

Jacek Malinowski

Logiki niemonotoniczne

Czy logika jest nauką kognitywną?

Celem tej rozprawy jest prezentacja pewnej dziedziny badań logicznych, która zarówno ze względu na swe zastosowania jak i na swój zakres należy do nauk kognitywnych stanowiąc próbę modelowania procesów praktycznego wnioskowania. Dziedziną tą są logiki niemonotoniczne — problematyka powstała na początku lat osiemdziesiątych, z niesłabnącą intensywnością rozwijana w wielu ośrodkach logicznych oraz centrach kognitywistyki na świecie.

Rozpocznijmy jednak od kilku uwag bardziej ogólnej natury. Co to jest kognitywistyka? Jak ma się ona do nauk szczegółowych, takich jak psychologia czy logika? Najogólniej rozumiana kognitywistyka to współczesna wersja teorii poznania — epistemologii. Tak rozumiana kognitywistyka stara się dać odpowiedź na postawione przez Kanta pytanie: Jak możliwe jest poznanie?, bazując, w odróżnieniu od Kanta, nie tylko na refleksji filozoficznej, lecz przede wszystkim na współczesnym dorobku tych nauk szczegółowych, które badają procesy składające się na poznanie. Do nauk tych zaliczyć należy przede wszystkim psychologię, logikę i lingwistykę. Baza empiryczna psychologii i lingwistyki, techniki formalne zaczerpnięte z matematyki i informatyki wypracowane przez logików i lingwistów stanowią dostatecznie potężne narzędzia badawcze, aby pokusić się o poważną próbę modelowania procesów poznawczych człowieka. Jako idealny i być może nieosiągalny cel owego modelowania można postawić stworzenie sztucznej inteligencji — maszyny do poznawania, która uczy się i wykorzystuje swą wiedzę w taki sam sposób jak człowiek.

Cel postawiony wyżej brzmi w takim sformułowaniu fantastycznie i pomimo licznych osiągnięć jest to idea niemożliwa dziś do zrealizowania w całej jej ogólności. Wiele jest jednak przykładów realizacji idei, które jeszcze całkiem niedawno wydawały się fantastyczną mrzonką. Odpowiedzi na pytanie, na ile cel ten jest osiągalny, udzielić mogliby jedynie (choć nie na pewno) specjaliści z różnych szczegółowych dziedzin wchodzących w skład kognitywistyki oraz tych, które — tak jak matematyka czy informatyka — choć w jej skład nie wchodzi, to udzielają jej swych narzędzi badawczych. Zauważmy, że nawet gdyby cel ten był nieosiągalny, wyznacza on charakter

badani kognitywistyki. Ideały służą nie tyle do tego, by do nich dążyć, ile do tego, by się według nich orientować. Postawiona jako cel budowa algorytmu realizującego funkcje poznawcze człowieka narzuca rozważaniom i badaniom reżim szczególnej ścisłości, wymuszając przy okazji zdecydowanie aplikacyjny charakter jej bardziej zaawansowanych działów — jeśli w danej dziedzinie zrobiono już wiele, a celem jest algorytm, to tworzymy go, nawet jeśli modeluje on jedynie z pewnym przybliżeniem tylko niektóre aspekty poznania.

Kognitywistyka jest zatem epistemologią ubraną w nowe szaty. Wydaje się, że dziedziny te mają się do siebie tak, jak oparta wyłącznie na refleksji filozoficznej fizyka arystotelesowska do fizyki nowożytnej utrzymywanej w ryzach wyznaczonych przez empirię. Szybko rosnąca liczba centrów kognitywistyki, znaczna ilość czasopism specjalistycznych jej poświęconych, wreszcie ogromna liczba prac i raportów badawczych pozwala posunąć tę analogię do przypuszczenia, że metody kognitywne mogą stanowić rewolucję w epistemologii podobną do tej, jaką w fizyce były doświadczenia Galileusza.

Skoro już mowa o empirii, słów kilka o empirycznych aspektach kognitywistyki. Z pewnością trudno byłoby doszukiwać się poważnych empirycznych aspektów w tradycyjnej epistemologii. Inaczej jest w przypadku kognitywistyki. Nie sposób nie doceniać empirycznego dorobku psychologii czy lingwistyki, również w ich kognitywnej części.

Zastanówmy się, w jakim stosunku do siebie pozostają logika i kognitywistyka. Rozpocznijmy od cytatu z przedmowy Kanta do *Uzasadnienia metafizyki moralności*: „Logika nie może mieć żadnej części empirycznej, tj. takiej, w której powszechne i konieczne prawa myślenia polegałyby na podstawach zaczerpniętych z doświadczenia; inaczej bowiem nie byłaby logiką, tj. kanonem dla intelektu i rozumu, który przy wszelkim myśleniu obowiązuje i musi być dowiedziony”¹. Tak rozumiana logika nie tylko nie mieści się w kognitywistyce, ale wręcz wydaje się nie mieć z nią wspólnych obszarów badawczych. Praktyczne rozumowania podlegają jedynie ocenie logiki pod względem poprawności, ale same nic do niej nie wnoszą. Definicja taka, jak sądzę, jest wynikiem dość skrajnego podejścia do logiki. Nie czuję się uprawniony ani do jej obrony, ani do polemizowania z nią. Potraktujmy ją jako wskazanie Kanta na pewne absolutne cechy logiki. Otóż nawet powszechna akceptacja danego sposobu wnioskowania, uznanie, że prowadzi on zawsze do poprawnych wniosków, o ile przesłanki były prawdziwe, nie może być nie tylko dowodem, ale nawet poważnie brany argumentem za jego logiczną poprawnością. Logika bowiem odnosi się do bytów absolutnych, takich jak prawda i fałsz. Jej prawa określające ważność rozumowań pozostałyby ważne nawet gdyby w praktyce nikt nie używał rozumu do wyciągania wniosków. Logika tak rozumiana mogłaby dalej się rozwijać. Z przymrużeniem oka można byłoby powiedzieć, że brak istot posiadających psychikę czyniłby pracę psychologa niemożliwą; aby uniemożliwić badania logiczne nie wys-

¹ I. Kant, *Uzasadnienia metafizyki moralności*, przeł. M. Wartenberg, Warszawa 1984, s. 3.

tarczy nieistnienie istot, które dokonują wnioskowań, trzeba by unicestwić więcej, bo pojęcia prawdy i fałszu.

Wydaje się, że logika w rozumieniu Kanta nie jest, a co najmniej nie w pełni jest częścią kognitywistyki, choć z całą pewnością jest kognitywistycie niezbędna. Bada kanony ważności rozumowań, uznając za poprawne te z nich, które ze względu na strukturę przesłanek i konkluzji prowadzą do poprawnych konkluzji, o ile poprawne są przesłanki.

Logicy opisali i zbadali wiele systemów tego rodzaju reguł. Wszystkie one, od klasycznej logiki zdań i predykatów poczynawszy, a na skomplikowanych nieklasycznych logikach różnego rodzaju silnej implikacji skończywszy, spełniają kryterium postawione przez Kanta — nie polegają na prawach zaczerpniętych z doświadczenia. Nawet systemy, które wywodzą się z refleksji nad paradoksami implikacji, nie polegają na empirycznych danych dotyczących jej rozumienia. Nie negują poprawności praw logiki klasycznej, lecz tworzą systemy, w których implikacja rozumiana jest inaczej.

Chciałbym w tej pracy wyjść poza cytowane kryterium Kanta, jak też poza fundamentalne zasady rządzące logiką od czasów Fregego i przedstawić najważniejsze syntaktyczne i semantyczne idee związane z logikami niemonotonicznymi, które zasad tych nie spełniają. Ta dość świeża, bo licząca sobie wszystkiego kilkanaście lat problematyka jest powszechnie uważana za część logiki, mimo iż uznaje ona za poprawne pewne rozumowania, które przy poprawnych przesłankach mogą czasami prowadzić do fałszywych konkluzji. Ściśle biorąc, logiki niemonotoniczne nie badają reguł poprawności rozumowań.

Przyjmijmy następujący plan. Rozpocniemy od omówienia fundamentalnych zasad logiki, abyśmy jasno zdali sobie sprawę, jakie reguły zamierzamy naruszać. Następnie zastanowimy się nad tym, co, skoro nie reguły poprawności rozumowań, należy do zakresu badań logik niemonotonicznych, aby dalej przedstawić syntaktyczne reguły najważniejszych klas logik niemonotonicznych. Na koniec postaramy się zaprezentować pewne ogólne idee semantyczne, definiując pewne klasy modeli za pomocą zbioru stanów oraz relacji preferencji.

Podstawowe zasady logiki

Zasada I. Klasycznie rozumiana inferencja logiczna jest relacją pomiędzy zdaniami lub sądami w sensie logicznym, a nie pomiędzy myślami, czy czymkolwiek innym związanym z procesem poznawczym.

Zgodnie z tradycyjnym podejściem rozumowanie logiczne jest relacją inferencji pomiędzy zbiorem zdań czy też sądów będących przesłankami a wnioskiem.

Zasada II. Według klasycznego podejścia ważność rozumowania zależy jedynie od logicznej struktury przesłanek i konkluzji, nie zależy natomiast od ich znaczenia, prawdziwości ani kontekstu.

Relacja inferencji może być definiowana na dwa główne, zasadniczo różne sposoby. Syntaktycznie — poprzez definicję dowodu zdania na gruncie przesłanek, w oparciu o przyjmowane reguły dowodzenia i tezy lub też semantycznie — poprzez określenie warunków prawdziwości konkluzji, jako funkcji prawdziwości przesłanek. Metody te opisują tę samą relację logicznej konsekwencji na dwa odmienne sposoby. Sposób syntaktyczny jest mechanizmem, który w oparciu jedynie o strukturę zdań pozwala nam generować te z nich, które są logicznie prawdziwe, nie odwołując się przy tym ani do prawdziwości zdań, ani do żadnej ich własności innej niż struktura logiczna. Semantyka natomiast nadaje zdaniom znaczenia, wiążąc je w ten sposób z fragmentem rzeczywistości opisywanym przez model.

Relację logicznego wynikania można w pełni opisać w sposób neutralny, to znaczy nie odnoszący się ani do syntaksy, ani do semantyki. Zrobili to jako pierwsi Tarski w terminach operacji konsekwencji² oraz niezależnie Gentzen i Jaśkowski w terminach dedukcji naturalnej³. Przytoczę obie te charakterystyki jako precyzyjną eksplikację intuicji, aby pokazać dalej, które z nich nie działają należycie w logice niemonotonicznej.

Rozważać będziemy język zdaniowy, rozumiany jako ustalony zbiór L zdań dobrze zbudowanych ze zmiennych zdaniowych za pomocą spójników. Zmienne zdaniowe oznaczać będziemy literami p, q, r , ewentualnie z indeksami. Zdania oznaczać będziemy greckimi literami P, Q, R , ewentualnie z indeksami.

Funkcję C , która dowolnemu zbiorowi zdań X przyporządkowuje zbiór zdań $C(X)$, który rozumieć będziemy jako zbiór logicznych konsekwencji zbioru X , nazywać będziemy operacją (logicznej) konsekwencji wtedy i tylko wtedy, gdy spełnia ona sformułowane poniżej zasady zwrotności, idempotentności i monotoniczności.

Zasada zwrotności: Dla dowolnego zbioru zdań X , $X \subseteq C(X)$. Mówi po prostu, że zdania, które przyjmujemy jako przesłanki, przyjmujemy też jako konkluzje. Zasady tej nie będziemy kwestionować również w logikach niemonotonicznych.

Zasada monotoniczności: Dla dowolnych zbiorów zdań X i Y , jeśli $X \subseteq Y$, to $C(X) \subseteq C(Y)$. Jej istotę można opisać w następujący sposób: jeśli dana konkluzja jest logiczną konsekwencją danego zbioru przesłanek, to jest ona także konsekwencją dowolnego szerszego ich zbioru. Rozszerzenie zbioru przesłanek pozwala zatem uznać co najmniej te same konkluzje, które daje się wywieść z pierwotnego ich zbioru.

Zasada idempotentności: $CC(X) = C(X)$. Mówi ona, iż to co jest logiczną konsekwencją zbioru logicznych konsekwencji przesłanek, jest również logiczną konsekwencją samego zbioru przesłanek.

² A. Tarski, *Über einige fundamentale Begriffe der Metamathematik*, „C. Séances Soc. Sci. Letters Varsovie”, vol. 23 (1930), s. 22-29.

³ G. Gentzen, *Über die Existenz unabhängiger Axiomensysteme zu unendlichen Sätzsystemen*, „Math. Ann.”, vol. 107, s. 329-350.

Jeśli C spełnia ponadto sformułowaną poniżej zasadę strukturalności, to nazywa się ją często strukturalną operacją konsekwencji lub też (której to konwencji nie możemy tutaj przyjąć) po prostu logiką.

Zasada strukturalności: dla dowolnego podstawienia e oraz dowolnego zbioru zdań X , $e(C(X)) \subseteq C(e(X))$.

Zasada ta wyraża kluczową dla logicznego wynikania ideę, iż jedyną cechą zdań istotną przy wnioskowaniu jest jego struktura logiczna. Jeśli mianowicie ze zbioru przesłanek X wynika zdanie P i w miejsce pewnych dowolnie wybranych zmiennych zdaniowych w zbiorze $X \cup P$ wstawimy inne zmienne (podstawienie e), to powstałe w ten sposób zdanie $e(P)$ będzie logiczną konsekwencją otrzymanego zbioru zdań $e(X)$.

Wielu logików zamiast pojęcia operacji konsekwencji woli używać pojęcia relacji logicznej konsekwencji, to znaczy relacji $\vdash$ pomiędzy zbiorami zdań a zdaniami, spełniającej warunki odpowiadające w naturalny sposób warunkom zwrotności, idempotentności i monotoniczności. Pojęcia relacji i operacji logicznej konsekwencji są nawzajem definiowalne w sposób opisany następującym wzorem:

$$P \in C_{\vdash}(X) \Leftrightarrow X \vdash_{\vdash} P$$

Wobec równoważności tych dwóch pojęć wybór jednego z nich, relacji lub operacji konsekwencji, wynika z przyczyn estetycznych bądź też dokonuje się go ze względu na większą łatwość i elegancję formułowania i uzasadniania danych twierdzeń w kategoriach jednego z nich.

Istnieje jeszcze inny sposób definiowania relacji konsekwencji logicznej. Pochodzi on od Gentzena, pewien jego wariant został jednak niezależnie skonstruowany przez Jaśkowskiego. Konstrukcja ta nosi nazwę metody dedukcji naturalnej, powstała bowiem w wyniku krytycznej refleksji nad standardowym pojęciem dowodu. Istotnie, metoda ta nieporównanie łatwiej pozwala dowodzić tez, na przykład klasycznego rachunku zdań, niż metoda standardowa. Dedukcja naturalna jest w swej istocie syntaktycznej natury. Nic nie stoi jednak na przeszkodzie traktowaniu jej jako kolejnego, a w tej pracy nawet najważniejszego sposobu definiowania relacji konsekwencji. W wersji oryginalnej nie służyła ona jako definicja relacji konsekwencji, lecz jako pewna formalizacja logiki klasycznej oraz intuicjonizmu. Trzy spośród reguł Gentzena ważne są jednak w dowolnej relacji konsekwencji i, w istocie, definiują ją.

Rozważmy następujące trzy sekweny:

Ref	$X, P \vdash P$	<i>Zwrotność</i>
Mon	$\frac{X \vdash P}{X \cup Y \vdash P}$	<i>Monotoniczność</i>

Cut

$$\frac{X, Q \vdash P, X \vdash Q}{X \vdash P}$$

Cięcie

Binarną relację $\vdash$ pomiędzy zbiorami zdań i zdaniami nazywamy relacją konsekwencji (logicznej) wtedy i tylko wtedy, gdy: należą do niej (relacja jest zbiorem) wszelkie pary postaci **Ref** oraz jeśli jest ona zamknięta na reguły **Mon** i **Cut**, to znaczy zawiera parę spod kreski, jeśli tylko zawiera odpowiednie pary nad kreski.

Łatwo zauważyć bliskie podobieństwo w konstrukcji warunków Tarskiego i Gentzena, nawet ich nazwy, poza cięciem, które odpowiada idempotentności, są identyczne.

Zauważmy, że relacja konsekwencji, którą właśnie zdefiniowaliśmy, nie musi być strukturalna. Istnieje jednak najmniejsza w sensie zawierania relacja konsekwencji spełniająca warunki 1, 2. Relacja ta jest strukturalna, to znaczy spełnia regułę:

Str

$$\frac{X \vdash P}{e(X) \vdash e(P)}$$

Strukturalność

gdzie e oznacza dowolne podstawienie, $e(X)$, $e(P)$ są odpowiednio wynikiem podstawienia e we wszystkich zdaniach zbioru X oraz podstawienia e w zdaniu P . Najmniejsza relacja spełniająca 1, 2 nie jest jednak zbyt interesująca, czemu trudno się dziwić, bo w jej definicji nie mówi się wcale o strukturze języka. Relacja ta zawiera jedynie inferencje postaci $X, P \vdash P$, odpowiada ona tak zwanej odtwarzającej operacji konsekwencji $C(X)=X$. Reguła strukturalności nie będzie grała w naszych rozważaniach szczególniejszej roli. Choć, co warto zauważyć, pewne bardzo szczególne logiki niemonotoniczne są strukturalne.

Struktura języka nie była, jak dotąd, istotna w naszych rozważaniach. W rzeczywistości, nie wyróżniając struktury przedmiotowego języka, nie dokona się wiele metodami logiki. Zaprezentowane definicje funkcjonują dla dowolnego języka. Począwszy od tej chwili rozważać będziemy jedynie klasyczny język zdaniowy ze spójnikami alternatywy $\vee$, koniunkcji $\wedge$, negacji $\neg$, implikacji $\rightarrow$ i równoważności $\leftrightarrow$ oraz zbiorem zmiennych zdaniowych V , który jest skończony lub przeliczalnie nieskończony.

Relację konsekwencji definiowaliśmy dotąd jako relację pomiędzy zbiorami (ewentualnie nieskończonymi) zdań a zdaniami. Rozważanie nieskończonych zbiorów zdań, w pewnych przypadkach istotne, dla kognitywnych zastosowań wydaje się w dużym stopniu nadmiarowe⁴. Ograniczymy się zatem do skończonych zbiorów przesłanek, co w naszym przypadku jest równoważne traktowaniu relacji konsekwencji logicznej jako binarnej relacji na zbiorze zdań.

Ostatecznie definicja relacji konsekwencji logicznej na użytek tej pracy przybierze postać:

⁴Nie do końca jednak. W niedawnych pracach Freund i Lehmann badają niefinitarne, tzw. supraklasyczne logiki niemonotoniczne.

Ref	$P \vdash P$	<i>Zwrotność</i>
Mon	$\frac{P \vdash R}{P \wedge Q \vdash R}$	<i>Monotoniczność</i>
Cut	$\frac{P \wedge Q \vdash R, P \vdash Q}{P \vdash R}$	<i>Cięcie</i>

Binarną relację – pomiędzy zdaniami nazywamy relacją konsekwencji (logicznej) wtedy i tylko wtedy, gdy: zawiera ona wszelkie pary postaci **Ref** oraz zamknięta jest na reguły **Mon** i **Cut**.

Niech $\models$ oznacza klasyczną relację wynikania logicznego, zadaną dla prostoty semantycznie. Zapis $P \models Q$ oznacza zatem, że zdanie Q przyjmuje wartości 1 w każdym (klasycznym) zero-jedynkowym wartościowaniu, w którym P przyjmuje wartość 1, natomiast zapis $\models P$ oznacza, że zdanie P przyjmuje wartości 1 w każdym zero-jedynkowym wartościowaniu.

Przytoczmy jeszcze pewną ważną własność logiki klasycznej:

Twierdzenie o zupełności. Jedyną logiką silniejszą od logiki klasycznej jest logika sprzeczna.

Naruszenie podstawowych zasad rządzących logiką, a w konsekwencji inne rozumienie terminu „logika”, pozwoli nam zajmować się w dalej kilkoma najważniejszymi spośród nieskończenie⁵ wielu niesprzecznych logik silniejszych od logiki klasycznej.

Co bada logika niemonotoniczna

Jak powiedzieliśmy wyżej, istnieją niesprzeczne logiki niemonotoniczne zawierające logikę klasyczną. Jasne jest zatem, że nie mogą one być logikami w zdefiniowanym powyżej sensie. Nie będziemy zajmować się dalej relacjami logicznej konsekwencji, lecz pewnymi innymi relacjami o znacznie słabszych własnościach. Relacje takie nazywa się czasem w literaturze logicznej relacjami inferencji. Takiego terminu będziemy właśnie używać dla określenia dowolnych binarnych relacji na zbiorze zdań języka L zawierających logikę klasyczną, a więc takich relacji $\sim$, że $P \sim Q$, jeśli $P \models Q$.

Co bada logika niemonotoniczna? Aby uniknąć wszelkich nieporozumień, należy wyraźnie powiedzieć, że nie neguje ona koncepcji relacji logicznej konsekwencji jako formalizacji pojęcia wynikania logicznego, dokładnie w tym samym sensie, w jakim badania logik relewantnych nie negują prawdziwości tautologii wyrażających paradoksy implikacji, a jedynie badają własności (innego) spójnika implikacji, wyrażającego inne intuicje, mającego inne znaczenie. Nie będziemy zatem negować

⁵ Jeśli dopuścimy do rozważań logiki infinitarne (tj. logiki z istotnie nieskończonymi zbiorami przesłanek), to logik takich będzie nieprzeliczalnie wiele.

tego, że pojęcie relacji konsekwencji jest właściwą formalizacją logicznego wnioskowania, a zatem właściwą precyzacją zwrotów o postaci:

Jeśli przyjmiemy że P , to musimy przyjąć Q .
 Jeśli prawdą jest P , to musi być prawdą, że Q .
 P , więc z konieczności Q .

Przedmiot badań logiki niemonotonicznej jest inny i, co warto podkreślić, znacznie bliższy potocznym wnioskowaniom. Rozważmy kilka „wyrażeń źródłowych” dla logiki niemonotonicznej:

- a) Jeśli P , to zwykle Q .
- b) Jeśli P , to w normalnych okolicznościach Q .
- c) Jeśli P jest akceptowalne, akceptowalne jest Q .
- d) Jeśli P jest prawdopodobne, to i Q jest prawdopodobne.
- e) Jeśli wiemy, że P , to powinniśmy przyjąć Q .
- f) Jeśli wierzę że P , to powinienem przyjąć Q .
- g) Jeśli w danym stanie badań uznajemy P , to powinniśmy przyjąć Q .
- h) Jeśli P i $\neg Q$ jest niesprzeczne, to Q .

W żadnym z powyższych schematów zasada monotoniczności nie jest ważna. Jako ilustrację dla punktu a) rozważmy naiwnie prosty przykład. Oznaczmy przez:

P : x prowadzi wykład.

Q : x stoi na rękach.

R : x nie jest bardzo czerwony na twarzy.

Jasne, że jeśli ktoś prowadzi wykład, to zwykle nie jest bardzo czerwony na twarzy. Traktując więc $\sim$ jako skrót dla relacji typu a), mieć będziemy $P \sim R$. Jeśli jednak niespodziewanie zauważymy, że prowadzący wykład stoi na rękach, to w zwykłych okolicznościach (już są niezwykle, więc raczej w najmniej zwykłych) skłonni będziemy oczekiwać, że jest on bardzo czerwony na twarzy. Zatem $P \wedge Q \sim R$, a zatem $P \wedge Q \not\sim R$. Dodanie nowej przesłanki Q powoduje, że nie wywnioskujemy już R .

W większości pozostałych przypadków przykłady przeczące spełnieniu zasady monotoniczności będą miały zbliżoną postać. W rzeczywistości niektóre z nich nie spełniają również *Cut*.

Nasuwać się dwa pytania. Czy lista powyższa jest wyczerpująca? Czy schematy na niej zamieszczone rzeczywiście są różne? Odpowiedź na pierwsze z nich jest z pewnością negatywna. Pytanie to jest jednak dla mnie raczej pretekstem dla znacznie poważniejszego pytania o możliwość stworzenia takiej listy lub co najmniej metody klasyfikacji potocznie realizowanych inferencji. Wydaje się, że badania tego rodzaju stanowią jedno z najważniejszych zadań logiki kognitywnej (tak chciałbym nazywać kognitywną część logiki).

Drugie z powyższych pytań ma pozornie oczywistą odpowiedź pozytywną. Ono jednakże jest również pretekstem dla pytania o podstawową jednostkę taksonomiczną w klasyfikacji inferencji potocznych. Wobec niezmierzonego bogactwa języka naturalnego trudno nawet byłoby znaleźć przykład dwóch różnych zwrotów, które we wszystkich poprawnych (nawet tylko typowych) użyciach wyrażałyby tę samą inferencję. Ale na tym nie koniec. Nawet ten sam zwrot użyty w różnych sytuacjach może wyrażać różne inferencje. To prowadziłoby do koncepcji poszczególnego egzemplarza jako podstawowej jednostki taksonomicznej. Z metodologicznego punktu widzenia jest to nie do przyjęcia. Każda potoczna inferencja byłaby jedyna w swym rodzaju i niepowtarzalna, a zjawisk tego typu naukowo badać nie sposób.

Badania w dziedzinie taksonomii inferencji potocznych wymagają niewątpliwie udziału lingwistów. Ograniczymy się do omówienia najważniejszych rodzajów inferencji, stanowiących punkt wyjścia dla systemów logik niemonotonicznych rozważanych przez logików. Można wyróżnić następujące typy: sytuacyjne, epistemiczne, probabilistyczne. Terminy tu użyte mają roboczy charakter i nie występują w tych znaczeniach w literaturze przedmiotu.

Inferencje sytuacyjne. Ich przykładami są a), b), c) oraz w pewnym sensie g). Ich interpretacja oparta jest na koncepcji zwykłych okoliczności. Jest to zatem specjalny rodzaj rozumowań entymematycznych, w których dodatkowe milczące założenia są nie tyle oczywiste, ile uznawane za akceptowalne w typowej hipotetycznej sytuacji, w której spełnione jest założenie (przesłanka). W tym podejściu inferencja jest uznana za poprawną, jeśli konkluzja jest prawdziwa w sytuacjach czy światach, które są najbardziej typowe spośród wszystkich światów spełniających przesłankę. Analiza tego typu zbliżona jest do typowej analizy okresów kontrfaktycznych za pomocą możliwych światów.

Inferencje epistemiczne. Występują one w przykładach e) i f). Przykład g) wydaje się pasować także i tutaj. Są one zwykle interpretowalne w terminach wiedzy mówiącego. W konsekwencji można je traktować jako pewne rodzaje ścisłej implikacji, przy epistemicznym rozumieniu operatora konieczności.

Inferencje probabilistyczne. Typowy przykład to d). Zachodzenie inferencji tego rodzaju oparte jest na szacowaniu prawdopodobieństwa prawdziwości przesłanek lub konkluzji. Inferencje probabilistyczne badane były wcześniej niż inne rodzaje niemonotonicznych rozumowań, tyle że raczej jako pewne rodzaje implikacji, a nie wynikania.

Poprzestaniemy na tym dalece niepełnym przedstawieniu i ograniczymy się dalej głównie do inferencji sytuacyjnych.

Pewne klasy logik niemonotonicznych

Istnieje jedno pojęcie relacji konsekwencji i nie spotyka się raczej poważnych obiekcji co do jego adekwatności dla opisu wynikania logicznego. W logice niemonotonicznej sytuacja nie jest tak prosta. Dotychczasowe badania wskazują raczej na funkcjonowanie w języku wielu różnych, choć podobnych pojęć, które mogą być

opisane w języku logiki niemonotonicznej. Zdefiniujemy teraz pewne ważne klasy logik niemonotonicznych. W zasadzie odnoszą się one do inferencji sytuacyjnych, choć w pewnych przypadkach mogą w poprawny sposób opisywać pewne inferencje innego rodzaju. Systemy, które opiszemy, pochodzą w tym sformułowaniu od Krausa, Lehmana i Magidora⁶. Pierwszym, który zasugerował skoncentrowanie badań w dziedzinie logik niemonotonicznych na badaniu własności relacji inferencji, był Gabbay.⁷ Badał on syntaktyczne własności systemów spełniających niżej zdefiniowane reguły zwrotności, cięcia i kumulatywności, sformułowane w słabym języku nie zawierającym spójników logiki klasycznej.

Jak już powiedziano wyżej, rozważać będziemy jedynie klasyczny język zdaniowy L ze spójnikami alternatywy $\vee$, koniunkcji $\wedge$, negacji $\neg$, implikacji $\rightarrow$ oraz równoważności $\leftrightarrow$, oraz zbiorem zmiennych zdaniowych V , który jest skończony lub przeliczalnie nieskończony. Przez waluację rozumiemy jak zwykle dowolne przyporządkowanie zmiennym zdaniowym jednej z wartości logicznych prawdy — 1 lub fałszu — 0.

Z językiem L stowarzyszymy semantykę zadaną przez zbiór U , którego elementy nazywać będziemy światami, oraz binarną relację spełniania $\models$ pomiędzy zdaniami a światami. Zakładamy, że semantyka $(U, \models)$ zachowuje się klasycznie wobec klasycznych spójników, to znaczy spełnia warunki:

- (1) $u \models \neg P$ wtw $u \not\models P$.
- (2) $u \models P \vee Q$ wtw $u \models P$ lub $u \models Q$.
- (3) $u \models P \wedge Q$ wtw $u \models P$ oraz $u \models Q$.
- (4) $u \models P \rightarrow Q$ wtw $u \models \neg P$ lub $u \models Q$.
- (5) $u \models P \leftrightarrow Q$ wtw $(u \models P \text{ wtw } u \models Q)$.

wtw jest tutaj skrótem dla zwrotu wtedy i tylko wtedy, gdy.

Zapisu $u \models P$ używać będziemy, gdy $u \models P$ dla dowolnego świata u . Jeśli $u \models P \rightarrow Q$, lub co jest równoważne, $u \models Q$ w każdym świecie u , w którym $u \models P$, to będziemy pisać: $P \models Q$. Oczywiście $P \models Q$ wtedy i tylko wtedy, gdy Q wynika z P w sensie logiki klasycznej.

Dowolny świat $u \in U$ możemy interpretować w naturalny sposób jako waluację. Będziemy tak postępować bez każdorazowych uzasadnień. Zbiór U możemy więc traktować jako zbiór klasycznych waluacji. Warto zauważyć, że U nie musi być zbiorem wszystkich waluacji, lecz ich właściwym podzbiorem. Zbiór zdań P takich, że $u \models P$, może wtedy być szerszy niż zbiór tautologii, będzie on po prostu pewną teorią klasyczną. Ograniczenie takie daje nam możliwości modelowania takich okoliczności, w których pewne zdania kontyngentne są zawsze prawdziwe. Jeśli przykładowo

⁶S. Kraus, D. Lehmann, M. Magidor, *Nonmonotonic Reasoning, Preferential Models and Cumulative Logics*, „Artificial Intelligence”, t. 44 (1990), s. 167-207. Tam też znaleźć można pełną bibliografię przedmiotu.

⁷D.M. Gabbay, *Theoretical Foundations for Non-Monotonic Reasoning in Expert Systems*, w: K.R. Apt (wyd.), *Proceedings NATO Advanced Institute on Logics and Models of Concurrent Systems*, La Colle-sur-Loup, France (Berlin 1985), s. 439-457.

chcemy ograniczyć się do tych jedynie światów, w których pingwiny są ptakami, to U jest zbiorem tych waluacji, w których spełnione jest zdanie „pingwin $\rightarrow$ ptak”. Zwykle symbol ten występuje w takim właśnie znaczeniu. Na potrzeby tej pracy ograniczymy się do węższego jego rozumienia jako zbioru wszystkich waluacji. W konsekwencji $=$ oznacza po prostu semantyczną relację wynikania w sensie logiki klasycznej.

Prezentację syntaksy systemów logiki niemonotonicznej rozpoczniemy od przedstawienia pewnych reguł.

Ref

$$P \sim P$$

Zwrotność

Zasada ta wydaje się uniwersalnie ważna dla wszelkich rozumowań związanych z wnioskowaniem. Relacje, które jej nie spełniają, opisują pewne pojęcia z teorii zmiany. Jeśli mianowicie dana teoria naukowa przeczy wynikom doświadczeń, to istnieją zdania należące do teorii, których jednakże nie akceptujemy. Nie będziemy jednak omawiać bliżej odstępstw od tej reguły.

LLE

$$\frac{\models P \leftrightarrow Q, P \sim R}{Q \sim R}$$

Lewostronna Równoważność

Zdania logicznie równoważne mają te same niemonotoniczne konsekwencje. Zauważmy, że znaku równoważności $\leftrightarrow$ nie da się zastąpić przez implikację. Tego rodzaju zastąpienie dałoby nam w efekcie regułę monotoniczności w sformułowaniu, które za chwilę omówimy. Z prawej strony znaku $\sim$ możemy jednak mieć implikację:

RW

$$\frac{\models P \rightarrow Q, P \sim R}{Q \sim R}$$

Prawostronne Osłabienie

Reguła ta wyraża prostą ideę, że wnioskowania niemonotoniczne są zamknięte względem klasycznego wynikania. Jeśli inferencja $\sim$ spełnia zarówno **LLE** jak i **RW**, to $P \sim Q$ wtedy i tylko wtedy, gdy $P' \sim Q'$, dla pewnych P', Q' równoważnych logicznie odpowiednio P i Q . Relacja $\sim$ jest więc obustronnie niezmiennicza względem zdań logicznie równoważnych. Fakt ten daje nam możliwość badania relacji inferencji jako relacji określonej na algebrze Lindenbauma zdań języka L . Powrócimy do tego tematu w części dotyczącej semantyki dla logik niemonotonicznych.

Dla uniknięcia nieporozumień zauważmy jednak, że relacja logicznej równoważności nie skleja ze sobą wszystkich niemonotonicznych przesłanek (ani też wszystkich konkluzji) danego zdania. Z tego, że $P \sim Q$ i $P \sim R$ nie wynika, że Q i R są równoważne. Również z $P \sim R$ i $Q \sim R$ nie wynika, że P i Q są równoważne. Co więcej, łatwo pokazać, że wynikanie takie zachodziłoby tylko wtedy, gdyby $\sim$ było identyczne z relacją $\models$ klasycznego wynikania.

Cut

$$\frac{P \wedge Q \sim R, P \sim Q}{P \sim R}$$

Cięcie

Ta omawiana wyżej, pochodząca od Gentzena reguła opisuje warunki, w których możliwe jest usunięcie (wycięcie — stąd nazwa) pewnych przesłanek. Otóż zbędne przesłanki to te, które dają się wywieść z innych przesłanek. Jest jasne, jak działa reguła cięcia i jak należy ją rozumieć dla relacji konsekwencji. Jednakże intuicje dotyczące relacji konsekwencji zwykle zawodzą przy próbie stosowania ich do niemonotonicznych inferencji. Przyjrzyjmy się bliżej regule cięcia w niemonotonicznym wydaniu. Otóż o ile nie mówi ona nic ciekawego o operacji konsekwencji — po prostu wszelkie konsekwencje są takie — o tyle w pewnych przypadkach rozumowań niemonotonicznych może ona być zawodna. W rezultacie nadaje ona rozważanym relacjom inferencji pewien określony sens. Rozważany przykład oznaczmy przez:

P: Pada deszcz,

Q: Zabieram parasol,

R: Nie zmoknę.

Przyjmijmy, że mam taki zwyczaj, że jeśli pada deszcz, to zabieram parasol, zatem $P \sim Q$. Jest jasne, że jeśli pada deszcz i zabieram parasol, to raczej nie zmoknę, zatem $P \wedge Q \sim R$. Z reguły cięcia wynika wtedy raczej paradoksalna konkluzja $P \sim R$ — jeśli pada deszcz, to nie zmoknę. Skąd ten paradoks? Otóż zauważmy, że w inferencji $P \sim Q$ przesłanka jest pewna. Jeśli widzę, że pada deszcz, to trudno o zasadne wątpliwości. Konkluzja jest znacznie mniej pewna, wynika jedynie z mojego zwyczaju. Mogłem się zamyślić, mogło mi coś przeszkodzić i nie zabrałem parasola. Paradoks jest właśnie rezultatem tego, że relacja $\sim$ rozważana w przykładzie dopuszcza inferencje, w których konkluzja jest mniej pewna niż przesłanki.

Zasada cięcia w wydaniu niemonotonicznym mówi więc po prostu, że konkluzje są co najmniej tak pewne jak przesłanki. Zakładając, że dane inferencje spełniają regułę cięcia, nie dopuszczamy po prostu do rozważań inferencji, w których konkluzje są mniej pewne niż przesłanki.

Przykład ten wyraźnie wskazuje na rodzaj trudności, które napotykają logicy w badaniach potocznych rozumowań. Co ciekawe, reguła cięcia jest raczej powszechnie przyjmowana za ważną również w przypadku niemonotonicznym. Niektórzy logicy (na przykład Gabbay) wyrażali wprawdzie co do niej pewne wątpliwości, jednak, o ile wiem, nie badano dotąd systematyczne rozumowań, w których reguła cięcia nie byłaby w całej swej ogólności ważna. My również przyjmujemy tę regułę jako ważną. Warto jednak pamiętać, że istotnie ogranicza to klasę rozważanych inferencji.

Oto probabilistyczna interpretacja relacji $\sim$, która nie spełnia reguły cięcia właśnie dlatego, że konkluzje są tutaj mniej pewne niż przesłanki. Niech $P \sim_q Q$ wtedy i tylko wtedy, gdy prawdopodobieństwo $p(P \mid Q)$, że *Q* jest prawdziwe pod warunkiem, że prawdziwe jest *P*, jest większe od pewnej zadanej z góry liczby $q \leq 1$. Reguła cięcia nie jest spełniona dla $\sim_q$.

Przejdźmy do reguły, która w naszych rozważaniach występuje w tytułowej roli — reguły monotoniczności. Przytoczyliśmy już wyżej przykład, w którym jej zastosowanie prowadzi do paradoksalnych konkluzji. Rozważmy jeszcze jeden przykład, znacznie mniej naiwny i choć praktyczny, to raczej nie potoczny.

Załóżmy, że od obserwatora znajdującego się na Ziemi z prędkością światła oddala się rakieta, która przed siebie, w kierunku swego ruchu wysyła promień światła. Jaka jest prędkość światła rakiety dla ziemskiego obserwatora? Oczywiście najlepiej byłoby zmierzyć. Jest to jednak niemożliwe. Nie ma takiej rakiety i może nigdy nie będzie, a i z samym pomiarem byłyby pewnie kłopoty. Obserwator musi więc obliczyć prędkość na podstawie dostępnej mu wiedzy. Odpowiedź zależeć będzie więc od przesłanek, które przyjmie. Jeśli jedyną dostępną dla obserwatora wiedzą jest mechanika Newtona (oznaczymy przez Q koniunkcję jej praw), to dowiedzie on dość łatwo, że:

R : prędkość światła rakiety jest dwa razy większa niż prędkość światła wysłanego z Ziemi.

Gdyby obserwator zapoznał się ze szczególną teorią względności (oznaczymy koniunkcję jej praw przez P), nie wywiódłby R jako wniosku, lecz zdanie z nim sprzeczne. Mechanika Newtona jest szczególnym przypadkiem teorii względności i na gruncie praw fizyki możemy wywieść implikację $P \rightarrow Q$. Zatem mamy $\models P \rightarrow Q$ — prawdziwość implikacji na gruncie praw fizyki (nie jako tautologię klasyczną). Przykład ten narusza zatem zasadę monotoniczności w brzmieniu:

$$\text{M} \quad \frac{\models P \rightarrow Q, Q \sim R}{P \sim R} \quad \text{Monotoniczność}$$

Przykład ten jest typowy dla wszelkich przypadków, gdy rozważane zdanie R wynika ze starej teorii, przecząc jednocześnie teorii nowej, zaakceptowanej jako lepiej i precyzyjniej opisująca świat fizyczny, która przy tym nie odrzuca starej teorii w całości, a jedynie ogranicza zakres jej stosowalności. Wypada dla uczciwości przyznać, że stwierdzając $\models P \rightarrow Q$ — prawdziwość implikacji na gruncie praw fizyki, przesłizgujemy się niejako nad istotnymi problemami. Budowa poprawnej teorii fizycznej, na gruncie której implikacja taka byłaby prawdziwa, nie jest sprawą łatwą. Poruszamy się jednak w obrębie pewnych intuicji poprawnych z punktu widzenia fizyki. Znacznie bardziej pogładowe byłoby ujęcie teorii P jako (w odpowiednim sensie) szerszej niż Q i zastosowanie reguły monotoniczności w brzmieniu sformułowanym poprzednio dla relacji konsekwencji.

Zasady monotoniczności nie można więc przyjąć w pełnym jej brzmieniu, jeśli mowa o praktycznych czy potocznych inferencjach. Zamiast niej przyjmiemy następującą, pochodzącą od Makinsona⁸, znacznie słabszą regułę:

$$\text{CM} \quad \frac{P \sim Q, P \sim R}{P \wedge Q \sim R} \quad \text{Kumulatywność}^9$$

⁸D. Makinson, *General Theory of Cumulative Inferences*, w: Reinfrank, (wyd.) *Proceedings Second International Workshop on Non-Monotonic Reasoning, Lecture Notes in Computer Science* (Springer, Berlin).

⁹Dz. cyt.; Makinson używa nazwy „ostrożna monotoniczność”. Gabbay, dz. cyt., używa „nazwy słaba monotoniczność”.

Zauważmy, że nasz przykład nie daje w jej przypadku paradoksalnych wyników. Mamy rzeczywiście $P \sim Q$, ale już nieprawda, że $P \sim R$. Mamy wprowadzić $P \sim R'$, gdzie:

R' : prędkość światła rakiety jest taka sama jak prędkość światła na Ziemi; ale inferencja $P \wedge Q \sim R'$ powstała przez zastosowanie reguły nie jest paradoksalna.

Regułę tę nazywano również regułą kumulatywności oraz regułą triangulacji. Nazwa, którą my przyjmujemy — ostrożna monotoniczność — pochodzi od Makinsona. Słowa „kumulatywny” używać będziemy dalej w nieco innym znaczeniu.

Jaki jest intuicyjny sens ostrożnej monotoniczności? Otóż zezwala ona na dodawanie nowych przesłanek, nie dowolnych jednak, lecz tylko tych, które dają się wywieść z przesłanek uprzednio przyjętych. Gabbay przekonująco przytacza następujące jej rozumienie: jeśli P jest dostateczną przyczyną by uwierzyć w Q , jak też i dostateczną przyczyną, by uwierzyć w R , to $P \wedge Q$ powinno wystarczać dla uwierzenia w R samo bowiem P było do tego wystarczające, a Q oczekiwaliśmy na podstawie P .

Z pragmatycznego punktu widzenia ostrożna monotoniczność jest niezwykle istotną zasadą, wyraża ona bowiem ważną własność procesu poznawania. Istotnie, ucząc się mamy wyraźną skłonność do minimalizowania ewentualnych zmian naszych przekonań. Ostrożna monotoniczność i cięcie razem wzięte mówią, że jeśli fakty nowo poznane były oczekiwane, nasze przekonania nie ulegają zmianie. Mamy bowiem:

Twierdzenie. Reguły **Cut** i **CM** mogą być wyrażone razem za pomocą własności: jeśli $P \sim Q$, to $P \sim R$ wtedy i tylko wtedy, gdy $P \wedge Q \sim R$.

Kolejna reguła, którą omówimy, to reguła wprowadzania koniunkcji.

And

$$\frac{P \sim R, P \sim Q}{P \sim Q \wedge R}$$

Koniunkcyjność

Nie jest ona niezależna od poprzednio wprowadzonych reguł, daje się wywieść z **LLE**, **RW**, **Cut** i **CM**. Nie daje się wywieść z **LLE**, **RW** i **Cut**, a ponadto jest na gruncie tych reguł istotnie słabsza od **CM**. Wydaje się ona bardzo słabą regułą, mimo to można podać dla niej całkiem, jak sądzę, sensowny kontrprzykład. Otóż jeśli inferencję $P \sim Q$ rozumieć będziemy jako: Skoro P , to wobec braku informacji że jest inaczej Q . Jeśli zatem za P przyjmimy koniunkcję aksjomatów Peano¹⁰, a za Q Wielkie Twierdzenie Fermata, to mamy zarówno $P \sim Q$ jak i $P \sim \neg Q$. Można istotnie wzmocnić ten przykład. Niech $\sim$ oznacza inferencję: jeśli P i nie da się z P wywieść $\neg Q$, to Q . Jeśli teraz P jest koniunkcją aksjomatów teorii mnogości¹¹, zaś Q oznacza zdanie w niej niedowodnione, na przykład hipotezę kontinuum, to zarówno $P \sim Q$ jak i $P \sim \neg Q$. Konkluzja $P \sim Q \wedge \neg Q$ wynikają z **And** jest raczej paradoksalna¹².

¹⁰Musiaby to być nieskończona koniunkcja lub jakaś arytmetyka drugiego rzędu, mniejsza jednak o szczegóły.

¹¹Patrz wyżej.

¹²Choć w tym wypadku może, mimo to, być prawdziwa. Rozważając teorię, o której wiemy że jest niesprzeczna, otrzymamy konkluzję jawnie fałszywą.

Reguła **And** wyraża zatem pewnego rodzaju racjonalność inferencji, mówiąc, że wyprowadzone konkluzje muszą być zgodne między sobą.

Zdefiniujemy teraz dwie ważne klasy systemów niemonotonicznych.

Relację $\sim$ nazwiemy *minimalną* wtedy i tylko wtedy, gdy spełnia regułę zwrotności i jest zamknięta na **LLE**, **RW**, **Cut** oraz **And**. Najmniejszy wśród systemów minimalnych oznaczać będziemy przez **A**.

Relację $\sim$ nazwiemy *kumulatywną* wtedy i tylko wtedy, gdy spełnia regułę zwrotności i jest zamknięta na **LLE**, **RW**, **Cut** oraz **CM**. Najmniejszy spośród systemów kumulatywnych oznaczać będziemy przez **C**.

Łatwo dowieść, że reguła **And** jest spełniona w każdej inferencji kumulatywnej. Każda inferencja kumulatywna jest więc minimalna, choć, jak pokażemy w części semantycznej, nie odwrotnie.

Logika minimalna jest słaba. Niewiele można w niej dowieść. Rozważmy jednak ważny sekwent:

$$\text{MPC} \quad \frac{P \sim Q \rightarrow R, P \sim Q}{P \sim R} \quad \text{Odrywanie}$$

Wyraża on stosowalność reguły odrywania we wnioskowaniach niemonotonicznych. Reguła **MPC** jest ważna w dowolnej logice minimalnej. Systemy kumulatywne są znacznie silniejsze. Rozważmy następujący sekwent ważny w dowolnej logice kumulatywnej:

$$\text{Eq} \quad \frac{P \sim Q, Q \sim P, P \sim R}{Q \wedge R} \quad \text{Równoważność}$$

Reguła ta jest istotnym wzmocnieniem reguły **LLE**. O ile **LLE** pozwala zastąpić przesłankę przez zdanie logicznie jej równoważne, o tyle **Eq** dopuszcza tego rodzaju zastępowanie w szerszym zakresie — można mianowicie zastąpić daną przesłankę dowolnym zdaniem jej równoważnym w sensie relacji inferencji $\sim$. Reguła ta ma ważne zastosowania przy budowaniu modeli kanonicznych dla logik niemonotonicznych.

W rozważanych dotąd regułach w istotnie niemonotoniczny sposób występował, w gruncie rzeczy, jedynie spójnik koniunkcji. Inne spójniki wprowadzane były poprzez reguły **LLE** i **RW** na podstawie tautologii klasycznych. Oto pierwsza z reguł wprowadzania alternatywy:

$$\text{IO} \quad \frac{P \vee Q \sim P, P \sim R}{P \vee Q \sim R} \quad \text{Wprowadzanie alternatywy}$$

Kryją się za nią interesujące intuicje. Niedostatek miejsca każe jednak nie zatrzymywać się nad tym. Zauważmy jedynie, że reguła **IO** jest, podobnie jak **Eq**, ważna w każdym systemie kumulatywnym.

Oto kilka reguł istotnie silniejszych niż reguły logiki kumulatywnej *C*. Stanowią one w istocie warianty zasady monotoniczności.

EHD	$\frac{P \sim Q \rightarrow R}{P \wedge Q \sim R}$	<i>Twierdzenie o dedukcji</i>
T	$\frac{P \sim Q, Q \sim R}{P \sim R}$	<i>Przechodność</i>
Cpt	$\frac{P \sim Q}{\neg Q \sim \neg P}$	<i>Kontrapozycja</i>

W dowolnym systemie kumulatywnym reguły *M*, *EHD* i *T* są równoważne, a ponadto każda z nich wynika z reguły kontrapozycji *Ctp*. W systemach kumulatywnych reguły *EHD*, *T* są po prostu odmiennymi formami zapisu reguły słabej formy monotoniczności *M*. Reguła kontrapozycji jest pewną wersją monotoniczności istotnie silniejszą niż *M*. Powróćmy do tej sprawy podczas rozważań semantycznych.

Zdefiniujemy dalej jeszcze pewne inne (węższe) klasy logik niemonotonicznych. Wygodnie będzie jednak wprowadzić wcześniej interpretacje semantyczne dla zdefiniowanych wyżej logik.

Modele

Zajmiemy się teraz semantycznymi interpretacjami niemonotonicznych inferencji. Aby znaleźć odpowiedni punkt odniesienia, rozpoczniemy od kilku uwag o prostej semantyce sytuacyjnej dla logik monotonicznych, a konkretnie dla logiki klasycznej. Semantyka tego rodzaju została przedstawiona na początku poprzedniej części jako punkt wyjścia dla definicji systemów inferencji niemonotonicznych. Przypomnijmy ją krótko.

Z językiem klasycznym *L* stowarzyszymy semantykę zadaną przez zbiór *U*, którego elementy nazywać będziemy światami, a rozumieć je będziemy jako klasyczne waluacje. W naturalny sposób zadana jest tym samym relacja spełniania = pomiędzy zdaniami a światami spełniająca sformułowane poprzednio klasyczne warunki. Powiemy, że z *P* wynika *Q*, symbolicznie $P \models Q$, wtedy i tylko wtedy, gdy *Q* spełnione jest w każdym świecie, w którym spełnione jest *P*.

Warunek, aby konkluzja prawdziwa była w każdym świecie, w którym prawdziwa jest przesłanka, daje w wyniku monotoniczność rozważanej logiki. Główna idea semantyki niemonotonicznych inferencji zasadza się na osłabieniu tego właśnie żądania. Omówione przykłady wskazują, że rozważanie wszystkich światów jest stanowczo zbyt mocne. W potocznych wnioskowaniach bierzemy pod uwagę poza światem aktualnym co najwyżej niektóre inne możliwe światy — te mianowicie, które najlepiej pasują do przesłanki i konkluzji. Można powiedzieć, w pierwszym przy-

bliżeniu, że $P \sim Q$ wtedy i tylko wtedy, gdy Q jest prawdziwe w tych spośród światów spełniających P , które są najbardziej normalne, mają własności najbardziej zbliżone do świata aktualnego. Jeśli podamy sposób znajdowania światów normalnych w dowolnym zbiorze światów, otrzymamy semantykę dla całkiem szerokiej, choć węższej niż klasa logik kumulatywnych, klasy logik — tak zwanych logik preferencyjnych.

Aby otrzymać semantykę dla możliwie szerokiej klasy inferencji, niezbędne jest rozważanie bardziej skomplikowanych modeli.

Model kumulatywny to uporządkowana trójka: $W = (S, l, <)$ taka, że S jest dowolnym ustalonym zbiorem zwanym zbiorem stanów, $l \longrightarrow 2^u$ — funkcją, która dowolnemu stanowi s przyporządkowuje pewien zbiór $l(s)$ waluacji. $<$ jest binarną relacją na S . Zakładamy ponadto, że model W spełnia pewne dodatkowe warunki, które sformułujemy dalej.

Jakie intuicje kryją się za tą definicją? S jest zbiorem dopuszczalnych stanów, w jakich może znajdować się podmiot poznający w czasie wyciągania wniosków. W zależności od zmiany sytuacji przeskakujemy od jednego stanu do drugiego, uznając za prawdziwe lub odrzucając te lub inne zdania, stanowiące dodatkowe przesłanki podczas wnioskowania. Z każdym stanem s związana jest pewna wiedza o świecie, dana przez waluacje (światy) ze zbioru $l(s)$.

Zilustrujmy pojęcie zbioru stanów za pomocą przykładu, który często można spotkać w literaturze przedmiotu.

Rozważmy zdania:

P Sx : x jest Szwedem.

Q Px : x jest protestantem.

R Rx : Rodzice x 'a są katolikami.

oraz:

1. Jeśli x jest Szwedem, to rodzice x 'a nie są katolikami.
2. Jeśli x jest Szwedem, to x jest protestantem.
3. Jeśli x jest Szwedem i rodzice x 'a są katolikami, to x nie jest protestantem.

Załóżmy P jako przesłankę wnioskowania. W typowym stanie wiedzy s spełnione są zdania 1, 2, 3. P pociąga zatem za sobą w tym stanie zdania Q . Dodatkowe informacje mogą zmienić stan naszej wiedzy, nie tylko wzbogacając go, lecz także zaprzeczając pewnym przyjmowanym dotąd danym. Jeśli więc na przykład dowiemy się, że R , zmuszeni będziemy przeskoczyć do innego stanu s' , w którym zdanie 1 i 2 nie będą prawdziwe. W tym nowym stanie zdanie P pociąga zdanie $\neg Q$. Mamy tu zatem rzeczywiście do czynienia z rozumowaniem niemonotonicznym, dodanie bowiem nowej przesłanki daje nam zaprzeczenie poprzedniej konkluzji.

Przykład ten nie opisuje oczywiście, w jaki sposób działa model kumulatywny, a jedynie wskazuje, jak należy rozumieć pojęcie stanu. Każdy zbiór waluacji wyznacza pewną teorię klasyczną, również dowolna teoria klasyczna jest wyznaczona przez pewien zbiór waluacji. Zależność ta nie jest wprawdzie wzajemnie jednoznaczna, niemniej jednak zbiór waluacji $l(s)$, stowarzyszony ze stanem s , może być traktowany

jako pewna teoria klasyczna, a mianowicie zbiór tych wszystkich zdań, które przyjmują wartość 1 na wszystkich waluacjach z $l(s)$. $l(s)$ jest zatem zasobem wiedzy w stanie s .

Wydaje się, że opisany mechanizm oddaje pewne istotne cechy tego, w jaki sposób wykorzystujemy naszą wiedzę. Nie posiadamy jednej teorii, dobrej na każdą okoliczność. Wiedza, którą przyjmujemy jako podstawę dla rozważań, to zestaw wyspecjalizowanych teorii opisujących wąskie fragmenty rzeczywistości. Teorie te mogą opisywać odrębne (rozłączne) fragmenty rzeczywistości, formalnie biorąc zdania kontyngentne z dwóch różnych teorii tego rodzaju nie mają wtedy wspólnych zmiennych. Może się również zdarzyć tak, że teorie te opisują na dwa różne (nawet sprzeczne) sposoby te same fragmenty rzeczywistości. Czego z pewnością brakuje w tym opisie, to mechanizmu wyboru stanu, który pozwalałby przechodzić w odpowiedni sposób od stanu do stanu. Modele kumulatywne nie posiadają takiego mechanizmu w pełnej postaci. Nie jest on potrzebny do spełniania funkcji, którą mają pełnić. Warto jednak zwrócić uwagę na problem zastosowania odpowiednio wzbogaconych modeli kumulatywnych do interpretacji modelowania procesów wykorzystywania i przetwarzania wiedzy.

Modele kumulatywne wyposażone są w relację preferencji $<$, która stanowi pewien mechanizm wyróżniania niektórych stanów. Relacja ta w swym zamyśle służy do porządkowania zbioru stanów od bardziej do mniej normalnych (typowych). $s < s'$ oznacza zatem, że s jest bardziej typowy (mniej nietypowy) niż s' . Słowo „porządkowanie” rozumieć tutaj należy w bardzo szerokim sensie, nie zakładamy bowiem, że relacja spełnia typowe warunki zwrotności, antysymetryczności i przechodniości. O relacji $<$ zakładamy jedynie, że jest przeciwwrotna, to znaczy $s \not< s$. Do intuicji związanych z relacją $<$ powrócimy jeszcze. Nie unikniemy jednakże pewnej dawki technicznych, choć prostych definicji, istotnych dla dalszych rozważań.

Definicja. Niech $X \subseteq S$ będzie dowolnym niepustym zbiorem stanów. Element $s_0 \in X$ nazywamy minimalnym w X , gdy $s_0 < s$ dla dowolnego $s \in X$ takiego, że $s_0 < s$ lub $s < s_0$ (jest mniejszy od dowolnego z nim porównywalnego stanu). Zbiór wszystkich elementów minimalnych w X oznaczamy przez $\min(X)$. Zbiór X nazywamy gładkim, gdy X jest pusty lub zawiera elementy minimalne.

Zdanie P jest spełnione w stanie s , symbolicznie $s \models P$, wtedy i tylko wtedy, gdy P jest spełnione przez dowolną waluację ze zbioru $l(s)$ lub, co na jedno wychodzi, należy do teorii $l(s)$. Zbiór wszystkich stanów, w których spełnione jest zdanie P , oznaczać będziemy przez $\hat{P}$.

Modelem kumulatywnym nazywaliśmy będziemy uporządkowaną trójkę: $W = (S, l, <)$, gdzie S jest dowolnym ustalonym zbiorem, $l: S \rightarrow 2^U$ — funkcją, $<$ jest binarną relacją na S . Jeśli ponadto dowolny zbiór stanów postaci $\hat{P}$ jest gładki, to model W nazywamy gładkim. Dowolny model kumulatywny W wyznacza następującą relację $\sim_w: P \sim_w Q$ wtedy i tylko wtedy, gdy $\min(\hat{P}) \subseteq Q$.

Semantyka dla logik niemonotonicznych zdefiniowana powyżej jest pewną modyfikacją konstrukcji Krausa, Lehmana i Magidora¹³. Wprowadzone różnice mają na celu uzyskanie takich modeli kanonicznych dla logik kumulatywnych, które będą

posiadały elegancką strukturę opartą na algebrze Lindenbauma jako zbiorze dopuszczalnych stanów. Wyniki zaprezentowane poniżej są modyfikacjami odpowiednich wyników uzyskanych przez Krausa, Lehmana i Magidora¹⁴

Pierwowzorem dla konstrukcji Krausa, Lehmana, Magidora były koncepcje Shohama¹⁵; pewną inną semantyczną interpretację rozumowań niemonotonicznych proponował wcześniej Makinson¹⁶.

Modele kumulatywne w zadowalający sposób charakteryzują dowolne logiki kumulatywne. Mamy bowiem:

Twierdzenie. Jeśli W jest gładkim modelem kumulatywnym, to $\sim_w$ jest logiką kumulatywną, relacja $\sim_w$ spełnia zatem reguły **Ref**, **LLE**, **RW**, **Cut**, **CM**. Dla dowolnej logiki kumulatywnej $\sim$ istnieje gładki model kumulatywny $W=(S, I, <)$ taki, że $\sim = \sim_w$.

Model kumulatywny, o którym mowa wyżej, będziemy nazywać modelem kanonicznym dla $\sim$. Nie jest znana jak dotąd charakterystyka logik minimalnych za pomocą modeli kumulatywnych, choć badanie klasy wszystkich modeli kumulatywnych (bez założenia gładkości) pozwala stosunkowo łatwo na znaczne wzmocnienie sformułowanego wyżej twierdzenia.

Konstrukcja modelu kanonicznego wymaga pewnych szczegółowych i raczej technicznych rozważań. Wydaje się ona jednakże na tyle istotna z filozoficznego punktu widzenia, że zaprezentujemy jej ideę pomijając szczegóły techniczne. Jak sądzę, pozwoli to na głębsze zrozumienie istoty przedstawianej tutaj interpretacji.

Niech dana będzie logika kumulatywna $\sim$ w języku L . Skonstruujemy model kanoniczny $W=(S, I, <)$ dla $\sim$. Relacja $\equiv$ klasycznej równoważności zdań $P \equiv Q$ wtedy i tylko wtedy, gdy $\models P \leftrightarrow Q$, jest, jak wiadomo, relacją równoważności na L . Jest to co więcej, relacja kongruencji, choć w naszych rozważaniach ten fakt nie będzie miał zastosowania. Utożsamiając zdania pozostające ze sobą w relacji $\equiv$, to jest zdania logicznie równoważne, otrzymamy algebrę Boole'a znaną jako klasyczna algebra Lindenbauma języka L . Za zbiór stanów S konstruowanego modelu kanonicznego W przyjmiemy tę właśnie algebrę: $S=L/\equiv$. Dla dowolnego $P \in L$, $\bar{P}$ oznacza klasę abstrakcji zdania P względem relacji $\equiv$, to znaczy zbiór wszystkich zdań klasycznie równoważnych zdaniu P . Bez obaw o nieporozumienie możemy zatem oznaczać stany modelu W przez $\bar{P}$, $\bar{Q}$ itd. Relację preferencji definiujemy w W następująco: $\bar{P} < \bar{Q}$ wtw $Q \sim P$ oraz $\bar{P} \neq \bar{Q}$. Łatwo sprawdzić, że definicja powyższa nie zależy od wyboru reprezentanta z $\bar{P}$ ani z $\bar{Q}$. Dla zakończenia konstrukcji pozostaje zdefiniować funkcję I .

Wymaga to zdefiniowania pojęcia świata normalnego. Świat $m \in U$ nazywamy normalnym dla P wtedy i tylko wtedy, gdy

¹³ Dz. cyt.

¹⁴ Dz. cyt.

¹⁵ Y. Shoham, *A Semantical Approach to Non-monotonic Logics*, w: *Proceedings Logics in Computer Sciences*, Ithaca (N.Y.) 1987, s. 275-279.

¹⁶ Dz. cyt.

$$\forall_{Q \in L} P \sim Q \Rightarrow m \models Q$$

Jeśli $\sim$ spełnia **Ref.**, **RW** and **And**, i $P, Q \in L$, wtedy $P \sim Q$ wtedy i tylko wtedy, gdy wszystkie światy normalne dla P spełniają Q . Połóżmy:

$$l(P) = \{m : m \text{ jest światem normalnym dla } P\}$$

Zreasumujmy. Stanami modelu są elementy klasycznej algebry Lindenbauma. Indeksowane są one światami normalnymi, relacja preferencji jest wyznaczona przez relację $\sim$. Stan $\hat{P}$ jest najmniejszy w P , co więcej, jeśli R jest najmniejszym stanem w $\hat{P}$, to $P \sim R$ i $R \sim P$. Łatwo sprawdzić, że model kanoniczny W dla logiki $\sim$ reprezentuje ją, to znaczy $\sim = \sim_W$.

Zdefiniujemy dalej kilka właściwych podklas klasy logik kumulatywnych. Dla każdej z nich przedstawimy twierdzenie o reprezentacji analogiczne do sformułowanego powyżej, w którym model gładki zastąpimy modelem o lepszych własnościach. Będziemy poprawiać model kumulatywny w dwóch kierunkach. Pierwszy z nich to poprawianie własności relacji preferencji $<$, drugi polega na nakładaniu warunków ograniczających dowolność w indeksowaniu stanów światami. Z konstrukcji modelu kanonicznego wynika, że dla dowolnej logiki kumulatywnej można znaleźć model kanoniczny z algebrą Lindenbauma jako zbiorem stanów. Negatywnym skutkiem poprawy własności zbioru stanów byłoby jednak znaczne pogorszenie własności relacji preferencji $<$. Nic za darmo. Dla dalszych rozważań będzie miała zastosowanie nieco zmodyfikowana konstrukcja modelu kanonicznego. Zastąpimy mianowicie relację klasycznej równoważności relacją niemonotonicznej równoważności $P \equiv Q$ wtedy i tylko wtedy, gdy $P \sim Q$ oraz $Q \sim P$. Relacja ta, jak łatwo pokazać, zawiera w sobie relację klasycznej równoważności - „skleja” z dowolnym zdaniem wszystkie zdania logicznie z nim równoważne oraz wiele innych zdań ponadto. Stany modelu kanonicznego to klasy abstrakcji tej relacji. Zbiór stanów będzie zatem inny dla każdej logiki niemonotonicznej - poprzednio był jeden uniwersalny dla wszystkich logik. Jego struktura jest bardzo trudna do określenia. Za tą cenę uzyskamy jednak poprawę własności relacji $<$. Już w przypadku logik kumulatywnych stanie się antysymetryczna, czego konsekwencją jest jedyność elementów minimalnych. Stan $\hat{P}$ jest wtedy jedynym minimalnym stanem w zbiorze stanów $\hat{P}$.

Zgodnie z intuicjami dotyczącymi relacji $<$ opisanymi wcześniej, świat „mniejszy” to świat bardziej normalny. Naturalne wydaje się pytanie, czy relacja ta jest relacją częściowego porządku na zbiorze S , czy raczej, co w tym wypadku na jedno wychodzi, pytanie o własności logik kumulatywnych wyznaczonych przez klasę modeli kumulatywnych, w których relacja $<$ jest częściowym porządkiem. Modele takie nazywać będziemy modelami uporządkowanymi.

Rozważmy następującą regułę:

$$\frac{P_0 \sim P_1, \quad P_1 \sim P_2, \dots, P_{k-1} \sim P_k, \quad P_k \sim P_0}{P_0 \sim P_k} \quad \text{Pętla}$$

Istnieją modele kumulatywne, w których nie jest ona spełniona. Zatem, na podstawie twierdzenia o reprezentacji, nie jest ona wywiedlna w systemie logiki kumulatywnej. Prawdą jest natomiast, że kumulatywna relacja $\sim$ spełnia regułę pętli wtedy i tylko wtedy, gdy jest ona wyznaczona przez pewien model uporządkowany.

Dwie uwagi na marginesie. Definicję modeli uporządkowanych potraktować możemy jako pretekst dla głębszej analizy intuicyjnych własności relacji preferencji oraz ich roli w modelu. W potocznym użyciu pojęcie „bardziej normalnego świata” raczej nie występuje. W każdym razie występuje na tyle rzadko i nieprecyzyjnie, że nie sposób czynić potocznego użycia arbitrem spełniania tych czy innych własności relacji preferencji. Wydaje się, że można znaleźć spektakularne kontrprzykłady zarówno dla przeciwwrotności, jak i przechodniości relacji $<$. Spełnienie reguły pętli istotnie ogranicza zatem zakres możliwych denotacji relacji $<$. Warto jednak zauważyć, że relacja $<$ jest nam potrzebna jedynie do wyróżnienia odpowiednich stanów z dowolnego zbioru postaci P — inaczej mówiąc, do określenia odpowiedniej funkcji wyboru. Relację $<$ definiujemy, aby wybrać stany minimalne. Wybór taki jest w zasadzie możliwy bez określania relacji $<$. Nie trzeba wiedzieć, czy 2^{327} jest większe od 3^{241} , aby stwierdzić, że 0 jest mniejsze od nich obu. W terminach funkcji wyboru można wyrazić znacznie więcej niż w terminach minimum i relacji $<$. Ten kierunek badań jest jak dotąd bardzo słabo wyeksploatowany. Funkcją wyboru na zbiorze stanów S nazywać będziemy dowolną funkcję c , która zbiorowi stanów X przyporządkowuje pewien jego podzbiór $c(X)$. W terminach własności funkcji wyboru c daje się wyrazić semantycznie reguły, które były przedmiotem naszych rozważań. A to w następujący sposób:

jeśli $c(X) \cap Y \subseteq X$, to $c(X) \subseteq c(Y)$, **Cut**

$c(X) \cap Y \subseteq c(X \cap Y)$, **AND**

jeśli $c(X) \cap Y \subseteq X$, to $c(Y) \subseteq c(X)$, **CM**

$c(X \cap Y) \subseteq c(X) \cap c(Y)$, **Or**

Wydaje się, że dowolna reguła może być semantycznie wyrażona w podobny sposób. Najciekawsze jest jednak tutaj to, iż dzięki funkcji wyboru otrzymujemy związek problematyki rozumowań niemonotonicznych z teorią wyboru społecznego. Co ciekawe, własności funkcji wyboru przedstawione powyżej zostały po raz pierwszy sformułowane i zbadane przez specjalistów w tej właśnie dziedzinie, między innymi w związku ze sławnym Twierdzeniem Arrowa. Powiązanie to otwiera, jak się wydaje szczególnie duże i jak dotąd słabo zbadane możliwości dalszych badań.

Uwaga druga. Reguła pętli nie zakłada występowania spójników w języku L . Niewiele rozważa się tego rodzaju reguł, językowo uniwersalnych, to znaczy niezmienniczych względem zmiany języka. Z reguł rozważanych dotąd jedynie zwrotność, równoważność i przechodność nie zawierają spójników. Ciekawe wydaje się pytanie

o siłę wyrazu tego rodzaju reguł. O ile wiem, nie wiadomo do tej pory, jakie logiki kumulatywne (czy też ogólnie relacje inferencji) można zdefiniować przy użyciu reguł tego tylko rodzaju.

W regułach dotąd rozważanych jedynym spójnikiem charakteryzowanym również niemonotonicznie była koniunkcja. Zdefiniujemy teraz jedną z najważniejszych klas logik niemonotonicznych, klasę logik preferencyjnych, wykorzystując kolejny spójnik: alternatywę.

Logiką preferencyjną nazywamy dowolną logikę kumulatywną, spełniającą następującą regułę:

$$\text{Or} \quad \frac{P \sim R, Q \sim R}{P \vee Q \sim R}$$

System logiki preferencyjnej był rozważany przez Adamsa¹⁷, J. Pearl i H. Geffner¹⁸ zaproponowali ją jako minimalny rdzeń systemu rozumowań niemonotonicznych.

Reguła **Or** nie implikuje monotoniczności, stąd nasza skłonność do zaakceptowania jej. Jak to zwykle z regułami logik niemonotonicznych bywa, łatwo o przykłady potwierdzające ważność tej reguły. Reguła ta, jak się wydaje, działa należycie, gdy relację $\sim$ rozumiemy jako „jeśli..., to zwykle...”. Jednakże pewne interpretacje epistemiczne obalają regułę **Or**. Przykłady: „Jeśli P , to umiem przekonać innych, że R ,” czy też „Jeśli wiem, że P , to uzasadnię R ,” lub wręcz „Jeśli wiem, że P , to R .” Wiedza że P i Q wzięta z osobna może pozwalać dowieść R , podczas gdy $P \vee Q$ może być wiedzą pustą, z której nic nie wywnioskujemy. Jak widać, blisko stąd do problemów znanych z intuicjonizmu i logik konstruktywnych.

Logiki preferencyjne posiadają elegancką semantykę. Relacja $\sim$ jest bowiem logiką preferencyjną wtedy i tylko wtedy, gdy istnieje model kumulatywny $W=(S, I, <)$ taki, że $I(s)$ jest jednoelementowym zbiorem światów (czyli pojedynczym światem) natomiast $<$ jest częściowym porządkiem, a ponadto $\sim = \sim_w$. Model taki nazywamy modelem preferencyjnym. Oto dwie ważne konsekwencje reguły **Or**:

$$\text{S} \quad \frac{P \wedge Q \sim R}{P \sim Q \rightarrow R}$$

$$\text{D} \quad \frac{P \wedge \neg Q \sim R, P \wedge Q \sim R}{P \sim R}$$

¹⁷E.W. Adams, *Probability and the logic of conditional*, w: J. Hintikka, P. Suppes, (wyd.), *Aspects of Inductive Logic*, Amsterdam 1966.

¹⁸J. Pearl, H. Geffner, *Probabilistic semantics for a subset of default reasoning*, TR CSD-8700XX, R-93-III, Computer Science Department, University of California, Los Angeles, CA (1988).

Pierwsza z powyższych reguł jest wariantem twierdzenia o dedukcji w silną stronę, druga pochodzi od Makinsona i jest w swej istocie zasadą dowodu przez przypadki.

Stany modelu preferencyjnego indeksowane są pojedynczymi światami, czyli pojedynczymi waluacjami, które jak wiadomo odpowiadają teoriom maksymalnym. W modelach kumulatywnych stany indeksowane wieloma światami naraz mogą spełniając alternatywę nie spełniać żadnego z jej członów. W modelach preferencyjnych stan spełnia alternatywę wtedy i tylko wtedy, gdy spełnia jeden z jej członów. Fakt ten wiąże się ściśle ze znaną własnością teorii maksymalnych. Interpretacja rozumowań wyznaczonych przez modele preferencyjne jest zatem następująca. Podmiot poznający ma w każdym stanie modelu preferencyjnego pełną wiedzę o świecie. Oczywiście wiedza ta nie zawsze musi być prawdziwa. Powiedzmy tak: tam gdzie nie znamy prawdy przyjmujemy za prawdę coś co z jakichś względów wydaje nam się najbliższe prawdy.

Omówimy jeszcze dwie klasy logik. Logiką kumulatywnie monotoniczną (**CM**-logiką) nazywamy dowolną logikę kumulatywną spełniającą regułę monotoniczności. Logiką monotoniczną (**M**-logiką) nazywamy dowolną logikę kumulatywną spełniającą regułę kontrapozycji.

Zauważmy, że przy przyjętych założeniach jest tylko jedna **M**-logika — logika klasyczna. Logiki wyżej zdefiniowane różnią się regułą **Or** spełnioną przez **M**-logiki, a nie spełnioną przez **CM**-logiki. Obie klasy wyznaczone są przez pewne naturalne klasy modeli kumulatywnych. Wiemy już, jak w modelach realizować regułę **Or**. Kilka słów o regule monotoniczności. Jakie modele ją spełniają? Otóż jak widać z definicji klasycznego wynikania są to te modele, w których pojęcia świata (stanu) dowolnego i minimalnego pokrywają się. Wymóg, aby każdy stan był minimalny w danym zbiorze, łatwo realizuje się technicznie. Wystarczy zażądać, aby relacja $<$ była pusta, to jest, aby żadne dwa światy nie były porównywalne. Odpowiada to intuicji, w której każdy świat jest „specjalny”, nieporównywalny z żadnym innym pod względem normalności.

Podsumowanie. Każda **CM**-logika wyznaczona jest przez model kumulatywny z pustą relacją preferencji. **M**-logika wyznaczona jest przez model preferencyjny z pustą relacją preferencji.

Prawdę mówiąc omówiliśmy jedynie jedną z wielu, choć niewątpliwie ważną klasę logik niemonotonicznych. Pominęliśmy, spośród najważniejszych: logiki rozumowań zaocznych (*default logics*), logiki cyrumskrypcji, podejście probabilistyczne. Najważniejszym chyba z pominiętych zagadnień są okresy kontrfaktyczne oraz ich semantyczny i syntaktyczny związek z rozumowaniami niemonotonicznymi. Ale bo też i nie systematyczna prezentacja tej problematyki była naszym celem. Badania logik niemonotonicznych definiowalnych semantycznie za pomocą modeli kumulatywnych leżą w głównym nurcie problematyki rozumowań niemonotonicznych. Nie obejmując całości, pozwalają one zrozumieć istotę problemów, które bada logika niemonotoniczna.