

PHILOSOPHIA

czasopismo studenckie



Redaktor naczelny: Błażej Gębura
Redakcja: Marcin Czakon, ks. Krzysztof Jaworski, Anna Karczewska
Projekt okładki: Marcin Jabłonski
Skład: Marcin Czakon

Recenzenci tomu:
ks. dr Tomasz Huzarek
dr hab. Marek Lechniak
dr Piotr Lipski
ks. dr hab. Marcin Tkaczyk, prof. KUL

Philosophia, 2 (33) 2012
Koło Filozoficzne Studentów KUL
Al. Raławickie 14, 20-950 Lublin
philosophia@student.kul.pl
<http://www.kul.pl/philosophia>
ISSN 1895-9431

Spis Treści

ks. Marcin Ferdynus	
<i>Przedłużanie ludzkiego życia: czy możliwa jest biologiczna nieśmiertelność?</i>	5
Błażej Gębura	
<i>Stanowisko Hume'a w kwestii istnienia Boga.</i>	13
ks. Bartłomiej Krzos	
<i>O rodzajach i genezie norm.</i>	19
Anita Mira Sobczyk	
<i>Johna Rogersa Searle'a próba ukazania możliwości wyprowadzenia powinności z faktu.</i>	29
ks. Krzysztof Jaworski	
<i>Geneza pojęcia hiperzbioru.</i>	39
Paweł Zarosa	
<i>Futura contingentia a logika trójwartościowa.</i>	49
Marcin Czakon	
<i>Pewnego rodzaju logiki niemonotoniczne jako model dla wnioskowania zawodnego.</i>	55
Anna Karczevska	
<i>Paradoksy łańcuchowe a logika rozmyta.</i>	62

A R T Y K U Ł Y

Philosophia, 2 (33) 2012, s. 5-12

ks. Marcin Ferdynus
Katolicki Uniwersytet Lubelski Jana Pawła II

Przedłużanie ludzkiego życia: czy możliwa jest biologiczna nieśmiertelność?

Zamiast wprowadzenia

Bernard Williams, w artykule *Sprawa Makropulos: Refleksje nad nudą nieśmiertelności*, przedstawia sprawę kobiety, żyjącej na dworze XVI-wiecznego cesarza, na której ojciec - medyk testuje siłę działania eliksiru życia. Stan znudzenia, oziębłości i obojętności życiem sprawia, że Elina Makropulos osiągnąwszy 342 lata, rozczarowana perspektywą nieśmiertelności, odmawia ponownego zażycia eliksiru i umiera¹. Sytuacja, w której znalazła się owa kobieta, sugeruje tyle, że "(...) śmierć niekoniecznie jest złem, i to nie tylko w takim sensie, w jakim niemal każdy by się z tym zgodził, gdy kładzie ona kres wielkiemu cierpieniu, lecz także w bardziej swoistym znaczeniu, że może być dobrą rzeczą nie żyć zbyt długo"². Sugestie Williamsa nie zatrzymują się tylko i wyłącznie na tym stwierdzeniu. Brytyjski etyk twierdzi, że posiadamy pewne powody, dla których unikamy śmierci. Są to - jego zdaniem - kategoryczne pragnienia³. Dopóki te pragnienia istnieją, podtrzymują chęć dalszego życia. Kiedy jednak pragnienia kategoryczne znikają - jak to miało miejsce w przypadku Makropulos - wówczas nie ma najmniejszych powodów, dla których przedłużanie życia w nieskończoność posiadałoby jakiegokolwiek znaczenie⁴.

Do podobnych wniosków dochodzi także Leon Kass. Amerykański filozof sugeruje, że istnieją pewne korzyści, dla których biologiczna śmiertelność mogłaby okazać się pożądana. Retoryczne pytania, za pomocą których Kass chce przekonać do swoich racji, wiąże z argumentami mającymi przemówić na rzecz śmiertelności. Wśród nich wymienia cztery takie argumenty⁵: (1) "interes i zaangażowanie" (Czy rodzice chcieliby przedłużać doświadczenia związane z wychowaniem dzieci o kolejne dziesięć lat?), (2) "znaczenie i aspiracje" (Czy traktowalibyśmy nasze życie

¹ B. Williams, *Sprawa Makropulos: refleksje nad nudą nieśmiertelności*, tłum. T. Duliński, [w:] Tenże, *Ile wolności powinna mieć wola?*, Warszawa: Fundacja Aletheia 1999, s. 65-66.

² Tamże, s. 66.

³ Tamże, s. 72.

⁴ Tamże, s. 74, 86.

⁵ L. Kass, *Life, Liberty and the Defense of Dignity. The Challenge for Bioethics*, New York - London: Encounter Books 2004, s. 266-267; Zob. też: B. Chyrowicz, *Życie: długość, jakość i moralność, Wprowadzenie*, [w:] Tenże, *Przedłużanie życia jako problem moralny*, Lublin: Wydawnictwo TN KUL 2008, s. 13-15.

na serio, jeśli byłoby wiadomo, że posiadamy nieskończony czas jego trwania?), (3) „*piękno i miłość*” (Czy to nie nasza śmiertelność jest powodem wzmożonej aprecjacji piękna, wartości i troski wobec tych, których kochamy?), (4) „*cnota i moralna doskonałość*” (Czy śmiertelność nie daje nam okazji do tego, abyśmy docenili duchowy wymiar naszego życia?).

Intuicje Williamsa i Kassa w jakiejś mierze wydają się słuszne. Jednakże kiedy chcemy odnieść się do życiowego pragmatyzmu, a szczególnie do sytuacji stawiających nas w obliczu możliwości utraty życia, rzeczywistość wydaje się nieco inna. W okolicznościach grożącego nam niebezpieczeństwa staramy się raczej chronić nasze życie niż, nic nie robiąc, oczekiwać na niechybną śmierć. Czynimy tak, ponieważ uznajemy, że życie posiada wartość samą w sobie. Pisał już o tym Władysław Tatarkiewicz: „*Póki jesteśmy pewni życia, traktujemy je jako środek, ale gdy czujemy, że grozi nam niebezpieczeństwo, wtedy ono staje się celem i dla ratowania go gotowi jesteśmy wszystko poświęcić*”⁶. Zatem musi istnieć bardziej fundamentalna racja, aniżeli pragnienia kategoryczne, która sprawia, że podejmujemy wysiłki chroniące nasze życie, chcąc zapewnić sobie niekończący się czas trwania - biologiczną nieśmiertelność. Cóż to za racja?

Wydaje się, że tą racją jest inklinacja do zachowania swojego istnienia, a przez to, do zachowania swojej bytowej tożsamości. Bynajmniej nie chodzi tutaj jedynie o ochronę swego życia przed zagrożeniem zewnętrznym, ale o pewne zakodowane w naturze zarówno jednostki, jak i całego gatunku przyporządkowanie do zachowania swego istnienia jako takiego⁷. Inklinacja ta staje się niejako wewnętrznym „imperatywem”, który popycha człowieka, by ten starał się swoje życie ocalić. W tym sensie owa inklinacja może stanowić główny motyw różnorodnych zachowań zmierzających do przedłużania naszego życia, w tym także do działań usiłujących odkryć eliksir nieśmiertelności.

Póki co, istnieje wiele przeszkód, które uniemożliwiają pokonanie ziemskiej śmiertelności. Jednakże nie brakuje śmiałków (szczególnie wśród transhumanistów), którzy twierdzą, że niebawem śmiertelność zostanie przewyżczona i będziemy w stanie żyć tak długo, jak tylko będziemy chcieli⁸. Czy aby entuzjazm dotyczący zrealizowania projektu biologicznej nieśmiertelności nie jest przedwczesny?

Biologiczne teorie starzenia się

Dane demograficzne sugerują, że średnia długość ludzkiego życia zwiększyła się na przestrzeni wieków niemal trzykrotnie⁹. Z dużą dozą pewności możemy powiedzieć, że zostało to spowodowane wieloma czynnikami, które przede wszystkim opóźniły proces starzenia się poszczególnych organizmów, a w szerszym zakresie, całej ludzkiej populacji. Nasilające się widmo zgonów zostało odsunięte w czasie i dotyka obecnie głównie osób starszych¹⁰. Postęp w medycynie przyczynił się do zwalczania procesów patologicznych (chorobowych), pozwalając cieszyć się dłuższym życiem. Pomimo wyłożonych wysiłków nauk biomedycznych,

6 W. Tatarkiewicz, *O szczęściu*, Warszawa: Wydawnictwo PWN 1985, s. 505.

7 Zob. M. A. Krapiec, *Inklinacja*, [w:] (red.) A. Maryniarczyk, *Powszechna Encyklopedia Filozofii*, t. 4, Lublin: Polskie Towarzystwo Tomasza z Akwinu 2003, s. 852.

8 “Our mortality will be in our own hands. We will be able to live as long as we want (...)”, R. Kurzweil, *The Singularity is near: When humans transcend biology*, London: Penguin 2005, s. 9.

9 J. R. Wilmoth, *The future of human longevity: A demographer's perspective*, “Science”, 280 (1998), s. 395-397.

10 G. Barazzetti, *Looking for the fountain of youth. Scientific, ethical, and social issues in the extension of human lifespan*, [w:] (red.) J. Savulescu, R. ter Meulen, G. Kahane, *Enhancing human capacities*, Oxford: Wiley-Blackwell 2011, s. 335.

procesu starzenia nie udało się zatrzymać. Dlaczego? Żeby odpowiedzieć na to fundamentalne pytanie filozoficzne, trzeba spojrzeć najpierw na tzw. teorie starzenia się¹¹.

W literaturze przedmiotu funkcjonuje wiele konkurencyjnych teorii starzenia się. Choć mają one potwierdzenie empiryczne, to jednak żadna z nich nie wyjaśnia w sposób całościowy tego zjawiska¹². Zasadniczo teorie usiłujące wyjaśnić mechanizmy starzenia się można podzielić na dwie grupy¹³: (1) opierające się na założeniu, że istnieje z góry zaprogramowany schemat procesu starzenia się, który jest stymulowany przez istnienie zegara biologicznego (teorie te wywodzą się z biologii ewolucyjnej), (2) opierające się na założeniu, że proces starzenia się przebiega przypadkowo w wyniku nagromadzenia mutacji, zniszczeń i błędów powstających w komórkach (teorie te wywodzą się z biologii molekularnej).

Zgodnie z ewolucyjną teorią starzenia się autorstwa Thomasa Kirkwooda, sam proces starzenia nie jest genetycznie zaprogramowany, ale stanowi wynik kompromisu między zasobami zużywanymi na reprodukcję a utrzymanie organizmu. Według Kirkwooda posiadamy "ciało jednorazowego użytku"¹⁴. Proces dokonującej się selekcji ukierunkowany jest na jak najefektywniejsze przekazanie genów potomstwu. Poszczególne organizmy najlepiej przystosowane są do radzenia sobie z czynnikami negatywnymi (wewnętrznymi i zewnętrznymi) do czasu rozmnażania się. Kiedy jednak organizm wkroczy w wiek poreprodukcyjny, wówczas układ immunologiczny staje się słabszy, a organizm bardziej podatny na choroby i zniszczenie. Jakie są tego powody? Przyczyn, według Kirkwooda, należy szukać w¹⁵: (1) genach, które traktują ciało jako coś zbytecznego inwestując w naprawę tylko tyle, ile to konieczne, (2) ograniczeniach lansujących interesy młodego pokolenia kosztem dalszej jego trwałości, (3) doborze naturalnym, który nie "przejmuje się" specjalnie mutacjami gromadzącymi się w sposób niekontrolowany w późnym wieku. Nie jest to jedyna teoria, która wyklucza osiągnięcie biologicznej nieśmiertelności.

Tezę o niemożliwości życia wiecznego na ziemi wzmacnia teoria ograniczonej liczby podziałów komórki autorstwa Leonarda Hayflicka. Podczas prowadzonych badań uczony ten zauważył, że komórki somatyczne charakteryzują się ograniczoną liczbą potencjalnych podziałów. Ilość tych podziałów ulega zmniejszeniu w miarę upływu lat. Badane przez Hayflicka komórki zazwyczaj przestawały się dzielić po około 20 podziałach, żadne jednak nie przeżyły 50 podziałów (liczba ta to tzw. "limit Hayflicka")¹⁶. Kres długości życia komórek nie zależy zatem od czasu, ale od ilości podziałów, którym ulegają poszczególne komórki organizmu. One niejako "pamiętają" liczbę przebytych podziałów, których ilość zawarta jest w chromosomach. Zachodzi zatem zależność - zdaniem Hayflicka - między starzeniem się komórek a całym organizmem. Poszczególne organy ludzkiego organizmu umierają, ponieważ nie mogą się odnawiać. Skoro liczba podziałów komórki jest skończona dla poszczególnych gatunków, to każdy organizm posiada genetycznie

11 W artykule przywołuję tylko niektóre teorie starzenia się, aby pokazać, że najwyraźniej przeczą one możliwości uzyskania biologicznej nieśmiertelności.

12 S. Steuden, *Psychologia starzenia się i starości*, Warszawa: Wydawnictwo PWN 2012, s. 32.

13 L. Hayflick, *Jak i dlaczego się starzejemy*, tłum. M. Symonides, Warszawa: Wydawnictwo "Książka i Wiedza" 1998, s. 208; F. Fukuyama, *Koniec człowieka. Konsekwencje rewolucji biotechnologicznej*, tłum. B. Pietrzyk, Kraków: Wydawnictwo "Znak" 2008, s. 85.

14 T. Kirkwood, *Czas naszego życia. Co wiemy o starzeniu się człowieka*, tłum. M. Kowaleczko-Szumowska, Kielce: Wydawnictwo "Charaktery" 2005, s. 92; E. Sikora, *Opóźnić starość*, "Akademia", 22 (2010) 2, s. 20.

15 T. Kirkwood, *Czas naszego życia*, dz. cyt., s. 100.

16 I. Ziemiński, *Metafizyka śmierci*, Kraków: Wydawnictwo WAM 2010, s. 120.

określoną maksymalną długość życia¹⁷. W przypadku gatunku ludzkiego wynosi ona około 125 lat¹⁸.

Kolejną ważną teorią w badaniach nad procesami starzenia stanowi teoria programowanej śmierci komórki, tzw. apoptoza. Gdy poziom uszkodzeń komórki przekroczy pewien próg, wtedy uruchamia się proces, który powoduje samozniszczenie komórki¹⁹. Apoptoza okazuje się niezbędna, aby poszczególne organizmy mogły dalej żyć i rozwijać się. Stanowi ona narzędzie przetrwania poszczególnych organizmów²⁰. A skoro tak, to na podstawie przywołanych już przesłanek empirycznych można dojść do wniosku, że sens apoptozy jest taki, iż wyklucza możliwość uzyskania biologicznej nieśmiertelności²¹.

Teorii starzenia się, które zarazem zaprzeczają możliwość osiągnięcia biologicznej nieśmiertelności, jak i potwierdzają biologiczną skończoność ludzkiego bytu, jest znacznie więcej²². Niemniej jednak ważne w tym miejscu staje się dla nas następujące spostrzeżenie: każda biologiczna koncepcja starzenia się, w sposób pośredni lub bezpośredni, wyklucza nieskończone biologiczne trwanie na ziemi.

Współczesne metody przedłużania życia - komórki macierzyste

Rozwój medycyny przyniósł w ostatnim czasie niesłychane postępy w diagnostyce i w leczeniu chorób (np. intensywna terapia, transplantacje, zastępowanie niewydolnych części narządów). Dzięki temu przyczynił się do znacznego wydłużenia ludzkiego życia. To nie wszystko. Również rozwój nanotechnologii wspartej technikami inżynierii genetycznej pozwolił wnikać w nanostruktury naszego organizmu docierając do miejsc, do których dotarcie nie byłoby możliwe przy zastosowaniu metod sprzed "rewolucji nanotechnologicznej". Niemniej jednak spośród wszystkich współczesnych technik zmierzających do przedłużenia naszej biologicznej egzystencji, najbardziej obiecującą metodą jest wykorzystanie komórek macierzystych w medycynie translacyjnej. Czymże są te komórki?

Zgodnie z definicją komórki macierzyste mają zdolność do samoodnawiania oraz różnicowania się w komórki potomne²³. To one dają początek komórkom wątroby, serca, mięśni, kości, mózgu. Znajdują się w różnych miejscach naszego organizmu (tkanki, narządy). Najwięcej ich jest w szpiku kostnym. W nim żyją komórki macierzyste dając początek pozostałym komórkom krwi²⁴. Jednakże za najbardziej "użyteczne" uchodzą tzw. pluripotencjalne komórki macierzyste, czyli takie komórki, które odznaczają się dużą plastycznością, mogą przekształcić

17 T. Kirkwood, *Czas naszego życia*, dz. cyt., s. 108-110.

18 L. Hayflick, *The future of aging*, "Nature", 408 (2000) 9, s. 268; B. Chyrowicz, *Życie: długość, jakość i moralność*, dz. cyt., s. 8.

19 A. Paszewski, *Co biologia mówi o starzeniu się?*, [w:] (red.) B. Chyrowicz, *Przedłużanie życia jako problem moralny*, dz. cyt., s. 26.

20 W. Glannon, *Genes and future people. Philosophical issues in human genetics*, Colorado-Oxford: Westview Press 2001, s. 137. T. Kirkwood, *Czas naszego życia*, dz. cyt., s. 119.

21 I. Ziemiński, *Metafizyka śmierci*, dz. cyt., s. 118.

22 Zob. S. Steuden, *Psychologia starzenia się i starości*, dz. cyt., s. 33-34; Caleb E. Finch, Rudolph E. Tanzi, *Genetics of aging*, "Science", 278 (1997) 10, s. 407-411; B. Zegarska, M. Woźniak, *Przyczyny wewnątrzpochoдного starzenia się skóry*, "Gerontologia Polska", 14 (2006) 4, s. 153-159.

23 M. Z. Ratajczak, M. Kucia, *Komórki macierzyste - wyzwanie XXI wieku?*, "Postępy Biologii Komórki", 32 (2005), Suplement 23, s. 11.

24 A. Szczeklik, *Medycyna regeneracyjna a mit wiecznej młodości*, "Nauka", 2 (2004), s. 8-9.

się w każdy rodzaj komórek tkanek i narządów. Można je uzyskać na drodze tzw. klonowania "terapeutycznego" (z komórki jajowej usuwa się jądro i w jego miejsce wprowadza się inne jądro, z komórki dorosłego organizmu)²⁵, bądź też z wczesnych zarodków (najczęściej z tzw. zarodków "nadliczbowych")²⁶. Jednakże działania związane z pozyskiwaniem komórek macierzystych na drodze klonowania lub z embrionów ludzkich wzbudziły poważne opory etyczne i stały się przedmiotem regulacji w różnych krajach. Najbardziej liberalne prawo wobec wykorzystania tych komórek obowiązuje w Wielkiej Brytanii²⁷.

Ostatnie doniesienia naukowe pokazują, że istnieje alternatywa w pozyskiwaniu komórek macierzystych. Somatyczne komórki macierzyste można reprogramować bez pomocy oocyty²⁸. Tego odkrycia dokonał Japończyk Shinya Yamanaki. Pobrał on dojrzałe komórki (np. z twarzy dorosłej osoby) i wprowadził do nich cztery geny, tzw. czynniki transkrypcyjne (wirusy). Geny te przekształciły owe komórki w ludzkie komórki macierzyste (*induced Pluripotent Stem Cells* - iPSC). Dorosła komórka cofnęła się do najwcześniejszego okresu swojego istnienia. Wyniki te potwierdzono w wielu ośrodkach badawczych na całym świecie. Technikę udoskonalono tak dalece, że obecnie wystarczy tylko jeden czynnik transkrypcyjny, by móc reprogramować komórkę. Prawdopodobnie w niedalekiej przyszłości nie trzeba będzie sięgać do komórek macierzystych pozyskiwanych z ludzkich komórek rozrodczych lub zarodków, unikając tym samym aporii moralnych²⁹.

Współczesna medycyna wiąże ogromne nadzieje z komórkami macierzystymi, które jak się wydaje, stanowią klucz do przedłużania życia i osiągnięcia długowieczności. Nie ulega też wątpliwości, iż "(...) w przyrodzie istnieje proste równanie mówiące, że: komórki macierzyste = regeneracja = długowieczność. Istnieje również jego odwrotność, mówiąca, iż starzenie = upośledzenie regeneracji = zmniejszenie liczby/funkcji puli komórek macierzystych"³⁰. Należy jednak zauważyć, że komórki macierzyste, czasem nazywane "nieśmiertelnymi" (podobnie jak i komórki rakowe), wcale nie są nieśmiertelne. Także i one umierają wraz z obumarciem organizmu. Jeśli nawet komórki macierzyste staną się arsenałem części zamiennych³¹, czy będą w stanie sprawić, by człowiek już nigdy nie umierał?

Biologiczna nieśmiertelność?

Możemy sobie wyobrazić, co by się stało, gdybyśmy osiągnęli biologiczną nieśmiertelność. Pewnie część problemów, z którymi się borykamy, musiałoby odejść do lamusa. Zniknęłyby szpitale i ośrodki zdrowia, nie musielibyśmy czekać wielu miesięcy, czy też lat, w kolejce do specjalisty, nie trzeba by płacić składek zdrowotnych ani ubezpieczać się od nieszczęśliwych wypadków, w końcu także

25 Ch. T. Scott, *Czas komórek macierzystych. Krótki wstęp do nadchodzącej medycznej rewolucji*, tłum. M. Betley, Gliwice: Centrum Kształcenia Akademickiego 2008, s. 45.

26 A. Szczeklik, *Medycyna regeneracyjna a mit wiecznej młodości*, art. cyt., s. 14.

27 Tamże, s. 15.

28 J. Yu, *Induced Pluripotent Stem Cell Lines Derived from Human Somatic Cells*, w: <http://www.gs.washington.edu/news/Yu.pdf>, (dostęp: 10.04.2012).

29 A. Szczeklik, *Nieśmiertelność. Prometejski sen medycyny*, Kraków: Wydawnictwo "Znak" 2012, s. 127-128.

30 M. Z. Ratajczak, M. Kucia, *Komórki macierzyste - wyzwanie XXI wieku?*, art. cyt., s. 13.

31 B. Chyrowicz, *Klonowanie*, [w:] *Powszechna Encyklopedia Filozofii*, t. 5, (red.) A. Maryniarczyk, Lublin: Polskie Towarzystwo Tomasza z Akwinu 2004, s. 658.

toczące się spory wokół NFZ zostałyby zażegnane, itd. Słusznie możemy założyć, że nie zabrakłoby także drugiego bieguna - na pozór optymistycznej wizji biologicznej nieśmiertelności - nowych problemów. Jednym z nich mógłby być problem przeludnienia (przy założeniu, że w dalszym ciągu istniałaby prokreacja), innym większe zanieczyszczenie środowiska, albo - jak starał się wskazać Williams - nieznośna nuda. Czy jednak taki stan rzeczy - biologiczna nieśmiertelność - możliwy jest do osiągnięcia? Chociaż nie jest absurdem pytać, czy człowiek może wieść życie wieczne w sensie biologicznym, to jednak "(...) wiara w urzeczywistnienie takiej utopii nie wydaje się zasadna"³².

Wykorzystanie komórek macierzystych, terapia genowa, nanomedycyna, inżynieria genetyczna, klonowanie, kryptonika i inne współczesne metody zmierzające do przedłużenia ludzkiego życia nie dają pewności, że kiedykolwiek stan biologicznej nieśmiertelności zostanie osiągnięty. Faktem jest, że działania te istotnie wpływają na wydłużenie ludzkiego życia, jednakże nie zabezpieczają przed jego całkowitą utratą. Jak dotąd nie udało się nikomu przekroczyć udokumentowanego rekordu długości życia na ziemi, który został ustanowiony przez Francuzkę Jeanne Calment żyjącą 122 lata. Większość z żyjących dziś ludzi prawdopodobnie nawet nie zbliży się do tego wieku, bowiem szanse na osiągnięcie matuzalemych lat są niewielkie³³.

Nie możemy pominąć w tym miejscu przywołanych już przez nas biologicznych koncepcji starzenia się. One również nie napawają optymizmem, jeśli idzie o urzeczywistnienie projektu biologicznej nieśmiertelności. Ponadto doświadczenie śmierci "widzianej" zarówno z zewnątrz (człowiek jako byt rozpadający się), jak i fakt śmierci "widziany" od wewnątrz³⁴, nie dają możliwości usprawiedliwienia sądów jednoznacznie uznających konieczność wiecznego trwania na ziemi³⁵. Nie są to jedyne argumenty na rzecz niemożności osiągnięcia biologicznej nieśmiertelności. Także filozofia klasyczna w wydaniu św. Tomasza z Akwinu podpowiada, że rozkład ciała jest nieunikniony, bowiem w "(...) *dwojaki sposób coś ulega zniszczeniu: samo przez się albo z przyczyny przypadłościowej, tzn. dzięki temu, że coś innego powstaje lub niszczyje (...)*"³⁶. Rozkład ostateczny, czyli śmierć biologiczna uzależniona jest w organizacji materii od przyrody, a człowiek będący jej wytworem nie może nie podlegać prawom tej przyrody³⁷. Ostateczną konsekwencją podlegania prawom przyrody jest śmierć biologiczna (rozkład ciała).

Amerykański gerontolog Leonard Hayflick twierdzi, że nawet jeśli uda się wyeliminować wszystkie choroby nękające ludzkość, to śmierć i tak będzie nam ciągle towarzyszyła wskutek wypadków i przyczyn naturalnych³⁸. W tym sensie "(...) *człowiek jest skazany na śmierć i musi tę śmierć « przeżyć », tzn. jej doznać*"³⁹. Słuszne w analizowanym przez nas kontekście wydaje się określenie: "*Być nieśmiertelnym - to być niezniszczalnym*"⁴⁰. Niemniej jednak ciało ludzkie jest zniszczalne i w związku z tym nie może być nieśmiertelne.

Musimy jednak przyznać, że przywołane przez nas argumenty nie przesądzają całkowicie sprawy o bezzasadności biologicznej nieśmiertelności, gdyż, jak dowodzi

32 I. Ziemiński, *Metafizyka śmierci*, dz. cyt., s. 152-153.

33 B. Chyrowicz, *Życie: długość, jakość i moralność*, dz. cyt., s. 8.

34 M. A. Krąpiec, *Ja-człowiek*, Lublin: Wydawnictwo KUL 2005, s. 427-436.

35 Wysiwanie takich sądów jest charakterystyczne dla współczesnych transhumanistów.

36 Św. Tomasz z Akwinu, *Summa teologii*, I, q. 75, a. 6, *Traktat o człowieku*, tłum. S. Swieżawski, Kęty: "Antyk" 2000, s. 48.

37 M. A. Krąpiec, *Ja-człowiek*, dz. cyt., s. 446.

38 L. Hayflick, *Jak i dlaczego się starzejemy?*, dz. cyt., s. 313.

39 M. A. Krąpiec, *Ja-człowiek*, dz. cyt., s. 446.

40 E. Gilson, *Tomizm. Wprowadzenie do filozofii św. Tomasza z Akwinu*, Warszawa: Instytut Wydawniczy Pax 1998, s. 221.

Ireneusz Ziemiński, "(...) żadna biologiczna konieczność śmierci nie istnieje"⁴¹. Jeśli jednak zgodzimy się na słowa *Prefacji o zmarłych*: "(...) życie Twoich wiernych, o Panie, zmienia się, ale się nie kończy, i gdy rozpadnie się dom doczesnej pielgrzymki, znajdą przygotowane w niebie wieczne mieszkanie"⁴², nie będziemy musieli dłużej zaprzętać sobie głowy deliberacjami na temat możliwości osiągnięcia "wiecznego teraz" na ziemi, bowiem "nasza ojczyzna jest w niebie"⁴³. Refleksję nad biologiczną nieśmiertelnością pozostawmy raczej naturalistom. Wydają się, że to ich główne egzystencjalne zajęcie. Kto wie, być może pragnienie osiągnięcia jej, nie jednemu transhumanście spędza sen z powiek.

Bibliografia:

1. Barazzetti G., *Looking for the Fountain of Youth. Scientific, Ethical, and Social Issues in the Extension of Human Lifespan*, [w:] (red.) J. Savulescu, R. ter Meulen, G. Kahane, *Enhancing Human Capacities*, Oxford: Wiley-Blackwell 2011. s. 335-349.
2. *Biblia Jerozolimska*, wyd. pierwsze, Poznań: Wydawnictwo Pallottinum 2006.
3. Chyrowicz B., *Klonowanie*, [w:] *Powszechna Encyklopedia Filozofii*, t. 5, (red.) A. Maryniarczyk, Lublin: Polskie Towarzystwo Tomasza z Akwinu 2004, s. 655-658.
4. Chyrowicz B., *Życie: długość, jakość i moralność, Wprowadzenie*, [w:] Tenże, *Przedłużanie życia jako problem moralny*, Lublin: Wydawnictwo TN KUL 2008.
5. Finch Caleb E., Tanzi Rudolph E., *Genetics of Aging*, "Science", 278 (1997) 10, s. 407-411.
6. Fukuyama F., *Koniec człowieka. Konsekwencje rewolucji biotechnologicznej*, tłum. B. Pietrzyk, Kraków: Wydawnictwo "Znak" 2008.
7. Gilson E., *Tomizm. Wprowadzenie do filozofii św. Tomasza z Akwinu*, Warszawa: Instytut Wydawniczy Pax 1998.
8. Glannon W., *Genes and Future People. Philosophical Issues in Human Genetics*, Colorado-Oxford: Westview Press 2001.
9. Hayflick L., *Jak i dlaczego się starzejemy*, tłum. M. Symonides, Warszawa: Wydawnictwo "Książka i Wiedza" 1998.
10. Hayflick L., *The Future of Aging*, "Nature", 408 (2000) 9, s. 267-269.
11. Kass L., *Life, Liberty and the Defense of Dignity. The Challenge for Bioethics*, New York - London: Encounter Books 2004.
12. Kirkwood T., *Czas naszego życia. Co wiemy o starzeniu się człowieka*, tłum. M.

⁴¹ I. Ziemiński, *Metafizyka śmierci*, dz. cyt., s. 155.

⁴² *Prefacja o zmarłych*, nr 86, [w:] *Mszał Rzymski dla Diecezji Polskich*, Poznań: Wydawnictwo Pallottinum 1986.

⁴³ Flp 3,20; Sigla biblijne podaje za: *Biblia Jerozolimska*, wyd. pierwsze, Poznań: Wydawnictwo Pallottinum 2006.

-
- Kowaleczko-Szumowska, Kielce: Wydawnictwo "Charaktery" 2005.
13. Krąpiec M. A., *Inklinacja*, [w:] (red.) A. Maryniarczyk, *Powszechna Encyklopedia Filozofii*, t. 4, Lublin: Polskie Towarzystwo Tomasza z Akwinu 2003.
14. Krąpiec M. A., *Ja-człowiek*, Lublin: Wydawnictwo KUL 2005.
15. Kurzweil R., *The Singularity is Near: When Humans Transcend Biology*, London: Penguin 2005.
16. Paszewski A., *Co biologia mówi o starzeniu się?*, [w:] (red.) B. Chyrowicz, *Przedłużanie życia jako problem moralny*, Lublin: Wydawnictwo TN KUL 2008.
17. *Prefacja o zmarłych, nr 86*, [w:] *Mszał Rzymski dla Diecezji Polskich*, Poznań: Wydawnictwo Pallotinum 1986.
18. Ratajczak M. Z., Kucia M., *Komórki macierzyste - wyzwanie XXI wieku?*, "Postępy Biologii Komórki", 32 (2005), Suplement 23, s. 11-26.
19. Scott Ch. T., *Czas komórek macierzystych. Krótki wstęp do nadchodzącej medycznej rewolucji*, tłum. M. Betley, Gliwice: Centrum Kształcenia Akademickiego 2008.
20. Sikora E., *Opóźnić starość*, "Akademia", 22 (2010) 2, s. 20-23.
21. Steuden S., *Psychologia starzenia się i starości*, Warszawa: Wydawnictwo PWN 2012.
22. Szczeklik A., *Medycyna regeneracyjna a mit wiecznej młodości*, "Nauka", 2 (2004), s. 7-16.
23. Szczeklik A., *Nieśmiertelność. Prometejski sen medycyny*, Kraków: Wydawnictwo "Znak" 2012.
24. Św. Tomasz z Akwinu, *Summa teologii*, I, q. 75, a. 6, *Traktat o człowieku*, tłum. S. Swieżawski, Kęty: "Antyk" 2000.
25. Tatarkiewicz W., *O szczęściu*, Warszawa: Wydawnictwo PWN 1985.
26. Williams B., *Sprawa Makropulos: refleksje nad nudą nieśmiertelności*, tłum. T. Duliński, [w:] Tenże, *Ile wolności powinna mieć wola?*, Warszawa: Fundacja Aletheia 1999.
27. Wilmoth J. R., *The Future of Human Longevity: A Demographer's Perspective*, "Science", 280 (1998), s. 395-397.
28. Yu J., *Induced Pluripotent Stem Cell Lines Derived from Human Somatic Cells*, [w:] <http://www.gs.washington.edu/news/Yu.pdf>, (dostęp: 10.04.2012).
29. Zegarska B., Woźniak M., *Przyczyny wewnątrzpochodnego starzenia się skóry*, "Gerontologia Polska", 14 (2006) 4, s. 153-159.
30. Ziemiński I., *Metafizyka śmierci*, Kraków: Wydawnictwo WAM 2010.

Błażej Gębura
Katolicki Uniwersytet Lubelski Jana Pawła II

Stanowisko Hume'a w kwestii istnienia Boga.

*(...) nie wiem, jakie złe następstwa mogą w
dzisiejszych czasach wyniknąć z opinii
ateusza, zwłaszcza gdy pod innymi względami
nic człowiekowi nie można zarzucić.*

Hume w liście do Jamesa Oswalda

Uwaga wstępna

Motywy skłaniającym do napisania tego tekstu jest pewna kontrowersja dotycząca interpretacji treści i charakteru filozofii Hume'a. Dotyczy ona tego, czy autor *Dialogów o religii naturalnej* był agnostykiem, ateistą, fideistą, deistą, a może – jak sądzili niektórzy badacze – teistą?¹ Chodzi więc – inaczej mówiąc – o odpowiedź na pytanie: jakie było stanowisko Hume'a w kwestii istnienia Boga, a nie ustalenie czy Hume był człowiekiem religijnym, gdyż to jest zupełnie poza dyskusją (patrz. motto). Problem nie wydaje się błahy, choćby ze względu na to, że pojawiły się rozbieżne nań odpowiedzi, a przede wszystkim dlatego, że jego rozwiązanie w oczywisty sposób rzutuje na całościową wymowę rozważań szkockiego filozofa nad religią. Rozbieżność udzielanych na to pytanie odpowiedzi można wytłumaczyć tym, że sam Hume nie ułatwił badaczom sprawy poprzez unikanie w swoich oficjalnych pismach jednoznacznej deklaracji, za którym stanowiskiem się opowiadał².

Wskutek tego, na pierwszy rzut oka wydaje się, że żadna z pięciu interpretacji nie pasuje do stanowiska zajmowanego przez Hume'a (sic!). Trudność w sprowadzeniu stanowiska szkockiego myśliciela do tych, które dotyczą kwestii istnienia Boga – objawia się również tym, że poszczególni badacze dodają do danego stanowiska (które miałyby zajmować Hume) różne dodatkowe określenia (np. mitigated, mystical theism czy attenuated deism) sugerujące, że pozycja teoretyczna Hume'a nie do końca odpowiada klasycznym sformułowaniom tychże stanowisk. Sądzę jednak, że problem ten domaga się rozwiązania, a uważna analiza tekstów pozwoli przytoczyć

¹ Zwolennikiem interpretacji ateistycznej jest B.S. Condry, teistycznej T.S. Yoder, rozwiązania agnostycznego N.K. Smith i F.C. Copleston, a deistycznego J.C.A. Gaskin.

² Z listów Hume'a do Adama Smitha wiemy, że niejednoznaczna wymowa *Dialogów...* była zamierzona i podyktowana ostrożnością Autora. Zob. *The Letters of David Hume*, ed. J.Y.T. Greig, Oksford 1932, t. II, s. 332-324; 325-326. Podaje za: A. Hochfeldowa, *Dawida Hume'a <<Dialogi o religii naturalnej>>* w: D. Hume, *Dialogi o religii naturalnej. Naturalna historia religii*, PWN, Warszawa 1962, s. XV.

argumenty przemawiające na korzyść interpretacji ateistycznej (nieposiadającej żadnego dodatkowego określenia). Postaram się to pokazać szkicując pokrótce argumenty na rzecz każdej z pozostałych interpretacji, a potem spróbuję wykazać, że są nietrafne.

Interpretacja fideistyczna i deistyczna

W ujęciu fideistycznym należałoby wyakcentować wypowiedzi Hume'a dotyczące wiary religijnej. Szkielec takiej konstrukcji możemy ująć następująco: *Dialogi o religii naturalnej* negatywnie oceniają argumenty na rzecz istnienia Boga oparte na działaniu rozumu. Dzieło zostawia jednak pewną otwartą drogę do Boga, którą jest osobista wiara religijna. Spójrzmy na mający to ukazywać cytat:

„(...) każdy człowiek czuje prawdziwość religii niejako we własnym swoim sercu i, że nie tyle rozumowanie, co raczej świadomość własnego upośledzenia i niedoli szukać mu każe opieki u tej Istoty, od której zależy on sam i cała przyroda. Najlepsze nawet karty naszego życia pełne są takiego niepokoju albo takiej nudy, że przyszłość stanowi zawsze przedmiot wszystkich naszych nadziei i obaw. (...) Wśród niezliczonych nieszczęść żywota, w czym moglibyśmy szukać ucieczki, gdyby religia nie podsunęła nam myśli o sposobach kompensaty i gdyby nie łagodziła owych trwóg, które bezustannie prześladowają nas i dręczą?”³.

W tym miejscu należałoby stwierdzić, że powyższe stwierdzenie rzeczywiście mogłoby sugerować interpretację fideistyczną. Warto jednak mieć na względzie wskazówkę A. Hochfeldowej: „Religia jest pozbawiona podstawy rozumu, wszelako odwołać się tu należy nie do rozumu, lecz do wiary, oświadcza Filon w zakończeniu *Dialogów*. (...) Wiara i wszelkie związane z nią fenomeny pochodne (...) okazują się (...) echem obaw i nadziei <<nieszczęsnego, biedującego zwierzęcia>>, jakim był człowiek pierwotny, refleksem przesądów po dziś dzień żywionych przez społeczeństwo, zabobonem utrwalonym przez wychowanie (...)”⁴. Ważne jest również to, że według Hochfeldowej *Naturalna historia religii* jest dziełem komplementarnym w stosunku do *Dialogów* – i dopiero łączna ich lektura ukazuje w pełni poglądy Hume'a na religię jako taką. Idąc za tymi wskazówkami można więc powiedzieć, że nawet jeśli *Dialogi* pozostawiają niewielką przestrzeń dla interpretacji fideistycznej, to lektura *Naturalnej historii religii* (a także eseju *O zabobonie i egzaltacji*) skutecznie usuwa tę perspektywę⁵. Wiara religijna jest nie tyle przeciwna rozumowi, co wprost sprzeczna z jego ustaleniami. Nie niesie ona ze sobą również tak szczególnej wartości, jakiej chciałby w niej upatrywać fideista. Z pewnością bowiem zgodziłby się on na tezę, że wiara nie wspiera się na ustaleniach rozumu (a nawet musi tak twierdzić, żeby utrzymać swoje stanowisko), ale upatrywałby w tym raczej jej zaletę, niż wadę. Jest tak dlatego, że jej podłoże jest – według Hume'a – wyjaśnialne w kategoriach czysto naturalistycznych. Wiara jest produktem naszych nawyków, lęków i ignorancji co do

³D. Hume, *Dialogi o religii naturalnej* w: *Dialogi o religii naturalnej. Naturalna historia religii*, PWN, Warszawa 1962, s. 87.

⁴Z. Hochfeldowa, *Dawida Hume'a <<Dialogi o religii naturalnej>>*, w: D. Hume, *Dialogi o religii naturalnej. Naturalna historia religii*, PWN, Warszawa 1962, s. LI.

⁵Ujęcie fideistyczne stanowiska Hume'a odrzuca J.C.A. Gaskin. Zob. J.C.A. Gaskin, *Hume on Religion*, w: *The Cambridge Companion to Hume*, red. D.F. Norton i J. Taylor, Cambridge University Press, New York 2009, s.507.

potwierdzonych przez doświadczenie ustaleń nauki. Z tego względu nie można więc powiedzieć, że wiara religijna jest rzetelnym środkiem do uznania tezy o istnieniu Boga.

Przechodząc do zagadnienia interpretacji deistycznej należy powiedzieć, że jej źródło tkwi (podobnie jak w przypadku ujęcia teistycznego) w różnych wypowiedziach Hume'a, w których jest mowa o „przyczynie panującego we wszechświecie porządku” czy „Stwórcy”⁶. Najważniejsza w tym kontekście jest jednak ostatnia w *Dialogach* wypowiedź Filona⁷. Na podstawie analizy tej wypowiedzi J.C.A. Gaskin skłania się ku interpretacji deistycznej⁸. Nie wchodząc w dyskusję na temat tego fragmentu, można zwrócić uwagę na trzy elementy.

Po pierwsze: deiści odrzucali Opatrzność i Objawienie. Po drugie: przyjmowali, że Bóg stworzył świat i umożliwił mu funkcjonowanie poprzez obdarzenie go pewnymi stałymi prawami, dzięki czemu przypomina on sprawnie działający mechanizm (np. zegarek). Po trzecie: lukę powstałą po zanegowaniu Objawienia chcieli zapełnić teologią naturalną opartą wyłącznie na rozumie.

Przyłóżmy teraz te wyznaczniki do poglądów Hume'a. Jak wiemy, filozof odrzucał zarówno Objawienie (patrz. *O cudach*)⁹ jak i Opatrzność¹⁰ (patrz. *O szczególnej opatrności i życiu przyszłym*). W tym więc zgodziłby się z deistami. Nie zgodziłby się natomiast na tezę o celowym uporządkowaniu świata przez istotę ponad-naturalną. Widać to doskonale w jego krytyce argumentu teleologicznego (nazywanego przezeń argumentem z analogii). Według Hume'a konstruowanie wszelkich analogii między światem a artefaktami jest wysoce wątpliwe i narażone na szereg zarzutów¹¹. Co więcej, Hume odrzuca całą teologię naturalną. Istnienia Boga nie można poznać za pomocą rozumu, a więc postulat deistów okazuje się nie do zrealizowania.

Interpretacja agnostyczna i teistyczna

Poza interpretacją fideistyczną i deistyczną można podać jeszcze dwie inne: agnostyczną i teistyczną. Zwolennikiem pierwszej z nich jest między innymi F.C. Copleston:

„Czytelnik może spodziewać się jasnej odpowiedzi na pytanie, czy Hume ma być uważany za ateistę, agnostyka czy teistę. Nie jest jednak łatwo dać taką <<jasną odpowiedź>> (...) Dopóki jednak nie uda nam się odczytać z Hume'a znacznie więcej, niż chciał nam wyjawiać, trudno tu mówić o teizmie. Można by tu mówić o agnostycyzmie; trzeba jednak pamiętać, że Hume nie był agnostykiem w stosunku do istnienia osobowego boga z cechami moralnymi, takiego, jakim Go szczegółowo opisują Chrześcijanie. Wydaje się, że jest tak, iż Hume jako bezstronny obserwator zaczyna badać, jak można by uwierzytelnić teizm, utrzymując jednocześnie, że religia wspiera się na objawieniu, w które on sam z pewnością nie wierzył. Wynikiem tego badania było (...) dążenie do

6 Zob. D. Hume, *Naturalna historia religii*, w: D. Hume, *Dialogi o religii naturalnej. Naturalna historia religii*, Op. Cit., s. 215-216.

7 D. Hume, *Dialogi o religii naturalnej*, w: D. Hume, *Dialogi o religii naturalnej. Naturalna historia religii*, Op. Cit., s. 134-136.

8 Zob. J.C.A. Gaskin, *Hume's attenuated deism*, *Archiv für Geschichte der Philosophie*, Vol. 65, Nr. 2 (1983), Walter de Gruyter, s. 160-174., a także J.C.A. Gaskin, *Hume on Religion*, w: *The Cambridge Companion to Hume*, Op. Cit. s. 489-490.

9 D. Hume, *Badania dotyczące rozumu ludzkiego*, tł. J. Łukasiewicz i K. Twardowski, PWN, Warszawa 1977, s. 131-161.

10 Tenże, *Badania dotyczące rozumu ludzkiego*, Op. Cit. s. 161-181.

11 Tenże, *Dialogi o religii naturalnej*, Op. Cit. s. 25-27.

zredukowania <<hipotezy religijnej>> do tak nikłej treści, że nie wiemy jak ją nazwać. Jest to – jak dobrze wiedział Hume – ta pozostałość, na którą może zgodzić się każdy, kto nie jest dogmatycznym ateistą”¹².

Trzeba na początku powiedzieć, że taki punkt widzenia stanowiska Hume’a wiąże się bardzo mocno z obiegowym rozumieniem jego filozofii - określanej po prostu jako radykalny sceptycyzm albo pyrronizm. W myśl tej interpretacji, rozważania prowadzone przez niego w kwestii istnienia Boga i zjawiska religii, nawet jeśli na początku mają na celu wypracowanie jednoznacznych odpowiedzi na rodzące się pytania, w toku badania wklajają się w nieprzewidywane trudności. Jedynym wyjściem z zaistniałej sytuacji może być więc tylko sceptyckie epoche, zawieszenie sądu. Można tu przywołać słowa kończące *Naturalną historię religii*, które mocno sugerowałyby właśnie taką interpretację: ”Wszystko to jest zagadką, enigmatem, niedocieczoną tajemnicą. Wątpienie, niepewność, zawieszenie sądu – oto co wydaje się jedynym rezultatem najskrupulatniejszych dociekań w tej materii”¹³.

Wydaje się, że należy zacząć od stwierdzenia, że określanie Hume’a jako sceptyka (w ścisłym sensie tego terminu, czyli kogoś, kto powstrzymuje się od wydawania sądów dotyczących natury rzeczy) nie wydaje się do końca odpowiadać charakterowi jego filozofii, a przynajmniej budzi poważne wątpliwości. Wystarczy tylko zwrócić uwagę na fakt, że Hume przynajmniej w dwóch miejscach analizując pewne problemy metafizyczne zajmuje twarde, jednoznaczne stanowisko. Chodzi mianowicie o kwestię cudów i zagadnienie nieśmiertelności duszy. Zaprzeczenie możliwości zachodzenia zdarzeń sprzecznych z prawami przyrody sprzyja stanowisku naturalistycznemu, podobnie jak stwierdzenie, że dusza nie może być nieśmiertelna. Naturalizm zaś jest stanowiskiem *par excellence* metafizycznym, gdyż mówi o tym jaka jest rzeczywistość.

Te uwagi jednak same w sobie nie przesądzają, że w kwestii istnienia Boga Hume nie mógł być agnostykiem. Należy więc sformułować inny argument. Otóż łatwo zauważyć, że znakomita część *Dialogów* poświęcona jest dyskusji nad dowodem teleologicznym. Hume podaje wiele argumentów mających wykazać jego słabość. Zaś na końcu utworu zdaje się pomimo przedstawionej krytyki przyjmować w końcu ten argument¹⁴. B.S. Condry dowodzi, że tak nie jest¹⁵. Agnostykiem Hume mógłby być tylko wtedy, gdyby przy rozważaniu dowodu teleologicznego (i innych) wykazał, że racje obu stron sporu są równie silne, że zachodzi *izostenia*. Tego jednak Hume nie robi. Wszystkie argumenty teistyczne zostają w toku dyskusji zbite. Z tego powodu rozważając stanowisko Hume’a w tej kwestii, moim zdaniem agnostycyzm należy odrzucić jako ujęcie nie posiadające wystarczających podstaw.

Co do interpretacji teistycznej, należy zauważyć, że swoje źródło bierze ona z wielu fragmentów w pismach Hume’a, które wyrażają, jeśli można się tak wyrazić, „treści teistyczne”. Warto najpierw zauważyć, że w większości tych wypowiedzi słowo „natura” pojawia się o wiele częściej niż słowo „Bóg”. Wypowiedzi te można zatem łatwo uzgodnić z naturalizmem Hume’a, gdyż najczęściej ma on na myśli albo przyrodę, albo jej prawa. B.S. Condry sugeruje, że Hume konstruując takie wypowiedzi chciał jedynie pokazać, że można używać teistycznego języka do wyrażenia naturalistycznych tez. Język teisty jest na tyle „pojemny”, że można dzięki

12 F. Copleston, *Historia Filozofii: Od Hobbesa do Hume’a*, PAX, Warszawa 1997, s. 269.

13 Tenze, *Naturalna historia religii*, w: D. Hume, *Dialogi o religii naturalnej. Naturalna historia religii*, Op. Cit., s. 217.

14 Sprawa końcowej wypowiedzi Filona jest szerzej poruszana przez A. G. Vinka w artykule: *Philo's Final Conclusion in Hume's "Dialogues"*, *Religious Studies*, Vol. 25, No. 4 (Dec., 1989), Cambridge, s. 489-499.

15 B.S. Condry, *A more dangerous enemy? Philo's „confession” and Hume's soft atheism*, *International Journal of Philosophy of Religion*, Springer, 2010, s. 1-23.

niemu przekazywać intuicje zupełnie sprzeczne z teizmem. Według Condry'ego to mógł mieć na myśli Hume określając cały spór między teistami a ateistami jako „czysto werbalny”.

Zamiast cytować te fragmenty w większej ilości i dogłębnie analizować ich wymowę - można podać argument innego rodzaju świadczący przeciwko interpretacji teistycznej. Otóż, jest jasne, że teista uznaje istnienie Opatrzności. Twierdzi on, że Bóg nie tylko stworzył świat, ale także roztacza nad nim opiekę i posiada możliwość ingerowania w jego strukturę i przebieg wydarzeń w nim zachodzących. Wydaje się, że np. cuda można rozumieć jako szczególny przypadek zachodzenia Opatrzności. Hume zaś - jak wiadomo - przeczy istnieniu czegoś takiego jak Opatrzność. Stąd w żadnym razie nie może być teistą. Nie jest to jednak jedyny powód. To, że Hume nie może zajmować stanowiska teistycznego i agnostycznego stanie się jaśniejsze w momencie, gdy sformułuje się argument za interpretacją ateistyczną.

Interpretacja ateistyczna

Odrzuciwszy cztery dotychczasowe interpretacje dochodzimy w końcu do momentu, w którym w naturalny sposób narzuca się ujęcie ateistyczne. Ktoś mógłby wysunąć jednak zarzut, że fakt odrzucenia innych interpretacji nie wskazuje automatycznie na ujęcie ateistyczne jako prawidłowe. To, że wymienione stanowiska nie pasują do tezy Hume'a zamiast oznaczać, że to ateizm jest odpowiednim rozwiązaniem – może pokazywać tylko, że Hume nie miał wyrobionego zdania na temat tego, czy Bóg istnieje. Jednak nie posiadać zdania w kwestii istnienia Boga („istnieje czy nie?”) oznacza ni mniej, ni więcej tylko pozostanie na gruncie agnostycyzmu. A jak pokazano wyżej – tę interpretację z różnych względów należy odrzucić. Stąd, można teraz przejść do sformułowania argumentu za interpretacją ateistyczną.

Jak widzieliśmy, argumentami przeciwko poprzednim interpretacjom były najczęściej te tezy Hume'a, które negowały pewne zjawiska, czy elementy konieczne do utrzymania tychże stanowisk. Patrząc na całość rozważań Hume'a nad religią można sformułować następujące twierdzenie: **Hume odrzuca elementy istotne dla uzasadnienia przekonania o istnieniu Boga**. Argumentem za tą tezą będzie po prostu wyliczenie tych elementów i wykazanie, że po usunięciu ich nie ostaną się żadne podstawy dla religii i tezy o istnieniu Boga. Zatem wykaże się, że wszelkie interpretacje sugerujące jakąś formę pozytywnego stosunku Hume'a do religii (w tym tezy o istnieniu Boga) są bezpodstawne. Można sporządzić listę ustaleń Hume'a dotyczących istnienia Boga i religii jako takiej.

- A – Odrzucenie wszystkich argumentów za istnieniem Boga
- B – Zanegowanie Opatrzności i występowania cudów (które można potraktować jako szczególny przypadek działania Opatrzności)
- C – Ujęcie wiary jako skutku ubocznego niezabezpieczonych potrzeb i trapiących ludzi lęków. Wiara jako szkodliwy społecznie zabobon
- D – Argument za śmiertelnością duszy
- E – Czysto naturalna geneza religii
- F – Krytyka wiarygodności Objawienia i jego świadectw, a także dogmatów religijnych.

G – „Nierzetelna” idea Boga¹⁶

H – Szkodliwy wpływ religii na moralność

Wszystkie te ustalenia zestawione razem według mnie nie pozostawiają wątpliwości, że jedynym stanowiskiem, na jakim w ich świetle może stać Hume - jest ateizm. Wszystkie twierdzenia stanowią wsparcie dla stanowiska naturalistycznego. Nie pozwalają ponadto podjąć jakiegokolwiek projektu teologii naturalnej, w tym sformułować tzw. argumentu moralnego za istnieniem Boga. Uderzają one także w instytucjonalne formy religii (doktrynę, księgi objawione, formy kultu), i przekonują, że przekonania religijne mają czysto naturalistyczne wyjaśnienie. Wszystkie te ustalenia Hume’a rozpatrywane łącznie pokazują, że nie ma żadnych podstaw pozwalających żywić przekonanie, że Bóg istnieje. Nic nie wskazuje, aby mogło tak być. Raczej jest tak, że wszystko przeczy temu, żeby Bóg istniał. Co więcej, wszelkie nadzieje na ukazanie racjonalności wiary religijnej, bądź jej swoistej wartości – w świetle systemu Hume’a - okazują się płonne.

¹⁶W sprawie Humowskiego rozumienia idei Boga zobacz mój artykuł: *Dowód ontologiczny w świetle założeń filozofii Davida Hume’a*, *Philosophia*, Nr. 3 (31), Lublin 2011, s. 21 – 25.

*ks. Bartłomiej Krzos
Katolicki Uniwersytet Lubelski Jana Pawła II*

O rodzajach i genezie norm.

Wprowadzenie

Wszyscy spotykamy się na co dzień z wyrażeniami językowymi dotyczącymi działania ludzkiego. Wyrażenia te mogą być różnego rodzaju. Różna może być też rola wspomnianych tu wyrażen w stosunku do działań człowieka. Mogą one opisywać te działania, charakteryzować je lub oceniać, mogą wreszcie je projektować, kierować nimi czy też wprost je nakazywać. Wśród wielości możliwych wyrażen dotyczących działania, na uwagę zasługują te, które można by nazwać normami lub wyrażeniami normatywnymi. Tematem tego skromnego artykułu jest ich klasyfikacja i geneza. Treść artykułu to porównawcza analiza poglądów twórców najstarszych systemów logik norm na temat podziału i źródeł tych ostatnich. Na uwagę zasługuje także fakt związku genezy norm z filozoficznymi poglądami dotyczącymi działania. Omówienie tego faktu jest również częścią treści tej pracy.

Rozumienie i kwalifikacja działań, czyli czynności ludzkich bywają zróżnicowane. Źródłem tych różnic mogą być poglądy filozoficzne na temat natury i wartości działań. Tak rozumiana kwalifikacja czynności ludzkich, miała znaleźć swoje odzwierciedlenie w tworzonych w połowie XX-go wieku pierwszych systemach logiki norm. Twórcami tychże systemów byli Georg Henryk von Wright i Jerzy Kalinowski. Zagadnienie wyrażalności poglądów filozoficznych dotyczących działań za pomocą wyrażen językowych jest związane z tematem niniejszego artykułu. Kwestie związane z tym tematem mogą być określane jako zagadnienia źródła i pochodzenia norm, czyli – najogólniej mówiąc – ich filozoficznej genezy¹.

Niniejszy artykuł został podzielony na dwie części. Pierwsza z nich będzie poświęcona klasyfikacji norm prezentowanej w różnych pracach Kalinowskiego i von Wrighta. Przywołanie i przebadanie tych podziałów może być ważne przy próbie nakreślenia obszaru badawczego zagadnienia norm. Druga część z kolei poświęcona będzie poglądom filozoficznym wspomnianych tu uczonych na temat pochodzenia norm i ich genetycznych powiązań. Refleksja ta pozwoli uwypuklić pewne istotne różnice filozoficzne obecne w analizowanych poglądach, zwłaszcza te dotyczące zapatrywania na genetyczne związki sfery normatywnej ze sferą aksjologiczną.

¹ Niektóre z przedstawionych w niniejszym artykule rozważań ukazały się już w formie samodzielnego artykułu *O prawdziwości norm postępowania* w „Studiach Paradygmatycznych” 22(2012), red. G. Chojnacki, J. Stoś.

Klasyfikacja norm

Jednym z ważniejszych zagadnień pojawiających się przy okazji badania tematyki norm i niezbędnych do przebadania zagadnienia rodzajów i genezy zdań normatywnych jest klasyfikacja norm występujących w języku potocznym. Zarówno von Wright jak i Kalinowski poświęcali sporo uwagi podziałowi i znaczeniu norm. Najpierw przyjrzymy się klasyfikacji norm zaproponowanej przez Kalinowskiego, a następnie zestawimy i porównamy z nią klasyfikację norm podaną przez von Wrighta.

Powszechnie znany jest, zaproponowany przez Arystotelesa, podział nauk na teoretyczne, praktyczne i poetyczne. Analogicznie do tego podziału można by przeprowadzić podział zdań na zdania teoretyczne, praktyczne i poetyczne. Dwie ostatnie grupy zdań można by wówczas połączyć w jedną, określaną jako zdania praktyczne *sensu largo*. W tym wypadku wyróżnione zdania praktyczne należałoby nazwać zdaniami praktycznymi *sensu stricto*². Da się zauważyć, że ten podział zdań nie spełnia jednak kryteriów podziału logicznego. Aby doprecyzować znaczenie zdania praktycznego i dokonać należytej klasyfikacji tego typu zdań, Kalinowski proponuje skonstruowanie pewnej definicji, która określałaby zdanie praktyczne jako zdanie, którego znaczeniem jest sąd kierujący działaniem ludzkim na pewien określony sposób³.

Do takich zdań kierujących działaniami zaliczają się z oczywistych względów wszystkie zdania rozkazujące. Jak już było wspomniane, nie należą one do kategorii zdań w sensie logicznym. Z tego względu pomijamy tu analizę logiczną zdań rozkazujących, nawet jeśli pośrednio wyrażają one jakoś sądy normatywne⁴.

Drugą – w kontekście naszych rozważań – najważniejszą grupę zdań praktycznych tworzą normy zwane także wyrażeniami normatywnymi. Norma może być tu rozumiana trojako. Po pierwsze można tak nazywać po prostu zdanie normatywne, po drugie sąd, który jest znaczeniem tego zdania, zaś po trzecie istniejący realnie stan rzeczy, do którego to zdanie się odnosi. W ostatnim przypadku nie mówimy o desygnacie zdania lecz o jego desygnacji. Należy dodać, że norma w pierwszym znaczeniu jest przez Kalinowskiego za Amselekiem definiowana kontekstualnie oraz syntaktycznie. Definicja kontekstualna polega na ukazaniu wnioskowania praktycznego, którego przesłanką lub wnioskiem jest norma, np.:

przesłanka: Sąd powinien wydać wyrok skazujący,

a więc

wniosek: Sąd ma prawo wyrok skazujący.

Definicja syntaktyczna polega na podaniu listy terminów, za pomocą których budowane są normy. Są to terminy takie jak: „powinien”, „ma prawo”, „może”, itp., wraz z wszystkimi ich formami fleksyjnymi oraz ich synonimami⁵. W odniesieniu do norm w znaczeniu sądów logicznych Kalinowski stwierdza, że każda norma (w znaczeniu sądu) jest normą dlatego, że jest znaczeniem zdania normatywnego⁶.

Zdania praktyczne należące do grupy norm można za Kalinowskim podzielić najpierw ze względu na ich strukturę syntaktyczną na warunkowe i bezwarunkowe (hipotetyczne i kategoryczne). Podział ten pochodzi od Kanta. Okazuje się, że

2 J. Kalinowski, *Teoria poznania praktycznego*, Towarzystwo Naukowe KUL, Lublin 1960, s. 11.

3 J. Kalinowski, *Teoria poznania praktycznego*, Towarzystwo Naukowe KUL, Lublin 1960, s. 53.

4 J. Kalinowski, *Logika Norm*, Instytut Wydawniczy Daimonion, Lublin 1993, s. 21-23.

5 J. Kalinowski, *Logika Norm*, Instytut Wydawniczy Daimonion, Lublin 1993, s. 12-13.

6 J. Kalinowski, *Logika Norm*, Instytut Wydawniczy Daimonion, Lublin 1993, s. 21.

podobne rozróżnienie norm kategorycznych i hipotetycznych (zwanych również technicznymi) poczynił także von Wright. Stanowisko von Wrighta pociągnęło u Kalinowskiego dalszy podział norm kategorycznych i hipotetycznych, a mianowicie wyróżnienie w obu tych grupach norm w sensie zwyczajnym i nadzwyczajnym (filozoficznym). W sensie nadzwyczajnym normy kategoryczne to takie, które z definicji wyznacza bytowość i dobro bytu zakorzenione w nim przez jego twórcę, pozostałe zaś są normami kategorycznymi w sensie zwyczajnym⁷. Normy hipotetyczne w sensie nadzwyczajnym to takie, które są wyznaczane przez pewien cel podmiotu. W sensie zwykłym natomiast, normą hipotetyczną będzie każda z norm zapisana w trybie warunkowym⁸.

Trzecią grupą zdań kierujących działaniem ludzkim (zdań praktycznych *sensu largo*) obok rozkazów i norm są według Kalinowskiego oceny. Są one dla niego istotne, gdyż są potrzebne przy poszukiwaniu podstaw norm i rozkazów oraz przy uzasadnianiu prawdziwości norm. Zdania oceniające mają według Kalinowskiego swoją genezę w rozumieniu działania. Otóż działanie ludzkie ma zawsze pewną „wartość” w odniesieniu do swojego rezultatu. Chodzi o pewną doskonałość, którą w sobie ma zawierać dzieło podmiotu, a którą istotnie zawiera lub której nie zawiera. Im dzieło bliższe jest zamiarowi twórcy, tym jest „lepsze” (cehuje się wyższą „wartością”). Zdanie oceniające to takie, które zawiera przymiotnik estymatywny w stopniu równym lub innym oraz jego synonimy (również stopniowalne)⁹.

Ważny jest pogląd Kalinowskiego, głoszący, że wszystkie zdania praktyczne są konstytuowane relacją między terminami, które je tworzą. Ta relacja może być konieczna lub kontyngentna. Jeśli ma ona swoje fundamenty we własnościach istotnych bytów, to wówczas jest ona relacją konieczną. Jeśli podstawy te leżą we własnościach przypadłości bytów oznaczanych przez terminy, o których mowa, to wówczas jest to relacja kontyngentna. Jeśli jesteśmy świadomi realnego zachodzenia relacji normatywnej między bytami, które są oznaczane przez nazwy wchodzące w skład danego zdania praktycznego, to zdanie praktyczne jawi się w naszym umyśle z mocą prawdy.

Ostatnią z grup zdań praktycznych opisywanych przez Kalinowskiego są metazdania, czyli zdania o normach, ocenach i rozkazach, a są sformułowane w metajęzyku. Wśród nich wyróżniają się te, które są po prostu normami, ocenami czy rozkazami, ale drugiego (wyższego) stopnia (typu: „powiedziano ci: rób to!”), inną grupą są zdania, które mówią o normach, ocenach i poleceniach, ale nie mają charakteru metanormy, metaoceny czy metapolecenia (np. zdania głoszące o obowiązywaniu lub promulgowaniu danej normy). Istnieje także oddzielna grupa zdań wyjaśniających charakter zdania, którego dotyczą (np.: „to jest rozkaz!”) i wreszcie ostatnia grupa zdań rezolucji, obietnic, deklaracji, zaprzysiężenia. Kolejna grupa to rady, a wreszcie ostatnia grupa to metajęzykowe zdania oznajmujące grożące wyroki lub obiecane nagrody. Ostatnia grupa tych zdań, to zdania, które opisują pewne okoliczności, które nasuwają nam oceny, normy lub polecenia¹⁰.

Rodzi się jeszcze pytanie o zaklasyfikowanie norm wyższego rzędu, czyli norm kierujących specyficznym rodzajem działania, jakim jest kierowanie działaniem. Jeśli chodzi o normy wyższego rzędu u von Wrighta, to wydaje się, że norma nie może być sama treścią innej normy, natomiast zdanie o normie może nią być

7J. Kalinowski, *Le problème de la vérité en morale et en droit*, wyd. Emanuel Vitte, Lyon 1967, s. 176.

8J. Kalinowski, *Teoria poznania praktycznego*, Towarzystwo Naukowe KUL, Lublin 1960, s. 55-56.

9J. Kalinowski, *Le problème de la vérité en morale et en droit*, wyd. Emanuel Vitte, Lyon 1967, s. 184.

10J. Kalinowski, *Wstęp do Nauk Prawnych. Skrypt opracowany na podstawie wykładów dra Jerzego Kalinowskiego wygłoszonych dla studentów wydziału prawa KUL w roku 1948/49*, wyd. Koło Prawników KUL, Lublin 1949, s. 10.

jak najbardziej, bo jego treścią jest istnienie normy. Jest to więc możliwe, żeby istnienie normy było nakazane. Taka norma nakazująca istnienie normy będzie nazywana normą idealną. Dzieje się tak dlatego, że samo ustanawianie norm jest działaniem i jako takie również podlega normowaniu. Oczywiście można też działać odwołując normy. Istnienie normy to stan świata, w którym zdanie o danej normie jest prawdziwe. Odwołanie to albo anihilacja danej normy, albo wprowadzenie jej negacji (promulgacja). Podmiotem normy wyższego rzędu jest sprawca, który jest jednocześnie normodawcą normy niższego rzędu.

Teraz pokrótce przedstawione będą poglądy G. H. von Wrighta na temat klasyfikacji norm. Według von Wrighta punktem wyjścia tych rozważań jest pojęcie prawa. Należy wyróżnić dwa rodzaje prawa: deskryptywne (opisujące) i preskryptywne (nakazujące)¹¹. Dokonując przywołanego tu podziału Von Wright odwołuje się do historii filozofii praktycznej. Stwierdza on, że istniały poglądy filozoficzne, w myśl których prawa kosmosu mogą być rzutowane także na życie społeczne. Rozumiane w ten sposób prawo, to zespół zdań, którym bez wątpienia przysługuje wartość logiczna. Z prawami deskryptywnymi mamy do czynienia wówczas, jeśli opisujemy istniejące prawa, na które nie mamy żadnego wpływu. Przede wszystkim chodzi tu o opisy mechanizmów prawa natury, któremu są niejako podporządkowane wszystkie procesy zachodzące we Wszechświecie. Obok prawa natury prawami deskryptywnymi są wszystkie opisy istniejących kodeksów norm zarówno moralnych jak i prawnych. Jest to bardzo szeroka dziedzina praw obejmująca sfery legislacyjne, religijne czy społeczne. Można w tym zbiorze uplasować wszelkiego rodzaju prawa nadawane, które mają swoją genezę np. w społeczeństwie teokratycznym opisanym w Biblii, w kodeksach moralnych lub prawnych, w regulaminach, itp. Jako zdania w sensie logicznym, omawiane tu deskrypcje mają moment treści i moment asercji.

Prawom deskryptywnym przeciwstawia von Wright prawa preskryptywne. Są to takie sformułowania normatywne, które pochodzą bezpośrednio od normodawcy obdarzonego odpowiednim autorytetem. Do tej dziedziny należą wszelkiego rodzaju promulgacje, zarówno te oficjalne, prawne (np. ogłoszenie ustawy) jak i te prywatne (zezwoleń komuś na używanie mojego długopisu). Nie mają one wartości logicznej i nie należą do zdań logicznych, gdyż są sformułowane zawsze w trybie podobnym do rozkazującego. Jako takie, preskrypcje mają w sobie element treści i element bodźca.

Zarówno jedne jak i drugie mogą być nazywane normami, ale w bardzo przednaukowym i szerokim sensie. Odwołując się do postulowanego ograniczenia zainteresowania domniemanej teorii norm należy powiedzieć, że zdania pierwszej grupy nie zawierają w sobie elementu „bodźcowego” skierowanego na działanie, natomiast zdania drugiej grupy „wychodzą” poza zdania logiczne w ogólnie przyjętym sensie. Z tego powodu ani jedne ani drugie nie mogą pretendować do miana normatywnych zdań logicznych.

Von Wright przywołuje jeszcze trzecią grupę praw, a mianowicie prawa logiki. Nazywa je regułami. Są to prawa a priori i nie tyle odnoszą się do tego „jak jest”, ale raczej do tego „jak powinno być”. Mają więc w sobie element „treściowy” i element „bodźcowy”. Jako zdania w sensie logicznym mają one również swój moment asercji i cechują się wartością logiczną. Element bodźcowy odnosi się do poprawnego wnioskowania a więc ogólnie – myślenia. Można w tym poglądzie dostrzec typowe dla pozytywizmu elementy psychologizmu logicznego, a więc także elementy

¹¹ G. H. von Wright, *Norm and Action: A Logical Inquiry*, wyd. Routledge & Kegan Paul Nowy Jork 1963, s. 138.

relatywizmu norm. Funkcja „bodźcowa” schodzi jednak w normie tego typu na dalszy plan¹². Na pierwsze miejsce wysuwa się tu funkcja deskryptywna poprawnego myślenia. Normy logiczne opisują byty logiczne i matematyczne. Problemатyczne w tym ujęciu wydaje się zakładanie realizmu (platonizmu) matematycznego. W tym paradygmacie prawa przyrody byłby konieczne i niezmiennie, natomiast prawa rządzące myślą byłby kontyngentne¹³.

Logik fiński wymienia jeszcze kilka typów wyrażeń, które podpadają jego zdaniem pod miano normy. Są to po pierwsze prawa stanowione przez zwyczaje, w których nie da się jasno określić normodawcy, a wszyscy członkowie społeczności, w której prawa te funkcjonują, są niejako ich autorami, czy też normy (regulacje) techniczne. Nie mogą mieć one nic wspólnego z prawem natury, ponieważ tego ostatniego nie da się złamać. Nie mogą też być utożsamiane ze zwykłym relatywnym opisem („jeżeli dom ma być mieszkalny to musi być ogrzewany”), ani ze zdaniami mówiącymi o niezmiennym przeznaczeniu (*anankastic proposition*), ponieważ to ostatnie (zresztą wcześniejsze również) zakładają pewne uprzednie presupozycje, które są opisowe *par excellence*.

Ostatnią grupą norm, jaką zajmuje się autor systemów logiki deontycznej są normy moralne (von Wright używa na ich określenie słów *so-called*). Nie twierdzi jednak jakoby te, ujęte w całości, zachowywały się jak reguły gry, niemniej jednak brane pod uwagę pojedynczo mogą przejawiać niekiedy takie cechy. Niektóre normy moralne należą do zwyczajowych (taki jest też ich źródłosłów łaciński: *mores*). Na pierwszy rzut oka normy moralne wyglądają jak nakazy, tyle, że nie da się wyróżnić człowieka (*individual/physical person*), który mógłby uchodzić za ich dawcę. Wyjątek stanowi prawnie ustanowiona władza, rodzice i wychowawcy w stosunku do ich dzieci i wychowanków oraz Pozytywne Prawo Boże (czy też ogólnie rozumując religijne). Cały czas jednak von Wright waha się między przyznaniem normom moralnym statusu nakazów albo reguł technicznych. Najlepszą odpowiedzią wydaje się być przyznanie im statusu normy *sui generis*. Jednak i to nie jest satysfakcjonujące, ponieważ powszechnie narzuca się nam przy okazji norm moralnych odniesienie do jakiejś pozamoralnej rzeczywistości (choćby szczęścia, dobra itp.). W poglądach von Wrighta mamy więc do czynienia z jawnym bądź ukrytym utylityzmem¹⁴.

W ogólnym zarysie wyróżniamy trzy podstawowe rodzaje norm: preskrypcje, deskrypcje i reguły. Pewne „mniejsze” grupy to zwyczaje, zasady moralne i reguły idealne. Jak widać z powyższej analizy normy jako sądy (stwierdzenia) poprzedzające (kierujące, projektujące) działanie ludzkie mogą być traktowane w sensie szerokim albo ścisłym. Na tych w sensie szerokim dokonuje się operacji podziałów logicznych i typologii. Normy w sensie węższym, to jedna z grup powstała w wyniku tego podziału logicznego.

Filozoficzne poglądy na temat genezy norm

Na szczególną uwagę przy okazji omawianego w niniejszym artykule zagadnienia zasługuje filozoficznie doniosły temat genezy norm w znaczeniu logicznym oraz różnice w filozoficznych poglądach von Wrighta i Kalinowskiego, twórców chronologicznie pierwszych systemów logiki norm, związane z tą kwestią.

12 G. H. von Wright, *Freedom and determination*, wyd. North-Holland Publishing Company, Amsterdam 1980, s. 65; 78.

13 G. H. von Wright, *Norm and Action: A Logical Inquiry*, wyd. Routledge & Kegan Paul Nowy Jork 1963, s. 2-6.

14 G. H. von Wright, *Norm and Action: A Logical Inquiry*, wyd. Routledge & Kegan Paul Nowy Jork 1963, s. 12.

Przywołując filozoficzne zagadnienie genezy norm należy najpierw przypomnieć o tym, że von Wright mówił przy tej okazji o filozoficznym rodowodzie norm w znaczeniu szerszym i węższym. Normy w znaczeniu szerszym, są to dla niego wszystkie zdania kierujące działaniem ludzkim, poczynając od zdań rozkazujących aż po zdania informujące, które mają charakter quasi-norm. Są to zdania w stylu: „jest przeciąg”, „ktoś puka”, itp. Nie zawierają one wyraźnie sprecyzowanego polecenia ani normy, ale sugerują istnienie pewnego stanu rzeczy, w którym należy podjąć pewne określone działanie. Ogólnie można stwierdzić, że wszystkie te zdania, które von Wright określał jako normy w znaczeniu szerszym, nazywamy zdaniami praktycznymi. Podobnie określał je również J. Kalinowski. Obydwaj autorzy nie poświęcili uwagi analizie zdań praktycznych *sensu largo* ponieważ nie ma w tym temacie różnic w poglądach filozoficznych i do tego jest to zagadnienie za zbyt ogólne.

Różnice w założeniach filozoficznych między twórcami logik norm pojawiają się w związku z zagadnieniem norm nazywanych przez von Wrighta normami w sensie węższym, czyli norm z dziedziny etyki, prawa, itp. Normy tego typu von Wright porównuje do reguł kierujących zachowaniem graczy w grze. Należy zauważyć, że reguły te nie tylko kierują pewnymi działaniami ludzkimi, ale także konstytuują samą grę. Pytając zatem o źródło i początek norm w sensie węższym, a więc bezpośrednich nakazów, zakazów lub dozwoleń, należałoby, według von Wrighta, postawić pytanie o źródło i początek reguł gry. Jak wiadomo, reguły określają samą grę. Dla przykładu można podać pewną roboczą definicję np. gry w szachy, głoszącą, że szachy są grą, rozgrywaną według takich a takich reguł. Jako takie, reguły stanowią zbiór preskrypcji, czyli przepisów. Nie mogą być one przez to ani prawdziwe ani fałszywe, bo one nie informują o rzeczywistości, tylko ją w jakiś sposób projektują. Dla przykładu reguły gry w szachy mówią wyraźnie, że goniec porusza się o dowolną ilość pól po skosie. Jeśli np. regułą gry w szachy byłoby stwierdzenie, że goniec porusza się o dowolną ilość pól w linii prostej, to wówczas goniec istotnie by się tak poruszał¹⁵.

Jak już wspominaliśmy, oprócz norm etycznych i prawnych do zbioru preskrypcji należą, według von Wrighta, pewne normy specyficzne. Pierwszą ich grupę logik fiński określa mianem norm „anankastycznych”, czyli „koniecznościowych”. Do takich należy m.in. norma głosząca, że „powinienem teraz wyjść z domu na stację (ponieważ jeśli nie wyjdę teraz z domu to nie zdążę na pociąg)”, „nie powinno się dotykać gorącego żelaza”, itp. Drugą grupę norm specyficznych von Wright nazywa normami technicznymi. Są to wyrażenia typu: „należy włączyć urządzenie do kontaktu”, „należy rozgrzać piekarnik do temperatury 200 stopni Celsjusza”, itp. Wspomniane tu normy – preskrypcje są normami w pewien sposób uwikłanymi w zdania oznajmujące (deskrypcje). Są to reguły podobne w istocie do reguł gry, ale są one formułowane tak, że mają współgrać z innymi, już wcześniej ustalonymi regułami¹⁶. Dla przykładu, jeżeli mówię, że pozwalam komuś używać mojego długopisu i istotnie pozwalam gdzieś w moim umyśle, to nie są to zdarzenia niezależne: wypowiedzianym zdaniem po prostu wyrażam normę, którą właśnie ustanowiłem. Jeśli ktoś inny wypowiedziałby się, że ja pozwalam komuś używać mojego długopisu, to wówczas nie wypowiedziałby normy (bo nie ma takiego prawa) ale tylko zdanie o normie. Możemy wypowiadać zdania o normach i mogą to być albo promulgacje, albo deskrypcje. Jeśli jesteśmy do tego uprawnieni do stanowienia

¹⁵ G. H. von Wright, *Norm and Action: A Logical Inquiry*, wyd. Routledge & Kegan Paul Nowy Jork 1963, s. 102-103.

¹⁶ G. H. von Wright, *Explanation and Understanding*, wyd. Cornell University Press, Nowy Jork 1972, s. 85-86.

norm, co nazywa się *genuine legal sentence*, to wyrażone przez nas zdanie o normie jest jej promulgacją. Jeśli oznajmimy zdanie o istnieniu normy nadanej przez kogoś innego, będzie to tzw. rzekome zdanie o normie (*spurious legal sentence*) – innymi słowy stwierdzenie normatywne (*normative statement*). Istnienie uprawnionej władzy do tworzenia norm jest, z filozoficznego punktu widzenia, źródłem norm (preskrypcji). Tak rozumiane normy (preskrypcje) mają logiczny status quasi-sądów. Jest tak dlatego, że nie przysługują im wartość logiczna. Można by powiedzieć, że są one samą treścią normatywną bez jakiegokolwiek momentu asercji. Tego typu quasi-sądy wzbogacone o moment asercji stają się deskrypcjami, czyli zdaniami o normach, których znaczeniem są same normy – preskrypcje¹⁷.

Tymczasem według Kalinowskiego, wykształconego w duchu filozofii klasycznej, wyrażenia zdaniowe są formalnym zapisem sądu logicznego. Sąd taki, który może być prawdziwy lub fałszywy, jest znaczeniem zdania logicznego. Zdanie nie oznacza konkretnego bytu, nie ma więc desygnatów. Niemniej jednak ma swoje desygnacje, czyli odniesienia do rzeczywistości. Ta rzeczywistością są pewne stany rzeczy. Opis danego stanu rzeczy jest treścią sądu logicznego. Oprócz tejże treści, w skład sądu wchodzi asercja lub negacja, czyli stwierdzenie bądź odrzucenie zachodzenia danego stanu rzeczy.

Jak już wspominaliśmy wcześniej sądy normatywne u von Wrighta są bytami myślnymi, projektami, inaczej mówiąc preskrypcjami. Jako takie nie mają wartości logicznej, lecz bodźcową, czyli deontyczną. Przedmiotem tych norm są stany rzeczy. W danym stanie rzeczy zawiera się zespół warunków, a więc sytuacja niezbędna do zaistnienia normy. Warunki te von Wright dzieli na kategoryczne i hipotetyczne. Według tego podziału dokonuje on również podziału norm. Norma jest kategoryczna, jeśli warunek jej zastosowania jest warunkiem, który musi być spełniony, innymi słowy warunek ten jest okazją do wykonania tego, co jest treścią normy i nie istnieją inne warunki. Norma jest hipotetyczna, jeśli warunek jej zastosowania jest warunkiem, który musi być spełniony oraz istnieją jakieś inne, dalsze warunki. Jeśli norma jest kategoryczna, to zwykle warunek jest zawarty w jej treści. Normy kategoryczne są nakładane przez autorytet normodawczy. Normy heteronomiczne są nakładane zawsze przez kogoś na kogoś innego, normy autonomiczne to takie, które są nałożone na samego siebie. Te ostatnie wiążą się zawsze z normami koniecznościowymi lub technicznymi. Normy moralne, którymi kieruje się podmiot działający mogą być jedynie heteronomiczne, jeśli pochodzą z jakiegoś istniejącego kodeksu norm. Jeśli natomiast podmiot wytwarza (projektuje) swoje własne normy postępowania, to wówczas von Wrighta uważa, że nie są one narzucone przez nikogo.

Istnieje jednak pewien klasyczny pogląd filozoficzny głoszącym iż powinności, pozwolenia i zakazy wywodzą się z ocen¹⁸. Zdania oceniające w paradygmacie von Wrighta zawierałyby treść odpowiadającą na pytanie: w jaką grę dany podmiot gra. Można by mówić o istnieniu wielu gier i zdania estymatywne służyłyby do wskazywania (wyboru) którejs z nich. Problem filozoficzny, jaki w tym miejscu się podnosi dotyczy tego, czy któraś z gier jest w jakiś sposób uprzywilejowana. Jeśli odpowiedź na to pytanie miałaby być twierdząca, to wówczas relacje jakie mogłyby decydować o ewentualnej prawdziwości lub fałszywości normy opierałyby na triadzie: gracz, gra i jej reguły. A ponieważ gra i jej reguły są to dwie strony jednej i tej samej definicji, kryteria tych ocen są – według fińskiego logika – użyteczne. Decydującą

17 G. H. von Wright, *The Varieties of Goodness*, wyd. Humanities Press, Nowy Jork 1963, s. 157-159.

18 G. H. von Wright, *The Varieties of Goodness*, wyd. Humanities Press, Nowy Jork 1963, s. 155.

rolę odgrywa a niech osoba grająca. Dokonany wybór gry sprawia natomiast, że podmiot w nią grający jest takim a nie innym graczem. Jeśli normy miałyby być nośnikami jakichś wartości, to należy tutaj podkreślić wartości prakseologiczne. Do tej grupy koncepcji należą pojęcia sprawczości, działania, aktywności, zachowania, wyboru, decyzji i wolności. Według von Wrighta deontologia i aksjologia wymagają prakseologii, a nawet są jej gałęziami. Zwolennikom poglądu upatrującego genezę norm w ocenach, von Wright powiedziałby, że chodzi tu raczej o wszystkie zdania praktyczne, a więc normy w sensie szerszym, niż normy w sensie ścisłym¹⁹.

Zwolennikiem poglądu głoszącego że normy mają swoją genezę w ocenach był J. Kalinowski. Konsekwencją myślenia w duchu Arystotelesa i św. Tomasza z Akwinu było przyjęcie przez Kalinowskiego zasady *ens et bonum/pulchrum conventuntur*. Gosi ona przyznawanie bytom proporcjonalnych wartości. Byt jest dobrem, i na tej podstawie wszystko, przez co potencjalności danego bytu są aktualizowane jest dlań dobre. Zdania normatywne są tym samym również zdaniem moralnym. Norma zaś jest moralna, kiedy zakazuje, nakazuje lub dozwala wykonanie świadomej i wolnej czynności, która jako taka jest mniej lub bardziej zgodna z bytowością podmiotu, a przez to cechuje się wartością²⁰.

Ze względu na źródło, normy można podzielić na autonomiczne, czyli te pochodzące od podmiotu i heteronomiczne, czyli te pochodzące spoza podmiotu. Normy autonomiczne mają określoną strukturę i są powiązane z konkretnym podmiotem: „x powinien/może/nie może czynić a”. Normy autonomiczne opierają się na sumieniu (świadomości moralnej). Normy moralne heteronomiczne nazywamy normami prawnymi. Jak już zostało wspomniane normy anankastyczne u von Wrighta same w sobie są autonomiczne, niemniej jednak w jakiś sposób wiążą się i wynikają ze zdań logicznych opisujących procesy mające oparcie w prawie natury. Kalinowski również poświęca sporo uwagi temu tematowi w swoich rozważaniach, powołując się w nich na myśli św. Tomasza. Jeśli chodzi o prawo naturalne, to Akwinata nazywa je za Awicenną *universale ante rem*. Jest to więc przyczyna, która niejako zdefiniowała naturę istniejąc wcześniej w umyśle Bożym. W czasie stworzenia stała się ona *universale in re* i jako pojęcie, które jest odczytane z natury Tomasz nazywa je *universale post rem*. Z tego względu tylko normy wchodzące w skład prawa natury są normami kategorycznymi. Wszystkie pozostałe, zarówno normy moralne jak i normy prawa stanowionego są normami hipotetycznymi. Te pierwsze są normami autonomicznymi, zaś te ostatnie – heteronomicznymi²¹.

W zbiorze norm autonomicznych Kalinowski, w przeciwieństwie do von Wrighta, wyróżnia normy autonomiczne pierwszoosobowe i trzecioosobowe. Te ostatnie zakresowo pokrywają się z autonomicznymi normami koniecznościowymi i technicznymi von Wrighta. wolno nam mniemać, że w podobny sposób postrzegał normy moralne również Leibniz²².

Pozytywne prawo stanowione to nie tylko inny rodzaj heteronomii, ale także inny rodzaj normy w sensie właściwym (językowym i logicznym). Prawo naturalne to pewna obiektywna rzeczywistość, z którą dane normy mogą być zgodne lub nie. Normy pozytywne, czyli stanowione są w tym paradygmacie wnioskami sylogizmów, w których przesłankach odnajdujemy zdania prawa naturalnego. Norma może być wówczas zgodna z zamysłem twórcy i na tej podstawie można dopatrywać się

19 G. H. von Wright, *An essay in deontic logic and the general theory of action*, wyd. North-Holland Publishing Company, Amsterdam 1968, s. 11-12.

20 J. Kalinowski, *Le problème de la vérité en morale et en droit*, wyd. Emanuel Vitte, Lyon 1967, s. 185-187.

21 J. Kalinowski, *Le problème de la vérité en morale et en droit*, wyd. Emanuel Vitte, Lyon 1967, s. 190.

22 J. Kalinowski, *Le bien, le morale et la justice*, *Archives de philosophie du droit*, tom 11, *La logique du droit*, s. 324.

prawdziwości normy, ponieważ twórca zawsze nadaje prawo naturalne temu, co tworzy, i to niezależne od tego czy twórcą jest człowiek, czy Absolut.

Warto wspomnieć tutaj jeszcze o moralnej genezie norm. Wszystkie normy zależne od działania moralnego, zarówno normy heteronomiczne jak i autonomiczne są nazywane przez Kalinowskiego normami moralnymi w sensie szerokim dlatego, że one określają działanie moralne i je regulują; ustawiają je względem celu ostatecznego i jemu wyznaczają właściwe środki do jego osiągnięcia. Moralność w szerokim sensie jest stanowiona przez prawo naturalne, prawo ludzkie i sumienie (jego nakazy), na które składają się częściowo oceny moralne. Stoją one u jego podstaw nakazów sumienia. Oprócz ocen moralnych sumienie człowieka wydaje także nakazy moralne, które wynikają ze wspomnianych tu ocen. Nazwa „prawo moralne” w sensie szerszym oznacza więc prawo naturalne i prawo ludzkie. Natomiast prawo moralne w sensie ograniczonym to prawo naturalne i reguły sumienia. Składają się na nie normy mające autora, ale nie promulgowane *explicite*. Z tej właśnie przyczyny nie wydaje nam się, że sądy o wartości utożsamiają się z sądami psychologicznymi („uważam to i to działanie za dobre”), czy socjologicznymi („większość uważa to i to za dobre”).

Kalinowski za, św. Tomaszem, stwierdza, że akt istnienia jest pierwszą doskonałością (wartością) bytu, i jest otrzymany z zewnątrz. Stanowi on wartość ontologiczną danego bytu. Niemniej jednak należy odróżnić wartości ontologiczne od moralnych. Te ostatnie realizują się przez działanie jako osoby, jako substancji i jako istoty rozumnej i wolnej. Termin „wartość” w pojęciu wartości ontycznej jest analogiczny, podobnie jak „istnienie”, ponieważ istnieje się na swój proporcjonalny sposób. Wartość ontyczną mierzy się właśnie „ilością” bytu. Pełność ontyczna musi być z natury pożądana przez każde stworzenie, ale człowiek może skierować swoją pożądaną miłość *amor concupiscentiae* w innym kierunku. W ten sposób cel moralny, który człowiek obiera może być więc prawdziwy lub fałszywy²³.

Zakończenie

Podjęty w niniejszym artykule temat jest jednym z wielu ważnych zagadnień filozoficznych wymagających zgłębienia przy okazji badania problematyki norm. Został on opracowany w dwóch zasadniczych częściach. Pierwsza z nich poświęcona była zestawieniu i porównaniu poglądów G. H. von Wrighta i J. Kalinowskiego – twórców systemów logiki norm – na temat językowych wyrażań kierujących działaniem ludzkim. Przedstawienie tych poglądów oraz porównanie podziałów zdań praktycznych proponowanych przez obu logików pozwoliło na określenie obecności norm w języku oraz charakterystykę ich podstawowych rodzajów. Druga część refleksji prowadzonej w tym artykule poświęcona była zagadnieniu filozoficznej genezy norm. Zostało w niej uwypuklone odniesienie wspomnianych tu autorów do tematu relacji norm do ocen. Pozwoliło to na wyprowadzenie kilku interesujących wniosków, jak również wskazało pewne obszary badawcze, które zasługują w naszym mniemaniu na dalszą analizę w osobnych artykułach. Chodzi tu między innymi o wyrażalność ocen i norm w języku formalnym oraz oparcie aksjomatów ewentualnych formalnych rachunków logiki norm o wcześniejsze, filozoficzne rozważania dotyczące ich genezy.

Warto w tym miejscu zauważyć, że zarówno oceny jak i normy nie koniecznie

23 J. Kalinowski, *Le problème de la vérité en morale et en droit*, wyd. Emanuel Vitte, Lyon 1967, s. 236-238.

muszą być wyrażone w języku zarówno naturalnym jak i formalnym, aby spełniać swoje funkcje. Na uwagę zasługuje jednak to, że istnieją oceny niewyraźne w języku, podczas gdy takich norm nie ma. Co do pochodzenia obu, von Wright twierdzi, że oceny są jedynie strukturalnie wcześniejsze od norm. Dość powiedzieć, że zagadnienie rodzaju norm, ich źródła oraz genezy nie znajduje odzwierciedlenia w rachunkach logiki deontycznej von Wrighta. Tymczasem Kalinowski uważa, że wszystkie te kwestie filozoficzne powinny znaleźć odzwierciedlenie w rachunku logicznym. Normy pochodzą od ocen na tej samej zasadzie, na jakiej normy hipotetyczne pochodzą od norm kategorycznych, a tezy systemu od jego aksjomatów. Mamy tutaj kolejną istotną różnicę między filozoficznymi założeniami rachunków logiki norm, która zasługuje na oddzielne opracowanie.

Bibliografia:

1. J. Kalinowski, *Le problème de la vérité en morale et en droit*, wyd. Emanuel Vitte, Lyon 1967;
2. J. Kalinowski, *Le bien, le morale et la justice*, *Archives de philosophie du droit*, tom 11, *La logique du droit*, s. 311-329;
3. J. Kalinowski, *Teoria poznania praktycznego*, Towarzystwo Naukowe KUL, Lublin 1960;
4. J. Kalinowski, *Logika Norm*, Instytut Wydawniczy Daimonion, Lublin 1993;
5. G. H. von Wright, *An essay in deontic logic and the general theory of action*, wyd. North-Holland Publishing Company, Amsterdam 1968;
6. G. H. von Wright, *The Varieties of Goodness*, wyd. Humanities Press, Nowy Jork 1963;
7. G. H. von Wright, *Explanation and Understanding*, wyd. Cornell University Press, Nowy Jork 1972;
8. G. H. von Wright, *Norm and Action: A Logical Inquiry*, wyd. Routledge & Kegan Paul Nowy Jork 1963;
9. G. H. von Wright, *Freedom and determination*, wyd. North-Holland Publishing Company, Amsterdam 1980.

Anita Mira Sobczyk
Katolicki Uniwersytet Lubelski Jana Pawła II

Johna Rogersa Searle'a próba ukazania możliwości wyprowadzenia powinności z faktu.

W artykule „*Jak wywieść powinien z jest*”¹ John Rogers Searle stara się obalić tezę o niewyprowadzalności zdań normatywnych z orzekających, funkcjonującą powszechnie pod nazwą *gilotyny Hume'a*². W tym celu podaje kontrprzykład dla niej, podkreślając jednocześnie, że odnosi się on do współczesnej wersji tej tezy. Istota owego kontrprzykładu polega na tym, iż w punkcie wyjścia przyjmuje twierdzenie (ewentualnie zbiór twierdzeń) o charakterze opisowym by następnie ukazać związek logiczny między nim a pewnym twierdzeniem wartościującym, występującym w punkcie dojścia³.

Co istotne, Searle uświadamia sobie wyraźnie fakt, że jeden kontrprzykład nie jest w stanie obalić tezy, która tak silnie zadomowiła się w teorii moralnej. Jak sam zauważa: „*nie należy oczywiście sądzić, że jeden kontrprzykład może obalić tezę filozoficzną, jeśli jednak uda nam się przedstawić przekonujący kontrprzykład i na dodatek wyjaśnić w jaki sposób i dlaczego stanowi on wystarczający kontrprzykład, a nadto jeśli potrafimy przedstawić teorię, która leży u podstaw tego kontrprzykładu i dostarczyć może nieskończonej liczby innych kontrprzykładów – uda nam się w najgorszym razie rzucić więcej światła na wyjściową tezę tych rozważań*”⁴.

Celem niniejszego artykułu jest zaprezentowanie stanowiska Johna Searle'a w omawianej kwestii. Aby to uczynić koniecznym jest przytoczenie w punkcie wyjścia szeregu twierdzeń, analizowanych przez interesującego nas autora. Przyjrzyjmy się zatem poniższym wypowiedziom:

1. Jones wypowiedział słowa: „Niniejszym obiecuję ci, Smith, zapłacić 5 dolarów”.
2. Jones obiecał zapłacić 5 dolarów Smithowi.
3. Jones wziął na siebie (podjął) zobowiązanie zapłacenia 5 dolarów Smithowi.
4. Jones jest zobowiązany zapłacić 5 dolarów Smithowi.
5. Jones powinien zapłacić 5 dolarów Smithowi⁵.

Teza Searla jest następująca: „*choć poszczególne zdania na tej liście nie*

¹ John R. Searle, *How to Derive Ought from Is*, „Philosophical Review” 73(1964), s. 43-58. W niniejszym artykule odwołuję się do przekładu polskiego: *Jak wywieść „powinien” z „jest”*, tłum. J. Hołówka, „Etyka” 16(1978), s. 163-177.

² Za autora tego określenia uznaje się Maxa Blacka, który jako pierwszy posłużył się nim dla określenia tezy Hume'a, mającej niczym gilotyna doprowadzić do zaprzestania stosowanej od stuleci praktyki wyprowadzania „znikąd” zdań w których występuje słowo „powinien” i inne jego odmiany. Por. M. Black, *The Gap Between „is” and „should”*, „The Philosophical Review”, 73(1964) s. 165-181. Sama zaś teza pochodzi z *Traktatu o naturze ludzkiej*. Por. David Hume, *Treatise on Human Nature*, III, 1,1; (wydanie polskie: D. Hume, *Traktat o naturze ludzkiej*, tłum. Cz. Znamierowski, Warszawa 2005, s. 531-547.

³ Por. *Jak wywieść „powinien” z „jest”*, s. 163-164.

⁴ Tamże, s. 163.

⁵ Tamże, s. 164.

*pociągają zdań bezpośrednio po nich następujących mocą wynikania logicznego, to jednak nie łączy ich wyłącznie związek przypadkowy, a co więcej dodatkowe twierdzenia, które zamieniają relację istniejącą między nimi na wynikanie logiczne, nie muszą zawierać twierdzeń wartościujących, zasad moralnych ani innych wyrażen tego typu*⁶.

Chcąc udowodnić powyższą tezę Searle wyjaśnia na jakiej podstawie jest możliwe przejście z jednego twierdzenia do następnego w rozważanym zbiorze wypowiedzi.

I tak, uzasadniając przejście od zdania (1) do zdania (2), autor ukazuje, że w pewnych okolicznościach wypowiedzenie słów: „*Niniejszym obiecuję ci ...*”, stanowi czynność złożenia obietnicy⁷. Co więcej, sam ten zwrot jego zdaniem jest „*paradygmatyczną formułą używaną dla dokonania czynu opisanego w zdaniu drugim jako obiecywanie*”⁸.

Aby jeszcze bardziej wyrazić zasadność przejścia od (1) do (2), Searle przyjmuje dwie dodatkowe przesłanki. Pierwszą z nich stanowi zdanie:

- 1.1. W pewnych okolicznościach C, każdy kto wypowiada słowa (zdanie): „*Niniejszym obiecuję ci, Smith zapłacić 5 dolarów*”, obiecuje zapłacić 5 dolarów Smithowi⁹.

W przekonaniu Searle’a powyższa przesłanka jest tautologią, która wprawdzie podnosi formalną elegancję samego wywodu, ale jednocześnie pod względem logicznym jest zbędna¹⁰. Wymienione okoliczności C zaś „*są to te okoliczności lub stany rzeczy, które stanowią konieczny i wystarczający warunek tego, by wypowiedzenie owych słów (zdania) było udanym przypadkiem złożenia obietnicy*”¹¹. A zatem, okoliczności C decydują o skutecznym złożeniu obietnicy.

Searle zauważa również, że wspomniane warunki C mają charakter empiryczny i na tej podstawie przyjmuje kolejną (empiryczną tym razem) przesłankę dodatkową:

- 1.2 Okoliczności C zachodzą¹².

Tym samym na podstawie zdań:

- W pewnych okolicznościach C, każdy kto wypowiada słowa (zdanie): „*Niniejszym obiecuję ci, Smith zapłacić 5 dolarów*”, obiecuje zapłacić 5 dolarów Smithowi. (1.1)
- Jones wypowiedział słowa: „*Niniejszym obiecuję ci, Smith, zapłacić 5 dolarów*”. (1)
- Okoliczności C zachodzą. (1.2)

Searle wyprowadza zdanie:

- Jones obiecał zapłacić 5 dolarów Smithowi¹³. (2)

W rezultacie, ze zdania (1) przechodzimy do zdania (2).

Wyjaśniając natomiast związek między zdaniem 2 a 3, Searle odwołuje się do definicji pojęcia *obiecywanie*. Jego zdaniem „*żadna analiza pojęcia obiecywania nie jest wyczerpująca, jeśli nie uwzględni, że obiecujący bierze na siebie pewne zobowiązanie, podejmuje je, i że zobowiązanie to dotyczy dokonania w przyszłości*

6 Tamże, s. 164.

7 Por. tamże, s. 164.

8 Tamże, s. 164.

9 Tamże, s. 164.

10 Por. tamże, s. 167.

11 Tamże, s. 164.

12 Tamże, s. 165.

13 Por. tamże, s. 165.

pewnego czynu, zazwyczaj na rzecz osoby, której obietnicę złożono”¹⁴. A następnie dodaje: „skłonny więc jestem uznać, że (2) pociąga za sobą bezpośrednio (3)”¹⁵.

Dla porządku formalnego, podobnie jak w przypadku wcześniejszego kroku, przyjmuje on dodatkową przesłankę tautologiczną, która wywodzi się z analizy definicji pojęcia obiecywanie:

- 2.1 Wszystkie obietnice są przypadkami wzięcia na siebie (podjęcia) zobowiązania wykonania rzeczy obiecaniej¹⁶.

Natomiast, wyjaśniając przejście od (3) do (4), to znaczy ukazując jak z faktu podjęcia zobowiązania wynika fakt bycia zobowiązanym, Searle przyjmuje kolejne założenia dodatkowe:

- 3.1 Nie zachodzą nadzwyczajne okoliczności.
- 3.2 Wszyscy, którzy biorą na siebie zobowiązanie, są *ceteris paribus* zobowiązani zobowiązaniem¹⁷.

Co ciekawe, tautologia wyjaśniająca relację między (3) a (4) ma również zastosowanie przy wyjaśnianiu związku pomiędzy (4) a (5). Mianowicie, zdaniem Searle’a „mamy tu tautologię stwierdzającą, że *ceteris paribus* każdy powinien robić to, co jest zobowiązany robić. Podobnie więc jak i w poprzednim przypadku potrzebna jest tu przesłanka mówiąca, że:

- 4.1 Nie zachodzą nadzwyczajne okoliczności”¹⁸.

W obu przytoczonych przypadkach klauzula *ceteris paribus*, wyklucza możliwość zakłóceń przez jakieś okoliczności (czynniki) zewnętrzne. Searle podkreśla: „sens wyrażenia *nie zachodzą nadzwyczajne okoliczności* jest w niniejszym przypadku z grubsza rzecz biorąc taki: jeśli nie mamy określonego powodu (to znaczy, jeśli nie jesteśmy gotowi podać określonego powodu), by przypuszczać, że zobowiązanie jest nieważne (zdanie 4) lub że składający obietnicę nie powinien jej dotrzymać (zdanie 5), wówczas zobowiązanie jest prawomocne i obiecujący powinien dotrzymać obietnicy”¹⁹. A zatem owe „nadzwyczajne okoliczności” są to najrozmaitsze stany rzeczy anulujące zobowiązanie osoby, która je podjęła.

Biorąc pod uwagę wszystkie dodatkowe założenia rozumowanie Searle’a przybiera następującą postać:

1. Jones wypowiedział słowa: „Niniejszym obiecuję ci, Smith, zapłacić 5 dolarów”.
 - 1.1 W pewnych okolicznościach C, każdy kto wypowiada słowa (zdanie): „Niniejszym obiecuję ci, Smith zapłacić 5 dolarów”, obiecuje zapłacić 5 dolarów Smithowi.
 - 1.2 Okoliczności C zachodzą.
2. Jones obiecał zapłacić 5 dolarów Smithowi.
 - 2.1 Wszystkie obietnice są przypadkami wzięcia na siebie (podjęcia) zobowiązania wykonania rzeczy obiecaniej.
3. Jones wziął na siebie (podjął) zobowiązanie zapłacenia 5 dolarów Smithowi.
 - 3.1 Nie zachodzą nadzwyczajne okoliczności.
 - 3.2 Wszyscy, którzy biorą na siebie zobowiązanie, są *ceteris paribus* zobowiązani zobowiązaniem²⁰.

¹⁴Tamże, s. 165.

¹⁵Tamże, s. 165.

¹⁶Tamże, s. 166.

¹⁷Tamże, s. 166.

¹⁸Tamże, s. 166.

¹⁹Tamże, s. 167.

²⁰Tamże, s. 166.

4. Jones jest zobowiązany zapłacić 5 dolarów Smithowi.

4.1 Nie zachodzą nadzwyczajne okoliczności.

5. Jones powinien zapłacić 5 dolarów Smithowi.

Searle podkreśla jednocześnie, że dodatkowe przesłanki, którymi się posłużył w dowodzie nie mają charakteru sformułowań moralnych lub wartościujących. Przesłanki te były bowiem założeniami empirycznymi, tautologiami lub opisami użycia słów. Dzięki temu jak sam podkreśla, udało mu się wywieść powinność z faktu. Co więcej, owa powinność nie ma charakteru hipotetycznego, ale kategoryczny²¹. Zdanie (5) bowiem w przekonaniu autora „nie stwierdza, że Jones powinien płacić Smithowi, jeśli chce realizacji tego lub innego celu. (5) stwierdza, że Jones ma płacić i kropka”²².

Jak zauważa Searle: „dowód ten ujawnia związek pomiędzy wypowiedzeniem pewnych słów i czynnością językową obiecywania, a następnie ujawnia, w jaki sposób obiecywanie polega na podjęciu zobowiązania, a zobowiązanie pociąga za sobą powinność”²³.

Searle swoje rozważania uzupełnia następującym stwierdzeniem: „Tendencja do przyjmowania ścisłego rozróżnienia między jest i powinien, między wyrażeniami opisującymi i wartościującymi, opiera się na pewnym wyobrażeniu sposobu, w jaki słowa odnoszą się do świata. Jest to bardzo ponętny obraz, tak bardzo (przynajmniej dla mnie), że nie jest całkiem jasne, w jakim stopniu ograniczenie się do przedstawienia kontrprzykładów może to wyobrażenie podważyć. (...) Mówiąc krótko, stwarzamy sobie ów obraz mniej więcej tak: najpierw przedstawiamy przykłady tak zwanych stwierdzeń opisowych (Mój samochód jedzie 80 mil na godzinę, Jones ma 6 stóp, Smith ma brązowe włosy), a następnie przeciwstawiamy je tak zwanym twierdzeniom wartościującym (Mój samochód jest dobry, Jones powinien zapłacić 5 dolarów Smithowi, Smith to łajdak). Każdy widzi, że są to różne wyrażenia. Podkreślamy tę różnicę wskazując, że stwierdzenia opisowe mogą być obiektywnie rozstrzygane pod względem wartości logicznej: ponieważ zrozumienie znaczenia wyrażen opisowych jest identyczne ze zrozumieniem, w jakich obiektywnie dających się ustalić warunkach stwierdzenia te są prawdziwe lub fałszywe. W przypadku stwierdzeń wartościujących sytuacja jest zupełnie inna. Samo zrozumienie znaczenia wyrażen wartościujących nie wystarcza do zrozumienia, w jakich warunkach są prawdziwe stwierdzenia, które je zawierają, ponieważ sens tych wyrażen jest taki, że zawierające je stwierdzenia nie mogą być w ogóle uznane za obiektywnie lub faktycznie prawdziwe bądź fałszywe. (...) Stwierdzenia opisowe są więc obiektywne, stwierdzenia wartościujące są subiektywne i różnica ta wynika z różnicy między typami wyrażen, które w obu przypadkach zostały użyte. U podstaw tej różnicy leży fakt, że stwierdzenia wartościujące wykonują zupełnie inne zadania niż stwierdzenia opisowe. Ich zadaniem nie jest opisywanie jakichkolwiek właściwości świata, lecz wyrażanie uczuć mówiącego, jego postaw, pochwał i nagan, zachwytów i przekleństw, zaleceń, poleceń, rad itd. Gdy spostrzeżemy, jak różne są to zadania, dochodzimy do przekonania, że musi istnieć przepaść logiczna między obydwoma typami stwierdzeń. Stwierdzenia wartościujące muszą być inne od stwierdzeń opisowych, jeśli mają wypełniać swą funkcję, bo gdyby miały charakter obiektywny, nie mogłyby służyć do oceniania. Używając języka metafizycznego można powiedzieć, że wartości nie mogą się znajdować w świecie, bo wtedy przestałyby

21 Por. tamże, s. 168.

22 Tamże, s. 168.

23 Tamże, s. 168.

być wartościami, a byłyby tylko pewną częścią świata. Mówiąc bardziej formalnie, nie można zdefiniować słów wartościujących przy pomocy słów opisowych, bo gdybyśmy to zrobili, nie moglibyśmy już używać słów wartościujących po to, by za ich pośrednictwem zalecać, lecz jedynie po to, by przy ich pomocy opisywać. Mówiąc jeszcze inaczej, każda próba wywiedzenia powinności z bytu jest stratą czasu, bo wszystko, co udało by się ustalić, o ile wywód zostałby przeprowadzony poprawnie, sprowadzało by się do tego, że użyte w nim jest nie byłoby prawdziwym jest, lecz jedynie zamaskowanym powinien i każde powinien nie byłoby prawdziwym powinien, lecz jedynie zamaskowanym jest²⁴.

Zdaniem Searle'a, konsekwencją przyjęcia powyższego stanowiska jest niemożność wytłumaczenia takich pojęć jak: *obowiązek*, *odpowiedzialność*, *zobowiązanie*. Ponadto, powyższe stanowisko nie rozróżnia typów stwierdzeń opisujących. Istnieje bowiem różnica między wypowiedziami: *Mój samochód jedzie 80 mil na godzinę*, *Jones na sześć stóp*, a zdaniami *Jones się ożenił*, *Jackson ma pięć dolarów*. Ostatnie dwa zdania są przykładami twierdzeń opisujących fakty, których zajście zakłada istnienie określonych instytucji. I tak, zdanie *Jones się ożenił* zakłada istnienie instytucji małżeństwa oraz obietnicy, zaś zdanie *Jackson ma pięć dolarów*, instytucji pieniądza. Na tej podstawie Searle odróżnia fakty instytucjonalne (zakładające istnienie instytucji jak to było opisane powyżej) od nagich faktów²⁵. Aby uwyraźnić swoje stanowisko autor podaje przykład: „*To, że ktoś ma w ręce kawałek papieru z zielonym nadrukiem jest nagim faktem, to że ma pięć dolarów jest faktem instytucjonalnym*”²⁶.

Searle wprowadza do swoich rozważań rozróżnienie na dwa rodzaje reguł. Pierwsze z nich to reguły regulujące wcześniej istniejące formy zachowania. Określa je mianem *reguł regulatywnych*. Przykładem takich reguł są zasady zachowania się przy stole. Z kolei *reguły konstytutywne*, stwarzają, bądź też definiują nowe, nie istniejące dotychczas formy zachowania (przykładem podanym w artykule są reguły szachowe)²⁷. Jak podkreśla Searle: „*Reguły regulatywne regulują czynności, których zachodzenie nie zależy od reguł. Reguły konstytutywne konstytuują (a nadto regulują) formy zachowania, których istnienie logicznie zależy od owych reguł*”²⁸. Poprzez powyższy zabieg odróżnienia reguł regulatywnych od konstytutywnych autor przejmuje (całkiem świadomie) terminologię Kantowskiego rozróżnienia na zasady regulatywne i konstytutywne²⁹.

Wspomniane wcześniej instytucje *małżeństwa*, *pieniędzy* oraz *obietnicy*, zdaniem Searle'a są systemami *reguł konstytutywnych*. Fakty instytucjonalne zakładają istnienie takich instytucji. Zatem, to czy ktoś złożył obietnicę bądź też przyjął na siebie określony obowiązek jest kwestią *faktów instytucjonalnych*, nie zaś *nagich faktów*. Searle wyprowadza dalej *powinien z jest* odwołując się do pojęcia faktu instytucjonalnego. W punkcie wyjścia pojawił się nagi fakt wypowiedzi pewnych słów. Ze zdania (1) dowiadujemy się, że Jones wypowiedział słowa: „*Niniejszym obiecuję ci, Smith, zapłacić 5 dolarów*”. Fakt ten, w przekonaniu Searle'a odwołuje się jednak do *instytucji obiecywania*, a tym samym prowadzi do *faktu instytucjonalnego*, jakim jest stwierdzenie: Jones powinien zapłacić 5 dolarów Smithowi. Cały zaś przytoczony dowód odwołuje się do *reguły konstytutywnej*

24 Tamże, s. 172-173.

25 Por. tamże, s. 173-174.

26 Tamże, s. 174.

27 Por. tamże, s. 174.

28 Tamże, s. 174.

29 Por. tamże, s. 174.

głoszącej, że składający obietnicę bierze na siebie zobowiązanie³⁰.

A zatem, zdaniem Searle'a „przesłanki faktyczne mogą pociągać za sobą konkluzje wartościujące mogą wynikania logicznego. (Jeśli mam rację to rzekome rozróżnienie między wypowiedziami opisującymi a wartościującymi można utrzymać jedynie jako rozróżnienie pomiędzy dwoma rodzajami siły wypowiedzi: opisującym i wartościującym.)”³¹.

Searle najwyraźniej świadomy tego, że jego dowód możliwości wyprowadzenia zdań normatywnych z orzekających może się spotkać z krytyką, w artykule „*Jak wywieść powinien z jest*” odparł trzy zarzuty, jakie jego zdaniem można sformułować przeciw wygłoszonej tezie. W niniejszym artykule pominięta zostanie prezentacja owych zarzutów. Czytelnik może się z nimi szczegółowo zapoznać sięgając do tekstu źródłowego³². W tym miejscu natomiast pojawi się rekonstrukcja zarzutu sformułowanego przeciw tezie Searle'a przez Jaakko Hintikka. W swojej książce „*Eseje logiczno-filozoficzne*”³³ przedstawia on błąd, jaki się kryje w Searle'a próbie ukazania możliwości wyprowadzenia powinności z faktu.

Hintikka ukazuje, że Searle z przesłanki będącej aktem obietnicy (a zatem posiadającej naturę czysto faktualną) wyprowadza obowiązek dotrzymania owej obietnicy. „*Krótko mówiąc, powinność wynika z faktu oraz dodatkowej tautologicznej, a więc pozbawionej treści, przesłanki*”³⁴.

W punkcie wyjścia swoich rozważań stwierdza on: „*Niech p będzie zdaniem mówiącym o złożeniu jakiejś konkretnej obietnicy, a q niech będzie stwierdzeniem jej dotrzymania. Analizę Searle'a można powiązać z naszą dotychczasową dyskusją, nadając jego drugiej, tautologicznej przesłance takie sformułowanie, że złożenie obietnicy, o której mowa, zobowiązuje do jej dotrzymania. Wówczas wyprowadzenie powinności z faktu Searle'a przybiera postać:*

$$\frac{p \quad p \text{ zobowiązuje do } q}{O q} \text{ }^{35}.$$

Hintikka zauważa ponadto, że istnieją dwie możliwe interpretacje rozumowania Searle'a. Przybierają one następujące postaci:

(a)

$$\frac{p \quad O(p \rightarrow q)}{O q}$$

³⁰ Por. tamże, s. 175.

³¹ Tamże, s. 177.

³² Tamże, s. 169-171.

³³ Jaakko Hintikka, *Eseje logiczno-filozoficzne*, Warszawa 1992, s. 360-370.

³⁴ Tamże, s. 361.

³⁵ Tamże, s. 361.

(b)

$$\begin{array}{c}
 p \\
 p \rightarrow Oq \\
 \hline
 Oq
 \end{array}$$

Hintikka wykazuje, że obowiązek o którym mowa w (a) jest obowiązkiem *prima facie*. Samo zaś rozumowanie nie jest poprawne. Natomiast obowiązek występujący w (b) jest obowiązkiem bezwzględny. Rozumowanie (b) jest rozumowaniem poprawnym. Jednak, jak zauważa autor druga przesłanka nie jest przesłanką analityczną. Podkreśla on, że ze złożenia obietnicy nie może wynikać analitycznie bezwzględny obowiązek jej dotrzymania. Searle zdaniem Hintikki przeoczył możliwość unieważnienia powstałego w wyniku obietnicy obowiązku *prima facie* przez jakiś konkurencyjny w stosunku do niego obowiązek. Skoro zaś przesłanka druga nie jest przesłanką analityczną, może być fałszywa. A zatem, oba sposoby zrekonstruowania rozumowania Searle'a są obarczone błędem, co zdaniem Hintikki udowadnia, że wniosek normatywny bezpośrednio nie wynika z przesłanek opisowych. Koniecznym jest dołączenie do rozumowania przesłanki o charakterze normatywnym³⁶.

Hintikka pokazuje, że zwolennicy rozumowania (b) mogą bronić jego prawomocności na dwa sposoby. Pierwszym jest nałożenie na pojęcie obietnicy dodatkowych wymogów. Tym samym mocne pojęcie obietnicy (autentyczna obietnica) wiąże się z wzięciem na siebie bezwzględnego obowiązku jej dotrzymania. Obowiązek ten nie może być zniesiony, bądź anulowany przez jakiegokolwiek wyższe względy. „Bez względu na to, jak wyraźnie wypowiada ona (osoba – przyp. A. S.) obiecuję, jak jest szczerą i do jakiego stopnia inne faktycznie potwierdzalne przesłanki skutecznego obiecywania są spełnione, nie da się w zamierzonym tutaj sensie orzec, że złożenie obietnicy miało miejsce (zamierzony sens jest tutaj taki, iż obiecywanie oznacza przyjęcie na siebie bezwzględnego zobowiązania dotrzymania słowa), jeżeli nie da się jednoznacznie ustalić, że nie ciąży na niej ważniejszy konkurencyjny obowiązek. Lecz wykazanie tego nie sprowadza się do ustalenia faktycznego stanu rzeczy. Polega ono na globalnej ocenie sytuacji normatywnej w celu stwierdzenia, czy pewna powinność *prima facie* nie została unieważniona. Krótko mówiąc, prawdziwość *p* nie jest kwestią tylko faktów, *p* nie jest przesłanką faktyczną, ale normatywną”³⁷.

Przy takiej interpretacji rozumowania (b) pierwsza przesłanka staje się przesłanką normatywną. A zatem, rozumowanie staje się zdaniem Hintikki bezużytecznym dla uzasadnienia możliwości wyprowadzenia obowiązku z faktu³⁸.

Drugi sposób obrony rozumowania (b) polega na dołączeniu do zdania *p* zdania *cp*, które stwierdza zachodzenie (spełnienie) określonych warunków *ceteris paribus*. Tym samym, zdanie *p* będzie zastąpione koniunkcją $p \wedge cp$ ³⁹. W rezultacie powinność bezwzględna dotrzymania obietnicy nie wynika z samego faktu jej złożenia, ale wynika jeśli zachodzi warunek bądź warunki *ceteris paribus*. Jednakże, Hintikka wykazuje, że to co zostało wypowiedziane wcześniej w odniesieniu do

³⁶ Por. tamże, s. 363.

³⁷ Tamże, s. 364.

³⁸ Por. tamże, s. 365.

³⁹ Por. tamże, s. 364.

zdania p ma również *mutatis mutandis* zastosowanie do koniunkcji $p \wedge cp$. Zatem, koniunkcja ta nie będzie zdaniem czysto faktualnym. Warunek cp będzie przynajmniej częściowo zdaniem normatywnym⁴⁰.

Powyższe interpretacje rozumowania (b) można zapisać w następujący sposób:

(b.1)

$$\frac{\begin{array}{l} p \text{ (faktualna)} \\ p \rightarrow Oq \text{ (nie-analityczna)} \end{array}}{Oq}$$

(b.2)

$$\frac{\begin{array}{l} p \text{ (normatywna)} \\ p \rightarrow Oq \text{ (analityczna)} \end{array}}{Oq}$$

(b.3)

$$\frac{\begin{array}{l} (p \wedge cp) \text{ (normatywna)} \\ (p \wedge cp) \rightarrow Oq \text{ (analityczna)} \end{array}}{Oq}$$

Zdaniem Hintikki „choć wszystkie trzy przedstawiają wnioskowania logicznie prawomocne, nie dostarczają nam jednak wyprowadzenia powinności z faktu w zamierzonym sensie”⁴¹. Najbardziej przekonującą jest interpretacja (b.3). Jednak i w tym przypadku Hintikka zarzuca Searle’owi przeoczenie faktu, że z warunkami *ceteris paribus* mogą łączyć się czynniki wartościujące. „Mówić, że pewien obowiązek *prima facie* nie został unieważniony przez żaden inny (...), to wypowiedzieć zdanie normatywne aczkolwiek negatywne”⁴².

Hintikka zauważa, że Searle ma rację w jednym miejscu: „wykazał, że z faktu i zasady analitycznej można prawomocnie wyprowadzić autentyczny obowiązek, mianowicie obowiązek *prima facie*. Z faktu nie wynika powinność, a tylko powinność *prima facie*. To spostrzeżenie nabiera w świetle naszych poprzednich spostrzeżeń tym większej wagi, że przecież nasze intuicje często dotyczą właśnie takiej powinności *prima facie*”⁴³.

Przeprowadzone rozważania ukazują, że zdaniem Hintikki mankament artykułu i rozumowania Searle’a nie polega na błędności przedstawionego argumentu, ale na nieukazaniu w jakim sensie należy rozumieć wyprowadzony wniosek⁴⁴. Sam zaś problem gilotyny Hume’a pozostaje tematem otwartym dla dalszej dyskusji.

⁴⁰ Por. tamże, s. 365.

⁴¹ Tamże, s. 366.

⁴² Tamże, s. 367.

⁴³ Tamże, s. 368-369.

⁴⁴ Por. tamże, s. 369.

Bibliografia:

1. Max Black, *The Gap Between „is” and „should”*, „The Philosophical Review”, 73(1964) s. 165-181.
2. Jaakko Hintikka, *Eseje logiczno-filozoficzne*, Warszawa 1992, s. 360-370.
3. David Hume, *Treatise on Human Nature*, III, 1,1; (wydanie polskie: D. Hume, *Traktat o naturze ludzkiej*, tłum. Cz. Znamierowski), Warszawa 2005, s. 531-547.
4. John Rogers Searle, *How to Derive Ought from Is*, „Philosophical Review” 73(1964), s. 43-58; (wydanie polskie: J. R. Searle, *Jak wywieść „powinien” z „jest”*, tłum. J. Hołówka, „Etyka” 16(1978), s. 163-177.

ks. Krzysztof Jaworski
Katolicki Uniwersytet Lubelski Jana Pawła II

Geneza pojęcia hiperzbioru.

Wstęp

Zbiór to jedno z najbardziej podstawowych pojęć matematyki. Niemal wszystkie obiekty współczesnej matematyki da się przedstawić w postaci zbiorów. Choć matematycy słowa „zbiór” używają chętnie, to nie dokonują oni nad tym słowem głębszej refleksji teoretycznej. Kiedy przychodzi moment, by zastanowić się nad tym, czym zbiór jest i jakie warunki musi spełniać dany twór, by można go było nazywać zbiorem, matematyka traci swe kompetencje, ustępując miejsca refleksji filozoficznej, zakorzenionej aż w metafizyce.

Zbiór według Georga Cantora

Za ojca teorii zbiorów (zwanej też teorią mnogości) uważa się niemieckiego matematyka Georga Cantora, żyjącego w latach 1845-1918. Z jego prac można wyczytać dwie definicje zbioru:

[Def. 1] zamieszczona w *Grundlagen einer allgemeinen Mannigfaltigkeitslehre* z 1883 roku:

„Unter einer «Mannigfaltigkeit» oder «Menge» verstehe ich nämlich allgemein jedes Viele, welches sich als Eines denken läßt, d.h. jeden Inbegriff bestimmter Elemente, welcher durch ein Gesetz zu einem Ganzen verbunden werden kann, und ich glaube hiermit etwas zu definieren, was verwandt ist mit dem Platonischen εἶδος oder ἰδέα, wie auch mit dem, was Platon in seinem Dialoge «Philebos oder das höchste Gut» μικτόν nennt”¹.

„Pod pojęciem «rozmaitości» czy «zbioru» rozumiem mianowicie ogólnie każdą wielość, która może być pomyślana jako jedność, tj. każdy ogół określonych elementów, które na mocy pewnego prawa mogą być złączone w jedną całość. Mam nadzieję, że definiuję w ten sposób coś, co jest spokrewnione z platońskim εἶδος czy ἰδέα, jak również z tym, co Platon w swym dialogu

¹ G. Cantor, *Grundlagen einer allgemeinen Mannigfaltigkeitslehre*, w: E. Zermelo (red.), G. Cantor. *Gesammelte Abhandlungen mathematischen und philosophischen Inhalts*, Verlag von Julius Springer, Berlin 1932 (reprint: Springer-Verlag, Berlin-Heidelberg-New York 1980) s. 204.

«Philebos albo najwyższe dobro» nazywa $\mu\kappa\tau\acute{o}\nu$ ”².

[Def. 2] pochodząca z *Beiträge zur Begründung der transfiniten Mengenlehre*, z roku 1895:

„Unter einer «Menge» verstehen wir jede Zusammenfassung M von bestimmten wohlunterschiedenen Objekten m unsrer Anschauung oder unseres Denkens (welche die «Elemente» von M genannt werden) zu einem Ganzen”³.

„Pod pojęciem «zbioru» rozumiemy każde zebranie w jedną całość M określonych, dobrze odróżnionych przedmiotów m naszego oglądu czy naszych myśli (które nazywane są «elementami» M)”⁴.

Jak widać, teoria mnogości rozważa zbiory w sensie dystrybutywnym, gdzie dany przedmiot x jest elementem zbioru A-ków wtedy i tylko wtedy, gdy x jest A-kiem. Należy zauważyć, że język potoczny podaje jeszcze jedno znaczenie słowa *zbiór*. Słowo to można rozumieć kolektywnie – jako zgromadzenie pewnych przedmiotów, gdzie x jest elementem zbioru A-ków wtedy i tylko wtedy, gdy x jest A-kiem lub x jest częścią A-ka. Refleksją nad zbiorami w sensie kolektywnym zajmuje się dziedzina nauki stworzona przez Stanisława Leśniewskiego – mereologia.

Metafory związane z pojęciem zbioru

W roku 1991 Jon Barwise i Larry S. Moss wydali krótki artykuł zatytułowany *Hypersets*, w którym, podążając drogą wyznaczoną wcześniej przez Petera Aczela, starali się opisać zjawisko tzw. hiperzbiorów, czyli zbiorów nieufundowanych (lub też, inaczej mówiąc, zbiorów nie-dobrze-ufundowanych). Wraz z książką P. Aczela *Non-Well-Founded Sets* artykuł Barwise’a i Moss’a uważa się za źródłowy, gdy chodzi o współczesną teorię hiperzbiorów. W artykule autorzy podają dwie metafory, które mają pomóc zrozumieć, czym jest hiperzbiór, demaskując jednocześnie naiwne i, być może, nieco skostniałe rozumienie pojęcia zbioru⁵.

Zbiór jako pudełko

Pierwszą myślą, która zwykle przychodzi do głowy człowiekowi, zadającemu sobie pytanie o istotę zbioru, jest namalowane na kartce papieru kółko oraz pewne umowne znaki w obrębie tego kółka. Jest to najprostszy chyba sposób wyobrażenia sobie zbioru, choć nie jest to sposób wolny od pewnych ograniczeń. Takie myślenie o zbiorach każe nam je traktować *pudełkowo*, tzn. wyobrażać sobie zbiór jako pudełko z przedmiotami w środku. Tworzenie takiego zbioru można by porównać z wkładaniem określonych przedmiotów do pudełka (tymi przedmiotami mogłyby

² Tłum. za R. Murawski, *Filozofia matematyki. Antologia tekstów klasycznych*, UAM, Poznań 2003, s. 175.

³ G. Cantor, *Beiträge zur Begründung der transfiniten Mengenlehre*, w: E. Zermelo (red.), G. Cantor. *Gesammelte Abhandlungen...*, dz. cyt., s. 282.

⁴ Tłum. za R. Murawski, *Filozofia matematyki. Antologia tekstów klasycznych*, dz. cyt., s. 176.

⁵ J. Barwise, L.S. Moss, *Hypersets*, “Mathematical Intelligencer”, 4(13) 1991, s. 36.

być także inne pudełka). Metafora zbioru-pudełka, choć intuicyjnie efektywna, obciążona jest jednak kilkoma trudnościami⁶:

- (1) Zbiór pusty byłby symbolizowany przez pudełko bez żadnego przedmiotu w środku. Możemy sobie wyobrazić nieskończenie wiele różnych pustych pudełek, ale zbiór pusty, jak wiemy, jest tylko jeden.
- (2) Jest fizycznie niemożliwe, aby przedmiot b znajdował się w kilku pudełkach jednocześnie w następujący sposób: $A = \{\{a, b\}, \{b, c\}\}$, tymczasem zbiór A określony jest w języku teorii mnogości ZFC poprawnie.
- (3) Aksjomat ekstensjonalności głosi, że zbiory, zawierające te same elementy, są równe:

$$\forall x(x \in A \equiv x \in B) \equiv A = B$$

Jeśli metaforę pudełka potraktować dosłownie, to żadne dwa zbiory nie byłyby sobie nigdy równe, bo, jak zaznaczyliśmy w (2), ten sam przedmiot nigdy nie może znajdować się w dwóch miejscach (pudełkach) jednocześnie. Wyjątkiem nie mogłyby być nawet dwa pudełka puste (posiadające te same elementy, czyli nie posiadające żadnego), ponieważ, co prawda, dwa puste pudełka mogą istnieć obok siebie, jednak nie istnieją dwa zbiory puste, co zaznaczyliśmy w (1).

- (4) W metaforze pudełkowej zbiór jest konstytuowany przez swoje elementy. Znaczy to, że aby zaistniał zbiór, muszą istnieć uprzednio jego elementy. Elementy są czymś pierwotnym w stosunku do zbioru, tzn. dopiero mając jakby *gotowe* elementy możemy utworzyć z nich zbiór. Jak zobaczymy niżej, teoria hiperzbiorów *łamie* czasem tę zasadę, przyznając pierwszeństwo istnienia zbiorowi, a nie jego elementom. Dopuszcza się tam bowiem zbiory typu $A = \{I, A\}$, które, jak widać, trudno skonstruować z komponentów uprzednio istniejących.

Zbiór jako struktura

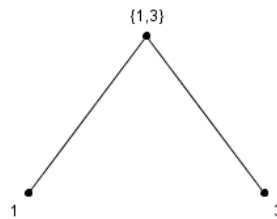
Inny sposób przedstawiania zbiorów jest związany ze wskazaniem na ich strukturę. Barwise i Moss w przytaczanym artykule pokazują, jak wyodrębnić taką strukturę: *Weźmy ustrukturyzowany obiekt dowolnego rodzaju, obiekt, zawierający «składniki», a następnie zapomnijmy poszczególne składniki i swego rodzaju «klej», który trzyma je razem. To, co nam pozostało, jest właśnie abstrakcyjną teoriomnościową strukturą*⁷. Każdy zbiór jest tworem o określonej strukturze. Strukturę tę wyznacza relacja należenia $\in$. Atomy, czyli elementy niebędące zbiorami, to podstawowy budulec owej struktury, natomiast zbiory, nabudowane na atomach, to kolejne piętra hierarchii. Ze zbiorów, których elementami są inne zbiory, można budować kolejne zbiory jeszcze bardziej złożone itd. Utworzoną w ten sposób strukturę można przedstawić za pomocą matematycznego narzędzia, zwanego grafem.

Dla wyjaśnienia tej idei spróbujmy utworzyć graf, będący obrazem zbioru $A = \{\{1, 3\}, 4\}$. Zbiór A posiada dwa elementy: $\{1, 3\}$, 4 . Pierwszy z tych elementów także jest zbiorem, którego elementami są: 1 , 3 , tak więc podstawowymi składnikami zbioru A są atomy: 1 , 3 i 4 . Odwracając ten proces, w następujących krokach tworzymy graf obrazujący zbiór wyjściowy $A = \{\{1, 3\}, 4\}$:

⁶Por. J. Paśniczek, *Filozoficzne znaczenie hiperzbiorów*, w: J. Perzanowski, A. Pietruszczak, C. Gorzka (red.) *Filozofia/Logika. Filozofia logiczna 1994*, UMK, Toruń 1995, s. 52-55.

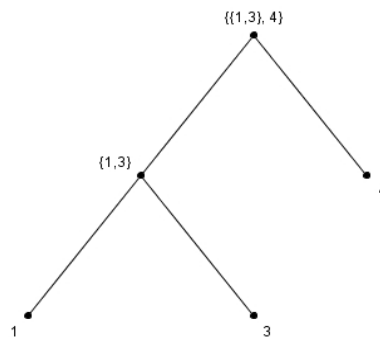
⁷Por. J. Barwise, L. Moss, dz. cyt., s. 36.

-
- (1) Rozpoczynając od atomów 1 i 3 tworzymy najpierw graf przedstawiający zbiór, którego elementami są te atomy:



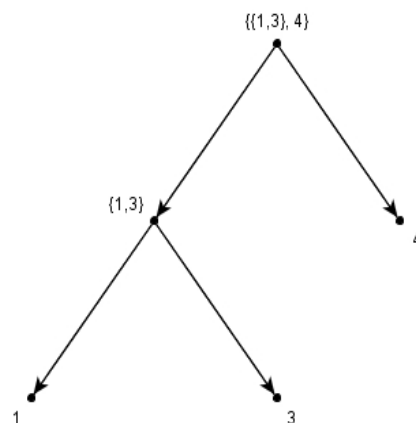
Graf 1

- (4) Posługując się Grafem 1 i atomem 4 tworzymy pożądany graf:



Graf 2

Jak widać, metoda tworzenia grafów obrazujących zbiory jest dość prosta: jeżeli dwa wierzchołki grafu połączone są krawędzią, to do wierzchołka znajdującego się *poniżej* przyporządkowujemy dokładnie jeden obiekt, będący elementem zbioru odpowiadającego wierzchołkowi *powyżej*. Aby uprościć sprawę, Aczel wprowadza do swej teorii grafy skierowane. Idąc tropem Aczela, Graf 2 można by przekształcić w następujący sposób:



Graf 3

[Def. 3] Graf skierowany G jest to taki graf, który składa się z niepustego zbioru $V(G)$ wierzchołków, zbioru $E(G)$ krawędzi oraz z funkcji γ przekształcającej zbiór $E(G)$ w zbiór $V(G) \times V(G)$. Jeżeli e jest krawędzią grafu G oraz $\gamma(e) = (p, p')$, to p nazywamy początkiem krawędzi e , a p' jej końcem. Mówimy wówczas, że krawędź biegnie od p do p' . Na rysunku kierunek krawędzi oznaczamy strzałką⁸:



Rysunek 1

W teorii Aczela jeśli strzałka biegnie od wierzchołka p do wierzchołka p' , to wierzchołek p nazywa się rodzicem (*parent*) wierzchołka p' , natomiast wierzchołek p' nazywa się dzieckiem (*child*) wierzchołka p . To, że od wierzchołka p do wierzchołka p' biegnie krawędź skierowana, zapisuje się w następujący sposób:

$$p \rightarrow p'$$

[Def. 4] Droga jest to skończona lub nieskończona sekwencja $p_0 \rightarrow p_1 \rightarrow p_2 \rightarrow \dots$ węzłów $p_0, p_1, p_2, \dots$ połączonych krawędziami $(p_0, p_1), (p_1, p_2), \dots$

[Def. 5] Graf punktowy jest to taki graf, który posiada wyróżniony węzeł zwany punktem. Graf punktowy nazywamy dostępnym wtedy, gdy do każdego węzła p istnieje droga $p_0 \rightarrow p_1 \rightarrow \dots \rightarrow p$ od punktu p_0 do węzła p .

Metafora struktury nawiązuje do tzw. funktora zapominania, występującego w teorii kategorii. Chodzi o to, aby myśleć o zbiorach jak o tworach, które powstają w procesie abstrakcji z grafów skierowanych w taki sposób, że jakby “zapomina się” o naturze samych wierzchołków, przypisując do każdego wierzchołka n zbiór $d(n)$ taki, że:

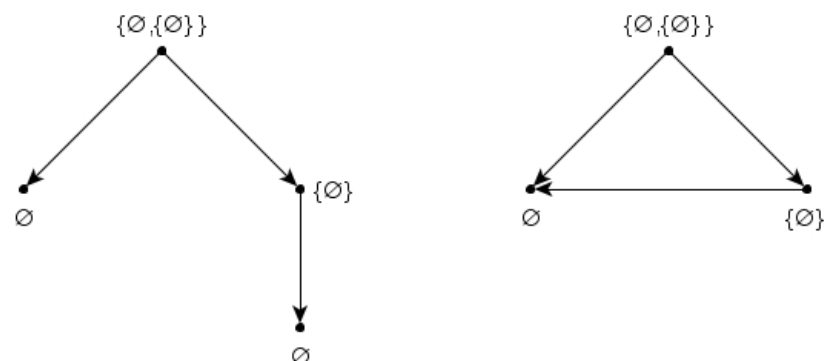
$$d(n) = \{d(m) \mid n \rightarrow m\}$$

Zbiór w takim rozumieniu to pewien rodzaj struktury, która może być zobrazowana przez dostępny graf punktowy (*accessible pointed graph* – w skrócie *apg*), gdzie funkcja oznaczania d pomiędzy wierzchołkiem grafu, a zbiorem, który ten wierzchołek reprezentuje, musi spełniać powyższy warunek⁹.

W tym kontekście warto zauważyć, że ten sam zbiór może być przedstawiony za pomocą kilku różnych grafów. Jako przykład weźmy zbiór $\{\emptyset, \{\emptyset\}\}$. Obrazem tego zbioru mogą być następujące grafy:

⁸ Zob. K. A. Ross, Ch. R. B. Wright, *Matematyka dyskretna*, PWN, Warszawa 2005, s. 142-151.

⁹ Por. J. Barwise, L.S. Moss, dz. cyt., s. 37.



Rysunek 2

Każdy zbiór można przedstawić za pomocą grafu. Postawmy jednak pytanie odwrotne: czy każdy graf obrazuje jakiś zbiór? Jeżeli poruszamy się w ramach standardowej teorii mnogości ZFC, to wystarczy wziąć pod uwagę np. Graf 4 (poniżej), aby negatywnie odpowiedzieć na to pytanie. Graf ten przedstawia krawędź, która wychodzi z jednego wierzchołka i biegnie z powrotem do tego wierzchołka, tworząc pętlę. W ten sposób powstaje nieskończona droga, która sprawia, że podążając za kierunkiem strzałki nigdy nie dojdzie się do elementu ostatecznego (atomu). Taki graf byłby obrazem zbioru nieufundowanego, tymczasem teoria ZFC, na mocy aksjomatu ufundowania, nie przewiduje istnienia zbiorów nieufundowanych¹⁰.



Graf 4

Aczel, natomiast, wypowiada twierdzenie, że każdy graf może być obrazem jakiegoś zbioru¹¹. Skoro zatem każdy graf jest obrazem zbioru, to również graf nieufundowany (czyli taki, który posiada nieskończoną drogę) jest obrazem zbioru. Wiemy jednak, że graf nieufundowany może być obrazem tylko zbioru nieufundowanego. Stąd wniosek, że teoria Aczela dopuszcza istnienie zbiorów nieufundowanych.

¹⁰Aksjomat ufundowania głosi wprost, że wszystkie zbiory w ramach teorii mnogości ZFC są dobrze ufundowane.

¹¹P. Aczel, *Non-Well-Founded Sets*, CSLI Lecture Notes, Stanford 1988, s. 5.

Charakterystyka hiperzbiorów

Aby pokazać, jakimi cechami charakteryzują się hiperzbiory, najpierw przywołamy najprostsze zjawisko, występujące w matematyce, gdzie bez skrępowania używa się pojęć opartych na zbiorach nieufundowanych, a następnie rozważymy kilka definicji hiperzbioru.

Ułamki nieskończone

Najprostszym przykładem samozwrotności, przywołanym z sentymentem przez Barwise'a w artykule *Hypersets*¹², jest ułamek o postaci:

$$x = 1 + \frac{1}{1 + \frac{1}{1 + \dots}}$$

Wspominając swojego nauczyciela matematyki, który mu pokazał ten ułamek, autor artykułu powiada: *Z jednej strony był wielokropek, zstępujący jak szlaczek Lucyfera coraz głębiej ku wnętrzu demonicznego mianownika. Z drugiej strony, jednak, było światło i klarowność rozwiązania [powyższego] równania, uzyskane dzięki obserwacji, że ułamek spełnia równanie:*

$$x = 1 + \frac{1}{x}$$

Rozwiązanie zaś takiego równania nie sprawiłoby już większej trudności nawet przeciętnemu uczniowi szkoły podstawowej.

Matematyk, zatem, widząc cyrkularność nie dostrzega żadnej sprzeczności. Matematyka zna bardzo wiele definicji cyrkularnych. Przykładem może tu być choćby rekurencyjna definicja silni, używana w matematyce bez żadnych formalnych zastrzeżeń:

$$n! = \begin{cases} 1 & \text{dla } n = 0 \\ n(n-1)! & \text{dla } n \geq 1 \end{cases}$$

Skąd zatem ta niechęć do cyrkularności w teorii mnogości? Czyżby odpowiedzialne za nią były tylko teoriomnogościowe antynomie, a remedium na antynomie, czyli m.in. zakaz używania formuł typu $x \in x$, okazał się wylaniem dziecka z kąpielą?

Hiperzbiór jako granica ciągu zbiorów dobrze ufundowanych

Niech A będzie następującym zbiorem:

$$A = \{I, \{A\}\}$$

Widać wyraźnie, że tak określony zbiór A jest hiperzbiorem, ponieważ jego postać można zapisać w następujący sposób:

$$A = \{I, \{I, \{I, \{I, \dots\}\}\}\}$$

¹² Zob. J. Barwise, L.S. Moss, dz. cyt., s. 31.

Wielokropek, który się pojawia w powyższym zapisie, wskazuje na to, że istnieje tzw. nieskończone zejście.

Niech b będzie ciągiem, którego elementami są następujące zbiory ufundowane:

$$b = (\{1\}, \{1, \{1\}\}, \{1, \{1, \{1\}\}\}, \{1, \{1, \{1, \{1\}\}\}\}, \dots)$$

Można powiedzieć, że nieufundowany zbiór A jest pewnym rodzajem granicy ciągu b , którego elementami są zbiory ufundowane.

$$\lim_{n \rightarrow \infty} b_n = A$$

Zauważmy, że im dany element ciągu b jest dalej od elementu pierwszego, tym trudniej jest go odróżnić od zbioru A .

Definicje hiperzbioru

W literaturze można spotkać wiele nazw, będących synonimami słowa hiperzbiór: zbiór niedobrze ufundowany (non-well-founded set), zbiór nieufundowany (słowo to pochodzi od aksjomatu antyufundowania), zbiór nienormalny. Słowo *hiperzbiór* jest dosłownym tłumaczeniem angielskiego wyrazu *hyperset*. Przedrostek *hyper* (*nad*), pochodzący z języka greckiego, oznacza pewien nadmiar, przesadę. Hiperzbiór, zatem, to swego rodzaju *nad-zbiór*, zbiór *przesadny*, zbiór nieokreślonej wielkości, jak naczynie bez dna. Co jest powodem tej głębi, w którą wciąga nas świat hiperzbiorów? Jest nim pewien pomysł, by spróbować wyzwolić się z *pudełkowej* koncepcji zbioru, i rozważyć istnienie na przykład takich zbiorów, których elementami są one same. Na pierwszy rzut oka klóci się to ze zdrowym rozsądkiem, jednak okazuje się, że świat pełen jest zjawisk o charakterze cyrkularnym, dających się opisać w języku teoriomnogościowym, który ową cyrkularność *uniewinnia*. Słowo *uniewinnić* zostało tu użyte celowo, bo to przecież cyrkularność, czy samozwrotność, była przyczyną zawału, który wstrząsnął podstawami matematyki w początkach XX wieku (chodzi o słynne odkrycie antynomii teoriomnogościowych).

Istnieje kilka definicji hiperzbioru, przy czym każda wskazuje na inny aspekt tego teoriomnogościowego pojęcia.

[Def. 6] Zbiór A nazywamy hiperzbiorem wtedy i tylko wtedy, gdy istnieje nieskończona sekwencja:

$$\dots \in B_{n+1} \in B_n \dots \in B_1 \in A$$

W przeciwnym razie zbiór A nazwiemy dobrze ufundowanym.

Na potrzeby kolejnej definicji hiperzbioru, trzeba najpierw podać definicję tzw. składnika zbioru:

[Def. 7] Zbiór X jest składnikiem zbioru Z (co zapisujemy $X \in^* Z$), jeśli X jest elementem zbioru Z lub elementem elementu zbioru Z lub elementem elementu elementu zbioru Z itd.¹³.

[Def. 8] Zbiór A nazywamy hiperzbiorem wtedy i tylko wtedy, gdy istnieje sekwencja $\langle B_n : n = 1, 2, \dots \rangle$ w następującej relacji bycia składnikiem $\in^*$:

$$\langle \dots \in^* B_{n+1} \in^* B_n \dots \in^* B_1 \in^* A \rangle$$

Kolejna definicja pochodzi od rosyjskiego matematyka Dmitrija Mirimanoffa. Definiuje on tzw. *zbiory normalne* i *zbiory nienormalne*:

[Def. 9] Niech A będzie zbiorem, A' jednym z jego elementów, A'' elementem

¹³ J. Barwise, L. Moss, dz. cyt., s. 34.

zbioru A' itd. Zejściem nazywamy sekwencję kroków od A do A' , od A' do A'' itd. Powiemy, że zbiór jest *normalny*, jeśli jego zejście jest skończone, natomiast zbiór jest *nienormalny*, jeśli pośród jego zejść istnieją takie, które są nieskończone¹⁴.

Peter Aczel wstęp do swojej książki *Non-Well-Founded Sets* rozpoczyna od zdefiniowania zbioru nieufundowanego. Czyni to używając właśnie powyższej definicji Mirimanoffa, stwierdzając krótko:

[Def. 10] Hiperzbiorem nazywamy zbiór nienormalny w sensie Mirimanoffa.

Inaczej mówiąc, hiperzbiór to taki zbiór, który posiada nieskończone zejście.

[Def. 11] Niepusty zbiór A jest nieufundowany wtedy i tylko wtedy, gdy nie posiada elementu minimalnego w relacji ϵ^* .

Zakończenie

Wydaje się, że kluczem, który otwiera drzwi do zrozumienia całej teorii mnogości, jest słowo *struktura*. Jakiegokolwiek systemu teoriomnogościowego byśmy nie rozważali, uniwersum obiektów jest tam zawsze jakoś ustrukturyzowane. To właśnie formalne doprecyzowanie struktury zbiorów pozwoliło dość skutecznie wyeliminować antynomie, które powstały na gruncie teorii mnogości. Ernst Zermelo, tworząc dla teorii mnogości pierwszy system aksjomatyczny, sformułował pojęcie hierarchii kumulatywnej, która jest niczym innym, jak właśnie ściśle uporządkowanym uniwersum zbiorów, w którym prawdziwe są aksjomaty. Ową hierarchiczność widać wyraźnie także w (konkurującej z teorią mnogości) teorii typów, stworzonej przez Bertranda Russella. Z kolei istotą dokonań Alfreda Tarskiego, jeśli chodzi o zagadnienie definicji prawdy, było odkrycie hierarchicznego charakteru języka.

Tam, gdzie jest hierarchia, tam jest i pewien porządek. Logika, będąca narzędziem do porządkowania świata, sama musiała najpierw zostać uporządkowana. Okazuje się, jednak, że owo uporządkowanie, które dokonało się na początku dwudziestego wieku, jest niewystarczające, ponieważ wyeliminowało ono z pola rozważań zjawiska cyrkularne (albo też nakazało zjawiska cyrkularne traktować z największą pogardą, jako coś nieprawidłowego, wręcz grzesznego). Tymczasem zjawiska takie są naturalnym elementem naszego świata: występują w filozofii, w matematyce, w naukach informatycznych, w języku. One wcale nie muszą być czymś błędnym – potrzebują tylko solidnego fundamentu teoretycznego. Wydaje się, że teoria hiperzbiorów, jako rozszerzenie teorii ZFC, powstała właśnie po to, by zjawiska cyrkularne opisać, zrehabilitować i przyznać im zaszczytne miejsce pośród tych naukowych zagadnień, o których warto mówić pozytywnie.

¹⁴Zob. D. Mirimanoff, *Les antinomies de Russell et de Burali-Forti et le probleme fondamental de la theorie des ensembles*, "L'enseignement mathematique", 19 (1917) s. 42.

Bibliografia:

1. Aczel P., *Non-Well-Founded Sets*, CSLI Lecture Notes, Stanford 1988.
2. Barwise J., Moss L.S., *Hypersets*, "Mathematical Intelligencer", 4(13) 1991, s. 31-41.
3. Cantor G., *Beiträge zur Begründung der transfiniten Mengenlehre*, w: Zermelo E. (red.) Georg Cantor. *Gesammelte Abhandlungen mathematischen und philosophischen Inhalts*, Verlag von Julius Springer, Berlin 1932 (reprint: Springer-Verlag, Berlin–Heidelberg–New York 1980) s. 282-356.
4. Cantor G., *Grundlagen einer allgemeinen Mannigfaltigkeitslehre*, w: Zermelo E. (red.) Georg Cantor. *Gesammelte Abhandlungen mathematischen und philosophischen Inhalts*, Verlag von Julius Springer, Berlin 1932 (reprint: Springer-Verlag, Berlin–Heidelberg–New York 1980) s. 165-209.
5. Mirimanoff D., *Les antinomies de Russell et de Burali-Forti et le probleme fondamental de la theorie des ensembles*, "L'enseignement mathematique", 19 (1917) s. 37-52.
6. Murawski R., *Filozofia matematyki. Antologia tekstów klasycznych*, UAM, Poznań 2003.
7. Pańniczek J., *Filozoficzne znaczenie hiperzbiorów*, w: Perzanowski J., Pietruszczak A., Gorzka C. (red.) *Filozofia/Logika. Filozofia logiczna 1994*, UMK, Toruń 1995, s. 49-65.
8. Ross K.A., Wright Ch.R.B., *Matematyka dyskretna*, PWN, Warszawa 2005.

Paweł Zarosa
Katolicki Uniwersytet Lubelski Jana Pawła II

***Futura contingentia* a logika trójwartościowa.**

Na początku XX w. Jan Łukasiewicz, wybitny profesor Uniwersytetu Warszawskiego, zbudował pierwszą logikę trójwartościową. Podjął się tego zadania, zainspirowany między innymi rozważaniami Arystotelesa na temat przyszłych zdarzeń przygodnych. W swoim artykule *O determinizmie*, będącym przeróbką mowy rektorskiej, z inauguracji roku akademickiego 1922/1923, Łukasiewicz kreśli filozoficzne podstawy pod swój system. Główny argument, jakim się posługuje dla uzasadnienia wprowadzenia do rachunku logicznego trzeciej wartości, ma tę mniej więcej treść: zastosowanie zasady dwuwartościowości do przyszłych zdarzeń przygodnych (*futura contingentia*) jest równoznaczne z opowiedzeniem się za ontologicznym determinizmem, zgodnie z którym przyczyny każdego faktu istnieją odwiecznie.

Problem pogodzenia przyszłych zdarzeń przygodnych z klasycznie pojętym wartościowaniem logicznym zdań jako pierwszy postawił Arystoteles w *Hermeneutyce*. Zastanawia się on nad prawdziwością *zdań szczegółowych i dotyczących przyszłości*¹. Tylko ten konkretny typ zdań wedle Arystotelesa nastręcza problemy, gdy zechce się stosować zawsze i wszędzie zasadę dwuwartościowości, inne nie powodują trudności. Jak pisze: *Co się tyczy tego, co jest i co było, twierdzenia lub przeczenia muszą być prawdziwe lub fałszywe*².

Jak natomiast mają się sprawy zdań mówiących o niekoniecznej przyszłości? Filozof tak prowadzi rozumowanie:

*„Jeżeli każde twierdzenie czy przeczenie jest bądź prawdziwe bądź fałszywe, to każdy orzecznik musi albo przysługiwać, albo nie przysługiwać swemu podmiotowi. Tak że jeżeli ktoś mówi, że coś się zdarzy, a ktoś inny zaprzeczy temu, to jasne, że jeden z nich musi mówić prawdę”*³.

Nietrudno zauważyć poważny kłopot, który powstaje, gdy do tego typu zdań zastosujemy zasadę dwuwartościowości i to kłopot metafizyczny, nie jedynie logiczny. Mianowicie problem zdeterminowania przyszłości. Bo jeśli w jakimś punkcie czasu istnieje zbiór wszystkich zdań prawdziwych (a z tego, co powiedziano, wynikałoby nawet, że jest tak w każdym punkcie czasu, skoro zasada dwuwartościowości ma obowiązywać zawsze), mówiących o każdym z nieskończenie wielu przyszłych zdarzeń, to przyszłość jest już wtedy określona, a wszystko, co się dzieje, dzieje się

¹ Arystoteles *Hermeneutyka*, 18a [w:] tegoż, *Dzieła wszystkie*, t. 1.; tłum. Kazimierz Leśniak; PWN, Warszawa 1990; str. 75.

² Tamże.

³ Tamże.

z konieczności i inaczej działać się nie może. Stagiryta jest tego świadomy, pisząc:

„A zatem nic nie jest i nic się nie zdarza ani przypadkiem, ani wskutek ślepego trafu, ani się zdarzy czy nie zdarzy jak tylko z konieczności”⁴,

i dalej:

„Jeżeli coś jest białe teraz, to można było prawdziwie twierdzić wcześniej, że będzie białe; tak że można było prawdziwie twierdzić o czymś, co się zdarzyło, że jest lub że będzie. A skoro można było zgodnie z prawdą twierdzić, że jest lub że będzie, to nie może być, żeby to nie było czy nie mogło być. Bo jeżeli coś nie mogło się nie zdarzyć, to niemożliwe, żeby się nie zdarzyło; a to, co nie może się nie zdarzyć, musi się zdarzyć. A więc wszystko, co ma być, musi się zdarzyć”⁵.

Trudności powstają, jeśli się przyjmie za konieczne dla każdego twierdzenia i przeczenia (...), że jedno ze zdań przeciwnych jest prawdziwe a drugie fałszywe i że nic, co się zdarzyło nie jest przypadkowe, ale że wszystko jest i powstaje z konieczności⁶. Jednak próba przewyciężenia wywnioskowanego powyżej determinizmu, na który Arystoteles wyraźnie się nie zgadza⁷ (w niejednym swym dziele), poprzez odmówienie takim zdaniom wartości logicznej także prowadzi do absurdu, co zauważa Filozof: *za przykład niech posłuży bitwa morska: ani by się jutro odbyła, ani nie odbyła*⁸. Z jednej strony: determinizm, z drugiej: złamana zasada niesprzeczności. Oba stanowiska obala doświadczenie. Jakie rozwiązanie można proponować?

Arystoteles nie rozwikłał zagadnienia do końca, miał pewne intuicje i jak się wydaje, mógłby się tu skłaniać w kierunku propozycji Łukasiewicza. W tym miejscu chcielibyśmy zwrócić uwagę na postawione w *Hermeneutyce* rozróżnienie rozumienia konieczności, które może być pomocne dla całego rozumowania. Konieczność może być pojęta na dwa sposoby: albo jako rzeczywiste istnienie danego stanu rzeczy, albo jako determinacja stanów przyszłych. Tak więc to, co jest, jest faktem, po prostu jest. Rzeczywistość nie daje pola dywagacjom, jest tak, jak jest. Co nie oznacza bynajmniej konieczności w drugim sensie, a więc że wcześniej już przyszłość była ustalona⁹. Poza tym, co konieczne, mamy do czynienia z mnóstwem faktów, które nie musiały się wydarzyć, ale zdarzyły się. Nie posiadały więc tego wcześniejszego zdeterminowania, ale stały się rzeczywistością i są, a skoro tak, to

⁴Tamże (18b).

⁵Tamże, str. 76.

⁶Tamże.

⁷Jak pisze niżej: *wiemy z własnego doświadczenia, że to, co ma zdarzyć, ma swój początek zarówno w zamierzeniach (a więc w tym, czego faktycznie nie ma – przyp. P. Z.), jak i w działaniu (w czymś dynamicznym i niezupełnie określonym – przyp. P. Z.) i że w ogóle w tym, co nie zawsze się urzeczywistnia, tkwi możliwość zarówno dojścia, jak i nie dojścia do skutku (...).* Jasne więc, że nie wszystko jest czy powstaje z konieczności; pewne fakty zdarzają się przypadkiem (...) inne natomiast zdradzają raczej ogólną tendencję w jednym bądź drugim kierunku, ale też mogą i w innym jeszcze dążyć; tamże, 19a, str. 76-77.

⁸Tamże, 18b, str. 76. *I tak jak nie do przyjęcia wydaje się luka, tak samo też i kolizja prawdziwościowa, bo wspomniana bitwa tak by się odbyła, jak i nie odbyła, co sprowadza się do tego samego.*

⁹To więc, co jest, jest konieczne, skoro jest, a tego, czego nie ma, konieczne, skoro nie ma. Ale nie wszystko, co jest, konieczne jest, i nie wszystkiego, czego nie ma, konieczne nie ma. Bo powiedzieć, że wszystko, co jest, jest konieczne, skoro jest, to nie to samo, co powiedzieć po prostu, że to jest z konieczności. Podobnie z tym, czego nie ma. Ta sama argumentacja odnosi się do zdań sprzecznych. Wszystko musi być albo nie być, i będzie albo nie będzie, ale nie zawsze można odróżnić i stwierdzić, który z tych członów jest konieczny twierdząc na przykład, że jutro odbędzie się bitwa morska, albo się nie odbędzie, ale nie jest konieczne, ażeby jutro odbyła się bitwa morska, ani też nie jest konieczne, żeby się jutro nie odbyła, chociaż jest konieczne, ażeby się bądź odbyła, bądź nie odbyła. Skoro więc prawdziwość zdań polega na zgodności z faktami, to jest też jasne, że jeżeli przyszłe zdarzenia są alternatywne i mogą się zdarzyć sytuacje przeciwne, to taki sam charakter będą miały zdarzenia sprzeczne. Przypadek ten dotyczy tego co nie zawsze istnieje albo nie zawsze nie istnieje. Bo jeden człon tej sprzeczności musi być prawdziwy, a drugi fałszywy. Ale nie możemy zdecydowanie stwierdzić, że ten czy ten jest fałszywy, lecz alternatywa musi pozostać nierozstrzygnięta. Jeden z członów może być bardziej prawdopodobny niż drugi, ale ni może już być prawdziwy lub fałszywy. Jasne więc, że nie jest konieczne, ażeby każde twierdzenie czy przeczenie mu przeciwne musiało być jedno prawdziwe, a drugi fałszywe. Tamże 19a-19b, str. 76

już nie może ich nie być.

Konieczność ma związek z aktem, podczas gdy niezdeterminowanie – z możliwością. Jeśli coś się już zaktualizowało albo jest pewnym aktem-formą, to jest faktem. Te metafizyczne rozważania prowadzą wprost na teren logiki: do zagadnienia prawdy. Nie tylko Arystoteles twierdził, że *prawdziwość zdań polega na zgodności z faktami*¹⁰. A teraz z tego, co powiedziano, można wyprowadzić kilka wniosków. Z pewnością prawdę trzeba przypisać zdaniom, opisującym fakty konieczne – w obydwu znaczeniach, fakty, można powiedzieć: aktualne. Te natomiast, które odnoszą się do nieaktualnych można (1) uznać za fałsz – jako, że nie istnieje to, o czym mówią – (2) bądź też przypisać im trzecią wartość – ponieważ nie istnieje rzeczywistość odpowiadająca sprzecznym z nimi zdaniom, a taka właśnie byłaby odpowiednikiem fałszu. Kolejnym problemem jest ocena, czy tak samo należy traktować fakty, które jeszcze nie nadeszły, jaki i te, które już przeminęły i czy kryterium do oceny ich aktualności ma być teraźniejsze istnienie, czy też prowadzące do nich przyczyny oraz pozostałe po nich skutki. Czy w ogóle i jaką rolę odgrywa tu czas? Możliwe jest też (3) przyznanie prawdy sądom mówiącym o tym, co jest, było lub będzie o ile to faktycznie jest, było czy będzie bez względu na relacje w czasie pomiędzy wypowiedzianym zdaniem, a zdarzeniem, do którego się odnosi. Prawda jest w tej koncepcji wieczna i odwieczna.

Łukasiewicz pierwszego rozwiązania nie rozważał, ale można sądzić, że odrzuciłby je. Przyjął drugie, a sprzeciwił się ostro trzeciemu, ponieważ, jak już wspomniano na początku, uważał je za uznające determinizm, który definiował następująco:

„Pogląd, który głosi, że jeśli A jest b w chwili t , to prawdą jest w każdej chwili wcześniejszej od t , że A jest b w chwili t ”¹¹.

Łukasiewicz rozumie też oczywiście determinizm jako pogląd ontologiczny, mówiący, iż wszelkie zdarzenia są z góry ustalone, czemu daje wyraz niżej:

„Determinista spogląda na fakty, dziejące się w świecie, jak na potworny dramat filmowy, sporządzony gdzieś we wszechświatowej wytwórni. Jesteśmy w środku przedstawienia, i choć każdy z nas nie jest tylko widzem, lecz i aktorem dramatu, to jednak końca filmu nie znamy. Ale ten koniec jest, istnieje od początku przedstawienia, bo cały film jest od wieków gotowy. W filmie tym ustalone są z góry wszystkie role nasze, wszystkie nasze przygody i koleje życiowe, wszystkie decyzje nasze, i złe, i dobre czyny, i przewidziany jest moment śmierci, Twojej i mojej. W dramacie wszechświatowym odgrywamy tylko rolę marionetek”¹².

Według niego te dwa ujęcia to w rzeczywistości jeden i ten sam determinizm, a to oznacza, że twierdzenie, iż wszystkie zdania mówiące o przyszłości są prawdziwe albo fałszywe jest równoznaczne z twierdzeniem, iż przyszłość jest z góry przewidziana. Dalej przedstawił bardzo uporządkowane, ścisłe dowodzenie, któremu warto się przyjrzeć. Rozważane jest zdanie postaci: „ A jest b w chwili t ”, wyrażone na sposób: „W dowolnej chwili prawdą jest, że A jest b w chwili t ”¹³. W konkretnym przykładzie A to Jan, b oznacza bytność w domu, chwila t jest zaś

¹⁰ Tamże.

¹¹ Łukasiewicz J. *O determinizmie* [w:] *Z zagadnień logiki i filozofii. Pisma wybrane*; PWN, Warszawa 1961; str. 116

¹² Tamże.

¹³ Co naszym zdaniem niepotrzebnie gmatwa sprawę i może w tym zagmatwaniu kryje się błąd. Przyjrzymy się temu poniżej.

jutrzejszym południem (godziną dwunastą konkretnego dnia, zatem bezwzględnie określonym punktem w czasie): *W dowolnej chwili prawdą jest, że Jan będzie jutro w południe w domu*¹⁴. Przesłanki, które *przyjmujemy bez dowodu jako intuicyjnie pewne*¹⁵ przedstawiają się następująco (1 i 3):

1. *Prawdą jest w chwili t ¹⁶, że Jan będzie jutro w południe w domu, lub prawdą jest w chwili t , że Jan nie będzie jutro w południe w domu*¹⁷.

Z alternatywy: $p \vee q$ wynika implikacja: $\neg p \rightarrow q$, zatem przekształcamy przesłankę (1) w:

2. *Jeśli nie jest prawdą w chwili t , że Jan nie będzie jutro w południe w domu, to prawdą jest w chwili t , że Jan będzie jutro w południe w domu*¹⁸.
3. *Jeśli prawdą jest w chwili t , że Jan nie będzie jutro w południe w domu, to Jan nie będzie jutro w południe w domu*¹⁹.

Na mocy prawa transpozycji prostej: $(p \rightarrow q) \rightarrow (\neg q \rightarrow \neg p)$ przekształcamy przesłankę (3) w:

4. *Jeśli Jan będzie jutro w południe w domu, to nie jest prawdą w chwili t , że Jan nie będzie jutro w południe w domu*²⁰.

Zestawiając ze sobą dwie implikacje – przesłanki (4) i (2) zauważamy, że następnik pierwszej (4) jest poprzednikiem drugiej (2). Ponieważ z przesłanek: $p \rightarrow q$ oraz $q \rightarrow r$ wynika logicznie wniosek: $p \rightarrow r$, otrzymujemy wynik:

5. *Jeśli Jan będzie jutro w południe w domu, to prawdą jest w chwili t , że Jan będzie jutro w południe w domu*²¹.

Łukasiewicz twierdzi, że owa konkluzja wyprowadzona z przyjętych intuicyjnie przesłanek sprowadza się w istocie do podstawowej tezy determinizmu:

„Chwila t jest dowolną chwilą, a więc bądź chwilą jutrzejszego południa, bądź chwilą od niej wcześniejszą lub późniejszą. Wynika z tego, że jeśli Jan będzie jutro w południe w domu, to prawdą jest (...) w każdej chwili, że Jan będzie jutro w południe w domu. Mówiąc ogólnie, wykazaliśmy na przykładzie, że jeśli A jest b w chwili t , to prawdą jest w każdej chwili, a więc i w każdej chwili wcześniejszej od t , że A jest b w chwili t ²².”

Tak przedstawia się wywodzenie determinizmu z bezwzględnej obowiązywalności zasady dwuwartościowości. Drugi argument, jaki by można przedstawić opiera się na zasadzie przyczynowości. Każdy fakt (F zaistniały w chwili t) ma swoją przyczynę, która jest od niego wcześniejsza. *Ponieważ ciąg coraz wcześniejszy faktów, będących przyczynami faktu F , jest nieskończony, przeto w każdej chwili wcześniejszej od t (...), dzieje się jakiś fakt będący przyczyną faktu F* ²³.

Łukasiewicz obala drugi argument, posługując się takimi cechami czasu jak *nieskończoność* i *ciągłość*²⁴. Zauważa, że nieskończone ciągi przyczynowe (dziejące się w nieskończenie wielu chwilach) mogą mieć swoją granicę w jakimś momencie,

¹⁴ Zob. tamże, str. 117.

¹⁵ Tamże.

¹⁶ Tutaj i w całym wnioskowaniu t oznacza dowolną chwilę.

¹⁷ Tamże.

¹⁸ Tamże, str. 118.

¹⁹ Tamże, str. 117.

²⁰ Tamże, str. 118.

²¹ Tamże, str. 119.

²² Tamże.

²³ Tamże, str. 120.

²⁴ Zob. tamże, str. 121.

zatem bynajmniej nie muszą istnieć w całości czasu²⁵.

Mimo, że twórca trójwartościowej logiki wspomina o tym, że dwa argumenty deterministyczne są różne od siebie, to jednak uznaje ich ścisły związek i uważa, że dopiero razem są one zrozumiałe, mają sens²⁶. I tu wracamy do wątku, na którym rozstaliśmy się z Arystotelesem, mianowicie do rozumienia *zgodności z faktami*. Dla Łukasiewicza fakt istnieje, czyli prawdziwie jest faktem, o ile istnieje w teraźniejszości – albo ten dany fakt, albo jego przyczyny, albo skutki²⁷. Przyszłość nieokreślona w swych przyczynach, „potencjalna”, „wolna”, przyszłość, której w teraźniejszości nic jeszcze nie determinuje, nie istnieje, *nie jest jeszcze w chwili obecnej przesądzona*²⁸. A więc zdania o niej mówiące *nie są w chwili obecnej prawdziwe, bo nie mają żadnego realnego odpowiednika, ani też nie są fałszywe, bo ich zaprzeczenia także nie mają realnego odpowiednika. Posługując się niezbyt jasną terminologią filozoficzną, można by powiedzieć, że zdaniom tym filozoficznie odpowiada ani byt, ani niebyt, lecz możliwość*²⁹.

Tak argumentował Łukasiewicz potrzebą wprowadzenia do rachunku zdań trzeciej wartości. Odpowiednim zdaniom modalnym według niego także odpowiada jakaś rzeczywistość, tyle że nie faktyczna, zatem ani prawdziwa, ani fałszywa, ale właśnie jakaś *możliwa* (choć, jak sam zauważył, ontologicznie jest to *niezbyt jasne*³⁰). Nie zamierzmy prowadzić tu naukowej dysputy z wybitnym logikiem na tematy ontologiczne, ale musimy postawić pytanie: czy dodatkowa wartość logiczna naprawdę jest potrzebna? Widać, że rozpoczynając od próby przewyciężenia determinizmu i ujęcia przyszłych niekoniecznych wydarzeń, dochodzi się do odmówienia prawdziwości także wszystkiemu, czego skutki się już wyczerpały, a to już budzi więcej wątpliwości. Być może sprawa jest metafizycznie poważniejsza i należało by rozważyć, czy faktycznie to, co było, może się całkowicie wyczerpać, a to, co będzie, czy rzeczywiście nie jest gdzieś ujęte. Jest to skomplikowane, niemniej nawet już bez zagłębiania się w zakamarki teorii bytu możemy pytać czy wartości logiczne potrzebujemy wiązać z czasem? Jeśli zdanie opisuje dany fakt, przypisuje daną cechę danemu przedmiotowi i sytuację tę umieszcza ponadto w czasoprzestrzeni, czy potrzeba jeszcze i wartość tego zdania odnosić do czasu? Naszym zdaniem zdecydowanie nie.

Co najważniejsze nie ma relacji **przyczynowej** między zdaniem (sądem) a opisywanym przez nie stanem rzeczy, ani w jedną, ani w drugą stronę. Relacja, jakiej się tu można dopatrzeć, jest – wedle naszej *niezbyt jasnej terminologii* – „odzwierciedlaniem”. Byt jest pewną istniejącą treścią, tak substancjalny konkret, jak i ogólnie: fakt. Poznający, stwierdzając coś w jego temacie, może uchwycić tę treść trafnie, lub się pomylić, może tę treść sobie uświadomić i pojąć, a może po prostu przypadkiem ją zgadnąć. Ale zawsze, gdy sąd zgodzi się z faktem, będzie prawdziwy i fałszywy, kiedy się z nim będzie kłócił. Mimo że porządek intelektu i bytu są ze sobą związane, to pozostają w ten sposób odrębne. Tak więc można sprzeciwić się determinizmowi, dostrzegając, że nie wszystko jest z góry ustalone³¹,

²⁵ Zob. tamże, str. 121-123.

²⁶ Zob. tamże, str. 122: *Argument ten* (tj. argument za determinizmem oparty na zasadzie przyczynowości – przyp. P. Z.), *jakkolwiek niezależny od poprzedniego* (tj. opartego na zasadzie wyłączonego środka, tu można powiedzieć: na zasadzie dwuwartościowości – przyp. P. Z.), *staje się dopiero wtedy naprawdę zrozumiałym, gdy przyjmujemy, że każdy fakt ma swe przyczyny istniejące od wieków* (a więc de facto zaprzeczona zostaje ich niezależność, czego też dowodzą dalsze rozważania – przyp. P. Z.).

²⁷ Zob. tamże, str. 126.

²⁸ Tamże, str. 123.

²⁹ Tamże, str. 125.

³⁰ Zob. powyżej.

³¹ Dostrzegając to jako fakt, a nie jedynie opcję do wyboru, nie gorzej uzasadnioną niż determinizm, jak to ostrożnie zrobił Łukasiewicz, zob. tamże, str. 126.

a jednocześnie twierdzić, że zdania mówiące o przyszłości są prawdziwe albo fałszywe. Kiedy nie ma jeszcze przyczyn lub już nie ma skutków, naszej wiedzy nie jest w żaden sposób dostępna prawda o fakcie, ale ten fakt albo będzie albo nie będzie – albo był albo nie był.

Jeśli będziemy się trzymać toku myślenia zaproponowanego przez Łukasiewicza, to powinniśmy też uznać, że im więcej i im „poważniejsze” (*terminologia niezbyt jasna*) są skutki, bądź przyczyny faktu, tym zdanie jest „bardziej prawdziwe”³², a to wprowadza nieskończenie wiele wartości logicznych, które trudno obliczać i przede wszystkim coraz bardziej odbiega od doświadczenia i zdroworozsądkowego myślenia, które było tak ważne w początku rozumowania. Rzeczywistość miała być przecież lepiej opisana w tym uzupełnionym systemie.

Jeśli, jak stwierdzono, w świecie ciągi przyczynowe mają początki i kończą się, to nie tylko nikt w świecie, ale nawet sam świat jako całość nie może o wszystkim „pamiętać”, czy wszystkiego „przewidzieć”. Wiedza ta jest w obrębie świata nie w pełni dostępna, zawsze wybrakowana. A co gdyby istniało coś poza światem? Otóż wydaje się, że transcendentny intelekt poznający rzeczywistość, mógłby (teoretycznie rzecz biorąc, nie będziemy się w te trudne kwestie zbytnio zagłębiać) ująć w jakiś sposób całość bytu i pozostając także poza czasem uporządkować swą wiedzę wedle następstw (czy czegoś, co poza czasem mogłoby jakoś oddać czasowość). W tej koncepcji prawda jest obiektywna, u Łukasiewicza zaś subiektywna, bo odnosi się jedynie do możliwości poznawczych wewnątrz świata.

Nie wiemy, czy przy założeniu, że poza światem już nic więcej nie ma, da się do końca udowodnić tę trójwartościową logikę. Być może tak, ale nawet jeśli, to ostrożniejsze wydaje nam się przyjęcie możliwości istnienia czegoś transcendentnego. Bez wnikania w istotę takiego Absolutu dla potrzeb logiki możemy po prostu nazwać go prawdą.

³²To, jak się zdaje, nie jest wcale konieczne, nie wypowiadamy się tu kategorycznie, ponieważ nie zbadaliśmy tej kwestii. Niemniej wyraźnie widoczne jest u Łukasiewicza wiązanie prawdy z wiedzą (czy z możliwością wiedzy, może bardziej), a to może łatwo prowadzić w tę właśnie stronę, tj. ku logikom nieskończenie wielowartościowym.

Marcin Czakon
Katolicki Uniwersytet Lubelski Jana Pawła II

Pewnego rodzaju logiki niemonotoniczne jako model dla wnioskowania zawodnego.

Wstęp

Wielu współczesnych autorów w logikach niemonotonicznych widzi od dawna poszukiwany formalny model wnioskowań zawodnych, w szczególności upatrują w nich formalizację wnioskowań zdroworozsądkowych i logikę sztucznej inteligencji¹. Przeanalizujemy pewną, jedną z wielu, metodę konstruowania logik niemonotonicznych, która opiera się na modyfikacji klasycznej teorii konsekwencji, a inspirowana jest wnioskowaniami entymematycznymi. Powstające w ten sposób logiki niemonotoniczne przebadamy jako model dla wnioskowań zawodnych.

Ustalenia terminologiczne

Przez wnioskowania zawodne rozumiemy wnioskowania, w których stopień uznania przesłanek jest większy niż stopień uznania wniosku². Wnioskowania te posiadają wiele bardzo ciekawych własności formalnych i merytorycznych, których szczegółowe omówienie można znaleźć w literaturze przedmiotu³. Z prowadzonych analiz wynika jednoznacznie, że w cecha niemonotoniczności jest warunkiem koniecznym, ale nie wystarczającym wnioskowania zawodnego.

Przez logikę niemonotoniczną rozumiemy teorię, która nie spełnia warunku monotoniczności wyrażanego w następujący sposób:

$$\text{Jeśli } A \subseteq B, \text{ to } Cn(A) \subseteq Cn(B).$$

Co należy odczytywać: *Jeżeli zbiór A jest podzbiorem zbioru B , to zbiór klasycznych konsekwencji zbioru A jest podzbiorem zbioru klasycznych konsekwencji zbioru B .*

¹ Zob. Karl Schlechta, *Nonmonotonic logics: A preferential approach*, w: *Handbook of the History of Logic, volume 8, The Many Valued and Nonmonotonic Turn in Logic*, edited by Dov M. Gabbay, John Woods, North-Holland, Elsevier, Amsterdam 2007, s. 451; Bartosz Brożek, *Nauka w poszukiwaniu logiki*, w: *Semina Scientiarum*, Numer 1 (2002), Kraków, s. 8; Grigoris Antoniou, Kewen Wang, *Default logic*, w: *Handbook of the History of Logic, volume 8, The Many Valued and Nonmonotonic Turn in Logic*, edited by Dov M. Gabbay, John Woods, North-Holland, Elsevier, Amsterdam 2007, s. 517.

² Zob. Kazimierz Ajdukiewicz, *Logika pragmatyczna*, Państwowe Wydawnictwo Naukowe, Warszawa 1965, s. 106.

³ Tamże.

Konsekwencja założeń osiowych.

Obecnie, wzorując się na wnioskowaniach entymematycznych, spróbujemy scharakteryzować operację konsekwencji, która korzystać będzie ze zbioru twierdzeń jakby „ukrytych w tle” i która posłuży nam do skonstruowania konsekwencji nie spełniającej warunku monotoniczności.

Konsekwencję założeń osiowych zdefiniujemy w następujący sposób.

Definicja 1⁴ :

$$A \vdash_K \phi \text{ wtw } K \cup A \vdash \phi,$$

gdzie A jest dowolnym zbiorem formuł, K jest ustalonym zbiorem formuł, ϕ jest pojedynczą formułą. Wyrażenie $A \vdash_K \phi$ odczytujemy: ϕ jest konsekwencją zbioru A modulo zbiór K . Można powiedzieć, że zbiór A to zbiór przesłanek (wypowiadanych wprost), natomiast zbiór K to zbiór założeń osiowych (przesłanek nie wypowiadanych wprost), ϕ odgrywa rolę wniosku.

Od razu widać, że relacji konsekwencji założeń osiowych jest tyle rodzajów, ile można utworzyć różnych zbiorów K . Zachodzi następująca zależność pomiędzy relacją założeń osiowych a relacją klasyczną: $\vdash \subseteq \vdash_K$. Makinson taki rodzaj relacji nazywa *nadklasycznymi*, ponieważ są silniejsze od relacji klasycznych⁵.

Konsekwencja założeń osiowych spełnia te same warunki, co konsekwencja klasyczna: *zwrotność*, *indempotencję*, a przede wszystkim *monotoniczność*. Posiada również własność *zwartości* oraz *łączenia przesłanek w alternatywę*⁶. Okazuje się, że nie spełnia ona własności domknięcia na podstawianie.

Konsekwencja założeń domyślnych

Konsekwencja założeń osiowych wykorzystamy jako pewnego rodzaju łącznik prowadzący do konsekwencji niemonotonicznej. Niemonotoniczną konsekwencję otrzymujemy zmieniając zbiór K przesłanek ukrytych w tle, tak, żeby był on niesprzeczny ze zbiorem A . W przypadku, gdy zbiory te są sprzeczne, szukamy maksymalnego podzbioru K' zbioru K niesprzecznego ze zbiorem A .

Konsekwencję założeń domyślnych zdefiniujemy w następujący sposób⁷:

Definicja 2:

$$A \approx_K \phi \text{ wtedy i tylko wtedy, gdy } K' \cup A \vdash \phi, \text{ dla każdego } K' \subseteq K, \text{ który jest} \\ \text{maksymalnie niesprzeczny z } A,$$

gdzie A jest zbiorem przesłanek, K' jest zbiorem założeń domyślnych, a ϕ jest konsekwencją zbioru A modulo zbiór założeń K' .

Konsekwencja ta różni się od konsekwencji założeń osiowych, ponieważ tym razem zbiorem przesłanek w tle jest zbiór K' maksymalnie niesprzeczny ze zbiorem A . Jak poprzednio, tak i tutaj, jest wiele konsekwencji założeń domyślnych, zależnie do ilości różnych zbiorów K' .

⁴Zauważmy, że definicję tę można zapisać również w notacji operacji na zbiorach, co wyglądałoby w następujący sposób:

$\phi \in Cn_K(A) \text{ wtw } \phi \in Cn(A \cup K).$

⁵Por. David Makinson, *Od logiki klasycznej do niemonotonicznej*, Wydawnictwo Naukowe Uniwersytetu Mikołaja Kopernika, Toruń 2008, s. 12.

⁶Por. tamże, s. 24.

⁷Symbolu „ $\approx$ ” używamy na oznaczenie konsekwencji niemonotonicznej.

Makinson zauważa, że konsekwencja założeń domyślnych jest *nadklasyczna*, spełnia warunek *kumulatywnej przechodniości* i *łączenia przesłanek w alternatywę*. Nie spełnia warunku *zwartości* oraz jest *niemonotoniczna*⁸. Niemonotoniczność tej konsekwencji można zobrazować następującym przykładem:

Przykład 1:

Ustalmy $K = \{p \rightarrow q, q \rightarrow r\}$, dla zbioru $A = \{p\}$ otrzymujemy $A \approx_K r$, ponieważ $\{p\} \cup \{p \rightarrow q, q \rightarrow r\} \vdash r$ i zbiór K jest niesprzeczny z A . Jednak dla $A' = \{p, \neg q\}$ nie zachodzi $A' \approx_K r$, ponieważ jedynym niesprzecznym podzbiorem zbioru K jest zbiór $K' = \{q \rightarrow r\}$, a oczywistym jest, że nieprawdą jest $\{p, \neg q\} \cup \{q \rightarrow r\} \vdash r$.

W tym momencie widać, że warunek monotoniczności nie jest zachowany. Po dodaniu nowej przesłanki do zbioru A , wniosek r , który wcześniej był konsekwencją, obecnie przestaje nią być.

Problemy z konsekwencją założeń domyślnych

Konsekwencja założeń domyślnych w takiej postaci napotyka na dwie zasadnicze trudności związane z rodzajem założeń, jakie używane są do tworzenia zbioru K . Makinson zauważa, że zbiór K może być domknięty na klasyczną konsekwencję, czyli $K = Cn(K)$ lub niedomknięty, czyli $K \neq Cn(K)$.

W pierwszym przypadku niemonotoniczna operacja konsekwencji założeń domyślnych daje identyczne lub podobne wyniki co konsekwencja klasyczna. W szczególności interesująca jest sytuacja, w której $K = Cn(K)$ oraz zbiór A jest sprzeczny z K , co prowadzi do tego, że konsekwencja założeń domyślnych jest identyczna z konsekwencją klasyczną⁹.

W przypadku, w którym zbiór K nie jest domknięty na klasyczną konsekwencję, może dochodzić do sytuacji, że dwa zbiory równoważne klasycznie prowadzą do różnych operacji konsekwencji. Co można w łatwy sposób udowodnić¹⁰.

Dowód 2:

Załóżmy zbiór $K = \{p \rightarrow q, p \rightarrow r, r \rightarrow \neg p\}$ oraz zbiór $K' = \{p \rightarrow (q \wedge r), r \rightarrow \neg p\}$. Ze względu na własności koniunkcji i implikacji obydwie te zbiory są równoważne klasycznie. Przyjmijmy teraz zbiór $A = \{p\}$. Rozważmy konsekwencję założeń domyślnych zbioru A modulo zbiór założeń K . Zbiór K ma 8 podzbiorów:

$$\begin{aligned} K_1 &= \emptyset; \\ K_2 &= \{p \rightarrow q\}; \\ K_3 &= \{p \rightarrow r\}; \\ K_4 &= \{r \rightarrow \neg p\}; \\ K_5 &= \{p \rightarrow q, p \rightarrow r\}; \end{aligned}$$

⁸ Makinson zauważa, że konsekwencja założeń domyślnych spełnia warunek, nazywany ostrożną monotonicznością, który wyraża w ten sposób: „[...] jeśli dla dowolnego $x \in B$, $A \approx_K x$ oraz $A \approx_K y$, wtedy $A \cup B \approx_K y$ ” lub inaczej „jeśli $A \subseteq B \subseteq C_K(A)$, to $C_K(A) \subseteq C_K(B)$ ”. Por. David Makinson, *Od logiki klasycznej do niemonotonicznej*, Wydawnictwo Naukowe Uniwersytetu Mikołaja Kopernika, Toruń 2008, s. 33.

⁹ Interesujący jest dowód tego twierdzenia zaprezentowany przez Makinsona, w którym korzysta on z własności zwartości konsekwencji klasycznej oraz lematu Kuratowskiego-Zorna. Zob. David Makinson, *Od logiki klasycznej do niemonotonicznej*, Wydawnictwo Naukowe Uniwersytetu Mikołaja Kopernika, Toruń 2008, s. 35-36.

¹⁰ Por. tamże, s. 35.

$$\begin{aligned}
K_6 &= \{p \rightarrow q, r \rightarrow \neg p\}; \\
K_7 &= \{p \rightarrow r, r \rightarrow \neg p\}; \\
K_8 &= \{p \rightarrow q, p \rightarrow r, r \rightarrow \neg p\};
\end{aligned}$$

Odrzucamy zbiory K_8 , K_7 jako sprzeczne ze zbiorem A oraz zbiory K_1, K_2, K_3, K_4 pozostawiając zbiory K_5, K_6 jako maksymalne zbiory niesprzeczne z A . Zachodzą następujące relacje $A \cup K_5 \vdash q$ oraz $A \cup K_6 \vdash q$, w związku z tym, na mocy definicji 2, zachodzi konsekwencja założeń domyślnych $A \approx_K q$. Rozważmy następnie operację konsekwencji założeń domyślnych zbioru A modulo zbiór K' . Zbiór K' ma 4 podzbiory:

$$\begin{aligned}
K_1' &= \emptyset \\
K_2' &= \{p \rightarrow (q \wedge r)\}; \\
K_3' &= \{r \rightarrow \neg p\}; \\
K_4' &= \{p \rightarrow (q \wedge r), r \rightarrow \neg p\};
\end{aligned}$$

Odrzucamy podzbiór K_4' ponieważ jest on sprzeczny z A , odrzucamy również zbiór pusty, bierzemy pod uwagę tylko dwa jednoelementowe zbiory K_2' oraz K_3' . Widzimy, że w tym przypadku mamy następującą sytuację, $A \cup K_2' \vdash q$ oraz nie zachodzi $A \cup K_3' \vdash q$, w związku z tym, na mocy definicji, nie zachodzi konsekwencja założeń domyślnych $A \approx_K q$. W ten sposób widzimy, że dwa równoważne klasycznie zbiory założeń domyślnych K i K' pozwalają wyciągnąć inne konsekwencje z tego samego zbioru przesłanek A .

QED

Próba rozwiązania problemów konsekwencji założeń domyślnych

W takim razie, w jaki sposób można poradzić sobie z tymi trudnościami. Przypadek drugi, w którym dwa zbiory równoważne klasycznie prowadzą do różnych konsekwencji, Makinson proponuje rozwiązać pozwalając zapisywać wyrażenia w zbiorze założeń domyślnych, tylko jako alternatywy elementarne pochodzące z rozbicia kanonicznych koniunkcyjnych postaci normalnych na poszczególne czynniki¹¹. Rozwiązanie to sprawdza się tylko w przypadku języków ze skończoną ilością liter zdaniowych, co zauważ również Makinson.

Niestety, wydaje się, że rozwiązanie to nie jest satysfakcjonujące. Z jednej strony nie wiadomo, dlaczego mielibyśmy wyróżniać kanoniczną postać normalną, jako jedyny możliwy (właściwy) sposób zapisu przesłanek. Z drugiej strony nie wiadomo, dlaczego mielibyśmy ograniczać się do języków zawierających skończoną liczbę liter zdaniowych. Ciekawa jest uwaga Makinson, który z pewną ironią pisze o logikach zajmujących się językami nieskończonymi, że „*odczuwają perwersyjną przyjemność w obcowaniu z subtelnościami nieskończoności*”. Praktyka wnioskowań zawodnych pokazuje, a w szczególności wnioskowanie indukcyjne, że wnioskujący w wielu przypadkach nie jest w stanie wymienić wszystkich przesłanek, na których opiera swój wniosek ogólny. W takiej sytuacji nie jest spełniony warunek korzystania z języków ze skończoną liczbą liter zdaniowych, a w konsekwencji nie

¹¹ Por. tamże, s. 37.

jest możliwe zapisanie przesłanek, jako kanonicznych postaci normalnych. Wydaje się więc, że w związku z tym, tego konsekwencja założeń domyślnych nie będzie właściwym modelem formalnym dla wnioskowań zawodnych, w szczególności dla wnioskowania indukcyjnego.

Jak zauważyliśmy wyżej problemem jest również sprowadzenie konsekwencji założeń domyślnych do konsekwencji klasycznej, które dokonuje się w sytuacji, gdy zbiór K jest domknięty na klasyczną konsekwencję. Rozwiązanie tego problemu wymaga modyfikacji definicji konsekwencji założeń domyślnych. W literaturze istnieją co najmniej trzy typy takich modyfikacji.

Przypomnijmy, że konsekwencja założeń domyślnych polega na wyborze maksymalnie niesprzecznych z A podzbiorów K' zbioru K i w efekcie daje zbiór, będący iloczynem operacji konsekwencji klasycznej na sumach zbioru A ze zbiorami K' . Modyfikując tę konsekwencję bez zmiany pozostawimy tylko warunek niesprzeczności, a zrezygnujemy lub osłabimy warunki maksymalności zbiorów K' lub kwantyfikacji konsekwencji. Tak więc możemy dokonać trzech zasadniczych zmian w konsekwencji założeń domyślnych:

- (1) Zrezygnować z kwantyfikatora ogólnego w definicji 2, tzn. pozwolić, żeby operacja konsekwencji założeń domyślnych była iloczynem tylko niektórych wyników uzyskiwanych na podstawie zbiorów K' . Makinson ten rodzaj nazywa *operacją częściowego przecięcia*¹².
- (2) Całkowicie zrezygnować z wymogu występowania iloczynu na zbiorze wyników uzyskiwanych na podstawie zbiorów K' , tzn. wyróżnić tylko pojedyncze wyniki. Makinson ten rodzaj nazywa *operacją wolną od przecięcia*¹³.
- (3) Zrezygnować z wymogu maksymalności zbiorów K' , tzn. pozwolić, żeby operacja konsekwencji założeń domyślnych była iloczynem wyników uzyskiwanych na podstawie zbiorów K' niesprzecznych z A , ale niekoniecznie maksymalnych. Makinson ten rodzaj nazywa *operacją podmaksymalną*¹⁴.

Rezygnacja z kwantyfikatora ogólnego może odbywać się na bardzo wiele różnych sposobów. Można traktować pewne elementy zbioru K jako bezwarunkowo wymagane w podzbiorach K' , można ustalić pewnego rodzaju porządek na podzbiorach zbioru K i preferować wybrane podzbiory, można również po prostu preferować pewne elementy zbioru K . Wszystkie te sposoby ostatecznie sprowadzają się do konsekwencji modulo jakaś funkcja selekcji, której przeciwdziałaniem jest pożądana rodzina podzbiorów zbioru K . Przykładem takiej konsekwencji jest *relacyjna konsekwencja częściowego przecięcia*, która polega na ustaleniu porządku, za pomocą relacji $<$, na rodzinie, maksymalnie niesprzecznych ze zbiorem A podzbiorów zbioru K , a następnie wybraniu podzbiorów K' minimalnych względem tej relacji. Konsekwencja ta jest rozwiązaniem problemu ukłasyzowania konsekwencji założeń domyślnych¹⁵.

Operacje konsekwencji nazywane przez Maknisona wolnymi od przecięcia polegają na wyborze jednego ustalonego podzbioru K' zbioru K lub inaczej mówiąc preferują tylko niektóre formuły ze zbioru K . W literaturze omawianych jest kilka rodzajów takich operacji: konsekwencja oparta na *łańcuchach epistemicznych*, *inferencja porównywalnych oczekiwań*, a także *konsekwencja bezpieczna*. Na

¹²Por. tamże, s. 43.

¹³Por. tamże.

¹⁴Por. tamże.

¹⁵Por. tamże, s. 43-48.

przykład konsekwencję bezpieczną definiujemy następująco: $A \approx_{SA} \phi \equiv A \cup S_A \vdash \phi$, gdzie S_A jest zbiorem przesłanek ukrytych w tle, który powstaje poprzez przyłączenie przesłanek ze zbioru K w skutek badania składników zbioru A w odpowiedni sposób¹⁶.

Operacje podmaksymalne, których zasadniczą cechą jest rezygnacja z warunku maksymalności podzbiorów K' , również mogą powstawać na różne sposoby. Makinson wspomina o trzech takich sposobach: konsekwencji opartej o *dodatkowe warunki ukryte w tle*, konsekwencji bazującej na *podziorach maksymalnie bogatych informacyjnie* oraz konsekwencji korzystającej z *pojęcia stanu epistemicznie najlepszego*. Dla przykładu, konsekwencja oparta o dodatkowe warunki ukryte w tle definiowana jest w następujący sposób: $A \approx_{KJ} \phi \equiv A \cup K'$ dla każdego $K' \subseteq K$, takiego, że K' jest maksymalnie niesprzeczne ze zbiorem $A \cup J$. Co należy rozumieć w ten sposób, że poszukujemy takich podzbiorów K' , które są niesprzeczne z pewnymi warunkami znajdującymi się w zbiorze J , który ukryty jest w tle, tak jak zbiór K . Warto zauważyć, że wyrażenia występujące w zbiorze J , nie biorą bezpośredniego udziału w wyciąganiu konsekwencji, a pozwalają tylko ustalić podzbiory zbioru K ¹⁷.

Wydaj się, że zaprezentowane sposoby rozwiązania problemu ukłasyzczenia konsekwencji założeń domyślnych nie są satysfakcjonujące. Makinson rozwiązując problem ukłasyzczenia konsekwencji założeń domyślnych, przedstawia ponad dziesięć różnych sposobów modyfikacji tej konsekwencji, jednocześnie ani razu nie wspominając o kryterium lub racjach, jakie stoją za tym lub innym rozwiązaniem. Problem sprowadzenia konsekwencji założeń domyślnych do konsekwencji klasycznej, zostaje tutaj rozwiązany generując problem, który wydaje się być poważniejszy, polegający na uzasadnieniu wyboru przesłanek (ukrytych w tle), z których będziemy korzystać przy wnioskowaniu (lub uzasadnieniu wyeliminowania innych przesłanek). Wręcz groteskowa wydaje się być uwaga Makinsona odnośnie tego problemu: „Z perspektywy praktycznej wszystkie te operacje powinny być traktowane jako element skrzynki z narzędziami, które stosujemy kiedy jest to odpowiednie i wygodne. Pewne z nich mogą być w jednym celu bardziej odpowiednie niż drugie, a w innym zaś może być odwrotnie. Czasami wybór ma znaczenie. Innym razem żadne z nich może nie być tym narzędziem, którego szukamy”¹⁸. I jeszcze pisze: „Chcielibyśmy zasugerować, że żaden [rodzaj konsekwencji domyślnych założeń] z nich nie jest tym jedynym poprawnym, i że żaden z nich nie jest zawsze najlepszy do zastosowania”¹⁹. Wydaje się, że od logiki należałoby oczekiwać, że dzięki niej człowiek będzie jaśniej myślał, ściślej wypowiadał się i poprawnie uzasadniał twierdzenia. Nawet jeżeli uznamy, że różne teorie logiczne opisują różne fragmenty rzeczywistości, to wydaje się, że zadanie logika polega nie tyle na tym, żeby stworzyć jak najwięcej takich teorii, ale żeby stworzyć teorie i podać uzasadnienie dla stosowania jej w danym fragmencie rzeczywistości.

Bez wątpienia możemy stwierdzić, że konsekwencja złożeń domyślnych spełnia warunek niemonotoniczności, który jest warunkiem koniecznym wnioskowania zawodnego. Wydaje się, że ta konsekwencja nie spełnia innych ważnych dla wnioskowania zawodnego własności formalnych i merytorycznych. Chociażby nie uwzględnia w żaden sposób zmiany stopnia uznania wniosku wraz z nową informacją. Wszelkie wnioskowania zawodne, a w szczególności wnioskowanie

¹⁶ Zob. tamże, s. 51-55.

¹⁷ Zob. tamże, s. 48-51.

¹⁸ Tamże, s. 56.

¹⁹ Tamże, s. 56.

indukcyjne, z każdym kolejnym prawdziwym zdaniem, jako przesłanką, zwiększają stopień uznawania wniosku. Własności tej w ogóle nie ujmuje konsekwencja założeń domyślnych. Może ona co najwyżej stwierdzić, że dany wniosek, na podstawie nowej informacji, nie może być już uznawany. Wydaje się, że przedstawiona konsekwencja nie jest właściwym modelem dla wnioskowań zawodnych, ponieważ nie uwzględnia w wystraczający sposób związku uprawdopodobniania zachodzącego pomiędzy przesłankami a wnioskiem takich wnioskowań.

Warto zauważyć, że opisana konsekwencja niemonotoniczna oddaje raczej ten rodzaj wnioskowania, który znany jest, zwłaszcza w logice prawniczej, jako *domniemanie usuwalne*, czyli *wzruszalne* (*praesumptio iuris tantum*). Wnioskowanie tego rodzaju opiera się na normie nakazującej uznawanie określonych zdań za prawdę dopóki coś przeciwnego nie zostanie dowiedzione. Przykładem domniemania usuwalnego jest domniemanie ojcostwa, określone w *Kodeksie rodzinnym i opiekuńczym* z 25 lutego 1964 r. (art. 62 §1). W świetle przywołanego przepisu stąd, że mężczyzna x jest w chwili narodzin dziecka y mężem matki tego dziecka, wynika, że mężczyzna x jest ojcem dziecka y , ale tylko dopóki skądinąd nie okaże się, że x nie jest ojcem dziecka y . Zbiór K , charakteryzujący omówiony w tym rozdziale typ niemonotonicznych operacji konsekwencji, wyznacza zakres domniemanych wniosków. Jak powiedzieliśmy, wnioskowania zawodne mają podobną własność, mianowicie swoista tymczasowość wniosków, ale jest ona tylko jedną z istotnych cech wnioskowania zawodnego.

Osobną kwestią są omówione tutaj poważne trudności rachunkowe, które napotyka sposób definiowania niemonotonicznych operacji konsekwencji jako konsekwencji modulo zbiór założeń domyślnych. W wielu punktach konstrukcji mamy do czynienia z praktycznie arbitralnymi rozstrzygnięciami, zmierzającymi wyłącznie do osiągnięcia pożądanego rezultatu formalnego. Trudno sobie wyobrazić, w jaki sposób te rozstrzygnięcia mogłyby zostać ugruntowane w teorii wnioskowań zawodnych.

Warto zwrócić również uwagę na moment przejścia od monotonicznej konsekwencji założeń osiowych do konsekwencji założeń domyślnych, polega on na pozwoleniu na taką modyfikację zbioru przesłanek, żeby uniknąć sprzeczności. Wydaje się, że jest to jedna z ważniejszych motywacji stojących za konstruowaniem kolejnych niemonotonicznych konsekwencji. B. Brożek pisze, że jedną z logik, która radzi sobie ze sprzecznością jest właśnie logika niemonotoniczna²⁰. Logikę niemonotoniczną wymienia on w jednym ciągu z różnego rodzaju logikami parakonsystentnymi, podkreślając, że wszystkie te logiki są *nieeksplozywne* (nie obowiązują w nich prawo Dunsza Szkota).

20 Zob. Bartosz Brożek, *Nauka w poszukiwaniu logiki*, w: *Semina Scientiarum*, Numer 1 (2002), Kraków, s. 7-11.

Anna Karczewska
Katolicki Uniwersytet Lubelski Jana Pawła II

Paradoksy łańcuchowe a logika rozmyta.

Powstawanie paradoksów łańcuchowych jest koronnym argumentem za tezą, że zwroty nieostre, będące źródłem tych paradoksów, leżą poza zasięgiem logiki klasycznej i że potrzebna jest w związku z tym jakaś inna, nieklasyczna logika. Do miana logiki nieostrości pretenduje m.in. logika nieskończenie wielowartościowa, zwana logiką rozmytą. Celem niniejszego artykułu jest zbadanie, czy logika ta jest satysfakcjonującym rozwiązaniem paradoksów łańcuchowych.

Paradoksy nieostrości

Mówiąc najogólniej, nieostre są te terminy, których sposób użycia nie jest całkowicie wyznaczony przez ich znaczenie. Z tego względu, posiadają one przypadki brzegowe, tzn. o niektórych przedmiotach nie sposób rozstrzygnąć, czy podpadają pod zakres tych zwrotów, czy też nie podpadają. [Ajdukiewicz, s. 35] Ta nieokreśloność ma przy tym charakter istotny, to znaczy nie zależy od naszych aktualnych możliwości dyskryminacji. Stąd nie jest, na przykład, nieostrym zwrot „mierzy 167 cm i 3mm”, chociaż na ogół nie możemy rozstrzygnąć, czy jest prawdziwy o danym człowieku, czy nie. Wymienia się trzy naczelne przyczyny nieostrości: (a) stopniowalność [Keefe, s. 134] („wysoki”, „bogaty”), (b) złożoność charakterystyk („piękny”, „dzieło sztuki”) i (c) alternatywność charakterystyk („dopływ” - krótsza rzeka jest dopływem dłuższej, węższa jest dopływem szerszej, w sytuacji gdy jedna rzeka jest dłuższa, ale węższa od drugiej, nie wiadomo, o której można powiedzieć, że jest dopływem tej drugiej).¹ [Haack, s.111]

Głównym problemem logicznym związanym z nieostrością są paradoksy łańcuchowe. Diogenes Laertios przypisuje autorstwo paradoksów łańcuchowych (obok kilku innych) uczniowi Euklidesa, Eubulidesowi z Miletu. [Diogenes Laertios, s. 137] Często przyjmuje się, że były one wymierzone przeciwko poglądom Arystotelesa, w szczególności przeciwko jego doktrynie środka, ale nie jest do końca jasne, jaki był rzeczywisty zamiysł Eubulidesa. Sam Arystoteles nie odnosi się nigdzie do domniemyanych zarzutów, niemniej takie rozumienie paradoksów - jako trudności stanowiska perypatetyckiego - pojawia się już we wczesnych komentarzach do pism Arystotelesa. [Mignucci, s. 231] Paradoksy łańcuchowe były później żywo

¹ S. Haack obok (b) i (c) wymienia niemożliwość rozstrzygnięcia, czy spełniony jest warunek (lub jeden z kilku warunków) należenia do zakresu predykatu: „[...] in certain cases it is indeterminate whether the condition, or one of the conditions, is satisfied”. Jako przykład tego rodzaju predykatu wymienia określenia kolorów: “how closely, e.g. does an object have to resemble an English pillar-box if it is to count as red?”. Haack zastrzega przy tym, że nie chodzi tu o niemożliwość, której źródłem jest brak dostatecznej informacji na temat przedmiotu. [Haack, s.111] Przypadki tego rodzaju nieostrości sprowadzają się jednak, jak sądzimy, do punktów (a)-(c).

dyskutowane w Szkole Stoickiej, zwłaszcza przez Chryzypa z Soloi. Niestety, zachowały się tylko pośrednie ślady zainteresowania stoików paradoksami. Treść paradoksów referują przede wszystkim Sekstus Empiryk (polemicznie) w „Przeciw uczonym” oraz Marek Tulliusz Ciceron w „Księgach akademickich” (*Academica*).

Eubulides przedstawił dialektyczne argumenty Soryt (od gr. *soros*, czyli stos) i Łysy (gr. *falakros*), jako sekwencje pytań: Czy jedno ziarno tworzy stos? Czy dwa ziarna tworzą stos? Czy trzy ziarna tworzą stos? ... Czy dziesięć tysięcy ziaren to stos? oraz: Czy człowiek z dziesięcioma tysiącami włosów na głowie jest łysy? Czy człowiek z dziewięcioma tysiącami dziewięćset dziewięćdziesięcioma dziewięcioma włosami jest łysy? ... Czy człowiek bez żadnego włosa na głowie jest łysy? Odpowiedzi na pierwsze i ostatnie z pytań są równie oczywiste i wzajemnie przeciwstawne. Jeżeli jednak zgodzimy się dodatkowo, że różnica jednego elementu (ziarna czy włosa) nie może mieć wpływu na zmianę odpowiedzi z pozytywnej na negatywną ani z negatywnej na pozytywną, dojdziemy do zupełnie kontrintuicyjnych rezultatów, mianowicie, że, dowolnie duża liczba ziaren nie stanowi stosu albo, że nawet człowieka bez jednego włosa na głowie nie można uznać za łysego.

Starożytni byli świadomi, że analogicznemu traktowaniu poddaje się wiele innych wyrażen i z powodzeniem to wykorzystywali, głównie w odniesieniu do określeń stopniowalnych, takich jak „niewiele”, „stary” i tym podobnych. „Soryt” (a przez jakiś czas także „falakros”) przyjął się jako ogólna nazwa paradoksów o podobnym kształcie. Wszystkie soryty opierają się na następującym schemacie: rozważamy n -elementowy ciąg przedmiotów: $a_1, a_2, \dots, a_n$, dla odpowiednio dużego n . Pewna własność F , na przykład bycie biednym, z założenia przysługuje pierwszemu elementowi tego ciągu – komuś, kto nie posiada nawet jednej złotówki (ani żadnej innej waluty) – i nie przysługuje ostatniemu – posiadaczowi miliona złotych. Własność F musi być przy tym taka, żeby minimalna różnica pomiędzy sąsiadującymi ze sobą elementami ciągu nie miała znaczenia dla możliwości orzekania F o tych elementach. Dowolny soryt można zatem ująć jako wnioskowanie, o niewątpliwie prawdziwych przesłankach i równie niewątpliwie fałszywym wniosku. Zwyczajowo dokonuje się tego na trzy, traktowane jako równoważne, sposoby: dwa z przesłankami w formie zdań warunkowych, jeden – z przesłankami w formie zanegowanej koniunkcji.

W pierwszej, najbardziej klasycznej wersji, soryt jest łańcuchem inferencji, w którym konkluzja każdego z występujących w nim wyprowadzeń (wyjawszy ostatnie) staje się racją wyprowadzenia następującego bezpośrednio po nim. Pierwsza przesłanka stwierdza przysługiwanie własności F początkowemu elementowi ciągu, a druga jest przesłanką warunkową, zgodnie z którą, jeżeli F przysługuje pierwszemu elementowi, to przysługuje też drugiemu. To, na mocy reguły *Modus Ponens* (MP) pozwala wnosić, że F przysługuje drugiemu elementowi. Analogicznie wyprowadza się przysługiwanie własności F pozostałym elementom ciągu. Jeżeli F oznacza bycie biednym, to mamy dla przykładu:

a_1 jest biedny.

Jeżeli a_1 jest biedny, to a_2 jest biedny.

A zatem a_2 jest biedny.

Jeżeli a_2 jest biedny, to a_3 jest biedny.

A zatem a_3 jest biedny.

...

A zatem a_{n-1} jest biedny.

Jeżeli a_{n-1} jest biedny, to a_n jest biedny.

a_n jest biedny.

Drugi sposób zapisu sorytu jest skróconą wersją pierwszego. Zamiast poszczególnych przesłanek warunkowych występuje tu jako założenie generalizacja

$\forall i: (Fa_{i-1} \rightarrow Fa_i),$

gdzie i przebiega zbiór liczb naturalnych. Założenie to, zwane *twierdzeniem o nieodróżnialności* (*Indiscriminability Thesis*), wyraża przekonanie, że minimalna różnica jaka występuje pomiędzy kolejnymi elementami rozważanego ciągu nie ma znaczenia dla możliwości orzekania cechy F , albo, mówiąc inaczej, że dowolna para sąsiadujących ze sobą elementów ciągu jest od siebie nieodróżnialna ze względu na własność F :

Ten, kto nie ma ani złotówki (a_i) jest biedny.

Każdy, kto posiada tylko złotówkę więcej od kogoś, kto jest biedny, też jest biedny.

A zatem: posiadacz miliona złotych jest biedny (albo: Każdy jest biedny).

Zauważmy, że powyższe wnioskowanie ma w istocie kształt wnioskowania rekurencyjnego. Pierwsza przesłanka (założenie wnioskowania) – orzekająca o pierwszym elemencie ciągu, że spełnia on zadany warunek – jest początkiem rekurencji, a druga (warunkowa – twierdzenie o nieodróżnialności) krokiem rekurencyjnym – ustala zależność, wedle której spełnianie warunku przenosi się na sąsiadujący element ciągu. Z tych dwóch przesłanek można wyprowadzić wniosek, że każdy element ciągu spełnia zadany warunek, czyli, że każdemu elementowi ciągu przysługuje własność F , a w tym przypadku, że biedny jest nawet posiadacz powiedzmy, miliona złotych.

Jedynym funktorem zdaniowym występującym w obu powyższych wariantach sorytu jest implikacja. Znaczenie zdań warunkowych było w starożytności przedmiotem żywego sporu w Szkole Megarejsko-Stoickiej. W ramach tego sporu zaproponowano przynajmniej cztery różne definicje implikacji [Tkaczyk], [Mates], ale do powstania paradoksu wystarcza najsłabsze pojęcie implikacji, przedstawione przez Filona z Megary. Implikacja Filona jest tożsama ze współczesną implikacją materialną definiowaną za pomocą negacji koniunkcji poprzednika i negacji następnika. Niektóre przedstawione przez stoików przykłady wnioskowań łańcuchowych, uwzględniają ten fakt. [Mignucci, s. 329], [Williamson, ss. 23-36] Stąd, za całkowicie równoprawne sformułowanie uważa się to, w którym zdania warunkowe zastępuje się równoważnymi im zdaniami z funktorami negacji i koniunkcji:

Ten, kto nie ma ani jednej złotówki jest biedny.

Nieprawda, że zarazem: ktoś, kto nie ma ani jednej złotówki jest biedny, a ten, kto ma [tylko] jedną złotówkę nie jest biedny.

Nieprawda, że zarazem: ktoś, kto ma [jedną] złotówkę jest biedny, a ten, kto ma dwa złote nie jest biedny.

...

Nieprawda, że zarazem: ktoś, kto ma $n-1$ (999999) złotych jest biedny, a ten, kto ma n (milion) złotych nie jest biedny.

A zatem: ten, kto ma n (milion) złotych jest biedny.

We wszystkich trzech wariantach sortyt jest wnioskowaniem konkluzywnym zarówno wedle wymogów logiki stoików (pierwsze i ostatnie sformułowanie), jak i wymogów współczesnej logiki standardowej. Tymczasem, choć jego przesłanki są ewidentnie prawdziwe, to wniosek jest nie mniej ewidentnie fałszywy.

Można wyróżnić cztery główne strategie wyjaśnienia przyczyn powstawania paradoksów i możliwości ich uniknięcia. [Keefe, ss. 19-20] Ogólnie, paradoks zobowiązuje do poddania kontroli formalnej i materialnej poprawności wnioskowania. Najbardziej oczywistą strategią, wydaje się podważenie prawdziwości założeń wnioskowania – któreś z poszczególnych przesłanek warunkowych (lub odpowiadającej jej przesłanki z negacją koniunkcji) bądź przesłanki indukcyjnej (tutaj na przykład epistemicyzm, którego przedstawicielami są T. Williamson, R. Sorensen). Druga, najmniej popularna strategia, polega na zakwestionowaniu przesłanki kategorycznej, przyznającej cechę F pierwszemu przedmiotowi ciągu lub, alternatywnie, uznaniu za prawdziwą konkluzji (P. Unger). Trzecia strategia sprowadza się do przyznania, że faktycznie z prawdziwych przesłanek wynika dedukcyjnie fałsz i że jest to dowód wewnętrznej sprzeczności pojęć nieostrych (G. Frege, B. Russell, M. Dummett). Ostatnia, najbardziej radykalna odpowiedź na paradoksy nieostrości polega na zaprzeczeniu konkluzywności wnioskowania.

Wielowartościowe podejście do paradoksów nieostrości jest wyjściem łączącym, w pewnym stopniu, pierwszą i ostatnią strategię – wysuwa zastrzeżenia zarówno co do poprawności materialnej, jak i co do konkluzywności wnioskowania. U podstaw takiego podejścia leży przekonanie, że logika klasyczna nie może adekwatnie reprezentować wnioskowań z zastosowaniem wyrażeń nieostrych właśnie dlatego, że takie wyrażenia nie zawsze są doskonale prawdziwe lub zupełnie fałszywe. Zauważmy na marginesie, że zarzuca się w tym kontekście logice klasycznej nadmierną precyzję, tymczasem wydaje się, że jest ona raczej nie dość dokładna – lekceważy fakt, że pewne przedmioty należą do zakresu nieostrej nazwy w mniejszym lub większym stopniu. Logika rozmyta natomiast aspiruje do tego, żeby odzwierciedlić tę fundamentalną własność wyrażeń nieostrych.

Logika rozmyta

Rozmyta teoria nieostrości opiera się na uogólnionej teorii zbiorów L. Zadeha. Uogólniona teoria zbiorów, w odróżnieniu od teorii klasycznej, dopuszcza stopniowalność relacji należenia do zbioru. Od strony matematycznej to dopuszczenie polega na przyporządkowaniu elementom niepustej przestrzeni X wartości z przedziału $[0,1]$. [Machina, s. 65] Wartości te wskazują na stopień, w jakim elementy X należą do rozmytego zbioru: 0 oznacza nienależenie do zbioru,

1 całkowite, a wartości pośrednie jakieś częściowe należenie. Zbiorem rozmytym nazywa się zatem funkcję f ze zbioru X w zbiór liczb rzeczywistych od 0 do 1:

$$f: X \rightarrow [0, 1]$$

W ujęciu mnogościowym zbiór rozmyty jest po prostu podzbiorem iloczynu kartezjańskiego $X \times [0, 1]$. Zauważmy jeszcze, że klasyczne zbiory, mające jedynie 0 i 1 jako możliwe stopnie przynależności, są szczególnymi przypadkami zbiorów rozmytych. [Malinowski, s. 557]

Dla rozmytych zbiorów definiuje się działania będące uogólnieniem działań na zbiorach klasycznych. Intuicyjnie, dodawanie dwóch zbiorów rozmytych X i Y polega na wybraniu dla każdego z ich argumentów tej wartości, przypisywanej mu w X lub Y , która jest większa:

$$f_{(X \cup Y)}(x) = \max(f_X(x), f_Y(x)).$$

Z kolei mnożenie zbiorów rozmytych przyporządkowuje argumentom mniejszą z dwu wartości:

$$f_{(X \cap Y)}(x) = \min(f_X(x), f_Y(x)).$$

Suma i iloczyn rozmytych zbiorów są zatem zbiorami wybranych par uporządkowanych z X lub Y . Stopień przynależności x do sumy zbiorów rozmytych X i Y jest równy wyższemu ze stopni przynależności x do X i do Y , a stopień przynależności do iloczynu tych zbiorów – niższemu. [Machina, s. 69] Dopełnienie zbioru rozmytego X o funkcji charakterystycznej $f_X(x)$ jest funkcją $1 - f_X(x)$, czyli element, który należy do X w stopniu i należy do dopełnienia X w stopniu $1 - i$.

Tak określony zbiór rozmyty ma być modelem zakresu predykatów nieostrych. [Machina, s. 65], [Malinowski, s. 557] Wartości, które funkcja przyporządkowuje elementom dziedziny są interpretowane jako miara spełniania przez nie (tj. elementy) nieostry predykatu. Dla przykładu, człowiek o 150 tysiącach włosów na głowie spełniałby predykat „jest łysy” w najwyższym stopniu (1), człowiek o 3 włosach w ogóle nie spełniałby tego predykatu (0), a człowiek o 50 tysiącach włosów w jakimś stopniu pośrednim. W przypadku takich określeń stopniowalnych jak „jest łysy”, „jest bogaty”, sytuacja jest więc stosunkowo prosta. Skojarzenie predykatu ze skalą pozwala oddać myśl, że spełnianie tego predykatu przez jakiś przedmiot jest zależne od intensywności określonej cechy w tym przedmiocie. Kazus nieostrości typu (b) i (c), tzn. tej, której źródłem jest złożoność bądź alternatywność charakterystyk, jest jednak bardziej zawiły, gdyż nie ma tam ścisłego powiązania pomiędzy nasileniem własności, a należeniem do zakresu predykatu – przeciwnie, niekiedy ten sam stopień nasilenia własności – jak na przykład szerokość rzeki, czy to, że w ogóle jest ona szersza od tej, z którą się łączy – przesądza o tym, że pewne przedmioty spełniają dany nieostry predykat, ale nie przesądza w przypadku innych. Przynależność do zakresu tego predykatu nie jest kwestią stopnia, jak nieostrość wyrażenia „dopływ” nie polega na tym, że można być bardziej lub mniej dopływem, ale na tym, że w niektórych przypadkach wyznaczniki bycia dopływem wchodzą ze sobą w konflikt. Wydaje się zatem, że omawiane tu rozwiązanie nie stosuje się jednakowo do wszystkich typów nieostrości. Zdecydujemy się jednak przejść do porządku nad tą, rysującą się już na wstępie, trudnością logiki rozmytej, gdyż jak powiedzieliśmy, paradoksy łańcuchowe powstają przy zastosowaniu właśnie określeń stopniowalnych.

U podstaw logiki rozmytej leży przekonanie, że zwrotem nieostrym jest też samo

określenie „prawdziwy” i że jest tak dlatego, że prawda jest stopniowalna. Innymi słowy, tak jak gradacji podlega przynależność przedmiotów do zakresu predykatu, tak zdania o tych przedmiotach mogą być bardziej lub mniej fałszywe, więc oprócz zwykłych wartości – prawdy i fałszu – należy przyjąć kontinuum pośrednich wartości logicznych. Przy tym przeważnie utożsamia się stopień prawdziwości zdania orzekającego predykat o jakimś przedmiocie x , ze stopniem przynależenia x -a do zakresu tego predykatu.² [Malinowski, s. 557]

Jest kwestią dyskusyjną, czy przyjmowane w logice rozmytej wartości można zasadnie traktować jako stopnie prawdziwości. Z omówionych powyżej względów, oczywistym źródłem wątpliwości są ponownie przypadki (b) i (c), jakkolwiek nie mniej problematyczne okazuje się przypisanie stopni prawdziwości zdaniom z określeniami stopniowalnymi. Prawdziwość zdania zawierającego takie określenia zależy wprawdzie od stopnia nasilenia pewnej mierzalnej cechy, jak prawdziwość wyrażenia „ x jest wysoki” zależy od wzrostu x -a, ale to, że można być bardziej lub mniej wysokim, nie znaczy jeszcze, że zdanie mówiące o tym, że ktoś jest wysoki może być mniej lub bardziej prawdziwe. R. Keefe rozważając to zagadnienie, podkreśla, że mimo, że intensywność własności rzeczywiście decyduje o spełnianiu lub niespełnianiu przez przedmiot nieostrego predykatu, nie znaczy to, że miara intensywności może być uznana za skalę prawdziwości. Świadczy o tym chociażby sposób użycia predykatów takich jak np. „ostry” (o kącie), „ciepły” (w znaczeniu „posiada temperaturę wyższą od 0 absolutnego”), które mimo, że mierzalne, nie są nieostre. Nieostrym predykatom rzeczywiście odpowiada pewien obszar na tej skali intensywności cechy, ale o nieostrości predykatu decyduje to, że jest to obszar jedynie przybliżony, niedoprecyzowany, a przyporządkowanie zdaniom wartości tej skali nie pozwala oddać tej nieokreśloności. Zdaniem Keefe, dowolna skala liczbowa, mająca odzwierciedlić nieostrość związaną ze stopniowalnością, jest raczej miarą własności (np. wzrostu) niż prawdziwości. [Keefe, ss. 134-135]

Okazuje się więc, że model rozmyty nie odpowiada nieostrości typu (b) i (c), a w odniesieniu do typu (a), ujmuje raczej to, co mierzalne i precyzyjne, niż typową dla nieostrości niedokładność. Argumenty przytaczane za rozwiązaniem wielowartościowym są przekonujące dla tych, którzy i tak wierzą, że prawda jest stopniowalna, ale nie muszą przemówić do tych, którzy sądzą inaczej. Wobec tego, pogląd, że przedział $[0,1]$ odpowiada stopniom prawdziwości wyrażen okazuje się opierać głównie na akcie wiary, nie jest bardziej uzasadniony niż przekonanie o dwu-, trój- lub inno- wielowartościowości. Z drugiej strony nie świadczy to jeszcze o tym, że mylne jest samo rozwiązanie – mimo, że słabo uzasadnione, może okazać się prawdziwe. W takim razie warto zbadać, jak sprawdza się ono w interesującym nas kontekście paradoksów łańcuchowych. Przyjmiemy więc na próbę, że wartości logiczne w logice rozmytej są, tak jak chcą jej zwolennicy, stopniami prawdziwości. Ponieważ jednak nieporęczne jest mówić o stopniach prawdziwości wyrażonych liczbowo, będziemy w dalszym ciągu posługiwać się głównie przybliżonymi określeniami opisowymi – „doskonale prawdziwy”, „nieomal prawdziwy” itd., co jest zresztą bliższe pomysłowi Zadeha. Omówimy teraz podstawy teorii logiki rozmytej.

W języku logiki rozmytej występują dokładnie te same stałe symbole logiczne, co w języku logiki klasycznej, a więc: „¬”, czyli funktor negacji, funktor koniunkcji „∧”, funktor alternatywy „∨”, funktor implikacji „→” oraz funktor równoważności

²Zauważmy jednak, że zbiór wartości logicznych logiki rozmytej faktycznie nie musi być identyczny ze zbiorem wskaźników. [Machina, s.65].

„ $\equiv$ ”, które odczytywane są tak, jak w logice klasycznej. Jednak, jak się przekonamy, zmienia się znaczenie tych symboli. Spójniki logiki rozmytej można definiować na wiele różnych sposobów. Często przyjmuje się na przykład charakterystyki alternatywy, koniunkcji i negacji odpowiadające przedstawionym uprzednio działaniom na zbiorach rozmytych, odpowiednio dodawania, mnożenia i dopełnienia. Będziemy posługiwali się literami A , B oraz C na oznaczenie dowolnych wyrażeń logiki rozmytej. Alternatywa rozmyta przyjmuje wyższą z wartości jej argumentów

$$V(A \vee B) = \max(V(A), V(B)),$$

koniunkcja (dwóch zdań) – niższą:

$$V(A \wedge B) = \min(V(A), V(B)),$$

a wartość negacji zdania jest równa różnicy między pełną prawdą i wartością tego zdania

$$V(\neg A) = 1 - V(A).$$

Widać, już że funktory te dla argumentów o klasycznych wartościach (0 lub 1) dają zawsze klasyczny wynik, są zatem uogólnieniem funktorów klasycznej, dwuwartościowej logiki. Niemniej dla wartości nieklasycznych, pośrednich wynik może być nieklasyczny. Z tego względu tautologiami logiki rozmytej nie są np. jedno z najbardziej charakterystycznych tautologii logiki klasycznej, jak prawo wyłączonego środka ($p \vee \neg p$) i prawo niesprzeczności ($\neg(p \wedge \neg p)$), które zgodnie z powyższymi wzorami, dla $V(p)=0,5$ są tylko w połowie prawdziwe. Niektórzy uważają ten fakt za wadę logiki rozmytej, mimo to funktory negacji, koniunkcji i alternatywy uchodzą za stosunkowo nieproblematyczne. Nieco więcej uwagi chcemy poświęcić natomiast spójnikowi implikacji.

Klasyczny funktor implikacji, w przeciwieństwie do pozostałych funktorów dwuargumentowych, nie jest symetryczny, wartość logiczna wyrażenia utworzonego za pomocą tego funktora nie zależy wyłącznie od wartości logicznych jego argumentów, ale również od tego, w jakim porządku te argumenty są ze sobą zestawione, [Kiczuk, s. 138] tzn. implikacja o fałszywym poprzedniku i prawdziwym następniku jest prawdziwa, a implikacja o prawdziwym poprzedniku i fałszywym następniku fałszywa. Tę właściwość implikacji (klasycznej) rozszerza się na implikację rozmytą, która zachowuje dwuwartościową matrycę dla klasycznych wartości (jest prawdziwa gdy fałszywy jest jej poprzednik lub prawdziwy jej następnik, a fałszywa gdy zarazem prawdziwy jest jej poprzednik i fałszywy następnik), a oprócz tego ma być malejąca względem pierwszego argumentu:

$$\text{dla dowolnych } A, B \text{ i } C, \text{ jeżeli } V(A) \leq V(C), \text{ to } V(A \rightarrow B) \geq V(C \rightarrow B),$$

oraz rosnąca względem drugiego:

$$\text{dla dowolnych } A, B \text{ i } C, \text{ jeżeli } V(B) \leq V(C), \text{ to } V(A \rightarrow B) \leq V(A \rightarrow C), \text{ [Baczyński,}$$

$$\text{Jayaram, s. 2]}$$

czyli, mówiąc mniej precyzyjnie, wartość implikacji jest wprost proporcjonalna do wartości następnika a odwrotnie proporcjonalna do wartości poprzednika.

Sformułowane właśnie wymagania nie wyznaczają jednak jednej implikacji. Wybór konkretnej definicji implikacji musi być podyktowany spełnianiem przez nią jakichś dodatkowych wymogów. Zwykle oczekuje się od implikacji rozmytej, żeby zachowywała jak najwięcej klasycznych tautologii, dla przykładu, w opinii S. Kundu

i J. Chena powinna zachowywać prawo transpozycji prostej $((p \rightarrow q) \rightarrow (\neg q \rightarrow \neg p))$ i prawo komutacji $((p \rightarrow (q \rightarrow r)) \rightarrow (q \rightarrow (p \rightarrow r)))$. [Januszewski, s. 116] Trudno oprzeć się jednak wrażeniu, że każde zaproponowane tu rozwiązanie byłoby naznaczone pewną arbitralnością. Najważniejsze implikacje rozmyte – implikacja J. Łukasiewicza, J. Goguena i K. Gödla – są scharakteryzowane następująco: we wszystkich przypadkach, implikacja jest zupełnie prawdziwa, o ile jej następnik jest prawdziwy co najmniej w tym samym stopniu, co poprzednik. Jeżeli natomiast wartość następnika jest niższa od wartości poprzednika, to cała implikacja przybiera wartość:

- $1 - V(A) + V(B)$ (implikacja Łukasiewicza),
- następnika - $V(B)$ (implikacja Gödla),
- ilorazu wartości następnika i wartości poprzednika - $(V(B))/V(A)$ (implikacja Goguena).

Przyjmuje się ponadto, że rozmyte wynikanie, jak wynikanie klasyczne, dziedziczy doskonałą prawdziwość bądź, że dziedziczy najniższą z wartości przesłanek, zatem: (a) wyrażenie B wynika ze zbioru wyrażeń $A_1, \dots, A_n$ wtedy i tylko wtedy, gdy B jest doskonale prawdziwe w każdej interpretacji, w której doskonale prawdziwe są przesłanki $A_1, \dots, A_n$, lub (b) B wynika ze zbioru wyrażeń $A_1, \dots, A_n$ wtedy i tylko wtedy, gdy B w żadnym wartościowaniu nie przyjmuje wartości mniejszej niż najniższa z wartości przysługujących przesłankom. Zauważmy, że przy tej pierwszej definicji wynikania będzie obowiązywała ważna klasyczna reguła, tj. MP:

$$\frac{A, A \rightarrow B}{B},$$

gdyż reguła ta zachowuje doskonałą prawdziwość przesłanek – dla dowolnej implikacji rozmytej (Łukasiewicza, Gödla i Goguena), jeżeli ta implikacja jest doskonale prawdziwa oraz doskonale prawdziwy jest jej poprzednik, to doskonale prawdziwy musi być też następnik.

Jeżeli natomiast od wynikania logicznego będziemy wymagali, żeby konkluzja była nie mniej prawdziwa od najmniej prawdziwej przesłanki (jak w (b)), to reguła MP będzie obowiązywała dla implikacji Gödla. Załóżmy bowiem, że w pewnej interpretacji V_I , $V_I(B)$ jest mniejsze zarówno od $V_I(A)$, jak od $V_I(A \rightarrow B)$. Skoro $V_I(B)$ jest mniejsze od $V_I(A)$, to $V_I(A \rightarrow B)$ jest równe $V_I(B)$, co jest sprzeczne z założeniem, że $V_I(B)$ jest mniejsze od $V_I(A \rightarrow B)$. Pokazaliśmy zatem, że dla implikacji Gödla obowiązuje reguła MP. Reguła ta nie obowiązuje tymczasem ani dla Łukasiewicza ani Goguena definicji implikacji, co można łatwo sprawdzić przyjmując dowolne nieklasyczne wartości $V(A)$ i $V(B)$, takie, że $V(B)$ jest mniejsze od $V(A)$ – w każdej takiej interpretacji obie przesłanki będą bardziej prawdziwe niż konkluzja.

Rozmyta implikacja w rozumowaniach łańcuchowych

Rozważymy teraz, jak narzędzia logiki rozmytej stosują się do wnioskowań, w których kluczową rolę odgrywają zwroty nieostre. W tym celu zrekonstruujemy wnioskowanie łańcuchowe w jego najbardziej podstawowej wersji, czyli z serią przesłanek warunkowych. Nasza rekonstrukcja musi przy tym uwzględniać

poczynione wcześniej intuicyjne ustalenia: po pierwsze, że przesłanki są prawdziwe, a wniosek fałszywy i po drugie, że mimo to wnioskowanie jest konkluzywne. Pamiętajmy też, że rozpatrujemy ciąg przedmiotów $a_1, a_2, \dots, a_n$, z których każdy kolejny w mniejszym stopniu spełnia dany nieostry predykat F . Zatem, zgodnie z tym, co powiedzieliśmy wcześniej na temat wartości logicznych, pierwsze zdanie – Fa_1 – jest doskonale prawdziwe, a każde kolejne zdanie postaci Fa_{i+1} będzie mniej prawdziwe od zdania Fa_i o pewną znikomą wartość, oznaczmy ją jako ε , aż do ostatniego, zupełnie fałszywego.

Korzystając z podanych uprzednio wzorów, możemy już obliczyć wartości przesłanek warunkowych: ponieważ następnik (Fa_{i+1}) będzie zawsze nieco mniej prawdziwy od poprzednika (Fa_i) , żadna implikacja nie będzie doskonale prawdziwa, niemniej dla bardzo małego ε , implikacja Łukasiewicza będzie nieomal doskonale prawdziwa (dokładnie $(1-\varepsilon)$) podczas gdy wartość implikacji Gödla będzie stopniowo malała, tak, że ostatnia przesłanka, o zupełnie fałszywym następniku, będzie zupełnie fałszywa. Przesłanka ta będzie też fałszywa, jeżeli jej wartość będziemy obliczać według wzoru Goguena – 0 podzielone przez dowolną liczbę daje 0. Okazuje się wobec tego, że implikacje Gödla i Goguena nie spełniają przyjętych przez nas wymogów. Uznaliśmy bowiem, że minimalne różnice pomiędzy kolejnymi wyrazami nie mają znaczenia dla stosowalności predykatu, czyli że sąsiadujące ze sobą przedmioty są od siebie nieodróżnialne ze względu na własność F . Jeżeli więc pewien przedmiot należy do zakresu predykatu, to przedmiot bezpośrednio po nim następujący w ciągu również. Tymczasem implikacja te nie pozwalają uchwycić tego powiązania.

Zauważmy przy okazji, że w logice rozmytej wnioskowanie łańcuchowe w postaci implikacyjnej nie jest równoważne temu z zanegowaną koniunkcją w przesłankach. Kolejne przesłanki o kształcie $\neg(A \wedge \neg B)$ będą, podobnie jak implikacja Gödla, przyjmowały wartość B , ale tylko dotąd, dokąd $V(B)$ nie będzie mniejsze od 0,5, oraz wartość $1-V(A)$, dla $V(B)$ mniejszego od 0,5. Skoro zarówno $V(A)$, jak i $V(B)$ będą sukcesywnie malały, to wartość przesłanek będzie najpierw stopniowo malała, do wartości 0,5, a następnie rosła, aż do nieomal prawdziwej ostatniej przesłanki. Wynik ten nie zgadza się z naszymi wcześniejszymi ustaleniami, gdyż kolejne przesłanki, które intuicyjnie uważamy (na gruncie logiki rozmytej) za równie prawdziwe i to prawdziwe w wysokim stopniu, przyjmują tu różne wartości logiczne. Przybliża natomiast podejście wielowartościowe do epistemicyzmu, zgodnie z którym istnieje gdzieś wyraźna granica pomiędzy przedmiotami spełniającymi nieostry predykat, a tymi, które go nie spełniają, a nieostrość polega wyłącznie na tym, że nie wiemy gdzie ta granica leży.

Zastanówmy się teraz nad kwestią konkluzywności wnioskowania. Jak pamiętamy, w grę wchodzi tylko reguła Modus Ponens. Powiedzieliśmy, że reguła ta obowiązuje dla wszystkich rozważanych definicji implikacji, jeśli przyjmiemy klasyczną definicję wynikania (a), a tylko dla implikacji Gödla, przy bardziej restryktywnej definicji, tj. definicji (b), która wymaga zachowania także stopnia prawdziwości. Gdyby więc zgodzić się na definicję (b) wraz z implikacją Łukasiewicza lub Goguena, wnioskowanie okazałoby się niekonkluzywne. Jednak odrzucenie reguły MP należałoby uznać za poważny mankament wielowartościowego rozwiązania paradoksu, bo, jak się wydaje, reguła ta oddaje najbardziej elementarny sens okresu warunkowego. Jeśli zaś z wyrażen $(A \rightarrow B)$ oraz A nie można wywnioskować B , to nie ma podstaw, żeby symbol „ $\rightarrow$ ” rozumieć jako odpowiednik wyrażenia „jeżeli ..., to ...” z języka potocznego. Takie rozwiązanie należy uznać za zbyt radykalne,

pozostaje więc przyjąć definicję (b) wraz z implikacją Gödla lub definicję (a) dla dowolnej implikacji. Ponieważ implikacje Gödla i Goguena odrzuciliśmy jako nieadekwatne, w świetle powyższych ustaleń właściwa logika nieostrości powinna przyjmować implikację Łukasiewicza wraz z klasyczną definicją wynikania.

Przeciwko wielowartościowemu rozwiązaniu paradoksów wysuwa się wiele zastrzeżeń natury filozoficznej, większość z nich dotyczy założenia o stopniowalności prawdy. Z kolei naczelnym argumentem zwolenników logiki rozmytej jest to, że takie ujęcie po prostu „działa”. [Haack, s. 230] Chcemy obecnie rozważyć pewne zarzuty związane z kontrintuicyjnymi własnościami implikacji Łukasiewicza. Zarzuty te mogłyby wskazywać na to, że logika rozmyta nie sprawdza się jednak w kontekście nieostrości.

Williamson zwraca uwagę na to, że zgodnie z charakterystyką implikacji, musimy przypisać (doskonałą) prawdziwość zdaniom, które intuicyjnie byłibyśmy raczej skłonni uznać za fałszywe. [Williamson, s.138] Przykładem takiego zdania miałyby być: *Jeżeli jest obudzony, to śpi* (*If he is awake, then he is asleep*). W sytuacji, gdy poprzednik tego zdania jest półprawdziwy ($V(A)=0,5$) (tzn. ten, o kim mowa jest tylko na wpół obudzony), półprawdziwy będzie też następnik, będący negacją poprzednika ($V(\neg A)=1-0,5=0,5$). Ponieważ zaś, jak powiedzieliśmy, gdy wartości poprzednika i następnika są sobie równe, to implikacja jest doskonale prawdziwa, doskonale prawdziwa – jak podkreśla Williamson, równie prawdziwa co prawo tożsamości ($A \rightarrow A$) – będzie implikacja ($A \rightarrow \neg A$). Williamson prawdziwość tej ostatniej uważa za poważny błąd.

Podkreślmy przede wszystkim, że ranga implikacji ($A \rightarrow A$) i ($A \rightarrow \neg A$) nie jest jednakowa: wartość logiczna drugiej z implikacji zależy od wartościowania, podczas gdy pierwsza jest tautologią, jest prawdziwa bez względu na wartościowanie. Z drugiej strony, wątpliwości Williamsona nie są całkiem bezpodstawne, niemniej trzeba uświadomić sobie, że dotyczą one w równym stopniu logiki klasycznej: zgodnie z klasyczną matrycą dla funktora implikacji prawdziwa jest implikacja zbudowana z dowolnych dwóch zdań o tej samej wartości logicznej. Co więcej, kiedy A jest zdaniem fałszywym, prawdziwa jest też problematyczna implikacja ($A \rightarrow \neg A$). Źródłem omawianej przez Williamsona trudności jest ekstensjonalność języka logiki. Wiadomo, że spójnik języka potocznego „jeżeli ... to...” co najmniej w niektórych użyciach nie jest ekstensjonalny i że klasyczny funktor implikacji nie do końca odtwarza znaczenie okresu warunkowego języka potocznego. Tę słabość „dziedziczy” implikacja rozmyta. Wydaje się jednak, że klasyczny funktor implikacji ujmuje przynajmniej rudymentalne znaczenie zwrotu „jeżeli... to...” i można go uważać za dostateczne przybliżenie języka potocznego [Haack, ss.119 i 125], podczas gdy, jak się przekonamy, implikacja Łukasiewicza podlega innym jeszcze zarzutom, których nie można zignorować.

Keefe omawia trudności związane z ujęciem stosunków zachodzących pomiędzy przedmiotami należącymi do zakresu nieostrego predykatu. Załóżmy, proponuje Keefe, że Zenon i Cezary należą do brzegowych przypadków predykatu „jest łyśy”, ale Zenon ma nieco mniej włosów niż od Cezarego i spełnia ten predykat w stopniu 0,5, podczas gdy Cezary – w stopniu 0,4 [Keefe, ss. 96-97]. Zatem, zgodnie z wcześniejszymi ustaleniami, wartość logiczna zdania „Zenon jest łyśy” jest równa 0,5, a zdania „Cezary jest łyśy” – 0,4. Zgodzimy się wówczas, że jeżeli można nazwać łyśym Cezarego, to tym bardziej powinniśmy nazwać łyśym Zenona. Z drugiej strony, jeżeli nie nazwiemy łyśym Zenona, to nie możemy też nazwać łyśym Cezarego. Tymczasem wartość implikacji „Jeżeli Cezary jest łyśy,

to Zenon też jest łysy” będzie równa wartości implikacji o tym samym poprzedniku i zaprzeczonym następniku – „Jeżeli Cezary jest łysy, to Zenon nie jest łysy”, oba zdania warunkowe będą mianowicie doskonale prawdziwe, skoro wartość poprzednika jest mniejsza od wartości następnika. Prawdziwość pierwszego ze zdań całkowicie czyni zadość naszym oczekiwaniom, podczas gdy drugie zdanie intuicyjnie zdaje się być całkowicie fałszywe. Analogicznie dzieje się w przypadku zdań „Jeżeli Zenon nie jest łysy, to Cezary jest łysy” i „Jeżeli Zenon jest łysy, to Cezary też jest łysy”. Można łatwo obliczyć, że przy przyjętych założeniach wartość obu implikacji będzie równa 0,9, choć pierwsze ze zdań należy uznać za fałszywe, skoro bowiem Zenon jest bardziej łysy od Cezarego, to jeżeli Zenon nie jest łysy, to tym bardziej łysy nie jest Cezary.

Wprowadzenie nieskończenie wielu wartości logicznych miało umożliwić modelowanie pewnych związków, polegających na tym, że niektóre przedmioty w większym, a inne w mniejszym stopniu należą do zakresu nieostrego predykatu. W świetle zarysowanych tu trudności okazuje się, że logika rozmyta wprawdzie spełnia to zadanie na elementarnym poziomie, ale zawodzi w bardziej złożonych przypadkach.

Literatura

1. Ajdukiewicz K., *Język i poznanie*, t. II, Warszawa 1985.
2. Baczyński M., Jayaram B., *An Intruduction to Fuzzy Implications*, STUDFUZZ 231 (2008), s.1-35.
3. Diogenes Laertios, *Żywoty i poglądy słynnych filozofów*, Warszawa 1984.
4. Haack S., *Deviant Logic, Fuzzy Logic*, Chicago 1996.
5. Januszewski E., *Logiczne i filozoficzne problem związane z logiką rozmytą*, RF LV (2007) 1, s. 109-128.
6. Keefe S., *Theories of Vagueness*, Cambridge 2000.
7. Malinowski G., *Many-Valued Logic*, w: Jacqueline D. (red.), *Companion to Philosophical Logic*, Oxford 2002.
8. Machina K. F., *Truth, Belief, and Vagueness*, w: *Journal of Philosophical Logic* 5 (1976) s.47-78.
9. Mates B., *Stoic Logic*, Berkeley, University of California Press 1961.
10. Mignucci M., *The Stoic Analysis of the Sorites*, w: *Proceedings of the Aristotelian Society, New Series*, Vol. 93 (1993), s. 231-245.
11. Tkaczyk M., *Zmienna czasowa w starożytnej i średniowiecznej teorii zdań warunkowych*, RF VL 2(2007).
12. Williamson T., *Vagueness*, London, New York 1998.