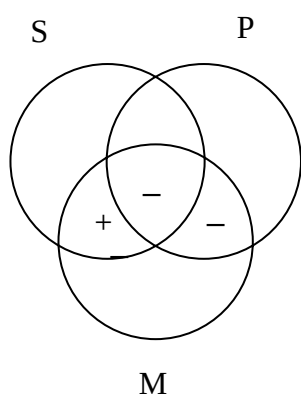
S o P

S – *uczony*, P – *mędrzec*, M – *fanatyk jednej idei*.

Pierwsza przesłanka stwierdza, że pusty musi obszar wspólny zbiorów P oraz M:



Zgodnie z drugą przesłanką coś musi znajdować się we wspólnej części zbiorów S oraz M. Teoretycznie obszar ten składa się z dwóch fragmentów. Ponieważ jednak w jednym z nich mamy już wpisany znak „-” na wpisanie „+” pozostaje nam tylko jedno miejsce. W takim wypadku „+” wpisujemy oczywiście bez znaku zapytania – mamy bowiem pewność, że musi być on w tym właśnie miejscu.



Obecnie musimy sprawdzić, czy taki rysunek gwarantuje nam prawdziwość wniosku sylogizmu, a więc zdania S o P. Aby zdanie to było prawdziwe, coś powinno się znajdować w części zbioru S leżącej poza P. Na diagramie w obszarze tym (w jego dolnej części) znajduje się znak „+”, a więc mamy pewność, że nie jest on pusty. Badany **sylogizm jest zatem poprawny**.

Gdy jedna przesłanka mówi „+”, a druga „-”.

Często zdarza się sytuacja, że zgodnie z jedną przesłanką w jakieś miejsce należy wpisać znak „+”, a zgodnie z drugą „-”. Poniższy przykład pokazuje, jak należy postąpić w takim przypadku.

Przykład:

Sprawdzimy poprawność następującego sylogizmu: *Niektórzy politycy są nacjonalistami. Każdy nacjonalista jest ograniczony. Zatem niektórzy politycy są ograniczeni.*

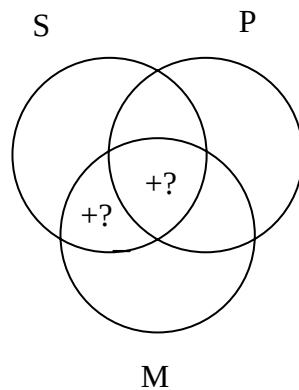
Reguła na której opiera się badany sylogizm wygląda następująco:

S i M

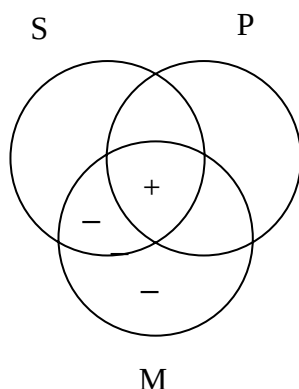
M a P

S i P

Pierwsza przesłanka stwierdza, że coś musi się znajdować we wspólnym obszarze zbiorów S oraz M, chociaż nie określa w której części tego obszaru (w jednej, drugiej, czy obydwu). Mamy więc:



Druga przesłanka mówi, że pusty musi być obszar zbioru M leżąca poza P. Jednakże w jednej części tego obszaru mamy już wpisany znak „+”. W takiej sytuacji należy zauważyć, że symbol „+” opatrzonej jest znakiem zapytania, co oznacza, że wcale nie jest konieczne, aby tam był. Ponieważ „-” wynikający z drugiej przesłanki jest „pewny”, jemu należy przyznać pierwszeństwo i wpisać go w sporny obszar. Jednocześnie modyfikacji ulec musi drugi z „+” wpisany na mocy pierwszej przesłanki. Ponieważ „skasowaniu” uległ pierwszy z nich, a przesłanka S i M stwierdza, że o obszarze wspólnym zbiorów S oraz M coś musi się znajdować, to drugi z plusów staje się „pewny” i należy zlikwidować stojący przy nim znak zapytania. Po prostu informacje z drugiej przesłanki pokazały nam, który z „niepewnych” plusów, o których informowała pierwsza przesłanka jest tym „właściwym”. Po wpisaniu informacji z obu przesłanek, diagram wygląda więc następująco:



Pozostaje nam teraz sprawdzić, czy taki rysunek gwarantuje prawdziwość konkluzji sylogizmu, czyli zdania S i P. Widać, że we wspólnym obszarze zbiorów S oraz P faktycznie coś się na pewno znajduje, a więc konkluzja ta jest prawdziwa. W związku z tym badany sylogizm jest poprawny.



WARTO ZAPAMIĘTAĆ:

Aby uniknąć kłopotliwego wymazywania symboli w diagramie i zastępowania ich innymi, najlepiej jest po prostu zaczynać wypełnianie diagramu od tej przesłanki, która daje nam „pewne” informacje (a więc zdania typu „a” bądź „e”, niezależnie, czy jest ono pierwsze, czy drugie w sylogizmie. Gdybyśmy tak postąpili w powyższym przykładzie, rozpoczynając od przesłanki M a P, przy wpisywaniu przesłanki S i M mielibyśmy już tylko jedną możliwość wpisania znaku „+”

Puste miejsce nie oznacza, że niczego w nim nie ma!

Przy sprawdzaniu, czy wypełniony według przesłanek diagram gwarantuje prawdziwość konkluzji, mogą powstać wątpliwości co do interpretacji miejsc, w których nie ma żadnego znaku.

Przykład:

Zbadamy poprawność następującego sylogizmu: *Niektórzy wykładowcy są dobrymi fachowcami. Każdy dobry fachowiec dużo zarabia. Zatem każdy wykładowca dużo zarabia.*

Reguła, na której oparty jest badany sylogizm, przedstawia się następująco:

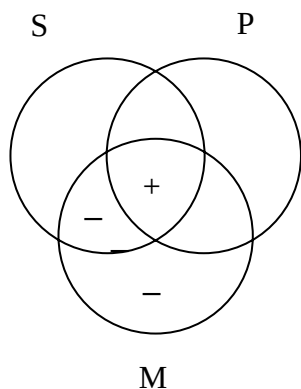
S i M

M a P

S a P

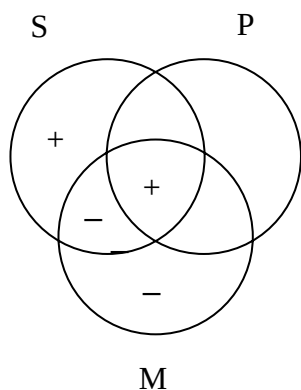
S – wykładowca, P – ktoś, kto dużo zarabia, M – dobry fachowiec.

Wypełnianie diagramu dobrze jest zacząć od wpisania informacji niesionych przez drugą przesłankę – a więc minusów w obszarze zbioru M leżącym poza zbiorem P. Gdy tak postąpimy, nie będziemy mieli wątpliwości, gdzie należy wpisać plus w części wspólnej S oraz M, co nakazuje nam pierwsza przesłanka. Diagram wygląda zatem następująco:



Czy tak wypełniony diagram gwarantuje nam prawdziwość konkluzji sylogizmu? Konkluzja ta ma postać $S \text{ a } P$, a więc stwierdza, że nic nie może się znajdować w obszarze zbioru S leżącym poza zbiorem P. Na rysunku w jednej części tego obszaru mamy minus (a więc tam faktycznie na pewno niczego tam nie ma), natomiast w części drugiej nie znajdujemy żadnego znaku. To, że w danej części nie wstawiliśmy żadnego symbolu, nie oznacza jednak, że niczego tam być nie może, a jedynie, że nie posiadamy żadnych informacji odnośnie tego obszaru. Tak więc wypełniony w ten sposób diagram nie gwarantuje nam wcale, że część zbioru S leżąca poza zbiorem P jest na pewno pusta. W związku z tym **sylogizm należy uznać za niepoprawny.**

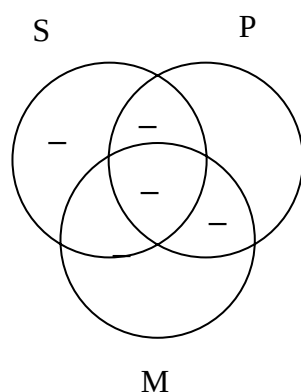
Graficzny kontrprzykład do reguły, na której opiera się badany sylogizm wygląda następująco:



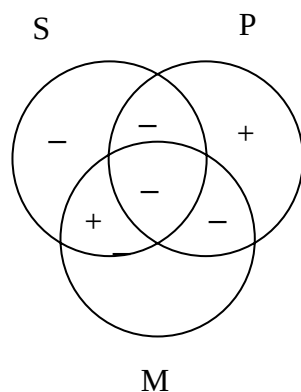
Inny kontrprzykład uzyskać można podstawiając za zmienne nazwy: S – *pies*, P – *jamnik*, M – *jamnik* (pamiętamy, że za różne zmienne wolno podstawiać te same nazwy). Otrzymamy wtedy przesłanki: *niektóre psy są jamnikami*, *każdy jamnik jest jamnikiem* (prawda) oraz wniosek: *każdy pies jest jamnikiem* (fałsz).

Nazwy nie mogą być puste.

Jak dotąd nie powiedzieliśmy jeszcze o jednej ważnej sprawie związanej ze sprawdzaniem poprawności sylogizmów. Otóż zawsze należy przyjąć milczące założenie, że terminy oznaczane symbolami S, P oraz M nie są tak zwanymi nazwami „pustymi”. **Nazwa pusta**, to mówiąc najprościej taka, która nie posiada ani jednego desygnatu, czyli taka, że nie istnieje ani jeden oznaczany przez nią obiekt. Nazwami pustymi są więc na przykład: *jednorożec*, *człowiek o wzroście 3 m*, *obecny król polski* itp. W sylogizmach takich nazw nie wolno nam stosować. Fakt ten niesie ze sobą istotną konsekwencję jeśli chodzi o wypełnianie diagramów Venna. Załóżmy na przykład, że na podstawie przesłanek sylogizmu otrzymaliśmy taki rysunek:



Spójrzmy teraz na obszary odpowiadające zbiorom S oraz P. Każdy z tych obszarów składa się z czterech części, z których w trzech są znaki „-” świadczące o tym, że nic w nich nie ma. Jaki można stąd wyciągnąć wniosek w połączeniu z faktem, że wykorzystane w sylogizmie nazwy na pewno nie są puste? Oczywiście taki, że z całą pewnością coś musi się znajdować w czwartej części każdego z tych obszarów. A zatem w części te możemy, a nawet powinniśmy wpisać znaki „+”:



Założenie o niepustości terminów nie jest wykorzystywane zbyt często, jednak czasami jest ono konieczne, aby właściwie ocenić poprawność sylogizmu.

Przykład:

Zbadamy poprawność sylogizmu: *Każdy pies jest ssakiem. Każdy ssak jest kręgowcem. Zatem niektóre kręgowce są psami.*

Reguła na której opiera się powyższy sylogizm wygląda następująco:

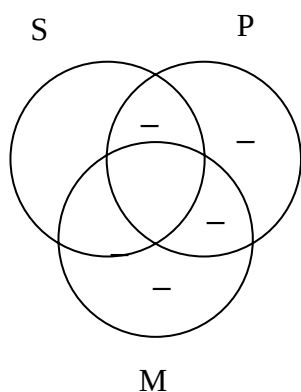
P a M

M a S

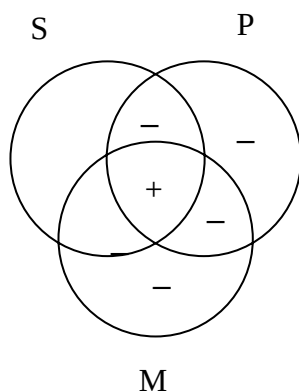
S i P

S – kręgowiec, P – pies, M – ssak.

Po wpisaniu do diagramu informacji z przesłanek mamy rysunek:



Zanim przystąpimy do sprawdzenia, czy taki rysunek gwarantuje nam prawdziwość konkluzji, powinniśmy jeszcze skorzystać z założenia o niepustości nazw użytych w sylogizmie, a konkretnie o niepustości nazwy P. Ponieważ w trzech częściach zbioru skupiającego obiekty określane przez P nic na pewno nie ma, jakieś elementy muszą znajdować się w czwartej części tego zbioru:



Konkluzja badanego sylogizmu stwierdza, że coś znajduje się w części wspólnej zbiorów S oraz P. Na rysunku widzimy, że w obszarze tym znajduje się plus, a więc wniosek ten jest na pewno prawdziwy. **Sylogizm ten jest zatem poprawny.** Aby tę poprawność wykazać, musieliśmy jednak skorzystać z założenia o niepustości terminu P. Gdybyśmy tego nie uczynili, wynik sprawdzania poprawności sylogizmu byłby nieprawidłowy.

Czy ten sylogizm jest na pewno poprawny?

Czasem wynik sprawdzenia poprawności sylogizmu może wydać się dość dziwny lub nawet ewidentnie sprzeczny ze zdrowym rozsądkiem.

Przykład:

Zbadamy poprawność następującego sylogizmu: *Żaden ptak nie jest ssakiem. Niektórzy ludzie są ptakami. Zatem niektórzy ludzie nie są ssakami.*

Sylogizm powyższy opiera się na następującej regule:

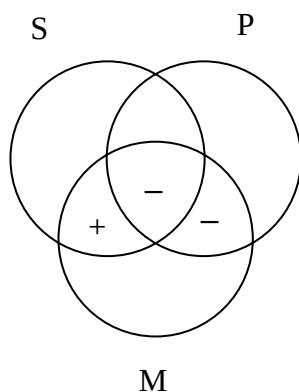
$M \text{ e } P$

$S \text{ i } M$

$S \text{ o } P$

S – człowiek, P – ssak, M – ptak.

Diagram wypełniony według przesłanek wygląda następująco:



Jak widać, diagram ten gwarantuje nam prawdziwość wniosku stwierdzającego, iż niektóre S nie są P, czyli, że coś powinno się znajdować w części zbioru S leżącej poza P. Tak więc sylogizm powyższy należy uznać za **poprawny**.

Odpowiedź taka może jednak budzić pewne opory: jak można uznać za poprawne wnioskowanie, które doprowadziło do jawnie fałszywego wniosku? Oto krótkie wyjaśnienie tego problemu.

Sylogizm powyższy jest poprawny pod tym względem, że jego wniosek wynika logicznie z przesłanek. Tak określona poprawność nazywana jest poprawnością formalną – i jest to ten rodzaj poprawności, jaka interesuje logików. Jednakże badane wnioskowanie nie jest tak całkiem bez zarzutu. Został popełniony w nim błąd polegający na przyjęciu fałszywej przesłanki, co w konsekwencji doprowadziło do otrzymania fałszywego wniosku. Błąd taki nazywany jest **błędem materialnym**. Tak więc odpowiedź do powyższego zadania, mówiąca, że badany sylogizm jest formalnie (logicznie) poprawny, możemy uzupełnić dodając, iż jest on jednak niepoprawny materialnie.



Sylogizm niepoprawny materialnie

Prawdziwość wniosku to jeszcze nie wszystko.

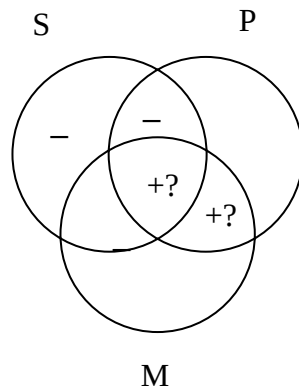
Niejako odwrotność poprzedniego przykładu stanowić może rozumowanie prowadzące do wniosku w sposób oczywisty prawdziwego.

Przykład:

Zbadamy poprawność następującego sylogizmu: *Każdy pies jest ssakiem. Niektóre ssaki mają czarną sierść. Zatem niektóre psy mają czarną sierść.*

Powyższy sylogizm na pierwszy rzut oka mógłby się wydać poprawny: zarówno przesłanki jak i wniosek są na pewno zdaniem prawdziwymi. Czy jednak wnioskowanie to jest na pewno prawidłowe? Reguła na której się ono opiera i wypełniony na jej podstawie diagram wyglądają następująco:

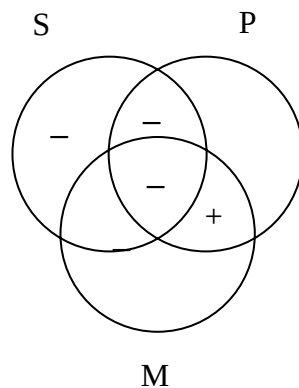
$$\begin{array}{l} S \text{ a } M \\ M \text{ i } P \\ \hline S \text{ i } P \end{array}$$



Powyższy rysunek nie gwarantuje prawdziwości wniosku, czyli tego, że w części wspólnej S oraz P coś się na pewno znajduje. Tak więc badany **sylogizm jest niepoprawny**.

Sylogizm ten jest niepoprawny, ponieważ pomimo prawdziwości przesłanek i wniosku, wniosek nie wynika logicznie z przesłanek. To, że wszystko są to zdania prawdziwe, jest pewnego rodzaju zbiegiem okoliczności, a nie zachodzących pomiędzy nimi związków logicznych.

Graficzny kontrprzykład stanowi następujący rysunek:



Kontrprzykład wykazujący zawodność powyższej reguły uzyskać można również podstawiając za zmienne następujące nazwy: S – *jamnik*, P – *pudel*, M – *pies*.



2.2.4. CZĘSTO ZADAWANE PYTANIA.

Czy kolejność wpisywania do diagramu przesłanek jest dowolna?

Tak, ponieważ ostatecznie i tak zawsze musimy wpisać wszystko co wiemy z obu przesłanek. Dobrze jest jednak zaczynać od przesłanki będącej zdaniem ogólnym („a” lub „e”), która daje nam „pewne” informacje odnośnie znaków „–” w diagramie.

2.3. SPRAWDZANIE POPRAWNOŚCI SYLOGIZMÓW PRZY POMOCY METODY 5 REGUŁ.

2.3.1. ŁYK TEORII.



ŁYK TEORII

Metoda diagramów Venna nie jest jedynym sposobem, w jaki można badać poprawność sylogizmu. Obecnie przedstawimy metodę opartą na pięciu regułach jakie spełniać musi każdy prawidłowy sylogizm. Sprawdzenie poprawności sylogizmu będzie polegało na zbadaniu, czy spełnia on wszystkie warunki sformułowane w owych regułach. Jeżeli tak, należy go uznać za poprawny; jeśli nie spełnia on choć jednego warunku – świadczy to o jego niepoprawności.

Zanim przedstawimy reguły poprawnego sylogizmu, konieczne będzie wprowadzanie nowego pojęcia – mianowicie tak zwanego **terminu rozłożonego** w zdaniu kategorycznym. Otóż, jeżeli zdanie udziela nam informacji o całym zakresie jakiejś nazwy (czyli o jej wszystkich desygnatach), to nazwa ta jest właśnie terminem rozłożonym w tym zdaniu.

W zdaniu *każde S jest P* mowa jest o wszystkich S, a zatem to właśnie S jest w nim terminem rozłożonym. Zdanie *żadne S nie jest P* informuje nas, że ani jeden desygnat nazwy S nie jest desygnatem nazwy P, ani też żaden desygnat P nie jest desygnatem S – a więc

stwierdza fakt dotyczący całych zakresów obu tych nazw. W zdaniu $S e P$ rozłożone są zatem oba terminy. W zdaniu *niektóre S są P* mowa jest o tylko niektórych S, które są „niektórymi” P – w zdaniu tym żaden z terminów nie jest więc rozłożony. Zdanie *niektóre S nie są P* stwierdza, że niektórych desygnatów nazwy S nie ma w całym zakresie nazwy P, a więc rozłożony jest tu termin P.

W skrócie:

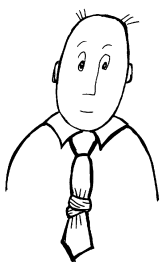
$S a P$ – rozłożony termin S

$S e P$ – rozłożone obydwa terminy – S oraz P

$S i P$ – żaden termin nie jest rozłożony

$S o P$ – rozłożony termin P.

Do sprawdzania sylogizmów metodą pięciu reguł trzeba też pamiętać, które zdania są ogólne ($S a P$ oraz $S e P$), a które szczegółowe ($S i P$ oraz $S o P$), które są twierdzące ($S a P$ oraz $S i P$), a które przeczące ($S e P$ oraz $S o P$), a także to, że M nazywany jest terminem średnim sylogizmu.



DO ZAPAMIĘTANIA:

A oto pięć reguł jakie musi spełniać poprawny sylogizm:

1. Termin średni musi być przynajmniej w jednej przesłance rozłożony.
2. Przynajmniej jedna przesłanka musi być zdaniem twierdzącym.
3. Jeśli jedna z przesłanek jest zdaniem przeczącym, to i wniosek musi być zdaniem przeczącym.
4. Jeśli obie przesłanki są zdaniami twierdzącymi, to i wniosek musi być twierdzący.
5. Jeśli jakiś termin ma być rozłożony we wniosku, to musi być i rozłożony w przesłance.

Sprawdzenie poprawności sylogizmu według powyższych reguł jest bardzo proste: jeżeli choć jeden z wymienionych w nich warunków został złamany, sylogizm należy odrzucić jako błędny; w przeciwnym wypadku jest on poprawny.

2.3.2. PRAKTYKA: ZASTOSOWANI METODY 5 REGUŁ.

Zbadamy przy pomocy omawianej metody kilka sylogizmów sprawdzonych już poprzez diagramy Venna. Nie będziemy przy tym przytaczać całej treści przesłanek i wniosku, a jedynie odpowiednią regułę.

Przykład:

Sprawdzimy poprawność sylogizmu badanego już wyżej przy pomocy diagramów Venna: *Żaden mędrzec nie jest fanatykiem jednej idei. Niektórzy uczeni są fanatykami jednej idei. Zatem niektórzy uczeni nie są mędrcami.* Reguła na której opiera się ten sylogizm przedstawia się następująco:

P e M

S i M

S o P

1 warunek jest spełniony, ponieważ termin M jest rozłożony w pierwszej przesłance;

2 warunek jest spełniony, ponieważ druga przesłanka jest zdaniem twierdzącym;

3 warunek jest spełniony – pierwsza przesłanka i wniosek są zdaniem przeczącymi;

4 warunek nie ma zastosowania do badanego sylogizmu, ponieważ mówi on, co powinno nastąpić, gdyby obie przesłanki były twierdzące. Jako że jedna przesłanka jest zdaniem przeczącym, złamanie czwartej reguły jest w przypadku powyższego sylogizmu niemożliwe;

5 warunek jest spełniony. We wniosku rozłożony jest termin P, a równocześnie jest on rozłożony w pierwszej przesłance.

Ponieważ żaden z warunków nie został złamany, sylogizm należy uznać za poprawny.

Przykład:

Zbadamy poprawność innego rozpatrywanego już sylogizmu: *Niektórzy politycy są nacjonalistami. Każdy nacjonalista jest ograniczony. Zatem niektórzy politycy są ograniczeni.*

S i M

M a P

S i P

1 warunek jest spełniony – termin M jest rozłożony w drugiej przesłance;

2 warunek jest spełniony – obie przesłanki są twierdzące;

3 warunek nie ma zastosowania do badanego przykładu, a więc nie mógł zostać złamany;

4 warunek jest spełniony – obie przesłanki są twierdzące i wniosek także;

5 warunek nie ma zastosowania, ponieważ w badanym sylogizmie żaden termin nie jest rozłożony we wniosku.

Ponieważ żaden warunek nie został złamany, sylogizm jest poprawny.

Przykład:

Zbadamy poprawność kolejnego rozpatrywanego wcześniej sylogizmu: *Niektórzy wykładowcy są dobrymi fachowcami. Każdy dobry fachowiec dużo zarabia. Zatem każdy wykładowca dużo zarabia.*

S i M

M a P

S a P

Warunki 1, 2, 3 i 4 są spełnione (przy czym warunek 3 dzięki temu, że nie ma on bezpośredniego zastosowania). W powyższym sylogizmie złamana została jednakże piąta reguła – termin S pomimo tego, że jest rozłożony we wniosku, nie jest rozłożony w przesłance. Ponieważ jeden z warunków nie został spełniony, sylogizm należy uznać za niepoprawny.

Przykład:

Na koniec sprawdzimy poprawność sylogizmu: *Każdy milioner jest bogaty. Niektórzy bogaci ludzie nie są szczęśliwi. Zatem niektórzy milionerzy nie są szczęśliwi.*

S a M

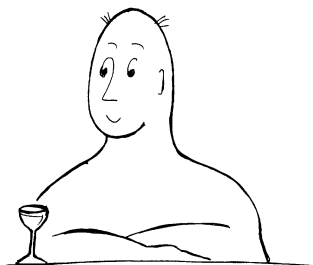
M o P

S o P

W powyższym sylogizmie złamana została już pierwsza reguła – termin średni nie jest rozłożony w żadnej przesłance. W związku z powyższym możemy już w tym momencie odrzucić sylogizm jako błędny, nie sprawdzając dalszych warunków. Dla porządku tylko dodajmy, że pozostałe reguły nie zostały złamane.

2.4. KWADRAT LOGICZNY.

2.4.1. ŁYK TEORII.



ŁYK TEORII

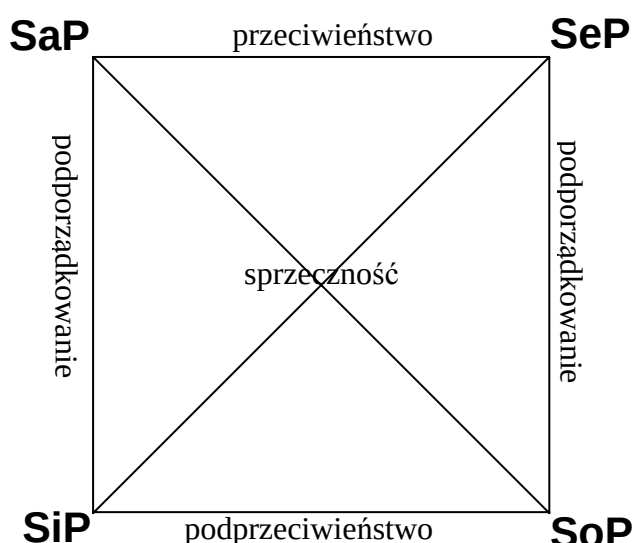
Omawiane w poprzednich paragrafach sylogizmy to wnioskowania mające zawsze dwie przesłanki. Jednakże zdania kategoryczne (*każde S jest P*, *żadne S nie jest P*, *niektóre S są P* oraz *niektóre S nie są P*) wykorzystuje się też czasem w tak zwanych wnioskowaniach bezpośrednich – rozumowaniach, w których występuje tylko jedna przesłanka, na podstawie której wyciąga się pewną konkluzję. Poprawność tego rodzaju wnioskowań badać

można przy pomocy tak zwanego kwadratu logicznego (omówionego w niniejszym paragrafie) oraz innych praw logiki tradycyjnej (przedstawionych w paragrafie 2.5).

Kwadrat logiczny pokazuje związki logiczne zachodzące pomiędzy zdaniami kategorycznymi. Znajomość tych zależności pozwala stwierdzić, jaka jest wartość logiczna pewnego zdania, na podstawie wartości innego zdania. Przykładowo, wiedząc, że prawdziwe jest zdanie SaP możemy z całkowitą pewnością stwierdzić, że prawdziwe jest również zdanie SiP, natomiast fałszywe SeP oraz SoP.

Zależności w kwadracie logicznym przedstawiane są przy pomocy linii. Każda z tych zależności ma swoją nazwę, która zostanie podana przy odpowiedniej linii.

Kwadrat logiczny wygląda następująco:



Zależności kwadratu logicznego – podporządkowanie, przeciwieństwo, podprzeciwieństwo i sprzeczność, przedstawimy w postaci odpowiednich wzorów, które, dla wygody w dalszych rozważaniach, ponumerujemy. Znak negacji w tych wzorach, postawiony przed danym zdaniem, będzie wskazywał, że zdanie to jest fałszywe. Przykładowo, wzór: $SaP \rightarrow \sim (SeP)$ (jeśli SaP to nieprawda, że SeP) odczytamy – prawdziwość zdania SaP implikuje fałszywość SeP (jeśli SaP jest prawdziwe, to SeP jest fałszywe).

Aby prawa kwadratu logicznego miały sens, należy pamiętać o specyficznym rozumieniu zdań SiP oraz SoP. Zdanie *niektóre S są P* oznacza w tym rozumieniu *istnieje (przynajmniej jedno) S będące P*. Natomiast *niektóre S nie są P* – *istnieje (przynajmniej jedno) S nie będące P*.

Należy również nadmienić, że prawa kwadratu logicznego obowiązują jedynie dla nazw niepustych. Oznacza to, że terminy S oraz P muszą mieć jakieś desygnaty. Nie mogą być to wyrażenia typu: *żonaty kawaler*, *niebieski krasnoludek* itp.

Podporządkowanie.

Pionowe linie reprezentują to **podporządkowanie**. Zależność ta polega na tym, że gdy prawdziwe jest zdanie „górne”, to prawdziwe jest też „dolne”. Symbolicznie:

- 1) $SaP \rightarrow SiP$,
- 2) $SeP \rightarrow SoP$

Na przykład, gdy prawdziwe jest zdanie *każda kura jest ptakiem*, to prawdziwe jest też *niektóre kury są ptakami* (lub lepiej: *istnieją kury będące ptakami*). Gdy prawdziwe jest *żadna krowa nie jest ptakiem*, to prawdziwe jest też *niektóre krowy nie są ptakami* (lub lepiej: *istnieją krowy nie będący ptakami*).

Możemy też powiedzieć, że zdanie „dolne” wynika ze zdania, któremu jest podporządkowane.

Przeciwieństwo.

Pozioma linia na górze pomiędzy SaP oraz SeP to **przeciwieństwo**. Polega ono na tym, że wymienione zdania nie mogą być zarazem prawdziwe. Czyli, gdy jedno jest prawdziwe, to drugie musi być fałszywe. Symbolicznie:

- 3) $SaP \rightarrow \sim (SeP)$,
- 4) $SeP \rightarrow \sim (SaP)$

Na przykład gdy prawdziwe jest zdanie *każda papuga jest ptakiem* to fałszywe musi być *żadna papuga nie jest ptakiem*. Natomiast, gdy prawdziwe jest *żadna krowa nie jest ptakiem*, to fałszywe musi być *każda krowa jest ptakiem*.

Zdania przeciwne mogą być jednak jednocześnie fałszywe. Przykładowo fałszywe jest zarówno zdanie *każda krowa jest czarna* oraz *żadna krowa nie jest czarna*.

W przypadku zdań przeciwnych możemy też powiedzieć, że zdania te się wykluczają.

Podprzeciwieństwo.

Pozioma linia na dole, łącząca zdania SiP oraz SoP, to **podprzeciwieństwo**. Zdania podprzeciwne nie mogą być zarazem fałszywe. Czyli, gdy jedno jest fałszywe, to drugie musi być prawdziwe. Symbolicznie:

$$5) \sim (\text{SiP}) \rightarrow \text{SoP}$$

$$6) \sim (\text{SoP}) \rightarrow \text{SiP}$$

Przykładowo, gdy fałszywe jest zdanie *niektóre kanarki są niedźwiedziami*, to prawdziwe jest *niektóre kanarki nie są niedźwiedziami* (lub lepiej: *istnieją kanarki nie będące niedźwiedziami*). Gdy natomiast fałszywe jest zdanie *niektóre żaby nie są płazami*, to prawdziwe musi być *niektóre żaby są płazami* (lub lepiej: *istnieją żaby będące płazami*).

Zdania podprzeciwne mogą być jednak jednocześnie prawdziwe, przykładowo: *niektórzy Polacy są katolikami* i *niektórzy Polacy nie są katolikami*.

W przypadku zdań podprzeciwnych możemy też powiedzieć, że zdania te się dopełniają.

Sprzeczność.

Linie skośne, łączące zdanie SaP z SoP oraz SeP z SiP, reprezentują **sprzeczność**. Sprzeczność oznacza, że zdania te nie mogą być zarazem ani prawdziwe, ani fałszywe. Mówiąc inaczej, mają one zawsze różną wartość logiczną; gdy jedno prawdziwe, to drugie fałszywe, a gdy jedno fałszywe, to drugie prawdziwe. Symbolicznie:

$$7) \text{SaP} \rightarrow \sim (\text{SoP})$$

$$8) \sim (\text{SaP}) \rightarrow \text{SoP}$$

$$9) \text{SoP} \rightarrow \sim (\text{SaP})$$

$$10) \sim (\text{SoP}) \rightarrow \text{SaP}$$

$$11) \text{SeP} \rightarrow \sim (\text{SiP})$$

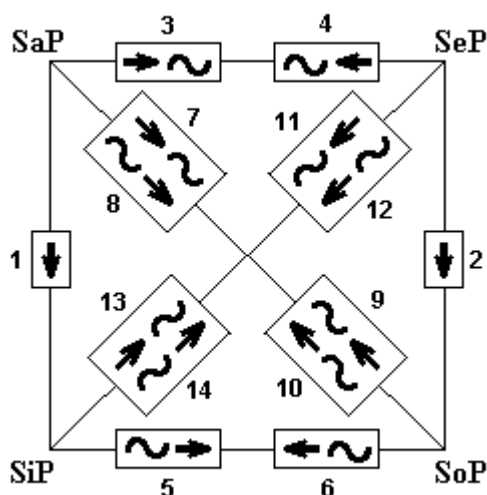
$$12) \sim (\text{SeP}) \rightarrow \text{SiP}$$

$$13) \text{SiP} \rightarrow \sim (\text{SeP})$$

$$14) \sim (SiP) \rightarrow SeP$$

Przykładowo, jeśli prawdziwe jest zdanie *każdy słoń jest ssakiem*, to fałszywe musi być *niektóre słonie nie są ssakami*. Gdy natomiast fałszywe jest zdanie *każdy słoń żyje w Afryce*, to prawdziwe musi być *niektóre słonie nie żyją w Afryce* (wzory 7 i 8). Podobne przykłady łatwo podać również w odniesieniu do pozostałych wzorów.

Poniższy rysunek może pomóc w zapamiętaniu wzorów kwadratu logicznego:



2.4.2. PRAKTYKA: WYKORZYSTANIE KWADRATU LOGICZNEGO.

Zadania związane z kwadratem logicznym polegają zwykle na tym, że na podstawie prawdziwości lub fałszywości podanego zdania kategorycznego, należy określić wartość logiczną pozostałych zdań, w których występują te same terminy S oraz P.

Przykład:

Prawdziwe jest zdanie: *Każdy struś jest ptakiem*.

Co można powiedzieć na podstawie kwadratu logicznego o innych zdaniach kategorycznych mających ten sam podmiot i orzecznik?

Aby rozwiązać to zadanie, musimy sprawdzić, co wynika z prawdziwości zdania typu SaP, a więc, w praktyce, poszukać wzorów rozpoczynających się od SaP. Wzór 1) mówi, że prawdziwe musi być również zdania podporządkowane SaP, czyli SiP – *niektóre strusie są ptakami* (lub lepiej: *istnieją strusie będące ptakami*). Wzór 3) stwierdza, że fałszywe musi być zdanie przeciwne do SaP, a więc SeP – *żaden struś nie jest ptakiem*. Wzór 7) stanowi, że fałszywe musi być zdanie sprzeczne z SeP, czyli SoP – *niektóre strusie nie są ptakami*.

Przykład:

Fałszywe jest zdanie: *Niektórzy goście dotrwali do końca imprezy.*

Sprawdzimy wartość logiczną pozostałych zdań kategorycznych o tym samym podmiocie i orzeczniku.

Szukamy wzorów, które mówią, co wynika z fałszywości zdania SiP. Zgodnie ze wzorem 5) widzimy, że prawdziwe musi być zdanie SoP – *niektórzy goście nie dotrwali do końca imprezy* (lub lepiej: *istnieją goście, którzy nie dotrwali do końca imprezy*). Wzór 14) stwierdza natomiast, że prawdziwe musi być zdanie sprzeczne z SiP, czyli SeP – *żaden z gości nie dotrwał do końca imprezy*.

Nie mamy więcej wzorów zaczynających się od $\sim$ (SiP). Jednakże mamy kolejne dane: dowiedzieliśmy się przed chwilą, że prawdziwe są zdania SoP i SeP. Musimy więc sprawdzić, czy z tych faktów nie da się jeszcze czegoś wywnioskować. Wzór 2) stwierdza coś, co już wiemy – że prawdziwe jest SoP. Natomiast wzory 4) i 9) dają nam nową informację: fałszywe jest zdanie SaP – *każdy gość dotrwał do końca imprezy*.

Przykład:

Fałszywe jest zdanie: *Każdy polityk jest uczciwy.*

Co można powiedzieć na podstawie kwadratu logicznego o innych zdaniach kategorycznych z tymi samymi terminami S oraz P?

Szukamy wzorów, które mówią, co wynika z fałszywości zdania SaP, czyli tych, które zaczynają się od $\sim$ (SaP). Znajdujemy tylko jeden taki wzór – 8). A zatem możemy stwierdzić, że prawdziwe jest zdanie SoP, czyli *niektórzy politycy nie są uczciwi*. Więcej z fałszywości zdania SaP nie da się wywnioskować. Szukamy więc, czy może czegoś więcej dowiemy się na podstawie informacji o prawdziwości SoP. Wzór 9) stwierdza to, co już wiemy, że fałszywe jest SaP. Widzimy więc, że na podstawie kwadratu logicznego nie jesteśmy zatem w żaden sposób określić wartości logicznej zdań SiP oraz SeP, czyli: *niektórzy politycy są uczciwi* oraz *żaden polityk nie jest uczciwy*. Możemy co najwyżej stwierdzić, że, ponieważ są to zdania sprzeczne, mają one różne wartości logiczne; które jest jednak prawdziwe, a które fałszywe, tego z kwadratu logicznego się nie dowiemy.

Przykład:

Prawdziwe jest zdanie: *Niektórzy złodzieje nie są politykami.*

Sprawdzimy, co możemy powiedzieć na podstawie kwadratu logicznego o innych zdaniach kategoriycznych z tym samym podmiotem i orzecznikiem.

Znajdujemy tylko jeden wzór zaczynający się od SoP. Wzór 9) stwierdza, że fałszywe musi być zdanie sprzeczne z SoP, czyli SaP – *każdy złodziej jest politykiem.*

O pozostałych zdaniach, czyli SiP oraz SeP, nic nie możemy powiedzieć.

DO ZAPAMIĘTANIA:



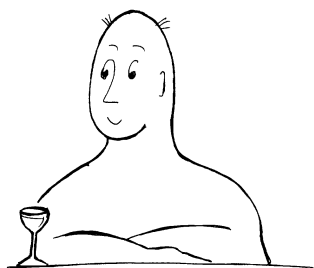
Znając wartość logiczną jakiegokolwiek zdania kategoriycznego, jesteśmy w stanie określić prawdziwość lub fałszywość przynajmniej jednego zdania o tym samym podmiocie i orzeczniku – zdanie sprzeczne z badanym zawsze będzie miało inną wartość.

Najwięcej jesteśmy w stanie powiedzieć na podstawie informacji o prawdziwości zdań ogólnych, czyli SaP i SeP oraz fałszywości szczegółowych SiP oraz SoP. Możemy wtedy zawsze określić wartości wszystkich pozostałych zdań.

Najmniej możemy wywnioskować z prawdziwości zdań szczegółowych (SiP oraz SoP) oraz fałszywości zdań ogólnych (SaP i SeP) – jedynie to, że odwrotną wartość posiada zdanie sprzeczne z badanym zdaniem.

2.5. INNE PRAWA WNIOSKOWANIA BEZPOŚREDNIEGO.

2.5.1. ŁYK TEORII.



ŁYK TEORII

Zależności kwadratu logicznego nie są jedynymi prawami wnioskowania bezpośredniego. Poniżej omówimy pozostałe.

W przedstawionych niżej prawach występować będą często tak zwane nazwy negatywne typu *nie-student*, *nie-pies*, *nie-wydra*, itp. Nazwy te będziemy oznaczać przy pomocy znaku „prim”. Przykładowo, jeśli przez S oznaczmy nazwę *człowiek*, to *nie-człowiek* zapiszemy S'. Zbiór desygnatów (denotację)

nazwy S' stanowić będzie zbiór dopełniający się ze zbiorem desygnatów S . Czyli, przykładowo, jeśli S to nazwa *książka*, to denotacją S' będzie zbiór wszystkich obiektów nie będących książkami.

Zakres nazwy negatywnej można rozumieć na dwa sposoby. Na przykład, dla jednej osoby *nie-pies* może oznaczać tylko zwierzęta nie będące psami (czyli bobry, chomiki, dziecioty, foki itp.), natomiast dla kogoś innego wszystkie obiekty nie będące psami, a więc oprócz zwierząt również np. książki, samochody, telefony itp. W naszych rozważaniach nie będziemy zwykle precyzować, o jakie znaczenie nam chodzi, przyjmując domyślnie takie, które wydaje się bardziej właściwe w danym kontekście.

Przy rozwiązywaniu niektórych zadań istotna będzie czasami znajomość oczywistego faktu, iż dwa przeczenia się znoszą. Przykładowo *nie-nie-ptak*, to to samo, co po prostu *ptak*. A zatem $(S')' \equiv S$

Przedstawione poniżej prawa wnioskowania bezpośredniego obowiązują, podobnie jak prawa kwadratu logicznego, jedynie dla nazw niepustych, czyli takich, które mając jakieś desygnaty. Dodatkowo, nie mogą być to też tak zwane nazwy uniwersalne – czyli obejmujące swym zakresem wszystkie przedmioty.

Konwersja.

Konwersja polega na zmianie miejsc podmiotu i orzecznika zdania bez zmiany jego jakości (czyli zdanie przeczące ma zostać przeczącym, a twierdzące – twierdzącym). Poniższe wzory pokazują, jaki rodzaj zdania wtedy otrzymujemy.

1) SeP $\rightarrow$ PeS

2) SiP $\rightarrow$ PiS

3) SaP $\rightarrow$ PiS

Zdanie SoP nie podlega konwersji.

Przykładowo, ze zdania *żadna krowa nie jest strusiem*, możemy na mocy konwersji wywnioskować, że *żaden struś nie jest krową*; ze zdania *niektórzy ministrowie są przestępcami* – *niektórzy przestępcy są ministrami*; a ze zdania *każdy kij ma dwa końce*, zdanie *niektóre przedmioty mające dwa końce są kijami*.

Obwersja.

Obwersja polega na dodaniu negacji do orzecznika zdania z jednoczesną zmianą (tylko) jego jakości. Tak więc ze zdania twierdzącego otrzymujemy przeczące, a z przeczącego twierdzące.

$$4) SaP \rightarrow SeP'$$

$$5) SeP \rightarrow SaP'$$

$$6) SiP \rightarrow SoP'$$

$$7) SoP \rightarrow SiP'$$

Przykładowo, ze zdania *każdy tygrys jest drapieżnikiem*, wynika, na mocy obwersji zdanie *żaden tygrys nie jest nie-drapieżnikiem*; ze zdania *żadna mrówka nie jest słoniem*, zdanie *każda mrówka jest nie-słoniem*, ze zdania *niektórzy posłowie są idiotami*, zdanie *niektórzy posłowie nie są nie-idiotami*, a ze zdania *niektórzy bogacze nie są skąpcami*, zdanie *niektórzy bogacze są nie-skąpcami*.

Kontrapozycja.

Mówimy o kontrapozycji częściowej (zamiana miejscami podmiotu i orzecznika oraz zanegowanie tego drugiego) oraz zupełnej (zamiana miejscami podmiotu i orzecznika oraz zanegowanie obu). Kontrapozycji nie podlega zdanie SiP.

Kontrapozycja częściowa:

$$8) SaP \rightarrow P'eS$$

$$9) SeP \rightarrow P'iS$$

$$10) SoP \rightarrow P'iS$$

Kontrapozycja zupełna:

$$11) SaP \rightarrow P'aS'$$

$$12) SeP \rightarrow P'oS'$$

$$13) SoP \rightarrow P'oS'$$

Przykładowo, ze zdania *każdy śledź jest rybą* wynika zdanie *żadna nie-ryba nie jest śledziem* (kontrapozycja częściowa) oraz *każda nie-ryba jest nie-śledziem* (kontrapozycja zupełna), ze zdania *żaden wieloryb nie jest rybą* wynika *niektóre nie-ryby są wielorybami* (k. cz.) oraz *niektóre nie-ryby nie są nie-wielorybami* (k. z.), a ze zdania *niektóre torbacze nie są kangurami* wynika *niektóre nie-kangury są torbaczami* (k. cz.) oraz *niektóre nie-kangury nie są nie-torbaczami* (k. z.).

Inwersja.

Inwersja, podobnie jak kontrapozycja, może być częściowa lub zupełna. Podlegają jej tylko zdania ogólne.

Inwersja częściowa:

$$14) SaP \rightarrow S'oP$$

$$15) SeP \rightarrow S'iP$$

Inwersja zupełna:

$$16) SaP \rightarrow S'iP'$$

$$17) SeP \rightarrow S'oP'$$

Przykładowo, ze zdania *każda mysz jest gryzoniem* wynika zdanie *niektóre nie-myszy nie są gryzoniami* (inwersja częściowa) oraz *niektóre nie-myszy są nie-gryzoniami* (inwersja zupełna). Natomiast ze zdania *żaden indyk nie jest żółwiem*, wynika zdanie *niektóre nie-indyki są żółwiami* (i. cz.) oraz *niektóre nie-żółwie nie są nie-indykami*.

2.5.2. PRAKTYKA: ZASTOSOWANIE PRAW WNIOSKOWANIA BEZPOŚREDNIEGO.

Prawa konwersji, obwersji, kontrapozycji i inwersji wykorzystujemy do sprawdzania, co wynika z danego zdania kategorycznego.

Przykład:

Zobaczmy, co wynika, na mocy poznanych praw, ze zdania: *Żaden demokrat nie jest faszystą*.

Ponieważ nasze zdanie ma postać SeP, możemy z niego wyciągnąć następujące wnioski:

Żaden faszysta nie jest demokratą (konwersja, wzór 1).

Każdy demokrat jest nie-faszystą (obwersja, wzór 5).

Niektórzy nie-faszyści są demokratami (kontrapozycja częściowa, wzór 9).

Niektórzy nie-faszyści nie są nie-demokratami (kontrapozycja zupełna, wzór 12).

Niektórzy nie-demokraci są faszystami (inwersja częściowa, wzór 15).

Niektórzy nie-demokraci nie są nie-faszystami (inwersja zupełna, wzór 17).

Przykład:

Sprawdzimy, co wynika, na mocy poznanych praw, ze zdania: *Każda dobra kochanka jest dyskretna.*

Nasze zdanie ma postać SaP. Widzimy więc, że możemy z niego wyciągnąć następujące wnioski:

Niektóre osoby dyskretne są dobrymi kochankami (konwersja, wzór 3).

Żadna dobra kochanka nie jest kimś niedyskretnym (obwersja, wzór 4).

Żadna osoba nie będąca dyskretną nie jest dobrą kochanką (kontrapozycja częściowa, wzór 8).

Każda osoba niedyskretna jest niedobłą kochanką (kontrapozycja zupełna, wzór 11).

Niektóre osoby nie będące dobrymi kochankami nie są dyskretne (inwersja częściowa, wzór 14).

Niektóre osoby nie będące dobrymi kochankami są niedyskretne (inwersja zupełna, wzór 16).

Czasem już w zdaniu, które poddajemy konwersji, obwersji itd. występują nazwy negatywne. W takich przypadkach, przy dokonywaniu niektórych operacji należy pamiętać o prawie znoszenia się podwójnego przeczenia, a więc: $(S')' \equiv S$.

Przykład:

Sprawdzimy, co na mocy poznanych praw wynika ze zdania: *Żaden nie-ptak nie jest wróblem.*

Nasze zdanie ma postać S'eP. Wynikają z niego następujące zdania:

Żaden wróbel nie jest nie-ptakiem (1).

Każdy nie-ptak jest nie-wróblem (5).

Niektóre nie-wróble są nie-ptakami (9).

Niektóre nie-wróble nie są ptakami (12 po zastosowaniu prawa: $(S')' \equiv S$).

Niektóre ptaki są wróblami (15 po zastosowaniu prawa: $(S')' \equiv S$).

Niektóre ptaki nie są nie-wróblami (17 po zastosowaniu prawa: $(S')' \equiv S$).

Przykład:

Sprawdzimy, co na mocy poznanych praw wynika ze zdania: *Niektóre ptaki są nie-kanarkami*.

Nasze zdanie ma postać SiP'. Wynikają z niego następujące zdania:

Niektóre nie-kanarki są ptakami (2).

Niektóre ptaki nie są kanarkami (6 po zastosowaniu prawa: $(P)' \equiv P$).

SŁOWNICZEK.

Błąd formalny – błąd polegający na tym, że wniosek rozumowania nie wynika logicznie z przesłanek.

Błąd materialny – błąd polegający na użyciu we wnioskowaniu przynajmniej jednej fałszywej przesłanki.

Denotacja nazwy (zakres nazwy) – zbiór wszystkich desygnatów danej nazwy. Przykładowo zbiór wszystkich studentów jest denotacją (zakresem) nazwy *student*.

Desygnat nazwy – obiekt oznaczany przez daną nazwę. Na przykład każdy z nas jest desygnatem nazwy *człowiek*.

Nazwa pusta – nazwa nie posiadająca ani jednego desygnatu. Na przykład *centaur*, *jednorożec*, *człowiek o wzroście 3 m*, *żonaty kawaler* itp.

Przesłanka mniejsza – przesłanka zawierająca termin mniejszy sylogizmu.

Przesłanka większa – przesłanka zawierająca termin większy sylogizmu.

Termin mniejszy sylogizmu – nazwa występująca jako podmiot we wniosku sylogizmu. Termin mniejszy oznacza się zwykle symbolem S.

Termin rozłożony – nazwa, o której całym zakresie (wszystkich desygnatach) jest mowa w zdaniu kategorycznym. W zdaniu $S \text{ a } P$ rozłożone jest S , w $S \text{ e } P$ zarówno S jak i P , w $S \text{ o } P$ – jedynie P . W zdaniu $S \text{ i } P$ żaden termin nie jest rozłożony.

Termin średni sylogizmu – nazwa nie występująca we wniosku sylogizmu, za to obecna w obu jego przesłankach. Termin średni oznacza się zwykle symbolem M .

Termin większy sylogizmu – nazwa występująca jako orzecznik sylogizmu. Termin większy oznacza się zwykle symbolem P .

Zdanie kategoryczne – zdanie mające jedną z następujących postaci (gdzie S i P reprezentują nazwy): *każde S jest P , żadne S nie jest P , niektóre S są P , niektóre S nie są P .*

Rozdział III

KLASYCZNY RACHUNEK PREDYKATÓW.

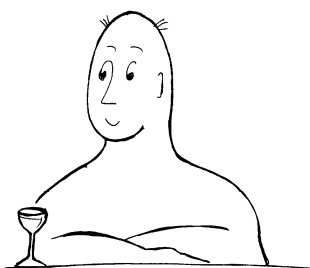
Wstęp.

W niniejszym rozdziale omówiony zostanie kolejny system logiczny, który może służyć do analizy rozumowań – klasyczny rachunek predykatów (KRP), nazywany również klasycznym rachunkiem kwantyfikatorów (KRK). System ten, będąc bardziej złożonym od rachunku zdań czy sylogistyki, nadaje się do analizy takich rozumowań, wobec których tamte systemy są bezradne.

Szerokie pole zastosowania rachunku predykatów okupione zostaje jednakże poważną wadą – system ten jest o wiele bardziej skomplikowany od dotychczas poznanych. Sprawne posługiwanie się nim wymaga znacznej wiedzy i uważane jest czasem za wyższy stopień wtajemniczenia logicznego. W obecnym rozdziale rachunek predykatów przedstawiony zostanie w postaci możliwie najprostszej, jednakże, nawet mimo tego, jego opanowanie będzie wymagało większego wysiłku, niż to było konieczne w przypadku poprzednich systemów. Zrozumienie rachunku predykatów wymaga w miarę sprawnego posługiwania się rachunkiem zdań. Przede wszystkim konieczna jest dobra znajomość spójników logicznych oraz tabelk zero-jedynkowych.

3.1. SCHEMATY ZDAŃ.

3.1.1. ŁYK TEORII.



ŁYK TEORII

Poznanie rachunku predykatów rozpoczniemy, tradycyjnie, od tłumaczenia zdań języka naturalnego na język tego systemu. Schematy zdań na gruncie rachunku predykatów przypominać będą w pewnym stopniu schematy zapisywane w ramach rachunku zdań. Podobieństwo to wynika z obecności w języku rachunku predykatów spójników logicznych – negacji, koniunkcji, alternatywy, implikacji i równoważności. Znaczenia tych spójników oraz reprezentujące je symbole ($\sim$, $\wedge$, $\vee$, $\rightarrow$, $\equiv$)

są tu dokładnie takie same jak w rachunku zdań. W rachunku predykatów mamy jednak również nowe elementy – **predykaty** oraz **kwantyfikatory**. Do pisania schematów będziemy też wykorzystywali tak zwane zmienne indywiduowe, które będą oznaczały dowolne obiekty (indywidua).

Predykaty pełnią w KRP rolę analogiczną do zmiennych zdaniowych w KRZ. To właśnie one, w połączeniu ze zmiennymi indywiduowymi, są tu najprostszymi wyrażeniami, z

których, za pomocą spójników, możemy budować dłuższe zdania. Predykaty symbolizować będziemy przy pomocy dużych liter, np.: P, Q, R, S itd., po których, w nawiasie, będą znajdowały się zmienne indywiduowe, reprezentowane przez małe litery x, y, z itd. Tak więc najprostszymi poprawnymi wyrażeniami na gruncie rachunku predykatów są takie zapisy jak np.: $P(x)$, czy $R(x,y)$. Pierwsze z nich odczytujemy jako *P od x*, a drugie jako *R od x, y*. Wyrażenia złożone otrzymujemy poprzez użycie spójników logicznych. Schemat $P(x) \wedge \sim Q(x)$ odczytamy jako *P od x i nieprawda, że Q od x*. Natomiast $R(x,y) \rightarrow (P(x) \vee P(y))$ – jako *jeśli R od x,y to P od x lub P od y*.

Predykaty są wyrażeniami opisującymi własności lub relacje. Własność to nic innego, jak pewna cecha posiadana przez jakiś obiekt. Własnością jest, na przykład, „bycie inteligentnym” (cecha jakiegoś człowieka), „bycie parzystą” (cecha liczby), „bycie smacznym”, „bycie drogim” itp. itd. Umówmy się, że predykat opisujący jaką cechę oznaczać będziemy zwykle, dla wygody, przy pomocy pierwszej litery tej cechy. I tak, na przykład, fakt, że jakiś obiekt posiada cechę bycia mężczyzną, oznaczmy $M(x)$, bycia bogatym – $B(x)$, bycia zaradczym – $Z(x)$ itp. Gdy w jakimś złożonym wyrażeniu pojawiają się dwie własności zaczynające się na tę samą literę, to oczywiście jedną z nich będziemy musieli oznaczyć inaczej.

Relacje to pewne związki łączące kilka obiektów. Nas będą przede wszystkim interesowały tak zwane relacje dwuargumentowe, będące związkami występującymi pomiędzy dwoma obiektami. Relacją taką jest na przykład „lubienie” (jedna osoba lubi drugą osobę), „bycie wyższym” (ktoś lub coś jest wyższe od kogoś lub czegoś), „okradzenie” (ktoś okradł kogoś) itp. Predykaty oznaczające takie relacje będziemy zapisywali odpowiednio: $L(x,y)$, $W(x,y)$, $O(x,y)$.

Relacjami z większą ilością argumentów nie będziemy się zajmować. Dla porządku podajmy jednak przykłady relacji łączących trzy obiekty. Może być to na przykład „relacja znajdowania się pomiędzy” ($P(x,y,z)$ – obiekt x znajduje się pomiędzy obiektem y a obiektem z), czy też relacja „zdradzania z kimś” ($Z(x,y,z)$ – osoba x zdradza osobę y z osobą z).

Uwaga na marginesie.

Ściśle rzecz biorąc własności też są relacjami – tak zwanymi relacjami jednoargumentowymi. Jednakże, dla większej jasności, w dalszych rozważaniach termin „relacja” zarezerwujemy dla relacji dwuargumentowych, natomiast relacje jednoargumentowe będziemy nazywali „własnościami”.

Kwantyfikatory to wyrażenia określające ilość przedmiotów, o których jest mowa. Z kwantyfikatorami zetknęliśmy się już w sylogistyce, choć tam nie wspominaliśmy, że tak je właśnie nazywamy. W rachunku predykatów będziemy mieli do czynienia z dwoma kwantyfikatorami. Pierwszy z nich odpowiada wyrażeniu *dla każdego* i jest najczęściej oznaczany symbolem $\forall$. Kwantyfikator ten bywa nazywany „dużym kwantyfikatorem” lub „kwantyfikatorem ogólnym”. Drugi z kwantyfikatorów odpowiada wyrażeniu *niektóre*, w znaczeniu *istnieje przynajmniej jedno takie*. Kwantyfikator ten, oznaczany symbolem $\exists$, nazywany jest „małym kwantyfikatorem”, „kwantyfikatorem szczegółowym” lub „kwantyfikatorem egzystencjalnym”.



DO ZAPAMIĘTANIA:

Osoby znające język angielski mogą łatwo zapamiętać znaczenie kwantyfikatorów. Kwantyfikator ogólny to odwrócona litera „A” od angielskiego słowa *All* – czyli *wszystkie*, natomiast kwantyfikator szczegółowy, to odwrócone „E” od słowa *Exists* – *istnieje*.

W schematach zdań, po kwantyfikatorach będą znajdowały się (bez nawiasów, a więc inaczej niż przy predykatach) symbole zmiennych, do których dany kwantyfikator się odnosi, na przykład $\forall x$ oznacza *dla każdego x*, natomiast $\exists y$ – *istnieje takie y lub niektóre y*

Zapis taki jak $\exists x P(x)$ – odczytamy jako *istnieje takie x, że P(x)* lub (mniej formalnie) *istnieje x mające własność P, niektóre x mają własność P* itp.

Kwantyfikatory, inaczej niż predykaty, mogą występować obok siebie nie połączone żadnymi spójnikami. Zapis $\forall x \exists y R(x,y)$ odczytamy *dla każdego x istnieje y, takie że R od x, y* lub *dla każdego x istnieje takie y, że x i y są w relacji R*.

Kwantyfikatory możemy poprzedzać spójnikiem negacji. Przykładowo, wyrażenie $\sim \exists x P(x)$ odczytamy – *nie istnieje takie x, że P od x (nie istnieje x mające własność P, żadne x nie ma własności P)*, natomiast $\exists x \sim \forall y R(x,y)$ – *istnieje x, takie że nie dla każdego y, R (x,y) (istnieje takie x, że nie dla każdego y, x jest do niego w relacji R, istnieje takie x, które nie do wszystkich y jest w relacji R)*.



DO ZAPAMIĘTANIA:

Przedstawmy w skrócie symbole konieczne przy pisaniu schematów zdań na gruncie rachunku predykatów

Spójniki zdaniowe:

$\sim, \wedge, \vee, \rightarrow, \equiv$

Zmienne indywiduowe:

$x, y, z, \dots$ itd.

Symbole predykatów:

$P, Q, R, S, \dots$ itd.

Symbole kwantyfikatorów:

$\forall$ – oznaczający *dla każdego* (tak zwany „duży kwantyfikator” lub „kwantyfikator ogólny”)

$\exists$ – oznaczający *istnieje* lub *niektóre* (tak zwany „mały kwantyfikator”, „kwantyfikator szczegółowy” lub „kwantyfikator egzystencjalny”)

Należy pamiętać, że predykaty występować będą zawsze razem z, ujętymi w nawiasach, zmiennymi np.:

$P(x)$ – zapis oznaczający, że x ma własność P ,

$R(x,y)$ – zapis oznaczający, że x i y są ze sobą w relacji R ,

Kwantyfikatory w praktyce występować będą razem ze zmiennymi nazwowymi, np.: $\forall x$, $\exists y$... itp.

Przy pisaniu schematów będziemy w rachunku predykatów korzystali również z nawiasów, które, podobnie jak w rachunku zdań, pełnią pomocniczą rolę, pokazując co się z czym łączy i likwidując możliwe wieloznaczności.

Do pisania schematów może przydać się jeszcze jedna istotna informacja. Dotyczy ona pojęcia tak zwanej zmiennej związanej przez kwantyfikator oraz zmiennej wolnej (niezwiązanej). Każdy kwantyfikator „wiąże” zmienną, która się przy nim znajduje – np. kwantyfikator $\exists x$ wiąże zmienną x , a $\forall y$ – zmienną y . Kwantyfikatory wiążą jednak nie wszystkie danego typu, ale tylko te, które znajdują się w ich zasięgu – czyli w nawiasie otwartym bezpośrednio po kwantyfikatorze lub, w przypadku braku nawiasu, w wyrażeniu najbliższym kwantyfikatorowi. Najłatwiej wyjaśnić to na przykładzie: w schemacie $\forall x (P(x) \rightarrow Q(x))$ związane są zmienne x w całej formule, natomiast w schemacie $\forall x P(x) \rightarrow Q(x)$ jedynie zmienna znajdująca się przy predykanie P (zmienna przy Q jest w takim razie zmienną wolną). W schemacie $\exists x (P(x) \wedge Q(x,y)) \rightarrow \forall z R(z,x)$ zmienna x jest związana przy predykanie P oraz Q , natomiast wolna przy R ; zmienna y jest wolna (nie ma w ogóle wiążącego jej kwantyfikatora); zmienna z jest związana (przez kwantyfikator $\forall$)

Pojęcie zmiennej wolnej i związanej będzie dla nas istotne, gdyż w prawidłowo zapisanych schematach zdań języka naturalnego nie mogą występować zmienne wolne (mówiąc inaczej wszystkie zmienne muszą być związane jakimś kwantyfikatorem). Z faktu tego wynika istotny wniosek – każdy schemat będzie musiał zaczynać się jakimś (przynajmniej jednym) kwantyfikatorem, który będzie wiązał występujące dalej zmienne. Żadna zmienna nie będzie mogła się pojawić, zanim nie wystąpi wiążący ją kwantyfikator.

Jeśli w schemacie nie ma zmiennych wolnych, to można go zawsze tak odczytać, aby nie wypowiadać słów *iks*, *igrek*, *zet* itp., których przecież w zdaniach języka naturalnego nie używamy. Przykładowo, gdy przyjmiemy, że predykat F oznacza własność bycia filozofem, to schematy $\exists x F(x)$ oraz $\forall x F(x)$ możemy wprawdzie odczytać kolejno: *istnieje x będący filozofem*, oraz *dla każdego x , x jest filozofem*, ale o wiele zgrabniej jest powiedzieć *istnieją filozofowie (niektórzy są filozofami)* oraz *każdy jest filozofem*. Zabieg „pozbycia” się zmiennych nie jest możliwy, gdy są one wolne; schemat $F(x)$ musimy odczytać: *x jest*

filozofem. To ostatnie wyrażenie nie jest na pewno, przynajmniej z punktu widzenia logiki, zdaniem języka naturalnego, a jedynie tak zwaną „formą zdaniową”.

Uwaga na marginesie.

To, że w schematach zdań języka naturalnego nie może być zmiennych wolnych, nie oznacza, że zmiennych takich w ogóle nie może być w formułach rachunku predykatów. W rachunku predykatów mogą istnieć bowiem formuły (m.in. te, które zawierają zmienne wolne) nie będące schematami żadnego zdania języka naturalnego.

3.1.2 PRAKTYKA: BUDOWANIE SCHEMATÓW ZDAŃ NA GRUNCIE KRP.

Przystępując do budowania schematów zdań w ramach rachunku predykatów, musimy sobie przede wszystkim uświadomić, jakie w naszym zdaniu występują własności i/lub relacje i zastąpić je odpowiednimi symbolami predykatów. Następnie powinniśmy się zastanowić, jakie kwantyfikatory będą nam w schemacie potrzebne. Ostatecznie musimy połączyć wszystko w całość przy pomocy spójników i nawiasów, tak aby otrzymać schemat danego zdania.

Pisząc schemat zdania należy pamiętać, że ma to być zawsze tak zwany schemat główny, czyli możliwie najdłuższy, najgłębiej wnikający w strukturę zdania; taki w którym obecne są wszystkie możliwe do wyodrębnienia spójniki, predykaty i kwantyfikatory.

Rozpoczniemy od budowania bardzo prostych schematów zdań, w których występują jedynie własności.

Przykład:

Zapiszemy schemat zdania: *Niektórzy złodzieje są politykami.*

W zdaniu tym jest mowa o dwóch własnościach – byciu złodziejem oraz byciu politykiem; oznaczmy je odpowiednio literami Z i P. Zdanie zaczyna się od zwrotu *niektórzy*, będącego odpowiednikiem kwantyfikatora $\exists$, a więc od tego symbolu

powinien rozpocząć się nasz schemat. Nasze zdanie stwierdza, że istnieją obiekty, które są zarówno złodziejami, jak i politykami (posiadają obie te cechy jednocześnie), w związku z



NIKTÓRZY ZŁODZIEJE SĄ POLITYKAMI

czym potrzebny nam będzie jeszcze spójnik koniunkcji. Ostateczny schemat przedstawia się następująco:

$$\exists x (Z(x) \wedge P(x))$$

Nawias w powyższym schemacie jest konieczny, aby pokazać, że kwantyfikator wiąże zmienną x znajdującą się zarówno przy predykanie Z , jak i przy P .

Przykład:

Zapiszemy schemat zdania: *Każdy rasista jest ograniczony*.

W powyższym zdaniu mowa jest o dwóch własnościach – bycia rasistą i bycia ograniczonym. Mamy tu też słowo *każdy*, będące odpowiednikiem kwantyfikatora ogólnego. Pewnym problemem dla początkujących może być znalezienie odpowiedniego spójnika łączącego predykaty R oraz Q . Gdybyśmy jednak wstawili tu koniunkcję, tak jak w poprzednim przykładzie, otrzymalibyśmy schemat $\forall x (R(x) \wedge O(x))$, czyli wyrażenie mówiące: *każdy jest rasistą i jest ograniczony (każdy jest ograniczonym rasistą)* – a więc na pewno nie zdanie, którego schemat mamy napisać. Nasze zdanie, *Każdy rasista jest ograniczony*, stwierdza, że **jeśli** ktoś jest rasistą, **to** jest on ograniczony, a więc prawidłowy schemat powinien wyglądać:

$$\forall x (R(x) \rightarrow O(x))$$



WARTO ZAPAMIĘTAĆ.

W schematach zdań języka naturalnego rzadko się zdarza, aby w formule związanej przez kwantyfikator $\forall$ głównym spójnikiem była koniunkcja. Na ogół jest to implikacja lub ewentualnie alternatywa. Koniunkcja występuje natomiast zwykle jako główny spójnik formuł związanych przez kwantyfikator $\exists$. Czyli:

$$\forall x (... \rightarrow ...) \text{ lub } \forall x (... \vee ...)$$

$$\exists x (... \wedge ...)$$

Powyższe stwierdzenia nie stanowią jednak w żadnym razie jakichkolwiek praw logicznych. Jest to po prostu użyteczna obserwacja, która sprawdza się w zdecydowanej większości (choć nie wszystkich!) przypadków.

Przykład:

Zapiszemy schemat zdania: *Nie każdy logik jest abstynentem.*

W powyższym zdaniu występują własności bycia logikiem oraz bycia abstynentem. Jest też odpowiednik kwantyfikatora *dla każdego*, jednak poprzedzony słowem *nie*. Tak więc schemat powinien zacząć się od zwrotu: $\sim \forall x$. Jako spójnika łączącego predykaty należy użyć implikacji (wykorzystanie koniunkcji dałoby schemat zdania: *Nie każdy jest logikiem i abstynentem*). Mamy więc:

$$\sim \forall x (L(x) \rightarrow A(x))$$

Przykład:

Zapiszemy schemat zdania: *Niektórzy studenci nie są pilni.*

W zdaniu mowa jest o własnościach bycia studentem i bycia pilnym. Ta druga jest jednak zanegowana. Zdanie stwierdza, że są osoby posiadające własność bycia studentem i jednocześnie nie posiadające własności bycia pilnym. A zatem:

$$\exists x (S(x) \wedge \sim P(x))$$

Przykład:

Zapiszemy schemat zdania: *Żaden dziennikarz nie jest obiektywny.*

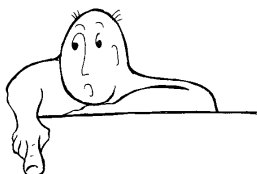
W powyższym zdaniu mamy na pewno do czynienia z własnością bycia dziennikarzem oraz bycia obiektywnym. Kłopot sprawić może wybór odpowiedniego kwantyfikatora. Czemu odpowiadać może słowo *żaden* w rozważanej wypowiedzi? Z jednej strony jest to „negatywny” sposób powiedzenia czegoś o *wszystkich* dziennikarzach – o każdym dziennikarzu zdanie stwierdza, że nie jest obiektywny. Z innego punktu widzenia można jednak również powiedzieć, iż zdanie stwierdza, że *nie istnieje* taki dziennikarz, który posiadałby cechę bycia obiektywnym. Czy schemat zacząć należy zatem wyrażeniem $\forall x$, czy też $\sim \exists x$? Obie odpowiedzi na to pytanie są dobre! Otóż, w przypadku powyższego zdania, napisać możemy dwa równie dobre schematy:

$$\forall x (D(x) \rightarrow \sim O(x)), \text{ oraz}$$

$$\sim \exists x (D(x) \wedge O(x))$$

Oba te schematy są logicznie równoważne; mówią one dokładnie to samo. Dyskusje budzić może, który z nich uznać należy za bardziej pierwotny; lepiej, w sposób bardziej

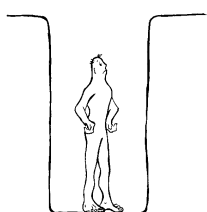
naturalny, oddający strukturę rozpatrywanego zdania. Wielu logików twierdzi, że zdanie typu *żaden... nie jest...* jest zdaniem ogólnym (więcej na ten temat w rozdziale o sylogizmach), a więc jego schemat powinien zaczynać się od kwantyfikatora $\forall$. Inni dopuszczają jednak również drugi schemat, jako w równym stopniu właściwy.



Uwaga na błędy!

Nie zawsze, tak jak w przypadku powyższego przykładu, dwa schematy można uznać za równie dobre, na podstawie tego, że są one logicznie równoważne. Przykładowo do schematu zdania w przykładzie *Nie każdy logik jest abstynentem* można utworzyć równoważny mu schemat: $\exists x (L(x) \wedge \sim A(x))$. W tym jednak przypadku wielu (choć również, nie wszyscy) logików nie uznałoby tego schematu za właściwy. Pomimo, że zdania *Nie każdy logik jest abstynentem* oraz *Niektórzy logicy nie są abstynentami* (literalne odczytanie drugiego schematu) są logicznie równoważne i wyrażają tę samą treść (opisują ten sam fakt), to trudno uznać, że są to te same zdania.

W wielu podobnych przypadkach nie ma zgody, które schematy należy uznać za poprawne, a które nie. Najlepiej kierować się wskazówką, że schemat powinien w sposób najbardziej intuicyjny odzwierciedlać strukturę danego zdania. Jeśli zdanie zaczyna się od zwrotu *nie każdy*, to schemat powinien zacząć się od $\sim \forall$, jeśli zdanie zaczyna się od *niektóre*, to schemat rozpoczynamy od $\exists$.



3.1.3. UTRUDNIENIA I PUŁAPKI.

Obecnie zajmujemy się bardziej złożonymi schematami. Często zdarza się tak, że w przypadku dłuższych zdań istnieje wiele możliwości zbudowania poprawnych schematów. Dopuszczalne są różne możliwości, szczególnie w zakresie stosowania nawiasów i ustawienia kwantyfikatorów. Na omówienie wszystkich tych możliwości i związanych z nimi niuansów nie starczyłoby tu miejsca – wspomniana zostanie tylko część z nich. Dlatego podane niżej rozwiązania należy traktować w niektórych przypadkach jako przykładowe, nie wykluczające innych poprawnych odpowiedzi.

Więcej predykatów.

Oczywiście w formule może znajdować się więcej predykatów niż jeden lub dwa.

Przykład:

Napiszemy schemat zdania: *Nie każdy znany muzyk jest artystą.*

W zdaniu powyższym mamy do czynienia z trzema własnościami – byciem muzykiem, byciem znanym oraz byciem artystą. Zdanie stwierdza, że nie każdy kto posiada dwie pierwsze, posiada również trzecią, czyli, mówiąc bardziej formalnie, nie każdy x , jeśli posiada własność M oraz Z , to posiada też własność A . Schemat będzie wyglądał zatem następująco:

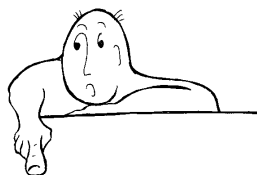


NIE KAŻDY ZNANY MUZYK JEST ARTYSTĄ

$$\sim \forall x [(M(x) \wedge Z(x)) \rightarrow A(x)]$$

W powyższym schemacie koniunkcja $M(x) \wedge Z(x)$ znajduje się w nawiasie, aby wyraźnie było widoczne, że głównym spójnikiem jest tu implikacja. Jeśli chodzi o zastosowanie nawiasów w złożonych formułach, to w rachunku predykatów obowiązują wszystkie zasady znane z rachunku zdań.

Wątpliwości może budzić, czy prawidłowa jest kolejność, w jakiej umieszczone zostały człony koniunkcji, czyli cechy bycia muzykiem i bycia znanym. Kolejność ta jest jednak całkowicie bez znaczenia. Koniunkcja, w jej rozumieniu przyjętym w logice, ma tę własność, że jej człony możemy umieszczać w dowolnej kolejności i nie zmienia to w niczym sensu wyrażenia. Tak więc równie dobry byłby schemat: $\sim \forall x [(Z(x) \wedge M(x)) \rightarrow A(x)]$



Uwaga na błędy!

Nie zawsze jest tak, że dwa określenia (tak jak *znany* i *muzyk* w poprzednim przykładzie) odnoszące się do pewnego obiektu dają się rozłożyć na dwie osobne cechy. Przykładowo, gdybyśmy mieli do czynienia ze zdaniem, w którym znalazłoby się stwierdzenie, że ktoś jest „dobrym rewolwerowcem”, to nie moglibyśmy rozbić tego określenia na cechy bycia dobrym i bycia rewolwerowcem, gdyż wypaczyło by

to sens zdania. Wymienione cechy tworzą całość – jej rozbitcie zmieniłoby znaczenie jednej z nich – bycia dobrym.

Nie istnieje żadna metoda pozwalająca jednoznacznie stwierdzić, kiedy wymienione w zdaniu cechy można i należy rozłożyć, a kiedy jest to niemożliwe. Zawsze będą istniały przypadki graniczne i dyskusyjne. Trudno na przykład ustalić, czy własność bycia „małym słoniem” możemy rozbić na dwie osobne własności – bycia słoniem i bycia małym, czy też trzeba tę własność traktować jako nierozkładalną całość.

Wiele kwantyfikatorów.

W schemacie może oczywiście występować więcej niż jeden kwantyfikator.

Przykład:

Zapiszemy schemat zdania: *Wszystkie inteligentne kobiety mają powodzenie, ale niektóre z kobiet mających powodzenie nie są inteligentne.*

W zdaniu powyższym widzimy trzy własności: bycia kobietą, bycia osobą inteligentną i posiadania powodzenia. Zdanie to składa się jednak z dwóch części połączonych słowem *ale*, czyli odpowiednikiem koniunkcji. Każda z tych części zaczyna się innym kwantyfikatorem – pierwsza ogólnym, druga szczegółowym.

$$\forall x[(K(x) \wedge I(x)) \rightarrow P(x)] \wedge \exists x[(K(x) \wedge P(x)) \wedge \sim I(x)]$$

Pamiętać należy, że, z uwagi na przemienność koniunkcji, równie poprawne byłyby schematy, w których człony koniunkcji znalazłyby się w odwrotnej kolejności.

Co znaczy „tylko”?

Przykład:

Zapiszemy schemat zdania: *Tylko kobiety są matkami.*

W zdaniu tym mamy oczywiście dwie własności: bycia matką i bycia kobietą. Problem stanowić może określenie kwantyfikatora i układu własności w formule. Z podobną trudnością spotkaliśmy się już przy pisaniu schematów na gruncie sylogistyki. Być może niektórzy pamiętają, że zdania typu *Tylko S są P* określiliśmy wtedy jako ogólnotwierdzące, a zatem zaczynające się od kwantyfikatora ogólnego – $\forall$. Jeśli jednak napisalibyśmy schemat: $\forall x (K(x) \rightarrow M(x))$ to otrzymalibyśmy fałszywe zdanie *Każda kobieta jest matką*. Nasze zdanie stwierdza natomiast coś odwrotnego: to, że tylko kobiety są matkami, oznacza, że każda matka jest kobietą. A zatem schemat powinien wyglądać:

$$\forall x (M(x) \rightarrow K(x))$$



DO ZAPAMIĘTANIA.

Schematy zdań typu *Tylko A są B* rozpoczynamy od kwantyfikatora ogólnego a następnie piszemy implikację zamieniając kolejność A i B. Czyli $\forall x (B(x) \rightarrow (A))$.

Co znaczy „tylko niektórzy”?

Rozpatrywane powyżej zdania typu *Tylko A są B* należy koniecznie odróżnić od zdań *Tylko niektóre A są B*.

Przykład:

Zapiszemy schemat zdania: *Tylko niektórzy studenci uczą się systematycznie*.

Zwrot *tylko niektórzy* w powyższym zdaniu oznacza, że istnieją studenci, którzy posiadają cechę U (uczą się systematycznie), ale są również tacy, którzy cechy takiej nie posiadają. Lub inaczej: istnieją studenci mający cechę U, lecz jednocześnie nie wszyscy cechę tę posiadają. Dwa równoprawne schematy powyższego zdania, to zatem:

$\exists x (S(x) \wedge U(x)) \wedge \exists x (S(x) \wedge \sim U(x))$, lub

$\exists x (S(x) \wedge U(x)) \wedge \sim \forall x (S(x) \rightarrow U(x))$

Pojawiają się relacje.

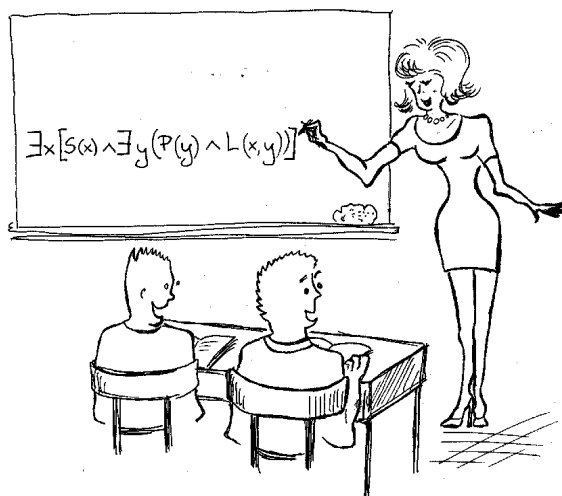
Dotąd rozpatrywaliśmy bardzo proste zdania, w których mieliśmy do czynienia jedynie z predykatami jednoargumentowymi, opisującymi własności. Więcej kłopotów sprawić mogą zdania w których obecne będą predykaty oznaczające relacje. Początkowo zapisywanie takich schematów może wydawać się niezmiernie skomplikowane, między innymi dlatego, że nie ma na to jakiejś jednej, sprawdzającej się zawsze metody. Przerobienie kilku przykładów powinno jednak wiele wyjaśnić.

Po nabraniu pewnej wprawy, zapisywanie schematów zdań w języku predykatów może stać się ciekawą rozrywką intelektualną, podobną np. do rozwiązywania krzyżówek.

Przykład:

Zapiszemy schemat zdania: *Niektórzy studenci lubią niektóre przedmioty.*

W zdaniu powyższym jest mowa o dwóch własnościach – bycia studentem oraz bycia przedmiotem (oznaczymy je literami S i P). Obok nich mamy tu jeszcze do czynienia z relacją, która zachodzi pomiędzy studentem i przedmiotem – relacją lubienia (x lubi y). Relację tę oznaczymy przy pomocy predykatu L, po którym, w nawiasie, będą znajdowały się dwie



NIEKTÓRZY STUDENCI LUBIĄ NIEKTÓRE PRZEDMIOTY

zmienne, czyli $L(x,y)$. W rozpatrywanym zdaniu występuje również, dwukrotnie, zwrot odpowiadający kwantyfikatorowi szczegółowemu (*niektóre*).

Przystępując do pisania schematu powyższego zdania dobrze jest spróbować na początku wypowiedzieć je przy pomocy wyrażeń używanych w języku predykatów. Zdanie to mogłoby wyglądać na przykład następująco: *Istnieje pewien obiekt (oznaczymy go x), który ma własność bycia studentem; istnieje też inny „obekt” (oznaczymy go y), który jest przedmiotem i pomiędzy tymi obiektami zachodzi relacja lubienia.* Teraz powyższe zdanie możemy zapisać przy pomocy symboli:

$$\exists x [S(x) \wedge \exists y (P(y) \wedge L(x,y))]$$

Przykład:

Zapiszemy schemat zdania: *Każdy student przeczytał jakąś książkę.*

W zdaniu powyższym jest mowa o dwóch własnościach – bycia studentem (S) i bycia książką (K) oraz o relacji przeczytania (P) zachodzącej pomiędzy studentem a książką. Zdanie zaczyna się od zwrotu odpowiadającego kwantyfikatorowi ogólnemu, a więc nasz schemat będziemy musieli zacząć od $\forall x$. Zdanie mówi o każdym obiekcie będącym studentem, a więc $\forall x S(x)$. Po predykatcie musi nastąpić jakiś spójnik. Zgodnie z opisaną wcześniej nieformalną zasadą, gdy zdanie rozpoczyna się kwantyfikatorem ogólnym, to spójnikiem tym będzie zapewne implikacja. Mamy więc: $\forall x S(x) \rightarrow$, czyli *dla każdego x , jeśli jest on studentem* (lub prościej *dla każdego studenta*). Zdanie, którego schemat piszemy, mówi, że ów „każdy

student” przeczytał jakąś książkę. Nie możemy jednak na razie wstawić predykatu oznaczającego relację przeczytania – $P(x,y)$, gdyż występuje w nim zmienna y , o której nie wiemy, co miałyby oznaczać i która, co ważniejsze, nie jest związana żadnym kwantyfikatorem (a jak powiedzieliśmy, w prawidłowo napisanych schematach zdań języka naturalnego, zmienne wolne (nie związane) nie mogą występować). Gdybyśmy wstawili teraz predykat oznaczający relację przeczytania, otrzymalibyśmy $\forall x (S(x) \rightarrow P(x,y))$, czyli *każdy student przeczytał y* . Aby można było użyć predykatu $P(x,y)$ musimy najpierw umieścić w schemacie kwantyfikator wiążący zmienną y . Ponieważ w dalszej części zdania mowa jest o *jakiejś* książce, będzie to zapewne kwantyfikator szczegółowy. Mamy więc $\forall x S(x) \rightarrow \exists y$, czyli *dla każdego studenta istnieje jakiś y* . Teraz aż się prosi, żeby napisać czym jest ten y : $\forall x S(x) \rightarrow \exists y K(y)$ – *dla każdego studenta istnieje y będący książką*, czyli *dla każdego studenta istnieje jakaś książka*. Teraz musimy jedynie dodać, że jest to książka, którą ten student przeczytał, czyli zachodzi jeszcze pomiędzy studentem i książką relacja P : $\forall x S(x) \rightarrow \exists y K(y) \wedge P(x,y)$. Należy jeszcze oczywiście pamiętać o nawiasach, dzięki którym będziemy wiedzieli, że kwantyfikatory wiążą wszystkie „swoje” zmienne. Aby było to widoczne, po każdym kwantyfikatorze otwieramy nawias i zamykamy go na końcu schematu – dzięki temu wszystkie zmienne pozostaną związane:

$$\forall x [S(x) \rightarrow \exists y (K(y) \wedge P(x,y))]$$

Po napisaniu schematu dobrze jest go sobie „odczytać”, aby sprawdzić, czy faktycznie oddaje on treść zdania, które ma reprezentować. Nasz schemat mówi, że *dla każdego x , jeśli jest on studentem, istnieje jakiś y , który jest książką i ten x (student) przeczytał y (książkę)*. Mówiąc prościej: *dla każdego studenta istnieje książką, którą on przeczytał*, czyli dokładnie to, że *każdy student przeczytał jakąś książkę*.

Przykład:

Napiszemy schemat zdania: *Niektórzy wykładowcy lubią wszystkich studentów*.

W powyższym zdaniu mamy do czynienia z własnościami bycia wykładowcą i bycia studentem, oraz z relacją lubienia. Oznaczmy je kolejno predykatami W , S i L . Zdanie zaczyna się ewidentnie od kwantyfikatora szczegółowego $\exists x$. Oczywiście ten „istniejący x ” to wykładowca, czyli $\exists x W(x)$. Teraz musimy dopisać, że ów wykładowca lubi wszystkich studentów. Czyli, oprócz posiadania własności W , o naszym x możemy powiedzieć, że dla

każdego obiektu y , jeśli ten y posiada własność S , to pomiędzy x i y zachodzi relacja lubienia. Pamiętajmy oczywiście o nawiasach.

$$\exists x [W(x) \wedge \forall y (S(y) \rightarrow L(x,y))]$$

Przykład:

Zapiszemy schemat zdania: *Niektórzy studenci nie lubią żadnego wykładowcy.*

W powyższym zdaniu występują predykaty takie same jak w poprzednim przykładzie. Początek schematu będzie na pewno wyglądał $\exists x S(x)$. Problem sprawić może ustalenie, jak oddać w schemacie stwierdzenie, że ów obiekt posiadający cechę S nie lubi *żadnego* obiektu o cesze W . Podobnie, jak w jednym z pierwszych omawianych przykładów, słowo *żaden* możemy oddać na dwa równoważne sobie sposoby. Można stwierdzić, że nie istnieje obiekt y , taki że posiada cechę W i jednocześnie pomiędzy x i y zachodzi relacja L . Można też powiedzieć, że dla każdego obiektu, jeśli ma on cechę W , to pomiędzy x i y nie zachodzi L .

Dwa równoprawne schematy naszego zdania to:

$$\exists x [S(x) \wedge \sim \exists y (W(y) \wedge L(x,y))]$$

$$\exists x [S(x) \wedge \forall y (W(y) \rightarrow \sim L(x,y))]$$

Czy można być w relacji do siebie samego?

Pomimo że relacje (dwuczłonowe) z natury łączą dwa obiekty, to może się zdarzyć, że obiekty te są w rzeczywistości jednym i tym samym; mówiąc inaczej, jakiś obiekt może być w pewnej relacji do siebie samego.

Przykład:

Zapiszemy schemat zdania: *Pewien bokser znokautował siebie samego.*

W zdaniu powyższym jest mowa o relacji znokautowania ($Z(x,y)$ – x znokautował y). Stwierdza ono jednakże, że pewien obiekt posiadający własność bycia bokserem, jest w tej relacji do siebie samego. Schemat zdania, to zatem:

$$\exists x (B(x) \wedge Z(x,x))$$



PEWIEN BOKSER ZNOKAUTOWAŁ SIEBIE SAMEGO

Czy jest tu jakaś własność?

Czasem przy pisaniu schematu musimy uwzględnić własność, która nie jest w zdaniu wprost wypowiedziana.

Przykład:

Napiszemy schemat zdania: *Każdy kogoś kocha.*

Na pierwszy rzut oka wydaje się, że w zdaniu powyższym występuje jedynie relacja kochania, nie ma w nim mowy natomiast o żadnej własności. W takim wypadku schemat mógłby wyglądać: $\forall x \exists y K(x,y)$ – dla każdego obiektu x , istnieje obiekt y , taki, że x kocha y . Czasem faktycznie dopuszczalne jest napisanie takiego „skrótowego” schematu. Czy jednak w powyższym zdaniu faktycznie jest mowa o *dowolnych* obiektach x i y ? Słowa *każdy* i *kogoś* wyraźnie wskazują, że nie chodzi tu o wszelkie możliwe do pomyślenia obiekty, ale tylko i wyłącznie o ludzi. Mamy więc tu do czynienia z cechą bycia człowiekiem, która nie jest wprost wypowiedziana. Zdanie *Każdy kogoś kocha* należy traktować jako skrót zdania *Każdy człowiek kocha jakiegoś człowieka*. W wersji bardziej pomocnej do przełożenia na język rachunku predykatów można powiedzieć: *Dla każdego obiektu, jeśli obiekt ten jest człowiekiem, istnieje inny obiekt, który też jest człowiekiem, i ten pierwszy kocha tego drugiego*. A zatem:

$$\forall x [C(x) \rightarrow \exists y (C(y) \wedge K(x,y))]$$

Przykład:

Napiszemy schemat zdania: *Są tacy, którzy nie czytają żadnych gazet.*

W powyższym zdaniu, podobnie jak w poprzednim przykładzie, mamy „ukrytą” cechę bycia człowiekiem. Druga cecha, bycia gazetą, jest już jednak wprost wypowiedziana. Relację czytania oznaczmy przez R , ponieważ predykat C oznacza już bycie człowiekiem. Fakt, że żadna gazeta nie jest przez pewnych ludzi czytana, oddać można na dwa sposoby. A zatem dwa możliwe schematy tego zdania to:

$$\exists x [C(x) \wedge \sim \exists y (G(y) \wedge R(x,y))]$$

$$\exists x [C(x) \wedge \forall y (G(y) \rightarrow \sim R(x,y))]$$

I znowu „tylko”...

Zdaniami ze zwrotem *tylko* zajmowaliśmy się już, gdy były w nich obecne jedynie własności. Bardzo podobnie postępujemy pisząc schematy takich zdań, w których występują również relacje.

Przykład:

Napiszemy schemat zdania: *Niektóre partie wspierane są tylko przez frustratów.*

W zdaniu powyższym musimy użyć predykatów oznaczających własności bycia partią, bycia frustratem oraz relację bycia wspieranym przez kogoś (x jest wspierany przez y). Schemat oczywiście rozpoczniemy od zwrotu: $\exists x P(x)$. Jak pamiętamy, zwrot *tylko* możemy oddać przy pomocy kwantyfikatora ogólnego. Jednakże trzeba uważać w jakiej kolejności nastąpią człony implikacji w formule związanej przez ten kwantyfikator. Gdybyśmy napisali schemat następująco: $\exists x [P(x) \wedge \forall y (F(y) \rightarrow W(x,y))]$, to otrzymalibyśmy schemat zdania mówiącego, że *niektóre partie wspierane są przez wszystkich frustratów (każdy frustrat wspiera taką partię)*. Nie jest to więc dokładnie schemat naszego zdania. To, że partia wspierana jest *tylko* przez frustratów, nie oznacza, że wspiera ją *każdy* frustrat, ale to, że każdy kto ją wspiera, ten jest frustratem (jeśli ją wspiera to jest frustratem). A zatem w schemacie musimy zamienić kolejność predykatów F i W . Prawidłowy schemat to:

$$\exists x [P(x) \wedge \forall y (W(x,y) \rightarrow F(y))]$$

Co jest x, a co y?

Czasami musimy zwrócić baczną uwagę na właściwą kolejność zmiennych x i y przy predykanie oznaczającym relację.

Przykład:

Napiszemy schemat zdania: *Istnieją podręczniki, z których korzystają wszyscy studenci.*

Przyjmujemy predykaty P , S i K oznaczające własności bycia podręcznikiem i studentem oraz relację korzystania z czegoś. Schemat: $\exists x [P(x) \wedge \forall y (S(y) \rightarrow K(x,y))]$ nie jest jednak prawidłowy, ponieważ po jego odczytaniu otrzymalibyśmy zdanie mówiące, że istnieją podręczniki, które korzystają ze wszystkich studentów. Ponieważ własność bycia studentem przypisaliśmy zmiennej y , a bycia podręcznikiem, zmiennej x , to aby oddać prawidłowo fakt, że to student korzysta z podręcznika, a nie na odwrót, musimy napisać $K(y,x)$. A więc właściwy schemat naszego zdania to:

$$\exists x [P(x) \wedge \forall y (S(y) \rightarrow K(y,x))]$$

W wielu przypadkach to, w jakiej kolejności powinny znaleźć się zmienne x i y w relacji, uzależnione jest od tego, w jaki sposób określimy naszą relację.

Przykład:

Napiszemy schemat zdania: *Niektóre programy lubią wszyscy widzowie.*

W schemacie powyższego zdania musimy użyć predykatów oznaczających własności bycia programem i bycia widzem oraz relację lubienia. Relację tę jednak możemy zinterpretować albo jako relację lubienia – x lubi y , albo jako relację bycia lubianym – x jest lubiany przez y . W zależności od tej interpretacji prawidłowe byłyby schematy, kolejno:

$$\exists x [P(x) \wedge \forall y (W(y) \rightarrow L(y,x))]$$

(L oznacza relację lubienia)

$$\exists x [P(x) \wedge \forall y (W(y) \rightarrow L(x,y))]$$

(L oznacza relację bycia lubianym)

Dłuższe schematy.

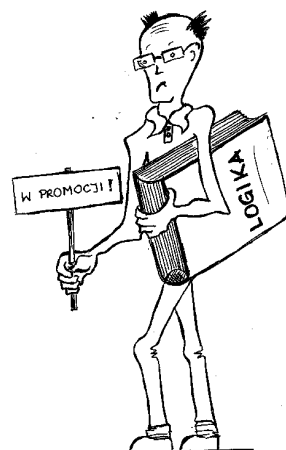
W schematach może pojawić się większa ilość kwantyfikatorów i predykatów.

Przykład:

Napiszemy schemat zdania: *Niektórzy filozofowie piszą niektóre książki, których nikt przy zdrowych zmysłach nie kupuje.*

Zdanie zaczyna się stwierdzeniem, że istnieje ktoś, kto jest filozofem. Dalej dowiadujemy się, że ów filozof pisze książki, czyli istnieje coś, co jest książką i ten filozof pozostaje do książki w relacji napisania. Następna informacja, to stwierdzenie, że nie ma nikogo, kto miałby cechę bycia przy zdrowych zmysłach i jednocześnie pozostawał w relacji kupowania do wymienionej wcześniej książki. Ten ostatni fakt możemy oddać na dwa sposoby; drugi sposób, to powiedzenie, że każdy, jeśli jest przy zdrowych zmysłach, to nie kupuje danej książki. A zatem:

$$\exists x \{F(x) \wedge \exists y [(K(y) \wedge P(x,y)) \wedge \sim \exists z (Z(z) \wedge R(z,y))]\}$$



NIKTÓRZY FILOZOFOWIE PISZĄ KSIĄŻKI,
KTÓRYCH NIKT PRZY ZDROWYCH ZMYŚŁACH NIE KUPUJE

$$\exists x \{F(x) \wedge \exists y [(K(y) \wedge P(x,y)) \wedge \forall z (Z(z) \rightarrow \sim R(z,y))]\}$$

Przy tego rodzaju dłuższych schematach należy zwracać szczególną uwagę na nawiasy (pamiętamy, aby wszystkie zmienne były związane przez kwantyfikatory) oraz o tym, aby przy własnościach i relacjach umieszczać właściwe zmienne. Przykładowo, gdy mamy na końcu napisać, że w pewnej relacji pozostaje ktoś przy zdrowych zmysłach oraz książka, to musimy sprawdzić, jakimi zmiennymi wcześniej oznaczyliśmy obiekty mające wymienione własności.



3.1.4. CZĘSTO ZADAWANE PYTANIA.

Czy błędem byłoby zapisanie schematu zdania w którym nie wszystkie własności lub relacje byłyby potraktowane osobno, na przykład napisanie schematu zdania: „Nie każdy znany muzyk jest artystą” jako $\sim \forall x (Z(x)) \rightarrow A(x)$ gdzie Z oznaczałby własność bycia znany muzykiem?

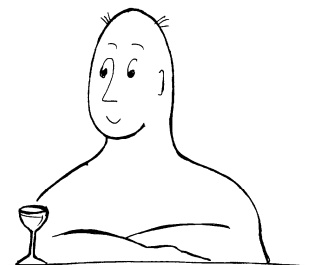
Nie jest to błąd w ścisłym tego słowa znaczeniu, jednakże tworząc schemat, należy zwykle pisać tak zwany schemat główny, możliwie najgłębiej wnikający w strukturę zdania, w którym obecne są wszystkie możliwe do wyodrębnienia predykaty i spójniki. Jednakże faktem jest, że nie zawsze do końca wiadomo, kiedy w zdaniu mamy do czynienia z dwiema osobnymi własnościami, a kiedy nie.

Kiedy możemy przyjąć domyślnie, że zmienne reprezentują jeden określony typ obiektów i nie podkreślać tego dodatkowo w schemacie, a kiedy musimy cechę bycia takim obiektem w schemacie umieścić? Przykładowo, kiedy pisząc schemat zdania „Każdy kogoś kocha”, powinniśmy uwzględnić w nich własność bycia człowiekiem i napisać $\forall x [C(x) \rightarrow \exists y (C(y) \wedge K(x,y))]$, a kiedy możemy przyjąć, że zmienne reprezentują tylko ludzi i napisać: $\forall x \exists y K(x,y)$?

Na powyższe pytanie nie ma jednoznacznej odpowiedzi. Rozwiązując tego typu przykłady najlepiej spytać wykładowcy, jakie odpowiedzi uznaje on za poprawne. Niektórzy mogą wymagać, na przykład, napisania obu wersji schematów.

3.2. DODATEK: STAŁE INDYWIDUOWE I ZNAK „=”

3.2.1. ŁYK TEORII.



ŁYK TEORII

Jak dotąd omawialiśmy rachunek predykatów w podstawowej, najbardziej ubogiej, wersji. W niektórych wypadkach wygodnie jest wzbogacić go o kilka dodatkowych elementów, które czasem mogą ułatwić zapisywanie schematów zdań.

Obecnie do słownika, z którego składa się język rachunku predykatów, dodamy dwa rodzaje elementów: tak zwane stałe indywiduowe, które będziemy oznaczać małymi literami: a, b, c, d, ...itd. oraz szczególny predykat oznaczający relację identyczności dwóch obiektów, czyli znany wszystkim z matematyki znak „=”. Gdy wprowadzimy znak równości, będziemy mogli również korzystać ze znaku „≠”, stwierdzającego nieidentyczność. Stanowiąc on będzie skrót wyrażenia *nieprawda, że obiekty są identyczne*, czyli $x \neq y \equiv \sim (x = y)$

Tak jak zmienne indywiduowe (x,y,z...) oznaczały dowolne obiekty, tak stałe indywiduowe (a,b,c...) oznaczają określone, konkretne obiekty. Stała może reprezentować np. Mikołaja Kopernika, Statuę Wolności, Kubusia Puchatka, Zenka, itp. Stałe wykorzystujemy w schematach, gdy zdanie mówi o takich właśnie, jednoznacznie określonych, obiektach.

Przykładowo zdanie Zenek jest starszy od Wacka możemy zapisać jako S (a,b), gdzie „a” oznacza Zenka, „b” – Wacka, a S reprezentuje relacje starszeństwa. Zasadniczą różnicę pomiędzy zmiennymi a stałymi stanowi to, że stałe nie mogą być wiązane przez kwantyfikatory. Nie wolno pisać np. $\exists a$ lub $\forall b$. W związku z powyższym, schematy, w których występują stałe indywiduowe, nie muszą rozpoczynać się od kwantyfikatora, choć oczywiście mogą – gdy oprócz stałych, w schemacie obecne są również zmienne.

Symbol identyczności przydaje się, gdy w zdaniu, którego schemat piszemy, mowa jest o pewnej określonej liczbie przedmiotów posiadających daną własność lub będących do czegoś w relacji, na przykład *Tylko jeden student oblał egzamin*, czy też *Przynajmniej dwóch posłów przyłapano na oszustwie*. Jak postępować w takich przypadkach pokażą przykłady poniżej.

Jeśli komuś pisanie schematów z wykorzystaniem stałych oraz, w szczególności, znaku „=” wyda się zbyt zagmatwane, a wykładowca nie wymaga od niego opanowania tej sztuki, może ten rozdział pominąć. Nie jest on konieczny do zrozumienia dalszej części, dotyczącej tautologii i reguł.

3.2.2. PRAKTYKA: BUDOWANIE SCHEMATÓW ZDAŃ Z WYKORZYSTANIEM STAŁYCH INDYWIDUOWYCH I SYMBOLU IDENTYCZNOŚCI.

Rozpoczniemy od zapisywania schematów zdań, w których wykorzystamy stałe indywiduowe.

Przykład:

Napiszemy schemat zdania: *Mieczysław kocha Karolinę, ale Karolina nie kocha Mieczysława.*

Zdanie powyższe stwierdza, że pomiędzy dwoma konkretnymi obiektami (Mieczysławem i Karoliną) zachodzi relacja kochania w jedną stronę, natomiast nie zachodzi ona w drugą. Oznaczając Mieczysława przez „a”, a Karolinę przez „b”, otrzymujemy schemat:

$$K(a,b) \wedge \sim K(b,a)$$

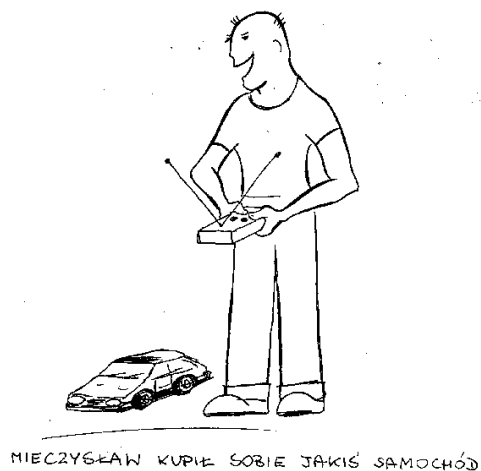
W schematach ze stałymi indywiduowymi mogą też pojawić się zmienne, a wraz z nimi kwantyfikatory.

Przykład:

Napiszemy schemat zdania: *Mieczysław kupił sobie jakiś samochód.*

Zdanie powyższe stwierdza, że istnieje pewna rzecz, mająca własność bycia samochodem i Mieczysław (oznaczony za pomocą stałej „a”) pozostaje do tej rzeczy w relacji kupienia.

$$\exists x (S(x) \wedge K(a,x))$$



Przykład:

Napiszemy schemat zdania: *Karolina lubi wszystkich bogatych mężczyzn.*

Powyższe zdanie stwierdza, że Karolina pozostaje w relacji lubienia do każdego, kto posiada dwie cechy – bycia mężczyzną i bycia bogatym. Mówiąc inaczej, jeśli ktoś posiada

wymienione własności, to Karolina pozostaje do niego w relacji lubienia. Oznaczając Karolinę przy pomocy stałej „a”, mamy schemat:

$$\forall x [(M(x) \wedge B(x)) \rightarrow L(a,x)]$$

Przykład:

Napiszemy schemat zdania: *Karolina lubi wszystkich, którzy lubią Mieczysława.*

Zdanie to stwierdza, że Karolina (którą oznaczymy przez „a”) pozostaje w relacji lubienia do wszystkich, którzy pozostają w tej samej relacji do Mieczysława (oznaczonego przez „b”). Mówiąc inaczej – dla każdego obiektu, jeśli obiekt ten znajduje się w relacji L do „b”, to „a” znajduje się do niego w L. Przyjmując domyślnie, że obiektami, o których jest mowa, są ludzie, mamy schemat:

$$\forall x (L(x,b) \rightarrow L(a,x))$$

Gdybyśmy chcieli wyraźnie zaznaczyć w schemacie, że w zdaniu chodzi o ludzi, otrzymalibyśmy schemat:

$$\forall x [(C(x) \wedge L(x,b)) \rightarrow L(a,x)]$$

Teraz zajmiemy się schematami zdań, w których będziemy musieli wykorzystać symbol identyczności.

Przykład:

Napiszemy schemat zdania: *Tylko jeden student zdał.*

Powyższe zdanie możemy rozbić na dwie części. Po pierwsze, mówi ono, że istnieje ktoś kto jest studentem i zdał, a po drugie, że nie ma innej osoby, która by miała te własności. Schemat pierwszej części jest oczywisty: $\exists x (S(x) \wedge Z(x))$. Część drugą można oddać na dwa sposoby. Można stwierdzić, że nie istnieje taki obiekt y, który byłby różny od x i posiadał te same własności lub też, że każdy obiekt, który te własności posiada, to właśnie x. A zatem: $\sim \exists y [(S(y) \wedge Z(y)) \wedge y \neq x]$ lub $\forall y [(S(y) \wedge Z(y)) \rightarrow y = x]$

Tak więc ostatecznie schemat naszego zdania może przedstawiać się następująco:

$$\exists x \{(S(x) \wedge Z(x)) \wedge \sim \exists y [(S(y) \wedge Z(y)) \wedge y \neq x]\} \text{ lub}$$

$$\exists x \{(S(x) \wedge Z(x)) \wedge \forall y [(S(y) \wedge Z(y)) \rightarrow y = x]\}$$

Przykład:

Zapiszemy schemat zdania: *Przynajmniej dwóch pasażerów było trzeźwych.*

Powyższe zdanie stwierdza, że istnieją na pewno dwa różne obiekty, które posiadają dwie cechy jednocześnie – bycia pasażerem i bycia trzeźwym. Można zatem powiedzieć, że istnieje jeden obiekt mający wymienione cechy, istnieje też drugi mający te cechy, przy czym obiekty te nie są ze sobą identyczne. A zatem:

$$\exists x \{ (P(x) \wedge T(x)) \wedge \exists y [(P(y) \wedge T(y)) \wedge x \neq y] \}$$

Uwaga na marginesie.

W powyższym schemacie jedyny spójnik, to koniunkcja, której człony możemy umieszczać w dowolnej kolejności. W związku z powyższym, dozwolone są inne warianty schematu; możemy z tego powodu również zrezygnować z niektórych nawiasów, zostawiając jedynie te, które wskazują na zasięg kwantyfikatorów. Na przykład:

$$\exists x \{ P(x) \wedge T(x) \wedge \exists y [P(y) \wedge T(y) \wedge x \neq y] \}$$

Możemy również rozpocząć schemat dwoma kwantyfikatorami, po których, w jednym nawiasie umieścimy (w dowolnej kolejności) wszystkie człony koniunkcji:

$$\exists x \exists y (P(x) \wedge T(x) \wedge P(y) \wedge T(y) \wedge x \neq y)$$

Tego typu uproszczenia można oczywiście stosować, ale lepiej tego nie robić jeśli nie ma się pewności, że jest to dozwolone.

Oczywiście stałe indywiduowe i symbol identyczności mogą występować jednocześnie w tym samym schemacie.

Przykład:

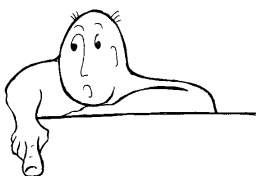
Napiszemy schemat zdania: *Nie tylko Zenek dotrwał do końca imprezy.*

Powyższe zdanie stwierdza, że po pierwsze, Zenek (którego oznaczmy przez „a”) posiada własność D (dotrwał do końca imprezy) i, po drugie jest jeszcze jakiś inny obiekt, różny od Zenka, który posiada wymienioną własność. A zatem otrzymujemy schemat:

$$D(a) \wedge \exists x (D(x) \wedge x \neq a)$$



NIE TYLKO ZENEK DOTRWAŁ DO KOŃCA IMPREZY



Uwaga na błędy!

Często się zdarza, że ktoś, pisząc schemat powyższego zdania, zapomina o jego pierwszej części. Jednakże schemat: $\exists x (D(x) \wedge x \neq a)$ nie byłby prawidłowy. Byłby to schemat zdania mówiącego, że jakaś osoba różna od Zenka dotrwała do końca imprezy, bez zaznaczenia, że Zenek również wykazał się taką umiejętnością.

Przykład:

Napiszemy schemat zdania: *Tylko jeden świadek rozpoznał „Marmoladę”*.

Oznaczmy przez $\acute{S}$ własność bycia świadkiem, przez R relacje rozpoznania, a przez stałą „a” obiekt zwany „Marmoladą”. Powyższe zdanie stwierdza, że istnieje pewien obiekt, mający własność $\acute{S}$, który znajduje się w relacji R do obiektu a , i nie ma jednocześnie nikogo innego (czyli obiektu różnego od x) mającego $\acute{S}$ i będącego w R do „a”. A zatem:

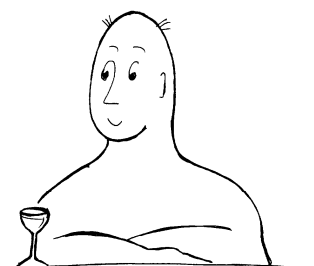
$$\exists x \{(\acute{S}(x) \wedge R(x,a)) \wedge \sim \exists y [(\acute{S}(y) \wedge R(y,a)) \wedge y \neq x]\}$$

Powyższe zdanie można również przedstawić:

$$\exists x \{(\acute{S}(x) \wedge R(x,a)) \wedge \forall y [(\acute{S}(y) \wedge R(y,a)) \rightarrow y = x]\}$$

3.3. TAUTOLOGIE I KONTRTAUTOLOGIE.

3.3.1. ŁYK TEORII.



ŁYK TEORII

W rachunku zdań mieliśmy do czynienia z prostą metodą zero-jedynkową, która pozwalała na szybkie, w zasadzie mechaniczne, stwierdzenie, czy dany schemat jest tautologią bądź kontrtautologią. W przypadku rachunku predykatów, niestety, nie ma takiej metody. Wykazanie tautologiczności lub kontrtautologiczności formuły wymaga dość zaawansowanych technik, wykraczających poza ramy niniejszego opracowania. O wiele prostsze jest zadanie odwrotne – udowadnianie, że dana formuła nie jest tautologią, lub nie jest kontrtautologią. I tylko tym – wykazywaniem, czym dany schemat **nie jest**, będziemy się dalej zajmować.

Zanim przejdziemy do tautologii i kontrtautologii musimy uświadomić sobie od czego zależy prawdziwość formuły rachunku predykatów. Rozpatrzmy bardzo prosty schemat: $\exists x P(x)$. Czy jest to schemat zdania prawdziwego czy fałszywego? To oczywiście zależy, przede wszystkim od tego, jaką własność podstawimy za predykat P . Podstawmy zatem za P własność bycia w wieku 200 lat ($P(x) - x$ ma 200 lat). Jeśli nasze rozważania ograniczymy do świata ludzi, to otrzymamy zdanie fałszywe – żaden człowiek nie ma bowiem dwustu lat. Jeśli jednak schemat odniesiemy, na przykład, do świata drzew, będziemy mieli do czynienia ze zdaniem prawdziwym – oczywiście istnieją drzewa mające dwieście lat. Prawdziwość naszej formuły zależy zatem od dziedziny, tak zwanego uniwersum, w którym ją umieścimy, oraz od interpretacji predykatu w tym świecie.

Układ złożony ze zbioru stanowiącego uniwersum (oznaczanego zwykle literą U) oraz dowolnej ilości własności i relacji będziemy określać mianem **struktury**. A zatem możemy powiedzieć, że prawdziwość formuły rachunku predykatów zależy od struktury, w której formułę tę będziemy rozpatrywać.

Strukturę oznaczać będziemy przy pomocy podkreślonej litery $\underline{U}$. Elementy struktury umieszczamy w nawiasach $\langle \rangle$. Obecne w strukturze własności i relacje, odpowiadające obecnym w formułach KRP predykatom będziemy oznaczać przy pomocy takich samych liter jak predykaty, jednakże podkreślonych. Na przykład podkreślone $\underline{R}$ będzie oznaczało konkretną relację w konkretnej strukturze, stanowiącą odpowiednik abstrakcyjnie pojętego predykatu R w formule.

Przykładowo struktury, o których była mowa wyżej, możemy zapisać następująco:

$$\underline{U}_1 = \langle U = \text{zbiór ludzi; } \underline{P}(x) \equiv x \text{ ma 200 lat} \rangle$$

$$\underline{U}_2 = \langle U = \text{zbiór drzew; } \underline{P}(x) \equiv x \text{ ma 200 lat} \rangle$$

Inne struktury, w których możemy rozpatrywać formułę $\exists x P(x)$, to na przykład:

$$\underline{U}_3 = \langle U = \text{zbiór ludzi; } \underline{P}(x) \equiv x \text{ jest studentem} \rangle$$

$$\underline{U}_4 = \langle U = \text{zbiór drzew; } \underline{P}(x) \equiv x \text{ jest studentem} \rangle,$$

W $\underline{U}_3$ nasza formuła reprezentować będzie zdanie prawdziwe, natomiast w $\underline{U}_4$ – fałszywe.

Strukturę, w której formuła rachunku predykatów jest prawdziwa, nazywamy **modelem** tej formuły, natomiast strukturę, w której jest fałszywa – **kontrmodelem**. Tak więc możemy powiedzieć, że dla formuły $\exists x P(x)$, $\underline{U}_2$ oraz $\underline{U}_3$ stanowią modele, natomiast $\underline{U}_1$ i $\underline{U}_4$ – kontrmodele.

Przejdźmy teraz do zdefiniowania pojęcia tautologii w rachunku predykatów. Jak pamiętamy z rachunku zdań, tautologia, to formuła, która jest zawsze prawdziwa. Skoro w rachunku predykatów prawdziwość formuły zależy od struktury, w jakiej formułę interpretujemy, możemy powiedzieć, iż tautologia KRP to formuła, która jest prawdziwa w każdej strukturze. Patrząc na to samo z drugiej strony możemy powiedzieć również, iż w przypadku tautologii nie istnieje struktura, w której formuła ta byłaby fałszywa. Mówiąc krótko, tautologia nie ma kontrmodelu.

Podobnie określić możemy kontrtautologię. Jest to formuła fałszywa w każdej strukturze. Mówiąc inaczej, nie istnieje struktura, w której formuła będąca kontrtautologią byłaby prawdziwa; kontrtautologia nie ma modelu.

3.3.2. PRAKTYKA: WYKAZYWANIE, ŻE FORMUŁA NIE JEST TAUTOLOGIĄ LUB KONTRTAUTOLOGIĄ.

Wykazanie, że dana formuła nie jest tautologią, teoretycznie jest bardzo proste. Skoro tautologia musi być prawdziwa w każdej strukturze, to aby udowodnić, że formuła tautologią nie jest, wystarczy wskazać strukturę, w której jest ona fałszywa (zbudować kontrmodel dla tej formuły). Analogicznie, aby wykazać, że formuła nie jest kontrtautologią, trzeba pokazać strukturę, w której jest ona prawdziwa (zbudować model formuły).

W praktyce trudność może czasem sprawić wymyślenie odpowiedniej struktury. Nie ma bowiem na to jakiejś jednej, sprawdzającej się zawsze, metody

Przykład:

Wykażemy, że formuła $\forall x (P(x) \rightarrow Q(x))$ nie jest tautologią ani kontrtautologią.

Najpierw zbudujemy kontrmodel formuły, a więc strukturę, w której jest ona fałszywa. W ten sposób wykazemy, że nie jest ona tautologią. Aby zbudować odpowiednią strukturę, zacząć musimy od odczytania tego, co mówi nasza formuła. Otóż stwierdza ona, że każdy obiekt, który ma własność P, ma również własność Q. Aby zbudować kontrmodel, musimy więc dobrać własności P i Q w taki sposób, aby w jakimś zbiorze nie było to prawdą. Weźmy przykładowo zbiór ludzi jako uniwersum i własność bycia kobietą jako odpowiednik predykatu P oraz bycia matką jako odpowiednik Q. Formalnie:

$$\underline{U}_1 = \langle U = \text{zbiór ludzi}; \underline{P}(x) \equiv x \text{ jest kobietą}, \underline{Q}(x) \equiv x \text{ jest matką} \rangle$$

W strukturze U_1 , nasza formuła stwierdza, że *dla każdego człowieka, jeśli człowiek ten jest kobietą, to jest również matką*, czyli w skrócie *każda kobieta jest matką*, co jest oczywiście zdaniem fałszywym. U_1 jest zatem kontrmodelem dla formuły $\forall x (P(x) \rightarrow Q(x))$

Aby zbudować model, musimy dobrać własności P i Q tak, aby otrzymać zdanie prawdziwe. W powyższym przykładzie możemy to łatwo uczynić zamieniając własności miejscami, czyli:

$$\underline{U}_2 = \langle U = \text{zbiór ludzi; } \underline{P}(x) \equiv x \text{ jest matką, } \underline{Q}(x) \equiv x \text{ jest kobietą} \rangle$$

W strukturze U_2 , nasza formuła stwierdza, że *każda matka jest kobietą*, co jest oczywiście zdaniem prawdziwym. U_2 jest zatem modelem dla formuły $\forall x (P(x) \rightarrow Q(x))$.

Skoro zbudowaliśmy dla formuły kontrmodel i model, oznacza to, że nie jest ona tautologią ani kontrtautologią.

Podane wyżej rozwiązanie jest oczywiście jednym z nieskończonej ilości właściwych odpowiedzi. Ktoś mógłby przykładowo zbudować takie struktury:

$$\underline{U}_3 = \langle U = \text{zbiór polityków; } \underline{P}(x) \equiv x \text{ jest posłem, } \underline{Q}(x) \equiv x \text{ jest uczciwy} \rangle, \text{ oraz}$$

$$\underline{U}_4 = \langle U = \text{zbiór liczb; } \underline{P}(x) \equiv x \text{ jest podzielne przez 4, } \underline{Q}(x) \equiv x \text{ jest parzyste} \rangle.$$

Struktura $\underline{U}_3$ stanowiłaby wtedy kontrmodel, gdyż umieszczona w niej formuła stwierdzałaby, że *każdy polityk, który jest posłem, jest uczciwy*, natomiast $\underline{U}_4$ byłaby modelem, ponieważ umieszczona w tej strukturze formuła głosiłaby, iż *każda liczba podzielna przez 4, jest liczbą parzystą*.

To, jaki model i kontrmodel zostanie stworzony, zależy tylko od wyobraźni budowniczego.

Przykład:

Wykażemy, że tautologią ani kontrtautologią nie jest formuła: $\forall x R(x,x)$

Formuła powyższa stwierdza, że każdy obiekt jest w pewnej relacji do samego siebie.

Jako kontrmodel dla naszej formuły posłużyć może struktura

$$\underline{U}_1 = \langle U = \text{zbiór ludzi, } \underline{R}(x,y) \equiv x \text{ jest starszy od } y \rangle$$

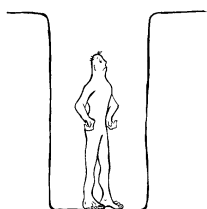
W $\underline{U}_1$ formuła reprezentowałaby fałszywe zdanie – *Każdy człowiek jest starszy do siebie samego*.

Jako model dla formuły wybierzemy strukturę

$$\underline{U}_2 = \langle U = \text{zbiór liczb, } \underline{R}(x,y) \equiv x \text{ jest równe } y \rangle$$

Umieszczając schemat w powyższej strukturze, otrzymujemy zdanie prawdziwe – *Każda liczba jest równa sobie samej*.

Ponieważ udało nam się znaleźć kontrmodel i model, wykazaliśmy, że badana formuła nie jest tautologią ani kontrtautologią.



3.3.3. UTRUDNIENIA I PUŁAPKI.

Największa trudność, jaka może powstać przy wykazywaniu, że schemat nie jest tautologią, ani kontrtautologią, wiąże się z prawidłową oceną, czy w strukturze, którą zbudowaliśmy, formuła jest prawdziwa, czy fałszywa, a więc to, co faktycznie zbudowaliśmy – model czy kontrmodel.

Aby nie popełnić przy tym błędu, kluczowa jest umiejętność właściwego odczytywania schematów w danej strukturze – stwierdzania, co mówi zdanie powstałe ze schematu przy zaproponowanej interpretacji predykatów i zmiennych.

Przykład:

Wykażemy, że nie jest tautologią, ani kontrtautologią formuła: $\forall x \forall y (R(x,y) \rightarrow R(y,x))$

Powyższy schemat stwierdza, że dla każdych dwóch obiektów, jeżeli jeden jest w relacji R do drugiego, to drugi jest w relacji R do pierwszego. Innymi słowy: dla dowolnych dwóch obiektów, jeśli R zachodzi pomiędzy nimi w jedną stronę, to zachodzi również w drugą.

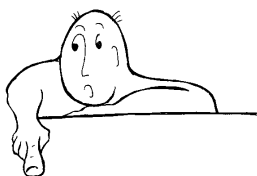
Za kontrmodel dla powyższej formuły może posłużyć struktura złożona ze zbioru ludzi i relacji kochania. Nie jest bowiem tak, że dla każdej pary ludzi, jeśli jedna osoba kocha drugą, to ta druga również kocha pierwszą.

Model stanowić może struktura, w której w zbiorze ludzi określimy relację bycia w tym samym wieku. Prawdą jest bowiem, że zawsze, jeśli jeden człowiek jest w tym samym wieku co drugi, to ten drugi jest w tym samym wieku co pierwszy. A zatem mamy:

$\underline{U}_1 = \langle U = \text{zbiór ludzi}, \underline{R}(x,y) \equiv x \text{ kocha } y \rangle$

$\underline{U}_2 = \langle U = \text{zbiór ludzi}, \underline{R}(x,y) \equiv x \text{ jest w tym samym wieku, co } y \rangle$

Ponieważ udało nam się znaleźć kontrmodel i model, wykazaliśmy, że badana formuła nie jest tautologią ani kontrtautologią.



Uwaga na błędy!

Ktoś mógłby błędnie sądzić, że w $\underline{U}_2$ formuła $\forall x \forall y (R(x,y) \rightarrow R(y,x))$ jest fałszywa, ponieważ „nie jest prawdą, że wszyscy ludzie są w tym samym wieku”. Trzeba jednak zauważyć, że wyrażenie w nawiasie nie mówi, że wszyscy są w danej relacji, ale że **jeśli** są w relacji w jedną stronę, **to** są i w drugą. Taka zależność zachodzi właśnie w przypadku relacji bycia w tym samym wieku.

Przykład:

Wykażemy, że nie jest tautologią ani kontrtautologią formuła: $\forall x [P(x) \rightarrow \exists y R(x,y)]$

Powyższy schemat stwierdza, że dla każdego obiektu jest tak, że jeśli posiada on własność P, to istnieje jakiś obiekt, że ten pierwszy jest w relacji R do tego drugiego.

Zdanie prawdziwe możemy z formuły tej otrzymać podstawiając w zbiorze ludzi za P własność bycia kobietą, a za R relację bycia czyjąś córką.

$\underline{U}_1 = \langle U = \text{zbiór ludzi}, P(x) \equiv x \text{ jest kobietą } R(x,y) \equiv x \text{ jest córką } y \rangle$

W $\underline{U}_1$ z naszej formuły powstaje prawdziwe zdanie: *Każda kobieta jest czyjąś córką*, a więc $\underline{U}_1$ jest dla tej formuły modelem.

Kontrmodel możemy zbudować podstawiając za R relację bycia żoną.

$\underline{U}_2 = \langle U = \text{zbiór ludzi}, P(x) \equiv x \text{ jest kobietą } R(x,y) \equiv x \text{ jest żoną } y \rangle$

W $\underline{U}_2$ otrzymujemy z naszej formuły zdanie fałszywe: *Każda kobieta jest czyjąś żoną*.

Ponieważ udało nam się znaleźć kontrmodel i model, wykazaliśmy, że badana formuła nie jest tautologią ani kontrtautologią.

W dotychczasowych przykładach, wszystkie formuły, dla których budowaliśmy modele i kontrmodele, były ostatecznie związane kwantyfikatorami; kwantyfikatory, od których zaczynała się formuła, miały zasięg do samego jej końca. Może się jednak zdarzyć, że formuła powstanie w wyniku powiązania jej części spójnikami logicznymi. W takich przypadkach do określenia, czy formuła reprezentuje w danej strukturze zdania prawdziwe, czy fałszywe, konieczna jest znajomość tabel zero-jedynkowych dla tych spójników.

Przykład:

Wykażemy, że nie jest tautologią, ani kontrtautologią formuła:
 $\forall x(P(x) \vee Q(x)) \rightarrow (\forall xP(x) \vee \forall xQ(x))$.

W powyższej formule należy koniecznie zauważyć, że jej głównym spójnikiem jest implikacja. Będzie to miało ogromne znaczenie dla określenia, czy pewna struktura jest jej modelem, czy kontrmodelem. Badany schemat możemy odczytać: *Jeśli każdy obiekt ma przynajmniej jedną z dwóch własności: P lub Q, to każdy obiekt ma P lub każdy obiekt ma Q*.

Na początek zajmiemy się poszukiwaniem kontrmodelu. Ponieważ formuła ma postać implikacji, to aby uzyskać z niej zdanie fałszywe, musimy tak dobrać własności, aby prawdziwy był poprzednik implikacji, a fałszywy jej następnik. Poprzednik mówi, że każdy obiekt ma własność P lub Q. Przykładowo, w zbiorze ludzi każdy człowiek ma własność bycia mężczyzną lub bycia kobietą. Zobaczmy, jaką wartość logiczną miałby w takiej strukturze następnik implikacji. Następnik ten mówi, że każdy obiekt ma własność P lub każdy ma własność Q. Przy zaproponowanej interpretacji predykatów, fałszem jest pierwszy człon alternatywy (bo nie jest prawdą, że każdy człowiek jest mężczyzną) i fałszem jest również drugi jej człon (bo nie jest prawdą, że każdy człowiek jest kobietą). Skoro oba człony alternatywy są fałszywe, to również, zgodnie z tabelkami zero-jedynkowymi, cała alternatywa jest fałszywa.

W strukturze:

$\underline{U}_1 = \langle U = \text{zbiór ludzi}; \underline{P}(x) \equiv x \text{ jest mężczyzną}, \underline{Q}(x) \equiv x \text{ jest kobietą} \rangle$

formuła $\forall x(P(x) \vee Q(x)) \rightarrow (\forall xP(x) \vee \forall xQ(x))$ jest zatem fałszywa. Fałszem jest zdanie: *Jeśli każdy człowiek jest mężczyzną lub kobietą, to każdy człowiek jest mężczyzną lub każdy człowiek jest kobietą*. Jest to zdanie fałszywe, gdyż ma ono postać implikacji, której poprzednik jest prawdziwy, a następnik fałszywy. Następnik jest fałszywy, gdyż jest on alternatywą, której każdy człon jest fałszywy.

Teraz musimy zbudować model dla naszej formuły. Ponieważ cała formuła ma postać implikacji, to, zgodnie z tabelkami zero-jedynkowymi może być ona prawdziwa na trzy sposoby. Pierwszy, gdy zarówno poprzednik, jak i następnik implikacji będą zdaniami prawdziwymi, drugi, gdy oba będą zdaniami fałszywymi, i trzeci, gdy poprzednik będzie fałszywy, a następnik prawdziwy. Z powyższej obserwacji można wysnuć bardzo pomocny wniosek: gdy sprawimy, że fałszywy będzie poprzednik implikacji, to bez względu na następnik, cała formuła stanie się schematem zdania prawdziwego.

Poprzednik naszej implikacji mówi, że każdy obiekt ma własność P lub Q. Aby otrzymać z tego zdanie fałszywe, możemy na przykład w zbiorze ludzi wstawić za P własność bycia nauczycielem, a za Q bycia studentem. Tworzymy więc strukturę:

$\underline{U}_2 = \langle U = \text{zbiór ludzi}; \underline{P}(x) \equiv x \text{ jest nauczycielem}, \underline{Q}(x) \equiv x \text{ jest studentem} \rangle$

$\underline{U}_2$ stanowi model dla naszej formuły. Umieszczona w nim, daje zdanie *Jeśli każdy człowiek jest nauczycielem lub studentem, to każdy człowiek jest nauczycielem lub każdy człowiek jest studentem*. Ponieważ zdanie to, mając postać implikacji, ma fałszywy poprzednik (*każdy człowiek jest nauczycielem lub studentem*) i fałszywy następnik (*każdy człowiek jest nauczycielem lub każdy człowiek jest studentem*), to jest to zdanie prawdziwe.

Ponieważ zbudowaliśmy kontrmodel i model dla naszej formuły, nie jest ona tautologią, ani kontrtautologią.

Oczywiście wcale nie musimy budować w przypadku formuły będącej implikacją modelu w powyższy sposób. Możemy spróbować stworzyć taki, w którym zarówno poprzednik implikacji, jak i jej następnik, byłyby zdaniami prawdziwymi. Jednakże nie zawsze jest to proste (na przykład w powyższym przykładzie). Przystąpienie do budowy modelu dla takiej formuły od próby uczynienia fałszywym poprzednika implikacji ułatwia nam pracę w ten sposób, że, bez względu na wartość logiczną następnika, otrzymamy w takiej strukturze zdanie prawdziwe. Na mocy tabelki zero-jedynkowych implikacja z fałszywym poprzednikiem jest bowiem zawsze prawdziwa.



DO ZAPAMIĘTANIA.

Niezwykle istotne jest odróżnienie, czy mamy do czynienia ze zdaniem, w którym główną rolę pełni kwantyfikator, czy też takim, w którym rola ta przypada spójnikowi logicznemu.

Jeśli wszystko związane jest kwantyfikatorem $\forall$ (np. $\forall x (P(x) \rightarrow Q(x))$), to odpowiedź, czy zdanie jest prawdziwe, czy fałszywe, uzależniona jest od tego, czy dana zależność zachodzi w stosunku do wszystkich obiektów. Jeśli jest to kwantyfikator $\exists$ (np. $\exists x (P(x) \vee Q(x))$), to wartość logiczna zdania zależy od tego, czy faktycznie istnieje dany obiekt.

Jeśli natomiast zdanie składa się z części powiązanych ostatecznie którymś ze spójników logicznych (np. $\forall x P(x) \rightarrow \exists x Q(x)$), to prawdziwość lub fałszywość takiego zdania oceniamy korzystając z tabelki zero-jedynkowych.

Przykład:

Wykażemy, że nie jest tautologią ani kontrtautologią formuła:

$$\forall x \exists y R(x,y) \rightarrow \exists y \forall x R(x,y)$$

Powyższa formuła ma postać implikacji. Zaczniemy od poszukiwania kontrmodelu, a więc takiej struktury, w której poprzednik implikacji stanie się zdaniem prawdziwym, a następnik fałszywym. Poprzednik stwierdza, że dla każdego obiektu istnieje jakiś obiekt, taki że ten pierwszy jest w relacji do drugiego. W zbiorze ludzi (nie tylko aktualnie żyjących!) zależność taka zachodzi w przypadku relacji bycia dzieckiem. Dla każdego człowieka istnieje jakiś człowiek, taki że ten pierwszy jest dzieckiem drugiego. Mówiąc po prostu, prawdą jest, że każdy jest czyimś dzieckiem. Zobaczmy teraz, co przy takiej interpretacji będzie mówił następnik naszej implikacji. Stwierdza on, że istnieje jakiś obiekt, taki że wszystkie inne są do niego w relacji. Czyli, istnieje człowiek taki, że wszyscy ludzie są jego dziećmi. Oczywiście jest to fałsz. W strukturze złożonej ze zbioru ludzi i relacji bycia dzieckiem otrzymamy zatem z naszej formuły fałszywe zdanie *Jeśli każdy jest czyimś dzieckiem, to istnieje ktoś, dla kogo wszyscy ludzie są jego dziećmi*. Jest to zdanie fałszywe, bo jego poprzednik jest prawdziwy, a następnik fałszywy. Mamy zatem kontrmodel:

$$\underline{U}_1 = \langle U = \text{zbiór ludzi}, R(x,y) \equiv x \text{ jest dzieckiem } y \rangle$$

Model w powyższym przypadku, podobnie jak w poprzednim przykładzie, najłatwiej będzie zbudować w taki sposób, aby uczynić fałszywym poprzednik naszej implikacji. Możemy to zrobić wstawiając na przykład za R relację bycia mężem.

$$\underline{U}_2 = \langle U = \text{zbiór ludzi}, R(x,y) \equiv x \text{ jest mężem } y \rangle$$

W $\underline{U}_2$ z naszej formuły otrzymamy zdanie: *Jeśli każdy jest czyimś mężem, to istnieje ktoś taki, że wszyscy są jego mężem*. Ponieważ poprzednik i następnik implikacji są tu fałszywe, całe zdanie jest prawdziwe. $\underline{U}_2$ stanowi zatem model dla naszej formuły.

Ponieważ zbudowaliśmy kontrmodel i model dla badanej formuły, nie jest ona tautologią, ani kontrtautologią.



3.3.4. CZĘSTO ZADAWANE PYTANIA.

Czy budując model i kontrmodel dla jednej formuły musimy korzystać z takiego samego uniwersum?

Nie jest to w żaden sposób konieczne. Może być na przykład tak, że uniwersum dla modelu stanowić będzie zbiór ludzi, a dla

kontrmodelu zbiór liczb. Rozwiązanie takie nie będzie w niczym gorsze od takiego, w którym uniwersum dla modelu i kontrmodelu byłoby takie same.

Czy jeśli nie mogę znaleźć dla jakiejś formuły kontrmodelu, to czy oznacza to, że formuła jest tautologią?

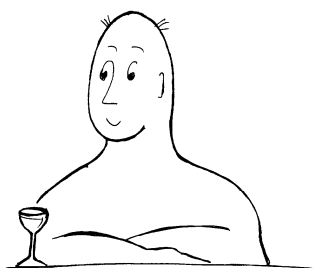
Fakt, że nie można znaleźć kontrmodelu, może być spowodowany tym, że formuła jest tautologią, jednak nie stanowi w żaden sposób na to dowodu. Być może kontrmodel istnieje, a my po prostu źle szukaliśmy. (Zobacz też odpowiedź na następne pytanie).

Czy budując model lub kontrmodel można wykazać, że formuła jest tautologią lub kontrtautologią?

Nie. Przy pomocy modeli i kontrmodeli możemy udowodnić jedynie rzecz „negatywną” – fakt, że formuła czymś **nie jest**. Wykazanie, że formuła **jest** tautologią, wymagałoby pokazania, że jest ona prawdziwa w każdej strukturze (każda struktura jest jej modelem). Z powodu nieskończonej ilości struktur, w jakich rozpatrywać można każdą formułę, nie jest to możliwe. Podobnie, wykazanie, że formuła jest kontrtautologią wymagałoby rozpatrzenia wszystkich struktur i pokazanie, że w każdej z nich jest ona fałszywa.

3.4. REGUŁY W RACHUNKU PREDYKATÓW.

3.4.1. ŁYK TEORII.



ŁYK TEORII

W sposób podobny do tego, w jaki wykazywaliśmy, że dana formuła nie jest tautologią lub kontrtautologią, można udowadniać zawodność reguł wnioskowania.

Jak pamiętamy z rachunku zdań, reguła jest to schemat wnioskowania – układ przynajmniej dwóch schematów, z których ostatni reprezentuje wniosek rozumowania, a poprzednie – przesłanki. Reguły będziemy zapisywać w ten sposób, że nad poziomą kreską będziemy umieszczać schematy przesłanek, natomiast pod kreską schemat wniosku.

Mówimy, że reguła jest dedukcyjna, a w związku z tym oparte na niej rozumowanie logicznie poprawne, jeśli nie jest możliwe, aby przesłanki stały się schematami zdań prawdziwych, a jednocześnie wniosek schematem zdania fałszywego.

Wykazanie, że dana reguła rachunku predykatów jest dedukcyjna, jest dość skomplikowane i, podobnie jak wykazywaniem, że formuła KRP jest tautologią bądź kontrtautologią, nie będziemy się tym obecnie zajmować. Ograniczymy się do, o wiele prostszego, udowadniania, że dana reguła nie jest dedukcyjna (czyli, mówiąc inaczej, jest zawodna).

Ponieważ to, czy formuły rachunku predykatów reprezentują zdania fałszywe czy prawdziwe, zależy od struktury, w której formuły te będziemy rozpatrywać, udowodnienie zawodności reguły polega na znalezieniu takiej struktury, w której wszystkie przesłanki staną się schematami zdań prawdziwych, a wniosek – schematem zdania fałszywego. W ten sposób wykazujemy, że możliwa jest sytuacja, aby przesłanki były prawdziwe, a wniosek fałszywy, a więc reguła jest zawodna – posługując się nią, możemy, wychodząc z prawdziwych przesłanek, dojść do fałszywego wniosku.

3.4.2. PRAKTYKA: WYKAZYWANIE ZAWODNOŚCI REGUŁ.

W praktyce, udowadnianie zawodności reguł przebiega tak samo, jak wykazywanie że formuła nie jest tautologią lub kontrtautologią.

Przykład:

Wykażemy, że zawodna jest reguła:

$$\sim \forall x P(x)$$

$$\forall x \sim P(x)$$

Jedyna przesłanka badanej reguły stwierdza, że nie każdy obiekt posiada własność P , natomiast jej wniosek głosi, iż żaden obiekt jej nie posiada. Zawodność powyższej reguły można wykazać budując strukturę $\underline{U} = \langle U = \text{zbiór ludzi}, \underline{P}(x) \equiv x \text{ jest Chińczykiem} \rangle$. W strukturze tej przesłanka stwierdza prawdziwie, iż *nie każdy człowiek jest Chińczykiem*, zaś wniosek, fałszywie, że *żaden człowiek Chińczykiem nie jest*.

Przykład:

Wykażemy, że zawodna jest reguła:

$$\exists x P(x), \exists x Q(x)$$

$$\forall x (P(x) \vee Q(x))$$

Pierwsza przesłanka reguły stwierdza, iż istnieje obiekt mający własność P, druga, że istnieje obiekt mający własność Q, natomiast wniosek, iż każdy obiekt ma przynajmniej jedną z tych własności. Zawodność reguły możemy wykazać budując strukturę:

$$\underline{U} = \langle U = \text{zbiór studentów}, \underline{P}(x) \equiv x \text{ ma 5 z logiki}, \underline{Q}(x) \equiv x \text{ ma 4 z logiki} \rangle$$

Przykład:

Wykażemy, że zawodna jest reguła:

$$\frac{\forall x \exists y R(x,y)}{\forall x \exists y R(y,x)}$$

Przesłanka powyższej reguły stwierdza, że każdy obiekt uniwersum pozostaje do czegoś w relacji R, natomiast wniosek, iż do każdego obiektu uniwersum coś pozostaje w R. Jako przykład struktury, w której przesłanka stanie się zdaniem prawdziwym, a wniosek fałszywym posłużyć może:

$$\underline{U}_1 = \langle U = \text{zbiór ludzi}, \underline{R}(x,y) \equiv x \text{ jest dzieckiem } y \rangle$$

Prawdą jest bowiem, że *każdy człowiek jest czyims dzieckiem*, fałszem natomiast, że *każdy człowiek dziecko posiada*.

SŁOWNICZEK

Kontrmodel – kontrmodelem formuły rachunku predykatów nazywamy strukturę, w której formuła ta jest fałszywa.

Kwantyfikator – wyrażenie określające ilość przedmiotów, o których mówi zdanie zawierające to wyrażenie. Kwantyfikatorami są wyrażenia *każdy* (oznaczany często symbolem $\forall$) oraz *niektóre (istnieje)* (oznaczany $\exists$).

Model – modelem formuły rachunku predykatów nazywamy strukturę, w której formuła ta jest prawdziwa.

Predykat – wyrażenie opisujące własność lub relację. Predykatami są na przykład takie wyrażenia jak *jest człowiekiem*, *jest wysoki* (własności), lub *kocha*, *jest wyższy od* (relacje).

Stała indywiduowa – symbol oznaczający pewien konkretny obiekt. Stałe indywiduowe oznaczamy zwykle literami a, b, c... itd. Nie podlegają one kwantyfikacji.

Struktura – układ złożony z pewnego uniwersum (zbioru) oraz dowolnej liczby własności i/lub relacji.

Zmienna indywiduowa – symbol oznaczający dowolny obiekt (indywiduum). Zmienne indywiduowe oznaczamy zwykle literami: x, y, z... itp. Można je wiązać kwantyfikatorami, np. $\forall x$, $\exists y$ itp.

Rozdział IV

NAZWY I DEFINICJE.

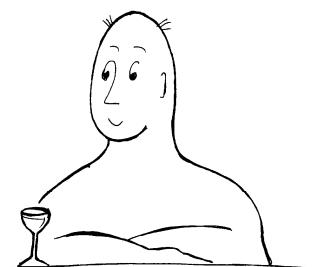
WSTĘP.

Obecny rozdział wiąże się z logiką rozumianą szerzej niż tylko jako nauka zajmująca się badaniem poprawności rozumowań. Poświęcony jest on problematyce zdecydowanie mniej skomplikowanej niż rachunek zdań, sylogistyka, czy też rachunek predykatów. Omówione są w nim kolejno: rodzaje nazw, zależności między nazwami, rodzaje definicji oraz niektóre błędy, jakie mogą w definicjach wystąpić.

Zadania, jakie pojawiają się w podręcznikach do logiki w związku z powyższą tematyką, są o wiele prostsze od zawartych w poprzednich rozdziałach. Dlatego też omówieniu ich rozwiązywania poświęcone zostało stosunkowo mało miejsca.

4.1. NAZWY I ICH RODZAJE.

4.1.1. ŁYK TEORII.

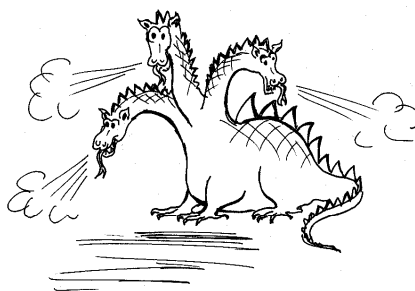


ŁYK TEORII

Nazwy są to wyrażenia służące do oznaczania przedmiotów. Nazwami są więc na przykład *człowiek*, *krzesło*, *książka* itp.

Rozważając problematykę nazw musimy pamiętać o dwóch ważnych sprawach. Po pierwsze, nazwa nie musi składać się z tylko jednego wyrazu. Nazwami są zatem takie złożone wyrażenia jak, przykładowo, *zły człowiek*, *drewniane krzesło z trzema nogami*, *niezwykle interesująca książka, którą przeczytałem w zeszłym tygodniu* itp. Każde z powyższych wyrażeń wskazuje nam pewien przedmiot, jest więc nazwą.

Drugą istotną sprawą, o jakiej należy pamiętać, gdy mówimy o nazwach, jest fakt, że owe „przedmioty” oznaczane przez nazwy musimy rozumieć bardzo szeroko, nie tylko jako obiekty materialne. Nazwy mogą bowiem odnosić się również, na przykład, do uczuć, pewnych procesów zachodzących w czasie, a także obiektów, które w ogóle nie istnieją w żaden sposób. Nazwami są więc również takie wyrażenia jak *miłość*, *śmiech*, *wykład z logiki*, *trójgłowy smok*, *niebieski krasnoludek* a nawet *żonaty kawaler*, czy też *kwadratowe koło*.



TRÓJGŁOWY SMOK

W obecnym rozdziale posługiwać się będziemy często dwoma pojęciami poznanymi w paragrafach poświęconych sylogizmom: **desygnat** nazwy oraz **zakres** (inaczej: **denotacja**) nazwy. Przypomnijmy, że desygnat jest to obiekt oznaczany przez daną nazwę (na przykład to, co trzymasz teraz przed sobą Czytelniku, jest desygnatem nazwy *książka*), natomiast zakres nazwy jest to zbiór jej wszystkich desygnatów (przykładowo zbiór wszystkich książek stanowi zakres nazwy *książka*). Zakres (denotację) nazwy A symbolicznie będziemy oznaczać $D(A)$.

Obecnie różnego rodzaju nazwy przedstawimy w sposób bardziej systematyczny. Podzielimy je na cztery różne sposoby.

1. Podział ze względu na ilość desygnatów.

Ze względu na ilość desygnatów nazwy podzielić możemy na trzy grupy:

a) Nazwy puste.

Nazwa pusta, to nazwa nie mająca ani jednego desygnatu. Nazwami pustymi są więc na przykład takie wyrażenia jak: *krasnoludek*, *dwustupiętrowy wieżowiec w Warszawie*, *uczciwy złodziej* itp.



KRASNOLUDEK

b) Nazwy jednostkowe.

Są to nazwy mające dokładnie jeden desygnat, na przykład: *Pałac Kultury i Nauki w Warszawie, Mieszko I, najdłuższa rzeka w Polsce* itp.

c) Nazwy ogólne.

Są to nazwy mające więcej niż jeden desygnat, przykładowo: *książka, poseł na sejm, medalista olimpijski* itp.

2. Podział ze względu na sposób istnienia desygnatów.

a) Nazwy konkretne.

Są to nazwy, których desygnaty są przedmiotami materialnymi (zajmują miejsce w przestrzeni, można je zobaczyć, dotknąć, zmierzyć itp.), lub byłyby takimi, gdyby istniały. W powyższym określeniu nazw konkretnych szczególnie istotny jest zwrot: „*lub byłyby takimi, gdyby istniały* [desygnaty]”. Tak więc oprócz takich wyrażen jak: *książka, człowiek, Adam Mickiewicz*, do nazw konkretnych zaliczamy również na przykład wyrażenia: *Smok Wawelski, uczciwy i inteligentny polityk, człowiek o wzroście 3 m, jednorożec* itp. Przedmioty oznaczane przez te nazwy wyobrażamy sobie bowiem jako obiekty materialne i gdyby istniały, to takimi by właśnie były.

b) Nazwy abstrakcyjne.

Do grupy tej zaliczamy wszystkie nazwy nie będące konkretnymi. A więc nazwy uczuć, relacji, własności, zdarzeń, procesów itp. Do grona nazw abstrakcyjnych zaliczamy również nazwy liczb i figur geometrycznych. Abstrakcyjnymi są więc takie nazwy jak: *miłość, podobieństwo, uczciwość, hałas, polityka, mecz piłkarski*, a także *liczba parzysta, trzynastcie, trójkąt*.

3. Podział ze względu na sposób wskazywania desygnatów.

a) Nazwy indywidualne.

Do grona nazwa indywidualnych zaliczamy imiona własne: nazwiska, nazwy geograficzne, nazwy statków itp., a także nazwy utworzone niejako przez „wskazanie palcem”, na przykład *ten oto człowiek*. Nazwy indywidualne przyporządkowane są danemu przedmiotowi na mocy arbitralnej decyzji, niezależnie od przysługujących temu przedmiotowi cech. Nazwami indywidualnymi są na przykład: *Adam Mickiewicz, Giewont, Warszawa, ta książka, którą trzymam w ręce* itp.

b) Nazwy generalne.

Są to nazwy, które przysługują przedmiotom ze względu na jakieś cechy, które tym przedmiotom przypisujemy. Nazwy generalne to na przykład: *poeta romantyczny*, *szczyt w Tatrach*, *stolica Polski*, a także *naukowiec*, *samochód*, *miasto* itp.

Nazwy indywidualne i generalne rozróżnić można jeszcze w jeden sposób. Otóż nazwy generalne w zdaniach podmiotowo-orzecznikowych typu *A jest B* nadają się zarówno na podmiot, jak i na orzecznik, a więc mogą wystąpić tak w miejscu zmiennej A, jak i B. Natomiast nazwy indywidualne nadają się jedynie na podmiot takich zdań. Możemy na przykład powiedzieć *Kraków* (nazwa indywidualna) *jest miastem nad Wisłą* (nazwa generalna), natomiast *miasto nad Wisłą jest Krakowem*, już nie.

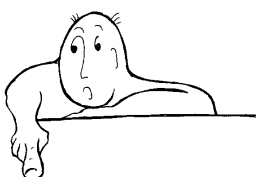
4. Podział ze względu na jednoznaczność (ostrość) zakresu.

a) Nazwy ostre.

Są to nazwy, w przypadku których da się jednoznacznie określić ich zakres, a więc oddzielić ich desygnaty od przedmiotów nimi nie będących. Nazwy ostre to na przykład: *tautologia KRZ*, *minister rządu RP*, *napój o zawartości alkoholu powyżej 4,5%*.

b) Nazwy nieostre.

W przypadku nazw nieostrych nie istnieje jednoznaczna, obiektywna granica oddzielająca przedmioty będące ich desygnatami od przedmiotów desygnatami takimi nie będących. Mówiąc inaczej, oprócz obiektów na pewno pod daną nazwę podpadających (desygnatów) oraz niewątpliwie niepodpadających (nie-desygnatów) istnieją też i takie, co do których nie bardzo wiadomo, do której grupy je zaliczyć. Nazwami nieostrych są na przykład: *piękna kobieta*, *ciekawa książka*, *geniusz*, *nudny wykładowca*, *tłum*, *pornografia*.



Uwaga na błędy!

Odróżniając nazwy ostre od nieostrych należy pamiętać, iż fakt, że ja osobiście nie wiem, czy jakiś przedmiot jest czy też nie jest desygnatem danej nazwy, nie powoduje jeszcze, że dana nazwa jest nieostra. Przykładowo, widząc idącego ulicą człowieka, nie wiem, czy jest on studentem, czy też nie jest. Jednakże nazwa *student* jest ostra, ponieważ, to, czy dany osobnik jest jej desygnatem, da się obiektywnie i ściśle ustalić, gdyby zaszła taka potrzeba. Inaczej będzie w przypadku nazwy, na przykład, *pijak* – tu na pewno znajdą się takie osoby, co do których nie będzie się dało w żaden obiektywny sposób stwierdzić, do której grupy należą: desygnatów, czy też nie-desygnatów. Pomiędzy zbiorem pijaków i nie-pijaków nie istnieje ostra i jednoznaczna granica.

4.1.2. PRAKTYKA: KLASYFIKOWANIE NAZW.

Zadania związane z klasyfikacją nazw są niezwykle proste. Polegają one na zaliczeniu danej nazwy do odpowiedniego członu każdego podziału.

Przykład:

Skłasyfikujemy kilka nazw:

a) *Student.*

Jest to nazwa ogólna (istnieje więcej niż jeden student), konkretna (desygnaty nazwy są obiektami fizycznymi), generalna (nazwa podaje pewną cechę desygnatu) i ostra (istnieje jednoznaczna granica oddzielająca studentów i nie-studentów).

b) *Obecna stolica Polski.*

Nazwa jednostkowa (jest tylko jedna obecna stolica Polski), konkretna (jest to „obiekt” fizyczny), generalna (podajemy pewną cechę desygnatu; gdyby chodziło o nazwę *Warszawa*, byłaby to nazwa indywidualna) i ostra.

c) *Wielka miłość.*

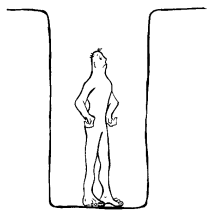
Nazwa ogólna, abstrakcyjna, generalna i nieostra (nie istnieje ścisła granica oddzielająca to, co jest wielką miłością, od tego, co nią nie jest).

W przypadku nazwy *wielka miłość*, podobnie jak i w związku z innymi nazwami abstrakcyjnymi, mogą powstać wątpliwości odnośnie ilości desygnatów. Kłopot polega na tym, że gdy desygnaty nazwy nie są obiektami materialnymi i nie można ich fizycznie „zobaczyć” trudno jest czasem powiedzieć, ile tych desygnatów faktycznie jest. I tak, na przykład, pesymista mógłby powiedzieć, że nazwa *wielka miłość* jest pusta, niektórzy filozofowie stwierdziliby, że jest to nazwa jednostkowa (bo istnieje tylko jedna idea Wielkiej Miłości), zaś ktoś jeszcze inny powiedziałby że jest to na pewno nazwa ogólna (bo sam przeżywa kolejną wielką miłość średnio co miesiąc).

W związku z tym, że logika nie dostarcza jednoznacznego rozwiązania tego typu problemów, może się zdarzyć, że różne odpowiedzi w tego typu zadaniach zostaną uznane za prawidłowe przez różne osoby.

d) *Obecny król Polski.*

Jest to nazwa pusta (przynajmniej w roku 2002 Polska nie ma króla), konkretna (bo gdyby król istniał, bo byłby zapewne człowiekiem, a więc obiektem materialnym), generalna i ostra.



4.1.3. UTRUDNIENIA I PUŁAPKI.

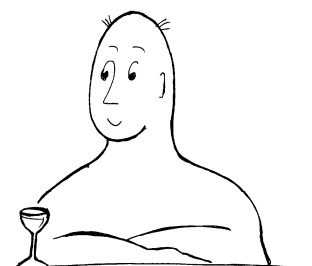
W przypadku klasyfikacji nazw trudno mówić o jakichkolwiek większych utrudnieniach lub pułapkach. W zasadzie jedyne poważne błędy, jakie można popełnić przy tego typu zadaniach, wynikają z niedokładnego zrozumienia lub zapamiętania charakterystyki różnych rodzajów nazw.

Najczęściej mylone bywają nazwy puste z abstrakcyjnymi, jednostkowe z indywidualnymi oraz ogólne z generalnymi. Dlatego zrozumieniu tych właśnie pojęć oraz różnic między nimi należy poświęcić szczególną uwagę.

Pewną trudność w klasyfikacji nazw sprawić może również fakt, że niektóre nazwy są ze swej natury wieloznaczne, jak na przykład *zamek*, które to wyrażenie może oznaczać zarówno budowlę, jak i zamek w drzwiach. Przed przystąpieniem do klasyfikacji takiej nazwy należy oczywiście najpierw ustalić o jakie znaczenie chodzi nam w danym wypadku, gdyż wzięta w różnych znaczeniach ta sama nazwa może mieć różne własności. Przykładowo nazwa *Mars* może być jednostkowa w znaczeniu planety, pusta w znaczeniu mitologicznego boga wojny, a ogólna w znaczeniu popularnego batonika. Należy też pamiętać, aby wieloznaczności nazwy nie mylić z jej nieostrością.

4.2. STOSUNKI MIĘDZY NAZWAMI.

4.2.1. ŁYK TEORII.



ŁYK TEORII

Dowolne dwie nazwy mogą znajdować się względem siebie w różnych zależnościach wynikających z ich zakresów (denotacji).

Ponieważ zakres nazwy jest to zbiór jej desygnatów, do omówienia stosunków zakresowych między nazwami konieczne jest przyswojenie sobie elementarnych wiadomości dotyczących zbiorów.

Gdy weźmiemy dwa dowolne zbiory X i Y , to mogą one pozostawać w następujących zależnościach.

$X = Y$ (**zbiór X jest równy zbiorowi Y**) – oznacza to, że zbiory X i Y mają dokładnie te same elementy. Na przykład: X – zbiór liczb parzystych, Y – zbiór liczb podzielnych przez 2.

$X \subset Y$ (**zbiór X zawiera się w zbiorze Y**) – oznacza to, że każdy element zbioru X jest również elementem zbioru Y , ale nie odwrotnie. Na przykład: X – zbiór wielbłądów, Y – zbiór ssaków.

$X \cap Y = \emptyset$ (**zbiór X jest rozłączny ze zbiorem Y**) – zbiory X i Y nie mają żadnego wspólnego elementu. Na przykład: X – zbiór ludzi, Y – zbiór samochodów.

$X \cap Y \neq \emptyset$ (**zbiór X krzyżuje się ze zbiorem Y**) – oznacza to, że zbiory X i Y mają jakieś elementy wspólne, ale oprócz tego każdy ma też takie, które nie są elementami drugiego zbioru. Na przykład: X – zbiór studentów, Y – zbiór osób palących; istnieją bowiem elementy wspólne – palący studenci, ale też elementy znajdujące się tylko w X – studenci niepalący, oraz elementy należące tylko do Y – osoby palące nie będące studentami.

Zależności między nazwami to nic innego, jak stosunki zachodzące między ich zakresami. Mogą być one następujące:

$D(A) = D(B)$ – mówimy wtedy, że **nazwy A i B są równoważne**. Na przykład: A – *Wisła*, B – *najdłuższa rzeka w Polsce* lub A – C_2H_5OH , B – *alkohol etylowy*.

$D(A) \subset D(B)$ – mówimy wtedy, że **nazwa A jest podrzędna względem nazwy B**, lub, jak kto woli, że **nazwa B jest nadrzędna względem A**. Na przykład: A – *dziesięć*, B – *ptak* lub A – *zdolny student*, B – *student*.

$D(A) \cap D(B) = \emptyset$ – mówimy, że **nazwy A i B się wykluczają**. Na przykład: A – *słoń*, B – *mrówka* lub A – *człowiek uczciwy*, B – *złodziej*.

$D(A) \# D(B)$ – mówimy, że **nazwy A i B się krzyżują** (lub że są niezależne). Na przykład: A – *człowiek bogaty*, B – *człowiek inteligentny* lub A – *blondynka*, B – *studentka*.

Uwaga na marginesie.

Pełna ścisłość nakazywałaby mówić o zależnościach między zakresami nazw, a nie samymi nazwami, a więc np.: *zakres nazwy A jest podrzędny wobec zakresu nazwy B*, czy też *zakres nazwy A wyklucza się z zakresem nazwy B*, jednak zwykle, dla uproszczenia, mówi się po prostu o stosunkach między nazwami.

4.2.2. PRAKTYKA: SPRAWDZANIE ZALEŻNOŚCI MIĘDZY NAZWAMI.

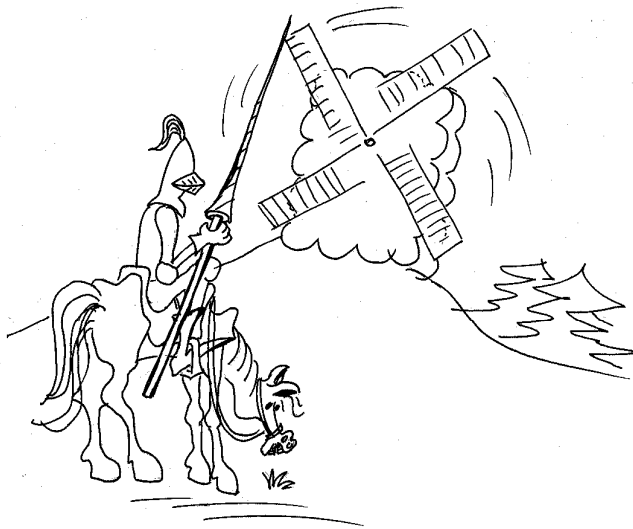
Jeden z typów zadań związanych ze stosunkami między nazwami polegać może na zbadaniu zależności pomiędzy dwiema podanymi nazwami.

W wielu prostych przypadkach zadania takie można rozwiązać bez uciekania się do jakichkolwiek wyrafinowanych sposobów. W przypadku niewielkich wątpliwości można spróbować określić zależności między nazwami drogą eliminacji. Przykładowa procedura będzie wtedy wyglądać następująco. (1) Najpierw oceniamy, czy nazwy mają takie same zakresy, co zwykle widać już na pierwszy rzut oka. Jeśli nie (a więc nie są równoważne), (2) patrzymy, czy w ogóle mają jakiekolwiek wspólne desygnaty. Jeśli nie mają, to znaczy się one wykluczają, jeśli mają, musimy szukać dalej – w takiej sytuacji (3) zadajemy sobie pytanie czy może każdy desygnat nazwy A jest desygnałem nazwy B, lub może, odwrotnie, każdy desygnat B jest desygnałem A. Jeśli tak, to znaczy że jedna nazwa (ta, której zakres zawiera się w zakresie drugiej) jest podrzędna względem drugiej. Jeśli nie, pozostaje nam ostatnia możliwość, a zatem (4) nazwy muszą się krzyżować.

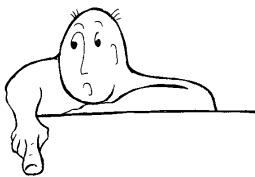
Przykład:

Zbadamy zależność między nazwami A – *piernik* B – *wiatrak*.

Jako że na pierwszy rzut oka widać, że nazwy *piernik* i *wiatrak* nie są równoważne, na początek pytamy więc, czy mają one jakiekolwiek wspólne desygnaty, a więc czy istnieje coś, co byłoby jednocześnie piernikiem i wiatrakiem. Ponieważ oczywiście nie ma takiej rzeczy, możemy zakończyć zadanie odpowiedzią, że badane nazwy się wykluczają.



CO MA PIERNIK DO WIATRAKA ???



Uwaga na błędy!

Należy pamiętać, że pytając o to, czy nazwy mają wspólne desygnaty, nie chodzi nam o to, czy istnieje jakaś cecha łącząca obiekty wskazywane przez badane nazwy, a więc na przykład czy istnieje piernik zrobiony z mąki wyprodukowanej w wiatraku, czy też piernik w kształcie wiatraka, albo wiatrak w kolorze piernika. Pytając o wspólne desygnaty pytamy, czy istnieje coś, co byłoby jednocześnie i jednym i drugim, a więc, w naszym przykładzie, coś będącego zarazem piernikiem i wiatrakiem.

Przykład:

Zbadamy zależności między nazwami A – *karp*, B – *ryba*.

Ponieważ widać, że nie są to nazwy równoważne, ale jakieś desygnaty wspólne posiadają, patrzmy, czy może zakres jednej z nazw zawiera się w zakresie drugiej. Oczywiście każdy karp jest rybą, czyli $D(A) \subset D(B)$. Tak więc nazwa *karp* jest podrzędna względem nazwy *ryba* (lub *ryba* nadrzędna względem *karp*).

Przykład:

Zbadamy zależności między nazwami A – *poseł na sejm*, B – *ograniczony nacjonalista*.

Po odrzuceniu pierwszej i drugiej możliwości, sprawdzamy, czy może jest tak, że każdy poseł na sejm jest ograniczonym nacjonalistą lub każdy ograniczony nacjonalista posłem. Ponieważ tak nie jest, wynika z tego, że badane nazwy muszą się krzyżować.

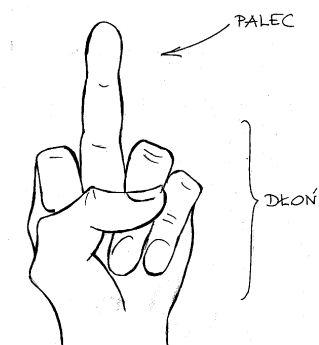
Przykład:

Zbadamy zależności między nazwami A – *palec* B – *dłoń*.

Na pierwszy rzut oka mogłoby się wydawać, że nazwy te mają coś wspólnego. W pewnym sensie jest to racja, jednak tym co je łączy, nie są na pewno wspólne desygnaty.

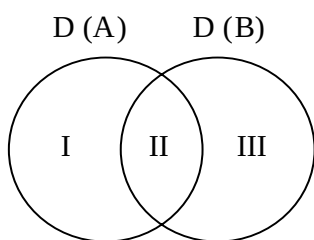
Wprawdzie palec jest częścią dłoni – nie oznacza to jednak, że istnieje taki palec, który byłby jednocześnie dłonią lub dłoń będącą palcem. Pamiętać należy, że sprawdzając zależności między nazwami pytamy, czy istnieją obiekty będące desygnatami jednej i drugiej nazwy, a nie czy istnieją pewne cechy łączące te nazwy lub ich desygnaty.

Nazwy *palec* i *dłoń* wykluczają się więc wzajemnie, podobnie jak rozpatrywane wyżej *piernik* i *wiatrak*.

**4.2.3. PRAKTYKA: ZASTOSOWANIE DIAGRAMÓW VENNA.**

Zależność między dwiema nazwami nie zawsze da się odkryć w tak prosty sposób, jak w powyższych przykładach. W niektórych przypadkach, szczególnie gdy mamy do czynienia z nazwami złożonymi, dobrze jest się posłużyć bardziej wyrafinowanym sposobem – metodą diagramów Venna. Diagramy te omawiane były już przy okazji sprawdzania poprawności sylogizmów. Obecnie ich wykorzystanie będzie na pewno o wiele prostsze.

Badanie zależności między dwiema nazwami przy pomocy diagramów Venna rozpoczynamy od narysowania dwóch kół reprezentujących zakresy rozważanych nazw:



Jak widać, diagram taki składa się z trzech obszarów. W obszary te będziemy musieli wpisać znaki „+” lub „-” w zależności od tego, czy coś się w nich znajduje, czy też są one puste.

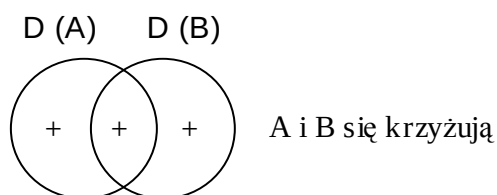
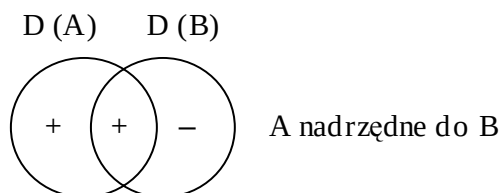
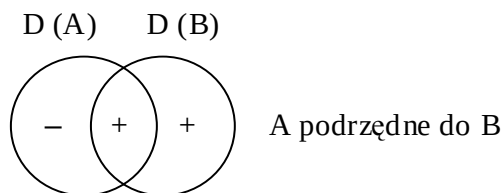
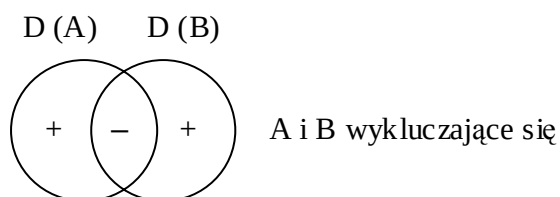
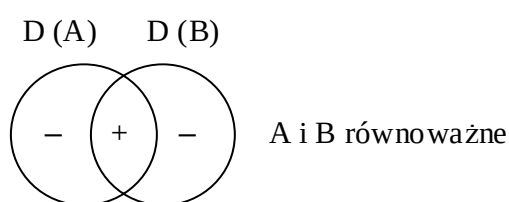
To, czy w danych obszarach diagramu znajdują się jakieś elementy odkrywamy odpowiadając na trzy proste pytania:

I – czy istnieje A, które nie jest B?

II – czy istnieje A, które jest B?

III – czy istnieje B, które nie jest A?

Przy założeniu, że żadna z nazw nie jest nazwą pustą, możemy otrzymać jeden z następujących rysunków świadczących o zależnościach między badanymi nazwami.





WARTO ZAPAMIĘTAĆ!

Gdyby ktoś miał problemy z zapamiętaniem, który rysunek świadczy o nadrzędności nazwy A względem B, a który o podrzędności, może to sobie utrwalić przy pomocy prostego skojarzenia. Gdy mamy rysunek ze znakiem „+” z jednej strony, a „-” z drugiej, to nadrzędna jest ta nazwa, przy której znajduje się „+”, a podrzędna ta, gdzie mamy „-”.

Powyższe rysunki ilustrują zależności pomiędzy nazwami przy założeniu, że żadna nazwa nie jest pusta. Nazwy puste rzadko bywają wykorzystywane w tego typu zadaniach. Dla porządku jednak dodajmy, że każda nazwa pusta jest podrzędna względem dowolnej nazwy niepustej, natomiast dwie nazwy puste są sobie zawsze równoważne.

Przykład:

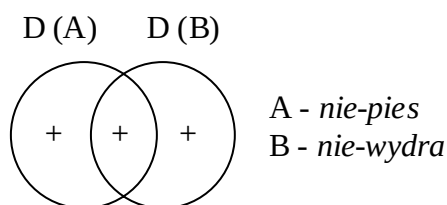
Zbadamy zależności między nazwami A – *nie-pies*, B – *nie-wydra*.

Po narysowaniu diagramu, w którym jedno koło symbolizuje zakres nazwy *nie-pies*, a więc zbiór wszystkich obiektów nie będących psami, natomiast drugie zakres nazwy *nie-wydra* (zbiór wszystkich nie-wydr), zadajemy trzy pytania:

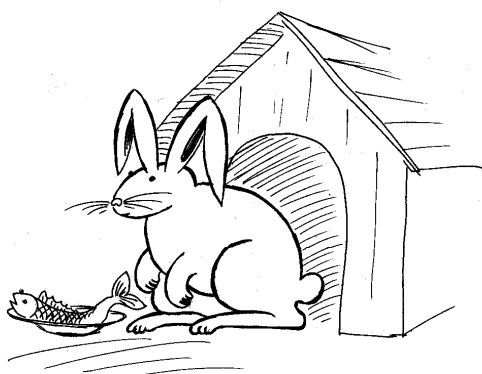
I – czy istnieje *nie-pies*, który nie jest *nie-wydrą*? Pytanie to początkowo wydaje się dość zagmatwane, możemy je jednak znacznie uprościć, korzystając z prawa mówiącego, że dwa przeczenia się znoszą. Tak więc, jeśli coś nie jest *nie-wydrą*, oznacza to, iż jest to po prostu wydrą. W ostatecznej, uproszczonej wersji nasze pytanie brzmi zatem: czy istnieje *nie-pies*, który jest wydrą? Oczywiście istnieje coś takiego i jest to po prostu wydra. W odpowiednim polu diagramu wpisujemy zatem znak „+”.

II – czy istnieje *nie-pies*, który jest *nie-wydrą*? Mówiąc inaczej, czy istnieje coś, co nie jest psem i jednocześnie nie jest wydrą. Oczywiście istnieje bardzo wiele takich rzeczy, na przykład może być to zając, tak więc w środkowym obszarze diagramu wpisujemy znak „+”.

III – czy istnieje *nie-wydra*, która nie jest *nie-psem*? Po uproszczeniu tego pytania w taki sam sposób jak w przypadku pytania I otrzymujemy: czy istnieje *nie-wydra*, która jest psem. Oczywiście istnieje coś takiego – jest to pies. W ostatnią część diagramu również wpisujemy zatem „+”.



Otrzymany rysunek świadczy, iż nazwy *nie-pies* i *nie-wydra* się krzyżują.



NIE PIES – NIE WYDRA

Przykład:

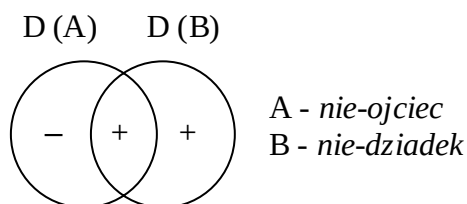
Zbadamy zależności między nazwami A – *nie-ojciec*, B – *nie-dziadek*.

Pytania konieczne do wypełnienia diagramu przedstawiają się następująco:

I – czy istnieje *nie-ojciec*, który nie jest *nie-dziadkiem*, a więc: czy istnieje *nie-ojciec*, który jest dziadkiem? Takiej osoby nie ma, ponieważ jeśli ktoś nie jest ojcem, nie może w żaden sposób zostać dziadkiem. W pierwszej części diagramu wpisujemy zatem znak „–”.

II – czy istnieje *nie-ojciec*, który jest *nie-dziadkiem*? Taka osoba istnieje, na przykład mężczyzna nie mający dzieci. W środkowej części diagramu wpisujemy znak „+”.

III – czy istnieje *nie-dziadek*, który nie jest *nie-ojcem*, a więc: czy istnieje *nie-dziadek*, który jest ojcem? Taka osoba istnieje – jest to mężczyzna mający dzieci, ale nie mający wnuków. W ostatnie pole diagramu wpisujemy „+”.



Otrzymany rysunek wskazuje, że nazwa *nie-ojciec* jest podrzędna względem nazwy *nie-dziadek* lub, jak kto woli, nazwa *nie-dziadek* jest nadrzędna do *nie-ojciec*.

4.2.4. PRAKTYKA: DOBIERANIE INNYCH NAZW DO NAZWY PODANEJ.

Inny rodzaj zadań związanych z zależnościami pomiędzy nazwami polegać może na poszukiwaniu nazwy podrzędnej, nadrzędnej, wykluczającej się i krzyżującej do podanej nazwy A (nazwy równoważnej często nie sposób podać, więc nie będziemy jej szukać w zadaniach).

W przypadku takich zadań nie istnieje ścisła metoda ich rozwiązywania; zwykle nie mają też one jednej odpowiedzi – niemal wszystko zależy tu od inwencji rozwiązującego.

Przykład:

Dobierzemy nazwę nadrzędną, podrzędną, wykluczającą się i krzyżującą w stosunku do nazwy A – *słoń*.

Nazwa nadrzędna do A to posiadająca szerszy zakres niż nazwa A. W przypadku słonia może więc być to na przykład *ssak* (każdy słoń jest ssakiem, ale nie na odwrót).

Nazwa podrzędna do A to taka, która posiada węższy zakres. Najprostszym sposobem utworzenia nazwy podrzędnej jest zwykle dodanie do nazwy wyjściowej jakiegoś przymiotnika zawężającego jej zakres – w naszym przypadku może być to na przykład *słoń afrykański* (każdy słoń afrykański jest słoniem, ale nie na odwrót).

Utworzenie nazwy wykluczającej się z A nie sprawi na pewno żadnego kłopotu – przykładowo może być to *mysz*. Nazwę wykluczającą można też zawsze utworzyć przez zaprzeczenie nazwy A – na przykład *nie-słoń*.

Najtrudniejsze może być początkowo utworzenie nazwy krzyżującej się z podaną. Musimy znaleźć taką nazwę B, żeby miała wspólne desygnaty z A, ale żeby również istniały A nie będące B oraz B nie będące A. W naszym przypadku musi być to takie B, że niektóre słonie tym są, ale też takie, że niektóre słonie owym B nie są, oraz niektóre B nie są słoniami. Nazwą spełniającą takie warunki jest na przykład *zwierzę żyjące w Afryce*. Są bowiem oczywiście słonie żyjące w Afryce, ale są też słonie mieszkające gdzie indziej (np. w Indiach), a także zwierzęta żyjące w Afryce, nie będące słoniami.

Mamy więc:

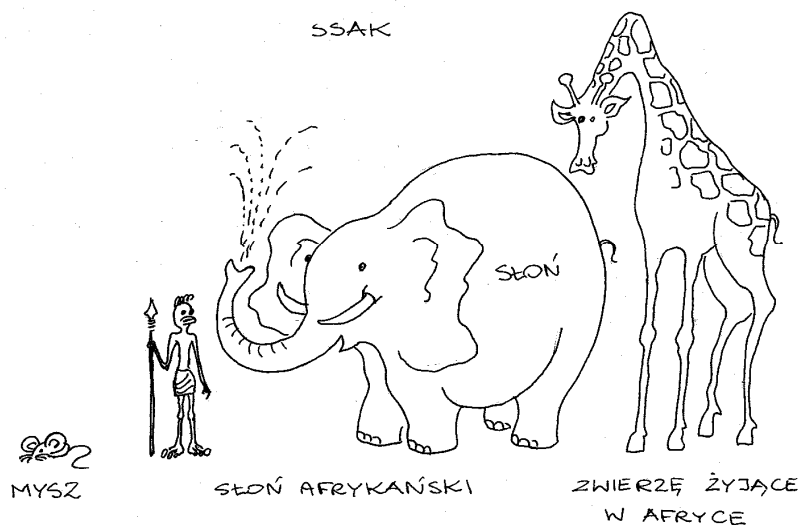
A – *słoń*

nadrzędna do A – *ssak*

podrzędna do A – *słoń afrykański*

wykluczająca się z A – *mysz*

krzyżująca się z A – *zwierzę żyjące w Afryce*



WARTO ZAPAMIĘTAĆ!

Istnieje prosty nieformalny sposób pozwalający niemal automatycznie stworzyć nazwę krzyżującą się z dowolną podaną nazwą. Aby utworzyć nazwę krzyżującą się z A należy:

- 1) Wziąć nazwę nadrzędną do A.
(Na przykład zwierzę do nazwy *słoń*)
- 2) Do nazwy tej dodać przymiotnik oznaczający cechę, którą niektóre (ale nie wszystkie!) desygnaty A posiadają.
(Niektóre (choć nie wszystkie) słonie żyją w Afryce, więc cechę tę dodaliśmy do nazwy *zwierzę*)

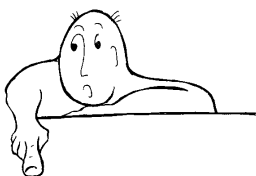
Otrzymamy zapewne nazwę krzyżującą się z A. W razie wątpliwości można to sprawdzić przy pomocy diagramów Venna.

Przykład:

Dobierzemy nazwę nadrzędną, podrzędną, wykluczającą się i krzyżującą z nazwą A – *nieuczciwy polityk*.

Nazwą o szerszym zakresie do A, a więc do niej nadrzędną będzie na pewno *polityk*.

Tworząc nazwę podrzędną do A możemy dodać do A jakąś zawężającą cechę – na przykład *amerykański nieuczciwy polityk*.



Uwaga na błędy!

Tworząc nazwę podrzędną do A poprzez dodanie przymiotnika zawężającego zakres, musimy dodać ten przymiotnik do całej nazwy A, a więc na przykład do *nieuczciwy polityk*, a nie tylko do samego *polityk*. W przeciwnym razie dostaniemy zapewne nazwę krzyżującą się zamiast podrzędnej.

Jako przykład nazwy wykluczającej się z A posłużyć może *uczciwy polityk*.

Nazwę krzyżującą się spróbujemy utworzyć w sposób podany wyżej. Weźmiemy więc nazwę nadrzędną do A, na przykład *człowiek* i dodamy do niej cechę, jaką zapewne niektórzy nieuczciwi politycy posiadają, na przykład wiek powyżej 40 lat. Otrzymujemy zatem nazwę *człowiek mający ponad 40 lat*. Innymi nazwami krzyżującymi się utworzonymi w ten sposób mogłyby być: *polityk angielski* lub *człowiek noszący okulary*.

Mamy więc:

A – *nieuczciwy polityk*

nadrzędna do A – *polityk*

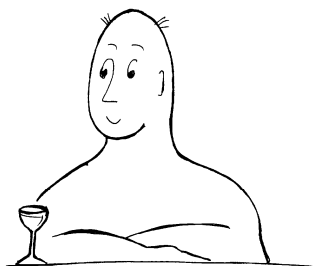
podrzędna do A – *amerykański nieuczciwy polityk*

wykluczająca się z A – *uczciwy polityk*

krzyżująca się z A – *człowiek mający ponad 40 lat*

4.3. DEFINICJE.

4.3.1. ŁYK TEORII.



ŁYK TEORII

Definicja to wyrażenie podające informacje o znaczeniu jakiegoś słowa lub zwrotu. Najczęściej spotykane są tak zwane definicje równościowe (nazywane również normalnymi). Definicja taka składa się z trzech części: **terminu definiowanego** (tak zwanego **definiendum**), **terminu definiującego** (tak zwanego **definiensa**) oraz zwrotu łączącego te dwa terminy – **łącznika definicyjnego**.

Jako przykład definicji równościowej może posłużyć wyrażenie: *Zegar jest to urządzenie do pomiaru upływu czasu*. Nazwa *zegar* jest tu terminem definiowanym, *urządzenie do pomiaru upływu czasu* – terminem definiującym, natomiast zwrot *jest to* – łącznikiem definicyjnym.

W skrócie możemy powiedzieć, że definicja normalna przyjmuje postać $A = B$, gdzie A i B są nazwami.

Rodzaje definicji ze względu na ich zadania.

Ze względu na to, jaki cel przyświecał autorowi tworzącemu daną definicję, możemy wyróżnić trzy rodzaje definicji:

a) Sprawozdawcze (analityczne).

Zadaniem takiej definicji jest wierne oddanie znaczenia terminu definiowanego, tak jak funkcjonuje ono w danym języku. Definicja taka stanowi „sprawozdanie” z ogólnie przyjętej treści danego terminu. Ogromną ilość definicji sprawozdawczych znaleźć można w dowolnym słowniku języka polskiego. Definicją taką jest również podane wyżej określenie słowa *zegar*.

b) Regulujące.

Zadaniem definicji regulującej jest precyzacja jakiegoś terminu nieostrego. Konieczność zastosowania takich definicji występuje najczęściej w prawodawstwie. Przykładowo w celu umożliwienia wpisywania do dowodów osobistych w rubryce „wzrost” słów: *niski, średni, wysoki*, konieczne stało się podanie definicji regulujących znaczenie tych nieostrych terminów. Tak powstać mogła definicja: *Przez wysokiego mężczyznę rozumieć będziemy mężczyznę mierzącego ponad 175 cm wzrostu*. Podobny rodowód może posiadać definicja – *Człowiek pełnoletni to osoba, która ukończyła osiemnasty rok życia*.

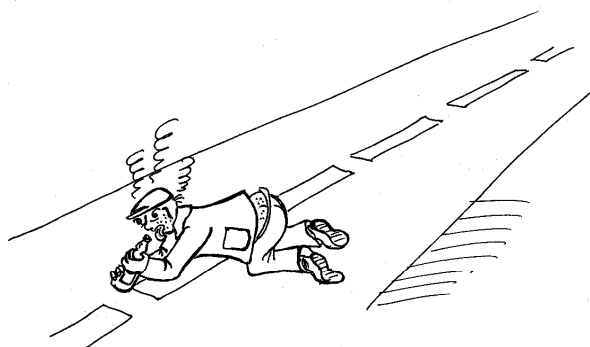
Czasem, gdy przyjęte w definicji regulującej znaczenie danego terminu staje się powszechne, definicja taka może przekształcić się w sprawozdawczą.

c) Konstrukcyjne (arbitralne).

Zadaniem takiej definicji jest wprowadzenie do języka nowego terminu lub nadanie już istniejącemu nowej treści, ignorującej dotychczasową. Definicje takie występują najczęściej w nauce, na przykład gdy wynalazca nadaje nazwę zbudowanemu przez siebie urządzeniu i określa, co należy pod tą nazwą rozumieć. Z czasem utworzone w ten sposób definicje konstrukcyjne, podobnie jak regulujące, mogą stać się sprawozdawczymi.

Definicje konstrukcyjne występują również na początku różnego rodzaju zbiorów przepisów lub zawieranych umów i określają, co dane słowa będą oznaczać w dalszym ciągu tekstu. Na przykład: *Pieszy – osoba, znajdująca się poza pojazdem na drodze i nie wykonująca na niej robót lub czynności przewidzianych odrębnymi przepisami*, lub: *Wartość polisy jest to wartość obliczana jako suma wartości jednostek funduszy przypisanych do danego rachunku po zarachowaniu z tytułu składki regularnej oraz dokonaniu stosownych*

odliczeń i potrąceń, gdzie środki zgromadzone w danym funduszu ustala się jako iloczyn liczby jednostek tego funduszu zarachowanych z tytułu składki regularnej znajdujących się na odpowiednim rachunku oraz wartości jednostki tego funduszu.



PIESZY – OSOBA, ZNAJDUJĄCA SIĘ
POZA POJAZDEM NA DRODZE ...

Warunki poprawności definicji sprawozdawczych.

Obecnie zajmujemy się warunkami poprawności definicji oraz tym, jak tę poprawność zbadać. Przedstawione niżej warunki odnoszą się zasadniczo do definicji sprawozdawczych. Definicje regulujące oraz arbitralne (jak już sama nazwa wskazuje) mogą być tworzone w sposób bardziej dowolny i nie podlegają tak ścisłym rygorom jak definicje sprawozdawcze, których zadaniem jest wierne oddanie znaczenia definiowanego terminu.

Jak już powiedzieliśmy definicja o normalnej (równościowej) budowie składa się z dwóch nazw (definiendum i definiens) połączonych spójnikiem definicyjnym; w skrócie: $A = B$. Ponieważ definicja sprawozdawcza ma na celu ścisłe oddanie znaczenia terminu definiowanego przy pomocy terminu definiującego, to aby można było uznać ją za w pełni poprawną, zakresy tych terminów powinny się pokrywać. Innymi słowy, w poprawnej definicji sprawozdawczej definiendum i definiens powinny być nazwami równoważnymi. Każdy inny stosunek zakresowy pomiędzy tymi terminami to błąd definicji. Błędy te charakteryzujemy następująco:

W definicji sprawozdawczej typu $A = B$:

Gdy definiendum (A) jest nadrzędne do definiensa (B), to mówimy, że definicja jest **za wąska**;

Gdy definiendum (A) jest podrzędne do definiensa (B), to mówimy, że definicja jest **za szeroka**;

Gdy definiendum (A) krzyżuje się z definiensem (B), to mówimy, że definicja obarczona jest **błędem krzyżowania zakresów**;

Gdy definiendum (A) wyklucza się z definiensem (B), to mówimy, że definicja obarczona jest **błędem wykluczania zakresów**.

W praktyce najczęściej występują w definicjach pierwsze dwa błędy (definicja za szeroka lub za wąska); natomiast ostatni z błędów (wykluczania zakresów) nie występuje prawie nigdy (poza specjalnie w tym celu spreparowanymi przykładami w podręcznikach do logiki).

4.3.2. PRAKTYKA: BADANIE POPRAWNOŚCI DEFINICJI SPRAWOZDAWCZYCH.

Sprawdzanie poprawności definicji sprawozdawczych jest niezwykle proste. Sprowadza się ono do określenia co stanowi definiendum oraz definiens, a następnie zbadania stosunków między nimi.

Przykład:

Sprawdzimy poprawność definicji: *Termometr jest to przyrząd do mierzenia.*

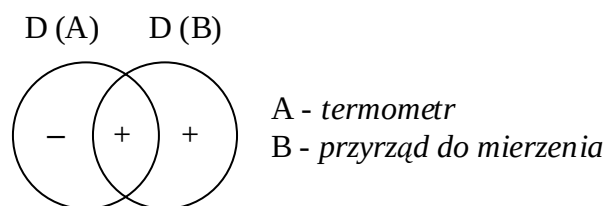
W definicji tej termin definiowany (definiendum) stanowi nazwa *termometr*, natomiast termin definiujący (definiens) – *przyrząd do mierzenia*.

Po narysowaniu diagramu możemy zadać trzy pytania, na które odpowiedzi są oczywiste:

I – czy istnieje termometr, który nie jest przyrządem do mierzenia – nie,

II – czy istnieje termometr, który jest przyrządem do mierzenia – tak,

III – czy istnieje przyrząd do mierzenia, który nie jest termometrem – tak (np. linijka).



Otrzymany rysunek wskazuje, że definiendum jest podrzędne względem definiensa, a zatem badana definicja jest za szeroka.

To, że badana definicja jest za szeroka widać w zasadzie już na pierwszy rzut oka – zbyt szeroko definiuje ona termometr.

Przykład:

Zbadamy poprawność definicji: *Termometr jest to przyrząd do mierzenia temperatury ludzkiego ciała.*

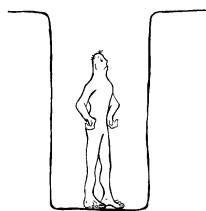
Odpowiedzi na odpowiednio zadane pytania są następujące:

I – tak (np. termometr okienny),

II – tak,

III – nie.

Wypełniony zgodnie z tymi odpowiedziami diagram wskazuje na nadrzędność definiendum względem definiensa, a więc badana definicja jest za wąska.

**4.3.3. UTRUDNIENIA I PUŁAPKI.**

Trudno mówić o jakichkolwiek pułapkach przy tak prostych zadaniach, jak sprawdzanie definicji sprawozdawczych. Jedyne kłopot może tu polegać na konieczności wykorzystania czasem wiedzy pozalogicznej potrzebnej do odpowiedzi na pytanie: czy istnieje pewna rzecz A będąca (lub nie będąca) B. Wiedza ta może czasem dotyczyć dziedzin specjalistycznych, obcych osobie badającej poprawność definicji.

SŁOWNICZEK.

Definiendum (termin definiowany) – termin, którego znaczenie podaje definicja.

Definiens (termin definiujący) – człon definicji wyjaśniający znaczenie terminu definiowanego.

Denotacja nazwy (zakres nazwy) – zbiór wszystkich desygnatów danej nazwy. Przykładowo zbiór wszystkich studentów jest denotacją (zakresem) nazwy student.

Desygnat nazwy – obiekt oznaczany przez daną nazwę. Na przykład każdy z nas jest desygnatem nazwy człowiek.

Łącznik definicyjny – zwrot łączący definiendum i definiens. Na przykład: *jest to, znaczy tyle co* itp.

Nazwa abstrakcyjna – nazwa, której desygnaty nie są przedmiotami materialnymi. Na przykład: *nienawiść, śmiech, egzamin*.

Nazwa generalna – nazwa, która przysługuje przedmiotowi ze względu na jakieś cechy, które temu przedmiotowi przypisujemy. Na przykład: *poeta romantyczny, miasto nad Wisłą, student*.

Nazwa indywidualna – nazwa przyporządkowana danemu przedmiotowi na mocy arbitralnej decyzji, niezależnie od przysługujących temu przedmiotowi cech. Na przykład: *Adam Mickiewicz, Kraków, ta oto książka*.

Nazwa jednostkowa – nazwa mające dokładnie jeden desygnat. Na przykład: *Pałac Kultury i Nauki w Warszawie, najwyższy szczyt w Tatrach*.

Nazwa konkretna – nazwa, której desygnaty są przedmiotami materialnymi lub byłyby takimi, gdyby istniały. Na przykład: *książka, krasnoludek*.

Nazwa nieostra – nazwa, której zakresu nie da się jednoznacznie i obiektywnie wyznaczyć. Na przykład: *wysoki mężczyzna, długie przemówienie, tłum*.

Nazwa ogólna – nazwa mająca więcej niż jeden desygnat. Na przykład: *człowiek, samochód*.

Nazwa ostra – nazwa, której zakres da się jednoznacznie określić. Na przykład: *medalista olimpijski, liczba parzysta, student*.

Nazwa pusta – nazwa nie mająca ani jednego desygnatu. Na przykład: *jednorożec, człowiek o wzroście 3 m*.

Rozdział V

ZBIORY.

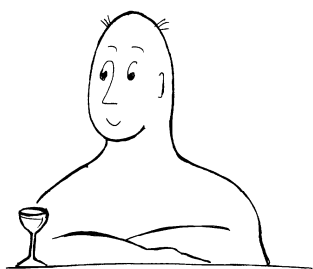
WSTĘP.

Obecny rozdział wraz z kolejnym – poświęconym relacjom, pełnią rolę w pewnym sensie pomocniczą. Omawiane w nich problemy nie dotyczą bezpośrednio logiki w jej tradycyjnym rozumieniu, jako dziedziny zajmującej się badaniem poprawności wnioskowań. Ponieważ jednak w XX wieku logika została silnie związana z matematyką, takie dziedziny jak teoria zbiorów i relacji uważane są współcześnie za jej pełnoprawne działy.

Ze zbiorami i relacjami spotkaliśmy się już we wcześniejszych rozdziałach. Obecnie pojęcia te zostaną omówione w sposób bardziej ścisły i systematyczny. Będzie się to wiązało, niestety, z większą ilością koniecznej to opanowania teorii. Jednakże, jak zwykle, największy nacisk położony zostanie na rozwiązywanie typowych zadań, spotykanych w podręcznikach do logiki w częściach poświęconych zbiorom i relacjom.

5.1. PODSTAWOWE WIADOMOŚCI O ZBIORACH.

5.1.1. ŁYK TEORII.



ŁYK TEORII

Zbiór to pewna kolekcja obiektów. Mówimy, na przykład, o zbiorze znaczków pocztowych, zbiorze liczb nieparzystych, zbiorze nudnych książek, zbiorze studentów itp., itd. Zbiory oznaczamy najczęściej dużymi literami, na przykład X , Y , Z lub A , B , C , D itd. Jeśli wypisujemy elementy jakiegoś zbioru, to zwykle umieszczamy je w nawiasach klamrowych, oddzielając od siebie przecinkami, na przykład: $\{a, b, c, d\}$. W zbiorze nie jest istotna kolejność, w jakiej elementy zostały przedstawione. Na przykład poniższe zbiory A i B są sobie równe (identyczne): $A = \{a, b, c\}$, $B = \{c, a, b\}$.

Również fakt, że jakiś element zostaje, z jakichś powodów, wymieniony kilkakrotnie, nie zmienia w niczym zbioru. Przykładowo zbiór $C = \{a, a, c, b, a, b, c, a\}$ jest identyczny z wymienionymi wcześniej A i B ; każdy z tych zbiorów (również C !) zawiera trzy elementy – a , b , oraz c .

Szczególnym zbiorem jest tak zwany **zbiór pusty**, który nie zawiera żadnego elementu. Zbiór pusty oznaczamy zwykle symbolem $\emptyset$ – bez żadnych nawiasów klamrowych.

Fakt, że jakiś obiekt jest elementem pewnego zbioru oznaczamy symbolem: $\in$. Symbol ten odczytujemy jako „należy” lub „jest elementem”. W odniesieniu do powyższego przykładu

możemy więc napisać: $a \in A$, $b \in A$ oraz $c \in A$. To, że obiekt nie jest elementem zbioru, zapisujemy przy pomocy znaku: $\notin$. Powiemy, na przykład: $d \notin A$.

Wypisanie elementów w klamrowych nawiasach nie jest jedyną metodą przedstawienia zbioru. Można to uczynić również podając swego rodzaju „przepis” według którego ktoś, gdyby chciał, mógł elementy zbioru wypisać. „Przepis” taki może być mniej lub bardziej formalny. Zbiór $D = \{1, 2, 3, 4\}$ możemy przedstawić na przykład: D – zbiór liczb naturalnych mniejszych od 5; lub bardziej formalnie: $D = \{x: x \in \mathbb{N} \wedge x < 5\}$ (gdzie $\mathbb{N}$ oznacza zbiór liczby naturalnych). Zapis typu $\{x: \dots\}$ odczytujemy: „zbiór takich x , które (elementów), że...”, a więc, w naszym wypadku, powiedzielibyśmy: *zbiór takich x , które są liczbami naturalnymi i są jednocześnie mniejsze od 5*.

Elementami jakiegoś zbioru mogą być nie tylko „zwykłe” obiekty, ale również inne zbiory. Na przykład $X = \{ \{a, b\}, \{c\}, \{d, e, f, g\} \}$. Zbiór X ma trzy elementy, które z kolei same też są zbiorami. To, że te „pomniejsze” zbiory też mają swoje elementy, nie ma żadnego wpływu na ilość elementów X . X ma trzy elementy, ponieważ w jego „głównych” nawiasach klamrowych znajdują się trzy obiekty oddzielone przecinkami.

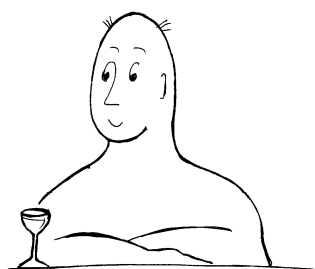
Oczywiście zbiory mogą mieć elementy różnego typu: zarówno „zwykłe” przedmioty, jak i inne zbiory. Na przykład: $Y = \{ \{a, b\}, c, d, \{e, f, g, h\} \}$; zbiór Y ma cztery elementy: c , d , $\{a, b\}$ i $\{e, f, g, h\}$.

Określając elementy zbiorów trzeba bardzo uważnie przyglądać się nawiasom klamrowym. Przykładowo zupełnie różne są zbiory: $A = \{a, b, c\}$ oraz $E = \{ \{a, b, c\} \}$. Zbiór A ma trzy elementy, natomiast E jeden, sam będący zbiorem.

Trzeba również koniecznie zdać sobie sprawę, że różne od siebie są następujące zbiory: $F = \{a\}$ oraz $G = \{ \{a\} \}$. Wprawdzie obydwa mają po jednym elemencie, jednak elementem F jest po prostu „zwykły” obiekt a , natomiast elementem zbioru G jest zbiór, którego elementem jest a .

5.2. STOSUNKI MIĘDZY ZBIORAMI.

5.2.1. ŁYK TEORII.



ŁYK TEORII

Zbiory mogą pozostawać względem siebie w różnych zależnościach.

Identyczność.

Mówimy, że dwa zbiory są sobie równe lub że są identyczne, gdy mają dokładnie te same elementy. Identyczność dwóch zbiorów oznaczamy symbolem: $=$. Posługując się znanymi z rachunku zdań i predykatów symbolami, możemy identyczność zbiorów zdefiniować:

$$A = B \equiv \forall x (x \in A \equiv x \in B)$$

(To, że A i B są równe, oznacza, że dla każdego x to, że x należy do A jest równoważne temu, że x należy do B)

Przykładowo identyczne są zbiory A – zbiór liczb parzystych oraz B – zbiór liczb podzielnych przez 2. Równe są też zbiory $A = \{a, b, c, d\}$ i $B = \{b, d, c, a\}$.

Inkluzja (zawieranie się zbiorów).

Mówimy, że zbiór A zawiera się w zbiorze B (A pozostaje w stosunku inkluzji do B), gdy każdy element A jest jednocześnie elementem B (choć niekoniecznie na odwrót). Inkluzję oznaczamy symbolem: $\subseteq$. Zawieranie się zbiorów możemy przedstawić wzorem:

$$A \subseteq B \equiv \forall x (x \in A \rightarrow x \in B)$$

Inkluzja zachodzi na przykład pomiędzy zbiorami: $A = \{a, b\}$, $B = \{a, b, c, d\}$ lub A – zbiór krokodyli, B – zbiór gadów.

Jeśli zbiór A zawiera się w zbiorze B, to możemy też powiedzieć, że A jest podzbiorem B.

Rozłączność.

Zbiory A i B są rozłączne, gdy nie mają żadnego elementu wspólnego. Rozłączność oznaczamy: $\cap$. Symbolicznie:

$$A \cap B \equiv \forall x (x \in A \rightarrow x \notin B) \text{ lub } \sim \exists x (x \in A \wedge x \in B)$$

Przykładowo, rozłączne są zbiory $A = \{a, b, c\}$ i $B = \{d, e\}$ lub A – zbiór ssaków, B – zbiór płazów.

Krzyżowanie.

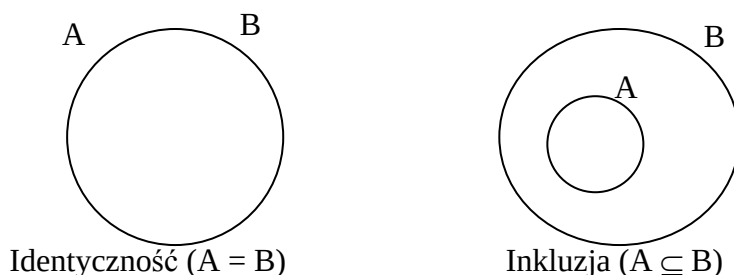
Zbiory się krzyżują gdy mają one pewne elementy wspólne, ale oprócz nich w każdym zbiorze znajdują się również takie obiekty, których nie ma w drugim. Krzyżowanie zbiorów oznaczamy najczęściej przy pomocy dwóch zazębiających się nawiasów, jednakże z przyczyn technicznych (brak takiego symbolu w edytorze tekstu) będziemy na oznaczenie krzyżowania używali obecnie znaku: #. Symbolicznie krzyżowanie zbiorów definiujemy:

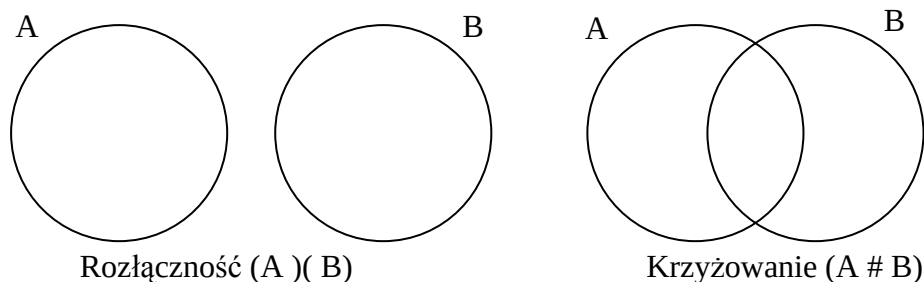
$$A \# B \equiv \exists x (x \in A \wedge x \in B) \wedge \exists x (x \in A \wedge x \notin B) \wedge \exists x (x \notin A \wedge x \in B)$$

Krzyżują się na przykład zbiory: $A = \{a, b, c, d\}$ i $B = \{a, b, e\}$ lub A – zbiór ssaków, B – zbiór drapieżników (istnieją ssaki będące drapieżnikami, ale też ssaki nie będące drapieżnikami oraz drapieżniki nie będące ssakami).

Odnosnie przedstawionych zależności pomiędzy zbiorami dobrze jest zauważyć, że stosunki identyczności, rozłączności oraz krzyżowania się zbiorów są symetryczne. Oznacza to, że jeśli taka zależność zachodzi „w jedną stronę”, to zachodzi również „w drugą”. Jeśli $A = B$, to również $B = A$, jeśli $A \cap B$, to również $B \cap A$, a jeśli $A \# B$, to również $B \# A$. A zatem w przypadku tych stosunków nie jest istotna kolejność, w jakiej wypiszemy pozostające w nich zbiory. Inaczej ma się sytuacja w przypadku inkluzji. Tu fakt, że $A \subseteq B$, nie oznacza, że $B \subseteq A$.

Zależności między zbiorami można przedstawić graficznie:





5.2.2. PRAKTYKA: OKREŚLANIE ZALEŻNOŚCI MIĘDZY ZBIORAMI.

Zadania związane ze stosunkami między zbiorami polegają zwykle na określeniu zależności pomiędzy kilkoma podanymi zbiorami. Po nabraniu pewnej wprawy, zadania tego typu są bardzo łatwe i rozwiązywać je można „od ręki”, bez stosowania jakichkolwiek systematycznych metod. Na początku można posłużyć się metodą eliminacji, po kolei sprawdzając, czy zachodzi dany stosunek, zaczynając od tych, które najłatwiej jest stwierdzić i ewentualnie odrzucić. Przykładowa procedura może wyglądać następująco:

1. Najpierw sprawdzamy, czy zbiory mają te same elementy. Jeśli tak, to znaczy, że są one identyczne, jeśli nie, szukamy dalej.
2. Sprawdzamy wtedy, czy badane zbiory mają choć jeden wspólny element. Jeśli nie mają, znaczy to, że są one rozłączne.
3. Jeśli natomiast zbiory mają jakieś wspólne elementy, to pytamy, czy może jest tak, że każdy element pierwszego jest elementem drugiego lub każdy element drugiego elementem pierwszego. Jeśli tak jest, to znaczy to, że jeden ze zbiorów zawiera się drugim (zachodzi inkluzja).
4. Jeśli tak nie jest, to zbiory muszą się krzyżować – jest to ostatnia możliwość, która nam została. Dla sprawdzenia, możemy zadać sobie pytanie, czy oprócz elementów wspólnych dla obu zbiorów są też takie, które są tylko w jednym i takie, które są tylko w drugim. Jeśli nigdzie wcześniej nie popełniliśmy błędu, to odpowiedź na to pytanie musi być twierdząca.

Przykład:

Sprawdzimy, jakie zachodzą stosunki między następującymi zbiorami:

$$A = \{4\}, B = \{2, 3\}, C = \{1, 2, 3, 4\}, D = \{1, 2, 4\}.$$

Zaczynamy od sprawdzenia, w jakich stosunkach do innych zbiorów pozostaje A. Zbiory A i B nie mają żadnego wspólnego elementu, więc są one rozłączne. W przypadku A i C zachodzi sytuacja przedstawiona w punkcie 3) – każdy element A jest elementem C, a więc A zawiera się w C. Z podobną sytuacją mamy do czynienia w przypadku zbiorów A i D – A zawiera się w D.

Następnie przechodzimy do zbadania, w jakich zależnościach do innych zbiorów pozostaje B. Ponieważ stosunek pomiędzy B i A już znamy, zaczynamy od B i C. Po odrzuceniu dwóch pierwszych możliwości widzimy, że każdy element B jest również elementem C, a zatem B zawiera się w C. W przypadku zbiorów B i D widzimy, że nie są one na pewno identyczne ani rozłączne; nie jest też tak, aby każdy element jednego był elementem

drugiego. A zatem zbiory te muszą się krzyżować. Faktycznie mają one element wspólny – 2, ale jest też taki element który jest tylko w B – 3 oraz elementy będące tylko w D – 1 i 4. Pozostało nam jeszcze określenie stosunku pomiędzy zbiorami C i D. Tutaj widzimy, że każdy element D jest elementem C. A więc zbiór D zawiera się w C. Pamiętajmy, że w przypadku inkluzji istotne jest, który zbiór zawiera się w którym, a więc piszemy: $D \subseteq C$. Ostateczne rozwiązanie zadania wygląda następująco:

$A \cap B, A \subseteq C, A \subseteq D, B \subseteq C, B \not\subseteq D, D \subseteq C$

Przykład:

Określimy stosunki pomiędzy następującymi zbiorami:

A – zbiór studentów prawa,

B – zbiór studentów,

C – zbiór studentów dziennych,

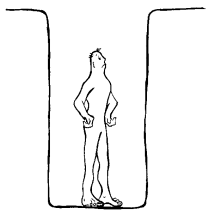
D – zbiór studentów matematyki.

W przypadku zbiorów A i B już na pierwszy rzut oka widać, że każdy element A jest elementem B (każdy student prawa jest studentem), a więc A zawiera się w B. W odniesieniu do zbiorów A i C odrzucamy pierwsze trzy możliwości, co świadczy, że zbiory te się krzyżują. Faktycznie mają one elementy wspólne: dziennych studentów prawa, ale są też obiekty będące elementami tylko zbioru A (zaoczni studenci prawa) oraz będące elementami tylko C (studenci dzienni innego niż prawo kierunku – np. filozofii). W przypadku zbiorów A oraz D z powodu braku danych empirycznych trudno dać jednoznaczną odpowiedź. Albo jest tak, że zbiory te są rozłączne (jeśli żaden student prawa nie studiuje jednocześnie matematyki), albo też, jeśli znajdzie się choć jedna osoba studiująca oba te kierunki, zbiory te się krzyżują. Zauważmy, że jeśli będziemy rozpatrywać wszystkich studentów na całym świecie, to zapewne zbiory te się krzyżują, jeśli natomiast ograniczymy nasze rozważania do jakiegoś wybranego niewielkiego uniwersytetu, to mogą być one rozłączne. Na pewno natomiast nie są to zbiory identyczne, ani też jeden z nich nie zawiera się w drugim.

Jeśli chodzi o zbiór B i C oraz B i D, to w każdym z tych przypadków zachodzi inkluzja. Pamiętajmy jednak o właściwej kolejności: to C zawiera się B (każdy student dzienny jest studentem) oraz D zawiera się w B (każdy student matematyki jest studentem) a nie na odwrót. W przypadku zbiorów C i D zachodzi krzyżowanie – istnieją dzienni studenci matematyki, a także dzienni studenci innych kierunków, oraz zaoczni studenci matematyki.

Ostateczna odpowiedź, to zatem:

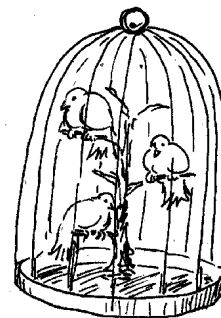
$A \subseteq B, A \not\subseteq C, A \cap D \text{ lub } A \not\subseteq D, C \subseteq B, D \subseteq B, C \not\subseteq D$.



5.2.3. UTRUDNIENIA I PUŁAPKI.

Pomiędzy zbiorami może zachodzić jeszcze jeden stosunek, trochę innego typu niż omówione wyżej. Może się mianowicie zdarzyć tak, że jeden zbiór sam jest elementem innego zbioru, czyli: $A \in B$. Aby tak było, zbiór B musi szczególnym rodzajem zbioru – takim, którego elementy (przynajmniej niektóre) są zbiorami. Sytuacja taka zachodzi na przykład w stosunku do następujących zbiorów: A – zbiór kanarków, B – zbiór, którego elementami są zbiory ptaków poszczególnych gatunków.

Bardzo istotne jest, aby nie mylić zawierania się zbiorów, czyli zależności $A \subseteq B$ oraz bycia elementem (należenia), czyli $A \in B$. Pierwsza zależność, inkluzja ($\subseteq$), oznacza, że każdy element zbioru A jest również elementem zbioru B. Należenie ($\in$) natomiast, oznacza, że sam zbiór A, jako całość, jest elementem zbioru B. W przypadku przedstawionych wyżej zbiorów A nie zawiera się w B, bo nie jest tak, aby każdy kanarek (elementy A) był jednocześnie zbiorem ptaków jakiegoś gatunku (elementy B). Natomiast A jako całość (czyli zbiór kanarków), jest jednym z elementów B.



Stosunek należenia (jeśli zachodzi), jest zależnością, która występuje niejako obok „zwykłych”, omawianych wyżej relacji między zbiorami. Znaczy to, że pewien zbiór A należąc do zbioru B (będąc elementem B) może jednocześnie być z nim rozłączny, zawierać się w nim lub krzyżować.

ZBIÓR KANARKÓW

Przykład:

Zobaczmy w jakich stosunkach pozostają do siebie zbiory:

$$A = \{a, b\},$$

$$B = \{ \{a, b\}, \{c, d, e\} \},$$

$$C = \{a, b, c, d, e\},$$

$$D = \{a, b, \{a, b\} \},$$

$$E = \{a, d, e, \{a, b\} \}$$

Zbiory A i B nie mają wspólnych elementów, ponieważ elementami A są „zwykłe” obiekty a oraz b, natomiast elementami B są zbiory. Tak więc A i B są rozłączne. Jednocześnie jednak zbiór A sam jest jednym z elementów zbioru B. W przypadku A oraz C sprawa jest oczywista: każdy element A jest elementem C, a zatem A zawiera się w C. Porównując A oraz D widzimy, że każdy element A jest elementem D. D ma jednak również trzeci element będący zbiorem; a ten zbiór, to nic innego, jak A. A zatem A zawiera się w D i jednocześnie należy do D. Jeśli chodzi o zbiory A i E, to mają one jeden element wspólny (a), ale też każdy z nich ma też takie elementy, których nie ma w drugim (b w zbiorze A oraz d, e i $\{a, b\}$ w zbiorze E. Tak więc zbiory te się krzyżują. Równocześnie jednak A sam jest jednym z elementów E.

Porównując B oraz C już na pierwszy rzut oka widzimy, że nie mogą mieć one żadnego wspólnego elementu, ponieważ elementami B są zbiory, natomiast elementami C „zwykle” obiekty. Tak więc B i C są rozłączne. Zbiory B i D mają jeden wspólny element: zbiór {a, b}. Jednocześnie w B jest element, którego nie ma D – zbiór {c, d, e}, natomiast w D elementy, których nie ma w B – a, b. Zbiory B i D się zatem krzyżują. Analogiczna sytuacja zachodzi w przypadku B i E.

Nie powinno nikomu sprawić trudności zauważenie, że krzyżują się również zbiory C i D, C i E oraz D i E.

Ostateczne rozwiązanie, to zatem:

A) ($B \cap A \in B, A \subseteq C, A \subseteq D \wedge A \in D, A \# E \wedge A \in E,$

B) ($C, B \# D, B \# E,$

C) ($\# D, C \# E,$

D) ($\# E.$

Zadanie:

Określmy zależności pomiędzy następującymi zbiorami.

A – zbiór studentów, którzy zdali logikę na 5,

B – zbiór studentów, którzy zdali logikę na 3,

C – zbiór studentów leniwych,

D – zbiór, którego elementami są zbiory studentów, którzy zdali logikę na taką samą ocenę.

Zbiory A i B są rozłączne (oczywiście przy założeniu, że nikt nie zdawał logiki dwukrotnie, na przykład „za kolegę”). A i C się krzyżują: na pewno są studenci, którzy zdali logikę na 5, będąc jednocześnie leniwymi, ale też są tacy, którzy otrzymali 5 i są pracowici, a także i tacy, którzy są leniwi i nie dostali 5. Zbiory A i D nie mogą mieć żadnego wspólnego elementu z tej prostej przyczyny, że elementami A są „zwykli” studenci, natomiast elementami D zbiory studentów; A i D mają więc elementy różnych typów. Oprócz tego, że są to zbiory rozłączne, zachodzi jednak między nimi jeszcze jeden stosunek: zbiór A sam jest jednym z elementów zbioru D. Gdybyśmy bowiem wypisali sobie elementy zbioru D, to byłyby to: zbiór studentów, którzy zdali logikę na 5, zbiór studentów, którzy zdali logikę na 4, zbiór studentów, którzy zdali logikę na 3 i zbiór studentów, którzy zdali logikę na 2. Zbiór A zatem należy do D.

Zbiory B i C się krzyżują, podobnie jak A i C. Natomiast w przypadku B i D, analogicznie jak w A i D, zachodzą dwa stosunki: rozłączności i należenia.

W przypadku C i D mamy do czynienia tylko z rozłącznością. Zbiory te nie mają wspólnych elementów, gdyż elementami pierwszego są studenci, a drugiego zbiory. Jednocześnie jednak C sam nie jest jednym z elementów D.

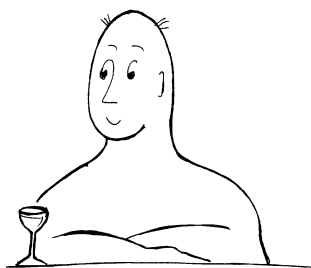
Ostateczne rozwiązanie:

A) ($B, A \# C, A$) ($D \wedge A \in D,$

$B \# C, B)(D \cap B \in D,$
 $C)(D.$

5.3. DZIAŁANIA NA ZBIORACH.

5.3.1. ŁYK TEORII.



ŁYK TEORII

Na zbiorach można wykonywać różne działania, w wyniku których powstają nowe zbiory.

Poniżej omówimy najważniejsze z nich.

Suma.

Suma zbiorów A i B , to zbiór powstały ze wszystkich elementów A i B . Obrazowo tworzenie sumy zbiorów możemy sobie wyobrazić, jako wsypywanie elementów dodawanych zbiorów do jednego dużego worka, który reprezentuje ich sumę. Przykładowo sumą zbiorów

mężczyzn i zbioru kobiet jest zbiór wszystkich ludzi. Sumę zbiorów oznaczamy symbolem $\cup$.

Warto zauważyć, że gdy jeden zbiór zawiera się w drugim, to ich sumą jest zbiór „większy”.

Iloczyn.

Iloczyn zbiorów A i B to po prostu część wspólna tych zbiorów; zbiór utworzony z tych elementów, które należą jednocześnie do A i B . Przykładowo, iloczynem zbioru kobiet oraz osób palących jest zbiór palących kobiet.

Iloczyn zbiorów nazywany bywa również ich przekrojem. Oznaczamy go symbolem $\cap$.

Warto zapamiętać, że gdy jeden zbiór zawiera się w drugim, to ich iloczynem jest zbiór „mniejszy”. Jeśli natomiast zbiory są rozłączne, to ich iloczynem jest zbiór pusty.



ZBIÓR PALĄCYCH KOBIET

Różnica.

Różnica zbiorów A i B to zbiór utworzony z elementów, które należą do A i jednocześnie nie należą do B . Obrazowo tworzenie różnicy zbiorów A i B możemy sobie wyobrazić jako wykreślanie ze zbioru A elementów, które są również w B ; to co pozostanie, to właśnie różnica A i B . Przykładowo różnicą zbioru mężczyzn oraz osób palących jest zbiór niepalących mężczyzn. Różnicę oznaczamy symbolem kreski poziomej lub skośnej, czyli: $-$ lub $/$.

Warto zapamiętać, że jeśli od dowolnego zbioru A odejmujemy jakiś zbiór z A rozłączny, to A pozostaje „nienaruszony”; wynikiem takiego działania jest A .

Dopełnienie.

Dopełnienie, to działanie trochę inne od dotychczas omawianych. Dotyczy ono bowiem nie dwóch zbiorów, ale tylko jednego. Dopełnienie pewnego zbioru A to zbiór tych obiektów, które nie należą do A . Dopełnienie wykonujemy zawsze w odniesieniu do tak zwanego uniwersum, czyli dziedziny, w której się poruszamy. Przykładowo, jeśli uniwersum stanowi zbiór ludzi, to dopełnieniem zbioru ludzi palących jest zbiór ludzi niepalących (a nie zbiór wszystkich istot i przedmiotów niepalących). Dopełnienie oznaczamy symbolem „prim”

Warto zapamiętać, że dopełnieniem uniwersum jest zawsze zbiór pusty, a dopełnieniem zbioru pustego uniwersum: $U' = \emptyset$, $\emptyset' = U$.

Ponadto suma dowolnego zbioru oraz jego dopełnienia da nam zawsze uniwersum ($A \cup A' = U$), natomiast iloczyn dowolnego zbioru i jego dopełnienia to zawsze zbiór pusty ($A \cap A' = \emptyset$).

5.3.2. PRAKTYKA: WYKONYWANIE DZIAŁAŃ NA ZBIORACH.

Obecnie wykonamy kilka przykładowych działań na podanych zbiorach.

Przykład:

Przyjmijmy uniwersum $U = \{1, 2, 3, 4, 5\}$, oraz następujące zbiory:

$A = \{4\}$, $B = \{2, 3\}$, $C = \{1, 2, 3, 4\}$, $D = \{1, 2, 4\}$.

Na zbiorach tych wykonamy kilka działań.

a) $B \cup D$

Suma dwóch zbiorów powstaje przez połączenie ich elementów w jednym zbiorze. Jeśli jakiś element występuje w obu zbiorach, to wypisujemy go tylko raz, a więc $B \cup D = \{1, 2, 3, 4\}$.

b) $D \cap B$

Iloczyn zbiorów to ich część wspólna. Zbiory D i B mają tylko jeden wspólny element – 2. A więc, $D \cap B = \{2\}$.

c) D'

Dopełnienie zbioru to zbiór złożony z tych elementów uniwersum, które nie należą do rozpatrywanego zbioru. A zatem: $D' = \{3, 5\}$.

d) $C - B$

Różnica C i B to zbiór z tych elementów C , których nie ma w B . Warunek ten spełniają: 1 i 4. A zatem: $C - B = \{1, 4\}$.

e) $B - C$

Różnicę B i C tworzymy biorąc zbiór B i „wykreślając” z niego te elementy, które znajdziemy również w C . Okazuje się, że postępując w ten sposób, pozbywamy się wszystkich elementów. Czyli: $B - C = \emptyset$.

Zauważmy, że wynik różnicy (podobnie jak odejmowania liczb) zależy od kolejności zbiorów; $B - C$, to co innego niż $C - B$.

f) $B' - A$

W powyższym przykładzie mamy dwa działania. Najpierw musimy wykonać B' , a potem od tego odjąć zbiór A . Dopełnienie B to zbiór: $\{1, 4, 5\}$. Gdy odejmiemy od niego A , czyli $\{4\}$, zostanie zbiór złożony z 1 i 5. A zatem: $B' - A = \{1, 5\}$.

g) $C \cap D'$

Dopełnienie zbioru D , to $\{3, 5\}$. Z elementów tych jedynie 3 jest elementem C , a więc: $C \cap D' = \{3\}$.

h) $D - (A \cap C)$

W tym przypadku musimy najpierw wykonać działanie w nawiasie. Iloczyn A i C to zbiór $\{4\}$. A więc ostatecznie wykonujemy działanie: $D - \{4\}$. Tak więc $D - (A \cap C) = \{1, 2\}$

i) $D - (C \cup B)$

Suma C i B , które to działanie musimy wykonać najpierw, to zbiór: $\{1, 2, 3, 4\}$. Gdy odejmiemy go od zbioru D , otrzymamy zbiór pusty. Zatem: $D - (C \cup B) = \emptyset$

Przykład:

Przyjmując uniwersum U – zbiór wszystkich kwiatów, oraz zbiory:

A – zbiór tulipanów,

B – zbiór róż,

C – zbiór kwiatów czerwonych,

D – zbiór białych róż,

wykonamy na tych zbiorach kilka działań.

a) $B \cap C$

Część wspólna zbiorów róż oraz kwiatów czerwonych to niewątpliwie zbiór czerwonych róż.

b) $B \cup D$

Do B, czyli zbioru róż, dodajemy zbiór białych róż, a więc zbiór w nim już zawarty. W takim przypadku wynikiem działania jest B – zbiór róż.

c) $A' \cap C$

Dopełnienie A to zbiór kwiatów nie będących tulipanami. Część wspólna tego zbioru ze zbiorem kwiatów czerwonych to, mówiąc najkrócej, czerwone nie-tulipany.

d) $A - C'$

Dopełnienie C to zbiór kwiatów we wszystkich innych kolorach, oprócz czerwonego. Jeśli od zbioru tulipanów, takie nie-czerwone kwiaty odejmiemy, pozostaną nam jedynie czerwone tulipany.

e) $B' - A$

B' to zbiór wszystkich kwiatów nie będących różami. Od takiego zbioru odejmujemy jeszcze zbiór tulipanów. Zostaje nam więc zbiór wszystkich kwiatów za wyjątkiem róż i tulipanów.

f) $D - B'$

Od zbioru białych róż odejmujemy zbiór kwiatów nie będących różami. Obrazowo rzecz ujmując, od zbioru D próbujemy odjąć coś, czego w nim nie ma. W takim przypadku D pozostaje „nienaruszony”. Wynikiem działania jest więc zbiór białych róż.

g) $B \cup C$

Do zbioru róż dodajemy wszelkie czerwone kwiaty. Otrzymujemy więc zbiór składający się ze wszystkich róż (bez względu na kolor) oraz pozostałych kwiatów będących jednak tylko czerwonymi.

h) $(A \cup B) - C$

Suma w nawiasie, to zbiór złożony z róż i tulipanów. Jeśli od takiego zbioru odejmiemy kwiaty czerwone, otrzymamy zbiór róż i tulipanów w innych niż czerwonym kolorze.

i) $(B - D) \cup A$

Wynikiem działania w nawiasie jest zbiór róż, które nie są białe. Jeśli dodamy do niego zbiór A, otrzymamy zbiór składający się z takich nie-białych róż oraz (wszystkich) tulipanów.

j) $(D - B) \cap C'$

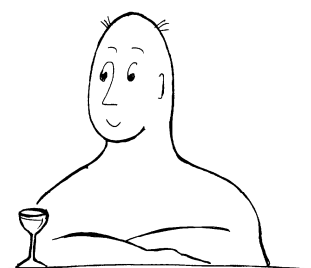
Gdy od zbioru białych róż odejmiemy róże, pozostanie nam zbiór pusty. Iloczyn (czyli część wspólna) zbioru pustego z dowolnym zbiorem, to też zbiór pusty, a zatem wynikiem całego działania jest $\emptyset$.

k) $D \cup D'$

Do białych róż musimy dodać pozostałe kwiaty. Suma jakiegokolwiek zbioru i jego dopełnienia to zawsze całe uniwersum, a więc, w tym wypadku, zbiór wszystkich kwiatów.

5.4. PRAWA RACHUNKU ZBIORÓW TYPU BEZZAŁOŻENIOWEGO.

5.4.1. ŁYK TEORII.



ŁYK TEORII

Badając teorię zbiorów odnaleźć możemy wyrażenia będące zawsze prawdziwymi, niezależnie od tego, do jakich zbiorów się one odnoszą. Przykładem takich wyrażeń mogą być wzory: $(A \cup B) = (B \cup A)$, $(A \cap B) \subseteq A$ czy też $(A \subseteq B \wedge B \subseteq C) \rightarrow A \subseteq C$. Takie zawsze prawdziwe wyrażenia nazywamy prawami rachunku zbiorów. Pierwsze dwa z powyższych wzorów mają postać równości oraz inkluzji pewnych zbiorów; stwierdzają one bezwarunkowe zachodzenie pewnego związku. Trzeci z przedstawionych wzorów ma postać implikacji; mówi on, że pewna zależność zachodzi, jeśli zachodzi inna. Tego typu, założeńiowymi prawami, zajmiemy się w kolejnym paragrafie. Obecnie

natomiast omówimy wyrażenia mające postać równości bądź inkluzji zbiorów.

Do wykrywania omawianych, bezzałożeńiowych, praw rachunku zbiorów posłużyć się możemy metodą wykorzystującą klasyczny rachunek zdań i pojęcie tautologii. W miarę dobra znajomość KRZ jest więc konieczna do dalszych rozważań.

Ponieważ wyrażenia, które będziemy badali, mają postać równości bądź zawierania się zbiorów, rozpoczniemy od uświadomienia sobie, co dokładnie oznaczają te dwie zależności. Fakt, że jeden zbiór zawiera się w drugim, przedstawić możemy przy pomocy stwierdzenia, że jeśli jakiś obiekt jest elementem zbioru A, to również jest on elementem B (dla dowolnego obiektu x, jeśli x należy do A, to x należy do B). Można zapisać to wzorem:

$$1) A \subseteq B \equiv \forall x (x \in A \rightarrow x \in B)$$

To, że dwa zbiory są równe, oznacza, że jeśli dowolny obiekt jest elementem A, to jest on również elementem B, a jeśli jest elementem B, to jest też elementem A. Innymi słowy, to, że dowolny x należy do A, jest równoważne temu, że należy on do B. Formalnie:

$$2) A = B \equiv \forall x (x \in A \equiv x \in B)$$

W naszych prawach, które będziemy badali, występują również pojęcia iloczynu, sumy, różnicy i dopełnienia zbiorów. Dlatego też powinniśmy zdać sobie sprawę, co oznacza fakt, że jakiś obiekt należy do iloczynu, sumy lub różnicy dwóch zbiorów, czy też dopełnienia jakiegoś zbioru.

To, że pewien obiekt x jest elementem iloczynu (czyli części wspólnej) zbiorów A i B oznacza, że należy on zarówno do A jak i do B. Symbolicznie:

$$3) x \in (A \cap B) \equiv (x \in A \wedge x \in B)$$

Fakt, że jakiś obiekt x należy do sumy zbiorów A i B , oznacza, że należy on do A lub też należy do B . Formalnie:

$$4) x \in (A \cup B) \equiv (x \in A \vee x \in B)$$

To, że pewien x należy do różnicy zbiorów A i B oznacza, że należy on do zbioru A i jednocześnie nie jest prawdą, że należy do B . Symbolicznie:

$$5) x \in (A - B) \equiv (x \in A \wedge \sim (x \in B))$$

Należenie jakiegokolwiek obiektu x do dopełnienia pewnego zbioru A oznacza po prostu, że nie jest prawdą, iż ów x należy do A :

$$6) x \in A' \equiv \sim (x \in A)$$

Znajomość powyższych wzorów 1) – 6) będzie konieczna, aby móc sprawdzić, czy jakieś wyrażenie jest prawem rachunku zbiorów.

5.4.2. PRAKTYKA: WYKRYWANIE PRAW RACHUNKU ZBIORÓW PRZY POMOCY RACHUNKU ZDAŃ.

Wykrycie, czy dane wyrażenie, mające postać równości bądź inkluzji zbiorów, jest ogólnie obowiązującym prawem, będzie polegało na przekształceniu formuły rachunku zbiorów na formułę rachunku zdań, a następnie sprawdzeniu, czy otrzymany schemat jest tautologią. Jeśli otrzymana formuła okaże się tautologią, będzie to oznaczało, że wyjściowy wzór jest prawem rachunku zbiorów; jeśli formuła nie będzie tautologią, to znak, że badane wyrażenie nie jest takim prawem.

Przekształcanie formuły rachunku zbiorów na rachunek zdań polegać będzie na systematycznym stosowaniu poznanych wyżej wzorów, aż otrzymamy wyrażenie, w którym nie będzie symboli oznaczających inkluzję, równość, sumę, iloczyn, różnicę i dopełnienie zbiorów, zamiast których pojawią się symbole rachunku zdań (implikacja, równoważność, alternatywa, koniunkcja, negacja). Przekształcenia takie będziemy wykonywali krok po kroku. W pierwszym ruchu będziemy zawsze stosowali wzór 1) lub 2), aby zamienić inkluzję bądź równość zbiorów na implikację lub równoważność. Następnie, w zależności od potrzeb, będziemy korzystali ze wzorów 3) – 6).

Przykład:

Sprawdzimy, czy prawem rachunku zbiorów jest wyrażenie:

$$(A - B) \subseteq (A \cup B)$$

Ponieważ całe wyrażenie ma postać inkluzji, rozpoczniemy od zastosowania wzoru 1), dzięki któremu otrzymamy:

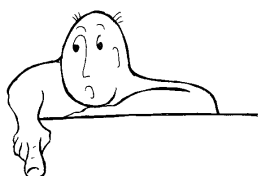
$$\forall x [x \in (A - B) \rightarrow x \in (A \cup B)]$$

Kwantyfikator na początku formuły, informujący nas, że implikacja powinna zachodzić dla każdego obiektu x , możemy w następnych krokach pomijać. Skoro bowiem implikacja ma być prawdziwa dla każdego (dowolnego) x , to w tym również dla naszego konkretnego obiektu x , który sobie wybraliśmy. A zatem możemy zapisać:

$$x \in (A - B) \rightarrow x \in (A \cup B)$$

Teraz możemy przystąpić do kolejnych przekształceń. Poprzednik implikacji stwierdza, że nasz obiekt x należy do różnicy zbiorów A i B ; musimy tam zatem zastosować wzór 5). W odniesieniu do następnika powinniśmy natomiast skorzystać ze wzoru 4). Otrzymamy wtedy:

$$(x \in A \wedge \sim (x \in B)) \rightarrow (x \in A \vee x \in B)$$



Uwaga na błędy!

Dokonując przekształceń należy bardzo uważać, aby nie zmienić struktury nawiasów. Jeżeli wzór mówi, że nasz x należy do pewnej całości umieszczonej w nawiasie, to po wykonaniu przekształcenia nawias ten musi pozostać. Można obrazowo powiedzieć, że x „wchodzi w głąb” nawiasu, nie naruszając go jednak.

Po przekształceniu symboli związanych ze zbiorami (poza $\in$) na symbole rachunku zdań możemy naszą formułę zmienić całkowicie na schemat KRZ, podstawiając na przykład za wyrażenie $x \in A$ zmienną p , natomiast za $x \in B$ zmienną q . Otrzymamy wtedy:

$$(p \wedge \sim q) \rightarrow (p \vee q)$$

Teraz pozostaje nam sprawdzenie, czy otrzymana formuła jest tautologią. Uczynienie tego skróconą metodą zero-jedynkową nie powinno sprawić nikomu najmniejszych trudności.

$$(p \wedge \sim q) \rightarrow (p \vee q)$$

1 1 1 0 0 **1 0 0**

Otrzymana sprzeczność (która mogła komuś również wyjść w poprzedniku implikacji) wskazuje, że formuła KRZ nie może być fałszywa, a zatem jest ona tautologią. Na tej podstawie możemy stwierdzić, że badane przez nas wyrażenie jest prawem rachunku zbiorów.

Przykład:

Sprawdzimy, czy prawem rachunku zbiorów jest wyrażenie:

$$(A' \cap B') = (A \cup B)'$$

W pierwszym kroku musimy zastosować wzór 2) i pozbyć się znaku równości:

$$\forall x [x \in (A' \cap B') \equiv x \in (A \cup B)']$$

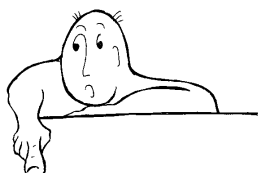
Po opuszczeniu kwantyfikatora otrzymujemy:

$$x \in (A' \cap B') \equiv x \in (A \cup B)'$$

Teraz możemy przystąpić do dalszych przekształceń. W każdym nawiasie mamy jednak dwa różne działania: iloczyn i dopełnienie w pierwszym oraz sumę i dopełnienie w drugim. Dokonując przekształceń, zajmujemy się zawsze najpierw działaniem w danym

momencie głównym, „najszerszym” w danej formule. W pierwszym członie równoważności działaniem takim jest iloczyn; nasz x należy tam do iloczynu A' oraz B' . W związku z tym najpierw zastosujemy tam wzór 3). Natomiast w drugim członie równoważności głównym działaniem jest dopełnienie; x należy do dopełnienia sumy A oraz B . Dlatego też w pierwszej kolejności zastosujemy tam wzór 6); skoro x należy do dopełnienia sumy A i B , to znaczy, iż nie jest prawdą, że należy on do tej sumy. Dokonując przekształceń pamiętamy o zachowaniu struktury nawiasów. Otrzymujemy:

$$(x \in A' \wedge x \in B') \equiv (\sim (x \in (A \cup B)))$$

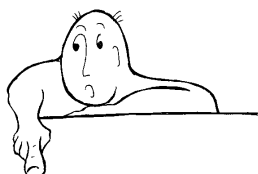


Uwaga na błędy!

W omawianym przykładzie szczególnie istotne jest właściwe umieszczenie nawiasów z prawej strony równoważności. Musimy tam dodać jeden (wynikający ze wzoru 6)) nawias, który wskazuje że całe wyrażenie: $x \in (A \cup B)$ jest nieprawdziwe. Błędne byłoby dodanie samej negacji, bez nawiasu, czyli: $\sim x \in (A \cup B)$

Teraz możemy dokonać dalszych przekształceń. Z lewej strony musimy zastosować (dwukrotnie) wzór 6), natomiast z prawej wzór 4). Otrzymamy wtedy:

$$(\sim (x \in A) \wedge \sim (x \in B)) \equiv (\sim (x \in A \vee x \in B))$$



Uwaga na błędy!

Jeśli w pewnym miejscu mamy znak negacji przed nawiasem (tak jak w prawej części naszej równoważności) to negację taką zostawiamy w tym miejscu. Nie wolno jej pod żadnym pozorem „wciskać” do środka nawiasu lub robić z niej dwóch negacji. Błędne byłyby następujące przekształcenia prawej strony naszej formuły: $(\sim x \in A \vee x \in B)$ lub $(\sim x \in A \vee \sim x \in B)$.

Doświadczenie wskazuje, że takie błędy są bardzo często popełniane przez studentów, dlatego warto dobrze sobie zapamiętać powyższą uwagę.

W tym momencie możemy ostatecznie przekształcić naszą formułę na wyrażenie rachunku zdań podstawiając p za $x \in A$ oraz q za $x \in B$. Otrzymamy wzór:

$$(\sim p \wedge \sim q) \equiv \sim (p \vee q)$$

Łatwo sprawdzić, że powyższa formuła jest tautologią (pamiętamy, że w przypadku równoważności musimy rozpatrzyć dwie możliwości):

1 0 1 1 0 0 0 0 1 0

$$(\sim p \wedge \sim q) \equiv \sim (p \vee q)$$

1 0 0 1 0 0 1 0 0 0

Ponieważ otrzymana formuła jest tautologią, badane wyrażenie jest prawem rachunku zbiorów.

Przykład:

Sprawdzimy, czy prawem rachunku zbiorów jest wyrażenie:

$$[(A \cap B) - C] \subseteq [(A - B) \cap (B - C)]$$

Po zastosowaniu wzoru 1) i opuszczeniu kwantyfikatora otrzymujemy:

$$x \in [(A \cap B) - C] \rightarrow x \in [(A - B) \cap (B - C)]$$

W poprzedniku implikacji głównym działaniem jest odejmowanie, dlatego najpierw musimy zastosować tam wzór 5). W następniku główne działanie, to iloczyn, więc wykorzystujemy wzór 3). Otrzymujemy:

$$[x \in (A \cap B) \wedge \sim (x \in C)] \rightarrow [x \in (A - B) \wedge x \in (B - C)]$$

Teraz w poprzedniku implikacji musimy jeszcze skorzystać ze wzoru 3), a w następniku, dwukrotnie, ze wzoru 5):

$$[(x \in A \wedge x \in B) \wedge \sim (x \in C)] \rightarrow [(x \in A \wedge \sim (x \in B)) \wedge (x \in B \wedge \sim (x \in C))]$$

Po podstawieniu zmiennej p za $x \in A$, q za $x \in B$ oraz r za $x \in C$ otrzymamy:

$$[(p \wedge q) \wedge \sim r] \rightarrow [(p \wedge \sim q) \wedge (q \wedge \sim r)]$$

Po sprawdzeniu formuły skróconą metodą zero-jedynkową okazuje się, że może ona stać się schematem zdania fałszywego, a więc nie jest ona tautologią:

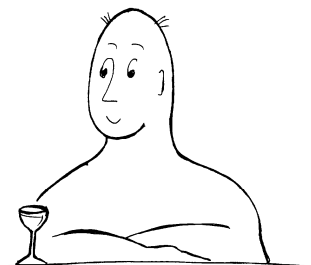
$$[(p \wedge q) \wedge \sim r] \rightarrow [(p \wedge \sim q) \wedge (q \wedge \sim r)]$$

1 1 1 1 1 0 0 1 0 0 1 0 1 1 1 0

Skoro otrzymana formuła nie jest tautologią, to badane wyrażenie nie jest prawem rachunku zbiorów.

5.5 ZAŁOŻENIOWE PRAWA RACHUNKU ZBIORÓW.

5.5.1. ŁYK TEORII.



ŁYK TEORII

Przedstawiona w poprzednim paragrafie metoda nie nadaje się do sprawdzania wszystkich praw rachunku zbiorów. W przypadku praw mających postać implikacji, stwierdzających, że jeśli mają miejsce pewne zależności, to występuje również inna zależność, będziemy posługiwać się, znanymi już z rozdziału o sylogizmach, diagramami Venna.

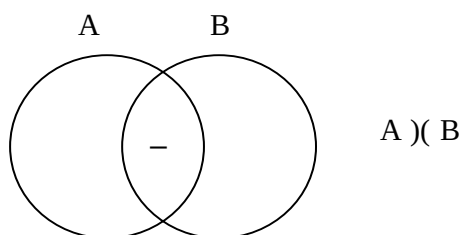
Osoby, które przy okazji przerabiania sylogistyki opanowały posługiwanie się diagramami, nie powinny mieć żadnych trudności ze zrozumieniem dalszego ciągu tego rozdziału, a wiele zawartych tu informacji i szczegółowych komentarzy wyda im się zbędnymi. Ponieważ jednak zapewne nie wszyscy czytelnicy obecnego rozdziału przerabiali wcześniej teorię sylogizmów, niektóre wiadomości odnośnie diagramów Venna będą się musiały powtarzać.

Diagramy Venna przybierają kształt kół symbolizujących zbiory, w które to koła wpisujemy znak „+”, gdy wiemy, że w danym obszarze na pewno znajduje się jakiś element lub „-”, gdy mamy pewność, że nic tam nie ma. Wypełniając diagramy musimy pamiętać, że wpisujemy znaki „+” lub „-” jedynie tam, gdzie wiemy, że na pewno coś jest lub na pewno nic nie ma. Jeśli w stosunku do jakiegoś obszaru nie mamy żadnych informacji, zostawiamy go pustym.

Poniżej przedstawimy sposoby zaznaczania na diagramach przykładowych zależności mogących występować w prawach rachunku zbiorów. Najpierw będziemy nanosić je na diagramy reprezentujące dwa zbiory, a następnie rozszerzymy nasze rozważania na diagramy składające się z trzech kół.

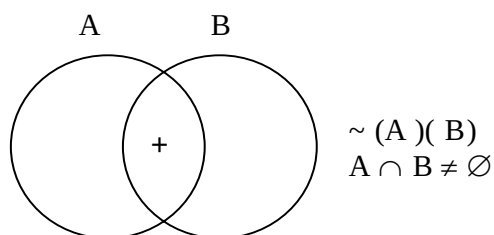
$A \cap B$

Aby zaznaczyć, że zbiory A i B są rozłączne wpisujemy znak „-” w obszar reprezentujący część wspólną tych zbiorów. W części „boczne” nie wolno nam jednak wpisać „+”, bo nie mamy pewności, czy A lub B nie są przypadkiem zbiorami pustymi. Jedyne, co wiemy na pewno, to to, że skoro A i B mają być rozłączne, to nie ma nic w ich części wspólnej.



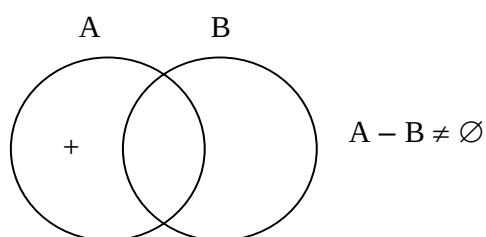
$$\sim (A \cap B) \text{ lub } A \cap B \neq \emptyset$$

Fakt, że zbiory A i B nie są rozłączne lub, ujmując rzecz inaczej, ich iloczyn nie jest zbiorem pustym, oznaczamy wpisując znak „+” w część wspólną tych zbiorów.



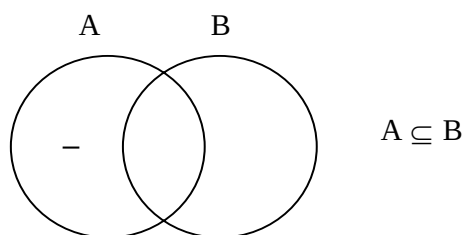
$$A - B \neq \emptyset$$

Fakt, że różnica zbiorów A i B nie jest zbiorem pustym, zaznaczamy wpisując „+” w część diagramu reprezentującą zbiór A - B, a więc obszar A leżący poza B.



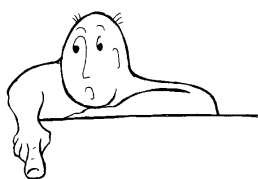
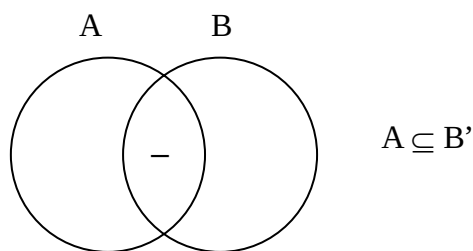
$$A \subseteq B$$

Fakt, że zbiór A zawiera się w B zaznaczamy, wpisując „-” w obszar A znajdujący się poza B. Skoro bowiem A ma się zawierać w B, to żadna część A nie może „wystawać” poza B. W część wspólną, wbrew pozorom, nie możemy wpisać „+”, gdyż nie można wykluczyć, że A jest zbiorem pustym. Jedyne, co wiemy na pewno, to fakt, że nic nie ma w części A leżącej poza B.



$$A \subseteq B'$$

Jeśli zbiór A zawiera się w dopełnieniu B, to znaczy to, że cały A znajduje się poza B, a więc, że żadna część A nie może znajdować się w B. Oznacza to nic innego, jak rozłączność tych zbiorów.

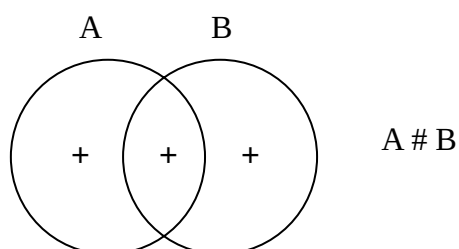


Uwaga na błędy!

Najczęściej popełnianie błędy przy wypełnianiu diagramów Venna polegają na wpisywaniu znaków „+” tam, gdzie nie możemy ich wpisać z uwagi na to, że nie można wykluczyć, iż rozpatrywany zbiór jest pusty. Dobrze zatem zapamiętać, że **zaznaczając inkluzję oraz rozłączność zbiorów nigdy nie wpisujemy żadnych plusów**. Stosunki te oddajemy jedynie przy pomocy minusów.

$A \# B$

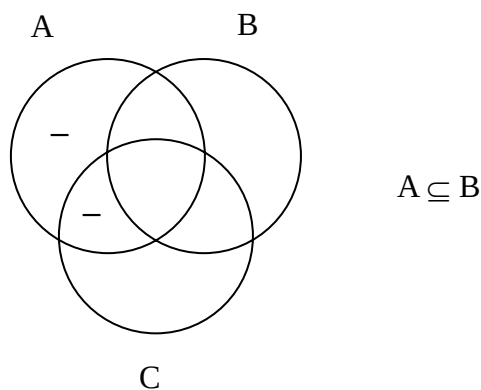
To, że zbiory się krzyżują, oznacza, że na pewno coś znajduje się w ich części wspólnej, a także na pewno jest coś w części A leżącej poza B oraz części B leżącej poza A.



Teraz kilka przykładowych zależności między zbiorami przedstawimy na diagramach reprezentujących trzy zbiory.

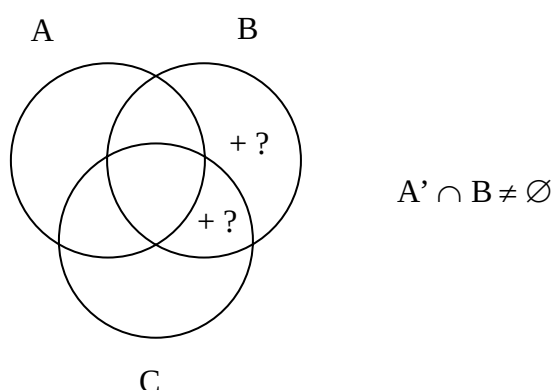
$A \subseteq B$

Jeśli zbiór A zawiera się w B, oznacza to, że A „nie wystaje” poza B. Pusta musi być zatem część A leżąca poza B. Obszar ten składa się teraz jednak z dwóch kawałków. Ponieważ ma on być cały pusty, musimy postawić „-” w każdej jego części.



$$A' \cap B \neq \emptyset$$

Część wspólna dopełnienia zbioru A oraz zbioru B, to ten obszar B, który znajduje się poza A – prawy półksiężyc zbioru B. Musimy przedstawić fakt, że część ta nie jest pusta. Ponieważ obszar ten składa się z dwóch części, ktoś mógłby pomyśleć, że w obie te części musimy wpisać znak „+”. Tak jednak nie jest. Już jeden „+” w którymkolwiek, dolnym lub górnym kawałku rozważanego półksiężyca, sprawi, że iloczyn A' oraz B nie będzie pusty. W związku z powyższym, nie mamy pewności, gdzie znak plusa postawić. Niepewność tę wyrazimy dodając znak zapytania przy każdym z plusów. Pytajniki te będą oznaczać, iż wiemy, że w którymś z rozważanych obszarów (a być może i w obydwu), coś się znajduje, jednak całkiem możliwe jest również, że jeden z nich jest pusty.



Uwaga na marginesie.

W praktyce, gdy będziemy wykorzystywali diagramy do rozwiązywania zadań, często będzie się zdarzać, że dysponując innymi informacjami, będziemy wiedzieli, który z „wątpliwych” plusów na pewno nie może wystąpić w danym miejscu. Wtedy drugi plus będziemy wpisywali „na pewno”, bez żadnego znaku zapytania. Dopóki jednak nie możemy żadnego z plusów wykluczyć, pytajniki muszą pozostać.



DO ZAPAMIĘTANIA:

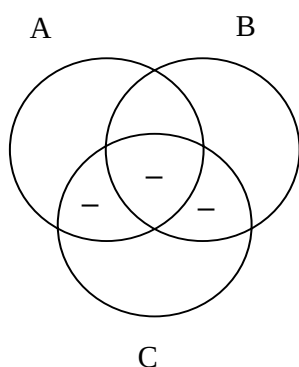
Minusy są zawsze „pewne”. Wynika to z tego, że jeśli pusty jest jakiś obszar składający się z kilku części, to oczywiście pusta musi być każda z tych części; w każdą z nich możemy zatem wpisać minus.

Jeśli natomiast wiemy tylko, że w jakimś obszarze składającym się z kilku części coś się znajduje, to wcale nie daje nam to pewności, w której z tych części postawić plus. Jakiś element znajdować się może w dowolnej z nich.

Sytuację tę można przedstawić bardziej obrazowo. Jeśli wiemy, że w mieszkaniu składającym się z kilku pomieszczeń nie ma nikogo, to wiemy, że na pewno pusty jest pokój, kuchnia, łazienka itd. Jeśli natomiast dowiadujemy się, że w mieszkaniu tym ktoś jest, to informacja ta nie daje nam pewności, w którym pomieszczeniu osoba ta się znajduje. Być może pusta jest kuchnia i łazienka, a człowiek, o którym mowa, jest w pokoju, ale może też być zupełnie inaczej.

$$A \cup B \subseteq C'$$

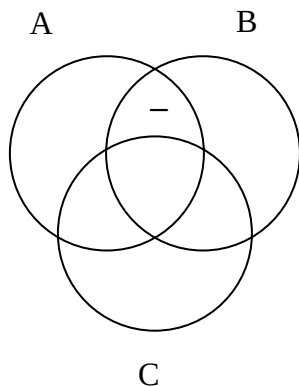
To, że suma zbiorów A i B zawiera się w dopełnieniu C, oznacza, iż suma A i B znajduje się poza C, a zatem żadna część tej sumy nie może znajdować się w C. Musimy więc wpisać minusy we wszystkich częściach zbiorów A oraz B leżących jednocześnie w C.



$$A \cup B \subseteq C'$$

$$A \cap B \subseteq C$$

To, że iloczyn A i B zawiera się w C, oznacza, że żadna część tego iloczynu (czyli części wspólnej A i B) nie może znajdować się poza C. W praktyce daje to tylko jeden minus w „górnej” części iloczynu A i B.



$$A \cap B \subseteq C$$

Zobrazowane wyżej zależności pomiędzy zbiorami nie wyczerpują oczywiście wszystkich możliwych przypadków, jakie mogą pojawić się w zadaniach. Jednakże powinny one pozwolić na zrozumienie istoty posługiwania się diagramami i umożliwić każdemu samodzielne zaznaczenie nawet takich stosunków pomiędzy zbiorami, z jakimi się nigdy nie zetknął.

5.5.2. PRAKTYKA: SPRAWDZANIE PRAW TEORII ZBIORÓW PRZY POMOCY DIAGRAMÓW VENNA.

Wyrażenia, które będziemy obecnie badali pod kątem tego, czy stanowią one prawa rachunku zbiorów, będą miały formę implikacji. Wykazanie, że implikacja taka stanowi ogólne prawo, będzie polegało na pokazaniu, że jest ona zawsze prawdziwa. Ponieważ implikacja, zgodnie z tabelkami zero-jedynkowymi, fałszywa jest tylko w jednym przypadku – gdy jej poprzednik jest prawdziwy a następnik fałszywy, to udowodnienie, że jest ona zawsze prawdziwa, polegać może na wykazaniu niemożliwości zajścia takiej sytuacji.

W praktyce będzie to wyglądało tak, że będziemy wpisywali do diagramu to, co mówi poprzednik implikacji, a następnie sprawdzali, czy gwarantuje nam to prawdziwość następnika. Jeśli bowiem wypełnienie diagramu według poprzednika implikacji zagwarantuje nam prawdziwość jej następnika, będzie stanowiło to dowód, że nie jest możliwa sytuacja, aby poprzednik był prawdziwy i następnik fałszywy, a zatem implikacja jest zawsze prawdziwa; przedstawia ona prawo rachunku zbiorów. Jeśli natomiast na diagramie będzie się dało przedstawić fałszywość następnika, będzie to świadczyć, że implikacja może być fałszywa, a więc nie opisuje ona ogólnie obowiązującego prawa.

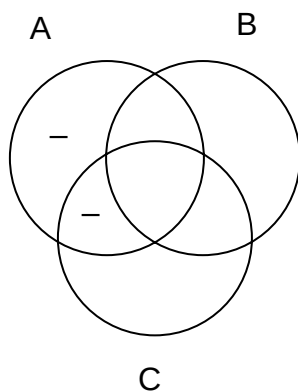
Cała procedura w skrócie:

- wpisujemy do diagramu wszystkie informacje z poprzednika implikacji,
- **nie wpisujemy** informacji z następnika implikacji, a jedynie sprawdzamy, czy wykonany według poprzednika rysunek daje nam gwarancję ich prawdziwości,
- jeśli rysunek gwarantuje nam prawdziwość następnika, oznacza to, że badane wyrażenie jest prawem rachunku zbiorów; jeśli nie mamy takiej gwarancji (diagram da się wypełnić tak, aby następnik był fałszywy), wyrażenie nie jest prawem.

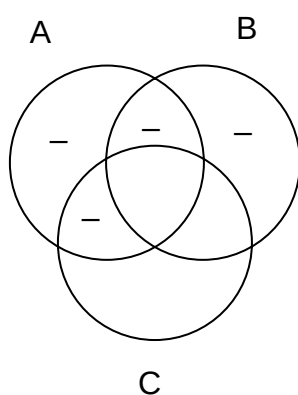
Przykład:

Sprawdzimy, czy prawem rachunku zbiorów jest wyrażenie: $(A \subseteq B \wedge B \subseteq C) \rightarrow A \subseteq C$.

W poprzedniku naszej implikacji mamy dwie zależności. Najpierw zaznaczymy to, że zbiór A zawiera się w B (czyli A nie może „wystawać” poza B):



Teraz dopiszemy jeszcze, że B zawiera się w C:



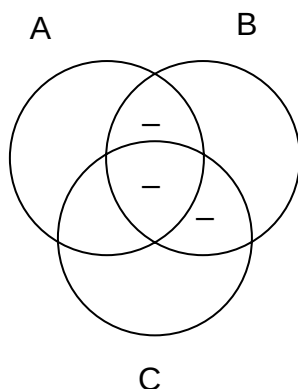
Po zaznaczeniu na diagramie wszystkich informacji zawartych w poprzedniku musimy sprawdzić, czy gwarantuje nam to prawdziwość następnika. Widzimy, że faktycznie następnik musi być prawdziwy. Mamy pewność, że zbiór A zawiera się w C, gdyż w odpowiednim obszarze mamy minusy świadczące, że A nie może „wystawać” poza C; są to dwa minusy w „górnej” części zbioru A.

Skoro w badanym wyrażeniu, mającym postać implikacji, nie jest możliwe aby poprzednik był prawdziwy, a następnik fałszywy (mówiąc inaczej, prawdziwość poprzednika gwarantuje prawdziwość następnika), oznacza to, że wyrażenie to jest prawem rachunku zbiorów.

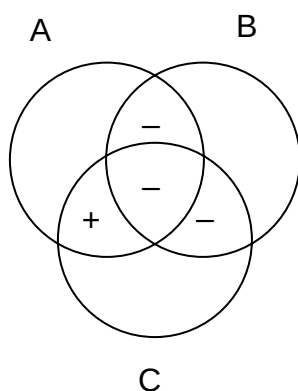
Przykład:

Sprawdzimy, czy prawem rachunku zbiorów jest wyrażenie: $(A) \cap (B \cap B) \cap (C) \rightarrow A \cap (C)$.

Po zaznaczeniu na diagramie informacji z obu członów poprzednika otrzymujemy rysunek:

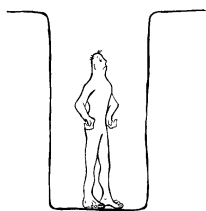


Rysunek ten nie daje nam jednak gwarancji, że następnik jest prawdziwy. Aby mieć pewność, że zbiory A i C są rozłączne, musielibyśmy mieć minusy w całym obszarze wspólnym tych zbiorów. Tymczasem w jednej części tego obszaru nie ma żadnego znaku. Oznacza to, że nic nie stoi na przeszkodzie, aby coś tam mogło się znajdować. Poniższy rysunek pokazuje wyraźnie, że da się zaznaczyć na diagramie prawdziwość poprzednika implikacji i jednocześnie fałszywość jej następnika.



Badane wyrażenie nie jest zatem prawem rachunku zbiorów.

Powyższy rysunek stanowi graficzny kontrprzykład, pokazujący, że badana implikacja nie jest zawsze prawdziwa. Możemy również podać kontrprzykład „materialny” to znaczy wymyślić takie zbiory A, B i C, aby pokazać, że badane wyrażenie może być fałszywe. Mogą być to przykładowo takie zbiory: A – zbiór drzew liściastych, B – zbiór drzew iglastych, C – zbiór dębów. Prawdą jest, że $A \cap (B \cup C) = \emptyset$, natomiast A i C wcale rozłączne nie są.



5.5.3. UTRUDNIENIA I PUŁAPKI.

Oczywiście nie zawsze badane wyrażenia są tak łatwe do zaznaczenia na diagramie jak w dwóch rozważanych dotąd zadaniach. Poniżej omówimy kilka przykładów nieco bardziej skomplikowanych.

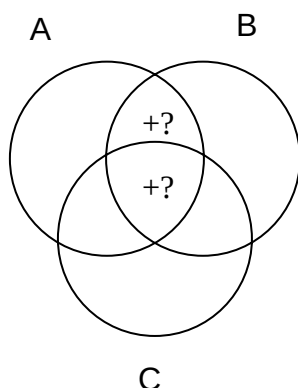
Czy tam ma być plus czy minus?

Przykład:

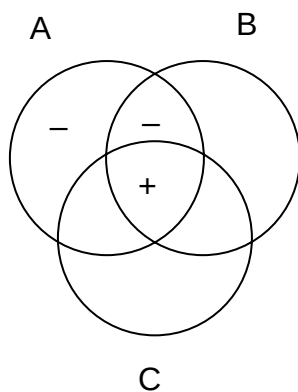
Sprawdźmy, czy jest prawem rachunku zbiorów wyrażenie:

$$(B \cap A \neq \emptyset \wedge A \cap C' = \emptyset) \rightarrow B \cap C \neq \emptyset$$

Pierwszy człon poprzednika implikacji informuje nas, że coś się musi znajdować w części wspólnej (iloczynie) zbiorów B i A. Ponieważ obszar ten składa się z dwóch kawałków, nie wiemy dokładnie, w którym z nich jakiś element się znajduje; być może w obydwu, ale może tylko w jednym z nich. Dlatego też możemy wpisać tu jedynie plusy ze znakiem zapytania.



Drugi człon poprzednika implikacji informuje nas, że pusty musi obszar wspólny A oraz C', czyli ta część zbioru A, która znajduje się poza zbiorem C. Na naszym rysunku są to dwa „górne” kawałki zbioru A. Widzimy, że w jednej z tych części znajduje się znak „+?”. Ponieważ jednak znak zapytania świadczy, że coś w tym obszarze może się znajdować, ale nie jest to konieczne, a teraz otrzymujemy informację, że na pewno nic tam nie ma, to wynikający stąd minus jest „silniejszy” od plusa ze znakiem zapytania i dlatego właśnie „–” powinien się tam ostatecznie znaleźć. Jeśli jednak jeden ze znaków „+?” zamienimy na minus, to wynika stąd, że drugi z plusów staje się „pewny” i widniejący przy nim pytajnik powinniśmy zlikwidować. Skoro bowiem dotąd wiedzieliśmy, że w jednym z dwóch obszarów coś jest, lecz nie mieliśmy pewności w którym, a teraz dowiedzieliśmy, że pierwszy jest pusty, to w takim razie na pewno wypełniony musi być obszar drugi. A zatem, po wpisaniu całego poprzednika implikacji, diagram będzie się przedstawiał następująco:



Teraz musimy sprawdzić, czy tak wykonany rysunek daje nam gwarancję prawdziwości następnika implikacji, a więc czy na pewno $B \cap C \neq \emptyset$. W jednym kawałku części wspólnej zbiorów B i C mamy plus, który daje nam gwarancję, że obszar ten z pewnością nie jest pusty. Badane wyrażenie jest zatem prawem rachunku zbiorów.

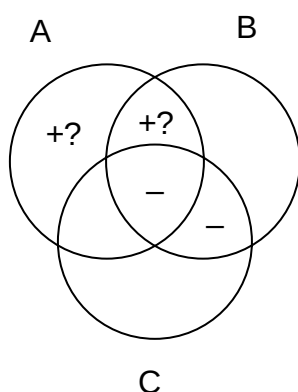
Plus ze znakiem zapytania nie daje pewności!

Przykład:

Sprawdźmy, czy jest prawem rachunku zbiorów wyrażenie:

$$(B \subseteq C' \wedge A - C \neq \emptyset) \rightarrow C' \cap B \neq \emptyset$$

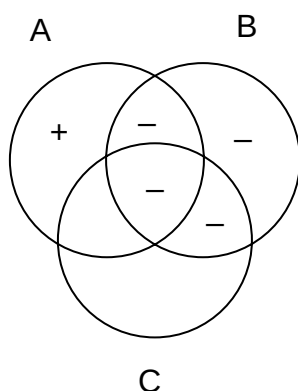
Fakt, że zbiór B zawiera się w dopełnieniu C, oznacza, że cały B znajduje się poza C, czyli żadna część B nie może znajdować się w C; zbiory te są prostu rozłączne. W całej części wspólnej B i C musimy zatem wpisać minusy. Jeśli $A - C$ ma być niepuste, to coś musi znajdować się w obszarze zbioru A leżącym poza C. Cały czas mamy jednak do wyboru dwie części tego obszaru i nie wiemy do końca, w którą z nich wpisać „+”. Wypełniony według poprzednika implikacji diagram wygląda zatem następująco:



Musimy teraz sprawdzić, czy powyższy rysunek gwarantuje nam, że $C' \cap B \neq \emptyset$. Część wspólna dopełnienia C oraz B to obszar zbioru B leżący poza C, czyli „górny” półksiężyc zbioru B. W jednej części tego obszaru znajduje się wprowadzie plus, jednak jest on z

pytajnikiem, co oznacza, iż nie mamy gwarancji, że jest on tam na pewno. Nie mamy zatem pewności, że następnik badanej implikacji jest prawdziwy, a więc nie jest ona prawem rachunku zbiorów.

Rysunek pokazujący, że pomimo prawdziwości poprzednika, następnik implikacji może być fałszywy, wygląda następująco:



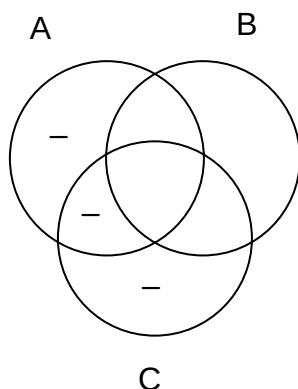
Zależności trudniejsze do zaznaczenia na diagramie.

Przykład:

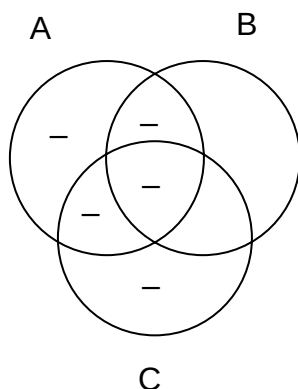
Sprawdźmy, czy jest prawem rachunku zbiorów wyrażenie:

$$[A \cup C \subseteq B \wedge A \subseteq (B \cup C)] \rightarrow B - C = \emptyset$$

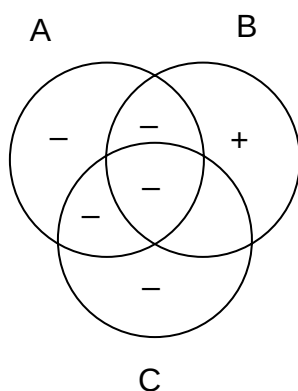
W powyższym przykładzie pewne trudności sprawić może właściwe zaznaczenie na diagramie informacji z poprzednika implikacji. Fakt, że suma zbiorów A i C zawiera się w B oznacza, że żadna część A oraz żadna część C nie może znajdować się poza zbiorem B. We wszystkich fragmentach zbiorów A i C leżących poza B musimy więc wpisać minusy.



Z kolei to, że A zawiera się w dopełnieniu sumy B i C znaczy, że żadna część zbioru A nie może znajdować się w zbiorze B lub C. W związku z tym w częściach wspólnych A i C oraz A i B musimy wpisać minusy. W jednym fragmencie wymienionego obszaru minus już się znajduje, zatem dodajemy jeszcze dwa:



Teraz musimy sprawdzić, czy mamy pewność, że obszar zbioru B leżący poza C (czyli $B - C$) jest pusty. W jednej części tego obszaru mamy znak „-”, w drugiej natomiast nie ma nic. To, że nie mamy tam wpisanego żadnego symbolu, nie oznacza jednak, że nic tam nie ma, a jedynie, że nie mamy na temat tej części żadnych informacji. Pewność, że obszar jest pusty, mielibyśmy jedynie wtedy, gdyby umieszczony był w nim minus. Tymczasem nic nie stoi na przeszkodzie, aby w wolne miejsce wpisać plus:



Ponieważ, jak widać na powyższym rysunku, da się wypełnić diagram w ten sposób, aby poprzednik implikacji był prawdziwy, a następnik fałszywy, badane wyrażenie nie jest prawem rachunku zbiorów.

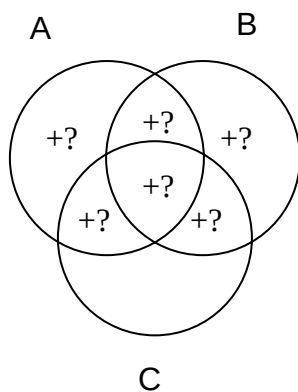
Czasem trzeba zacząć od końca.

Przykład:

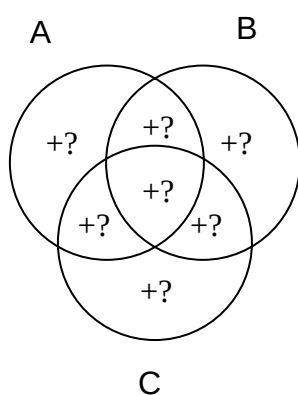
Sprawdźmy, czy jest prawem rachunku zbiorów wyrażenie:

$$(A \cup B \neq \emptyset \wedge B \cup C \neq \emptyset) \rightarrow A \cup C \neq \emptyset$$

Fakt, że nie jest pusta suma zbiorów A i B, oznacza, że coś musi znajdować się w którejkolwiek części obszaru składającego się aż z sześciu części:

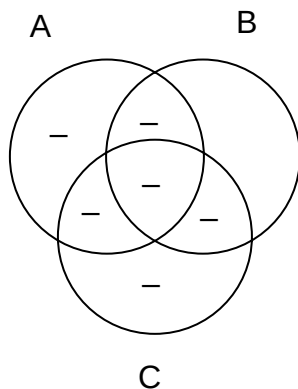


Gdy dodamy to tego informację, że niepusta jest również suma B i C otrzymamy rysunek:

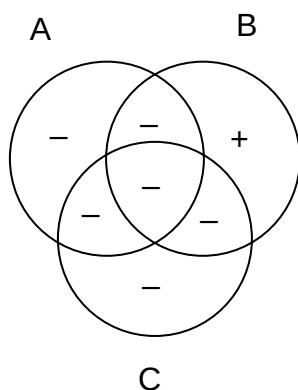


Pozostaje nam teraz odpowiedzieć na pytanie, czy mamy pewność, że coś znajduje się w którejkolwiek części sumy zbiorów A oraz C. Udzielić jednoznacznej odpowiedzi na to pytanie przy pomocy powyższego rysunku może wydawać się trudne – w wymienionej części znajduje się wprawdzie sześć plusów, ale wszystkie z pytajnikiem. W takiej sytuacji możemy spróbować rozwiązać zadanie niejako od drugiej strony, zaczynając od budowy kontrprzykładu. Zobaczmy, czy da się stworzyć rysunek, na którym następnik implikacji byłby fałszywy, a potem sprawdzimy, czy poprzednik może być wtedy równocześnie prawdziwy.

Fałszywość następnika naszego wyrażenia oznacza, że pusta jest suma zbiorów A i C, czyli:



Czy możemy teraz sprawić, aby prawdziwe były oba człony poprzednika implikacji? Stanie się tak, gdy wpiszemy znak „+” w jedyne wolne pole:



Powyższy rysunek pokazuje, że można zaznaczyć na diagramie jednocześnie prawdziwość poprzednika implikacji (coś znajduje się zarówno w sumie zbiorów A i B jak i w sumie B i C), jak i fałszywość jej następnika (pusta jest suma A i C). Badane wyrażenie nie jest więc prawem rachunku zbiorów.

SŁOWNICZEK:

Dopełnienie zbioru – dopełnienie zbioru A to zbiór zawierający te elementy uniwersum, które nie należą do A.

Identyczność zbiorów – zbiory A i B są identyczne ($A = B$), gdy mają dokładnie te same elementy.

Iloczyn zbiorów – iloczyn (przekrój) zbiorów A i B ($A \cap B$) to zbiór zawierający elementy należące zarówno do A jak i do B.

Inkluzja (zawieranie się zbiorów) – zbiór A zawiera się w zbiorze B ($A \subseteq B$), gdy każdy element A jest elementem B.

Krzyżowanie zbiorów – zbiory A i B się krzyżują ($A \# B$), gdy mają one wspólne elementy, ale równocześnie do każdego z nich należą takie elementy, które nie należą do drugiego.

Rozłączność zbiorów – zbiory A i B są rozłączne ($A \nmid B$), gdy nie mając ani jednego wspólnego elementu.

Różnica zbiorów – różnica zbiorów A i B to zbiór zawierający te elementy A , które nie należą do B .

Suma zbiorów – suma zbiorów A i B ($A \cup B$) to zbiór powstały z połączenia elementów A i B .

Zbiór pusty – zbiór nie zawierający żadnego elementu. Zbiór pusty oznaczamy symbolem $\emptyset$.

Rozdział VI

RELACJE.

WSTĘP.

Obecny rozdział poświęcony będzie relacjom. Z relacjami zetknęliśmy się już w części poświęconej rachunkowi predykatów. Obecnie zostaną one omówione o wiele dokładniej. Będzie to rozdział najbardziej teoretyczny ze wszystkich; zadania będą stanowiły niewielki procent całości. Wynika to z faktu, iż związane z relacjami zadania polegają zwykle na wykrywaniu pewnych własności podanych relacji. Aby móc je rozwiązać, trzeba przede wszystkim posiadać teoretyczną wiedzę o tych własnościach. Gdy wiedza ta zostanie zdobyta, rozwiązanie takiego zadania jest zwykle niemal oczywiste.

6.1. CO TO JEST RELACJA.

6.1.1. ŁYK TEORII.



ŁYK TEORII

Relacja to pewien związek łączący obiekty. Mówiąc „relacja” mamy zwykle na myśli tak zwaną relację dwuczłonową, czyli związek łączący dwa obiekty. Taką relacją jest na przykład bycie starszym – pewna osoba x jest starsza od innej osoby y ; inne przykłady to bycie żoną – osoba x jest żoną osoby y , lubienie – osoba x lubi osobą y itp.

Dla porządku dodajmy, że relacje mogą mieć dowolną ilość członów. Relacje jednoczłonowe (odnoszące się do jednego obiektu) nazywamy własnościami – tego typu relacje, to na przykład bycie wysokim, bycie w wieku 25 lat, bycie kobietą itp. Przykładem własności trójczłonowej jest słuchanie rozmowy – osoba x słucha rozmowy osoby y z osobą z . Relacjami innymi niż dwuczłonowe nie będziemy się jednak zajmować; mówiąc relacja – bez dodania do niej żadnego przymiotnika, będziemy mieli na myśli zawsze relację dwuczłonową.

Symbolicznie relacje możemy oznaczać na różne sposoby. Zwykle fakt, że dwa obiekty x i y są ze sobą w relacji R zapisujemy $R(x,y)$ lub xRy . Spotyka się też zapis $(x,y) \in R$ (para x, y należy do relacji R).

Do lepszego zrozumienia relacji potrzebne nam będzie pojęcie tak zwanej pary uporządkowanej oraz iloczynu kartezjańskiego dwóch zbiorów.

Para uporządkowana.

Jak pamiętamy z poprzedniego rozdziału, w zwykłym zbiorze nie jest istotna kolejność elementów, w jakiej je wypiszemy. I tak na przykład zbiór $X = \{a, b\}$ jest równy zbiorowi $Y = \{b, a\}$. Inaczej ma się rzecz w przypadku tak zwanych par uporządkowanych, w skrócie zwanych po prostu parami. Elementy par wypisujemy w nawiasach skośnych, np. $\langle a, b \rangle$ lub, czasem, zwykłych – (a, b) . W przypadku pary kolejność elementów ma kluczowe znaczenie. I tak para $\langle a, b \rangle$ nie jest równa parze $\langle b, a \rangle$; są to zupełnie różne obiekty.

Iloczyn kartezjański.

Iloczyn kartezjański to pewne specyficzne działanie na zbiorach, o którym jednak nie było mowy w rozdziale poświęconym zbiorom. Iloczyn kartezjański symbolicznie oznaczamy znakiem $\times$. Zbiór, który powstaje w wyniku wykonania takiego działania, nie jest zwykłym zbiorem, ale zbiorem, którego elementy stanowią pary. Dokładniej, iloczyn kartezjański zbiorów X i Y (czyli $X \times Y$) to zbiór wszystkich par, takich, w których na pierwszym miejscu jest element zbioru X , a na drugim element zbioru Y . Przykładowo, jeśli $X = \{a, b, c\}$, natomiast $Y = \{1, 2\}$, to iloczyn kartezjański $X \times Y = \{\langle a, 1 \rangle, \langle a, 2 \rangle, \langle b, 1 \rangle, \langle b, 2 \rangle, \langle c, 1 \rangle, \langle c, 2 \rangle\}$.

Kwadrat kartezjański jakiegoś zbioru X , oznaczany symbolicznie X^2 , to nic innego, jak iloczyn kartezjański zbioru X z sobą samym, czyli $X \times X$. Jeśli zatem $X = \{a, b, c\}$, to $X^2 = \{\langle a, a \rangle, \langle a, b \rangle, \langle a, c \rangle, \langle b, a \rangle, \langle b, b \rangle, \langle b, c \rangle, \langle c, a \rangle, \langle c, b \rangle, \langle c, c \rangle\}$

Pojęcia pary uporządkowanej oraz iloczynu kartezjańskiego łączy się z teorią relacji w ten sposób, że każdą relację możemy przedstawić (przynajmniej teoretycznie) jako zbiór par. Jeśli relacja określona jest w pewnym uniwersum, to możemy powiedzieć, że relacja ta zawiera się w kwadracie kartezjańskim tego uniwersum (stanowi podzbiór kwadratu kartezjańskiego tego uniwersum). Mówiąc po prostu, relacja to niektóre (a czasem wszystkie) pary, jakie można utworzyć z elementów uniwersum.

Najlepiej wyjaśnić to na przykładzie. Weźmy uniwersum złożone z czterech liczb $U = \{1, 2, 3, 4\}$ i określmy w tym zbiorze relację większości. Bardziej formalnie relację tę możemy zapisać tak: $xRy \equiv x > y$. Relacja nasza zawiera się w kwadracie kartezjańskim uniwersum (symbolicznie $R \subseteq U^2$), ponieważ należą do niej niektóre z par liczb, które to pary możemy utworzyć z uniwersum. Relację naszą możemy przedstawić jako następujący zbiór par, w których pierwsza liczba jest większa od drugiej: $R = \{\langle 2,1 \rangle, \langle 3,1 \rangle, \langle 3,2 \rangle, \langle 4,1 \rangle, \langle 4,2 \rangle, \langle 4,3 \rangle\}$.

Gdybyśmy w uniwersum złożonym z ludzi chcieli utworzyć relację bycia żoną, to relację tę moglibyśmy przedstawić jako zbiór takich par, gdzie pierwsza osoba jest żoną drugiej osoby: $R = \{\langle \text{Maria Kowalska}, \text{Jan Kowalski} \rangle, \langle \text{Danuta Wałęsa}, \text{Lech Wałęsa} \rangle, \langle \text{Hilary Clinton}, \text{Bill Clinton} \rangle \dots \text{ itd. } \}$. Oczywiście wypisanie wszystkich par należących do naszej relacji nie jest praktycznie możliwe, jednak nie ulega wątpliwości, że jest to podzbiór kwadratu kartezjańskiego zbioru ludzi, czyli $R \subseteq U^2$.

6.2. DZIEDZINY I POLE RELACJI.

6.2.1. ŁYK TEORII.



ŁYK TEORII

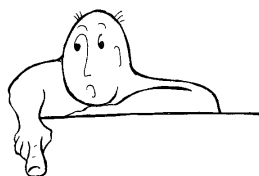
W każdej relacji możemy określić tak zwaną dziedzinę lewostronną, nazywaną czasem po prostu dziedziną, dziedzinę prawostronną, nazywaną również przeciwdziedziną oraz pole. Dziedzinę lewostronną relacji R oznaczamy symbolicznie $D_L(R)$, dziedzinę prawostronną – $D_P(R)$, natomiast pole – $P(R)$.

Dziedzina lewostronna relacji R , to zbiór takich przedmiotów, które pozostają w R do jakiegoś (przynajmniej jednego) przedmiotu. Symbolicznie możemy to zapisać: $D_L(R) = \{x: \exists y (xRy)\}$ (dziedzina lewostronna relacji R to zbiór takich x , w stosunku do których istnieje jakiś y , taki że x jest w relacji R do tego y). W praktyce możemy sobie bardzo łatwo uzmysłować, co jest dziedziną danej relacji, wypisując (lub wyobrażając sobie) pary tworzące tę relację. Dziedzinę lewostronną stanowić będzie zawsze zbiór tych obiektów, które przynajmniej raz znalazły się na pierwszym miejscu w jakiejś parze. Gdy weźmiemy, wspomnianą wcześniej relację większości określoną w zbiorze $U = \{1, 2, 3, 4\}$, to po przedstawieniu tej relacji jako zbioru

par: $R = \{\langle 2,1 \rangle, \langle 3,1 \rangle, \langle 3,2 \rangle, \langle 4,1 \rangle, \langle 4,2 \rangle, \langle 4,3 \rangle\}$, łatwo zauważymy, że $D_L(R) = \{2, 3, 4\}$. W przypadku relacji bycia żoną dziedzinę lewostronną stanowiłby zbiór kobiet zamężnych.

Dziedzina prawostronna (przeciwdziedzina) relacji R to, jak łatwo się domyślić, zbiór tych przedmiotów, do których jakiś przedmiot pozostaje w R . Symbolicznie: $D_P(R) = \{y: \exists x (xRy)\}$. W przypadku naszej relacji większości $D_P(R) = \{1, 2, 3\}$, natomiast przeciwdziedzinę relacji bycia żoną stanowiłby (jeśli ograniczymy się do małżeństw heteroseksualnych) zbiór żonatych mężczyzn.

Pole relacji to po prostu suma dziedziny lewej i prawej. Symbolicznie $P(R) = D_P(R) \cup D_L(R)$. W naszej relacji większości $P(R) = \{1, 2, 3, 4\}$. W tym przypadku pole pokryło się z uniwersum, jednak nie jest to wcale konieczne. Widać to na przykładzie relacji bycia żoną, gdzie pole to zbiór ludzi pozostających w związkach małżeńskich (będących żoną lub mających żonę), a więc nie całe uniwersum.



Uwaga na błędy!

Za błąd może zostać uznane powiedzenie, że polem relacji bycia żoną jest zbiór małżeństw. Zbiór małżeństw to bowiem zbiór, którego elementami są małżeństwa, a nie pojedyncze osoby (ma on w przybliżeniu dwa razy mniej elementów niż zbiór osób pozostających w związkach małżeńskich). Natomiast pole relacji bycia żoną musi być zbiorem złożonym z osób.

6.2.2. PRAKTYKA: OKREŚLANIE DZIEDZIN I POLA RELACJI.

Zadania związane z dziedzinami i polem relacji polegają na określeniu tych własności dla zadanej relacji. Rozwiązywanie takich przykładów nie jest trudne, jeśli tylko pamiętamy, że każdą relację możemy, przynajmniej teoretycznie przedstawić jako zbiór par. Dziedzina lewostronna będzie każdorazowo zbiorem tych elementów, które przynajmniej raz znalazły się w naszych parach na pierwszym miejscu, natomiast dziedzina prawostronna, zbiorem elementów, które przynajmniej raz wystąpiły na drugim miejscu. Po określeniu dziedziny lewej i prawej, wyznaczenie pola jest już banalnie proste.

Przykład:

Określimy dziedzinę, przeciwdziedzinę i pole relacji bycia matką ($xRy \equiv x$ jest matką y) określonej w zbiorze wszystkich ludzi (żyjących kiedykolwiek, a nie tylko aktualnie).

Gdybyśmy chcieli przedstawić naszą relację w postaci zbioru par, to na pierwszym miejscu byłaby każdorazowo kobieta posiadająca przynajmniej jedno dziecko, natomiast na drugim osoba będąca dzieckiem tej kobiety. Oczywiście więc jest, że dziedzinę lewostronną naszej relacji stanowić będzie zbiór kobiet mających dzieci. Dziedzina prawa to zbiór osób, które mają matkę. Ponieważ nasze uniwersum stanowi zbiór wszystkich ludzi kiedykolwiek żyjących, to o każdym człowieku można powiedzieć, że ma on (aktualnie lub kiedyś żyjącą) matkę; każdy więc znajdzie się jako element jakiejś pary z prawej strony. A zatem przeciwdziedzina naszej relacji to zbiór wszystkich ludzi. Skoro jedna z dziedzin stanowi już całe uniwersum, to oczywiście jest, że również pole naszej relacji będzie równe uniwersum, czyli zbiorowi wszystkich ludzi.

Przykład:

Określimy dziedziny i pole określonej w zbiorze liczb naturalnych relacji bycia dwukrotnością ($xRy \equiv x$ jest dwukrotnością y).

Do naszej relacji należeć będą takie pary złożone z liczb naturalnych, gdzie pierwsza liczba będzie dwukrotnością drugiej, a zatem $R = \{\langle 2, 1 \rangle, \langle 4, 2 \rangle, \langle 6, 3 \rangle, \langle 8, 4 \rangle, \dots\}$. Po wypisaniu kilku przykładowych par, widać jasno, że dziedzina lewa relacji, to zbiór liczb parzystych, natomiast dziedzina prawa (i jednocześnie pole) to zbiór wszystkich liczb naturalnych, czyli uniwersum.

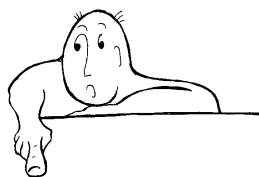
Przykład:

Określimy dziedziny i pole określonej w zbiorze wszystkich ludzi relacji bycia przeciwnej płci ($xRy \equiv x$ jest przeciwnej płci niż y).

Gdybyśmy chcieli wypisać niektóre z par należących do naszej relacji otrzymalibyśmy $R = \langle \text{Jan}, \text{Maria} \rangle, \langle \text{Maria}, \text{Mieczysław} \rangle, \langle \text{Karolina}, \text{Zenon} \rangle, \langle \text{Zenon}, \text{Karolina} \rangle, \langle \text{Zenon}, \text{Maria} \rangle, \dots$

Widać wyraźnie, że każdy człowiek znajdzie się (wielokrotnie) zarówno z lewej strony jakiejś pary, jak i z prawej strony; do każdego można bowiem dobrać kogoś przeciwnej płci.

A zatem w tym przypadku dziedzina prawa, równa się dziedzinie lewej, równa się polu relacji i stanowi całe uniwersum, czyli zbiór wszystkich ludzi.



Uwaga na błędy!

W przypadku powyższej relacji częstymi odpowiedziami na pytanie o którąś z dziedzin są dość dziwnie brzmiące stwierdzenia na przykład: „ludzie przeciwnej płci”, „ludzie określonej płci”, czy też „ludzie jednej płci”. Nie są to jednak dobre odpowiedzi – cóż to bowiem są na przykład „ludzie przeciwnej płci”, jaki dokładnie jest to zbiór?

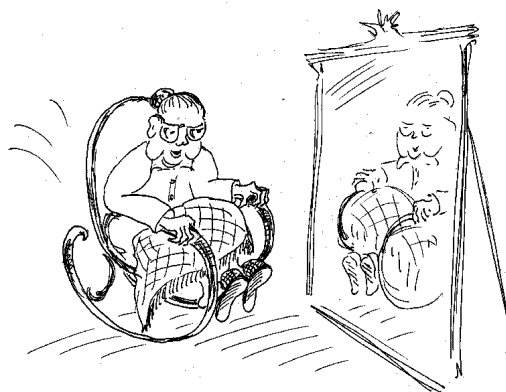
Przykład:

Określimy dziedziny i pole określonej w zbiorze wszystkich ludzi relacji bycia w tym samym wieku ($xRy \equiv x$ jest w tym samym wieku co y).

Gdybyśmy wypisali pary należące do powyższej relacji, łatwo zauważylibyśmy, że do człowieka mającego np. 20 lat łatwo dobrać kogoś będącego w tym samym wieku; podobnie w stosunku do kogoś mającego np. 15 lat, 23 lata, 35 lat, 78 lat itd. Wątpliwości może budzić fakt, czy jesteśmy w stanie stworzyć parę z kimś mającym przykładowo 108 lat, zakładając że jest to jedyny człowiek na świecie w tym wieku.

Otóż zawsze możemy to uczynić, tworząc parę złożoną z tego człowieka występującego zarówno na pierwszym miejscu, jak i na drugim; w przypadku tej relacji bowiem każdy, oprócz możliwości bycia w niej w stosunku do innych osób, pozostaje w niej również do samego siebie. Każdy jest bowiem w tym samym wieku, w którym jest on sam. Każdy więc, przynajmniej w tym jednym przypadku, wystąpi zarówno na pierwszym, jak i na drugim miejscu w pewnej parze.

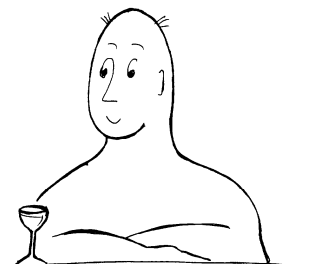
Podobnie jak w poprzednim przykładzie, dziedzina lewa naszej relacji równa się dziedzinie prawej, równa się polu i stanowi całe uniwersum, czyli zbiór wszystkich ludzi.



RELACJA BYCIA W TYM SAMYM WIEKU

6.3. WŁASNOŚCI FORMALNE RELACJI.

6.3.1. ŁYK TEORII.



ŁYK TEORII

Relacje możemy charakteryzować pod względem pewnych własności. Obecnie omówimy najważniejsze z tych własności, grupując je w związku z istotnymi dla nich pojęciami.

Uwaga na marginesie.

Omawiane własności relacji dotyczą zawsze jakiegoś konkretnego uniwersum. Relacja posiadająca daną własność w jednym uniwersum, może nie posiadać jej w innym. Dlatego, ściśle rzecz biorąc, w poniższych wzorach wyrażenia $\forall x$ (dla każdego x) powinny przybierać formę $\forall x \in U$ (dla każdego x należącego do danego uniwersum); podobnie $\exists x$ (istnieje takie x) – $\exists x \in U$ (istnieje takie x należące do U). Aby zbytnio wzorów nie komplikować, nie będziemy tak jednak pisać, domyślnie przyjmując, że każdorazowo chodzi nam jedynie o elementy z danego uniwersum.

Zwrotność.

Mówimy, że relacja jest **zwrotna**, gdy każdy element uniwersum jest w tej relacji do siebie samego. Symbolicznie:

R jest zwrotna $\equiv \forall x (xRx)$

Przykładem relacji zwrotnej jest bycie w takim samym wieku (w zbiorze ludzi) lub bycie sobie równą (w zbiorze liczb). Każdy człowiek jest bowiem w takim samym wieku w stosunku do siebie samego, a każda liczba jest równa sobie samej.

Relacja jest **przeciwzwrotna**, gdy żaden element uniwersum nie jest w relacji do siebie samego. Symbolicznie:

R jest przeciwzwrotna $\equiv \forall x (\sim (xRx))$

Przeciwzwrotna jest przykładowo relacja bycia matką w zbiorze ludzi lub relacja mniejszości w zbiorze liczb.

Relacja może być **ani zwrotna, ani przeciwzwrotna**. Oznacza to, że są w uniwersum elementy będące w relacji do siebie samego, ale są też i takie, które do siebie samego w niej nie są. Relację taką nazywamy czasem **niezwrotną**. Symbolicznie:

R nie jest zwrotna ani przeciwwrotna $\equiv \exists x (xRx) \wedge \exists x \sim (xRx)$

Przykładem takiej relacji może być relacja kochania – są ludzie, którzy kochają samych siebie, a są też i tacy, którzy siebie nie kochają.

Symetria.

Mówimy, że relacja jest **symetryczna**, gdy jest tak, że jeśli relacja zachodzi pomiędzy dwoma elementami w jedną stronę, to zachodzi i w drugą (jeśli zachodzi pomiędzy x i y, to zachodzi też pomiędzy y i x). Symbolicznie:

R jest symetryczna $\equiv \forall x \forall y (xRy \rightarrow yRx)$

Symetryczną jest na przykład relacja bycia tej samej płci – jeśli osoba x jest tej samej płci, co osoba y, to również osoba y jest na pewno tej samej płci co osoba x.

Relacja jest **asymetryczna (antysymetryczna, przeciwsymetryczna)**, gdy jest tak, że jeśli zachodzi w jedną stronę, to nie zachodzi w drugą. Symbolicznie:

R jest asymetryczna $\equiv \forall x \forall y [xRy \rightarrow \sim (yRx)]$

Asymetryczna jest na przykład relacja bycia ojcem – jeśli jedna osoba jest ojcem drugiej, to druga na pewno nie jest ojcem pierwszej.

Relacja jest **słabo asymetryczna (słabo antisymetryczna)** gdy dla wszystkich różnych od siebie elementów uniwersum jest tak, że jeśli relacja zachodzi w jedną stronę, to nie zachodzi w drugą. Symbolicznie:

R jest słabo asymetryczna $\equiv \forall x \forall y [(x \neq y \wedge xRy) \rightarrow \sim (yRx)]$

Relacją słabo asymetryczną jest na przykład relacja „ $\geq$ ” w zbiorze liczb. Gdy weźmiemy bowiem dwie różne od siebie liczby i nasza relacja zachodzi między nimi w jedną stronę, to na pewno nie zachodzi między nimi w drugą.

Odróżnienie „mocnej” asymetrii od słabej jest dla niektórych dość trudne. Można sobie tę różnicę zapamiętać w taki praktyczny sposób: przy zwykłej („mocnej”) asymetrii żadne elementy nie mogą być w relacji do siebie samego (relacja taka musi być jednocześnie przeciwwrotna), natomiast słaba asymetria tym tylko różni się od zwykłej, że w jej przypadku niektóre (bądź wszystkie) elementy mogą być w relacji do siebie samego.

Uwaga na marginesie.

Odnośnie nazewnictwa własności związanych z symetrią w wielu podręcznikach panuje zamieszanie. To co u nas określone zostało jako asymetria w innych nazywane jest przeciwsymetrią lub antisymetrią; nasza słaba asymetria występuje natomiast jako słaba, ale również jako „zwykła” (bez żadnego przymiotnika), antisymetria.

Dlatego też, w celu uniknięcia nieporozumień dobrze jest zawsze sprawdzić, w jaki sposób autor danego podręcznika bądź zbioru zadań definiuje te własności.

Relacja może być też **ani symetryczna, ani asymetryczna** (czasem mówimy wtedy, że jest ona **niesymetryczna**). Oznacza to, że są w uniwersum takie pary, że relacja zachodzi pomiędzy nimi w jedną stronę i nie zachodzi w drugą, ale są też takie, w przypadku których zachodzi ona w obie strony. Symbolicznie:

R nie jest ani symetryczna ani asymetryczna $\equiv \exists x \exists y [xRy \wedge \sim (yRx)] \wedge \exists x \exists y (xRy \wedge yRx)$

Relacją ani symetryczną ani asymetryczną jest określona w zbiorze ludzi relacja kochania. Są bowiem takie pary ludzi, gdzie jedna osoba kocha drugą a druga pierwszą, ale są też i takie, gdzie relacja zachodzi tylko w jedną stronę.



RELACJA NIESYMETRYCZNA

Przechodność.

Relacja jest **przechodnia**, jeśli zachodząc pomiędzy jakimiś elementami x i y , a także elementem y i elementem z , zachodzi również pomiędzy x i z . Symbolicznie:

R jest przechodnia $\equiv \forall x \forall y \forall z [(xRy \wedge yRz) \rightarrow xRz]$

Przechodnia jest na przykład relacja bycia starszym. Jeśli jedna osoba jest starsza od drugiej, a druga od trzeciej, to na pewno pierwsza jest również starsza od trzeciej.

Fakt, że dana relacja **nie jest przechodnia** oznacza, że istnieją takie trzy elementy, że pierwszy jest w relacji do drugiego, drugi do trzeciego, natomiast pierwszy nie jest w relacji do trzeciego. Symbolicznie:

R jest nieprzechodnia $\equiv \exists x \exists y \exists z [xRy \wedge yRz \wedge \sim (xRz)]$

Nieprzechodnia jest relacja bycia znajomym. Jeśli osoba x jest znajomym osoby y , a osoba y znajomym osoby z , to wcale nie jest konieczne, aby x był również znajomym z .

Spójność.

Relacja jest spójna, jeśli dla dowolnych dwóch różnych elementów uniwersum zachodzi ona przynajmniej w jedną stronę, czyli x jest w relacji do y lub y do x . Symbolicznie:

R jest spójna $\equiv \forall x \forall y [x \neq y \rightarrow (xRy \vee yRx)]$

Spójna jest na przykład relacja mniejszości w zbiorze liczb. Jeśli weźmiemy dwie liczby i będą one różne od siebie, to na pewno jedna będzie większa od drugiej albo druga od pierwszej.

Relacja nie jest spójna, gdy istnieją w uniwersum dwa różne od siebie elementy, takie że ani jeden nie jest w relacji do drugiego, ani drugi do pierwszego. Symbolicznie:

R jest niespójna $\equiv \exists x \exists y [x \neq y \wedge \sim (xRy) \wedge \sim (yRx)]$

Niespójna w zbiorze ludzi jest na przykład relacja bycia żoną – łatwo znaleźć dwie osoby, takie że ani jedna nie jest żoną drugiej, ani druga żoną pierwszej.

W związku z trzema z wymienionymi wyżej własnościami określa się pewien szczególny typ relacji – tak zwaną **równoważność**. Mówimy, że relacja jest równoważnością, gdy jest ona jednocześnie zwrotna, symetryczna i przechodnia. Typu równoważności jest na przykład relacja bycia w tym samym wieku.

6.3.2. PRAKTYKA: OKREŚLANIE WŁASNOŚCI FORMALNYCH RELACJI.

Zadania związane z własnościami formalnymi relacji polegają najczęściej na określeniu wszystkich własności podanej relacji. W związku z każdym wyróżnionym wyżej pojęciem – zwrotnością, symetrią, przechodniością i spójnością każda relacja musi posiadać jakąś własność. Trzeba więc po prostu sprawdzić, która z możliwych sytuacji zachodzi w danym przypadku – czy relacja jest zwrotna, przeciwzwrotna, czy też ani taka, ani taka; następnie czy jest symetryczna, asymetryczna (mocno lub słabo), czy też ani symetryczna ani asymetryczna, itd.

Przykład:

Zbadamy własności formalne określonej w zbiorze wszystkich ludzi relacji bycia matką ($xRy \equiv x$ jest matką y).

Oczywiście nikt nie jest swoją własną matką, a więc jest to relacja przeciwzwrotna. Jeśli jedna osoba jest matką drugiej, to na pewno druga nie jest matką pierwszej – jest to więc relacja asymetryczna. Jeśli jedna osoba jest matką drugiej, a druga matką trzeciej, to ta pierwsza na pewno nie jest matką trzeciej (jest bowiem jej babcią), czyli nasza relacja jest nieprzechodnia. Nie jest to też relacja spójna, ponieważ nie jest tak, że dla dowolnych dwóch różnych osób jedna jest matką drugiej lub druga matką pierwszej.

Przykład:

Zbadamy własności formalne relacji bycia tej samej płci, określonej w zbiorze ludzi.

Każdy jest tej samej płci co on sam, a więc jest to relacja zwrotna. Jeśli jedna osoba jest tej samej płci co druga, to ta druga jest tej samej płci co pierwsza, a więc jest to relacja symetryczna. Jeśli osoba A jest tej samej płci co B, a B tej samej co C, to również zawsze A jest tej samej płci co C, a więc jest to relacja przechodnia. Nie jest to relacja spójna, ponieważ nie każde dwie różne osoby są tej samej płci.

Ponieważ nasza relacja jest zwrotna, symetryczna i przechodnia, możemy również powiedzieć, że jest ona równoważnością.

Z omawianych własności największe problemy może sprawić przechodniość.

Przykład:

Określimy własności formalne relacji bycia w różnym wieku (w zbiorze ludzi).

Jest to oczywiście relacja przeciwwzrotna (nikt nie jest w różnym wieku od siebie samego) i symetryczna (jeśli jedna osoba jest w różnym wieku od drugiej, to i ta druga jest w różnym wieku od pierwszej).

Zajmijmy się teraz przechodniością. Oczywiście w większości przypadków bywa tak, że jeśli jedna osoba jest w różnym wieku od drugiej, a druga od trzeciej, to i ta pierwsza będzie w różnym wieku od trzeciej. Czy jest tak jednak zawsze? Łatwo wyobrazić sobie na przykład takie trzy osoby: A – mającą 20 lat, B – 25 lat i C – 20 lat. Wtedy A będzie w relacji bycia w różnym wieku do B, B w relacji do C, natomiast A do C już nie. Ponieważ relacja jest przechodnia, gdy zawsze jest tak, że jeśli x jest w relacji do y, a y do z, to również x jest w relacji do z, to wystarczy znaleźć choć jeden przypadek, kiedy tak nie jest, aby móc stwierdzić, że relacja nie jest przechodnia. Ponieważ taki przypadek znaleźliśmy, widzimy, że nasza relacja jest nieprzechodnia.

Pewne wątpliwości może też budzić to, czy omawiana relacja jest spójna. Czy gdy weźmiemy dwóch dowolnych różnych od siebie ludzi, to zawsze będą oni w różnym wieku? Odpowiedź na to pytanie zależy od dokładności, jaką przyjmujemy. Gdy uznamy na przykład, że jeśli różnica pomiędzy dwoma ludźmi jest mniejsza niż rok, to są oni w tym samym wieku, to wtedy nasza relacja nie będzie spójna – łatwo będzie znaleźć pary ludzi, pomiędzy którymi ona nie zachodzi (a więc nie są oni w różnym wieku). Gdybyśmy jednak uznali, że różnica nawet ułamka sekundy w momencie urodzenia sprawia, że ludzie są już w różnym wieku, to

naszą relację moglibyśmy uznać za spójną – zachodziłaby ona pomiędzy dowolnymi różnymi od siebie ludźmi.

Przykład:

Zbadamy własności formalne określonej w zbiorze ludzi relacji bycia bratem.

Nikt nie jest swoim własnym bratem, więc jest to relacja przeciwzrotna. Ponieważ może być tak, że jedna osoba jest bratem drugiej, a druga bratem pierwszej, ale może być też tak, że jedna jest bratem drugiej, a druga nie jest bratem pierwszej (bo jest siostrą), oznacza to, że nasza relacja nie jest ani symetryczna, ani asymetryczna.

Jeśli chodzi o przechodność, to na pierwszy rzut oka mogłoby się wydawać, że omawiana relacja jest przechodnia – zwykle jest tak, że jeśli A jest bratem B, a B bratem C, to również A jest bratem C. Jest tu jednak pewna pułapka. Wyobraźmy sobie dwie osoby A i B będące braćmi. Wówczas A jest w relacji bycia bratem do B, B jest w relacji do A, natomiast oczywiście A nie jest swoim własnym bratem. A zatem mamy sytuację, że jedna osoba jest w relacji R do drugiej, druga do trzeciej, a pierwsza do trzeciej nie. To, że pierwsza i trzecia osoba są faktycznie tym samym człowiekiem, nic tu nie zmienia, ponieważ w definicji przechodności nie ma mowy, że muszą występować tam różne obiekty. Nasza relacja nie jest więc przechodnia.

Relacja bycia bratem nie jest też oczywiście spójna.

Relacje w tego typu zadaniach mogą być też podawane jako zbiór par.

Przykład:

Zbadamy własności formalne relacji $R = \{\langle a, b \rangle, \langle b, c \rangle, \langle a, c \rangle, \langle c, d \rangle, \langle b, b \rangle\}$ określonej w uniwersum $U = \{a, b, c, d\}$.

Ponieważ jeden z elementów uniwersum (b) jest w relacji do samego siebie, natomiast pozostałe nie są, relacja ta nie jest ani zwrotna, ani przeciwzwrotna. Dla elementów różnych od siebie jest tak, że gdy jeden jest w relacji do drugiego, to drugi nie jest w relacji do pierwszego. Wskazywało by to na asymetrię; jednak jeden z elementów jest w relacji do siebie samego, co w „mocnej” asymetrii jest niemożliwe. A zatem mamy do czynienia ze słabą asymetrią. Udaje się znaleźć takie trzy elementy (są to a, c oraz d), że pierwszy jest w relacji do drugiego, a drugi do trzeciego, natomiast pierwszy nie pozostaje w relacji do

trzeciego; jest to więc relacja nieprzechodnia. Ponieważ istnieją różne od siebie elementy, takie że ani jeden nie jest w relacji do drugiego, ani drugi do pierwszego, jest to relacja niespójna.

6.4. DZIAŁANIA NA RELACJACH.

6.4.1. ŁYK TEORII.



ŁYK TEORII

Wiemy, że każdą relację możemy przedstawić jako zbiór par. Ponieważ relacje są zbiorami (zbiorami par), możemy wykonywać na nich działania, jakie wykonywaliśmy na „zwykłych” zbiorach: sumę, iloczyn, różnicę i dopełnienie. W przypadku relacji możemy wykonywać też pewne specyficzne działania, z których poznamy tak zwany konwers relacji. Najpierw jednak zajmiemy się działaniami poznanymi w rozdziale poświęconym zbiorom.

Suma dwóch relacji to zbiór par należących do jednej lub do drugiej relacji. Na przykład sumą relacji bycia ojcem i relacji bycia matką jest relacja bycia rodzicem.

Iloczyn dwóch relacji to zbiór par należących jednocześnie do jednej jak i do drugiej relacji. Iloczynem relacji bycia bratem oraz bycia starszym jest relacja bycia starszym bratem.

Różnica dwóch relacji to zbiór tych par, które należą do pierwszej z nich, lecz nie należą do drugiej. Jeśli od relacji bycia rodzicem odejmiemy relację bycia matką, otrzymamy relację bycia ojcem.

Dopełnienie jakiejś relacji to zbiór par, które do tej relacji nie należą. Na przykład dopełnieniem relacji bycia starszym jest relacja bycia w tym samym wieku lub młodszym.

Symbolicznie działania na relacjach przedstawiamy przy pomocy takich samych znaków, jak w przypadku „zwykłych” zbiorów, czyli: $\cup$, $\cap$, $-$, $'$.

Konwers relacji to działanie, z którym się dotąd nie spotkaliśmy, jednak jego zrozumienie nie powinno sprawić większych trudności. Konwers relacji R nazywany jest często relacją odwrotną do R i bywa oznaczany symbolicznie R^{-1} lub $\check{R}$. Konwers relacji R , czyli R^{-1} to relacja zachodząca pomiędzy elementami y i x wtedy i tylko wtedy, gdy pomiędzy x i y zachodzi R . Symbolicznie $yR^{-1}x \equiv xRy$. Przykładowo, konwersem relacji bycia rodzicem, jest relacja bycia dzieckiem (bowiem y jest dzieckiem x wtedy i tylko wtedy, gdy x jest rodzicem

y), natomiast konwersem relacji bycia młodszym jest relacja bycia starszym. Konwersem pewnej relacji może być też czasem ta sama relacja – na przykład konwersem relacji bycia w tym samym wieku jest ta sama relacja bycia w tym samym wieku (y jest w tym samym wieku co x wtedy i tylko wtedy, gdy x jest w tym samym wieku co y).

6.4.2. PRAKTYKA: WYKONYWANIE DZIAŁAŃ NA RELACJACH.

Zadań związanych z działaniami na relacjach nie ma sensu szczegółowo omawiać. Jeden przykład powinien w zupełności wystarczyć.

Przykład:

Wykonamy kilka działań na następujących relacjach: $xRy \equiv x$ jest bratem y, $xTy \equiv x$ jest rówieśnikiem y, $xSy \equiv x$ jest rodzeństwem y, $xQy \equiv x$ jest siostrą y.

$$R \cap T$$

Iloczyn relacji bycia bratem i bycia rówieśnikiem to relacja zawierająca pary należące zarówno do T jak i R, a zatem relacja bycia bratem rówieśnikiem (bratem bliźniakiem) (x jest bratem bliźniakiem y).

$$S - R$$

Gdy od relacji bycia rodzeństwem odejmiemy relację bycia bratem, otrzymamy relację bycia siostrą (x jest siostrą y).

$$S \cup R$$

Dodając do relacji bycia rodzeństwem relację bycia bratem, nie dodajemy do S w istocie niczego nowego – wszystkie pary należące do R już się w S znajdują – a zatem wynikiem działania jest S, czyli relacja bycia rodzeństwem (x jest rodzeństwem y).

$$T'$$

Dopełnienie relacji T to zbiór par, które do T nie należą, a zatem jest to relacja bycia w innym wieku (x jest w innym wieku niż y lub: x nie jest rówieśnikiem y).

$$(R \cup Q)'$$

W nawiasie mamy sumę relacji bycia bratem i bycia siostrą, a więc relację bycia rodzeństwem. Dopełnienie tej ostatniej relacji to relacja nie-bycia rodzeństwem (x nie jest rodzeństwem y)

$$S - T'$$

Dopełnienie relacji T to, jak już powiedzieliśmy wyżej, relacja bycia w różnym wieku. Gdy odejmiemy ją od relacji bycia rodzeństwem, otrzymamy relację bycia rodzeństwem

będącym w tym samym wieku (x jest rodzeństwem y i są w tym samym wieku, lub: x jest bliźniaczym rodzeństwem y)

$$Q' \cap S$$

Dopełnienie relacji Q , to relacja nie-bycia siostrą. Cześć wspólna tej relacji z relacją bycia rodzeństwem to oczywiście relacja bycia bratem (x jest bratem y).

$$S^{-1}$$

Relacją odwrotną (czyli zachodzącą między y i x) do relacji bycia rodzeństwem jest ta sama relacja bycia rodzeństwem (y jest rodzeństwem x).

6.5. ZALEŻNOŚCI MIĘDZY RELACJAMI.

6.5.1. ŁYK TEORII.



ŁYK TEORII

Ponieważ relacje są zbiorami par, mogą one, podobnie jak inne zbiory, pozostawać do siebie w różnych stosunkach: inkluzji, krzyżowania i rozłączności. Zależności te zdefiniowane są tak samo jak w przypadku „zwykłych” zbiorów.

Relacja R zawiera się w relacji T ($R \subseteq T$), gdy każda para należąca do R należy również do T . Przykładowo relacja bycia kuzynem, zawiera się w relacji bycia krewnym.

Relacja R jest rozłączna z relacją T ($R \cap T = \emptyset$), gdy żadna para należąca do R nie należy równocześnie do T . Rozłączne są na przykład relacje bycia starszym i bycia młodszym.

Relacja R krzyżuje się z relacją T ($R \# T$), gdy istnieją pary należące zarówno do R jak i do T , ale są też takie, które należą jedynie do R i są takie, które należą wyłącznie do T . Przykładowo relacja bycia starszym krzyżuje się z relacją bycia bratem – może być tak, że ktoś jest starszy od kogoś innego, będąc jednocześnie jego bratem, ale można też być od kogoś starszym nie będąc jego bratem, oraz być czyimś bratem nie będąc od niego starszym.

6.5.2. PRAKTYKA: OKREŚLANIE ZALEŻNOŚCI POMIĘDZY RELACJAMI.

Zadania na określanie zależności pomiędzy relacjami są bardzo proste i jeden przykład powinien tu wystarczyć.

Przykład:

Określimy zależności pomiędzy następującymi relacjami R, S, T, Q : $xRy \equiv x$ jest matką y , $xSy \equiv x$ jest młodszym od y , $xTy \equiv x$ jest starszy od y , $xQy \equiv x$ jest rodzeństwem y .

Oczywiście niemożliwe jest, aby być jednocześnie czyjąś matką i być od tej osoby młodszym, a więc relacje R i S są rozłączne (nie ma par należących jednocześnie do nich obu). Jeśli x jest matką y , to na pewno x jest starszy od y (ale nie na odwrót), a więc relacja R zawiera się w relacji T (każda para należąca do R należy również do T). Nie można być jednocześnie czyjąś matką i rodzeństwem, a więc R jest rozłączna z Q . Z oczywistych powodów rozłączne są również relacje S i T . Rozpatrując relacje S oraz Q należy zauważyć, że można być od kogoś młodszym i być jednocześnie jego rodzeństwem, można być od kogoś młodszym i nie być jego rodzeństwem, a także można być czyimś rodzeństwem i nie być od niego młodszym; a zatem S i Q się krzyżują. Z podobnych powodów krzyżują się T i Q . A zatem, symbolicznie:

$$R \cap (S, R \subseteq T, R \cap (Q, S) \cap (T, S \# Q, T \# Q.$$

6.5.3. PRAKTYKA: DOBIERANIE RELACJI BĘDĄCYCH W RÓŻNYCH STOSUNKACH DO PODANEJ.

Zadania związane z zależnościami pomiędzy relacjami mogą też polegać na dobieraniu w stosunku do danej relacji R innych relacji: takiej żeby R się w niej zawierała, żeby ona zawierała się w R , rozłącznej z R i krzyżującej się z R . Zadania takie nie mają jednej odpowiedzi; można wymyślać wiele różnych, równie prawidłowych – wszystko zależy od wyobraźni rozwiązującego.

Przykład:

Do relacji R mieszkania w tym samym mieście ($xRy \equiv x$ mieszka w tym samym mieście co y), dobierzemy S – taką że $S \subseteq R$, T – taką że $R \subseteq T$, Q – taką że $Q \cap R$ oraz P taką że $P \# R$.

Relacja S ma się zawierać w R , a więc każda para należąca do S musi również należeć do R . Relacją taką jest na przykład relacja mieszkania na tej samej ulicy – jeśli x mieszka na tej samej ulicy co y , to na pewno x mieszka w tym samym mieście co y . Teraz musimy znaleźć relację T , taką żeby R się w niej zawierała; czyli każda para mieszkająca w tym samym

mieście musi również należeć do naszej nowej relacji T. Relacją taką może być, na przykład, relacja mieszkania w tym samym kraju. Za przykład relacji Q rozłącznej z R może posłużyć relacja mieszkania w innym mieście. Jako relację krzyżującą się z R możemy podać relację bycia bratem – jedna osoba może być bratem drugiej i mieszkać jednocześnie w tym samym mieście co ta druga, ale można też być czyimś bratem i mieszkać w innym mieście, a także mieszkać z kimś w tym samym mieście, ale nie być jego bratem. A zatem ostateczna, jedna z wielu możliwych, odpowiedź to:

$xSy \equiv x$ mieszka na tej samej ulicy co y ,

$xTy \equiv x$ mieszka w tym samym kraju co y ,

$xQy \equiv x$ mieszka w innym mieście niż y ,

$xPy \equiv x$ jest bratem y .

SŁOWNICZEK:

Iloczyn kartezjański – iloczyn kartezjański zbiorów A i B ($A \times B$) to zbiór wszystkich par, takich w których na pierwszym miejscu jest element zbioru A, a na drugim element zbioru B.

Kwadrat kartezjański – kwadrat kartezjański zbioru A to iloczyn kartezjański A z nim samym, czyli $A \times A$.

Dziedzina lewostronna relacji – zbiór tych obiektów, które pozostają w relacji do jakiegoś obiektu.

Dziedzina prawostronna relacji (przeciwdziedzina) – zbiór tych obiektów, do których jakiś obiekt pozostaje w relacji.

Pole relacji – suma dziedziny lewostronnej i prawostronnej relacji.