

STUDIA Z METODOLOGII I FILOZOFII JĘZYKOZNAWSTWA

2

Współczesna filozofia języka.
Inspiracje i kierunki rozwoju

pod redakcją Piotra Stalmaszczyka

Seria

STUDIA Z METODOLOGII I FILOZOFII JĘZYKOZNAWSTWA

2

Współczesna filozofia języka. Inspiracje i kierunki rozwoju

pod redakcją Piotra Stalmaszczyka

Katedra Językoznawstwa Angielskiego i Ogólnego Uniwersytetu Łódzkiego
& Primum Verbum

Łódź 2013

Seria STUDIA Z METODOLOGII I FILOZOFII JĘZYKOZNAWSTWA

Redaktor naczelny Piotr Stalmaszczyk

Rada redakcyjna

Ireneusz Bobrowski (Uniwersytet Jagielloński i Instytut
Języka Polskiego PAN)
Romuald Gozdawa-Gołębiowski (Uniwersytet Warszawski)
Henryk Kardela (Uniwersytet Marii Curie-Skłodowskiej
w Lublinie)
Joanna Odrowąż-Sypniewska (Uniwersytet Warszawski)
Renata Przybylska (Uniwersytet Jagielloński i Instytut
Języka Polskiego PAN)
Przemysław Tajsner (Uniwersytet im. Adama Mickiewicza
w Poznaniu)
Mieszko Tałasiewicz (Uniwersytet Warszawski)
Ewa Willim (Uniwersytet Jagielloński)

Sekretarze redakcji Anna Obrębska
Wiktor Pskit

Recenzent Tadeusz Szubka

Skład i korekta Anna Obrębska

© Katedra Językoznawstwa Angielskiego i Ogólnego UŁ
© Wydawnictwo Primum Verbum

ISBN 978-83-62157-72-3

Łódź 2013

Wydawnictwo PRIMUM VERBUM
ul. Gdańska 112, 90-508 Łódź
e-mail: wydawnictwo@primumverbum.pl
www.primumverbum.pl

Spis treści

Piotr Stalmaszczyk

Wstęp: Współczesna filozofia języka. Inspiracje i kierunki rozwoju 7

Andrzej Pawelec

Rozdział 1: Kontynentalna filozofia języka: wybrane stanowiska 15

Henryk Kardela

Rozdział 2: Formy symboliczne Ernsta Cassirera a gramatyka kognitywna 29

Marek Nowak

Rozdział 3: Teoria aktów mowy Adolfa Reinacha 50

Paweł Grabarczyk

Rozdział 4: Dyrektywalna teoria znaczenia w interpretacji behawioralnej 78

Marek Maciejczak

Rozdział 5: Kontekst znaczenia językowego 96

Filip Kawczyński

Rozdział 6: Czy filozofia języka potrzebuje pojęcia sądu logicznego? 111

Janusz Maciaszek

Rozdział 7: Niedosłowność i interpretacja 136

Joanna Odrowąż-Sypniewska

Rozdział 8: Debata między minimalizmem semantycznym a kontekstualizmem 161

Justyna Grudzińska

Rozdział 9: Nowe spojrzenie na referencję i kwantyfikację. Filozoficzne implikacje dynamicznego zwrotu w semantyce 185

Mieszko Tałasiewicz

Rozdział 10: Co to jest "termin jednostkowy"? 204

Joanna Klimczyk

Rozdział 11: Czy „kot” jest racją? Spór o normatywność znaczenia 213

Tadeusz Ciecierski, Katarzyna Kuś

Rozdział 12: Przyczynowe teorie nazw a argument Stephena Sticha 235

Elżbieta Chrzanowska-Kluczewska

Rozdział 13: Filozoficzne podstawy teorii gier — najnowsze tendencje 252

Jarosław Jakielaszek

Rozdział 14: Semantyka formalna i Program Minimalistyczny: od logicznej analizy języka do biolingwistycznego naturalizmu 271

Indeks 287

Wstęp

Współczesna filozofia języka. Inspiracje i kierunki rozwoju

Piotr Stalmaszczyk (Uniwersytet Łódzki)

Filozofia to przede wszystkim filozofowanie, a filozofowanie to bezspornie sposób życia, tak jak jest nim bieganie, bycie zakochanym, granie w golfa, oburzanie się na politykę, bycie damą. Wszystko to są to formy i sposoby życia.

José Ortega y Gasset *¿Qué es filosofía*¹

1. Wstęp

Studia zebrane w niniejszym tomie poświęcone są współczesnej filozofii języka, a także (choć nie zawsze bezpośrednio) związkom zachodzącym pomiędzy filozofią języka a językoznawstwem. Lingwistyka — jako jedna z nauk szczegółowych — uniezależniła się od filozofii, jednocześnie jednak podstawowe pytania językoznawcze, takie jak te dotyczące natury języka i natury znaczenia albo relacji zachodzących pomiędzy rzeczywistością, myśleniem i językiem, mają zdecydowanie filozoficzny charakter. Związki pomiędzy wspomnianymi tu dyscyplinami wciąż pozostają ścisłe, co wynika również z tego, że — jak zauważa André Comte-Spontville — lingwistyka powstała na skutek zerwania z filozofią języka, podobnie jak socjologia na skutek zerwania z filozofia społeczną. Powołując się na sformułowanie Gastona Bachelarda, Comte-Sponville (2007: 120) mówi w tym kontekście o „epistemologicznym cięciu”.

Doskonałym przykładem najwcześniejszego filozoficznego namysłu nad problemami językowymi jest *Kratylos* Platona, powszechnie uważany za najstarszy (w tradycji europejskiej) traktat z zakresu językoznawstwa i filozofii języka. Ważne miejsce w tym dialogu zajmują rozważania dotyczące nazw własnych oraz istoty znaczenia, a także pochodzenia języka. Badania dotyczące znaczenia i nazw własnych prowadzone są dzisiaj w obrębie zarówno semantyki i pragmatyki językoznawczej, jak i filozofii języka (natomiast kwestie

¹ José Ortega y Gasset, *¿Qué es filosofía?* (1957/2005: 173), przekład własny.

związane z pochodzeniem języka badane są przede wszystkim przez dynamicznie rozwijającą się teorię 'ewolucji języka'²).

Na potrzeby niniejszego wstępu można wymienić cztery (częściowo) niezależne dyscypliny:

- językoznawstwo (ang. *linguistics*)
- filozofia języka (ang. *philosophy of language*)
- filozofia językoznawcza (ang. *linguistic philosophy*)
- filozofia językoznawstwa (ang. *philosophy of linguistics*)

Pytanie o (nie)zależności czterech powyższych dyscyplin ma siłą rzeczy charakter filozoficzny. Terminy angielskie zostały dodane, ponieważ większość cytowanych we wstępie prac jest po angielsku oraz ze względu na wieloletnią dominację anglo-amerykańskiej literatury na temat filozofii języka i istotnej dla jej rozwoju filozofii analitycznej³.

Ian Mackenzie (1997) we wstępie do swojego podręcznikowego opracowania z zakresu filozofii lingwistycznej określa przedmiot wspomnianych powyżej dyscyplin w sposób następujący: językoznawstwo zajmuje się empirycznymi badaniami języka, zadaniem filozofii języka jest wgląd w podstawową istotę zjawisk badanych przez językoznawstwo, natomiast filozofia lingwistyczna stanowi pewnego rodzaju podejście do filozofii języka⁴. Natomiast Zeno Vendler (1974) określa cel filozofii lingwistycznej jako wszelkie badania prowadzone w związku ze strukturą i funkcjonowaniem języków zarówno naturalnych, jak i sztucznych (i przywołuje tu przykłady tak różne jak metafizyczne pisma Arystotelesa, teorię deskrypcji Russella czy prace Ryle'a dotyczące pojęć)⁵ i dodatkowo wyróżnia jeszcze filozofię językoznawstwa, zajmującą się filozoficznym namysłem nad takimi pojęciami, jak znaczenie, synonimia, parafraza, składnia, tłumaczenie, oraz badaniem logicznego statusu teorii lingwistycznych. Tak rozumiana filozofia językoznawstwa jest jedną z gałęzi bardziej ogólnej filozofii nauk. Według Vendlera filozofia języka zajmuje się filozoficznym namysłem nad językiem, problemami takimi jak związek pomiędzy językiem a rzeczywistością i innymi. Do prac z zakresu filozofii języka Vendler zalicza między innymi książkę Benjamina Lee Whorfa *Język, myśl i rzeczywistość*,

² Zob. na ten temat: Waciewicz (2008).

³ Należy jednocześnie podkreślić, że w niniejszym zbiorze znalazły się również teksty dotyczące dokonań autorów piszących po francusku (M. Merleau-Ponty, P. Ricoeur), niemiecku (E. Cassirer, H.-G. Gadamer, E. Husserl, A. Reinach), a także po polsku (K. Ajdukiewicz).

⁴ Mackenzie (1997: ix) twierdzi, że owo 'podejście' do filozofii języka najlepiej ilustrują *Dociekania filozoficzne* Wittgensteina, a zwłaszcza przedstawiony tam sposób analizowania języka w użyciu (prymat badania użycia nad analizowaniem struktur logicznych).

⁵ Do tak rozumianej filozofii lingwistycznej Zeno Vendler zalicza również swoją własną pracę, Vendler (1974).

a także — choć z pewnym zawahaniem — *Tractatus Logico-Philosophicus* Ludwiga Wittgensteina (zob. Vendler 1974: 5–6).

Nie przesądzając o słuszności tych rozróżnień, warto zauważyć, że według najnowszych opracowań dotyczących filozofii języka jej podstawowe pytania dotyczą problemów znaczenia, rozumienia i komunikacji (Davies 2006: 29), a filozofów interesują trzy szerokie obszary języka — składnia, semantyka i pragmatyka (Martinich 2009: 1). Według Scotta Soamesa do podstawowych pojęć filozofii języka (i całej filozofii) należą prawda, odniesienie, znaczenie, możliwość, sąd oraz implikatura (Soames 2010: 1), natomiast Peter Prechtl za podstawowe zagadnienie uznaje określenie stosunku: „język — myślenie — rzeczywistość” (Prechtl 2007: 10).

Na jeszcze inny charakter relacji zachodzących pomiędzy filozofią a językiem zwraca uwagę Ludwig Wittgenstein, mówiąc o opętaniu umysłu przez środki naszego języka i walce z tym opętaniem, którą prowadzi filozofia⁶. Według niego to właśnie filozofowie sprowadzają słowa „z ich zastosowań metafizycznych z powrotem do użytku codziennego” (*Dociekania filozoficzne*, § 116). Ale — jak przytomnie zauważyła Hannah Arendt — walka z owym opętaniem (dodajmy, iż domniemanym) „może się dokonać jedynie za pomocą samego języka” (Arendt 2002: 166)⁷. Z tego dwugłosu filozofów można wyprowadzić wskazówki, niekoniecznie zgodne z intencjami samych autorów (to zastrzeżenie dotyczy zwłaszcza Wittgensteina), co do konieczności wszechstronnego badania i analizy języka, zarówno jako systemu, jak i jego różnych użyć we wszelkich możliwych kontekstach, czyli do dopełnienia rozważań filozoficznych badaniami empirycznymi⁸.

Jednocześnie spektrum pytań filozoficznych w dyskusjach nad językiem znacznie się w ciągu ostatnich dziesięcioleci poszerzyło. Warto w tym kontekście przyjrzeć się dwóm niedawnym publikacjom o charakterze encyklopedycznym: *The Blackwell Guide to the Philosophy of Language* (Devitt i Hanley, red. 2006) oraz *The Oxford Handbook of Philosophy of Language* (Lepore i Smith, red. 2006). Przewodnik wydawnictwa Blackwell przynosi 20 rozdziałów podzielonych na trzy podstawowe części: zagadnienia podstawowe (w znacznej części poświęcone problemom badania znaczenia, kontrastujące program Davidsona z podejściem Grice’a), znaczenie oraz odniesienie. Natomiast w tomie wydawnictwa oksfordzkiego 41 rozdziałów składa się na siedem części poświęconych kolejno kontekstowi historycznemu (przede wszystkim kon-

⁶ „Filozofia jest walką z opętaniem naszego umysłu przez środki naszego języka” (*Dociekania filozoficzne*, § 109).

⁷ Warto w tym kontekście przytoczyć również obserwację Hilarego Putnama, poczynioną w kontekście omawiania późnej filozofii Wittgensteina i pojęcia ‘gier językowych’: „celów gry językowej nie da się [...] opisać bez użycia języka tej gry lub jakiejś gry pokrewnej” (Putnam 1999: 79).

⁸ Zob. np. teksty zebrane w Stalmaszczyk (red.) 2010.

cepcjom Fregego i Wittgensteina), naturze języka, naturze znaczenia, naturze odniesienia, teorii semantycznej, różnym zjawiskom lingwistycznym (takim jak czas, liczba, kwantyfikatory), różnorodności aktów mowy oraz epistemologii i metafizyce języka.

Wprawdzie John Searle, współtwórca teorii aktów mowy, twierdzi, że centrum zainteresowania we współczesnej filozofii przesunęło się obecnie z języka na umysł⁹, jednakże ilość nowych publikacji (w tym także podręczników) i konferencji poświęconych różnym aspektom filozofii języka zdaje się wskazywać na wzrost zainteresowania taką tematyką. Martinich wspomina, że przed współczesną filozofią języka (mimo że jest to dyscyplina trudna i techniczna) otwierają się jasne perspektywy (Martinich 2009: 1), a Soames zwraca uwagę na wręcz majeutyczny charakter filozofii języka, zwłaszcza w odniesieniu do naukowych badań nad językiem (zarówno językiem naturalnym, jak i językami formalnymi) i użyciem języka (Soames 2010: 1).

Wbrew dość kłśliwej uwadze Briana Magee, że „filozof, który poświęca swe życie trosce o język, jest jak cieśla, który cały swój czas poświęca ostrzeniu narzędzia, a tych nie używa do niczego innego, jak tylko do ostrzenia narzędzi właśnie” (Magee 1998: 55), współcześni filozofowie języka dysponują narzędziami, które pozwalają na precyzyjną analizę problemów dalece wykraczających poza rozważania metajęzykowe. Niektóre z tych narzędzi i obszarów badawczych są omówione w kolejnych rozdziałach.

2. Zawartość tomu

Autorami poszczególnych rozdziałów zebranych w niniejszym tomie są językoznawcy, filozofowie i logicy. Podejmowane zagadnienia obejmują przywoływane w podtytule zbioru inspiracje i kierunki rozwoju współczesnej filozofii języka.

Andrzej Pawelec przedstawia wybrane stanowiska w obrębie kontynentalnej filozofii języka, a zwłaszcza dokonania Maurice’a Merleau-Ponty’ego (reprezentującego podejście fenomenologiczne), Hansa-Georga Gadamera (reprezentanta hermeneutyki) i Paula Ricoeura (związanego z podejściem hermeneutyczno-semiotycznym). Pawelec omawia też relacje zachodzące pomiędzy podejściami analitycznymi i kontynentalnymi (również aspekty terminologiczne tego rozróżnienia) i postuluje przekroczenie istniejących horyzontów metodologicznych.

⁹ Według Searle’a „wiele problemów językowych to szczególne przypadki problemów umysłu. Użytek, jaki czynimy z języka, jest wyrazem bardziej fundamentalnych biologicznie niż sam język zdolności mentalnych, więc żeby w pełni zrozumieć, jak język funkcjonuje, musimy pokazać, jak jest on w tych zdolnościach zakorzeniony” (Searle 2010: 20).

Henryk Kardela omawia związki zachodzące pomiędzy badaniami form symbolicznych prowadzonymi przez Ernsta Cassirera a językoznawstwem kognitywnym, zwłaszcza w ujęciu Ronalda Langackera i jego koncepcji poznania oderwanego.

Marek Nowak przedstawia koncepcję aktu społecznego i teorię czynności mowy niemieckiego fenomenologa Adolfa Reinacha, a w dalszej części rozdziału również ontologiczne założenia teorii Reinacha. Nowak podkreśla pionierski charakter dokonań Reinacha, widoczny również poprzez porównanie z późniejszymi opracowaniami Austina, Searle'a i Vandervekena.

Rozdział Pawła Grabarczyka przybliży przedstawioną w latach trzydziestych ubiegłego wieku przez Kazimierza Ajdukiewicza dyrektywalną teorię znaczenia i proponuje jej nieco bardziej współczesną interpretację behawioralną, inspirowaną pracami Quine'a i Sellarsa.

Marek Maciejczak omawia kontekst znaczenia językowego i zauważa, że znaczenie językowe należy rozpatrywać w jego związkach z systemem pojęć, procesami percepcyjnymi, doświadczeniem i działaniem w świecie. Według Maciejczaka bliższe wyjaśnianiu zagadnienia znaczenia są referencyjne koncepcje znaczenia, wypracowywane przez Davidsona, Kripkego, Montague, Hintikę, natomiast różne koncepcje niereferencyjne (Husserla, Fregego, Russella i wczesnego Wittgensteina) należy raczej odrzucić.

Filip Kawczyński stara się odpowiedzieć na pytanie, czy filozofia języka potrzebuje pojęcia sądu logicznego. Pojęcie sądu logicznego (ang. *proposition*) należy do kluczowych w filozofii języka, a jednocześnie do najbardziej kontrowersyjnych (łącznie z negowaniem ich istnienia). Kawczyński opowiada się za akceptacją istnienia sądów logicznych, zauważa też, że dyskusja dotycząca tego zagadnienia powinna odwoływać się również do badań empirycznych z zakresu kognitywistyki i psychologii.

Janusz Maciaszek uzasadnia i wyjaśnia niekognitywną teorię metafory Donalda Davidsona, odwołując się przy tym do jego teorii interpretacji oraz teorii działania. Maciaszek rozszerza analizę metafory na pozostałe wyrażenia niedosłowne, zwane tradycyjnie figurami retorycznymi, przeprowadza analizę mechanizmu rozpoznawania zdań niedosłownych w oparciu o teorię implikowania konwersacyjnego Grice'a oraz teorię interpretacji i racjonalności Davidsona. Dodatkowo proponuje kryterium rozpoznawania metafory w wyrażeniach interpretowanych niedosłownie oraz potraktowanie ironii jako wypowiedzenie dosłowne.

Joanna Odrowąż-Sypniewska poświęca swój rozdział debacie prowadzonej między minimalizmem semantycznym a kontekstualizmem. Omawiana debata jest niezwykle istotna dla współczesnych sporów dotyczących przebiegu granicy pomiędzy semantyką a pragmatyką, a co za tym idzie pytania o to,

czemu przysługują warunki prawdziwości i wartość logiczna, a także o zakres zależności kontekstowej.

Justyna Grudzińska proponuje nowe spojrzenie na referencję i kwantyfikację, koncentrując się na dynamicznym zwrocie w semantyce i jego filozoficznych implikacjach. Podczas gdy podejście statyczne utożsamia znaczenie zdania z jego warunkami prawdziwości, w podejściu dynamicznym główną rolę przestaje odgrywać pojęcie prawdziwości w modelu, natomiast centralnym pojęciem staje się „potencjał do zmiany kontekstu”. Grudzińska opowiada się za wypracowaniem teorii semantycznej pozwalającej na jak najlepsze rozpoznanie i przeanalizowanie subtelności rzeczywistości lingwistycznej.

Rozdział Mieszka Tałasiewicza poświęcony jest terminom jednostkowym (ang. *singular terms*). Terminy jednostkowe bywają kontrastowane z jednej strony z predykatami, z drugiej zaś z terminami ogólnymi, według Tałasiewicza mamy tu do czynienia z dwoma odrębnymi sensami, których nie należy mieszać.

Joanna Klimczyk przedstawia argumenty występujące w toczonym od wielu lat sporze o normatywność znaczenia. Klimczyk stara się pokazać, że dyskusja między semantycznymi normatywistami i antynormatywistami utknęła w martwym punkcie, ponieważ obie strony sporu odmiennie interpretują podstawową tezę semantycznego normatywizmu, która głosi, że znaczenie językowe jest wewnętrznie normatywne. Mimo nierozstrzygalności sporu możliwa jest ocena poszczególnych argumentów, co w konsekwencji powinno ograniczyć chaos pojęciowy.

Tadeusz Ciecierski i Katarzyna Kuś omawiają przyczynowe teorie nazw w kontekście argumentu zaproponowanego przez Stephena Sticha, pokazującego pewne kłopotliwe konsekwencje przemieszania odniesień nazw¹⁰.

Rozdział Elżbiety Chrzanowskiej-Kluczewskiej przedstawia najnowsze tendencje w filozoficznych podstawach teorii gier. Autorka omawia szerokie tło teorii gier społecznych i językowych, zwracając uwagę na konieczność interdyscyplinarnego spojrzenia na gry. Oprócz tego Chrzanowska-Kluczewska pokazuje, w jakim sensie teoriogrowy paradygmat strategiczny może konkurować z paradygmatem obliczeniowym w kognitywnym podejściu do języka.

Jarosław Jakielaszek kontrastuje i porównuje semantykę formalną (w tradycji Fregego i Montague) z Programem Minimalistycznym Noama Chomsky'ego. Autor omawia podstawowe założenia logicznej analizy języka i biolingwistycznego naturalizmu oraz wskazuje na możliwość uzgodnienia współczesnej semantyki formalnej z koncepcjami minimalizmu, podkreślając równocześnie fundamentalne różnice metodologiczne.

¹⁰ Rozdział ten ukazał się wcześniej w *Przeglądzie Filozoficznym — Nowa Seria*, 2010, 3 (75), s. 35–52. Autorzy i redaktor niniejszego tomu dziękują Joannie Odrowąż-Sypniewskiej i Justynie Grudzińskiej za zgodę na przedruk tekstu.

3. Zakończenie

Jarosław Jakielaszek, w ostatnim rozdziale tomu, na marginesie innych rozważań wspomina o filozoficznych implikacjach współczesnych teorii językoznawczych. Ten temat z pewnością zasługuje na dokładne opracowanie i będzie mu poświęcony jeden z przyszłych tomów serii *Studia z Metodologii i Filozofii Językoznawstwa*. W tym miejscu należy tylko zauważyć, że w najnowszej literaturze polskiej poświęcono już nieco uwagi filozoficznym implikacjom zmieniającego się paradygmatu generatywnego (zob. Willim 2010, 2012) i Tajsner (2012) oraz kontrowersjom w paradygmatach generatywnych i kognitywnych (zob. Lewandowska-Tomaszczyk 2008).

Powracając na zakończenie tego wstępu do motta zaczerpniętego z wykładów hiszpańskiego filozofa, José Ortegi y Gasset, można stwierdzić, że uprawianie filozofii języka to forma i sposób życia, jak bycie zakochanym, jak bycie damą.

Bibliografia

- ARENDT H. (2002), *Myślenie* (przeł. H. Buczyńska-Garewicz). Warszawa: Czytelnik.
- COMTE-SPONVILLE A. (2007), *Filozofia* (przeł. E. Burska). Warszawa: Instytut Wydawniczy PAX.
- DAVIES M. (2006), *Foundational Issues in the Philosophy of Language*, [w:] Devitt M., Hanley R. (red.), 19–40.
- DEVITT M., HANLEY R. (red.) (2006), *The Blackwell Guide to the Philosophy of Language*. Oxford: Blackwell Publishers.
- LEPORE E., SMITH B. C. (red.) (2006), *The Oxford Handbook of Philosophy of Language*. Oxford: Oxford University Press.
- LEWANDOWSKA-TOMASZCZYK B. (2008), *Czym jest język? Dzisiejsze kontrowersje w paradygmatach generatywnych i kognitywnych*, [w:] Stalmaszczyk P. (red.), 9–26.
- MACKENZIE I. E. (1997), *Introduction to Linguistic Philosophy*. Thousand Oaks: SAGE Publication.
- MAGEE B. (1998), *Popper* (przeł. P. Dziliński). Warszawa: Prószyński i S-ka.
- MARTINICH A. P. (2009), *General Introduction*, [w:] Martinich A. P. (red.), *Philosophy of Language. Critical Concepts. Volume 1*. London & New York: Routledge, 1–18.
- ORTEGA Y GASSET J. (2005), *¿Qué es filosofía?* (wydanie 19). Madrid: Revista de Occidente.
- PUTNAM H. (1999), *Pragmatyzm. Pytania otwarte* (przeł. B. Chwedeńczuk). Warszawa: Aletheia.
- PRECHTL P. (2008), *Wprowadzenie do filozofii języka* (przeł. J. Bremer). Kraków: Wydawnictwo WAM.
- SEARLE J. (2011), *Umysł. Krótkie wprowadzenie* (przeł. J. Karłowski). Poznań: Rebis.
- SOAMES S. (2010), *Philosophy of Language*. Princeton and Oxford: Princeton University Press.

- STALMASZCZYK P. (red.) (2008), *Metodologie językoznawstwa. Współczesne tendencje i kontrowersje*. Kraków: Wydawnictwo Lexis.
- STALMASZCZYK P. (red.) (2010), *Metodologie językoznawstwa. Filozoficzne i empiryczne problemy w analizie języka*. Łódź: Wydawnictwo Uniwersytetu Łódzkiego.
- STALMASZCZYK P. (red.) (2012), *Współczesne językoznawstwo generatywne. Podstawy metodologiczne* (seria *Studia z Metodologii i Filozofii Językoznawstwa* 1). Łódź: Primum Verbum.
- TAJSNER P. (2012), *Minimalizm z perspektywy biolingwistyki*, [w:] Stalmaszczyk P. (red.), 87–116.
- VENDLER Z. (1974), *Linguistics in Philosophy*. Ithaca–London: Cornell University Press.
- WACEWICZ S. (2008), *Ewolucja języka: horyzont metodologiczny*, [w:] Stalmaszczyk P. (red.), 27–42.
- WILLIM E. (2010), *O sporach wokół formy i funkcji we współczesnym językoznawstwie. Formalizm kontra funkcjonalizm?*, „Linguistica Copernicana” 3, 81–127.
- WILLIM E. (2012), *O przyczynach zmian głównych kierunków badawczych w gramatyce generatywnej Noama Chomsky’ego (1957–2007)*, [w:] Stalmaszczyk P. (red.), 18–86.
- WITTGENSTEIN L. (2000), *Dociekania filozoficzne* (przeł. B. Wolniewicz). Warszawa: Wydawnictwo Naukowe PWN.

Rozdział 1

Kontynentalna filozofia języka: wybrane stanowiska

Andrzej Pawelec (Uniwersytet Jagielloński)

Słowa kluczowe: fenomenologia, Gadamer, Heidegger, hermeneutyka, Husserl, kontynentalna filozofia języka, Merleau-Ponty, Ricouer

1. Wprowadzenie

Tytuł niniejszego rozdziału jest z kilku powodów problematyczny, co wymaga wyjaśnienia na wstępie. Ogólnie rzecz biorąc, w powszechnym użyciu znajdujemy rozróżnienie między filozofią „analityczną” a „kontynentalną”. Już ono może budzić zdziwienie, gdyż — jak zauważa polski filozof analityczny — nie jest „rozłączne”. Wymienione określenia należą jednak do odmiennych rodzajów („sposób uprawiania filozofii” i „lokalizacja”) więc z logicznego punktu widzenia nie są przeciwstawne. Nowaczyk wskazuje na rodowód tego rozróżnienia: „początki [filozofii analitycznej] umiejscawia się w Cambridge”, gdzie pracował Bertrand Russell i G. E. Moore, a w roku 1939, po rezygnacji Moore'a, katedrę filozofii otrzymał Ludwig Wittgenstein. Wiemy zaś skądinąd, że z perspektywy angielskiej reszta Europy to „kontynent” (por. Nowaczyk 2008: 11–12).

Sytuacja historyczna jest wszelako bardziej skomplikowana. Analityczne podejście do filozofii, w wersji węższej określane mianem neopozytywizmu bądź empiryzmu logicznego (por. Kołakowski 1966: rozdz. 8), praktykowano od początku wieku XX również na kontynencie. Wittgenstein, jako autor *Traktatu*, stał się w drugiej połowie lat 20. ważnym interlokutorem dla członków Koła Wiedeńskiego. I tak jak wielu z nich — co dotyczy również przedstawicieli szkoły lwowsko-warszawskiej (por. Woleński 1985) — musiał w latach 30. uciekać z hitlerowskiej Europy (większość trafiła do Stanów Zjednoczonych, gdzie po wojnie filozofia analityczna rozwijała się najprężniej). W rezultacie, jak zauważa (jednak z pewną przesadą) autor monografii na temat historii interesującego nas rozróżnienia, w Europie po śmierci Husserla „pozostał tylko

jeden ważny (*major*) filozof, Heidegger; nic więc dziwnego, że to, co nazywamy dziś ‘tradycją kontynentalną’, bierze nieodmiennie swój początek właśnie od niego”. Friedman konkluduje: „Całkowite intelektualne odseparowanie się od siebie obydwu tradycji — niemal totalny brak wzajemnego zrozumienia — wynika z przejścia władzy przez Hitlera w 1933 roku” (2000: 156-157; przeł. A. P.).

Nie mniej kłopotliwa jest cała fraza tytułowa, gdyż w powszechnym użyciu znajdujemy najczęściej bezprzymiotnikowe określenie „filozofia języka” (ang. *Philosophy of Language*), które oznacza jednak wyłącznie gałąź „analityczną”. Ten osobliwy stan rzeczy wynika zapewne z tego, że kwestia języka jest w tym nurcie tematem sztandarowym. Jak to ujął Kołakowski: „wspólną cechą [różnych kierunków filozofii analitycznej] było przekonanie, że właściwym powołaniem filozofii jest praca nad analizą języka — zarówno potocznego, jak i naukowego — w celu dokładnego wyjaśnienia sensu zastanych pojęć, twierdzeń i sporów; dzięki temu dojrzeje w myśli jasność dostateczna, aby kwestie tradycyjne stały się wreszcie rozstrzygalne naukowo lub odrzucone jako niesensowne” (1966: 187). Słowem, nazwa „analityczna filozofia języka” to w zwykłym użyciu pleonazm, który pojawia się jedynie w sytuacji szczególnej, a mianowicie w opozycji do „nieanalitycznej” filozofii języka, zwanej wtedy „kontynentalną”.

Odkładając na bok kwestie terminologiczne, wypada więc zapytać, czy zjawisko określone negatywnie — czyli jako „to wszystko, co nie jest analityczną filozofią języka” — da się w ogóle spójnie scharakteryzować. Jak wylicza Nowaczyk, do tej reszty należą: „poststrukturalizm, postmodernizm, dekonstrukcjonizm, feminizm, filozofia dialogu” (2008: 12), a jego listę z pewnością należałoby poszerzyć o centralne przeciwieństwa dla tradycji kontynentalnej kierunki filozoficzne: „fenomenologię” (zob. np. Michalski 1998: rozdz. 1) i „hermeneutykę” (zob. np. Przyłębski 2005).

Na koniec wreszcie, sama filozofia analityczna, będąca naszym negatywnym punktem odniesienia, nie jest zjawiskiem homogenicznym i znacząco zmieniała się w ponadstuletnim okresie swojego istnienia (por. Szubka 2009, Burge 1995). Symptomatycznym przypadkiem tego wewnętrznego zróżnicowania, a zarazem ewolucji stanowisk, jest *casus* Wittgensteina: „późny” Wittgenstein — autor *Dociekań* — zdystansował się wobec poglądów „wczesnego” Wittgensteina — autora *Traktatu* — i bywa zazwyczaj umieszczany przez komentatorów w jednym szeregu z filozofami „kontynentalnymi” (np. Taylor 1995). Można też odnieść wrażenie, że w ostatnich dekadach filozofia analityczna wytraciła impet i coraz częściej poszukuje inspiracji u autorów należących do tradycji kontynentalnej, lokowanych wcześniej na antypodach własnej myśli, np. Hegla czy Heideggera (por. Redding 2007). Powstają też prace zbiorowe, w których nie tylko kwestionuje się aktualność opozycji między obiema

tradycjami (np. Boundas 2007), ale mówi się już o filozofii „postanalitycznej” oraz „metakontynentalnej” (np. Reynolds et al. 2010).

Spróbujmy jednak — mimo wszystko — scharakteryzować podstawową różnicę dzielącą obydwie nurty filozoficzne, kierując się przy tym maksymą Iris Murdoch, krytycznej komentatorki filozofii analitycznej z lat 60. ubiegłego stulecia: „uprawianie filozofii to poznawanie własnego temperamentu, a zarazem próba odkrycia prawdy” (1996: 67). Kwestię temperamentu na plan pierwszy wysuwają filozofowie analityczni (np. Woleński 1995, Nowaczyk 2006), posiłkując się popularnym epigramem Kotarbińskiego o śledziu:

Do jasnych dążąc głębin — nie mógł trafić w sedno
Śledź pewien, obdarzony naturą wybredną.
Dokądkolwiek wędrował, zawsze nadaremno:
Tu jasno, ale płytko — tam głębia, lecz ciemno.

Choć w czterowierszu nie znajdujemy jednoznacznej waloryzacji obydwu podstaw czy „temperamentów”, to zarówno jego autor, jak i ci, którzy go cytują, opowiadają się po stronie „jasności” (niekoniecznie „płytkiej” ich zdaniem), motywując ten wybór względami intersubiektywnej komunikatywności. Celem tak uprawianej filozofii ma być dostarczanie klarownych i uzasadnionych sądów, gdyż tylko w ten sposób odpowiedzialnie poszerzamy krąg racjonalnych przekonań. W związku z tym, jak powiada Nowaczyk: „Filozoficzna głębia nie musi z konieczności być ciemna, a kiedy rozjaśnić się nie daje, to zapewne po prostu jej nie ma” (2006: 8).

W tym miejscu możemy już pokusić się o schematyczny — a więc uproszczony, jednowymiarowy — opis interesującej nas opozycji i zaryzykować następujące stwierdzenie: filozofowie analityczni zajmują się tym, co racjonalne w wymiarze „poziomym”, podczas gdy filozofowie kontynentalni — tym, co racjonalne w wymiarze „pionowym” (por. Morris 2007: 531). Skoro ci pierwsi skupiają się na uzasadnionych sądach, to znaczy, że poszukują — za pomocą środków językowych — adekwatnych pojęć do uchwycenia rzeczywistości. Mamy tu więc do czynienia z jakimś opisem właściwych relacji w ramach triady „język, myśl, rzeczywistość”, która pojawia się w rozmaitych wariantach w licznych tytułach dzieł, w tym również reprezentantów filozofii analitycznej (por. Whorf 1982, Putnam 1975). W tym kontekście zauważmy dwie rzeczy. Po pierwsze, elementy triady są tu traktowane jako byty — bezsporne dane — a problematyczne są jedynie relacje między nimi. Po drugie, język odgrywa tu rolę instrumentalną: ma stanowić jak najbardziej przejrzyste medium pośredniczące między myślą a światem. W gruncie rzeczy istotna jest struktura myśli, którą język miałby — w miarę możliwości — jednoznacznie odzwierciedlać.

Filozofowie kontynentalni kwestionują ten punkt wyjścia na wiele sposobów. Jednym z nich jest argumentacja Heideggera, iż „obecność” takich czy

innych bytów jest wtórna wobec bardziej pierwotnego sposobu bycia, a mianowicie „poręczności” (1994: 101). Rzeczy objawiają się najpierw nie jako po prostu „istniejące”, lecz jako korelat ludzkich (a również zwierzęcych) działań nastawionych na jakieś cele. Zauważmy, że nie wynika z tego, iż świat nie istnieje bez świadomości (czy raczej praktycznej, a nie czysto intelektualnej, aktywności, bo o niej mówi Heidegger). Bez eksplorującej aktywności rzeczywistość pozostaje „nieodkryta”. Prawda bądź „nieskrytość” (grecka *aletheia*) to dla Heideggera nie „pozioma” zgodność myślenia i rzeczywistości, wyrażona w jakimś języku, lecz „pionowe” wydobywanie się bytu z jego skrytości (1994: 309). Te „pionowe” dzieje, zauważmy, nie są Heideggerowską parafrazą kosmicznej ewolucji od Wielkiego Wybuchu przez kolejne stany złożoności materii aż po świadomość, tak jak ją opisuje nauka. Naturalistyczna wizja ewolucji Wszechświata to łańcuch bytów o wzrastającym stopniu komplikacji. Tymczasem Heideggera interesuje „pionowa” historia „sposobów bycia”, dzięki którym byty się pojawiają (słynna „różnica ontologiczna” pomiędzy „byciem” a „bytami”). Poza „poręcznością” i „obecnością” Heidegger wymienia jeszcze trzeci główny sposób bycia zwany *Dasein* (przedmiot analizy w *Byciu i czasie*), który przysługuje człowiekowi i językowi (por. Dreyfus 1991: 218).

Nie wchodząc w szczegóły koncepcji Heideggera, zauważymy, że jeśli sposób bycia człowieka — *Dasein*, egzystencja — opiera się na bardziej pierwotnym sposobie „odkrywania” rzeczywistości poprzez aktywną eksplorację, ujawniającą rzeczy jako „poręczne” (*zuhanden*), to nie ma mowy o osiągnięciu pełnej jasności co do rzeczy jako „obecnych” (*vorhanden*), gdyż każdy teoretyczny, bezinteresowny ogląd dokonuje się w horyzoncie uprzednich, tzn. bardziej źródłowych, postaw praktycznych, działań. Dany „horyzont” możemy refleksyjnie „odstaniać” (normalnie dostrzegamy tylko to, co się pojawia w jego ramach) i „przekraczać” (poprzez dyscyplinę, czujność czy podejrzliwość wobec własnych motywacji i wyników z jednej strony oraz ekspresję i iluminację z drugiej strony), ale nie jesteśmy w stanie go całkiem „zawiesić”, by dotrzeć do „rzeczy samych”. Najdalej idącą próbę tego rodzaju podjął Husserl, postulując wzięcie w nawias (*epoche*) wszelkich sądów na temat rzeczywistości i ściśle badanie tego, co się pojawia („fenomenologia”). W rezultacie pokazał, że każdy akt świadomości ma swój przedmiot („intencjonalność”), a dany przedmiot jest zawsze wyróżniony w ramach jakiegoś horyzontu. Te odkrycia nie zostały zakwestionowane przez Heideggera, ale gruntownie pogłębione: pojawianie się nie jest zasadniczo rezultatem aktów świadomości, lecz dokonuje się na głębszym podłożu postaw czy zaangażowania w rzeczywistość („bycie-w-świecie”). W rezultacie sens tego, co się pojawia, nie jest przejrzysty — wymaga interpretacji, sztuki rozumienia („hermeneutyka”). Można też powiedzieć, że sens nie jest dany bezpośrednio, lecz jest ugruntowany w szeregu zapośredniczeń, mediacji — ma naturę znakową („semiotyka”).

Powyższe „hasła” nie wyczerpują, rzecz jasna, obszarów eksploracji w obrębie kontynentalnej filozofii języka. Brak wśród nich odniesień do ujęć psychoanalitycznych, wychodzących od Freuda i Lacana (por. Lang 2005), a także do rozmaitych form „filozofii różnicy”, w pierwszym rzędzie Derridy i Deleuze'a. Celem niniejszego szkicu nie jest jednak prezentacja pełnego spektrum stanowisk, lecz pokazanie, na czym polega swoistość kontynentalnej filozofii języka — odmienność zadań, które przed sobą stawia, w relacji do celów filozofii analitycznej. Dlatego chciałbym omówić poniżej poglądy na temat języka trzech głównych (a zarazem najbliższych mi) filozofów „kontynentalnych”: Merleau-Ponty'ego, Gadamera i Ricoeura (opieram się na mojej prezentacji — Pawelec 2005: rozdz. 7). Dla większej przejrzystości wywodu traktuję tych filozofów jako reprezentantów, odpowiednio, podejścia fenomenologicznego, hermeneutycznego oraz semiotycznego. Takie postawienie sprawy jest nieco arbitralne, gdyż twórczość każdego z nich wykracza poza ramy jednego nurtu, a w przypadku Ricoeura jest świadomą „drogą okrężną” poprzez dialog z propozycjami wszelkiego rodzaju w poszukiwaniu stanowiska jak najbardziej „ekumenicznego” (np. hermeneutycznego poststrukturalizmu, jeśli chcielibyśmy używać etykiet).

2. Fenomenologia: Merleau-Ponty

Filozofię języka Merleau-Ponty'ego można uznać za komentarz do jednego, na pozór banalnego zdania: „słowo *ma sens*” (2001: 198; wyróżnienie w oryginale — A. P.). Dociekania autora *Fenomenologii percepcji* mają status fenomenologiczny, gdyż z jednej strony biorą w nawias teoretyczne ujęcia kwestii „znaczenia” językowego, z drugiej zaś — kierują się w stronę samego fenomenu „sensu”. Zaczniemy od aspektu negatywnego — od krytyki stanowisk dogmatycznych, odległych od doświadczenia. Jak powiada Merleau-Ponty: „Przez zwykłe stwierdzenie, że *słowo ma sens*, przekraczamy [...] zarazem intelektualizm i empiryzm” (tamże).

To pierwsze stanowisko jest fałszywe, bo „(myślenie) zapoznaje samo siebie i umieszcza się pośród rzeczy” (s. 40); to drugie odnosi wprawdzie „łatwe zwycięstwo” nad tym pierwszym (s. 166), ale również nie zdaje sprawy z prawdziwej natury doświadczenia. Empiryzm wymaga „hipotezy stałości” — trwałej korelacji między bodźcem, impresją i elementarną percepcją (s. 25) — gdyż bez tego założenia sypie się fundament, na którym można by wznosić budowlę poznania. Tymczasem już gestaltyści pokazali, że percepcja jest zjawiskiem „polowym”: zależnie od ich miejsca w „polu wizualnym” czy „słuchowym” identyczne bodźce wywołują odmienne percepcje, a bodźce istotnie różne — percepcje takie same. Z tego zaś wynika, że fizyczny bodziec czy atomowa impresja nie może stanowić fundamentu dla teorii percepcji.

Intelektualizm (bądź idealizm subiektywny) przyjmuje tę krytykę doświadczenia obiektywnego i w punkcie wyjścia umieszcza subiektywność, aktywność

podmiotową. Jest to jednak podmiot anonimowy — „transcendentalny” — który wyprzedza wszelkie doświadczenie, tzn. współtworzy je dzięki apriorycznej aparaturze pojęciowej (system kategorii). Taki podmiot może rozpoznać w doświadczeniu tylko to, co sam w nie włożył.

W rezultacie, w obydwu przypadkach percypujący pomiot jest „podmiotem próżnującym”. W wersji empirystycznej działają tylko mechanizmy łączenia wrażeń, podczas gdy w wersji intelektualistycznej istnieje gotowy świat myśli, która jedynie rozpoznaje siebie w doświadczeniu (s. 47). Dokładniej rzecz biorąc, dla empiryzmu „poznanie jawi się jako system zastąpień, w którym jedna impresja zapowiada inną impresję”, a znaczenie słów sprowadza się do wrażeń i efektów ich kojarzenia przez podobieństwo i styczność. A zatem słowa nie mają własnego sensu — są tylko elementem „gry obiektywnej przyczynowości” (s. 33, 197). Owa gra jest jednak mirażem, gdyż wskazane zasady kojarzenia nie mogą mieć mocy sprawczej. „Podobieństwo” nie może być przyczyną skojarzenia, gdyż ujawnia się *post factum* — kiedy zakwalifikujemy już coś do jakiegoś rodzaju — podczas gdy „styczność” w ogóle nie pozwala niczego trwale skojarzyć, jeśli doświadczenie traktujemy jako „wiązkę” impresji, w której mogą się ze sobą „stykać” dowolne wrażenia.

Z kolei intelektualizm postuluje przejrzystą dla siebie świadomość odcieraną od tworzywa, w którym się ona urzeczywistnia (s. 144). Również w tym przypadku słowo *nie ma* sensu, własnej mocy sprawczej, gdyż jest tylko „zewnątrznym znakiem wewnętrznego rozpoznania, które mogłoby dokonać się bez słów i do którego słowa nic nie wnoszą” (s. 197–8). A zatem w samej wypowiedzi słownej sens się nie materializuje, nie urzeczywistnia — ona tylko odsyła do sensu uprzedniego, rezultatu niezależnej operacji myślowej. Związek mówienia z myśleniem jest w tym ujęciu równie automatyczny (i niepojęty) jak w empiryzmie.

Merleau-Ponty pokazuje, że mowy i myśli nie można od siebie — źródłowo, istotowo — oddzielać. Nieprzypadkowo mówi o „sensie” mowy w związku z wieloznacznością tego słowa w języku francuskim: *le sens* może oznaczać „prąd” nurtu wody, „znaczenie” zdania, „stronę” tkaniny, „zmysł” powonienia (s. 431, przyp. tłum.). Słowa nie odsyłają więc, lecz ujawniają, są wyrazem jakiegoś nakierowania na świat, postawy przystającej do sytuacji komunikacyjnej (jak przód i tył tkaniny). W przeciwieństwie do „znaczenia”, które jest gotowe i przejrzyste (jako uprzednie przedstawienie czy dokonany akt myślowy), „sens” musi się rozwinąć, odsłonić. Mowa nie jest sposobem przekazywania skryzalizowanej myśli, ale jest obszarem rodzenia się myśli (s. 199). Należy podkreślić, że chodzi tu o „myśl twórczą” — dopiero szukającą tworzywa, w którym mogłaby się uformować. Merleau-Ponty stwierdza *expressis verbis*, że jego opis dotyczy „autentycznej mowy, która formułuje coś po raz pierwszy” (s. 199, przyp.), „mowy źródłowej — mowy dziecka, które wypowiada swoje

pierwsze słowo, mowy zakochanego, który odkrywa swoje uczucie, mowy 'pierwszego człowieka, który przemówił', albo pisarza i filozofa, którzy pod warstwami tradycji odkrywają doświadczenie pierwotne" (s. 200, przyp.). Idzie więc o „mowę wysławiającą” w odróżnieniu od „mowy wysłowionej” — codziennej „automatycznej” gadaniny (s. 217).

Jak przebiega ów proces autentycznej ekspresji, źródłowej artykulacji? „Egzystencja polaryzuje się tu w pewien 'sens', którego nie można zdefiniować przez żaden przedmiot naturalny, bo egzystencja usiłuje dotrzeć do samej siebie ponad bytem i dlatego tworzy mowę jako empiryczne oparcie dla własnego niebytu. Mowa jest wykraczaniem naszej egzystencji poza naturalny byt. Ale akt ekspresji ustanawia świat językowy i świat kultury, w którym to, co wykraczało poza byt, znowu pogrąża się w bycie” (s. 218). Mamy tu do czynienia z „doświadczeniem otwartym”, w którym nie umiemy się jeszcze odnaleźć. Nie chodzi tu po prostu o sytuację problemową — określoną na tyle dobrze, by można ją było rozwiązać znanymi środkami. Punktem wyjścia jest raczej wewnętrzne rozdwojenie („polaryzacja”), poczucie nieadekwatności tego, czym dysponujemy, a zarazem przecucie nowych możliwości. Udana ekspresja jest niczym nowy zmysł (*le sens*): otwiera nowe pole eksploracji dla naszego doświadczenia (s. 203).

Istotną rolę w tej wizji odgrywa pojęcie „motywacji”. W ujęciu materialistycznym dominuje myślenie przyczynowe — byty są powiązane łańcuchami kauzalnymi; z kolei w ujęciu idealistycznym kluczowe jest pojęcie „racji” — zasady konstytucji świata zjawisk (s. 68). Tymczasem strumień fenomenów kształtuje się dzięki „racji w działaniu”, czyli motywacji. To źródło formy doświadczenia ujawnili już gestaltyci: obiekty są potencjalnie wieloznaczne, zyskują określony sens dopiero w *akcie* percepcji. W rysunku dzbana, który podczas oglądania „przeistacza się” w rysunek dwóch twarzy widzianych z profilu, dostrzegamy po fakcie równorzędne „motywy”, które równie dobrze „uzasadniają” obie percepcje. Podobnie dzieje się w języku: istniejące formy są „podejmowane” w twórczych aktach ekspresji jako nośniki, „wehikuły” nowych treści ze względu na swój potencjalny „sens” — moc „nadawania kierunku” uwadze, sposób „przylegania” do sytuacji.

W tym fenomenologicznym ujęciu — aspirującym do opisu czystego doświadczenia — percepcja i ekspresja to ustalanie relacji już jakoś łączących nas ze światem, w których ujawnia się jego sens. Podstawę tego procesu stanowi ciało wyposażone w pewne zdolności motoryczne i rozwijające wciąż nowe: świadomość to najpierw „mogę”, a nie „myślę” (s. 157). Motoryczność nie jest po prostu mocą fizyczną, ale intencjonalnością — nakierowaniem na świat, postawą wobec niego, odpowiedzią na apel z jego strony. W tym kontekście Merleau-Ponty rozróżnia ruch konkretny, przylegający do swojego tła, oraz ruch abstrakcyjny, rozwijający swoje tło. Dzięki ciału podmiot ma zdolność

„projekcji” — rzutowania projektów motorycznych na świat, który tym samym nabiera nowych sensów (czy raczej „ujawnia” je dla aktywnego, tzn. zajmującego jakieś postawy podmiotu). Innymi słowy: „życie świadomości — życie poznające, życie pragnieniami lub życie percepcyjne — opiera się na ‘łuku intencjonalnym’, który rzutuje wokół nas naszą przeszłość, naszą przyszłość, nasze ludzkie środowisko, naszą sytuację fizyczną, naszą sytuację ideologiczną, naszą sytuację moralną, a raczej który sprawia, że jesteśmy usytuowani pod wszystkimi tymi względami” (s. 155).

Normalnie nie zauważamy, że mowa jest przejawem motoryczności, gdyż w zwyczajnej komunikacji słowa są przezroczystym nośnikiem sensów („przystają” do naszego świata jak prawa i lewa strona tkaniny). Niemniej jednak — jak wszelkie inne projekty motoryczne — słowa pozostają mocami. W przypadku „mowy wysłowionej” — potocznej, habitualnej — słowa są mocami zajmowania stanowiska w świecie znaczeń już ukonstytuowanych; w przypadku „mowy wysławiającej” krystalizują, artykułują sens na razie tylko przeczuwany. Innymi słowy, mowa jest gestem, a jej znaczenie światem (s. 205), którego horyzont próbujemy stale poszerzać czy odsuwać. A zatem sens świata odsłania się na kolejnych „szczepkach” działania. W przypadku percepcji „pojawia się” w takiej czy innej formie w ramach osobniczych projektów motorycznych. W przypadku pierwotnej ekspresji, czyli gestów, jest efektem zajęcia postawy w świecie wspólnych działań. W przypadku mowy jest wynikiem gestykulacji słownej w intersubiektywnym świecie znaczeń ukonstytuowanych w dziejach wspólnoty językowej.

3. Hermeneutyka: Gadamer

Rozważania Gadamera dotyczą w dużej mierze problematyki rozmowy i rozumienia. Autor *Prawdy i metody* umieszcza ją w kontekście „dzieł efektywnych” (1993: 285). Dla naszych celów wystarczy powiedzieć, że Gadamer podkreśla historyczność wszelkiego poznania. Sytuacja hermeneutyczna, w ramach której zadajemy jakieś pytania, jest ukształtowana przez nasze dzieje: żeby coś zrozumieć, musimy zajmować już jakieś stanowisko. Gadamer odrzuca za Heideggerem mit o bezałożeniowości poznania obiektywnego. Rozumienie nigdy nie jest bezpośrednie, zawsze jest jakoś usytuowane (ma swój horyzont, by przywołać Husserla).

Jednym ze sposobów rozwinięcia tej intuicji jest opis „kręgu hermeneutycznego” (s. 265 i nast.). Jeśli chcemy zrozumieć jakieś dzieło, musimy zarysować wstępny „projekt sensu”: musimy założyć, „o czym to jest”. Z reguły dokonuje się to samoczynnie dzięki „przesądom”, czyli „przed-sądom”, bezrefleksyjnym sposobom ujmowania danej sytuacji. Wbrew tradycji oświeceniowej nie można ich po prostu odrzucić i stanąć poza historią. Do „gry o rozumienie” można bowiem wejść tylko dzięki przesądom, które zarazem są stawką

w tej grze. Wszelkie autentyczne rozumienie jest bowiem samorozumieniem, polega na odsłanianiu własnych przesądów w konfrontacji z tekstem czy przekazem, który stawia im opór. Koło hermeneutyczne nie jest więc „błędnym kołem”, gdyż „przed-projekt” sensu może ulec modyfikacji, a przesady mogą zostać po części ujawnione i skorygowane.

Zastrzeżenie wskazujące na częściowość naszej samowiedzy jest istotne. Rozumienie dokonuje się bowiem w jakiejś „sytuacji”, a pojęcie to oznacza, że „nie znajdujemy się na zewnątrz niej i dlatego nie możemy mieć o niej żadnej przedmiotowej wiedzy” (s. 286). Sytuacja to „miejsce ograniczające możliwość widzenia”, a zarazem posiadające „horyzont”. Myślący podmiot może mieć „wąskie horyzonty”, może też je „poszerzać” albo „otwierać” nowe. Horyzont „oznacza wolność od ograniczenia do tego, co najbliższe, i możliwość wyglądania poza to. Kto ma horyzont, umie trafnie ocenić co do bliskości i dali, wielkości i małości znaczenie wszelkich rzeczy w obrębie tego horyzontu” (s. 287).

Jak przebiega proces rozumienia innego? Gadamer pokazuje najpierw, że wymóg „wstawienia się w cudze położenie” nie jest adekwatny. Z jednej strony, postulat ten nie jest wystarczający, gdyż nie prowadzi do autentycznej rozmowy, tzn. służącej porozumieniu. Jeżeli celem jest tylko ustalenie stanowiska i horyzontu rozmówcy, to taka rozmowa przypomina „rozmowę egzaminacyjną lub pewne formy prowadzenia rozmowy przez lekarza” (s. 287). Z drugiej strony, ów postulat fałszuje samą sytuację hermeneutyczną, gdyż zakłada istnienie odrębnych horyzontów; sugeruje, że trzeba porzucić własną sytuację, żeby wstawić się w cudzą. Tymczasem, w związku z pracą dziejów efektywnych, horyzonty „tworzą [...] wspólnie jeden wielki, poruszany od wewnątrz horyzont, który przekraczając granice współczesności ogarnia dziejową głębię naszej samoświadomości” (s. 289). Dlatego „nie istnieje horyzont współczesności sam dla siebie, tak jak nie istnieją horyzonty historyczne, do których można by docierać. *Raczej rozumienie jest zawsze procesem stapiania się takich rzekomo dla siebie istniejących horyzontów*” (s. 290, wyróżnienie w oryginale).

Możliwość rozumienia — opisywana powyżej w kategoriach „przesądów” i „horyzontów” — opiera się w sposób konieczny na języku. Tylko dzięki językowi możemy się znaleźć w sytuacji hermeneutycznej, gdyż dzięki językowi mamy „świat”, uwalniamy się od presji „otoczenia” (s. 403). Dzięki językowemu „opanowaniu” doświadczenia jego „groźna, wręcz unicestwiająca bezpośredniość zostaje jakby oddalona, rozłożona na czynniki, uczyniona przekazywalną i przez to poskromiona” (s. 411). Język *sans phrase* różni się pod tym względem od języka nauki. Ten drugi ma na celu „stawianie do dyspozycji i czynienie przedmiotem rachunku”, zakłada pełne uprzedmiotowienie (tamże). Tymczasem język potoczny (jak pamiętamy z opisu „horyzontu”) pozwala coś wyodrębnić z tła, umieścić w perspektywie. Język umożliwia dystans, ale nie obiektywizuje. Ogólnie rzecz biorąc, w języku wyłania się świat, który „jest

sam z siebie otwarty na każde możliwe wejście i tym samym na każde rozszerzenie obrazu świata, a odpowiednio do tego jest dostępny dla innych” (s. 406). Dlatego Gadamer może stwierdzić, że „uniwersalizm [języka] kroczy obok uniwersalizmu rozumu” (s. 368).

Jeszcze mocniej wyraża tę myśl bodaj najczęściej cytowane zdanie z *Prawdy i metody*: „Bytem, który może być rozumiany, jest język” (s. 429, wyróżnienie w oryginale). Przekład Barana może zapewne sugerować laikowi, że język znajduje się w opozycji do reszty bytu, który jest „sam w sobie” niepojmowalny. Sens powyższego zdania jest jednak daleki od takiej konstruktywistycznej interpretacji i wyraża go lepiej przekład Przyłębskiego: „byt, który można zrozumieć, jest językiem” (2006: 50). Innymi słowy, rzeczywistość, w miarę jak staje się coraz bardziej zrozumiała, nabiera charakteru językowego — jest „wyartykułowana”. Ta „pionowa” racjonalność — w odróżnieniu od „poziomej” zgodności myśli i rzeczy, wyrażonej w jak najbardziej „przejrzystym”, tzn. jednoznaczny języku — jest procesem złożonym i nie polega głównie na „obiektywizacji” czy „deskrypcji”. Jest to raczej wysiłek na rzecz „udomowienia” świata, wydobywania z niego całego potencjału, który pozwala nam czuć się w nim jak „u siebie”. Ta metaforyka sugeruje, że „świat” nie jest „dany” — nie jest tylko uniwersum obiektów w ruchu, które mamy za zadanie opisać. Świat się staje, przybiera zrozumiałe formy dzięki naszym artykulacjom.

4. Semiotyka: Ricoeur

Myśl Ricoeura zasadza się na dialogu z wieloma propozycjami intelektualnymi i zmierza do wyłuskania tego, co w nich prawomocne (por. Drwięga 1998). To hermeneutyczne podejście *par excellence* wyraża się, na przykład, w krytyce fenomenologii języka Merleau-Ponty'ego: „'Powrót do osobnika mówiącego', głoszony i zapoczątkowany przez Merleau-Ponty'ego na przedłużeniu ostatnich pism Husserla, został tak pomyślany, że przeskakuje się etap obiektywnej wiedzy o znakach i przechodzi od razu do słowa” (1985: 256). Zarzut Ricoeura nie jest całkiem sprawiedliwy — Merleau-Ponty, zwłaszcza w późniejszych pismach (por. 1999), nawiązał dialog z semiotyką Saussure'a — niemniej jednak dobrze pokazuje różnicę podejścia między fenomenologiem a hermeneutą. Dla Merleau-Ponty'ego „tajemnica pierwszego słowa nie jest większa niż tajemnica wszelkiej udanej ekspresji” (1999: 153), w związku z tym jego myśl nieustannie krąży wokół tej zagadki — fenomenu „genezy sensu”. Tymczasem dla Ricoeura liczy się przede wszystkim droga myśli. „Droga okrężna”, dodajmy, gdyż próby rozumienia danej domeny powinny uwzględniać to wszystko, co zostało ujawnione dzięki rozmaitym, z natury cząstkowym, ujęciom naukowym: „obiektywizującym” czy „metodycznym”. Zawsze bowiem trzeba pamiętać o ich nieuchronnym perspektywizmie: „świadomość prawomocności jakiejś metody jest nieodłączna od świadomości jej granic” (1985: 164).

W punkcie wyjścia hermeneutyki Ricoeura znajdujemy pojęcie „symbolu” — struktury znaczącej, która na skutek swojej nieprzejrzystości „daje do myślenia” (jak autor z upodobaniem powtarza za Kantem). Symbol jest objawieniem — „logosem uczucia, które pozbawione go pozostałoby nieokreślone, zatarte, nieprzekazywalne” (1985: 72). Oznacza to zarazem, że symbol jest zakotwiczony w doświadczeniu przedjęzykowym, przedpojęciowym, które z jakichś powodów domaga się — wzywa do — artykulacji. Na przykład symbole *sacrum* opierają się na odczuwanych „prawach” korespondencji między „kreacją *in illo tempore* i teraźniejszym porządkiem zjawisk natury i ludzkich działań”. Innymi słowy, symbole są wyrazem naszego kontaktu z mocami, który daje impuls do dalszych artykulacji (Ricoeur 1989: 139–147). Z czasem owe artykulacje mogą ulec „mumifikacji” — symbol stacza się wtedy do poziomu „obrazu, nazwy lub idola” (1985: 132).

Symbol ma zatem dwa wymiary: semantyczny, który dzięki udostępnionej przez symbol „nadwyżce znaczenia” stanowi obszar kolejnych artykulacji i interpretacji („daje do myślenia”), oraz niesemantyczny, którym jest więź egzystencji z bytem. W tym ujęciu źródłem człowieczeństwa byłoby poczucie kontaktu z mocami: człowiek rozpoznaje siebie np. w niebie i ziemi jako męskość i żeńskość, rozpoczynając tym samym nieskończoną serię odczytań motywowanych pragnieniem zajęcia „lepszego miejsca” w bycie bądź „udomowienia” świata (jak powiedzieliśmy wcześniej).

Powyższy ogólny opis stanie się jaśniejszy, jeśli pojęcie „symbolu” umieścimy najpierw w zwykłym kontekście „egzegezy” — biblijnej czy filologicznej interpretacji niejasnych tekstów. Ta praktyka zakładała już pewną teorię znaczenia: „jeśli tekst może mieć wieloraki sens, na przykład historyczny i duchowy, trzeba odwołać się do pojęcia znaczenia [...] złożonego” (1985: 183), czyli właśnie „symbolicznego”. Mianem symbolu Ricoeur określa „każdą strukturę znaczeniową, w której sens bezpośredni, pierwotny, dosłowny wyznacza ponadto inny sens pośredni, wtórny, przenośny, który nie może być inaczej ujęty niż poprzez ten pierwszy”. Interpretacja tekstu symbolicznego polega z kolei na „odszyfrowaniu sensu ukrytego w sensie widocznym, na rozwinięciu poziomów znaczeniowych zawartych w znaczeniu dosłownym” (s. 191–2).

Na tym poziomie analizy — na poziomie „mowy” bądź „semantyki” — hermeneutyka filozoficzna ma dwa podstawowe zadania. Po pierwsze, wskazać pełne spektrum form symbolicznych (Ricoeur wymienia sferę symboli *sacrum*, symboli marzeń sennych i symboli poetyckich) oraz określić relacje między nimi. Po drugie, hermeneutyka powinna opracować metodologię interpretacji (s. 192–3). W związku z naturą symboli — które są nieprzejrzyste (znaczenie dosłowne odsyła do znaczenia symbolicznego przez analogię), przygodne (tak jak kultury, do których należą) oraz zależne od niepewnych interpretacji — nie może istnieć „hermeneutyka ogólna”; istnieją jedynie „odrębne i przeciwstawne”

teorii interpretacji (s. 107). Ricoeur pokazuje na przykład, że fenomenologia religii interpretuje symbole sacrum jako odsyłające do przedmiotu religijnego, podczas gdy psychoanaliza opisuje dowolne symbole jako rezultat stłumionego pragnienia. Ta pierwsza należy do „hermeneutyk rekonstrukcji i rekolekcji”, ta druga — do „hermeneutyk podejrzania”. Obydwie są konieczne: podejście demaskatorskie pozwala oczyszczać pole symboliczne z fałszywych „idoli” („purytanizm symbolu”); podejście „rekolekcyjne” pozwala powrócić do przedmiotu wskazywanego przez symbole (s. 133–4). Ostatecznie celem hermeneutyki na poziomie semantycznym jest uzasadnienie każdej metody interpretacji we właściwych jej teoretycznych granicach (s. 194).

Analizy na poziomie semantycznym stanowią tylko niezbędny etap przygotowawczy. Pozostając na płaszczyźnie „mowy”, zamykamy się bowiem w świecie „gier językowych”, absolutyzujemy język. W ten sposób „negujemy podstawową intencję znaku: odnosić się do czegoś, a więc przekraczać siebie i zatajać w tym, do czego się odnosi” (s. 195). Mowę symboliczną należy zatem powiązać z egzystencją, która próbuje dotrzeć do siebie przez ekspresję (jak pisał Merleau-Ponty). Ten etap rozumienia Ricoeur nazywa „refleksyjnym”, gdyż łączy on rozumienie znaków z rozumieniem siebie (s. 195). Powstaje pytanie: dlaczego do siebie można dotrzeć tylko drogą okrężną, poprzez interpretację pojmowanych szeroko „wypowiedzi” innego (działań, tekstów, dzieł symbolicznych)? Ricoeur wskazuje dwa powody: po pierwsze, czysta refleksja jest ślepa, jeśli pomija ekspresje, w których życie się obiektywizuje; innymi słowy, akt istnienia można przyswoić sobie jedynie „za pośrednictwem krytyki zastosowanej do dzieł i do aktów, które są znakami tego aktu istnienia” (s. 196); po drugie, refleksja nie może być bezpośrednia, gdyż w punkcie wyjścia świadomość jest z reguły fałszywa — trzeba ją wpierw zdemaskować, co postulują „hermeneutyki podejrzania” (s. 197).

Wszelkie autentyczne rozumienie jest więc — jak pokazywał Gadamer — samorozumieniem, prowadzi do przemiany świadomości. Zarazem jest wezwaniem do poprawy własnego położenia w bycie, gdyż symbol nie jest tylko bramą do samowiedzy, ale przede wszystkim jest wyrazem naszej więzi z bytem (s. 74). A zatem dzięki refleksji możemy próbować wspiąć się na poziom egzystencji. „Ruch refleksji” — w związku z konfliktem hermeneutyk — prowadzi wszakże do odsłonięcia egzystencji w sposób częściowy i przeciwstawny. Na przykład psychoanaliza odkrywa egzystencję w postaci pragnienia, które wyprzedza i fałszuje świadomość. Mamy tu więc do czynienia z „archeologią” podmiotu. Z kolei w Hegłowskiej fenomenologii ducha refleksja prowadzi nie w stronę pierwotnego pragnienia, lecz w kierunku egzystencji pełnej, pojednanej z sobą („teleologia” podmiotu). W tym ujęciu egzystencja staje się sobą — „ludzką i dojrzałą” — gdyż przyswaja sobie sens, „który przebywa najpierw ‘na zewnątrz’, w dziełach, instytucjach, pomnikach kultury, w których

obiektywizuje się życie ducha” (s. 201). A wreszcie w przypadku fenomenologii religii ujawnia się sacrum, które ustanawia egzystencję w sposób absolutny — mamy tu do czynienia z „eschatologią” egzystencji (s. 202).

Podsumowując, „semiotyczne” podejście Ricoeura mówi o tym, jak egzystencja dociera do siebie samej poprzez symboliczne zapośredniczenia, nawarstwione w dziejach kultury.

5. Zakończenie

Zaproponowana powyżej lektura kilku prac trzech głównych filozofów „kontynentalnych” nie daje oczywiście wyobrażenia ani o całości ich drogi myślowej, ani tym bardziej o stanie „kontynentalnej filozofii języka”. Może jednak wystarczy, aby pokazać odmienność postaw czy temperamentów przedstawicieli obydwu głównych nurtów filozofii współczesnej.

Wypada jeszcze wrócić na koniec do najważniejszego elementu wypowiedzi Iris Murdoch: „uprawianie filozofii to [...] próba odkrycia prawdy”. Gdyby wierzyć cytowanemu wcześniej polskiemu filozofowi analitycznemu, to Merleau-Ponty, Gadamer i Ricoeur przemawiają zapewne „z wnętrza metafor”, a ich myślenie w znacznej części „nie ma aspiracji poznawczych” (Nowaczyk 2006: 8). Tymczasem myślenie dyskursywne — zmierzające do odsłonięcia prawdy — postępuje drogą „literalizacji metafor”, czyli jak najbardziej dosłownych parafraz. Wierny ideałowi przejrzystości sensu Nowaczyk formułuje analityczne *credo*:

Co w ten sposób możemy osiągnąć? Nie sędzę, aby tą drogą można było ostatecznie przekształcić filozofię w dyscyplinę formalną na wzór matematyki bądź w jedną z nauk przyrodniczych. Natomiast wierzę, że jest to droga pozwalająca tym wypowiedziom filozoficznym, którym przypisuje się walory poznawcze, nadać taką postać, która sprawi, że mogą one tworzyć harmonijną całość z naszymi przekonaniami zdroworozsądkowymi i naukowymi, wchodząc z nimi w związki logiczne. Czy stan taki jest osiągalny i czy warto do niego zmierzać? Sędzę, że tak. Wielu filozofów, zarówno dawnych jak i współczesnych, stara się nie popadać w konflikt z aktualnym stanem wiedzy naukowej, o której słusznie się powiada, że jest krytycznym przedłużeniem zdrowego rozsądku. Wręcz przeciwnie, filozofowie ci starają się zbudować pewien ogólny schemat pojęciowy, w którym byłoby miejsce na wszelkie przekonania uchodzące za racjonalne. Ponadto stawiają pytania, które po odpowiedniej konkretyzacji mogą stać się przedmiotem dociekań naukowych lub dociekania takie inspirować (Nowaczyk 2006: 9–10).

Szanując rzetelność filozofa, który jasno określa swój horyzont badawczy, chciałbym jednak zadeklarować na koniec, że w imię prawdy (jako „nieskrytości”) należy ten horyzont przekroczyć i pójść dalej. Mam nadzieję, że przedstawione w tym szkicu poglądy na naturę języka i rzeczywistości deklarację tę uzasadniają.

Bibliografia

- BOUNDAS V. C. (red.), (2007), *The Edinburgh Companion to Twentieth-Century Philosophies*. Edinburgh: Edinburgh University Press.
- BRONK A. (red.) (1995), *Filozofować dziś. Z badań nad filozofią najnowszą*. Lublin: TN KUL.
- BURGE T. (1995), *Filozofia języka i umysłu (1950–1990)*, [w:] Bronk A. (red.), 93–139.
- DREYFUS H. (1991), *Being-in-the-World. A commentary on Heidegger's "Being and Time"*, Division I. Cambridge, Mass.: MIT Press.
- DRWIĘGA M. (1998), *Ricoeur daje do myślenia*. Bydgoszcz: Homini.
- GADAMER H.-G. (1993), *Prawda i metoda* (przeł. B. Baran). Kraków: inter esse.
- HEIDEGGER M. (1994), *Bycie i czas* (przeł. B. Baran). Warszawa: PWN.
- KOŁAKOWSKI L. (1966), *Filozofia pozytywistyczna: od Hume'a do Koła Wiedeńskiego*. Warszawa: PWN.
- LANG H. (2005), *Język i nieświadomość* (przeł. P. Piszczatowski). Gdańsk: słowo/obraz terytoria.
- MERLEAU-PONTY M. (1999), *Proza świata* (przeł. różni). Warszawa: PIW.
- MERLEAU-PONTY M. (2001), *Fenomenologia percepcji* (przeł. M. Kowalska, J. Migasiński). Warszawa: Aletheia.
- MICHAŁSKI K. [1978] (1998), *Heidegger i filozofia współczesna*. Warszawa: PIW.
- MORRIS D. (2007), *Philosophy of mind*, [w:] Boundas V. C. (red.), 531–544.
- MURDOCH I. (1996), *Prymat dobra* (przeł. A. Pawelec). Kraków: Znak.
- NOWACZYK A. (2006), *Poławianie sensu w filozoficznej głębi*. Łódź: Wydawnictwo Uniwersytetu Łódzkiego.
- NOWACZYK A. (2008), *Filozofia analityczna*. Warszawa: PWN.
- PAWELEC A. (2005), *Znaczenie ucieleśnione. Propozycje kręgu Lakoffa*. Kraków: Universitas.
- PRZYŁĘBSKI A. (2005), *Hermeneutyczny zwrot w filozofii*. Poznań: WNUAM.
- PRZYŁĘBSKI A. (2006), *Gadamer*. Warszawa: Wiedza Powszechna.
- PUTNAM H. (1975), *Mind, language and reality. Philosophical Papers vol. 2*. Cambridge: Cambridge University Press.
- REDDING P. (2007), *Analytic Philosophy and the return of Hegelian thought*. Cambridge: Cambridge University Press.
- REYNOLDS J. et al. (red.) (2010), *Postanalytic and metacontinental. Crossing philosophical divides*. London: Continuum.
- RICOEUR P. (1985), *Egzystencja i hermeneutyka* (przeł. różni). Warszawa: PAX.
- RICOEUR P. (1989), *Język, tekst, interpretacja* (przeł. P. Graf, K. Rosner). Warszawa: PIW.
- SZUBKA T. (2009), *Filozofia analityczna. Koncepcje, metody, ograniczenia*, Wrocław: WUW.
- TAYLOR C. (1995), *Philosophical arguments*. Cambridge, Mass.: Harvard University Press.
- WHORF B. (1982), *Język, myśl i rzeczywistość* (przeł. T. Hołówka). Warszawa: PIW.
- WITTGENSTEIN L. (1970), *Tractatus logico-philosophicus* (przeł. B. Wolniewicz). Warszawa: PWN.
- WITTGENSTEIN L. (1972), *Dociekania filozoficzne* (przeł. B. Wolniewicz). Warszawa: PWN.
- WOLEŃSKI J. (1985), *Filozoficzna szkoła lwowsko-warszawska*, Warszawa: PWN.
- WOLEŃSKI J. (1995), *Sens i nonsens w filozofii*, „Znak” 6 (481), 12–24.

Rozdział 2

Formy symboliczne Ernsta Cassirera a gramatyka kognitywna

Henryk Kardela (Uniwersytet Marii Curie-Skłodowskiej w Lublinie)

Słowa kluczowe: Ernst Cassirer, Ronald Langacker, Bronisław Malinowski, formy symboliczne, konstruktywizm komunikacyjny, subiektyfikacja, radykalna subiektyfikacja, wirtualny dokument, poznanie zaangażowane, poznanie oderwane, ucieleśnienie

1. Wstęp

Według Ronalda Langackera (2008/2010 — dalej: Langacker), aby zrozumieć, czym jest gramatyka, musimy wziąć pod uwagę dwa ściśle powiązane z gramatyką aspekty poznania. Są to (i) sensoryczne i motoryczne interakcje związane z tworzeniem się neurologicznie pobudzanych sekwencji czasowych leżących u podstaw poznania oraz (ii) sfera społeczna i kulturowa, w ramach której rozwija się poznanie.

Jeśli idzie o sferę społeczną i kulturową jest ona, zdaniem Langackera, „równie rzeczywist[a] i istotn[a] jak same wydarzenia fizyczne” (s. 666), ponieważ

wchodzimy [...] w interakcje ze światem za pośrednictwem pamięci, antycypacji, przewidywań, uogólnień i kontemplacji alternatywnych możliwości. Zjawiska te, które odwołują się do wielu odmian i poziomów konstrukcji mentalnej, wykraczają poza bezpośrednie doświadczenie cielesne. Ale i one są w nim zakorzenione.

Wspomniane w powyższym cytacie „wykroczenie zjawisk poza bezpośrednie doświadczenie cielesne” powstaje w rezultacie oddzielenia się poznania od bezpośredniego doświadczenia fizycznego; jest to wynik procesu mentalnego, w trakcie którego, jak czytamy (s. 714), „akt *zaangażowanego poznania* (*engaged cognition*), gdzie osoba wchodzi w bezpośrednią interakcję z czymś w świecie na poziomie fizycznym”, ustępuje miejsca tzw. *poznaniu oderwanemu* (*disengaged cognition*), w którym „niektóre aspekty przetwarzania zachodzące w mózgu

„zaczynają pojawiać się autonomicznie, przy nieobecności jakiegokolwiek bieżącej interakcji ze światem”.

Prawie sto lat wcześniej w studium *Philosophie der symbolischen Formen: Die Sprache* (1923), przełożonym na angielski jako *The Philosophy of Symbolic Forms: Language*, Ernst Cassirer tak ujmował wspomnianą przez Langackera interakcję świata doznań i języka, pozwalającą człowiekowi „wykroczyć poza bezpośrednie doświadczenie cielesne” (Cassirer 1955: 198 — dalej Cassirer)¹:

W językoznawstwie, podobnie jak w epistemologii nie jest możliwy podział na to, co zmysłowe i to, co umysłowe, na dwie niezależne od siebie sfery, z których każda odwoływałaby się do własnej samowystarczalnej modalności rzeczywistości. [...] Poznanie zmysłowe nigdy nie jest nam dane w czystej, „przed-formacyjnej” postaci — podstawowe spostrzeżenie zawsze odwołuje się do formy przestrzeni i czasu. Lecz w ustawicznym postępie wiedzy nieokreślone to początkowo odwołanie się nabiera cech bardziej konkretnych: sama możliwość zestawienia i następstwa owych sfer rozwija się w całościowość przestrzeni i czasu — w porządek, który jest tak jednostkowy, jak uniwersalny. Można sądzić, iż język jako manifestacja ducha także powinien odzwierciedlać w jakiś sposób ten fundamentalny proces. [...] Nawet w najbardziej abstrakcyjnych kategoriach języka można odkryć związek między kategoriami a prymarną intuicyjną podstawą, w której są zakorzenione. I znowu, „znaczenie” nie różni się tu od tego, co zmysłowe — obydwie te sfery się przenikają. A zatem przejście ze świata zmysłów do świata „czystej intuicji”, które to przejście epistemologia uważa za konieczny czynnik konstytutywny „ja” i czystego pojęcia rzeczy, ma dokładny odpowiednik w języku. [...] tylko poprzez medium intuicyjnych form, przez intuicje dotyczące przestrzeni, czasu i liczby język jest w stanie wykonać zasadniczo logiczne operacje — przekształcić doznania zmysłowe w reprezentacje.

Niniejszy artykuł jest próbą ustalenia perspektywy badawczej, która pozwoliłaby znaleźć punkty styczne między Cassirerowską filozofią form symbolicznych — filozofią „przekształceń doznań zmysłowych w reprezentacje” — i Langackerowską koncepcją *poznania oderwanego*, tzn. poznania pozwalającego człowiekowi „przekroczyć bezpośrednie cielesne doświadczenie przy nieobecności jakiegokolwiek bieżącej interakcji ze światem”.

Wydaje się, że najważniejszą perspektywą badawczą — i taką też przyjmujemy dla potrzeb naszej dyskusji — będzie *konstruktywizm komunikacyjny* — postawa metodologiczna, definiowana przez Michała Wendlanda (2011: 64–65) jako

zbiorcz[a] nazw[a] stanowisk filozoficznych (i językoznawczych), w myśl których podmiot (człowiek, ale głównie podmiot zbiorczy jako określona zbiorowość ludzka) (1) konstruuje wiedzę o świecie i (2) aktywnie uczestniczy w konstruowaniu świata, w którym żyje. Użycie przymiotnika „komunikacyjny” ma natomiast wskazywać,

¹ Cytaty w niniejszym tekście — o ile nie podano inaczej — pochodzą z anglojęzycznej wersji *The Philosophy of Symbolic Forms: Language*; przeł. H. K.).

że (1) procesy konstruowania wiedzy o świecie i samej rzeczywistości społecznej przebiegają głównie w ramach działań komunikacyjnych, a także, że (2) istnieje możliwość aplikacji konstruktywizmu jako perspektywy badawczej do rozważań nad komunikacją w ramach nauk społecznych.

Konstruktywizm komunikacyjny zakłada — precyzuje Wendland (s. 65) —

uprzedniość (pierwotność) aktywności podmiotu względem rzeczywistości (lub podważa zasadność, a przynajmniej efektywność dualistycznego podmiotowo-przedmiotowego modelu poznania, który określany jest także jako obiektywistyczny model poznania). W wersji słabszej zakłada on, że wiedza o świecie konstruowana jest przez człowieka, w wersji silnej zaś — że w ogóle nie można mówić o „obiektywnej” rzeczywistości poza podmiotem i jego działalnością (albo: podmiotowością zbiorową wyznaczaną przez działania komunikacyjne). Zarówno w słabszej, jak i silnej wersji perspektywa konstruktywistyczna przyjmuje nieuniknione uwikłanie wiedzy o świecie w kulturę albo — wytwarzanie rzeczywistości społecznej w ramach aktywności kulturowej człowieka.

2. Filozofia form symbolicznych

Filozofia Cassirera jest próbą ujęcia związku istniejącego między podstawowymi danymi w naszym postrzeganiu — prymarnymi przejawami „reprezentacji świata” a wyższego rzędu kulturowo zdeterminowanym znaczeniem, które powstaje w oparciu o owe prymarne formy w trakcie dialektycznych procesów rozwojowych.

Cassirer przeprowadza ważne rozróżnienie między substancją a formą — istotne z punktu widzenia dwóch podejść do relacji między pojęciem a rzeczą, jakie zasadniczo rozróżnia się w literaturze przedmiotu. Według tzw. podejścia *arystotelesowskiego* istnieje bezpośredni związek między przedmiotem a pojęciem, w tym sensie, że pojęcia reprezentują rzeczy istniejące w świecie rzeczywistym. Jak zauważa Sebastian Luft (2005), w tym ujęciu zarówno pojęcie, jak i rzeczy są bytami *substancjalnymi*. Pojęcia są bytami substancjalnymi, ponieważ odnoszą się do rzeczy, które istnieją w świecie realnym. Lecz także i rzeczy są substancjalne: istnieją samoistnie w tym sensie, że ich istnienie nie zależy od istnienia innych rzeczy.

Innym sposobem konceptualizowania rzeczy jest ujęcie relacji „rzecz — pojęcie” w kategoriach *funkcji*. I to właśnie ujęcie leży u podstaw form symbolicznych Ernsta Cassirera. Weźmy dla przykładu pojęcie „neutrino”. Neutrino nie jest pojęciem substancjalnym, lecz raczej pojęciem funkcjonalnym w tym sensie, że nabiera ono znaczenia tylko w kontekście dopuszczającym istnienie takich pojęć, jak „neutrino”, „elektron”, „cząstka gamma” itd. Pojęcie „neutrino” nie jest pozbawione sensu, ponieważ jest „widoczne” w ramach pewnych serii funkcjonalnych — w tym przypadku w ramach teorii cząstek elementarnych.

Realność neutrino nie jest zatem realnością substancjalną — neutrino to konceptualny schemat funkcjonujący w teorii cząstek elementarnych.

W istocie nie tylko takie pojęcia jak neutrino są pojęciami funkcjonalnymi, wszystkie reprezentacje przedmiotów tzw. „życia codziennego” mają charakter funkcjonalny; są one, jakby to ujął Cassirer, w stanie *symbolicznego nasycenia* (*symbolische Pragnanz*), ponieważ są widoczne tylko w pewnym świetle — w kontekście znaczeń związanych z językiem, religią, mitem, nauką czy sztuką. To, co doświadczamy, stwierdza Luft za Cassirerem, to „coś, co objawia się nam jako reprezentant pewnej formy nasyconej znaczeniem”. Dla Cassirera, Luft konstatuje, „każda rzecz ma swoje własne szczególne nasycenie, nasycenie, które jest zróżnicowane i zależy od formy, w jakiej jest postrzegane. Linia falista na przykład może być sinusoidą albo też artystycznym ornamentem”, czyli — mówiąc krótko — linia falista jest formą symboliczną.

Co to jest forma symboliczna? — Luft cytuje tu Cassirera (Luft 2005: 7 — przekład H. K; por. także Cassirer 2004: 21):

Pod terminem formy symbolicznej będziemy rozumieli jakąkolwiek energię ducha, poprzez którą duchowa treść znaczeniowa łączy się z konkretnym zmysłowym znakiem i jest z tym znakiem ściśle związana.

Czyli forma symboliczna

jest energią ducha wytwarzającą „przestrzeń znaczenia”, w której się znajdujemy wraz z tym, czego doświadczamy i co ma szczególny sens. [...] Forma symboliczna jest warunkiem możliwości doświadczenia. Jest to transcendentalna forma intuicji. [...] Nie mówimy tu o rzeczy, jaką doświadczamy, lecz o naszym doświadczeniu rzeczy, ze szczególnym naciskiem, iż to doświadczenie jest już schematyczne na wiele różnych sposobów. Formy symboliczne nie są „kontekstami” pochodzącymi z empirycznych danych, lecz formami apriorycznymi, jakie wypełnia życie (Luft, 8–9).

Cassirer wyróżnia trzy rodzaje znaczenia symbolicznego. Najbardziej podstawowym rodzajem znaczenia symbolicznego, leżącym u podstaw tzw. „mitycznej świadomości”, jest *funkcja wyrażania*. Funkcja ta dotyczy zdarzeń zachodzących w świecie takim, jakiego doświadczamy. Charakterystyczną cechą „świadomości mitycznej” jest brak podziału między pozorem a rzeczywistością. Świat mityczny — czytamy w *Stanford Encyclopedia of Philosophy* (<http://plato.stanford.edu/entries/cassirer/> (przeł. H. K.).

nie składa się ze stałych niezmiennych substancji, jakie objawiają się nam z różnych punktów widzenia i przy różnych okazjach, lecz raczej w przelotnych wiązkach zdarzeń powiązanych ze sobą emocjonalną i uczuciową charakterystyką „fizjonomiczną”. Świat mityczny przejawia swój własny rodzaj przyczynowości, gdzie każdy element przyczynowości jest już całością i wywiera efekt przyczynowości na całość. Nie ma tu istotnej różnicy pomiędzy żyjącymi i martwymi, między świadomym doświadczeniem a doświadczeniem danym we śnie, pomiędzy

nazwą rzeczy a samą rzeczą, i tak dalej. Fundamentalne kategorie kantowskie przestrzeni, czasu, substancji (lub rzeczy) i przyczynowości tworzą swoistą konfigurację reprezentującą aprioryczną formalną strukturę myśli mitycznej.

Z owych „przelotnych wiązek zdarzeń powiązanych ze sobą emocjonalną i uczuciową charakterystyką fizjonomiczną” wywodzi się drugi rodzaj znaczenia symbolicznego — znaczenie związane z *funkcją przedstawiania*. Bazując na naszej intuicji przestrzeni i czasu, konstruujemy nasz intuicyjny świat „zwykłego postrzegania zmysłowego”. W konsekwencji powstaje przestrzennie-czasowy porządek, w którym każda postrzegana rzecz przyjmuje określone położenie względem mówiącego, w stosunku do jego punktu widzenia i w relacji do szeregu potencjalnych pragmatycznych czynności, jakie mówiący podejmuje. Możemy rozróżnić teraz stałą *substancję-rzecz* z jednej strony i *pozór* z drugiej.

Rozróżnienie między pozorem a rzeczywistością, stwierdza Cassirer, wyrażone przy pomocy propozycjonalnej prawdziwościowej formuły językowej „X jest Y”, prowadzi „dialektycznie” ludzką myśl do wypełnienia nowego zadania — do systematycznego naukowego badania prawdy istniejącej w świecie. W tym wypadku mamy już jednak do czynienia z trzecią, *sensotwórczą* funkcją znaczenia. Odzwierciedla ona „czyste relacje kategorialne”, które wyzwolone z intuicji zmysłowych ujawniają się w badaniach naukowych — są one precyzyjnie ujmowane przez prawa matematyki, logiki i fizyki.

Jeśli idzie o formę symboliczną, to jej siłą sprawczą, formującą, jest wspomniana już „energia ducha”. Dzięki tej energii dane zjawisko nabiera określonego sensu, staje się „symbolicznie nasycone”. Zastosowanie pojęcia „symbolicznego nasycenia” lub — jak to ujmuje Krystyna Świąćicka (1996) — „symbolicznej pregnancji”,

umożliwia uniknięcie trudności związanych z wyjaśnieniem sposobu, w jaki wrażenie zostaje powiązane ze znaczeniem. Wszystko, z czym możemy mieć do czynienia, jest już jakoś uporządkowane, ma jakiś sens, jest niejako naładowane znaczeniem. To właśnie relacja między przeżyciem zmysłowym i zawartym w nim znaczeniem jest tym, co pierwotne.

Do powiązania wrażenia sensualnego ze znaczeniem powrócimy za chwilę, proponując „kognitywne spojrzenie” na „Cassirerowską metodę” — teraz przyjrzyjmy się nieco bliżej wspomnianemu przez Świąćicką pierwotnemu charakterowi relacji między przeżyciem sensualnym i zawartym w nim znaczeniem. Nawiążemy tu z jednej strony do Cassirerowskiego pojęcia *obiektywizacji* postrzegania zmysłowego, z drugiej do teorii języka Bronisława Malinowskiego — koncepcji „mowy w działaniu”, która — jak to ujmuje Malinowski — „przenosi słowa w określoną formę gramatyczną” i „samodzielną funkcję znaczenia”.

3. Obiektywizująca funkcja języka: Cassirer i Malinowski

Przyjmując poglądy Wilhelma von Humboldta na język i — ogólnie — na relację język — rzeczywistość, Ernst Cassirer stwierdza (Cassirer 1955: 156), iż według Humboldta

język reprezentuje opozycje między jednostką a „obiektywnym” duchem i jego działaniami. Każdy człowiek mówi własnym językiem — a jednak, pomimo nieograniczonego sposobu, w jakim go używa, jest świadomy jego wewnętrznego uwarunkowania. Język funkcjonuje zawsze jako ogniwo pośredniczące — pośredniczy najpierw między znaną i nieogarnioną przyrodą, następnie między jednym a drugim człowiekiem, umożliwia ów związek i jednocześnie wyrasta z niego.

Ową „mediacyjną” funkcję języka, „pośredniczącą między znaną i nieogarnioną przyrodą, następnie między jednym a drugim człowiekiem” najlepiej jest — zdaniem Cassirera — prześledzić, porównując relacje między dźwiękiem i znaczeniem w językach ludów prymitywnych i w językach rozwiniętych kulturowo. O ile w językach ludów prymitywnych wydaje się istnieć naturalny związek między dźwiękiem i znaczeniem, o tyle w językach rozwiniętych kulturowo związek ów zanika. Cassirer (2004: 25–26) cytuje tu D. H. Westermanna, który w swojej gramatyce języka Ewe stwierdza, iż Ewe

jest niezwykle bogaty w środki dla oddania za pomocą dźwięku odbieranego wrażenia: bogactwo wypływające z nieprzezwyciężalnej niemal chęci do naśladowania wszystkiego, co zasłyszane, widziane, w ogóle odczute, do oddania tego wszystkiego za pomocą jednego lub więcej dźwięków. Tu oraz w kilku pokrewnych językach istnieją na przykład przysłówki, które opisują tylko *jedną* czynność, *jeden* stan lub *jedną* właściwość (kursywa oryginalna) i które w związku z tym przynależą i mogą być związane tylko z jednym czasownikiem. Westermann przytacza dla jednego czasownika *chodzenie* nie mniej niż 33 tego rodzaju przysłówkowe obrazy dźwiękowe, z których każdy określa każdorazowo jeden szczególnie sposób, szczególnie odcień i swoistość chodzenia.

Cassirer kontynuuje (s. 26):

Jak widać, wyraz językowy nie oddzielił się tu jeszcze od czysto mimicznego i nie posiada w opozycji do niego jakiejś wyższej formy ogólności. Ów mimiczny charakter języków ludów prymitywnych występuje szczególnie w całej pełni zróżnicowanych form wyrazu, jakie wykazują one dla oznaczenia i dokładnego określania stosunków przestrzennych.

W językach kulturowo rozwiniętych nastąpiło — zdaniem Cassirera (Cassirer 2004: 26) — „uwolnienie właściwej i pierwotnej formy języka od treści zmysłowego oglądu”. Proces dokonał się dwuetapowo. Najpierw „wyrażenie mimiczne” zostało zastąpione „analogicznym sposobem oznaczania”, w którym

nie ma [...] już jakiejś szczególnej obiektywnej własności przedmiotu, zatrzymanej i odtwarzanej w dźwięku, lecz całą subiektywność myślenia czy odczuwania przenika tu stosunek pomiędzy dźwiękiem i znaczeniem. [...] Już nie sama „rzecz”, lecz wypływające z jej subiektywności jej wrażenie bądź forma czynności podmiotu jest tym, co ma znaleźć w dźwięku swoje przedstawienie i jakiś rodzaj „odpowiedniości”.

Takim „analogicznym sposobem przedstawienia”, „jakimś rodzajem odpowiedzi” jest według Cassirera m.in. reduplikacja. W przypadku reduplikacji (Cassirer 2004: 27) „określony zmysłowy środek brzmieniowy i dźwiękowy służy również wyrazowi różnorodnych myślowych odniesień i znaczeń”. Reduplikacja, precyzuje Cassirer (tamże),

łączy się najpierw znów z obiektywnym procesem i stara się go bezpośrednio odtworzyć: podwojenie i powtórzenie sylaby służy dla oznaczania czynności lub procesu, jaki faktycznie dokonuje się w kilku podobnych fazach. Z tego miejsca posuwa się ona dalej, aby oznaczać takie treści, które z fenomenem powtarzalności łączą się tylko na zasadzie dalekiej analogii. W przypadku rzeczownika służy ona do tworzenia wielości, w przypadku przymiotników do stopniowania, w przypadku czasownika tworzy obok form frekwencyjnych przede wszystkim formy intensywne i w dalszej kolejności staje się wyrazem dużej liczby rozróżnień, w szczególności czasowych. Istnieją języki, w których owa metoda reduplikacji opanowuje całą budowę gramatyczną. Z tego wszystkiego wynika jasno, jak język, uwolniwszy się nawet od zwykłego onomatopeicznego sposobu wyrazu, usiłuje wciąż jeszcze dopasować się do treści znaczenia, niejako podążać za nią i jej dotykać.

W drugim etapie „uwalniania się formy języka od treści zmysłowego oglądu” zachodzi „przeniesienie słowa w określoną kategorię gramatyczną” (s. 27). W tym przypadku — na najwyższych etapach rozwoju języka, stwierdza Cassirer, „rezygnuje się ze wszelkich form rzeczywistego naśladowania”, miejsce naśladowania zajmuje „samodzielna funkcja znaczenia”, przy czym (s. 28)

[w] im mniejszym stopniu forma językowa usiłuje dostarczać — czy to bezpośredniego, czy to pośredniego odbicia przedmiotowego świata, w im mniejszym zakresie identyfikuje się z bytem (*Sein*) tego świata, w tym większym stopniu dociera do swej swoistej sprawności (*Leistung*), do specyficznego sensu. Dopiero teraz osiąga ona, w miejsce wyrazu mimicznego lub analogicznego, stopień wyrazu symbolicznego, który odzegnując się od wszelkiego podobieństwa z przedmiotowością, właśnie w tym oddaleniu i odwrocie zyskuje nową duchową treść.

Podobnie, choć z innego punktu widzenia — przez pryzmat „mowy w działaniu” — widzi proces „odzegnania się formy językowej od wszelkiego podobieństwa” Bronisław Malinowski. To dzięki „mowie w działaniu” następuje — według Malinowskiego — „przeniesienie słowa w określoną formę gramatyczną” i „samodzielną funkcję znaczenia”. W pracy *Problem znaczenia w językach pierwotnych* Malinowski (1923/2000: 361) zauważa, na przykład, iż:

Świat zewnętrzny interesuje tubylca o tyle, o ile przynosi coś pożytecznego. Pożytek trzeba tu oczywiście rozumieć w najszerszym sensie tego słowa, obejmującym nie tylko to, co człowiek może spożyć, nie tylko schronienie lub narzędzie, lecz również to wszystko, co pobudza jego aktywność w zabawie, obrzędzie, wojnie lub twórczości artystycznej. Wszystkie w tym sensie istotne przedmioty dziki wyróżnia jako odrębne całości, oderwane od nieodróżnicowanego tła.

Owo „wyodrębnienie z tła” — konstatuje Malinowski — w szczególności wyodrębnienie przez tubylców kształtów zwierząt odpowiada zainteresowaniu dzieci zwierzętami. Według Malinowskiego tak umysł pierwotny, jak i umysł dziecka z nieodróżnicowanego tła wyróżniają ogólną kategorię osób i personifikowanych rzeczy, odpowiadającą Arystotelesowskiemu pojęciu *ousia*. Jest to „surowa, nieobrobiona, niezgrabna forma”, z której „mogły się rozwinąć różne pojęcia substancji” (s. 363). *Ousia* lub *protousia* — stwierdza Malinowski (s. 363):

wymaga artykułowanych dźwięków — i otrzymuje je — do oznaczania różnych rzeczy. Zbiór słów używanych do nazywania osób i personifikowanych przedmiotów tworzy pierwszą kategorię gramatyczną rzeczowników. Widać więc, że ta część mowy zakorzeniona jest w aktywnych sposobach zachowania i w czynnościach zastosowania mowy, które można zauważyć u dziecka, u dzikiego, a przypuszczalnie także i u pierwotnego człowieka.

Nie inaczej ma się sprawa z pozostałymi kategoriami gramatycznymi, np. czasownikami — także i ta kategoria „zakorzeniona jest w aktywnych sposobach zachowania”. Powiada Malinowski (s. 363–364):

Dziecko dowiadyuje się i musi się dowiadywać o pożywieniu oraz poznawać osobę, która go dostarcza, zanim jest w stanie wypłatać, gdy tego potrzebuje, czynności z podmiotu czynności i zanim staje się świadome swoich własnych działań. Stany, w których znajduje się ciało dziecka, również o wiele mniej wyraźnie zaznaczają się w jego sytuacji niż rzeczy, które wchodzą w jej zakres. Tak więc dopiero na następnym etapie rozwoju dziecka można dostrzec, iż oddziela ono zmiany w swoim otoczeniu od przedmiotów, które się zmieniają lub zmiany wywołują. Staje się to wówczas, kiedy dziecko zaczyna stosować artykułowane dźwięki. Na tym etapie zaczynają znajdować swój wyraz językowy takie czynności, jak jedzenie, picie, wypoczywanie czy chodzenie; takie stany ciała, jak sen, głód lub wytchnienie; oraz takie nastroje, jak upodobanie i niechęć. Można powiedzieć, iż ta rzeczywista kategoria działania (action), stanu (state) i nastawienia (mood) przydatna jest zarówno do nakazywania, jak i do wskazywania bądź opisywania; można też dodać, że jest ona kojarzona z elementem zmiany, tj. czasu, oraz, że znajduje się w szczególnie ścisłym związku z osobami mówiącego i słuchacza. W ramach światopoglądu dzikich daje się zauważyć taki sam charakter tej kategorii; duże zainteresowanie wszystkimi zmianami dotyczącymi istoty ludzkiej, fazami i typami ludzkich działań, stanami ludzkiego ciała i ludzkimi nastrojami.

Dalej Malinowski zauważa, iż (s. 344)

gdy przyglądamy się zbiorowi słów służących do oznaczenia elementów należących do tej rzeczywistej kategorii, zauważymy, że dokładnie odpowiada ona części mowy. We wszystkich językach słowo dotyczące działania, czyli czasownik, dopuszcza gramatyczne modyfikacje wyrażające stosunek do czasu i tryby wyrażania.

Czasownik z kolei ściśle wiąże się z inną klasą słów — z zaimkami, które odpowiadają „rzeczywistej kategorii ludzkich zachowań oraz pierwotnych zwyczajów językowych”, odzwierciedlając tym samym „sposoby ludzkiego działania” (s. 364–365):

Mowa jest [...] jednym z głównych sposobów ludzkiego działania, dlatego podmiot działający za pomocą języka — mówiący — znajduje się na pierwszym planie pragmatycznej wizji świata. Co więcej, skoro mowa wiąże się z uzgadnianiem postępowania, ten, kto mówi, musi stale porozumiewać się ze słuchaczem lub słuchaczami. W ten sposób mówiący i słuchacz zajmują — by tak rzec — dwa główne, niejako narożne miejsca w perspektywie, którą język udostępnia. Wobec tego pojawia się bardzo ograniczona, specjalna klasa słów odpowiadająca pewnej rzeczywistej kategorii: są one stale w użyciu, łatwo łączy się je ze słowami przeznaczonymi dla działań, ale swoim charakterem gramatycznym są podobne do rzeczowników — tworzą część mowy nazywaną zaimkiem, obejmującą tylko kilka stale używanych słów; z reguły są to słowa krótkie, z którymi łatwo sobie poradzic, pojawiające się z czasownikiem, lecz funkcjonujące niemal tak jak rzeczowniki. Jest oczywiste, że ta część mowy dokładnie odpowiada rzeczywistej kategorii, co można wyśledzić w wielu interesujących szczegółach — w osobliwym asymetrycznym położeniu trzeciej osoby zaimka, w problemie rodzajów i częstek klasyfikacyjnych, zwłaszcza gdy pojawia się on w tej trzeciej osobie.

Naturalnie można podać inne przykłady, gdzie język — poprzez konkretnie kategorie gramatyczne — odzwierciedla „sposoby ludzkiego działania” — np. działania człowieka w przestrzeni. W tym wypadku będziemy mieli do czynienia z wyrażeniami przyimkowymi funkcjonującymi jako okoliczniki miejsca — słowami „udostępnionymi przez język”, przy pomocy których kodowane są relacje przestrzenne „lokujące” substancję w określonym miejscu lub określające miejsce działania. Do tego zagadnienia powrócimy za chwilę, omawiając pojęcie tzw. *ścieżek naturalnych* leżących u podstaw Langackerowskiej koncepcji *poznania zaangażowanego*.

4. Zaangażowanie świata i subiektyfikacja

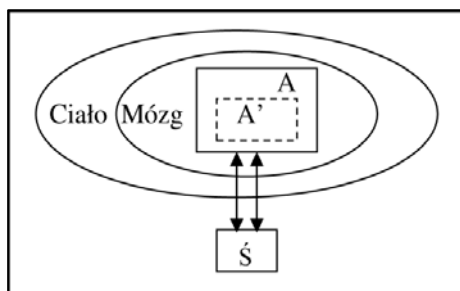
Zacznijmy od pojęcia *ucieleśnienia* — kluczowego terminu językoznawstwa kognitywnego. Powiada Langacker (2008/2010: 699):

W ostatecznym rozrachunku świat, który budujemy, jest zakorzeniony w naszym doświadczeniu stworzeń obdarzonych ciałami, które wchodzi w interakcje z oto-

czeniu poprzez fizyczne procesy, polegające na sensorycznej i motorycznej aktywności. W językoznawstwie kognitywnym nazywamy taki stan rzeczy ucieleśnieniem.

Owe interakcje z otoczeniem leżą u podstaw poznania zaangażowanego — procesu kognitywnego, w którym postrzegający podmiot jest bezpośrednio zaangażowany w interakcję ze światem. Proces poznania zaangażowanego możemy przedstawić następująco:

(1)



W poznaniu zaangażowanym interakcja podmiotu ze światem (Ś) odbywa się bezpośrednio — poprzez ciało, organy zmysłowe i motoryczne. „A” symbolizuje mózg i aktywność sensomotoryczną, składającą się na interaktywne doświadczenie; A’ to niektóre elementy A, które — jak to ujmuję Langacker — „immanentnie zawierają się w A” i są „symulacją A” (s. 714–715). Elementy te pojawiają się autonomicznie, nie wchodzą one w bezpośrednią reakcję ze światem. Jak zauważymy za chwilę, owe elementy immanentne — A’ — staną się kluczowymi elementami innego typu przetwarzania informacji związanego z poznaniem — poznania oderwanego.

Przykładem poznania zaangażowanego są tzw. *ścieżki naturalne*, czyli — jak definiuje je Langacker (s. 667) — „ciąg[i] koncepcji, w którym każda prowadzi do następnej”. Przykłady takich ścieżek są liczne; mogą to być wypowiedzane w pewnej kolejności wyrazy lub na przykład łańcuch relacji uszczegóławiających występujących w zdaniach złożonych, jak w zdaniu (2), gdzie kolejne zdanie podrzędne konkretyzuje (lub uszczegóławia) zdanie poprzedzające:

- (2) [Rzecznik rządu ujawnił] [że premier nie zgadza się] [aby żądania związków zawodowych były omawiane na posiedzeniu rady ministrów] [która zbiera się w piątek]

Każdą ścieżkę otwiera tzw. *punkt początkowy*. Takim punktem może być na przykład konceptualizator (mówiący/adresat) występujący w łańcuchu innych

konceptualizatorów, z „których każdy postrzega następnego i w jakiejś mierze symuluje jego doświadczenia mentalne” (Langacker 2008/2010: 668). I tak, w (2) takim ciągiem konceptualizatorów jest *mówiący > rzecznik rządu > premier > związki zawodowe*. Inne ścieżki mogą określać — w różnych wymiarach i w różnym stopniu — dystans dzielący konceptualizatora-nadawcę od innych konceptualizatorów lub uczestników dyskursu, np. *mówiący > adresat > inny* lub *ożywiony > nieożywiony* czy *rzeczywisty > wirtualny* (tamże).

Jak zauważa Langacker, ścieżki naturalne często są ściśle powiązane ze sobą. Na przykład w powyższym zdaniu nakładają się na siebie następujące ścieżki naturalne: kolejność prezentacji, następujące po sobie kolejne konkretyzacje, łańcuch konceptualizatorów oraz przywoływane przez to zdanie, zanurzone jedna w drugiej, przestrzenie mentalne.

Owa zbieżność ścieżek jest, zdaniem Langackera, „ułatwieniem w przetwarzaniu wyrażen językowych” (s. 669) i bezpośrednio wiąże się z gramatyką — w tym wypadku z szykiem wyrazów, który, jak konstatuje Langacker, „zazwyczaj skorelowany jest z rozlicznymi ścieżkami kontaktu mentalnego” (tamże). Poniższe przykłady dowodzą słuszności Langackerowskiej tezy o korelacji naturalnej ścieżki, szyku wyrazów i przetwarzania informacji (s. 669):

- (3) a. Artykuł jest w dzisiejszej gazecie, w sekcji sportowej, na ostatniej stronie, na dole.
b. ?? Artykuł jest w sekcji sportowej, na dole, w dzisiejszej gazecie, na ostatniej stronie.
- (4) a. Pora deszczowa zaczyna się w styczniu i kończy w marcu.
b. ?? Pora deszczowa kończy się w marcu i zaczyna w styczniu.

Zauważmy, iż struktura zdania (3a) odpowiada naturalnej ścieżce przeszukiwania, stąd zdanie to jest o wiele łatwiej zinterpretować niż zdanie (3b), w którym okoliczniki pojawiają się w przypadkowej kolejności. Podobnie rzecz się ma z (4a), które brzmi o wiele bardziej naturalnie niż (4b), ponieważ występująca w (4a), lecz nie w (4b) kolejność prezentacji odzwierciedla naturalną sekwencję zdarzeń.

Powyższe zdania, ilustrujące korelację ścieżek naturalnych z szykiem wyrazów, są przykładami poznania zaangażowanego — kodowana jest tu przy pomocy gramatyki (jednostek językowych i szyku wyrazowego) Cassirerowska „bezpośrednia interakcja podmiotu ze światem”. Jak jednak zauważa Langacker (s. 699):

nasze życie mentalne wykracza oczywiście poza granice bezpośredniego doświadczenia cielesnego. Poszczególne procesy poznawcze powołują do życia struktury mentalne na kolejnych poziomach organizacji, a ich powiązanie z doświadczeniem

cielesnym jest coraz bardziej odległe. Struktury te nie tylko pozwalają nam radzić sobie lepiej ze światem rzeczywistym, lecz również definiują — i niezwykle poszerzają — wszystko, co się na niego składa. [...]

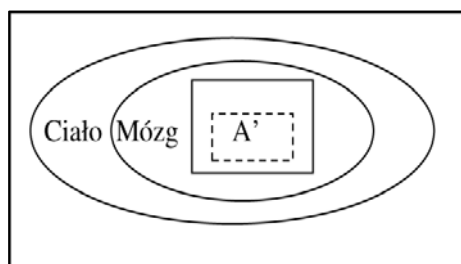
Zatem spora część naszej aktywności poznawczej ma charakter oderwany (w większym czy mniejszym stopniu) od bezpośredniego doświadczenia somatycznego.

W jaki jednak sposób wykraczamy poza to doświadczenie? — pyta Langacker? I odpowiada (tamże):

Jedną ścieżkę wyznacza pamięć, umożliwiającą częściowe odnowienie wcześniejszych doświadczeń. Inną otwiera antycypacja, dzięki której obserwowanie bieżącej sytuacji pozwala na dokonywanie projekcji w przyszłość. Może to mieć charakter zupełnie podstawowy, jak na przykład standardowe przewidywanie, że istniejący dziś przedmiot fizyczny będzie nadal istniał w przyszłości. Może się też opierać na wzorcach opanowanych przez wcześniejsze doświadczenia. Wzorców uczymy się poprzez ich abstrahowanie, co stanowi zasadniczy środek umożliwiający wykroczenie poza doświadczenie bieżącej chwili.

Wspomniane przez Langackera procesy „abstrahowania wzorców” i „wykroczenia poza doświadczenie bieżącej chwili” lub, używając Cassirerowskiej terminologii, „analogicznego sposobu przedstawienia” odbywają się w ramach *symulacji* — procesu, który „odrywa” poznanie człowieka od jego bezpośredniej, bieżącej interakcji ze światem. Proces symulacji można przedstawić następująco (por. Langacker, 714):

(5)



Zauważmy, iż w przeciwieństwie do poznania zaangażowanego (por. (1)), w którym elementy A' , jak to ujmuje Langacker, „immanentnie zawierają się w A ”, w poznaniu oderwanym elementy A' stają się elementami kluczowymi — są one „czymś w rodzaju cienia doświadczenia związanego z A ” — są zatem całkowitą „symulacją A ” (s. 714–715). Według Langackera jednym z najważniejszych procesów leżących u podstaw poznania oderwanego jest subiektywizacja — w szczególności tzw. *radykalna subiektywizacja*.

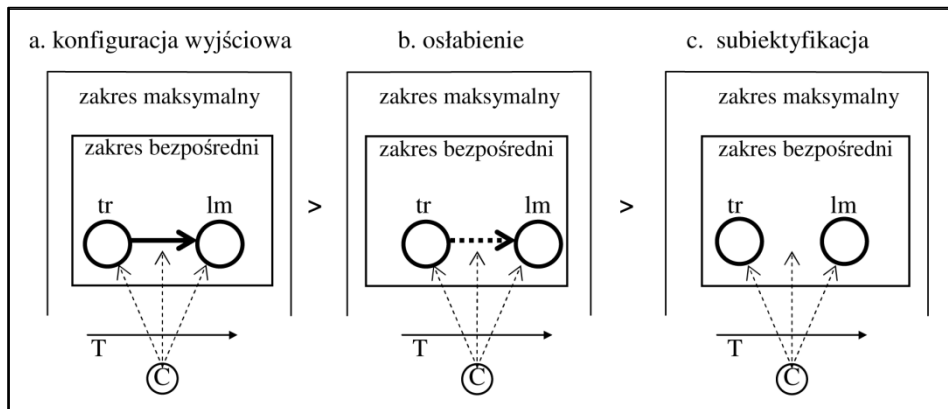
Subiektyfikację definiuje Langacker (2003: 5) jako proces konceptualizacyjny, w obrębie którego „zanika rzecz obiektywnie konstruowana, pozostawiając w wyniku mentalnych operacji swój subiektywnie konstruowany odpowiednik”². Proces subiektyfikacji, zwany także przez Langackera *ruchem abstrakcyjnym* (lub *wirtualnym*), ilustrują następujące przykłady:

- (6) Balon wzniósł się na wysokość 5000 metrów.
- (7) Skała wznosi się ponad 200 metrów nad poziom morza.

W zdaniu (6) mamy do czynienia z rzeczywistym ruchem balonu w górę, natomiast w zdaniu (7) ruch jest subiektywny — nie skała jest w ruchu, lecz *konceptualizator*, czyli mówiący, który przesuwał wzrok po krawędzi skały (czyli *skanując* krawędź skały), mentalnie porusza się po abstrakcyjnej ścieżce od podnóża owej skały na jej szczyt.

Subiektyfikacja jest procesem, w którym *obiektywnie istniejąca relacja OVA* (*objective viewing arrangement*) pomiędzy konceptualizatorem i konceptualizowaną przez niego sceną lub rzeczą, zastępowana jest przez tzw. relację EVA (*egocentric viewing arrangement*) (por. Langacker 2000: 298):

(8)



² Nieco inaczej definiuje subiektyfikację Elizabeth Traugott (1989). Jak zauważa Langacker (2006: 17), o ile Traugott używa terminu *subiektyfikacja* na określenie „tendencji znaczeń do stawania się bardziej subiektywnymi” i „dotyczy treści konceptualnych” (np. znaczenia „zaczynają odzwierciedlać subiektywne odczucia mówiącego, jak w przypadku zmiany semantycznej «temporalnego» przysłówka *while* na jego znaczenie «przyzwalające»”), o tyle subiektyfikacja à la Langacker dotyczy „punktu widzenia (*vantage point*), czyli jest elementem *konstruowania sceny* (*construal*)”. Nie rozstrzygając w tym momencie, czy, a jeśli tak, to jakie są punkty styczne między obydwoma koncepcjami subiektyfikacji, na potrzeby niniejszego wywodu przyjmuję koncepcję Langackerowską.

Rysunek (a) przedstawia konfigurację OVA. Strzałki symbolizują obiektywnie istniejącą relację między tzw. *trajektorem* (relacyjną figurą, która może pełnić funkcję agensa lub doświadczającego (*experiencer*) a *landmarkiem* (figurą, która może być np. *patienssem*). T to czas procesualny (*processual time*), tzn. czas, w którym dokonuje się proces konceptualizacji, natomiast C to znajdujący się *poza sceną* (*off-stage*) konceptualizator, który mentalnie skanuje relację między trajektorem a landmarkiem. Przerywane strzałki symbolizują mentalne operacje konceptualizatora. Rysunek (b) przedstawia osłabienie relacji OVA (*attenuation*). Konceptualizator w dalszym ciągu mentalnie skanuje scenę, lecz relacja między tr i lm jest motywowana w mniejszym stopniu obiektywnie istniejącą sytuacją. Obiektywnie istniejąca motywacja między tr i lm kompletnie zanika w końcowym stadium subiektyfikacji (por. rysunek (c)); związek między tr i lm opiera się w tym przypadku na całkowicie subiektywnej relacji ustanowionej w procesie mentalnego skanowania.

Stopniowy proces subiektyfikacji można zilustrować następującymi przykładami (Langacker 2000: 301; cytowane także w: Kardela 2006):

- (9) a. The child hurried across the street.
Dziecko przebiegało przez ulicę.
- b. The child is safely across the street.
Dziecko jest teraz bezpieczne po drugiej stronie ulicy.
- c. You need to mail a letter? There is a mailbox just across the street.
Chcesz wysłać list? Skrzynka pocztowa znajduje się po drugiej stronie ulicy.
- d. A number of shops are conveniently located just across the street.
Wiele sklepów jest dogodnie rozmieszczonych po drugiej stronie ulicy.
- e. Last night there was a fire across the street.
Zeszłej nocy ogień płonął w poprzek ulicy.

Zdanie (a) profiluje obiektywny ruch trajektoria (tutaj: dziecka), zdanie (b) opisuje statyczne położenie dziecka, które jest wynikiem wyprofilowania uprzedniego rzeczywistego ruchu trajektora, przykład (c) wskazuje na miejsce jako rezultat wyprofilowanego celu związanego z potencjalnym ruchem wysyłającego list, (d) odnosi się do miejsca jako wyniku profilowanego potencjalnego ruchu nieokreślonej osoby, zaś zdanie (e) profiluje lokalizację ognia przy nieobecności jakiegokolwiek ruchu.

5. Faktyczność i wirtualność

Jeśli, jak stwierdza Langacker, przeważająca część poznania związana jest nie tyle z poznaniem zaangażowanym, bazującym na bezpośredniej, sensomotorycznej reakcji człowieka ze światem, ile raczej z poznaniem oderwanym, cechującym się brakiem bieżącej interakcji ze światem, to powstaje pytanie — do którego typu poznania odnosi się przede wszystkim język i jego struktury gramatyczne?

Odnosząc się do tezy o związku między językiem a rzeczywistością, Langacker zauważa, iż istnieje powszechna tendencja, aby język naturalny traktować jako narzędzie używane do opisu otaczającego nas świata. Pogląd taki jest błędny zdaniem Langackera — język w wielu przypadkach nie opisuje bowiem rzeczywistych bytów i sytuacji, lecz byty i sytuacje wirtualne. Powiada Langacker (2005: 44):

Twierdzę [...], że w dużej mierze nie zdajemy sobie sprawy i z wielości, i z różnorodności odstępstw od bezpośredniego opisu faktycznych relacji i bytów. Jeśli mamy dokonać rzeczowej oceny języka, poznania i mentalnej konstrukcji naszego świata, musimy w pełni uświadomić sobie, na czym polegają i czym są owe odstępstwa. [...] Rozróżnienie, które chcę [...] wprowadzić [to rozróżnienie na] byty faktyczne i wirtualne [...]. Granica między faktycznością (actuality) a wirtualnością (virtuality) nie jest absolutna w tym sensie, że linię podziału można przeprowadzić w różnych miejscach. [...] Zajmę się zatem różnymi rodzajami niebezpośredniości (indirectness), jakie dostrzegam w związkach między wyrażeniami językowymi a faktycznymi bytami i relacjami, do jakich te wyrażenia się odnoszą.

Tak rozumiane pojęcie „faktyczności” odbiega od ogólnie przyjętego rozumienia „faktyczności” czy „autentyczności”, ponieważ dla Langackera (2005: 45)

[...] faktyczność jest [...] czymś odmiennym od takich pojęć jak ‘świat rzeczywisty’, ‘realność’ czy ‘prawda’. Opozycja między faktycznością a wirtualnością może zostać wprowadzona w obręb każdego całościowego ‘świata’, bez względu na to, czy rozumiemy przez to ‘świat realny’, czy tzw. ‘świat projektowany’, czyli ten, jaki znajdujemy w wyobrażeniowym świecie mitu, powieści, Biblii, filmu itd.

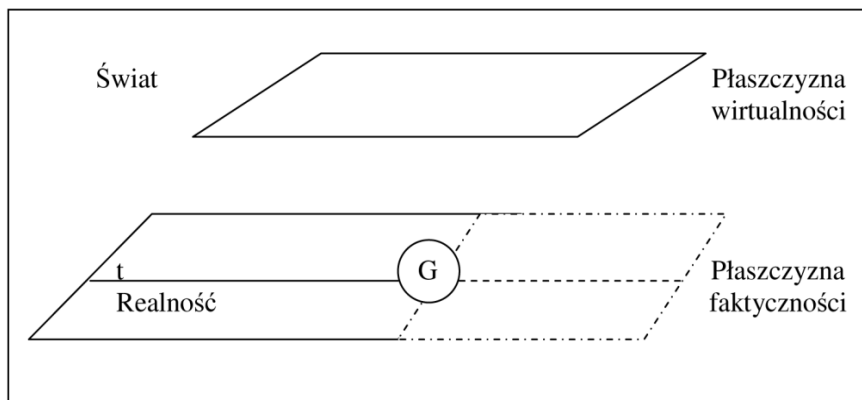
Opozycję „faktyczność — wirtualność” najlepiej ilustrują następujące przykłady (Langacker 2005: 45):

- (10) a. Adam ate an apple ‘Adam zjadł jabłko’ [bezpośredni opis konkretnego, chociaż mitycznego zdarzenia]
- b. Serpents seldom seem sincere ‘Wężę rzadko są szczerze’ [opis generyczny, który w sposób bezpośredni nie dotyczy żadnej konkretnej postaci]

Zdanie (a) jest bezpośrednim opisem faktycznego (jakkolwiek mitycznego) zdarzenia; zdanie (b) natomiast nie jest już takim bezpośrednim opisem faktyczności — nie odnosi się ono do faktycznych zdarzeń. Zdanie (b) opisuje zdarzenie generyczne — generalizację tego, co znajdujemy w faktyczności.

Obydwie płaszczyzny: wirtualności i faktyczności możemy przedstawić następująco (Langacker 2005: 46; adaptacja):

(11)



Według Langackera (2005: 46–47):

W obrębie danego świata realność (rzeczywistość) definiujemy jako historię tego, co się wydarzyło do chwili obecnej, zgodnie z tym, jak to widzi i konceptualizuje użytkownik języka [...] Realność ewoluuje w czasie (t) i „rośnie” w kierunku przyszłości w miarę akumulowania się zdarzeń. Punkt zakotwiczenia G (*ground*) symbolizuje dane zdarzenie mowne, jego uczestników oraz najbardziej bezpośrednie okoliczności zdarzenia. Ponieważ punkt zakotwiczenia znajduje się „na początku” realności, podążającej w stronę przyszłości, G, symbolizuje „teraz” lub terażniejszość. Realność jest zatem jednym z wymiarów faktyczności lub aktualności (*actual plane* — H. K). Faktyczność obejmuje również „modalne przekształcenia realności”, takie jak potencjalność (*potentiality*), przeniesienie elementów w przyszłość (*extrapolation*), przewidywalność (*predictions*), a nawet negację.

Dodajmy, że byty, które nie należą do faktyczności, umieszczone są w wirtualności, przy czym między obiema płaszczyznami — faktycznością i wirtualnością występują różnorakie związki i korelacje.

Z powyższej dyskusji wyłania się odpowiedź na pytanie, do czego odnosi się język: język nie tyle odnosi się bezpośrednio do jakiejś konkretnej sytuacji (choć także do takiej), ale przede wszystkim do poznania oderwanego. Język opisuje raczej *typy sytuacji i wzorce zachowań człowieka*, których egzemplifikacją jest konkretna sytuacja. Zauważmy, iż nawet w przypadku, gdy w pewnym

momencie się oparzemy i krzyknijemy: *Oparzyłem się — ależ ta woda to ukrop!*, to użyjemy skonwencjonalizowanej struktury, która może być także użyta w wielu innych przypadkach tego rodzaju. Możliwość użycia wyżej wymienionej struktury zdaniowej w wielu innych podobnych przypadkach, w oderwaniu od bieżącej interakcji człowieka ze światem (w tym przypadku gorącej wody i oparzonej ręki), jest wynikiem *radikalnej subiektywizacji* — „ekstremalnego” procesu subiektywizacji. Ostatecznym produktem radykalnej subiektywizacji jest *wirtualny dokument*.

6. Wirtualny dokument

Przyjrzyjmy się następującym przykładom (Langacker 2003: 22)

- (12) a. My mother will arrive next week.
Moja matka przyjedzie w przyszłym tygodniu.
- b. My mother arrives next week.
Moja matka przyjeżdża w przyszłym tygodniu.
- c. ?? My mother arrives.
Moja matka przyjeżdża.
- d. ?? An earthquake strikes next year.
Trzęsienie ziemi następuje w przyszłym roku.

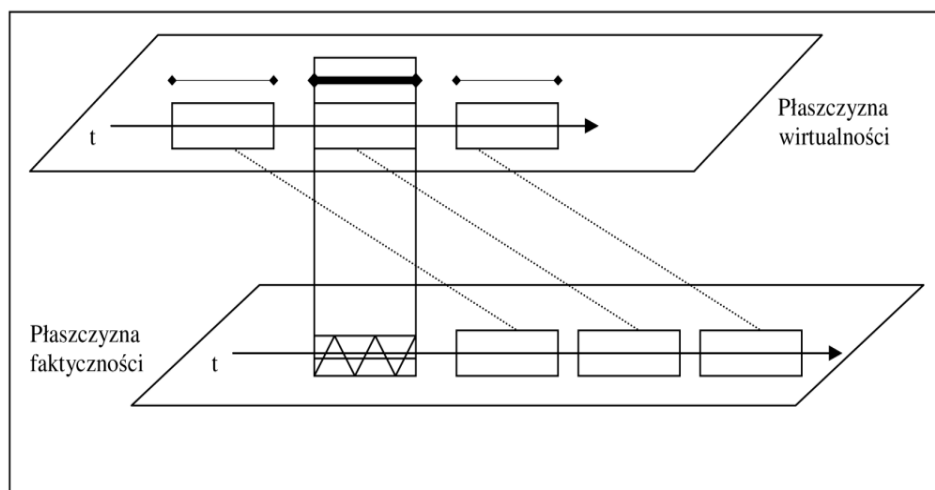
Jak zauważa Langacker, (a) jest bezpośrednim opisem konkretnego (*actual*) zdarzenia, pomimo użycia tu czasownika modalnego *will*, który sytuje zdarzenie przyjazdu matki w rzeczywistości (*actuality*) zanim akt wypowiedzenia miał miejsce (zdarzenie przyjazdu matki jeszcze nie zaistniało). Zdarzenie to należy jednak uznać za element ewoluującej rzeczywistości, ponieważ — zdaniem Langackera — „koncepcja rzeczywistości, jaka jest udziałem mówiącego, może objąć (przy pewnym mentalnym wysiłku z jego strony) także owo zdarzenie” (Langacker 2003: 22). Oznacza to, że zdarzenie jest zakotwiczone (*grounded*) w rzeczywistości (aktualności). Sytuacja opisywana w (b) jest inna. Według Langackera zdania takie jak (b) opisują nie zdarzenie rzeczywiste, lecz wirtualne. Profilowany tu proces dotyczy — używając Langackerowskiego sformułowania — „nie przyjazdu jako takiego, lecz *reprezentacji* objętej *wirtualnym planem*, tzn. planem lub projekcją odnoszącą się do oczekiwanego zdarzenia i czasu, w jakim zdarzenie zaistnieje w przyszłej aktualności”. Ów wirtualny plan, zwany przez Langackera także *wirtualnym dokumentem*, jest dostępny „w każdej

chwili”, w tym także w trakcie wypowiedzi³. Oznacza to, że (a) nie opisuje rzeczywistego zdarzenia (choć także i zdarzenie rzeczywiste), lecz raczej konceptualizację, która jest wytworem „czytania” wirtualnego dokumentu.

Za istnieniem wirtualnego dokumentu (lub planu) przemawia właśnie kontrast między akceptowalnym zdaniem (12b), w którym forma czasu teraźniejszego opisuje zdarzenie, które będzie miało miejsce w przyszłości (w przyszłym tygodniu), a nieakceptowalnymi przykładami (c) i (d). Zdanie (12c) jest nie do przyjęcia, ponieważ brak tu informacji, kiedy zdarzenie zaistnieje, natomiast w (12d) mamy wprowadzić podany czas zdarzenia (przyszły rok), lecz zdarzenia w rodzaju trzęsień ziemi nie są planowane, stąd (d) nie może być traktowane jako reprezentacja występująca w wirtualnym dokumencie.

Wirtualny dokument, używany do opisu przyszłych „planowanych” zdarzeń, „dostępny w każdej chwili”, możemy przedstawić następująco (Langacker 2003: 22; adaptacja; por. także (11)):

(13)



³ W angielskim nie jest możliwe użycie formy czasu teraźniejszego zwykłego (Simple Present) do opisu zdarzenia, które nie pokrywa się z czasem wypowiedzi. O ile po polsku powiemy: *Marysia ogląda telewizję codziennie* i *Marysia w tym momencie ogląda telewizję* (zdarzenie pokrywa się z czasem wypowiedzi), w angielskim użyjemy czasu zwykłego teraźniejszego tylko w sytuacji, gdy Marysia ogląda telewizję codziennie (*Mary watches television everyday*) natomiast użyjemy formy czasu teraźniejszego ciągłego (Present Continuous), gdy zdarzenie pokrywa się z czasem wypowiedzi (*Mary is watching television now*, ale nie: **Mary watches television now*). Wyjątkiem są tu czasowniki performatywne (*I promise to come tonight* 'Obiecuję przyjść dziś wieczorem'), czasowniki statyczne (*Peter resembles John* 'Piotr przypomina Jana' (por. **Peter is resembling John*), *Mary knows Peter very well* (por. **Mary is knowing Peter very well*) oraz formy czasu teraźniejszego w wypowiedzeniach traktowanych jako reprezentacje występujące w dokumencie wirtualnym.

Prostokąty oznaczają tu sytuujące się na osi czasu (t) zdarzenia, zakreskowany prostokąt występujący na płaszczyźnie faktyczności to zdarzenie mowne (wypowiedzenie), krótkie odcinki zaś, w tym pogrubiony (tj. *wyprofilowany*) odcinek, symbolizują procesy desygnowane przez odnośne czasowniki. Linie wykropkowane to tzw. *odpowiedniości* (*correspondences*), łączące zdarzenia „faktyczne” ze zdarzeniami „wirtualnymi”. Te ostatnie, kodowane przez zdania typu *My mother arrives next week*, powstają w wyniku posługiwania się przez konceptualizatora strukturami gramatycznymi umożliwiającymi mu „czytanie wirtualnego dokumentu”.

7. Formy symboliczne raz jeszcze

Przypomnijmy cytowaną już definicję formy symbolicznej: forma symboliczna to „energia ducha, poprzez którą duchowa treść znaczeniowa zostaje związana z konkretnym zmysłowym znakiem i temu znakowi wewnętrznie przyporządkowana”. W przypadku języka owym znakiem jest znak dźwiękowy, który

jest środkiem i warunkiem samego wewnętrznego różnicowania wyobrażeń. Artykulacja dźwięku wypowiada nie tylko gotową artykulację myśli, lecz sama przygotowuje jej dopiero drogę. Znak zmierza zawsze do tego, aby przyłgnąć możliwie blisko do określanego przedmiotu, aby go niejako wchłonąć i oddać możliwie najdokładniej i najpełniej (Cassirer 2004: 24).

Jak zauważa Cassirer (2004: 21), język, mityczno-religijny świat i sztuka to formy symboliczne, jako że (tamże)

nasza świadomość nie zadawała się odbieraniem wrażenia czegoś zewnętrznego, lecz że każde wrażenie łączy ona i przenika swobodną czynnością wyrazu. Naprzeciwko tego, co nazywamy obiektywną rzeczywistością rzeczy, pojawia się świat samodzielnie stworzonych znaków i obrazów i trwa wobec niej w samodzielnej pełni i pierwotnej sile.

„Samodzielnie stworzone znaki i obrazy”, precyzuje Cassirer, stają między nami i przedmiotami, pośrednicząc jako medium, „poprzez które staje się dla nas możliwy do ujęcia i zrozumiały jakikolwiek byt” (s. 22). W tym ujęciu słowo nie jest „odbiciem przedmiotu w sobie”, jest to raczej „wytworzony z przedmiotu przez duszę” obraz tegoż przedmiotu. Cassirer cytuje Humboldta (Cassirer 2004: 22):

Tak jak poszczególny dźwięk staje pomiędzy przedmiotem i człowiekiem, tak też cały język staje między nim i oddziaływującą na niego z wewnątrz i z zewnątrz przyrodą. Otacza się on światem dźwięków, aby przejść do swego wnętrza i opracować świat przedmiotowy. Poprzez ten sam akt, mocą którego człowiek wydobyla z siebie język, wzięła się on w tenże [...].

Owo „wewnętrzne opracowanie świata przedmiotowego” pozwala zamknąć — jak to ujmuje Cassirer — „przepaść między heraklitejskim prawem ‘stawiania się’ świadomości a stałą ciągłością i względną trwałością rzeczy przyrodniczych w ich obiektywno-realnym istnieniu”. W działalności bowiem ducha, powiada Cassirer (2004: 23–24),

dokonuje się nieustannie cud zamykania się przepaści; że to, co ogólne spotyka się, niejako w jakimś duchowym środku, z tym, co szczegółowe i splata się z nim w prawdziwą konkretną jedność. Proces ten przedstawia nam się wszędzie tam, gdzie świadomość nie zadawała się tym, aby jakąś zmysłową treść po prostu *posiadać*, lecz gdzie ona ją z siebie *wytwarza* (kursywa oryginalna). Siła tegoż wytwarzania jest tym, co zwykłą treść odczuć i wrażeń przekształca w treść symboliczną. W tej ostatniej obraz przestał być czymś po prostu odbieranym z zewnątrz; stał się on czymś wytworzonym od wewnątrz, gdzie panuje zasada swobodnego tworzenia. Każda z tych form posiada w zmysłowości nie tylko swój punkt wyjścia, lecz pozostaje ona także na stałe zamknięta w kręgu zmysłowości. Nie zwraca się ona przeciwko materiałowi zmysłowemu, lecz w nim samym żyje i tworzy.

8. Konkluzja

Powtórzmy za Cassirerem: żadna z form symbolicznych „nie zwraca się [...] przeciwko materiałowi zmysłowemu, lecz w nim samym żyje i tworzy”. Dodajmy w duchu kognitywistycznym, trawestując Cassirera: w przypadku poznania i języka żadna z form symbolicznych nie zwraca się [...] przeciwko ucieleśnionej myśli i strukturom poznawczym (por. Lakoff 1987), lecz w nich samych „żyje i tworzy” — żyje i tworzy m.in. dzięki (radikalnej) subiektyfikacji, umożliwiając konceptualizatorowi posługującemu się gramatyką i strukturami językowymi „odczytywanie wirtualnego dokumentu”.

Bibliografia

- CASSIRER E. (1955), *The Philosophy of Symbolic Forms*. Volume 1: *Language*. Yale.
- CASSIRER E. (2004), *Symbol i język* (wybór i przekład B. Andrzejewski). Poznań: Wydawnictwo Naukowe Wyższej Szkoły Pedagogiki i Administracji.
- “Cassirer”. *Stanford Encyclopedia of Philosophy*, <http://plato.stanford.edu/entries/cassirer/>.
- LAKOFF G. (1987), *Women, Fire and Dangerous Things. What Categories Reveal about the Mind*. Chicago: Chicago University Press.
- LANGACKER R. (1990), *Subjectification*, “Cognitive Linguistics” 1, 5–38.
- LANGACKER R. (2000), *Grammar and Conceptualization*. Berlin: Mouton de Gruyter.
- LANGACKER R. (2005), *Wykłady z gramatyki kognitywnej*. Lublin: UMCS.
- LANGACKER R. (2008/2010), *Cognitive Grammar. A Basic Introduction*. Oxford: OUP.
- Polski przekład: *Gramatyka kognitywna. Wprowadzenie*. Universitas.

- LUFT S. (2005), *Cassirer's Philosophy of Symbolic Forms: Between Reason and Relativism; a Critical Appraisal*, "Idealistic Studies", vol. 34, no 1.
- MALINOWSKI B. (2000), *Dzieła*. Tom 8. *Jednostka, społeczność, kultura*. Warszawa: PWN.
- ŚWIĘCICKA K. (1996), *Filozofia kultury Ernsta Cassirera*, [w:] Skarga B. (red.), *Przewodnik po literaturze filozoficznej XX wieku*, t. 4.
- TRAUGOTT E. (1989), *On the rise of epistemic meanings in English. An example of Subjectification in semantic change*, "Language" 65, 31–55.
- WENDLAND M. (2011), *Konstruktywizm komunikacyjny*. Poznań: Wydawnictwo UAM.

Rozdział 3

Teoria aktów mowy Adolfa Reinacha

Marek Nowak (Uniwersytet Łódzki)

Słowa kluczowe: J. L. Austin, A. Reinach, J. R. Searle, akt illokucyjny, akt społeczny, obiecywanie, wymaganie, zdanie analityczne, zdanie syntetyczne *a priori*, zobowiązanie

1. Wstęp

Tekst ten jest opracowaniem jednego tematu z zakresu historii filozofii języka, nie zaś współczesnej filozofii języka. Tym, co usprawiedliwia jego obecność w niniejszym tomie, jest inspiracyjne oddziaływanie tego tematu na zupełnie współczesne rozważania z zakresu teorii aktów mowy. Owym tematem jest dojrzała, rozwinięta i uzasadniona w szczegółach, pierwsza chronologicznie i właściwie niemająca ancestorów teoria czynności mowy, opublikowana w 1913 roku przez Adolfa Reinacha w pracy *A priori podstawy prawa cywilnego* (zob. Reinach [1913] 1983).

Reinach (1883–1917) nie zajmował się filozofią języka. Był fenomenologiem z monachijskiego kręgu zainteresowanych pismami Edmunda Husserla (J. Daubert, A. Pfänder, M. Geiger i inni). Zajmował się przede wszystkim filozofią prawa, również ontologią (m.in. dokonał przełomowego rozróżnienia między osądami, sądami w sensie logicznym, stanami rzeczy) — por. Schumann, Smith (1987) oraz Smith (1990). Interdyscyplinarną pracę z 1913 roku zaliczyć można do trzech dyscyplin: filozofii prawa, ontologii, filozofii języka. Celem pracy nie jest wcale opis czy analiza czynności dokonywanych podczas wypowiedzania zdań. To zostało w niej zrealizowane jedynie jako środek do osiągnięcia innego celu. Jest nim analiza statusu ontycznego pewnych fundamentalnych obiektów, których dotyczy jakiejkolwiek (kiedykolwiek i gdziekolwiek stanowiące) prawo cywilne, takich jak *własność*, *zobowiązanie* (się do czegoś w interesie kogoś), *wymaganie* (czegoś od kogoś). Reinach traktuje te przedmioty jako niefizyczne, bo pozaprzestrzenne oraz niementalne, choć trwające w czasie. Te dziwnej natury obiekty rządzą się dziwnej natury prawami.

Mianowicie prawa te są, według Reinacha, zdaniem *syntetycznymi a priori*, a więc takimi, których prawdziwość jest konieczna — niezależna od doświadczenia (zewnętrznego, czyli zmysłowego oraz wewnętrznego, czyli refleksji) i gwarantowana przez jakiś czysty ogląd pozazmysłowy stanów rzeczy opisywanych przez takie zdania. Oto przykłady praw tego rodzaju: *zobowiązanie* oraz *wymaganie* wygasają (przestają istnieć), gdy zobowiązanie jest dotrzymane, a wymaganie spełnione; *zobowiązanie* i *wymaganie* powstają jako skutek wywołany przez pewną przyczynę — jest nią pewna czynność psychiczna.

Konsekwencje tego drugiego prawa są istotne dla filozofii języka. Reinach bowiem jako przyczyny powstania *zobowiązania-wymagania* wymienia dwie czynności umysłowe: *obiecywanie* oraz *pożyczanie*. Pierwsza z nich, jak wiadomo, jest dla Searle'a wzorcową *pełną czynnością mowy* — *elementarnym aktem illokucyjnym* (w terminologii Austina *czynnością illokucyjną* — Austin [1962] 1993), którego analiza (zob. Searle [1969] 1987) umożliwiła wyodrębnienie początkowo siedmiu, a później sześciu parametrów (zob. Searle 1975, Searle, Vanderveken 1985) stosowanych do dzisiaj do klasyfikacji wszelkich typów *mocy illokucyjnych* (ang. *illocutionary forces*). Dla realizacji celów ontologicznych Reinacha *obiecywanie* stanowi najważniejszą z tzw. *czynności społecznych* (ang. *social acts*). Właśnie pojęcie *czynności społecznej* jest podstawowym przedmiotem badanym w teorii aktów mowy Reinacha. Jest bowiem odpowiednikiem *czynności illokucyjnej*.

Nazwy *social act* (oraz synonimicznie *social operation*) używał przed nim twórca Scottish School of Common Sense Philosophy, żyjący w XVIII wieku Thomas Reid, przez niektórych postrzegany jako prekursor współczesnych teorii aktów mowy czy współczesnej filozofii języka (zob. Schumann 1990) — według kryterium mówiącego, iż filozofem języka jest ten, kto jest przekonany, że uważna analiza języka potocznego pozwala usunąć czy uznać za nieważne pewne problemy filozoficzne, bo spowodowane niewłaściwym użyciem języka. Otóż w swojej pracy Reid ([1785] 2002) (por. również Schumann 1990 oraz Yaffe 2007) rozróżnia dwa typy czynności umysłowych: *social acts* (np. pytania, czyli prośby o informację lub radę, akty akceptacji i odmowy, zeznawania i akceptowania zeznań, rozkazywania komuś i otrzymywania rozkazu od kogoś, obiecywania i uzyskiwania obietnicy) oraz *solitary acts* (np. osądzanie, rozumienie, zamierzanie, chcenie, radowanie i smucenie się). Dla aktów solitarnych nie jest istotne, czy są one wyrażone w języku, a więc czy towarzyszy im wypowiedź słowna; ich wykonanie nie zakłada istnienia słuchacza. Tymczasem akty społeczne muszą być wyrażone w języku (tzn. ich wykonanie następuje dzięki wypowiedzianiu zdań) i muszą być skierowane do osób innych niż mówca. Nie jest tak, jakoby można było do danego aktu solitarnego dołączyć czynność wyrażenia go w języku, aby otrzymać w ten sposób akt społeczny. Na przykład rozkaz (akt społeczny) nie jest aktem życzenia (tj. aktem solitarnym)

wyrażonym w języku. Obietnica nie jest rodzajem chcenia czy zamiaru, który może być lub nie wyrażony w języku. Tylko akty solitarne mogą, lecz nie muszą być wyrażone językowo. Milczące zeznanie — to sprzeczność, lecz milczący osąd — jest możliwy.

Według Reida (zob. Schumann 1990) *rozumienie* i *zamiar* to dwie presupozycje dowolnego aktu społecznego. Jest on wykonywany dokładnie wówczas, gdy

- 1) następuje rozumienie wypowiedzi przez jej nadawcę, tzn. rozumienie tego, co nadawca zamierza zrobić,
- 2) nadawca ubierze swoje zamiary w słowa,
- 3) odbiorca właściwie zrozumie intencje nadawcy.

Zdarza się, że dwa różne akty — jeden społeczny a drugi solitarny — wykonywane są podczas wypowiadania tego samego zdania, a to, o jaki akt chodzi, zależy od kontekstu sytuacyjnego. Np. według Reida we wszystkich językach zeznanie i osąd są wyrażane przez tę samą formę mowy. Jednakże twierdzi on, iż podstawowym i bezpośrednim celem, jakiemu służy posługiwanie się językiem, jest wykonywanie aktów społecznych. Używanie języka do wyrażania aktów solitarnych jest czymś drugorzędnym i nieistotnym.

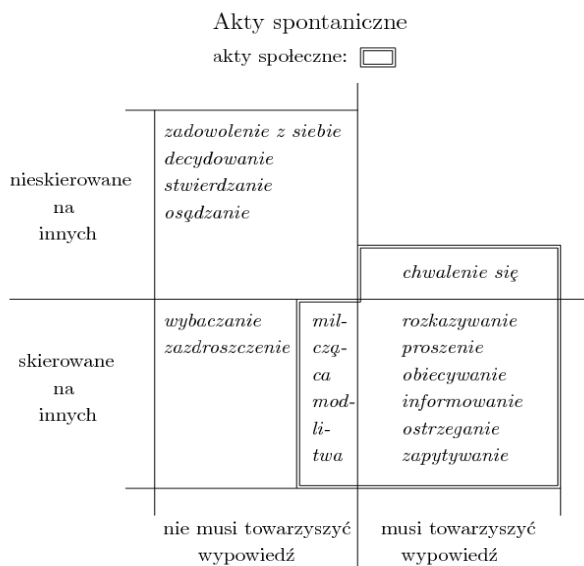
Według Schumanna (1990) koncepcja Reida *aktu społecznego* nie wywarła dużego wpływu na potomnych — nie poświęcił jej bowiem wiele miejsca w swoich pismach. W szczególności Reinach ponoć nie znał jej. Jak zobaczymy, jego pojęcie *aktu społecznego* właściwie niewiele różni się od pojęcia Reida, choć jest znacznie lepiej opracowane.

2. Koncepcja *aktu społecznego* według Reinacha

Reinach wyróżnia klasę doświadczeń umysłowych, które nazywa *spontanicznymi*. Na przykład zwrócenie na coś uwagi bądź uczynienie postanowienia to doświadczenia spontaniczne. Doświadczenia te Reinach przeciwstawia z jednej strony wrażeniom bólu czy dźwięków, które są *narzucane* umysłowi przez czynniki zewnętrzne wobec niego oraz w doświadczaniu których umysł zachowuje się pasywnie, z drugiej strony — takim doświadczeniom jak poczucie szczęścia, entuzjazmu lub oburzenia, które, choć umysł zachowuje się tu aktywniej, to jednak powodowane są przez zewnętrzne czynniki. Wewnętrzne, stosunkowo krótkie działanie się doświadczenia spontanicznego powodowane jest przez podmiot w sposób właśnie spontaniczny, bez przyczyny czy czynnika zewnętrznego i stąd bierze się nazwa dla całej klasy tych doświadczeń. *Intencjonalność* (używamy tego pojęcia w sensie Brentana ([1874] 1999) — jest cechą charakterystyczną każdego doświadczenia psychicznego, mianowicie każde doświadczenie czy zjawisko psychiczne, w przeciwieństwie do zjawisk fizycznych, jest skierowane na pewien obiekt czy dotyczy pewnego obiektu,

indywidualnego, stanu rzeczy itd.) przysługująca doświadczeniom spontanicznym nie jest czymś dla nich charakterystycznym. Na przykład uczucia żalu czy nienawiści, które nie są doświadczeniami spontanicznymi (bo są pewnymi długo trwającymi postawami czy nastawieniami psychicznymi, w przeciwieństwie na przykład do uczucia zazdrości, które jest doświadczeniem spontanicznym), także są intencjonalne, bo skierowane na pewien obiekt. Spontaniczność należy również odróżniać od *aktywności*. Na przykład *oburzenie* będące niespontanicznym doświadczeniem można nazwać doświadczeniem cechującym się aktywnością wychodzącą z podmiotu, w przeciwieństwie do doświadczenia *smutku* (również doświadczenia niespontanicznego), który jakby przychodzi do podmiotu z zewnątrz, bez jego aktywności. Również *posiadanie* jakiegoś postanowienia uznaje Reinach za doświadczenie aktywne, ponieważ to podmiot je posiada. Jednakże odróżnia wyraźnie posiadanie postanowienia — jako stan umysłu — od ustalania postanowienia jako czynności poprzedzającej ów stan i będącej właśnie doświadczeniem spontanicznym. Wymienia on inne czynności (doświadczenia) spontaniczne: decydowanie, preferowanie, przebaczenie, chwalenie, obwinianie, stwierdzanie, pytanie, rozkazywanie.

Akty spontaniczne Reinach najpierw dzieli na takie, którym nie musi towarzyszyć głośna wypowiedź słowna (np. decydowanie, zazdrośczenie czegoś komuś, przebaczenie komuś) oraz na takie, którym taka wypowiedź musi towarzyszyć (np. rozkazywanie czy zapytywanie). Następnie wyróżnia akty spontaniczne *skierowane na podmiot* je wykonujący (np. akt zadowolenia z siebie, akt nienawiści do siebie) oraz *skierowane na inne podmioty* (np. akt przebaczenia, zazdrośczenia, rozkazywania). Z kolei wyróżnia akty *adresowane* do innego podmiotu (np. rozkazywanie). Przebaczenie nie jest aktem adresowanym do innego podmiotu, choć jest aktem skierowanym na inny podmiot. Chwalenie się jest aktem spontanicznym skierowanym do osoby, która go wykonuje, lecz adresowanym do słuchacza. Wreszcie Reinach dokonuje podziału aktów spontanicznych: na takie, które nie muszą być słyszane i rozumiane przez inny podmiot (np. przebaczenie), oraz takie, które muszą być słyszane i rozumiane przez inny podmiot (np. rozkazywanie). Te ostatnie Reinach nazywa *aktami społecznymi*. Zatem wszystkie akty spontaniczne, którym musi towarzyszyć głośna wypowiedź słowna (bez względu na to, czy są one skierowane na inne podmioty czy nie), są aktami społecznymi (bowiem po to konieczna jest słowna głośna wyrażalność niektórych aktów spontanicznych, aby były one słyszane i rozumiane przez inny podmiot). Ponadto wydaje się i nie jest to pewne, że wszystkie akty społeczne to akty adresowane do innego podmiotu oraz odwrotnie — wszystkie akty adresowane do innego podmiotu to akty społeczne.



Reinach zwraca następnie uwagę na dwie strony czy części aktu społecznego: wewnętrzną, czyli duchową i zewnętrzną, czyli fizyczną. Owa wewnętrzna jest bardziej istotna — ostatecznie akt społeczny jest czynnością umysłu. Zewnętrzna czy fizyczna strona aktu społecznego jest całkowicie podporządkowana woli wykonującego akt społeczny, choć mówiący bierze pod uwagę możliwości zrozumienia swojej wypowiedzi przez słuchacza. Zewnętrzna strona aktu społecznego dokonuje się przez mówienie, robienie min i gestów, lecz wszystkie te fizyczne czynności są kontrolowane przez umysł wykonującego akt, w przeciwieństwie — na przykład — do niezależnych od woli oznak wstydu, zdenerwowania, miłości itd. Z drugiej strony należy odróżniać fizyczną część aktu społecznego od opisów doświadczeń psychicznych mających teraz miejsce, dokonywanych przy użyciu zdań: „Nie chcę tego robić”, „Boję się”. Te doświadczenia odbywają się, czy mają miejsce, bez takich opisów językowych, tymczasem akt społeczny musi posiadać swą część językową. Wewnętrzna część aktu społecznego nie może zaistnieć bez zewnętrznej wypowiedzi. Żadna z tych dwóch części nie jest aktem społecznym — takim aktem jest ich zespolenie. Zewnętrzna wypowiedź nie jest czymś opcjonalnym, co może, lecz nie musi być dodane do części wewnętrznej. Jest konieczna co do swojej formy. Naturalnie części fizycznej aktu społecznego nie wolno mylić z opisem takiego aktu, np. z wypowiedzią: „Właśnie wydałem rozkaz.”.

Istotą aktu społecznego jest bycie zrozumianym przez inny podmiot niż ten, który aktu dokonuje. To znaczy, że część zewnętrzna aktu społecznego nie byłaby konieczna, gdyby ów inny podmiot rozumiał doświadczenie psychiczne (czyli część wewnętrzną aktu) zachodzące w dokonującym akt społeczny bez jego słów mówionych głośno lub pisanych. Wobec tego na przykład milcząca

modlitwa *adresowana* do Boga jest aktem społecznym — mimo braku zewnętrznej (fizycznej) części.

Akt społeczny ma swoją *treść* oraz swój *cel*. Zobaczmy to na przykładzie aktu informowania. Mogę być przekonany co do zachodzenia pewnego stanu rzeczy i mogę zachować to przekonanie dla siebie. Mogę również wyrazić to przekonanie w postaci stwierdzenia — głośno przy użyciu zdania, lecz tylko dla siebie. Nie jest to informowanie. Zaadresowanie do kogoś innego tego zdania jest istotne dla czynności informowania. Do istoty informowania należy zaadresowanie wypowiedzi o moim przekonaniu do innej osoby i ogłoszenie jej tej osobie tak, aby mogła ona zrozumieć treść wypowiedzi, tzn. *treść czynności informowania*. Gdy ta osoba zrozumie tę treść, *cel czynności informowania* został osiągnięty i, jak mówi Reinach, obwód rozpoczęty wysyłaniem wypowiedzi zostaje zamknięty.

Proszenie i rozkazywanie są aktami społecznymi podobnymi do siebie ze względu na ich części zewnętrzne: tych samych słów używa się prosząc i rozkazując; różnica polega na sposobie mówienia, w emfazie i ostrości. Proszenie i rozkazywanie mają również swoje *treści*, tak jak informowanie. Lecz w przypadku informowania tylko prezentacja jego treści adresatowi jest istotna, natomiast w przypadku proszenia czy rozkazowania istotna jest nie tylko prezentacja treści, lecz całych tych aktów. Adresat musi być świadomy, że ktoś go prosi o coś lub coś mu rozkazuje; tymczasem nie musi być świadomy tego, że jest informowany. Co więcej, nawet po uświadomieniu adresata o charakterze aktu, w przypadku prośby czy rozkazu, obwód nie jest jeszcze zamknięty. Będzie zamknięty dopiero po wykonaniu pewnych czynności przez adresata (słuchacza) „przepisanych” prośbą bądź rozkazem. Zatem celem proszenia i rozkazowania jest wykonanie owych „przepisanych” czynności przez adresata.

Oprócz *treści* i *celu* Reinach wyróżnia jeszcze trzeci parametr charakteryzujący akty społeczne: *postawę psychiczną* w stosunku do pewnego obiektu, w terminologii Reinacha, *presuponowaną* przez akt społeczny. Z konieczności, *a priori*, każdy akt społeczny presuponuje jako swoją podstawę pewne doświadczenie psychiczne, którego *intencjonalny obiekt* (w sensie Brentana [1874] 1999) jest albo intencjonalnym obiektem samego aktu społecznego, albo jest jakoś z nim związany. Informowanie presuponuje *przekonanie* (co do treści tego aktu). Zapytywanie wyklucza przekonanie i wymaga niepewności co do sytuacji, o którą się pytamy. Prośba presuponuje *życzenie* tego, o co się prosi i aby było ono spełnione przez osobę, którą się prosi. Rozkazywanie oprócz *życzenia* presuponuje *wolę*, aby ten, komu rozkazano, wykonał rozkaz (tzn. wykonał to, co jest treścią rozkazowania). Obietnica presuponuje *zamiar* wykonania przez obiecującego tego, co zostało obiecanie.

Z kolei Reinach dzieli akty społeczne na *proste* (np. informowanie), następnie takie, które *presuponują* inne akty społeczne (np. odpowiadanie na

pytanie jest aktem presuponującym akt zapytywania), oraz takie, których celem jest wykonanie czynności zewnętrznych bądź innych aktów społecznych (np. proszenie, rozkazywanie; zapytywanie jest takim aktem społecznym, który domaga się od adresata wykonania czynności w postaci aktu społecznego odpowiadania na pytanie).

Przypadków tzw. konwencjonalnych pytań (gdy znana jest pytającemu odpowiedź) lub próśb przeciwnych do życzeń Reinach nie uznaje za akty społeczne pytań i próśb — są to modyfikacje takich aktów i są one *pseudo-wykonywane*. Efekty takich modyfikacji aktów nie presuponują odpowiednich postaw psychicznych i dlatego są wadliwe. Reinach nazywa je *pseudoaktami*. Na przykład pseudoinformacja nie presuponuje przekonania, nazywa się ona kłamstwem; pozostałe przypadki pseudoaktów nie mają swoich nazw.

Kolejne rozróżnienie wśród aktów społecznych dowodzi, iż to Reinach, nie Searle, jest pierwszym odkrywcą *złożonych warunkowych* aktów illokucyjnych (ang. *complex conditional illocutionary acts*). Otóż Reinach dzieli wszystkie akty społeczne na warunkowe i bezwarunkowe. Akt warunkowy to czynność społeczna wykonywana w wypadku (*przypadku*), *gdy...*, czy może lepiej: *pod warunkiem, że ...* (ang. *in the event that*). Reinach twierdzi, iż dla niektórych (nie wszystkich) typów aktów społecznych bezwarunkowych istnieją ich warunkowe odpowiedniki. Np. można rozkazywać w wypadku, *gdy...* lub prosić w wypadku, *gdy...* czy pod warunkiem, *że...* Nie można jednak informować pod warunkiem, *że...*

Warunkowe akty społeczne są rzeczywiście wykonywane, ale skuteczność (*efficacy*) ich wykonania jest związana z czymś, co się zdarzy później. Nikt oczywiście nie powinien pomylić warunkowego wykonania z ogłoszeniem wykonania w przeszłości. Nie ma żadnego problemu z tym rozróżnieniem. Gdy dany wypadek [warunek] wystąpi, skuteczność aktu jest — bez żadnych dalszych starań ze strony wykonawcy — taka, jakby bezwarunkowy akt był wykonywany. Zaś od chwili, gdy jest jasne, że ów wypadek [warunek] nie wystąpi, sprawy mają się tak, jakby żaden akt nigdy nie był wykonywany. Jest istotnie wymagane, aby wypadek [warunek], od którego zależy skuteczność aktu warunkowego, mógł się wydarzyć [mógł zachodzić], lecz jest niemożliwe, by musiał się wydarzyć [musiał zachodzić] (Reinach [1913]1983: 23).

Według Reinacha, gdyby ów wypadek, od którego zależy wykonanie aktu, nie mógł się wydarzyć, to niczego by nie warunkował, zatem mówienie o warunkowości aktu nie miałoby sensu. Jeśli zaś ów wypadek musiałby się zdarzyć, to mimo istnienia relacji uwarunkowania, akt nie byłby warunkowy. Byłby bezwarunkowym aktem, w którego *treści* występuje uwarunkowanie, mianowicie odniesienie do momentu, gdy wypadek nastąpi. Czyli wówczas mówiąc: „jeżeli ten oto wypadek nastąpi, to rozkazuję ci zrobić to lub tamto”, mówimy równoznacznie:

Rozkazuję ci (bezwarunkowo) zrobić to lub tamto, w momencie, gdy wypadek nastąpi. Tutaj nie mamy do czynienia z modyfikacją aktu [na warunkowy], lecz tylko zmianą treści aktu. Prócz bycia czasowo zależną, treść aktu może być okresem warunkowym. Odróżniamy, tak wyraźnie jak tylko to jest możliwe, warunkowość treści [aktu] od warunkowości aktu. Bezwarunkowy rozkaz o warunkowej treści bezpośrednio domaga się realizacji pewnego działania [słuchacza], gdy tylko przyszłe możliwe zdarzenie nastąpi. On bezpośrednio tworzy — przy pewnych założeniach — obowiązek zrobienia lub zaniechania czegoś, gdy zdarzenie wystąpi; wystąpienie zdarzenia po prostu uobecnia obowiązek. Przeciwnie, warunkowy rozkaz z bezwarunkową treścią czyni działanie obowiązkowym tylko wtedy, gdy zdarzenie nastąpi i tylko wówczas tworzy obowiązek bezpośredniego działania lub zaniechania (Reinach [1913]1983: 23).

Aby pogłębić rozumienie ostatniego rozróżnienia, warto użyć odpowiedniej notacji, por. Searle, Vanderveken (1985), dla aktu illokucyjnego warunkowego: $P \Rightarrow F(Q)$ oraz elementarnego aktu illokucyjnego, którego treść propozycjonalna jest okresem warunkowym: $F(P \rightarrow Q)$, gdzie P , Q reprezentują sądy w sensie logicznym, czyli treści oraz F — moc illokucyjną, np. prośby. Na przykład akt warunkowy: *Jeżeli wygrasz na loterii, to proszę cię o pożyczkę 100 tys. zł* jest różny od aktu bezwarunkowego: *Proszę cię [o spełnienie stanu rzeczy, że], jeżeli wygrasz na loterii, to pożyczysz mi 100 tys. zł*.

Następnie Reinach rozważa akty społeczne adresowane do wielu słuchaczy oraz takie, które są jednocześnie wykonywane przez wiele osób. Stwierdza on, że pojedynczy akt społeczny może mieć wielu adresatów, zaś jego skutki są inne niż skutki wielu aktów, z których każdy jest adresowany do jednego słuchacza. Na przykładzie rozkazu widać, iż w tym ostatnim przypadku powstaje tyle obowiązków, ilu jest słuchaczy. Tymczasem, gdy jeden akt rozkazu jest adresowany do wielu, powstaje *jeden* obowiązek dzielony przez wielu.

Trudniejszy i bardziej interesujący jest przypadek, gdy kilka osób razem wykonuje jeden akt społeczny. Każda z osób wykonuje akt, np. rozkazywania i każda wyraża [słownie] to wykonywanie. Ale każda wykonuje ten akt „razem z innymi”. Mamy tutaj bardzo specjalny rodzaj wspólnoty [*togetherness*]. [...] Mamy raczej do czynienia tutaj z przypadkiem, gdy każda z osób wykonuje ten akt „w sumie” [*in union*] z innymi osobami, gdzie każda wie o uczestnictwie innych, pozwala innym uczestniczyć i uczestniczy sama: mamy *jeden pojedynczy* akt, który jest wykonywany przez dwie lub więcej osób razem, jeden akt przez kilka podmiotów (Reinach [1913]1983: 24).

Okazuje się, że w przypadku takiego grupowego aktu rozkazywania pojawia się tylko *jedno pojedyncze wymaganie* (wypełnienia nakazanego obowiązku) dzielone przez członków grupy rozkazującej.

Kolejne, ostatnie już rozróżnienie Reinacha może wydawać się ważniejsze dla prawników niż dla językoznawców. Jednakże wyróżniona tu klasa aktów społecznych jest bardzo osobliwa i winna zainteresować tych drugich. Reinach

dzieli akty społeczne na wykonywane w swoim imieniu oraz na takie, które są dokonywane *przez pełnomocnika w czyimś imieniu*. „Pełnomocnik wykonuje czynność całkiem osobiście, ale w taki sposób, że czynność jest prezentowana jako ostatecznie pochodząca od innej osoby” (Reinach [1913] 1983: 25). Reinach wyraźnie odróżnia rozkaz wydany w czyimś interesie lub w zastępstwie kogoś od rozkazu wydanego w czyimś imieniu. Bowiem w tym ostatnim przypadku żądanie wypełnienia obowiązku przez słuchacza pojawia się nie po stronie osoby wykonującej psychofizycznie czynność rozkazywania, lecz po stronie osoby, w imieniu której wydano rozkaz. W przypadkach rozkazu wydanego w czyimś zastępstwie lub interesie owo żądanie pojawia się — tak jak w wypadku zwykłego rozkazu — po stronie wykonującego czynność rozkazywania.

3. Akty społeczne Reinacha w świetle współczesnych teorii aktów mowy

Spróbujmy najpierw porównać pojęcia: *aktu społecznego* (lub szerzej — *aktu spontanicznego*) z *aktem illokucyjnym*. Już w efekcie samego wymienienia przykładów różnych aktów społecznych, jak i różnych aktów spontanicznych niebędących aktami społecznymi widać, że nie można utożsamić *aktu społecznego* Reinacha z *aktem illokucyjnym* (ani w wersji Austina, ani Searle’a). Np. *przebaczenie*, *decydowanie*, *stwierdzanie* nie są aktami społecznymi, tymczasem według Searle’a są elementarnymi aktami illokucyjnymi, odpowiednio — *deklaracji* i *asercji* (zob. Searle, Vanderveken 1985), zaś według Austina — aktami illokucyjnymi (zob. Austin [1962] 1993). Tymczasem nie widać przykładu aktu społecznego, który by nie był aktem illokucyjnym. Taki pobieżny rzut oka zdaje się świadczyć o tym, iż zbiór wszystkich aktów społecznych jest właściwą podklasą klasy wszystkich czynności illokucyjnych. Natomiast opierając się wyłącznie na przykładach czynności, można stwierdzić, że klasa wszystkich czynności spontanicznych nie zawiera się w klasie wszystkich aktów illokucyjnych, np. skierowanie czy skupienie na czymś uwagi bądź uczucie zazdrości są doświadczeniami spontanicznymi według Reinacha, nie są jednak wymieniane nigdzie (tzn. ani przez Austina, ani przez Searle’a, ani przez ich uczniów) jako akty illokucyjne.

Chcąc porównać odpowiednie klasy aktów, bazując na ich definicjach, napotkamy istotne trudności. Rzecz nie dotyczy Austina, Searle’a i jego następców, którzy choć nigdzie wyraźnie nie definiują pojęcia *czynności illokucyjnej* (poza jego matematyczną reprezentacją w tzw. logice illokucyjnej — Vanderveken 1990–91), to jednak na podstawie ich opisów względnie łatwo jest takie definicje podać. Trudności porównawcze biorą się z braku precyzji i staranności pojęciowej Reinacha w niektórych obszarach jego rozważań. Zauważmy po pierwsze, że autor ten uznaje pewne czynności umysłowe niespontaniczne za akty społeczne (co jest sprzeczne z definicją aktu społecznego): oto akt społeczny odpowiadania na pytanie jest — w jego rozumieniu — *powodowany*

aktem zapytywania i wobec tego nie może być uznany za czynność spontaniczną (*czynność spontaniczna* nie została precyzyjnie nigdzie zdefiniowana, ale dysponując podanymi przez Reinacha przykładami i jej opisem, nikt chyba nie zaliczy *odpowiadania na pytanie* do doświadczeń umysłowych spontanicznych). Ponadto zazwyczaj akt wykonywany przez pełnomocnika w czyimś imieniu bywa zaplanowany i wcześniej przemyślany. Nie można go zaliczyć do czynności spontanicznych. Po drugie, Reinach postrzega pewne akty społeczne jako w ogóle niebędące doświadczeniami umysłowymi (co również jest sprzeczne z definicją aktu społecznego). Na przykład wypowiedzianie „grupowo”, w jednym czasie i miejscu przez wiele osób jednego zdania rozkazującego jest traktowane przez niego jako wykonywanie *jednego* aktu społecznego rozkazowania, mimo tego, że umysłów jest tu wiele i wobec tego wiele jest różnych doświadczeń umysłowych.

Przede wszystkim nie jest niczym uzasadnione, aby *spontaniczność* — w rozumieniu Reinacha — była cechą mającą przysługiwać wszystkim takim czynnościom umysłowym, które *muszą być rozumiane* przez inne umysły, tzn. czynnościom społecznym. Cóż takiego jest w tej cesze, że musi ona przysługiwać tym czynnościom, które przecież nieprojektująco nazywa Reinach *społecznymi* — skoro *konieczność rozumienia* przez inne umysły jest jedyną cechą istotną tych czynności. Jeśli *spontaniczność* nie jest *przygodną* cechą aktów społecznych (a tekst Reinacha wyraźnie akcentuje, że nie jest to cecha przygodna) i nie jest jednocześnie ich cechą istotną, to jest ona cechą, której przysługiwanie aktom społecznym można wyprowadzić z wiedzy o ich cechach istotnych. Jednakże przecież między *spontanicznością* a *koniecznością bycia zrozumiałym przez innych* nie ma żadnego związku logicznego — są to cechy zupełnie ze sobą niezwiązane. Innymi słowy, konieczność zrozumienia tego, na czym czynność polega, jaki jest jej cel, czym ona jest, to wyróżniki charakterystyczne czynności zwanej przez Reinacha *społeczną*. Tymczasem spontaniczność, a więc niezdeterminowanie, efekt wolnej woli, samoistność są czymś zupełnie przygodnym dla czynności społecznej. Po takim „ustawieniu” koncepcji Reinacha *aktu społecznego*, aby dokonać jej porównania, przejdźmy do określeń *czynności illokucyjnej* w teoriach Austina oraz Searle’a.

Ogólnikowo i nieprecyzyjnie można powiedzieć, że jest to czynność, do której wykonania używa się jako narzędzia jakiegoś zdania języka etnicznego, zaś użycie tego narzędzia polega na jego wypowiedzeniu (napisaniu). Akt illokucyjny nie jest tożsamy z czynnością wypowiedziania zdania, jest jednak czynnością, która jakby równolegle do tego wypowiedziania, wraz z nim i tylko dzięki niemu się dokonuje. Na ogół bezpośrednim (bynajmniej nie ostatecznym czy finalnym) celem wypowiedziania zdania jest jego zrozumienie przez słuchacza(y). Czynność illokucyjna jest wykonywana w trakcie wypowiedziania zdania nie dlatego, że oto dźwięki o określonej częstotliwości jej towarzyszą (bądź

znaki o odpowiednich kształtach na papierze), lecz przecież dzięki rozumieniu owego zdania. Zatem, na pierwszy rzut oka, akt illokucyjny byłby czynnością, której cechą istotną jest, sformułujmy to ostrożnie, *możliwość* (nie: *konieczność*, jak w przypadku aktu społecznego Reinacha) bycia rozumianym przez słuchacza. Jak zobaczymy po bardziej wnikliwej analizie tekstów Austina i Searle'a, cecha ta, choć niewątpliwie istotna, w tak oczywisty sposób przysługuje aktom illokucyjnym, że nie akcentuje się jej wcale w ich definicji.

Z pracy Austina ([1962] 1993) można wydobyć dwa różne określenia aktu illokucyjnego (por. Nowak 2003), pierwsze: jest to czynność wymieniona w definicji zdania performatywnego (dokonawczego, ang. *performative sentence*): zdanie nazywamy *performatywnym* wtedy, gdy jego wypowiedzanie jest wykonywaniem pewnej *czynności*, którego efektem jest powstanie nowego stanu rzeczy, którego nie można do świata wprowadzić inaczej niż tylko przez wypowiedzenie tego zdania (jest to druga, szersza definicja zdania performatywnego; w pierwszej Austin akcentował ponadto wystąpienie w zdaniu performatywnym nazwy *czynności* wykonywanej przez jego wypowiedzanie, nazwy w postaci tzw. czasownika performatywnego). Ponieważ, jak próbuje się to uzasadnić w pracy *Formalna reprezentacja pojęcia sądu (dla zastosowań w teorii aktów mowy)* (Nowak 2003), każde zdanie jest performatywne, innych zdań nie ma (Austin nigdy tak wyraźnie tego nie sformułował), czynność illokucyjna byłaby tą, której wykonywanie jest wypowiedzaniem jakiegoś zdania, przy czym rezultatem tego wypowiedzania byłoby zaistnienie w świecie pewnego stanu rzeczy, który w żaden inny sposób nie mógłby zaistnieć, gdyby tego zdania nie wypowiedziano. A więc wypowiedzenie zdania „Będę tam na pewno” jest wykonaniem czynności illokucyjnej *obietowania*, którego rezultatem jest zaistnienie stanu rzeczy, że „obiecano być tam”, który w żaden inny sposób nie mógłby zaistnieć bez tej wypowiedzi. Wypowiedzenie zdania „Piję teraz kawę” jest wykonaniem czynności illokucyjnej *stwierdzania*, którego efektem jest zaistnienie w świecie stanu rzeczy, że „stwierdzono picie kawy” (naturalnie różnego od stanu rzeczy „picia kawy” — zupełnie niezależnego od owego stwierdzania), niemogącego zaistnieć w żaden inny sposób, jak tylko dzięki tej wypowiedzi itd.

Drugie określenie aktu illokucyjnego: czynność illokucyjna jest czynnością lokucyjną (tzn. czynnością wypowiedzania zdania ze zrozumieniem), do której jest dołączony zamiar czy która kierowana jest zamiarem osiągnięcia pewnego celu. Czynność lokucyjna nie musi mieć celu — nie musi być czynnością celową. Tak jak podnoszenie do góry rąk nie musi być czynnością celową — wówczas, gdy nie jest, nazywa się potocznie „przeciąganiem się”. Wydłubywanie w bloku marmuru dziur i jego ociosywanie nie musi mieć żadnego celu — może być bezmyślną zabawą. Ale podnoszenie do góry rąk sterowane zamiarem akceptowania lub odrzucenia ustawy jest czynnością celową, zwaną głosowaniem,

zaś ociosywanie bloku marmuru z zamiarem uformowania go w pewien sposób jest już czynnością celową, zwaną rzeźbieniem. W taki sposób czynność lokucyjna, do której dołącza się zamiar osiągnięcia pewnego celu: jest nim zaistnienie w świecie owego nowego stanu rzeczy (rozważanego w związku z pierwszym określeniem czynności illokucyjnej), jest właśnie czynnością illokucyjną.

Jak widać, zarówno pierwsze, jak i drugie Austinowskie określenia *czynności illokucyjnej* zakładają *zrozumienie* zdania, którego wypowiedzenie jest konieczne, aby czynność illokucyjna była wykonana. W przypadku pierwszego określenia to założenie nie jest oczywiste, ale musi obowiązywać, inaczej bowiem ów nowy stan rzeczy pojawiający się w świecie bezpośrednio po wypowiedzeniu zdania nie miałby tego charakteru, jaki ma, np. nie byłoby stanów rzeczy, że obiecano czy stwierdzono, gdyby odpowiednio nie rozumiano zdań, których wypowiedzanie służy zaistnieniu tych faktów. W przypadku drugiej definicji czynności illokucyjnej *zrozumienie* wypowiedzi przez mówcę jest jawnie zakładane. Jest ono atrybutem czynności lokucyjnej i dlatego stanowi warunek konieczny, choć niewystarczający wykonania czynności illokucyjnej: jeżeli wykonano czynność illokucyjną, to tym samym wykonano czynność lokucyjną, a więc rozumiano wypowiedziane zdanie. Jednakże można wykonać czynność lokucyjną, tym samym rozumieć wypowiedź, nie wykonując illokucyjnej. Dla większej jasności zauważmy jeszcze, że Austin nie zakłada *konieczności* rozumienia zdania przez słuchacza, w każdym razie nie zakłada tego dla każdego aktu illokucyjnego. Można coś *nazwać*, coś *stwierdzić* czy coś komuś *wybaczyc* bez słuchacza. Jednakże zakłada *konieczność* rozumienia wypowiedzianego zdania przez mówcę.

Biorąc pod uwagę prace Searle'a ([1969] 1987; Searle, Vanderveken 1985), można zdefiniować (zob. Nowak 2003) *elementarny akt illokucyjny* jako *czynność wyrażania sądu*, do której dołącza się zamiar osiągnięcia *celu illokucyjnego*. Celem illokucyjnym (ang. *illocutionary point*) czynności wyrażania sądu jest ustanowienie *stosunku przystosowania* między tym sądem a reprezentowanym przezeń stanem rzeczy w świecie fizycznym.

Sąd (ang. *proposition*) jest tu rozumiany jak fregowska *Gedanke* (zob. Frege [1918] 1977), czyli myśl wyrażona zdaniem (Searle [1969] 1987) w ogóle pojęcia *sądu* nie definiuje, określa tylko *wyrażanie sądu*). Myśl zaś jest abstrakcyjnym wytworem czynności myślenia, reprezentującym stan rzeczy ze świata. Np. sąd wyrażony zdaniem „Obecnie w Warszawie pada deszcz” jest pewnym przedmiotem abstrakcyjnym, różnym od i reprezentującym fizyczny stan rzeczy (fakt) opisywany tym samym zdaniem. Tak rozumiany *sąd* należy odróżniać od *osądu*. W języku polskim, w wielu kontekstach („jaki jest twój sąd na ten temat?”) wyraz „sąd” występuje właśnie w znaczeniu słowa „osąd”. Tymczasem tutaj *osąd* jest wyrażaniem *sądu* z przekonaniem, z przyporządkowaniem mu

(sądowi) wartości logicznej prawdy albo fałszu. Osąd jest elementarną czynnością illokucyjną. Samo *wyrażanie sądu* to tylko pewne pomyślenie — zestawienie w myśli pewnych rzeczy wraz z towarzyszącym wypowiedaniem zdania. Według Searle’a każde zdanie języka etnicznego (nie tylko zdania oznajmujące) wyraża pewien sąd. Wypowiedzenie zdania jest więc wyrażeniem pewnego sądu (pewnej myśli). Jeśli zaś towarzyszy mu zamiar osiągnięcia celu illokucyjnego, a więc ustanowienia pewnego stosunku przystosowania między tym sądem a reprezentowanym przez ów sąd stanem rzeczy, to owo wypowiedzenie jest wykonywaniem elementarnego aktu illokucyjnego. Z logicznego punktu widzenia istnieją cztery możliwe stosunki przystosowania:

1) *sąd przystosowuje się do stanu rzeczy*: stan rzeczy faktycznie lub domniemanie istniejący w świecie jest obiektem, na który nakierowuje się czy do którego dobiera się sąd mający ten stan rzeczy reprezentować,

2) *stan rzeczy przystosowuje się do sądu*: sąd reprezentujący faktycznie lub domniemanie nieistniejący w świecie fizykalnym stan rzeczy jest obiektem, do którego ma się dopasować ów stan rzeczy przez faktyczne zaistnienie,

3) *sąd i stan rzeczy wzajemnie się do siebie przystosowują*: stan rzeczy reprezentowany przez sąd urzeczywistnia się w trakcie wyrażania tego sądu,

4) *brak przystosowania sądu i stanu rzeczy*.

Ustanowienie stosunku nr 1 to cel illokucyjny *asercji* (stwierdzeń, zeznawań, przypuszczeń, składania sprawozdań, przepowiadania, prognozowań, informowań, przypominania, sugerowań itd.). Ustanowienie stosunku nr 2 jest charakterystyczne dla dwóch typów czynności illokucyjnych, w zależności od tego, czy mówca, czy słuchacz po wykonaniu aktu będzie dzięki swojemu działaniu urzeczywistniał w świecie fizykalnym ten stan rzeczy, który jest reprezentowany przez sąd wyrażony w akcie. Searle rezygnuje tu z utożsamienia ustanowienia stosunku nr 2 z celem illokucyjnym. Cel illokucyjny staje się parametrem służącym do odróżnienia owych dwóch typów czynności illokucyjnych. Mamy zatem cel illokucyjny *zobowiązywać* (obiecywać, grozić, przyrzekać, przysięgać, poręczać, zatrudniać, ślubować, gwarantować itd.), gdy mówca będzie powodował zaistnienie stanu rzeczy reprezentowanego przez wyrażany w akcie sąd, oraz cel illokucyjny *dyrektyw* (prośb, rozkazów, pytań, błagań, proponowań, zapraszań, żądań, radzeń, rekomendowań, ponaglań itd.), gdy słuchacz(e) będzie wprowadzał do świata ów stan rzeczy. Ustanowienie stosunku nr 3 to cel illokucyjny *deklaracji* (nazywań, nominowań, definiowań, rezygnowań, wybaczań, ratyfikowań, decydowań, akceptowań itd.). Wreszcie realizacja stosunku nr 4 to cel illokucyjny *ekspresji* (dziękowań, przepraszań, gratulowań, protestowań, lamentowań, powitań itd.). W czynności ekspresji mówca nie ustosunkowuje wyrażanego sądu do reprezentowanego przezeń

stanu rzeczy, lecz prezentuje słuchaczowi swoją postawę psychiczną (ang. *propositional attitude*) w stosunku do tego sądu.

Można łatwo uzgodnić drugą definicję aktu illokucyjnego Austina z definicją elementarnego aktu illokucyjnego Searle'a, gdyby, po pierwsze, uznać, że czynność lokucyjna Austina jest tożsama z Searle'owską czynnością wyrażania sądu, po drugie, uznać osiągnięcie celu illokucyjnego, czyli odpowiedniego ustosunkowania, jak pisze Searle, „słów” i „świata” (tzn. sądu i stanu rzeczy), a więc że stwierdzono, obiecano, poproszono, ochrzczono, przeproszono itd. za odpowiedni stan rzeczy, którym dana czynność illokucyjna bezpośrednio się kończy — naturalnie, w większości przypadków, różny od stanu rzeczy reprezentowanego przez sąd wyrażany w czynności, choć w przypadku czynności typu deklaracji właśnie z nim tożsamy.

Opisu i wyjaśnienia zróżnicowania czynności illokucyjnych w ramach każdego z tych pięciu wielkich typów Searle dokonuje dzięki dodatkowym, poza celem illokucyjnym, pięciu parametrom. Wszystkie sześć parametrów składa się na tzw. *moc illokucyjną* (ang. *illocutionary force*), tzn. konkretne wartości wszystkich sześciu parametrów wyznaczają jednoznacznie konkretną moc illokucyjną elementarnego aktu illokucyjnego. Jest on jednoznacznie scharakteryzowany mocą illokucyjną F oraz wyrażanym w nim sądem P (Searle czasem zamiast wyrażenia „sąd” używa sformułowania: *sądowa zawartość* aktu illokucyjnego, ang. *propositional content of an illocutionary act*). Elementarny akt illokucyjny F(P) to czynność wyrażania sądu P z mocą illokucyjną F, aby osiągnąć cel illokucyjny mocy F.

Rozumienie zdania, którego wypowiedzanie służy wykonywaniu aktu postaci F(P), jest w sposób oczywisty zakładane (zarówno przez mówcę, jak i słuchacza), bo to z tego zdania muszą być wydobyte (przez mówcę i słuchacza) znaczniki mocy illokucyjnej (ang. *illocutionary force markers*) oraz znacznik sądu (ang. *clause*), aby ustalić (pojąć) moc illokucyjną F oraz sąd P. Można powiedzieć inaczej: sposób *rozumienia* zdania (czyli *znaczenie zdania* według jednej z koncepcji *znaczenia*) polega na ustaleniu tych dwóch składników: F, P, elementarnej czynności illokucyjnej. Jednakże, podobnie jak u Austina, nie w przypadku każdego elementarnego aktu illokucyjnego konieczna jest obecność słuchacza.

Ostatecznie, należy uznać, że wszystkie akty społeczne w rozumieniu Reinacha, czyli czynności cechujące się *koniecznością* rozumienia przez słuchacza towarzyszącej im wypowiedzi, są czynnościami illokucyjnymi. Tymczasem nie wszystkie czynności illokucyjne są społecznymi. Czynności illokucyjne cechuje bowiem *możliwość* rozumienia przez słuchacza wypowiedzi, przez które czynności te się dokonują. Zaś *konieczność* implikuje *możliwość*, lecz nie odwrotnie.

Jak widać, atrybut *rozumienia* przysługujący wypowiedzi służącej do wykonania aktu illokucyjnego jest zarówno w ujęciu Austina, jak i Searle'a banalną

charakterystyką tego aktu. Inne cechy tego aktu uznają ci autorzy za ważne. Można by rzec, iż Reinacha definicyjna charakterystyka *aktu społecznego* nie jest tak dojrzała. Jednakże Reinach poza samą definicją charakteryzuje akty społeczne w sposób zupełnie niebanalny.

Oto jest oczywiste na podstawie podanych przykładów, iż Reinacha *treść* aktu społecznego jest tożsama z Searle'owskim *sądem*. Z kolei Reinachowski *cel* aktu społecznego można postrzegać jako tzw. *cel finalny aktu illokucyjnego* — zob. Nowak (2003). Celem *finalnym* czynności asercji jest jej cel illokucyjny, w którego realizacji dopasowuje się do świata *prawdziwy* sąd. Jeśli sąd wyrażany w trakcie wykonywania asercji jest fałszywy, cel finalny tej czynności nie zostaje osiągnięty. Celem finalnym czynności deklaracji jest dokładnie jej cel illokucyjny. Celami finalnymi zobowiązywań i dyrektyw nie są ich cele illokucyjne, lecz zaistnienia w świecie stanów rzeczy dopasowanych do sądów wyrażanych w tych czynnościach i spowodowanych przez odpowiednie działania mówcy lub słuchacza. Celem finalnym aktu ekspresji jest prezentacja słuchaczowi postawy propozycjonalnej mówcy w stosunku do sądu wyrażonego w tym akcie.

Wreszcie, ostatni parametr — postawa psychiczna w stosunku do obiektu intencjonalnego aktu społecznego — może być interpretowany jako jeden z sześciu parametrów charakteryzujących *moc illokucyjną*, mianowicie *sincerity conditions of an illocutionary act*, czyli typy postaw psychicznych, jakie mówca przejawia w stosunku do sądu wyrażonego w trakcie wykonywania aktu illokucyjnego — tutaj chodzi nie o wszystkie owe typy, lecz tylko te, które są zdeterminowane przez *cel illokucyjny* (zob. np. Vanderveken 1990–91).

Z kolei Reinachowska analiza czynności społecznych warunkowych różni się od Searle'owskiej tylko pod jednym względem: Reinach uznaje akt warunkowy postaci: $P \Rightarrow F(Q)$ za będący tego samego typu co akt $F(Q)$. Mówi o warunkowych czynnościach prośby czy rozkazu. Tymczasem Searle wrzuca wszystkie akty postaci $P \Rightarrow F(Q)$ do jednego worka: tzw. *illokucyjnych aktów złożonych warunkowych* (jest to jeden z trzech, obok *koniunkcji* i *denegacji*, wyróżnianych przez niego typów aktów illokucyjnych złożonych, zob. Searle, Vanderveken (1985), Nowak (2000)). Żaden z nich nie jest aktem prośby czy rozkazu, bowiem te akty zaliczane są do *illokucyjnych aktów elementarnych*.

Reinach niestety nie zwraca uwagi na pewną charakterystyczną, wręcz definicyjną osobliwość aktów społecznych wykonywanych przez pełnomocnika w czyimś imieniu. Czy czynność wykonywana przez pełnomocnika w czyimś imieniu, będąc czynnością umysłową (bo społeczną), jest *tego typu* czynnością, którą wykonałby ów ktoś, w czym imieniu działa pełnomocnik, gdyby działał samodzielnie? Nie pytamy, czy czynność wykonywana przez pełnomocnika w czyimś imieniu jest właśnie *tą* czynnością, którą wykonałby ów ktoś, w czym imieniu działa pełnomocnik, gdyby działał samodzielnie. Bowiem odpowiedź na

to ostatnie pytanie jest oczywista: obie czynności są różne, skoro są to czynności psychiczne różnych umysłów. Pytamy, czy doświadczenie umysłowe pojawiające się w umyśle pełnomocnika wykonującego w imieniu kogoś czynność rozkazywania jest rzeczywiście rozkazywaniem? Czy czynność *obiecywania dokończona przez pełnomocnika w czyimś imieniu* można uznać za *obiecywanie*? Inaczej rzecz ujmując, czy sformułowanie: „przez pełnomocnika w czyimś imieniu” ma charakter *modyfikujący* czy *atrybutywny* w sensie Brentana [1874] (1999), Twardowskiego [1894](1965)? Dany przymiotnik w danym kontekście ma charakter *modyfikujący*, gdy zakres (zbiór obiektów odniesienia) nazwy złożonej powstałej przez jego dołączenie do danej nazwy jest rozłączny z zakresem tej danej nazwy. Natomiast przymiotnik ma w danym kontekście charakter *atrybutywny*, gdy zakres takiej nazwy złożonej jest właściwym podzbiorem zakresu owej nazwy, do której przymiotnik został dołączony. Np. „fałszywy” jest przymiotnikiem modyfikującym w wyrażeniu: „fałszywy diament” (żaden fałszywy diament nie jest diamentem), zaś atrybutywnym w kontekście: „fałszywe zdanie” (każde fałszywe zdanie jest zdaniem, lecz nie każde zdanie jest fałszywe). Jeśli wyrażenie „przez pełnomocnika w czyimś imieniu” ma charakter modyfikujący, to akty społeczne wykonywane przez pełnomocnika tworzą nowy typ — są niesprowadzalne do zwykłych aktów społecznych rozkazywania, obiecywania itd. W każdym razie — dla Reinacha wyrażenie to ma charakter atrybutywny.

4. Odrobina ontologii Reinacha: zobowiązanie-wymaganie a obiecywanie

Przedstawimy najpierw kilka ontologicznych Reinachowskich opisów ważnych dla niego obiektów, jakimi są *zobowiązanie* oraz *wymaganie*, jak również czynności *obiecywania* i związków między nimi. Naszym celem będzie sformułowanie pewnych wniosków na temat ontologii Reinacha, pokazujących, w jaki sposób stanowi ona bazę dla jego teorii aktów mowy. Chcemy wskazać odpowiedzi, jakie w tej ontologii znajdujemy, na fundamentalne pytania: *dłaczego (na jakiej zasadzie) możliwe są działania ludzkie dokonujące się dzięki wypowiedaniu i rozumieniu zdań? Na czym te działania polegają?* Następnie próbujemy zakwestionować ontologię Reinacha, w szczególności prezentując konkurencyjną, naszym zdaniem lepiej uzasadnioną, ontologię Searle’a.

Reinach uważa *zobowiązanie* (się do czegoś w interesie kogoś) oraz *wymaganie* (czegoś od kogoś) za obiekty pozaprzestrzenne, lecz trwające w czasie (powstające i ginące). Taką ogólną charakterystykę ontologiczną mają wszelkie obiekty psychiczne, tzn. doświadczenia czy zjawiska mentalne, dzielone na czynności (np. czynność przedstawiania, wyobrażania, osądzania, obiecywania itd.) i stany mentalne (zwane też stanami intencjonalnymi lub postawami psychicznymi, np. bycie przekonanym o czymś, chcenie czegoś, stan radości, stan żalu, uczucie bólu itd.). Jednakże *zobowiązanie* oraz *wymaganie*

nie są, według Reinacha, doświadczeniami mentalnymi — w przeciwieństwie do nich mogą przecież trwać latami. Trwają również w czasie snu osób zobowiązanej i mającej uprawnienie. Każe on odróżniać *poczucie* bycia zobowiązanym, a więc pewien stan psychiczny, od bycia zobowiązanym. Podobnie *wymaganie* nie jest tożsame z uczuciem czy stanem psychicznym *oczekiwania* na realizację zobowiązania.

Ponadto Reinach nie uznaje *zobowiązania* danej osoby jako jej *obowiązku moralnego* ani też nie utożsamia go z *obowiązkiem* innego rodzaju — narzuconego jej przez inne osoby, prawo lub instytucje. *Zobowiązanie* jest bowiem obowiązkiem nałożonym na siebie w wyniku działania wolnej woli i dlatego nie może być obowiązkiem moralnym, a więc takim, który jest nałożony na osobę niezależnie od jej woli, czy innym obowiązkiem narzuconym z zewnątrz (np. obowiązek nauki w szkole). Nie mogą zobowiązać się do czegoś (np. niesienia pomocy innym), czego realizacja jest moralnym obowiązkiem lub wynika z obowiązującego prawa. Ponadto zobowiązania mogą być przenoszone na inne osoby (np. jako efekt dziedziczenia), tymczasem moralne obowiązki — nie. Co więcej, z zobowiązania można się wycofać, np. ogłaszając to oficjalnie, wycofanie się z obowiązku moralnego jest niemożliwe. Wreszcie Reinach akcentuje rozróżnienie: *zobowiązanie* jako obowiązek narzucony przez samego siebie i *obowiązek dotrzymywania zobowiązań*, będący oczywiście obowiązkiem moralnym. Jednakże warto zauważyć, że w szerokim aspekcie ontologicznym *zobowiązania*, *obowiązki moralne* oraz *obowiązki prawne* nie różnią się wcale — mają te same własności: są obiektami niementalnymi, pozaprzestrzennymi, trwającymi w określonym czasie. Reinach analogicznie traktuje *wymaganie* jako pod pewnymi względami różne od *uprawnienia moralnego*, a w aspekcie ontologicznym identyczne z nim.

Zobowiązanie i wymaganie nie istnieją bez swoich nośników: osoby, która się zobowiązała oraz osoby, która wymaga (tzn. ma uprawnienie). Każde zobowiązanie jest zobowiązaniem do wykonania przyszłego działania przez nośnika zobowiązania. To działanie może być robieniem czegoś, zaniechaniem robienia czegoś bądź tolerowaniem robienia czegoś. Owa przyszła akcja nośnika zobowiązania to *treść* zobowiązania. Zobowiązanie o danej treści jest zawsze skierowane do osoby, która ma wymaganie o tej samej *treści*; jest ona adresatem zobowiązania. Jednocześnie to wymaganie jest skierowane do osoby, która się zobowiązała. Tymczasem przyszłe *działanie* osoby, która się zobowiązała, może być skierowane do adresata zobowiązania, lecz nie musi. Może być skierowane do innej osoby, może wcale nie być na nic skierowane. Należy odróżniać skierowanie *zobowiązania* od skierowania *działania*, do którego się zobowiązało.

Zobowiązaniami i wymaganiami rządzą prawa *a priori*. Np. wymaganie, aby coś było zrobione, ma swój koniec, gdy tylko to coś zostanie zrobione. Reinach

uważa, że takie oto prawo jest zasadą *a priori*, bowiem oparte jest ono na istocie wymagania, z konieczności. Jednocześnie, mimo tego, że nie jest ono wnioskiem indukcyjnym z obserwacji, to jednak jest zdaniem syntetycznym, tzn. nie jest analityczne, bowiem w pojęciu wymagania nie jest zawarte w żaden sposób rozwiązanie wymagania. Według Reinacha negacja tego prawa z pewnością jest fałszywa, lecz nie jest sprzeczna.

Według obiegowych opinii „obietcywanie jest deklaracją woli, lub bardziej dokładnie, wyrażeniem lub uczynieniem znanym intencji zrobienia lub zaniechania czegoś w interesie kogoś, do kogo wypowiedź jest skierowana” (Reinach [1913] 1983: 16). Nie wyjaśniano jednak, w jaki sposób wyrażenie intencji zrobienia czegoś w interesie słuchacza powoduje powstanie *zobowiązania* mówcy oraz *wymagania* słuchacza. Zauważmy, że przecież sam zamiar czy wola zrobienia czegoś w interesie kogoś, bez słownego wyrażenia tego zamiaru, nie powoduje powstania pary *zobowiązanie-wymaganie*. Reinach zapytuje: jak to jest możliwe, że obietnica tworzy (ang. *produces*) zobowiązanie i wymaganie, skoro sam zamiar wykonania czegoś dobrego dla słuchacza, zamiar niewypowiedziany, niczego takiego nie tworzy? Jak to jest możliwe, że zobowiązanie i wymaganie wprowadza się do świata za pomocą słów? Na pewno nie jest tak, że wyrażanie językowe jakiegoś mojego postanowienia tworzy moje zobowiązanie. Mogę powiedzieć o jakiejś mojej decyzji komuś, np. że nie będę czytał jego artykułu, wcale w ten sposób do niczego się nie zobowiązując. Czy zatem zobowiązanie powstaje wtedy, gdy mówię komuś o mojej decyzji zrobienia czegoś dobrego dla tego kogoś? Otóż również nie. Mogę bowiem wyrazić swój zamiar owemu komuś w formie informacji: „Informuję cię, że w najbliższym czasie przeczytam twój artykuł” (zakładam, że czynność czytania tego artykułu przez mnie jest czymś dobrym i pożądanym przez słuchacza). Żadne zobowiązanie moje ani wymaganie (żądanie) słuchacza nie zostało tu powołane do istnienia. Według Reinacha informowanie o swojej decyzji zrobienia czegoś dobrego dla kogoś oraz obietcywanie zrobienia tego czegoś może być dokonywane przy użyciu tego samego zdania. Uważa, iż w ogólności postawy psychiczne, a więc radość, smutek, żal, wewnętrzna zgoda czy niezgoda na coś, nie mają wpływu na społeczne stosunki prawne. Zatem również informacja słowna o tych stanach psychicznych nie ma takiego wpływu.

Związek *obietcywania z zobowiązaniem-wymaganie* przedstawia Reinach następująco. Efekt, jaki osiąga mówca, wyrażając obietnicę jest całkiem odmienny od rezultatu, jaki osiąga mówca, informując kogoś bądź pytając o coś. Otóż w efekcie wykonania obietnicy kreowany jest specyficzny stosunek czy związek między mówcą, który tę obietnicę złożył, oraz osobą, której obiecano. Mianowicie mówca jest *zobowiązany* do zrobienia czegoś, zaś słuchacz *oczekuje i wymaga* zrobienia tego czegoś przez mówcę. Ów związek, jak pisze Reinach, może trwać długi okres czasu, lecz do jego istoty należy jego *rozwiązanie*, które

może nastąpić na kilka sposobów: rzecz obiecana zostaje wykonana przez mówcę (jest to naturalne rozwiązanie owego związku), słuchacz rezygnuje z wymagania wobec mówcy, mówca wycofuje się z obietnicy.

Według Reinacha *obietnica* oraz *zobowiązanie-wymaganie* tworzą związek przyczynowo-skutkowy (zauważa on również, że obietnica nie jest jedynym źródłem powstawania zobowiązania oraz wymagania. Innym źródłem jest na przykład czynność pożyczania czegoś od kogoś — ten, kto od kogoś pożycza, jest zobowiązany to coś oddać, zaś ten, kto użycza, ma wymaganie w stosunku do tego, któremu pożyczył). Jednak „przyczyna” jest tu rozumiana w mocniejszym sensie niż zazwyczaj, a więc gdy np. za Hume’em mówimy, że przyczyną dymu jest ogień. Otóż w istocie ognia dym się nie mieści. Tymczasem w istocie tej przyczyny, jaką jest obiecywanie, mieści się zobowiązanie i wymaganie. Ponadto w przypadku zwykłych, występujących w przyrodzie związków przyczynowo-skutkowych można postrzegać skutek niezależnie od jego przyczyny, bez zwracania na nią uwagi. Tymczasem, aby uświadomić sobie zaistnienie zobowiązania i wymagania, umysł musi zlokalizować ich przyczynę (np. akt obiecywania). Po trzecie, tutaj tylko ustalając istnienie przyczyny, jesteśmy w stanie ustalić jej następstwa.

Bowiem — jak objaśnia Reinach — zobowiązanie i wymaganie są ugruntowane w obietnicy w podobny sposób jak abstrakcyjny, matematyczny stan rzeczy ugruntowany jest w innych matematycznych stanach rzeczy (tzn. twierdzeniach, z których wynika). Żeby go „powołać do istnienia”, muszę go udowodnić na podstawie innych stanów rzeczy, podobnie — aby zdać sobie sprawę z zaistnienia zobowiązania i wymagania — muszę odwołać się do obietnicy czy czynności pożyczania czegoś.

W lepszy sposób Reinach jakby nie umie, i ma tego świadomość, ustalić na czym polega specyfika ontologiczna związku *obiecywania* z *zobowiązaniem-wymaganiem*. Być może, czując intelektualny niedosyt, stosuje jeszcze dwie metody zaradzenia mu. Po pierwsze, akcentuje ten związek wielokrotnie, jak świadczą o tym cytaty: „zobowiązanie jest ugruntowane [ang. *grounded*] w istocie aktu społecznego” (Reinach [1913] 1983: 44); „A więc utrzymujemy naszą tezę, że wymaganie i zobowiązanie są ugruntowane z istotną koniecznością [ang. *are grounded with essential necessity*] w akcie obiecywania jako takim”, „Ściśle mówiąc, nie proponujemy żadnej teorii obiecywania. Bowiem jedynie kładziemy nacisk na pojedynczą tezę, że obiecywanie jako takie tworzy [ang. *produces*] wymaganie i zobowiązanie. Każdy może spróbować, i my rzeczywiście spróbowaliśmy, uczynić ją zrozumiałą [ang. *to bring out the intelligibility of this thesis*] dzięki analizie i objaśnianiu. Próbować ją wyjaśniać to tak jakby próbować objaśniać twierdzenie: $1 \times 1 = 1$ ” (Reinach [1913] 1983: 46).

Druga metoda wreszcie wystarczająco zrozumiale naświetla charakter ontologiczny związku *obiecywania* z *zobowiązaniem-wymaganiem*. Mianowicie

Reinach porównuje swoje rozumienie tego związku z rozumieniem innych, w szczególności D. Hume'a oraz T. Lippsa (promotora rozprawy doktorskiej Reinacha).

Hume [1739] (2006) krytykuje naturalny, tj. przyczynowo-skutkowy związek obiecywania z zobowiązaniem. Nie znajduje żadnego aktu umysłu towarzyszącego wypowiedzi obietnicy, który mógłby powodować, w sposób naturalny, powstanie zobowiązania. Samo mówienie — tożsame dla niego z obiecywaniem — niczego nie może spowodować. Wymienia możliwe doświadczenia umysłu, takie jak: *powzięcie decyzji* o wykonaniu czegoś, *życzenie* takiego wykonania, *wola* takiego wykonania, wreszcie *wola zobowiązania się* do czegoś. Wszystkie te doświadczenia umysłowe są niedostateczne, aby być przyczyną powstania obowiązku. Dlatego traktuje powstanie takiego obowiązku, po obiecaniu, jako efekt *konwencji* powziętej przez ludzi w ich interesie.

Reinach stawia Hume'owi dwa zarzuty. Po pierwsze, że nie postrzega on *obiecywania* jako samoistnego aktu umysłu, różnego od mówienia i wyrażania postaw psychicznych. Po drugie, że nie widzi on tego, iż to akt obiecywania w naturalny sposób jest przyczyną powstania obowiązku:

Akt obiecywania nie jest naturalnie tym samym, co wola kogoś do zobowiązania się. Ale mimo wszystko obowiązek w stosunku do innej osoby powstaje przez obiecywanie [*does come about through promising*], obiecujący może być świadomy tego [tzn. woli zobowiązania] w momencie obiecywania i w tym przypadku wola samozwiązania się może równie dobrze towarzyszyć obiecywaniu (Reinach [1913] 1983: 37).

Reinach przytacza koncepcję Theodora Lippsa: obiecywanie to słowne wyrażanie zamiaru wykonania czegoś w czyimś interesie. Gdy ów zamiar zostanie przez słuchacza, któremu się obiecuje, rozpoznany, „wraca” wzmocniony przez interes słuchacza do obiecującego jako jego obowiązek czy powinność.

Reinachowi podoba się w tej koncepcji tylko to, że oto *obowiązek* powstaje tu w naturalny sposób, nie na drodze sztucznej konwencji. Pozostałe elementy tego poglądu uważa za błędne. Przede wszystkim *obiecywanie* nie może być utożsamiane ze *słownym wyrażaniem zamiaru*. Ponadto nie można utożsamiać *obowiązku* z *poczuciem posiadania obowiązku*. Obowiązek nie jest stanem w niczym umyśle, a zamiar, wzmocniony przez interes słuchacza, takim stanem właśnie jest.

Podsumowując porównanie koncepcji związku *obiecywanie* a *zobowiązanie-wymaganie* Reinacha oraz Hume'a i Lippsa, zwróćmy uwagę na to, że podstawowym kryterium różnicującym te stanowiska jest to, czy ów związek ma charakter naturalny czy konwencyonalny. Następnym kryterium podziału jest ontologiczny charakter pierwszego członu związku — obiecywania:

(1) obowiązek (zobowiązanie) powstaje po obietnicy w sposób naturalny, dzięki związkowi przyczynowo-skutkowemu,

(a) przyczyną powstania obowiązku jest samoistny akt społeczny, zwany obiecywaniem (A. Reinach),

(b) przyczyną powstania obowiązku jest pewien stan psychiczny — zamiar, wola, zaś *obiecywanie* jest jego słownym wyrażaniem (T. Lipps).

(2) obowiązek powstaje po obiecywaniu dzięki konwencji, tzn., mówiąc współczesnym językiem Wittgensteina i Searle'a, dzięki konwencjonalnym *regułom konstytucyjnym* „gry w obiecywanie”,

(a) którą zapoczątkowuje (zaczyna grę) odpowiednie sformułowanie słowne, zwane obiecywaniem (D. Hume),

(b) którą zapoczątkowuje samoistny akt illokucyjny obiecywania (J. R. Searle).

Przedstawiciele stanowiska (1) uznają prawo postaci: *akt społeczny obiecywania jest naturalną przyczyną powstania zobowiązania* (postawa (a)) lub postaci: *pewien stan mentalny jest naturalną przyczyną powstania zobowiązania* (postawa (b)) za zdanie syntetyczne, niektórzy, tak jak Reinach, również za zdanie *a priori*. Tymczasem reprezentanci stanowiska (2) uznają zdanie: *obiecywanie powoduje powstanie zobowiązania* za zdanie analityczne.

Nie należy się dziwić, że w kwestii stosunku obietnicy do zobowiązania-wymagania Reinach stoi na stanowisku (1a), jak również uważa, że prawo opisujące ten stosunek ma charakter *a priori*. Jest on przedstawicielem specjalnego nurtu filozoficznego — fenomenologii, której analizy ontologiczne opierały się na tzw. *wglądzie w istotę rzeczy*. Jeżeli jakaś własność czy stosunek „są ugruntowane” w istocie danej rzeczy (jak *zobowiązanie w obiecywaniu*) lub są z tej istoty wyprowadzalne, to są z tą rzeczą związane w sposób naturalny i konieczny, a zdanie opisujące ich związek z istotą jest zdaniem syntetycznym *a priori*. Co prawda fenomenologowie, w szczególności Reinach, nie używali przymiotnika „syntetyczny”, mówili w tym kontekście tylko o zdaniach *a priori*. Prawdopodobnie dlatego, iż nie uznawali Kantowskiego uzasadnienia możliwości istnienia zdań syntetycznych *a priori*. Bowiem, inaczej niż u Kanta, konieczność twierdzeń *a priori* sformułowanych na podstawie wglądu w istotę jest oparta na ich bezpośredniej oczywistości, ujmowanej intuicyjnie. W każdym razie prawa *a priori* rozważane przez Reinacha nie mają według niego charakteru analitycznego, zatem z definicji są syntetyczne.

Podsumowując, można następująco streścić ontologiczne poglądy Reinacha w interesującej nas kwestii. Dla niego obietnica jest tworem naturalnym (zatem niewytworzonym przez człowieka), a jej związek z zobowiązaniem-wymaganiem ma taki status ontyczny, jaki mają twierdzenia matematyczne. Związek ten jest więc jakimś *abstrakcyjnym stanem rzeczy* (por. Nowak 2003) mającym miejsce (zachodzącym) z konieczności, a więc tym samym *a priori*. Co

prawda istnieje w filozofii matematyki stanowisko epistemologiczne (reprezentowane np. przez Ajdukiewicza [1949] (2004)), według którego twierdzenia matematyczne są zdaniami analitycznymi, ale Reinach nie jest jego przedstawicielem. Według niego twierdzenia matematyczne istnieją niezależnie od ludzkich umysłów jako prawa syntetyczne i taki status ma związek obietnicy z zobowiązaniem-wymaganiem. Kwalifikuje go jednak Reinach jako przyczynowo-skutkowy. Rozszerza on klasę związków przyczynowo-skutkowych poza sferę doświadczenia empirycznego.

W ten sposób, według Reinacha, mówienie, a więc wykonywanie aktów społecznych, to postępowanie podlegające naturalnym uniwersalnym rygom, mającym postać praw syntetycznych *a priori*. Podobnie wykonuje się czynność zliczania obiektów indywidualnych przy zastosowaniu dodawania i mnożenia, podlegających prawom o tym samym charakterze. Oczywiście pierwiastek konwencjonalny dla każdego aktu społecznego istnieje. Jest nim takie, a nie inne użycie konkretnego zdania w konkretnym języku etnicznym dla wykonania tego aktu. Lecz prawa, jakim to wykonanie podlega, mają charakter uniwersalny, niezależny od użycia konkretnego języka.

5. Czy ontologia Reinacha jest uzasadniona?

Syntetyczny charakter stosunku obiecywania do zobowiązania-wymagania, a tym samym stanowisko (1) dotyczące tego stosunku można, jak sądzę, łatwo zakwestionować i to przynajmniej na dwa sposoby. Pierwszy z nich nie sprzeciwia się metodzie *wglądu w istotę*, ale opiera się na bardziej prozaicznym niż w fenomenologii rozumieniu sformułowania „wgląd w istotę”.

Najbardziej dla nas oczywista interpretacja wyrażenia „wgląd w istotę rzeczy” jest zupełnie inna niż jego rozumienie przez fenomenologów. Według niej wgląd w istotę rzeczy polega na tworzeniu definicji *sprawozdawczej* nazwy owej rzeczy (zob. Pawłowski 1986). Sprowadza się to do ustalenia *zastanego* znaczenia nazwy (lub inaczej: jej *eksplikacji*). W *definiensie* takiej definicji podajemy właśnie *istotną* charakterystykę desygnatów tejże nazwy, ustaloną na podstawie analizy zastanego znaczenia nazwy. W ten sposób dokonać wglądu w istotę obietnicy to tyle co ustalić zastane znaczenie językowe nazwy „obietnica” — podając definicję nominalną sprawozdawczą tego terminu. Dokonać wglądu w istotę wody to tyle co ustalić czy sprecyzować zastane znaczenie językowe terminu „woda” — podając jego definicję sprawozdawczą.

Jeśli zgodzimy się, że wgląd w istotę rzeczy oznacza ustalenie znaczenia nazwy tej rzeczy przez podanie definicji sprawozdawczej, to pozostaje uznać, że ową istotą jest właśnie to znaczenie, a wreszcie, że to, co ugruntowane w istocie lub co wynika z istoty, to tyle co ugruntowane w znaczeniu lub z niego wynikające. Zatem z tak rozumianej *istoty* wynikać może (lub być w niej ugruntowane) wyłącznie zdanie analityczne, bo z charakterystyki znaczenia

jakiegoś wyrażenia wynikają wyłącznie zdania analityczne, czyli takie, których prawdziwość lub fałszywość jest gwarantowana sposobem rozumienia słów w nich występujących. Ostatecznie, przy takim oczywistym i banalnym rozumieniu wyrażenia „wgląd w istotę”, łatwo się zgodzić, że implikacje tego wglądu są zdaniami niewątpliwie *a priori*, lecz dlatego, że są to po prostu zdania analityczne, nie zaś — jak chce Reinach — syntetyczne. Na dodatkową obronę tego stanowiska zastosujmy jeszcze *argument z autorytetu*. Ajdukiewicz ([1949] 2004: 49) pisze:

To, co fenomenologowie nazywają wglądem w istotę, można też nazwać uważnym wzywaniem się w znaczenie wyrazów. W takim razie jednak twierdzenia zdobyte na drodze wglądu w istotę można też nazwać zdaniami opartymi na uważnym wzywaniu się w znaczenie wyrazów. Zdania zaś na takiej podstawie oparte wyłuszczają tylko znaczenie zawartych w nich terminów i jako takie są zdaniami analitycznymi.

Skądinąd kontrowersyjne jest zagadnienie, czy same definicje sprawozdawcze są zdaniami analitycznymi. Prawdziwe zdanie syntetyczne *a posteriori*: „woda jest to substancja chemiczna, złożona z pierwiastków wodoru i tlenu w stosunku wagowym 1:8” nie jest definicją sprawozdawczą nazwy „woda” (choć jest niewątpliwie definicją realną wody, por. Ajdukiewicz [1956] (1985)). Natomiast zdanie: „woda jest to ciecz, którą pijemy, służąca do rozpuszczania i do mycia, zamarza, zamieniając się w lód, oraz paruje, zamieniając się w gaz; jest w rzekach, jeziorach, oceanach, morzach” może, jak sądzę, uchodzić za definicję sprawozdawczą terminu „woda” — zdającą sprawozdanie z tego, jak należy rozumieć ten termin. Definicja ta mówi tylko tyle, że znaną wszystkim popularną substancję w języku polskim w taki oto sposób się nazywa. Zatem taka definicja daje wiedzę o sposobie rozumienia (czy użycia) definiowanego w niej wyrazu (Ajdukiewicz [1956] (1985) mówi o definiowaniu wyrażań, nie ich znaczeń). Mimo że wiedza ta jest w pewnym sensie zależna od danych zmysłowych sądzimy, że podana definicja jest zdaniem analitycznym. Ów sens jest następujący: trzeba dysponować danymi zmysłowymi o rzekach, morzach, picciu, myciu itd., aby powyższa definicja była zrozumiała. Podobnie jest z innymi definicjami obiektów postrzegalnych zmysłowo. Np. Ajdukiewicz uznaje definicję zera stopni Celsjusza: „zero stopni Celsjusza jest to temperatura, w której woda pod normalnym ciśnieniem zamarza” za zdanie analityczne. Jednakże definicja ta będzie niezrozumiała bez pewnej wiedzy o danych zmysłowych, tutaj o tym, że woda zamarza (zmienia stan skupienia na stały). Ponadto, gdybyśmy w powyższej definicji sprawozdawczej „wody” zamienili to słowo na słowo „alkohol”, wówczas uznamy nowe zdanie za analitycznie fałszywe, tzn. do uznania go za fałszywe nie będzie nam służyć żadne świadectwo zmysłów, lecz tylko wiedza o sposobie rozumienia słowa „alkohol”. Co więcej, dla obcokrajowca nieznającego terminu „woda”, a rozumiejącego wyrażenia

języka polskiego występujące w *definiensie*, definicja ta będzie projektująca, nie sprawozdawcza. Jest niewątpliwe, że każda definicja projektująca (w relatywizacji do kogoś) jest zdaniem analitycznym. Dowolna definicja nominalna, według autora tego pojęcia (zob. Ajdukiewicz [1956] 1985), może być dla jednych sprawozdawcza, dla innych projektująca. Jednakże cechy *analityczności* i *syntezy* przysługujące zdaniom nie podlegają takiej relatywizacji. Oznacza to, że wszystkie definicje sprawozdawcze byłyby zdaniami analitycznymi.

Drugi argument przemawiający za tym, że prawa Reinacha (np. *zobowiązanie mieści się w istocie obiecywania*) mają rzeczywiście charakter *a priori*, ale z tego powodu, że są zdaniami analitycznymi, znajdujemy w pracy Wittgensteina ([1953] 2000), a przede wszystkim w książce Searle'a ([1969] 1987). Pierwszy z tych wybitnych autorów dość luźno i jakby w formie aforyzmu stwierdził, że używanie języka etnicznego to uczestnictwo w wielości różnych gier. Miał na myśli specjalny typ gier, które nie kończą się ani wygraną, ani przegraną, ani remisem. Zatem takie jak np. dziecięca gra „w chowanego”, które być może lepiej byłoby nazwać „zabawami” niż „grami”. Uczestnictwo w takich grach-zabawach polega na postępowaniu według pewnych konwencjonalnych reguł. Tak więc na tym polega posługiwanie się językiem.

Searle [1969] (1987) najwyraźniej zwrócił uwagę na aforyzm Wittgensteina (nie mamy pewności, czy rzeczywiście punktem wyjścia rozważań ontologicznych Searle'a był ów aforyzm) i szczegółowo go rozwinął, fundując w ten sposób bazę ontologiczną dla swojej teorii aktów mowy. Searle bada precyzyjnie status ontyczny *gier*. *Gra* jest tworem umysłu ludzkiego, powstałym dzięki ustaleniu jej *reguł konstytucyjnych* (ang. *constitutive rules*), mających zawsze konwencjonalny charakter. Reguły te określają pojęcia występujące w grze, jak i sposób uczestniczenia w niej, a więc sposoby postępowania w ramach grania (potocznie: reguły gry). Gra jest więc bytem *a priori*, a jej reguły konstytucyjne, definiując ją, są wyrażane w postaci zdań *analitycznych*. Naturalnie konsekwencje takich reguł są również zdaniami analitycznymi. Zdanie „W meczu otwarcia mistrzostw Europy w piłkę nożną w 2012 roku w wypadku strzelenia bramki gra będzie wstrzymana, a jej ponowne rozpoczęcie zacznie się zaraz potem dotknięciem piłki w środku boiska przez zawodnika drużyny, która bramkę straciła.” jest niewątpliwie zdaniem *a priori* — jest to bowiem zdanie analityczne. Ktoś mógłby temu zaprzeczyć, twierdząc, że zdanie to jest *syntetyczne* jako zdanie obserwacyjne, którego prawdziwość będzie oparta na świadectwie zmysłów. Ów ktoś konsekwentnie musiałby wówczas zaprzeczyć analityczności zdania Ajdukiewicza o temperaturze zero stopni Celsjusza, twierdząc, że zdanie to jest prawdziwe dzięki obserwacji (dobrze wyskalowanego) termometru, co jest wyraźnie absurdalne.

Searle zauważa, że ten sam status ontyczny gier mają wszelkie *instytucje* — nie tylko te będące strukturami czy organizmami społecznymi, lecz wszelkie

zorganizowane przez człowieka twory powołane do realizacji jego różnorodnych celów. Instytucja *małżeństwa* ma swoje reguły konstytucyjne. Określają one po pierwsze — pojęcia *małżeństwa*, *męża*, *żony* itd., po drugie — miejsca i wszelkie istotne fizyczne okoliczności procedury zawierania związku małżeńskiego, po trzecie — odpowiednie zachowania i czynności wykonywane przez osoby uczestniczące w zawieraniu małżeństwa (w tym najważniejszą z czynności — akt illokucyjny *ślubowania*), po czwarte — uprawnienia i obowiązki małżonków (nawet być może skodyfikowane w kodeksie prawa cywilnego) i inne zasady dotyczące trwania małżeństwa, po piąte — rozwiązanie małżeństwa (chyba że reguła mówi, iż jest ono wieczne i nierozwiązywalne). Mimo to, że taka instytucja ma realizować pewne praktyczne cele ludzkie, jej reguły konstytucyjne są arbitralne i dlatego mają charakter analityczny.

Obiecywać to działać w ramach pewnej instytucji. W efekcie wypowiedzenia obietnicy, dzięki regule konstytucyjnej tej instytucji, pojawia się zobowiązanie obiecującego oraz odpowiednie wymaganie adresata obietnicy. Zatem Searle uznałby prawo Reinacha za zdanie analityczne.

W opinii Searle’a mówienie językiem etnicznym, czyli wykonywanie czynności illokucyjnych, to postępowanie zgodne z regułami konstytucyjnymi różnych instytucji czy gier-zabaw (językowych) realizujących różne praktyczne cele. Reguły te mają charakter konwencjonalny — są wytworem człowieka, zaś zdania je opisujące, czyli „prawa”, jakim mówienie podlega, są analityczne.

6. Podsumowanie

Na użytek ostatniej części pracy będziemy abstrahować od różnic między teoriami Austina i Searle’a, uznając te teorie za jedną, co często jest czynione w literaturze przedmiotu.

Choć wyraźnie krytycznie ustosunkowaliśmy się do teorii aktów mowy A. Reinacha, zarówno co do charakterystyki czynności społecznych (w części trzeciej), jak i strony ontologicznej tej teorii (w części piątej), naszym celem była jej prezentacja jako, nazwijmy to, równoprawnej teorii Austina-Searle’a. Równoprawnej, to znaczy mającej podobną wartość heurystyczną. Taki pogląd jest współcześnie odosobniony. Sądzimy, że jedynym tego powodem jest nieznamość teorii Reinacha wśród językoznawców i filozofów języka, wynikająca, jak przypuszczamy, z nieadekwatności tytułu pracy Reinacha [1913] (1983). Tytuł ten zachęca prawników i filozofów prawa, nie językoznawców. Część lingwistyczna rozprawy Reinacha mogłaby być wydana osobno pod adekwatnym tytułem. Reinach tego nie zrobił, choć zarówno on sam, jak i inni fenomenologowie dostrzegali ważność wątków lingwistycznych w filozofii. W każdym razie zainteresowanie filozofii językiem etnicznym czy komunikacją językową przed I wojną światową było nieporównywalnie mniejsze niż w drugiej połowie XX wieku.

Istnieją jednak filozofowie doskonale orientujący się w problematyce teorii aktów mowy Reinacha i oceniający tę teorię, jak się wydaje, równie wysoko jak teorię Austina-Searle'a. Oto autorzy Crosby (1983) (tłumacz rozprawy Reinacha) oraz Smith (1990), porównując teorię czynności mowy Reinacha z teorią Austina i Searle'a, dostrzegają jakoby wiele istotnych podobieństw i wyraźnie je akcentują. W tym względzie nie dzielimy jednak ich opinii. Nie mamy tu miejsca na szeroką polemikę. Powiedzmy więc krótko: obie teorie cechuje *oczywiste* podobieństwo, istotne są tu jednak różnice, nie zaś to, co oczywiste. *Oczywistym* podobieństwem obydwu teorii jest specjalne rozumienie tego, na czym polega posługiwanie się językiem etnicznym. Polega ono na wykonywaniu czynności psychicznych, w których narzędziami są zdania. Drugim *oczywistym* podobieństwem owych teorii jest ten sam typ ontyczny obiektów w nich opisywanych, mianowicie obiektów czy struktur *a priori*: obiektów pozaprzestrzennych trwających w czasie, takich jak zobowiązanie (Reinach), czy gier lub instytucji, jak nazywa je Searle. Oczywiście tego podobieństwa wynika z samej natury posługiwania się językiem etnicznym: mówienie jest rzeczywiście działaniem w ramach struktur apriorycznych. Jest całkowicie jasne, że gdyby takich oczywistych podobieństw nie było, nikt nigdy nie przystąpiłby do porównywania obu teorii. Tylko wtedy jest sensowne porównywanie pewnych bytów, gdy istnieje między nimi jakieś fundamentalne, a więc oczywiste podobieństwo, tzn. istnieją pewne oczywiste aspekty, ze względu na które owe byty nie różnią się wcale. Czy ktoś będzie porównywał liczby (rzeczywiste) z rybami? Chyba jedynym aspektem, ze względu na który liczby są identyczne z rybami, jest rodzaj żeński rzeczowników będących ich nazwami.

Tymczasem istotne różnice między porównywanymi teoriami aktów mowy są dwie:

1) Jedyną cechą istotną (definicyjną) aktu społecznego w ujęciu Reinacha jest *konieczność* rozumienia towarzyszącej mu wypowiedzi przez jej adresata. Tymczasem rozumienie przez słuchacza zdania, którego wypowiedzenie służy wykonaniu aktu illokucyjnego, według Searle'a, jest cechą przypadłościową tego aktu. W ogóle bowiem wypowiedziane zdanie może nie być przez słuchającego zrozumiane, a tymczasem akt, któremu wypowiedzenie tego zdania towarzyszy, np. stwierdzenia, zostaje w zupełności niewadliwie wykonany. Istotną cechą czynności illokucyjnej jest odpowiednie ustosunkowanie do siebie czy ustawienie względem siebie wyrażanego w niej sądu i reprezentowanego przez ten sąd stanu rzeczy lub inaczej: zaistnienie nowego stanu rzeczy w świecie, polegającego na odpowiednim ustosunkowaniu tych obiektów. Reinach nic o tym nie wspomina.

2) Struktury aprioryczne umożliwiające dokonywanie aktów społecznych (czyli umożliwiające komunikację językową) nie są, zdaniem Reinacha, wytworem umysłu ludzkiego, są bytami naturalnymi, podobnie niezależnymi od umy-

słu ludzkiego jak pojęcia matematyczne. Różnią się jednak od tych ostatnich tym, że trwają w czasie, tzn. mają swój początek i koniec. Zdania opisujące związki między tymi strukturami są zdaniami syntetycznymi *a priori*. Tymczasem w opinii Searle’a owe struktury aprioryczne są konwencjonalnymi wytworami ludzkiego umysłu, tak jak wszelkie ludzkie gry. Zdania podające związki między tymi strukturami są analityczne.

Bibliografia

- AJDUKIEWICZ K. [1949] (2004), *Zagadnienia i kierunki filozofii*. Kęty-Warszawa: Wydawnictwo Antyk, Fundacja Aletheia.
- AJDUKIEWICZ K. [1956] (1985), *O definicji*, [w:] Ajdukiewicz K., *Język i poznanie*, t. II. Warszawa: PWN, 226–247.
- AUSTIN J. L. [1962], *How to Do Things with Words*. Oxford: Clarendon Press; (1993), *Jak działać słowami*, [w:] Austin J. L., *Mówienie i poznawanie*. Warszawa: PWN, 543–713.
- BRENTANO F. [1874], *Psychologie vom empirischen Standpunkte*. Leipzig: Dunckner & Humblot; (1999), *Psychologia z empirycznego punktu widzenia*. Warszawa: PWN.
- CROSBY J. F. (1983), *Adolf Reinach's Discovery of the Social Acts*, "Aletheia" 3, 143–194.
- FREGE G. [1918], *Der Gedanke. Eine logische Untersuchung*, "Beiträge zur Philosophie des deutschen Idealismus" 1, 58–77; (1977), *Myśl — studium logiczne*, [w:] Frege G., *Pisma semantyczne*. Warszawa: PWN, 101–129.
- HUME D. [1739], *Treatise of Human Nature*. London: J. Noon; (2006), *Traktat o naturze ludzkiej*. Warszawa: Wydawnictwo Aletheia.
- NOWAK M. (2000), *Złożone akty illokucyjne*, „Zeszyty Naukowe WSHE w Łodzi. Logika i Filozofia” 2, 35–75.
- NOWAK M. (2003), *Formalna reprezentacja pojęcia sądu (dla zastosowań w teorii aktów mowy)*. Łódź: Wydawnictwo UŁ.
- PAWŁOWSKI T. (1986), *Tworzenie pojęć w naukach humanistycznych*. Warszawa: PWN.
- REID T. [1785] (2002), *Essays on the Intellectual Powers of Man*. University Park: Pennsylvania State University Press; (1975), *Rozważania o władzach poznawczych człowieka*. Warszawa: PWN, wydanie niekompletne.
- REINACH A. [1913], *Die apriorischen Grundlagen des bürgerlichen Rechts*, "Jahrbuch für Philosophie und phänomenologische Forschung" 1/2, 685–847; (1983), *The Apriori Foundations of the Civil Law*, "Aletheia" 3, 1–142.
- SCHUMANN K. (1990), *Elements of Speech Act Theory in the work of Thomas Reid*, "History of Philosophy Quarterly" 7, 47–66.
- SCHUMANN K., SMITH B. (1987), *Adolf Reinach: An Intellectual Biography*, [w:] Mulligan K. (red.), *Speech act and Sachverhalt: Reinach and the Foundations of Realist Phenomenology*. Dordrecht-Boston-Lancaster: Nijhoff, 1–27.
- SEARLE J. R. [1969], *Speech Acts*. Cambridge-New York-Melbourne-Madrid: Cambridge University Press; (1987), *Czynności mowy*. Warszawa: PAX.

- SEARLE J. R. (1975), *A Taxonomy of Illocutionary Acts*, [w:] Gunderson K. (red.), *Language, Mind, and Knowledge*. Minneapolis: University of Minnesota Press, 344–369.
- SEARLE J. R., VANDERVEKEN D. (1985), *Foundations of illocutionary logic*. Cambridge: Cambridge University Press.
- SMITH B. (1990), *Towards a History of Speech Act Theory*, [w:] Burkhardt A. (red.) *Speech acts, Meanings and Intentions. Critical Approaches to the Philosophy of John R. Searle*. Berlin–New York: de Gruyter, 29–61.
- TWARDOWSKI K. [1894], *Zur Lehre vom Inhalt und Gegenstand der Vorstellungen. Eine Psychologische Untersuchung*. Wien: Hölder; (1965), *O treści i przedmiocie przedstawień*, [w:] Twardowski K., *Wybrane pisma filozoficzne*. Warszawa: PWN, 3–91.
- VANDERVEKEN D. (1990–91), *Meaning and Speech Acts*, vol. I, II. Cambridge–New York–Port Chester–Melbourne–Sydney: Cambridge University Press.
- WITTGENSTEIN L. [1953], *Philosophical Investigations*, Blackwell; (2000), *Dociekania filozoficzne*. Warszawa: PWN.
- YAFFE G. (2007), *Promises, Social Acts and Reid's First Argument for Moral Liberty*, "Journal of the History of Philosophy" 45(2), 267–289.

Rozdział 4

Dyrektywalna teoria znaczenia w interpretacji behawioralnej

Paweł Grabarczyk (Uniwersytet Łódzki)

Słowa kluczowe: Ajdukiewicz, Quine, Sellars, dyrektywa znaczeniowa, behaviorizm, znaczenie, sensowność

1. Wstęp

Nie ulega wątpliwości, że przedstawiona w latach trzydziestych ubiegłego wieku przez Kazimierza Ajdukiewicza Dyrektywalna Teoria Znaczenia¹ (w dalszej części artykułu posługiwał się będę często skrótem DTZ) nie spotkała się z entuzjastycznym przyjęciem filozofów języka. Poświęcone jej opracowania są prawie wyłącznie pracami o charakterze historycznym i można odnieść wrażenie, że jest ona bardziej jednym z zamkniętych rozdziałów w historii namysłu nad znaczeniem wyrażeń, niż teorią, którą warto by się było zajmować. Powodów takiego stanu rzeczy jest zapewne kilka. Po pierwsze, oryginalne sformułowanie dyrektywnej teorii znaczenia zniechęca zawartością i nieco archaicznym językiem². Po drugie, teoria ta wydaje się nazbyt restrykcyjna — opisując język w terminach składni i pragmatyki programowo pomija aspekt semantyczny, co spowodowane jest nieaktualną już dziś obawą przed paradoksami, do których mogłoby doprowadzić powołanie się na pojęcia takie jak „prawda” czy „odniesienie”. Po trzecie, w wyniku dostrzeżonych usterek została zarzucona przez samego autora i to dość szybko po jej opublikowaniu. Nie miała więc szans na to, żeby okrzepnąć i obrosnąć w zastrzeżenia i objaśnienia, które obroniłyby ją przed nasuwającymi się zarzutami. Mimo to wydaje mi się, że warto do DTZ wrócić i przyjrzeć się możliwościom, jakie daje. Jej

¹ Teoria ta przedstawiona została w: Ajdukiewicz (1931) i Ajdukiewicz (1933). Uwagi rozjaśniające niektóre kwestie znaleźć też można w: Ajdukiewicz (1953).

² Na przykład wspomniane dalej macierze języka konstruowane są w oparciu o mało dziś popularną notację Łukasiewicza.

podstawowa idea, czyli to, że znacznie wyrażen daje się wyprowadzić ze zbioru zinternalizowanych przez użytkowników reguł uznawania zdań, jest oryginalna i inspirująca. Moim celem w tym artykule nie jest jednak ani egzegza tekstu Kazimierza Ajdukiewicza, ani nawet racjonalna rekonstrukcja DTZ. Chodzi mi raczej o wskazanie drogi rozwoju dla tej teorii i naprawienie niektórych jej usterek, tak aby w przyszłości mogła okazać się przydatna dla współczesnej filozofii języka.

Z powodu braku miejsca nie będę szczegółowo przedstawiał DTZ w jej oryginalnym sformułowaniu. Czytelników zainteresowanych szczegółami tej wersji odsyłam do: Maciaszka (2007), Zmyślonego (2009) i Nowaczyka (2006). Jednakże, jako że moim zamiarem jest między innymi wprowadzenie do DTZ nowego typu dyrektyw, przedstawię teraz trzy oryginalne rodzaje dyrektyw wprowadzone przez Ajdukiewicza.

2. Pierwotna wersja DTZ

Dyrektywami znaczeniowymi nazywa Ajdukiewicz reguły języka, które oddane być mogą przez schemat o następującej postaci:

Jeżeli **u** zna język *L*, to jeśli **u** znajduje się w sytuacji *S*, **u** uzna zdanie *Z*.

Gdzie **u** oznacza użytkownika języka, *L* język, *S* pewną sytuację, a *Z* pewne konkretne zdanie języka *L*. Pasujące do tego schematu dyrektywy są zatem regułami, których pogwałcenie oznacza wyłączenie mówiącego z grona osób ten język znających. Inaczej mówiąc, przestrzeganie dyrektyw rozumieć należy jako warunek konieczny znajomości języka, do którego są one przypisane.

Różnice pomiędzy typami wprowadzonych dyrektyw przedstawić teraz możemy, różnicując sytuację *S* w powyższym schemacie. Jeżeli sytuacją *S* jest jakieś inne zdanie (lub zbiór zdań), powiedzmy *Z'*, a zatem, gdy dyrektywa nakazuje uznanie zdania *Z* w sytuacji, w której użytkownik języka uznał wcześniej zdanie *Z'*, to ten rodzaj dyrektyw nazwiemy *dyrektywami dedukcyjnymi*. Jeżeli motywem uznania zdania *Z* jest pewna dana empiryczna, to odpowiednią dyrektywę nazwiemy *dyrektywą empiryczną*, a jeśli motyw jest dla uznania zdania zupełnie nieistotny, a więc dyrektywa nakazuje je po prostu uznać w każdych możliwych okolicznościach, to mamy do czynienia z *dyrektywą aksjomatyczną*. Zilustrujmy to przykładami:

Dyrektywy aksjomatyczne nakazują uznanie pewnych zdań niezależnie od okoliczności, na przykład: *Paweł Grabarczyk jest identyczny z Pawłem Grabarczykiem*.

Dyrektywy dedukcyjne nakazują uznanie pewnego zdania w wyniku przyjęcia innego zdania. Przykładem niech będzie dowolne podstawienie reguły *modus ponens*, na przykład *Jeżeli liczba dyrektyw w DTZ jest nieparzysta, to nie*

jest podzielna przez 2, a liczba dyrektyw w DTZ jest nieparzysta; A zatem: liczba dyrektyw w DTZ jest niepodzielna przez 2.

Dyrektywy empiryczne nakazują uznanie pewnego zdania w obliczu określonej danej empirycznej. Słynnym przykładem, którym posługuje się Ajdukiewicz (1933: 156) jest reguła nakazująca uznanie zdania *Boli* w sytuacji, w której ktoś dotyka obnażonego nerwu zębowego użytkownika języka.

Warto zwrócić uwagę na to, że we wszystkich tych przypadkach mówi się o *uznaniu*, a nie *wypowiedzeniu* zdania. Nikt nie oczekuje zatem od użytkowników języka ciągłego emitowania zdań w podanych w dyrektywach sytuacjach. Testowanie dyrektywy znaczeniowej rozumieć należy w ten sposób, że użytkownik reaguje odpowiednio na pytanie zadane przez kogoś, kto testuje jego kompetencję językową. Nie należy też mylić *uznania*, o którym tu mowa, z prawdziwością zdań, co szczególnie w przypadku dyrektyw aksjomatycznych natychmiast się nasuwa. DTZ zachowuje swój asemantyczny charakter, odwołując się jedynie do reakcji użytkowników (którzy, teoretycznie, mogliby się systematycznie mylić). Warto również pamiętać o tym, że DTZ nie postuluje znajomości reguł przez użytkownika — jego zadaniem jest postępowanie zgodnie z dyrektywami, a nie na przykład zdolność do ich werbalizowania.

Ten pragmatystyczny punkt wyjścia pozwala Ajdukiewiczowi na podanie definicji znaczenia wyrażeń, która nie odwołuje się do żadnych pojęć semantycznych, dalsza część DTZ jest już bowiem konstrukcją czysto syntaktyczną. W uproszczeniu — jeżeli wyobrazimy sobie listę wszystkich dyrektyw danego języka, to synonimiczne będą w tym języku te wyrażenia, które można wzajemnie zastąpić we wszystkich dyrektywach w taki sposób, że zmieni się co najwyżej kolejność zdań na liście (ale nic nie przybędzie i nic nie wypadnie). Jeżeli teraz rozłożymy wszystkie znajdujące się na liście dyrektyw zdania na składowe, docierając w ten sposób do wyrażeń elementarnych języka, to, zachowując kolejność odpowiadającą procedurze rozkładania zdań, możemy utworzyć coś, co Ajdukiewicz nazywa *macierzami języka*. Znaczeniem danego wyrażenia będzie wtedy zbiór wszystkich tych miejsc, które zajmuje ono w macierzy swojego języka (zbiór miejsc zajmowanych przez wyrażenie jest zatem zrelatywizowany do macierzy, w której się znajduje, co odpowiada intuicyjnej idei zrelatywizowania znaczenia do konkretnego języka)³.

Tak pomyślana teoria ma kilka niezaprzeczalnych zalet. Po pierwsze, wychodzi ona od niezwykle intuicyjnych i rzadko kwestionowanych założeń: opanowanie języka oznacza nabycie zdolności do działania zgodnie z pewnymi regułami (nawet jeśli użytkownik nie jest w stanie reguł tych podać), a część

³ Raz jeszcze chciałbym podkreślić, że jest to omówienie niezwykle skrótowe. Wyczerpujące czy choćby nawet zadowalające przedstawienie tych kwestii wymagałoby oddzielnego artykułu. Jednocześnie znajomość szczegółów teorii dyrektywnej, takich jak sposób budowania macierzy, nie jest do zrozumienia niniejszego artykułu wymagana, ponieważ przedstawiane dalej modyfikacje DTZ nie dotyczą budowy macierzy.

z tych reguł odgrywa rolę kluczową, co przejawia się w tym, że działanie wbrew nim owocuje wykluczeniem działającego ze społeczności użytkowników języka. Po drugie, pozwala ona na redukcję semantycznego pojęcia *znaczenia* (a także pojęcia *przekładu*, o czym z powodu wprowadzonych uproszczeń nie piszę) do kombinacji pojęć syntaktycznych (zbiór miejsc wyrażenia w macierzy języka) i pragmatycznych (uznawanie zdań). Po trzecie, mimo upływu lat teoria ta nie posiada żadnego odpowiednika we współczesnej filozofii języka, choć część z idei w niej zawartych została później przez innych badaczy podjęta (oprócz dwóch teorii wspomnianych przeze mnie poniżej warto odnotować podobieństwo do teorii znaczenia Davidsona. Dokładne omówienie tej kwestii znajdzie czytelnik w: Maciaszek (2007) i Marsonet (1997)).

3. Kontrprzykład Tarskiego i języki zamknięte

Nie da się jednak ukryć, że mimo tych zalet DTZ posiada sporo mankamentów. Najbardziej znanym jest trudność wskazana Ajdukiewiczowi przez Alfreda Tarskiego. Tarski podał Ajdukiewiczowi oparty o bardzo prosty język, zawierający jedynie dyrektywy aksjomatyczne, kontrprzykład, który dowodził, że DTZ dopuszcza przypadek, w którym dwa wyrażenia o odmiennym odniesieniu mają na gruncie DTZ to samo znaczenie, zob. Ajdukiewicz (1964: 397). Przykład ten jest jednak z paru powodów wysoce kontrowersyjny i ostatecznie nie wydaje się bardzo poważnym zarzutem. Po pierwsze, nic nie zmusza nas do tego, by przyjmować zasadę determinacji odniesienia wyrażenia przez jego znaczenie. Dysponując (wprowadzonym wiele lat po sformułowaniu DTZ) rozróżnieniem na znaczenie wąskie i szerokie, możemy uznać, że DTZ dotyczy jedynie znaczenia wąskiego, o którym dobrze wiadomo, że zasady determinacji odniesienia przez znaczenie nie spełnia, zob. Putnam (1975). Po drugie (bronił się w ten sposób sam Ajdukiewicz), trudno byłoby sobie wyobrazić coś podobnego w języku, w którym funkcjonują dyrektywy empiryczne (a zatem we właściwym języku, dla którego stworzona została teoria dyrektywalna). Musiałby to być przypadek, w którym nie umiając podać żadnej różnicy między obiektami, bylibyśmy jednocześnie zobligowani przez reguły języka do zastrzegania, że mimo to są różne. Po trzecie, kontrprzykład Tarskiego (i podobne) daje się zablokować przez drobną modyfikację DTZ — w miejsce wzajemnego zastępowania wyrażen w dyrektywach wystarczy posłużyć się zwykłym zastępowaniem (bez wymogu wzajemności), zob. Buszkowski (2010) i zmodyfikować definicję synonimiczności następująco:

Wyrażenia A i B są synonimiczne na gruncie języka L wtw, gdy zastąpienie jednego drugim w każdej dyrektywie da w efekcie istniejącą dyrektywę języka tego samego typu (zob. też Nowaczyk 2006).

Inną znaną trudnością DTZ jest to, że sformułowana ona została dla tzw. języków zamkniętych. Językami zamkniętymi nazywa Ajdukiewicz takie języki, w przypadku których każde nowo wprowadzone wyrażenie ma już w nich synonim. Jest to, oczywiście, idealizacja, ale jak się okazuje, dość kłopotliwa. Problemem jest nie tylko to, że języków takich nie ma, okazuje się, że nie da się ich nawet zbudować w sposób sztuczny, zob. Buszkowski (2010). Cena za rezygnację z tej idealizacji jest jednak dość wysoka. Ponieważ znaczenie wyrażen zdefiniowane zostało jako zbiór miejsc, które wyrażenie zajmuje w macierzy języka, a macierz ta zmienia się za każdym razem, gdy tylko wprowadzimy do języka nowe wyrażenie⁴, to znaczenie wszystkich wyrażen w języku zmienia się w momencie, w którym wprowadzimy do niego nowe wyrażenie. Problemowi temu można jednak starać się zaradzić na różne sposoby, łącznie z pogodzeniem się z tą nieintuicyjną konsekwencją i uznaniem, że teoria znaczenia zrelatywizowana jest zawsze do pewnego konkretnego etapu rozwoju języka⁵. Na dodatek nie należy tracić z oczu tego, że jest to jedynie nieintuicyjna konsekwencja, która zawsze może okazać się ceną wartą zapłacenia za inne zalety teorii.

Poza tymi dwoma problemami DTZ rodzi jeszcze sporo drobniejszych, ale nie mniej dotkliwych trudności. Wymieńmy niektóre z nich. W odróżnieniu od macierzy języka, których budowa opisana jest przez Ajdukiewicza bardzo dokładnie, o samych dyrektywach, ich rodzajach i charakterystyce dowiadujemy się zdecydowanie zbyt mało. Czy podana lista rodzajów dyrektyw jest listą wyczerpującą? Być może teoria języka naturalnego powinna uwzględniać więcej typów reguł?

Nie jest również jasne, kiedy możemy w ogóle powiedzieć o danym wyrażeniu, że posiada znaczenie — czy wystarczy, że występuje w jakiejś dyrektywie? Trudno jest też powiedzieć, jak mamy rozumieć zastrzeżenie, że dyrektywy podane są dla wszystkich wyrażen języka. Jest to spodziewane w przypadku predykatów, ale co z innymi kategoriami wyrażen — wyrażeniami złożonymi, zdaniami, nazwami własnymi? Czy mamy przyjąć, że dla każdego zdania istnieje dyrektywa, która reguluje jego użycie?

Na dodatek kluczowa, jeżeli tylko chcemy uczynić z DTZ narzędzie do badania języków naturalnych, kategoria dyrektyw empirycznych prowokuje wiele pytań. Po pierwsze — czy mamy uznać, że wyrażenia znajdujące się w tych dyrektywach odnoszą się do danych empirycznych? Czym właściwie są te dane empiryczne? Dowiadujemy się jedynie tego, że mają swoje nazwy w metajęzyku, w którym opisane są macierze języka, ale do czego się te nazwy w metajęzyku

⁴ Ponieważ musi ono zostać wprowadzone za pomocą dyrektyw, które zostaną rozłożone na wyrażenia elementarne i dopisane do macierzy.

⁵ Z powodu braku miejsca nie przedstawiam szczegółów tych rozwiązań. Czytelnik znajdzie je w: Maciaszek (2007) i Jedynak (2003).

odnoszą? Do sytuacji, przedmiotów, zdarzeń? Może do reprezentacji w umyśle użytkowników języka?

Poszukiwanie odpowiedzi na te pytania rozpocznę od przywołania dwóch teorii, które stanowiły dla mnie najważniejszą inspirację do wprowadzenia większości z przedstawionych poniżej modyfikacji DTZ.

4. Quine i problem dyrektyw empirycznych

Pierwszą z tych teorii jest behawioralna teoria znaczenia Quine'a⁶. Zbieżności pomiędzy nią a DTZ nietrudno zauważyć. Przedstawiając swój słynny eksperyment myślowy z przekładem radykalnym, Quine stwierdza, że badający język tubylców lingwista radykalny kolekcjonuje dane, którymi jest zapis potwierdzanych przez obserwowanych użytkowników języka zdań. Co więcej, lingwista ten grupuje swoje znaleziska, tworząc trzy kategorie zdań. Są to, po pierwsze — zdania, które potwierdzane są przez tubylców niezależnie od okoliczności, po drugie — zdania, które tubylcy potwierdzają zawsze wtedy, gdy potwierdzili już wcześniej inne zdania i po trzecie — zdania, które potwierdzają w obecności pewnego bodźca empirycznego. Jak widać, jest to struktura niezwykle podobna do podziału na dyrektywy aksjomatyczne, dedukcyjne i empiryczne. Między propozycjami Ajdukiewicza i Quine'a zachodzi jednak pewna poważna różnica. Zamiast o sytuacjach czy przeżyciach Quine mówi o bodźcach. Na dodatek są to bodźce rozumiane w sposób dość niestandardowy — bodźcem nie jest zdarzenie, które wywołuje w nas jakąś reakcję, ale samo pobudzenie receptorów w organizmie.

Proponuję zatrzymać się przy tej różnicy, ponieważ wiąże się ona z jedną z zasygnalizowanych niejasności teorii dyrektywalnej.

Zastanówmy się, co właściwie jest składnikiem dyrektyw empirycznych? Jak wspomniałem, mówi się tam o pewnej danej empirycznej czy wrażeniowej. Czym właściwie ów empiryczny składnik jest? Zaczniemy od jego kategorii ontologicznej. Na pierwszy rzut oka wydaje się, że najlepiej do tej roli nadają się zdarzenia, ponieważ same przedmioty, jeśli tylko nie oddziałują na użytkownika języka (na przykład nie są w ogóle obecne w jego polu percepcji), nie wymuszają na nas uznawania żadnych zdań. Ale jaka jest właściwie rozpiętość czasowa tych zdarzeń? Metaforycznie, ale, jak mi się wydaje, dość trafnie, moglibyśmy tę główną trudność związaną z odniesieniem empirycznego składnika dyrektyw oddać, mówiąc, że chodzi o to, jak drobno mamy to doświadczenie podzielić. Na jednym krańcu skali znajduje się tu pojedyncze, elementarne zdarzenie a na drugim niezwykle skomplikowana sytuacja, złożona ze wszystkich zdarzeń, które są w danej chwili w polu percepcji użytkownika języka. Weźmy słynny przykład Quine'a — założmy, że lingwista badający język tubylców

⁶ Choć propozycja Quine'a była przez niego modyfikowana, jej zasadnicze idee, do których odwołuję się w tym artykule, znalazły się już w: Quine (1960).

ustalił, że w obecności danej empirycznej, jaką jest pojawienie się królika, reguły rekonstruowanego języka nakazują uznanie zdania *Gavagai*. Co mielibyśmy tu wpisać do macierzy w charakterze danej empirycznej — zdarzenie, jakim jest zjawienie się królika w polu widzenia, ciąg zdarzeń, jakim jest zjawienie się królika w polu widzenia i jego pozostawanie w tym polu, czy może coś większego — jakiś wycinek całej sytuacji, jaką jest „natknięcie się na królika w dżungli”? Pojedyncze zdarzenia wydają się najczęściej zbyt drobne, wygląda więc na to, że powinniśmy brać pod uwagę raczej ich sekwencje albo też zbiory zdarzeń (czyli sytuacje). Co gorsza, nawet jeżeli na coś się już zdecydujemy, to jest to tylko początek kłopotów, ponieważ w odróżnieniu od przedmiotów, które, jak zauważyliśmy, na składowe dyrektyw empirycznych raczej się nie nadają, zdarzenia, ciągi zdarzeń i sytuacje mają o wiele mniej zrozumiałe warunki tożsamości. Czy to samo zdarzenie albo ta sama sytuacja może pojawić się wielokrotnie? Wydaje się, że raczej nie, zamiast zdarzeń i sytuacji musimy zatem w dyrektywach umieścić raczej nazwy typów zdarzeń, względnie nazwy typów sytuacji. Nie bardzo jednak wiadomo, w jaki sposób mielibyśmy te typy w metajęzyku wyróżniać.

Skoro zdarzenia i sytuacje przysparzają nam tylu trudności, to może lepiej byłoby skupić się na przeżyciach psychicznych⁷? Być może dyrektywy empiryczne powinny po prostu nakazywać uznanie danego zdania w obliczu tego, a nie innego przeżycia? Postawiwszy sprawę w ten sposób od razu widzimy jednak, że grożą nam tutaj te same trudności, co przed chwilą — trudno jest właściwie powiedzieć, jakie są warunki tożsamości przeżyć — czy można mieć na przykład kilka razy to samo wyobrażenie (w odróżnieniu od takiego samego wyobrażenia albo wyobrażenia tego samego przedmiotu)?

Dlatego też sądzę, że optymalnym rozwiązaniem będzie dla nas przyjęcie bodźców w sensie Quine’a jako składników dyrektyw empirycznych. Pytania o ich warunki tożsamości nie wprowadzają nas w zakłopotanie — zawsze potrafimy powiedzieć, czy pobudzony był ten sam receptor (albo zbiór receptorów) co ostatnio. Co więcej, ponieważ nie jest pewne, jak dużą część receptorów mamy wziąć w danej chwili pod uwagę, zawsze możemy po prostu uwzględnić wszystkie z nich, czyli tak zwany „bodziec całościowy”, nie ulega bowiem wątpliwości, że jest on o wiele lepiej poznany i łatwiejszym do opisanego obiektem niż „sytuacja całościowa”.

Skoro jesteśmy przy rozwiązywaniu trudności związanych z dyrektywami empirycznymi, to warto wspomnieć o innym kłopotcie, który komplikuje mocno obraz przedstawiany pierwotnie przez DTZ. Użytkownik języka może nie potwierdzić wymaganego zdania, ponieważ jest, dajmy na to, przekonany o tym, że zmysły go w danej chwili zawiodą. Na przykład, będąc przekonany,

⁷ Rozwiązanie to preferował Ajdukiewicz.

że właśnie zakropiono mi specjalne krople, które zmieniają sposób, w jaki postrzegam kolory albo że znajduję się w pomieszczeniu oświetlonym w jakiś bardzo nietypowy sposób, mogę nie uznawać zdania „to jest czerwone” w obecności czerwonego bodźca. Zauważmy, że nie możemy tego problemu rozwiązać, odwołując się do stanu faktycznego (czy to pomieszczenia, czy to moich oczu), ponieważ nie potwierdzę tego zdania również wtedy, gdy fałszywie będę przekonany o tym, że okoliczności są nietypowe. Ajdukiewicz dostrzegał ten problem, ale proponowane przez niego rozwiązanie można, w najlepszym razie, uznać za prowizoryczne. Wprowadził on jedynie dodatkowe zastrzeżenie, że użytkownik musi być przekonany o tym, że „sytuacja jest normalna”, zob. Ajdukiewicz (1934: 156). To prowizoryczne rozwiązanie prowokuje zbyt wiele pytań: do czego właściwie przekonanie to się sprowadza? Czy jest ono jedynie skrótem koniunkcji konkretnych przekonań, które przypisujemy użytkownikowi? Jeżeli tak, to czy potrafimy skróć ten rozwinąć i dopisać do dyrektyw?

Wyjściem z tej sytuacji jest przyjęcie, że przekonanie o normalności sytuacji sprowadza się do jakiegoś, zapewne dość licznego, zbioru przekonań niezwerbalizowanych. Należałoby to rozumieć tak, że osoba, która interpretuje jakieś doświadczenia, nie myśli w danym momencie o wszystkich szczegółach, które upewniają ją o normalności sytuacji, ale są one jej dostępne w postaci sieci przekonań, które ma, ale których niepytana, nie przywołuje. W jaki jednak sposób mielibyśmy taki niezwerbalizowany zbiór przekonań reprezentować w dyrektywach?

To jeszcze nie koniec trudności. Aby dyrektywy działały tak jak powinny testowany użytkownik języka musi żywić przekonania nie tylko na temat stanu swojego organizmu i otoczenia, ale i intencji rozmówcy. Jeżeli uzna, że padł właśnie ofiarą żartu lub prowokacji, może nie potwierdzić zdań wymienianych w dyrektywie, ale nie będzie to przecież oznaczało nieznamośności języka. Co więcej, choć zastrzeżenia te nasunęły się nam przy okazji omawiania dyrektyw empirycznych, to dotyczą one również dyrektyw aksjomatycznych i dedukcyjnych. Zdanie, którego uznanie nakazuje jakaś dyrektywa aksjomatyczna, mógłbym odrzucić dlatego, że jestem przekonany o tym, że ktoś ze mnie kpi (pytając o coś tak oczywistego) albo gdybym uznał, że mój rozmówca porozumiewa się jakimś szyfrem, który zmienia znaczenie słów, albo też gdybym był na przykład przekonany, że podano mi środek, który sprawia, że wszystko wydaje mi się pewne i oczywiste. Zastrzeżenia Ajdukiewicza w postaci wspomnianego już założenia o normalności sytuacji czy tego, że użytkownik języka działać musi na serio i w dobrej wierze (zob. Ajdukiewicz 1934: 150–151), są tak ogólne, że trudno je uznać za pełnoprawny składnik teorii, która w pozostałych aspektach jest bardzo drobiazgowa i jednoznaczna.

Sądzę, że odwołując się do proponowanej przeze mnie behawioralnej interpretacji dyrektyw, możemy te trudności przezwyciężyć. Wystarczy tylko

przyjąć, że częścią kompetencji językowej, którą testuje dyrektywa, jest zdolność rozpoznania „normalnego” oświetlenia czy stanu naszego organizmu. Nie jest to jednakże jakaś dodatkowa analiza danych zmysłowych, wymagająca zaangażowania przekonań czy rozumowań — całościowy bodziec, a więc w przypadku naszego przykładu pobudzenie wywołane kolorową plamą w tym, a nie innym oświetleniu i pytaniem o odpowiednim brzmieniu ma wywoływać reakcję potwierdzenia zdania.

Zauważmy teraz, że integralnym składnikiem bodźca całościowego jest samo zdanie, które ma zostać potwierdzone⁸. Chodzi tu o coś tak oczywistego, że mogłoby umknąć naszej uwadze — gdy na przykład mówimy, że dyrektywa aksjomatyczna nakazuje uznanie zdania niezależnie od bodźca, to nie możemy zapominać o tym, że użytkownik słyszy jednak pewne zdanie, które ma potwierdzić — to jest, rzecz jasna, rodzaj bodźca, który do niego dociera i nie można go zamienić na nic innego. A zatem również w dyrektywach aksjomatycznych istnieje specyficzne dla nich pobudzenie receptorów. Rozpoznanie, że jest to zdanie, które ma sprawdzić, czy kierujemy się dyrektywą, a więc zorientowanie się, że sprawdza się właśnie naszą podstawową kompetencję językową (a nie na przykład z nas żartuje), jest częścią tej kompetencji językowej (tak samo jak rozpoznanie, że coś jest ironią lub groźbą). Jeśli będziemy o tym pamiętać, to, na przekór Ajdukiewiczowi, moglibyśmy powiedzieć, że użytkownik powinien rozpoznawać występujące w dyrektywach bodźce całościowe jako wysoce nietypowe⁹ — w typowych sytuacjach nikt nie prosi go o potwierdzanie zdań tożsamościowych albo potwierdzanie zdania „to jest czerwone” w obliczu czerwonego bodźca. Robi to tylko w dość specyficznych okolicznościach — gdy chce rozwiązać wątpliwość co do zdolności komunikacyjnych rozmówcy. Zdania występujące w dyrektywach (albo pary zdań z bodźcami) odznaczają się tym, że nauczeni jesteśmy je rozpoznawać i reagować na nie odruchowo, bez analizy czy zaangażowania jakichkolwiek przekonań.

5. Sellars i dyrektywy imperatywne

Drugą teorią, którą chciałbym w charakterze inspiracji przywołać, jest teoria Sellarsa wyłożona w Sellars (1954). Sposób, w jaki rozumie on w tym artykule język, niezwykle przypomina podejście Ajdukiewicza. Punkt wyjścia jest właściwie identyczny — Sellars zwraca uwagę na to, że język można rozumieć jako zestaw reguł podobnych do gry czy dyscypliny sportowej¹⁰. Nawet jeśli część reguł można naginać, to istnieje taki ich podzbiór, który stanowi warunek

⁸ Zwraca na to uwagę również Maciaszek, zob. Maciaszek (2007: 200).

⁹ Uwaga ta nie stoi w sprzeczności z tym, co powiedziałem parę zdań wcześniej — częścią tego nietypowego bodźca całościowego może być „typowe oświetlenie”. Nietypowe jest właśnie zadanie zbyt oczywistego pytania przy typowym oświetleniu.

¹⁰ Sellars posługuje się głównie przykładami gry w szachy, a sama idea zaczerpnięta jest, rzecz jasna, od Wittgensteina.

konieczny grania w daną grę. Znaczeniem wyrażenia (które, w zależności od stopnia złożenia, może być analogonem posunięcia w grze albo pionka) jest rola, jaką odgrywa w grze, rolę, którą wyczerpuje zestaw reguł, które go dotyczą. Sellars korzysta z behawioralnego języka, co stanowi dla nas znaczne ułatwienie, jako że zdecydowaliśmy się na bodźce w roli danych empirycznych¹¹. Podążając za metaforą gry, proponuje on wprowadzenie czterech kategorii reguł językowych:

Reguły wolne — są to reguły, które nakazują potwierdzanie danego zdania niezależnie od bodźców.

Reguły wejścia¹² — chodzi tu o reguły, które informują nas o tym, w jakich pozajęzykowych okolicznościach należy uznać dane zdanie.

Reguły pomocnicze — które instruują nas, jakie zdanie należy uznać, jeżeli wcześniej uznało się jakieś inne zdanie.

Reguły wyjścia — reguły, które w wyniku uznania jakiegoś zdania wymagają od użytkownika pewnego zachowania.

Nie potrzeba chyba podkreślać tego, że mamy tu do czynienia ze znajomą konstrukcją. Pierwsze trzy rodzaje reguł odpowiadają kolejno: dyrektywom aksjomatycznym, empirycznym i dedukcyjnym. Ostatni typ reguł nie ma w koncepcji Ajdukiewicza odpowiednika, dlatego też warto rozważyć wprowadzenie do DTZ jakiejś jego wersji. Wydaje się, że brak tego rodzaju dyrektyw, nazwijmy je dyrektywami imperatywnymi, jest swego rodzaju niedopatrzeniem. Nie ulega bowiem wątpliwości, że istnieją w języku reguły, które nakazują wykonanie jakiejś czynności w wyniku usłyszenia i uznania danej komendy. Człowiek, który akceptuje zdanie „Stop!”, ale nie zatrzymuje się nawet na chwilę może oczekiwać, że jego rozmówcy uznają, że nie zna znaczenia tego słowa. Zauważmy, że korzystamy w tym miejscu z asemantyczności DTZ. Fakt, iż zdania rozkazujące nie mają wartości logicznej nie oznacza, że użytkownicy nie mogą ich akceptować czy uznawać — byłoby tak tylko wtedy, gdybyśmy przez potwierdzenie czy uznanie rozumieli przyjęcie czy uznanie za prawdę, a tak nie jest. „Uznanie” czy „potwierdzenie”, na które się w DTZ powołujemy, rozumieć można jako niesemantyczny substytut pojęcia „rozumienia”. W typowych przypadkach różnica pomiędzy uznaniem zdania a jego zrozumieniem polega choćby na tym, że istnieją zdania, które rozumiem, ale ich nie uznaję (ponieważ, na przykład, nie zgadzam się z nimi). Różnica ta z oczywistych powodów znika

¹¹ Choć wydaje się, że Sellars mówi raczej o bodźcach w znaczeniu „zewnętrznej przyczyny pobudzenia” niż samego pobudzenia.

¹² Sformułowanie to może nasuwać informatyczne skojarzenia, ale jest to jedynie przypadek (choć samo to powiązanie nie jest całkowicie nietrafne). W angielskim oryginale występuje termin „entry”, a nie „input”, a wejście rozumie się tu w sensie „wejścia do gry językowej”.

w przypadku zdań, z którymi nie można się nie zgodzić, a właśnie takie zdania występują w dyrektywach.

Jak zatem mogłaby wyglądać taka dyrektywa nakazująca działanie w przypadku uznania zdania „Stop!”? Wzorując się na Sellarsie, najwygodniej będzie ją chyba uznać za przeciwieństwo dyrektywy empirycznej. Jedyną trudnością może być to, że w odróżnieniu od dyrektyw empirycznych nie możemy odnieść się tu do stanu receptorów — nie na tym wykonanie rozkazu polega, że odbieramy ten, a nie inny bodziec¹³. Sądzę, że chcąc pozostać w zgodzie z przyjętym modelem behawioralnym, jako drugi człon dyrektywy najwygodniej będzie nam przyjąć na przykład reakcję motoryczną¹⁴. Ponieważ podciągnięcie wszystkich rozkazów pod taki schemat dyrektyw byłoby nadmiernym uproszczeniem, wygodnie będzie również wprowadzić nieco bardziej złożoną wersję dyrektyw imperatywnych — takich, w których podjęte działanie motywowane jest nie tylko uznanym zdaniem, ale i równoległe pewną daną doświadczenia. Przykładem takiej dyrektywy byłaby komenda „łap!” połączona z rzuceniem piłki.

6. Sensowność wyrażen

Przejdźmy teraz do problemu sensowności wyrażen. Jak wspomniałem, DTZ podaje definicję synonimiczności, dzięki czemu, przez abstrakcję, jesteśmy w następnej kolejności w stanie podać właściwą definicję znaczenia. Kłopot w tym, że teoria ta nie daje nam żadnych wskazówek dotyczących tego, co sprawia, że dane wyrażenie w ogóle posiada jakiekolwiek znaczenie. Inaczej mówiąc — jak odróżnić można na jej gruncie wyrażenia nonsensowne od sensownych? Choć Ajdukiewicz nie zajął się w ogóle tym problemem, to sądzą, że rozszerzenie DTZ o to użyteczne rozróżnienie jest możliwe. Wystarczy, że wyciągniemy ostateczne wnioski z idei, która za DTZ stoi, a mianowicie, że znaczeniem wyrażenia jest sieć powiązań z pozostałymi wyrażeniami języka. Jeśli tak jest, to za brak znaczenia uznać możemy po prostu kompletny brak powiązań danego wyrażenia z innymi wyrażeniami. Na pierwszy rzut oka wydaje się, że jest to kryterium nieużyteczne, ponieważ takich wyrażen w języku po prostu nie ma, a to dlatego, że niektóre z dyrektyw aksjomatycznych dotyczą wszystkich wyrażen w języku. Na przykład zasada tożsamości każe uznać wszystkie zdania o postaci $x=x$, a zatem jeżeli naszym kandydatem na nonsens będzie na przykład, by posłużyć się klasycznym przykładem Wittgensteina, wyrażenie „tralalala”, to można by zaprotestować, że przecież występuje ono w powiązaniu ze znakiem identyczności w regule aksjomatycznej „tralalala = tralalala”. Możemy temu zaradzić, wykorzystując dodatkowe rozróżnienie, które

¹³ Chyba że rozkazem tym jest na przykład „Patrz!”, ale tutaj też wykonaniem jest bardziej ruch głowy czy powiek niż samo odebranie bodźca.

¹⁴ Ostateczna wersja schematów dyrektyw podana jest na końcu artykułu.

Ajdukiewicz stosował — a mianowicie podział na występowanie wyrażenia w dyrektywie w sposób istotny i nieistotny:

Wyrażenie występuje w dyrektywie D w sposób istotny wtw, gdy nie jest tak, że może być zastąpione dowolnym innym wyrażeniem, a powstała w ten sposób dyrektywa D' będzie nadal dyrektywą języka.

Zgodnie z tą definicją wyrażenie „tralalala” występuje we wspomnianej dyrektywie aksjomatycznej w sposób nieistotny. Do podania definicji sensowności będą nam potrzebne jeszcze dwa inne pojęcia wprowadzone przez Ajdukiewicza:

Wyrażenie A pozostaje w bezpośrednim związku znaczeniowym z wyrażeniem B wtw, gdy oba występują w obrębie tej samej dyrektywy.

Wyrażenie A pozostaje z wyrażeniem B w pośrednim związku znaczeniowym wtw, gdy istnieje co najmniej trójelementowy ciąg wyrażen, w którym A jest wyrażeniem pierwszym, a B ostatnim i w którym wszystkie sąsiadujące wyrażenia są ze sobą powiązane bezpośrednio.

Możemy teraz, za Ajdukiewiczem, wprowadzić jeszcze jedno pojęcie — pojęcie języka spójnego:

Język L jest spójny wtw, gdy wszystkie jego wyrażenia są ze sobą powiązane pośrednio lub bezpośrednio.

Mając te pojęcia pomocnicze, możemy teraz pozwolić sobie na podanie ogólnej definicji wyrażenia posiadającego znaczenie, której w teorii Ajdukiewicza brakowało:

Wyrażenie posiada znaczenie wtw, gdy jest pośrednio lub bezpośrednio powiązane ze wszystkimi pozostałymi wyrażeniami spójnego języka za pomocą dyrektyw, w których występuje w sposób istotny.

Korzystając z przedstawionych rozróżnień pojęciowych, możemy dodatkowo wprowadzić węższe pojęcie sensowności, bliższe tradycji empirystycznej:

Wyrażenie ma znaczenie empiryczne wtw, gdy ma znaczenie, a ponadto występuje w sposób istotny w przynajmniej jednej dyrektywie empirycznej¹⁵.

Tak zdefiniowane „znaczenie empiryczne” może być rozumiane jako odpowiednik „znaczenia bodźcowego” Quine’a. Wolę jednak traktować „znaczenie empiryczne” jako osobne pojęcie (a nie na przykład eksplikację pojęcia znaczenia bodźcowego), ponieważ tak zdefiniowane na gruncie DTZ pojęcie jest

¹⁵ Co ciekawe, Ajdukiewicz posługuje się w artykule *Język i znaczenie* tym określeniem, ale w ogóle go nie precyzuje, zob. Ajdukiewicz (1934: 156–157).

niezależne od naszej decyzji o uznaniu empirycznego składnika dyrektyw za bodźce całościowe. Aby dopełnić obraz, możemy teraz wzbogacić DTZ o relację posiadania tego samego znaczenia empirycznego. Jak nietrudno się domyślić, będzie ono analogiczne do zwykłej synonimiczności:

Dwa terminy mają to samo znaczenie empiryczne wtw, gdy zamiana jednego na drugi w dowolnej dyrektywie empirycznej nie zmieni charakteru tej dyrektywy¹⁶.

Różnica pomiędzy tak zdefiniowaną tożsamością znaczenia empirycznego a zwykłą synonimicznością polega na tym, że nie każde dwa wyrażenia o tym samym znaczeniu empirycznym są synonimiczne, choć każde dwa synonimiczne wyrażenia mają to samo znaczenie empiryczne. Na przykład wyrażenia „bydło” i „wielość sztuk bydła” nie są synonimiczne, ale mają to samo znaczenie empiryczne¹⁷.

7. Znaczenie zdań i nazw własnych

Następnym problemem DTZ, którym chciałbym się teraz zająć, jest pytanie o znaczenie w przypadku różnych kategorii wyrażeń — czy DTZ równie dobrze radzi sobie z podaniem znaczenia wyrażeń dowolnego typu (predykatów, nazw własnych, zdań)? Niekontrowersyjnym przypadkiem są z pewnością predykaty — język zawiera ich skończoną ilość, dają się one intuicyjnie powiązać z danymi empirycznymi czy aksjomatami (które mogłyby wysławiać jakieś analityczne zależności pomiędzy nimi). A co z nazwami własnymi? Ajdukiewicz nie podjął tego zagadnienia, a odpowiedź na to pytanie nie wynika jednoznacznie z założeń teorii dyrektywnej. Można zatem powiedzieć, że mamy w tej kwestii wolną rękę. Sądzę, że najwygodniejszym rozwiązaniem będzie przyjęcie, że dyrektywy nie regulują ich znaczenia. Argumentów za takim rozwiązaniem jest co najmniej kilka. Po pierwsze, to, czy nazwy własne posiadają jakiegokolwiek znaczenie, jest kwestią sporną i wielu filozofom, w tym mnie, bliżej jest dziś do koncepcji, które im tego znaczenia odmawiają (nie trzeba być przy tym zwolennikiem przyczynowej teorii odniesienia, choć to, oczywiście, pomaga). Po drugie, kwestią równie sporną jest decyzja co do tego, czy nazwy własne są w ogóle częścią jakiegoś konkretnego języka — zauważmy choćby, że powiedzenie, iż używana na gruncie języka angielskiego nazwa „Ajdukiewicz” jest przekładem polskiej nazwy „Ajdukiewicz” pozostaje niezgodne ze zwykłym użyciem języka. Nawet w przypadkach, w których nie są to ciągi równokształtne, mówimy często o transliteracji, nie o przekładzie. Po trzecie

¹⁶ Określenie „charakter dyrektywy” może wydawać się nieco nieprecyzyjne, ale chodzi tu po prostu o to, że zamiana wyrażenia na inne nie sprawia, że dana dyrektywa przestaje być dyrektywą języka lub też staje się dyrektywą innego typu.

¹⁷ Oddzielnym ciekawym zagadnieniem jest to, czy na gruncie DTZ można w ogóle mówić o synonimiczności wyrażeń różniących się od siebie składniowo.

— zauważmy, że w odróżnieniu od predykatów trudno byłoby podać przykład dyrektywy, w której występuje (w sposób istotny) nazwa własna. Weźmy klasyczny, pochodzący od Fregego przykład — czy zdanie w rodzaju *Arystoteles jest nauczycielem Aleksandra* rzeczywiście chcielibyśmy podnieść do rangi reguły języka? Wydaje się, że naturalniej byłoby jednak powiedzieć, że ktoś, kto odrzuca to zdanie, ma po prostu błędne przekonania na temat Arystotelesa¹⁸.

Przejdźmy teraz do pytania o znaczenie zdań. Kwestii tej Ajdukiewicz również nie podjął, ale odpowiedź na nie nietrudno sformułować — DTZ nie daje żadnych wskazówek co do tego, w jaki sposób użytkownik interpretuje znaczenie większości zdań. W dyrektywach znajduje się jedynie pewien podzbiór zdań danego języka, nie możemy zatem powoływać się na „miejsce, które zdanie zajmuje w macierzy języka”. Większości zdań tam po prostu nie ma. Aby rozwiązać ten problem, należy przede wszystkim uświadomić sobie, że kryje się za nim założenie, że teoria znaczenia, która nie podaje sposobu, w jaki znaczenie zdań powstaje ze znaczeń ich składowych, nie spełnia jednej (być może nawet najistotniejszej) ze swoich ról, zob. Maciaszek (2007: 325). Założenie to nie jest jednak samo w sobie oczywiste. Wydaje się, że Ajdukiewicz reprezentował tradycyjny punkt widzenia, zgodnie z którym najważniejszą zagadką do wyjaśnienia jest znaczenie wyrażen elementarnych, bo w objaśnieniu, w jaki sposób ze znaczeń tych powstają znaczenia zdań, pomóc nam mogą inne teorie, takie jak choćby logika czy gramatyka (szczególnie w którejś z nietradycyjnych odmian, na przykład gramatyka kategorialna, której podwaliny Ajdukiewicz w końcu stworzył).

Mam wrażenie, że pewne drobne modyfikacje DTZ pozwolą nam jednak coś na temat znaczenia zdań powiedzieć. Na początek rozgraniczmy trzy poziomy wyrażen: zdania złożone, zdania elementarne i wyrażenia elementarne (rozumiane jako składniki zdań, które nie są zdaniami i nie dają się już rozłożyć na mniejsze elementy należące do słownika języka). Jak już wiemy DTZ podaje nam znaczenie tych ostatnich — jest nim zbiór miejsc zajmowanych przez te wyrażenia w macierzy języka. Z wcześniejszych ustaleń wiemy również, że wyrażeniami tymi nie są nazwy własne. Będą to zatem przede wszystkim predykaty zero- i więcej argumentowe. Zacznijmy od predykatów zeroargumentowych — w większości języków naturalnych jest to dość niewielka klasa wyrażen, ale warto o niej wspomnieć z dwóch powodów. Po pierwsze — do tej właśnie klasy należy jedyny przykład dyrektywy empirycznej, który podaje Ajdukiewicz (wspomniane już zdanie *Boli*). Po drugie, wychodząc od tej grupy wyrażen łatwo będzie nam objaśnić sposób, w jaki możemy wyprowadzić znaczenie zdań z DTZ (a ściślej — podać przy jej udziale). Zauważmy, że w przypadku predykatów zeroargumentowych nie ma żadnej różnicy pomiędzy

¹⁸ Wielu przykładów tego typu dostarczają, rzecz jasna, rozważania Saula Kripkego, zob. Kripke (1972).

znaczeniem wyrażenia elementarnego i zdania elementarnego zbudowanego z tego wyrażenia. Nawet na zupełnie intuicyjnym poziomie (którego nie należy tracić z oczu, ponieważ owo intuicyjne pojęcie znaczenia chcemy wyeksplikować) trudno byłoby podać uchwytą różnicę pomiędzy znaczeniem terminu „grzmi” i zdania „Grzmi”. Gdyby w języku nie było żadnych innych predykatów, poza zeroargumentowymi, problem znaczenia zdań byłby już na tym etapie rozwiązany. Teoria dyrektywalna podawałaby nam znaczenie zdań elementarnych (jako równoznacznych z terminami elementarnymi), a dalszą część pracy wykonywałyby już wspólnie gramatyka z logiką¹⁹. Bazując na tym rozwiązaniu, spróbujmy zbudować coś analogicznego dla predykatów jedno- i więcej argumentowych. Sądzę, że rozwiązaniem najmniej inwazyjnym w stosunku do oryginalnej wersji DTZ będzie przyjęcie następującej reguły:

Znaczenie predykatu jedno- lub więcej argumentowego nie różni się od znaczenia elementarnego zdania zbudowanego z tego predykatu i zmiennych wolnych.

Intuicyjny sens tej zasady jest następujący — skoro DTZ podaje nam znaczenie słowa „spi”, to podaje nam też znaczenie zdania „Coś spi”. Skoro podaje nam znaczenie wyrażenia „spogląda na”, to podaje nam również znaczenie zdania „Coś spogląda na coś”. Podaje, ponieważ nie ma pomiędzy tymi znaczeniami różnicy. Podobnie jak w przypadku zdań zbudowanych z predykatów zeroargumentowych dalszą pracę wykonują dla nas zewnętrzne teorie — to one informują nas, w jaki sposób interpretować należy różne sposoby, na jakie związane mogą zostać w tych zdaniach zmienne, to one dają nam wgląd w to, w jaki sposób interpretować mamy zdania złożone ze zdań elementarnych, spójników i operatorów. Pokażmy to na przykładzie. Weźmy pod uwagę przykładowe zdanie *Antoni spogląda na tęczę*. Rozszyfrowanie jego znaczenia odbywałoby się w następujących krokach. Najpierw gramatyka informuje nas o strukturze tego zdania. Dowiadujemy się zatem, że zdanie to głosi trzy rzeczy: jest ktoś taki, jak Antoni, jest coś, co jest tęczę, i to pierwsze spogląda na to drugie. Zdanie informujące nas o tym, że jest ktoś taki, jak Antoni, nie niesie ze sobą żadnego znaczenia, a jedynie ustanawia łańcuch komunikacyjny pomiędzy nami a Antonim. Sposób, w jaki się to dzieje, wyjaśnić może nam przyczynowa teoria odniesienia. Pozostałe dwa zdania objaśnia nam DTZ — to z niej dowiadujemy się, co to znaczy, że coś jest tęczę (jest to znaczenie słowa tęcza, czyli wszystkie miejsca, które to słowo zajmuje w macierzy języka) i co to znaczy, że coś na coś spogląda (jest to znaczenie wyrażenia „spoglądać na”).

¹⁹ Mówiąc to, pozwalamy sobie na dość istotne uproszczenie — jak wiadomo, teorie te nie sprawdzają się najlepiej w kontekstach nieekstensjonalnych, takich jak zdania o przekonaniach czy modalności, a przecież nie możemy kontekstów tych po prostu pominąć. Z tego, że jest to problem do rozwiązania, nie wynika jednak, że jest to problem do rozwiązania dla teorii znaczenia.

8. Schematy dyrektyw

Podsumujmy teraz zaproponowane modyfikacje DTZ. Po pierwsze, jako składnik dyrektyw empirycznych proponuję wprowadzenie bodźca całościowego. Po drugie, do istniejących już w DTZ trzech typów dyrektyw proponuję dodanie dyrektyw imperatywnych (w wersji prostej i złożonej), które nakazują pewne działanie (które rozumiem jako reakcję motoryczną) w wyniku uznania pewnego zdania (oraz, ewentualnie, odebrania jakiegoś bodźca). Po trzecie, proponuję wprowadzenie definicji brakujących w DTZ pojęć *posiadania znaczenia* i *posiadania znaczenia empirycznego*. Po czwarte — sugeruję sposób, w jaki możemy przy udziale DTZ wyprowadzać znaczenie zdań ze znaczeń wyrażen elementarnych.

Na zakończenie uporządkujmy rozważania, przedstawiając schematy pięciu wspomnianych typów dyrektyw (trzech pierwotnych i dwóch imperatywnych), zmodyfikowane zgodnie z zastrzeżeniami, które pojawiły się nam po drodze. Wszystkie dyrektywy wpisujemy w schemat BODZIEC — REAKCJA. Przez bodźce rozumiemy bodźce całościowe (BC), czyli pobudzenia wszystkich receptorów w danej chwili. Dla wygody w bodźcach tych wyróżniamy część językową i pozajęzykową. Jest to rozróżnienie czysto konwencjonalne, ponieważ na poziomie fizjologicznym żadnej różnicy tu nie ma — słowa muszą być tak samo słyszane czy widziane jak to, czego nie uważamy za słowa, a jedynie za dźwięki lub kształty. Rozróżnienie to jest jednak dla celów teoretycznych przydatne, dlatego też przez część językową bodźca rozumieć należy wyizolowane postrzeżenie (reprezentowane przez $P[Z]$) realizacji dźwiękowej czy graficznej danego zdania. Reakcją na bodziec może być reakcja motoryczna (reprezentowana w poniższej tabeli przez RM) lub uznanie danego zdania (na przykład Z, co reprezentujemy za pomocą oznaczenia $U[Z]$). Zapis konkretnych schematów rozumieć zaś należy następująco:

W przypadku dyrektyw aksjomatycznych bodźcem całościowym może być dowolny bodziec, którego częścią jest postrzeżenie pewnego zdania; reakcją może być dowolna reakcja motoryczna, jeśli tylko powiązana ona będzie z uznaniem tego zdania.

W przypadku dyrektyw dedukcyjnych bodźcem jest dowolne pobudzenie receptorów, którego częścią jest postrzeżenie ciągu zdań Z' ; Z. Reakcja motoryczna, podobnie jak w przypadku dyrektyw aksjomatycznych, może być zupełnie dowolna, byleby spełniony był warunek, że uznawszy zdanie Z' , użytkownik języka uznaje również zdanie Z.

W dyrektywach empirycznych pojawia się dodatkowo odwołanie do konkretnej części pozajęzykowej bodźca — ma to być jakieś określone pobudzenie receptorów — w poniższej tabeli reprezentuję je jako BC- $P[Z]$. Oznaczenie to rozumieć należy jako bodziec całościowy z wyłączeniem części językowej. Dyrektywy te nakazują, by w obecności pewnego konkretnego pobudzenia, którego

częścią jest postrzeżenie pewnego zdania Z, użytkownik uznał zdanie Z, niezależnie od towarzyszącej temu reakcji motorycznej.

Ostatnią kategorią są dyrektywy imperatywne proste, które nie precyzują, jak wyglądać ma pozajęzykowa część bodźca całościowego — mówią jedynie, że w sytuacji, w której częścią aktualnego bodźca jest postrzeżenie pewnego zdania, to — jeśli użytkownik zdanie to uznaje — nastąpi pewna konkretna reakcja motoryczna. Dyrektywy imperatywne złożone różnią się od prostych tym, że precyzują również pozajęzykową część bodźca całościowego. Zastrzeżenie, że reakcja motoryczna nastąpić ma dopiero po uznaniu zdania, wyklucza sytuację, w której ktoś na przykład uparcie odmawia wykonania poleceń — uznanie takiego zachowania za manifestację braków w kompetencji językowej byłoby oczywistym błędem.

DYREKTYWA	BODZIEC		REAKCJA
	CZĘŚĆ POZAJĘZYKOWA	CZĘŚĆ JĘZYKOWA	
AKSJOMATYCZNA		P[Z]	U[Z]
DEDUKCYJNA		P[Z'; Z]	U[Z'] → U[Z]
EMPIRYCZNA	BC-P[Z]	P[Z]	U[Z]
IMPERATYWNA PROSTA		P[Z]	U[Z] → RM
IMPERATYWNA ZŁOŻONA	BC-P[Z]	P[Z]	U[Z] → RM

Bibliografia

- AJDUKIEWICZ K. ([1931] 1985), *O znaczeniu wyrażen*, [w:] tenże, *Język i poznanie*, t. 1. Warszawa: PWN, 102–136.
- AJDUKIEWICZ K. ([1933] 1985), *Język i znaczenie*, [w:] tenże, *Język i poznanie*, t. 1. Warszawa: PWN, 145–174.
- AJDUKIEWICZ K. ([1953] 1985), *W sprawie artykułu prof. A. Shaffa o moich poglądach filozoficznych*, [w:] tenże, *Język i poznanie*, t. 2. Warszawa: PWN, 155–191.
- AJDUKIEWICZ K. ([1964] 1985), *Zagadnienie empiryzmu a koncepcja znaczenia*, [w:] tenże, *Język i poznanie*, t. 2. Warszawa: PWN, 388–400.
- BUSZKOWSKI W. (2010), *O równoznaczności wyrażen w ujęciu Ajdukiewicza*, [w:] Grad J., Sójka J., Zaporowski A. (red.), *Nauka — Kultura — Społeczeństwo. Księga jubileuszowa dedykowana Profesor Krystynie Zamiarze*. Poznań: Wydawnictwo Naukowe Uniwersytetu im. Adama Mickiewicza.
- JEDYNAK A. (2003), *Ajdukiewicz*. Warszawa: Państwowe Wydawnictwo Wiedza Powszechna.
- KRIPKE J. ([1972] 1988), *Nazywanie i konieczność* (przeł. B. Chwedeńczuk). Warszawa: Instytut Wydawniczy PAX.

- MACIASZEK J. (2007), *Holizm Znaczeniowy Kazimierza Ajdukiewicza*. Łódź: Wydawnictwo WSHE.
- MARSONET M. (1997), *Davidson, Ajdukiewicz and conceptual schemes*, "Axiomathes", t. 8, nr 1–3, 261–280.
- NOWACZYK A. (2006), *Dyrektywalna teoria znaczenia, czyli dramat Filozofa*, [w:] tenże, *Poławianie sensu w filozoficznej głębi*. Łódź: Wydawnictwo Uniwersytetu Łódzkiego, 176–185.
- PUTNAM H. ([1975] 1998), *Znaczenie wyrazu 'znaczenie'*, [w:] tenże, *Wiele twarzy realizmu i inne eseje* (przeł. A. Grobler). Warszawa: PWN, 93–184.
- QUINE W. V. O. ([1960] 1999), *Słowo i przedmiot* (przeł. C. Cieśliński). Warszawa: Fundacja Aletheia.
- SELLARS W. ([1954] 1999), *Some Reflections On Language Games*, [w:] tenże, *Science, Perception and Reality*. Atascadero, California: Ridgeview Publishing Company, 321–358.
- ZMYŚLONY I. (2009), *Kazimierza Ajdukiewicza pojęcie aparatury pojęciowej*, „Filozofia Nauki” 2009/1, 85–106.

Rozdział 5

Kontekst znaczenia językowego

Marek Maciejczak
(Politechnika Warszawska)

Słowa kluczowe: Pojęciowy model świata, znaczenie językowe, wprowadzanie nowych znaczeń językowych, przedjęzykowy poziom doświadczenia, anomalie semantyczne, referencja, prawda

Badanie procesów poznawczych — funkcji i struktury umysłu — pozwala lepiej zrozumieć język i odwrotnie, im większa wiedza o języku, tym lepsze rozumie umysłu. Badania akwizycji języka u dzieci i genezy kategorii językowych wyraźnie wskazują, że procesy poznawcze mają miejsce w dynamicznej całości — „modelu świata”, którego elementem jest system pojęć. W jego świetle znaczenie językowe nie jest gotową treścią, ideą, składnikiem semantycznym itd. To nie lista słów i reguł ich łączenia w zdania, lecz system procedur, operacji ustalania sensu wypowiedzi — twórczej aktywności nadawców i odbiorców¹.

Określenie „model świata” wprowadził wybitny psychiatra i filozof francuski Henry Ey ([1956] 1968) w swojej pracy *La conscience*. Obejmuje ono szereg aspektów składających się na fenomen bycia świadomym, stanowiąc systemową i zhierarchizowaną całość. Jednym z aspektów modelu jest język. Język bierze udział w spostrzeganiu i poznawaniu jako obiektywizujące medium, w którym formułowane są rezultaty poznania i tworzony indywidualny obraz świata:

Żaden przedmiot, dany nam w tradycyjnym znaczeniu tego słowa, nie doszedłby do prezentacji bez funkcjonowania tego modelu w tle naszego życia świadomego i bezustannie musimy się do niego odwoływać dla rozpoznania przedmiotów, z którymi mamy do czynienia; inaczej nie moglibyśmy odróżnić jakichkolwiek elementów czy momentów naszego pola świadomości, pole to pozostawałoby całkowicie puste (Półtański 1996: 110).

¹ W niniejszym tekście rozwijam idee naszkicowane w artykule *Znaczenie językowe jako fragment pojęciowego modelu świata* (zob. Maciejczak 2010).

Język i jego znaczenia są częścią indywidualnego systemu konceptualnego, który odzwierciedla werbalne i niewerbalne doświadczenia jednostki — modelu świata, który umożliwia efektywność języka, powstawanie przekonań i wprowadzanie nowych znaczeń².

W hierarchicznej strukturze modelu świata doświadczenie zmysłowe pełni rolę podstawy dla orientacji podmiotu w świecie i formowania pojęć i znaczeń, jakie dla nas mają obiekty doświadczenia, w tym też kwalifikacji „prawdziwy lub fałszywy”. Do modelu świata należy się odwołać, by zrozumieć, jak powstaje doświadczenie i wiedza, prezentacja tego, co bezpośrednio doświadczane i reprezentacja semantyczna (pojęciowa) doświadczenia. Ponadto, jak tworzy się swoista autonomia bytu świadomego i jego osobowy charakter w trakcie osobistego doświadczenia i, co szczególnie ważne dla zagadnienia znaczenia, w społecznym kontekście komunikacji. Co więcej, wyjaśnia też, jak powstaje ukierunkowane odnoszenie się człowieka do świata, uczestników komunikacji i samego siebie jako kryteriów subiektywnej ważności poznania.

Pojęcie modelu świata zbliża opozycje — naturalistyczne i komputacyjne teorie umysłu, ukazując związek percepcji z językiem, znaczeń wyrażen językowych ze schematami pojęciowymi, reprezentacji mentalnej z intencjonalnością i racjonalnością. Jest szansą przezwyciężenia jednostronności wymienionych stanowisk i aktualną potrzebą. Wychodząc od założeń i zasad — najbardziej ogólnych i powszechnych pojęć umożliwiających wielość aspektów ludzkiego obrazu świata — zrozumiemy funkcjonowanie języka, ponieważ połączone pojęcia podstawowe stanowią ramę przekonań i teorii w każdej dziedzinie naszych zainteresowań i działań. Zestaw pojęć potrzebnych do opisu doświadczenia jest dokładnie taki sam jak zestaw pojęć potrzebnych do opisu świata. Strawson np. zalicza do nich pojęcia: czasu, przestrzeni, przedmiotu, ciała, zdarzenia, faktu, tożsamości, znaczenia, prawdy, przyczynowości, wyjaśniania (Strawson 1994).

Stwierdzenie, że znaczenie, jakie nadaje się znakom (treść), określone jest przez zawartość systemu konceptualnego, w jakim dokonuje się interpretacji, kwestionuje mocno zakorzeniony kartezjański obraz umysłu jako dziedziny oczywistości i prawdy, wrodzonych pojęć jako podstawy wiedzy o świecie. Umysł ujmuje rzeczy wprost, przez jej idee, nie przez słowo; powodem różnic w sądach jest nie tyle niejasny język, co niewłaściwe użycie umysłu. Kartezjusz pomija to, w jakim stopniu myślenie związane jest z percepcją, działaniem i językiem oraz swoim wyrazem w mowie. Nie jest prywatną dziedziną całko-

² Wyniki tych prac i badania znajdziemy w: E. Holenstein, *Von der Hintergebarkeit der Sprache. Kognitive Unterlagen der Sprache. Anhang: Zwei Vorträge von Roman Jakobson*, Frankfurt 1980; red. G. Shugar, M. Smoczyńska, *Badania nad rozwojem języka dziecka*, Warszawa 1980; J. Piaget, *Psychologia i epistemologia*, tłum. Z. Zakrzewska, Warszawa 1977; red. E. Dąbrowska, W. Kubiński, *Akwizycja języka w świetle językoznawstwa kognitywnego*, Kraków 2003.

wicie niedostępną dla zewnętrznego obserwatora. Nawet najbardziej prywatne myśli związane są z językiem i cielesną ekspresją i w tym znaczeniu są dostępne innym, a nie tylko temu, kto je żywi. Co więcej, cielesne zachowanie w takim stopniu stanowi kryterium zachodzenia pewnego procesu umysłowego, intencji czy też pragnień osoby działającej, że nie znając cielesnego wyrazu stanów umysłu opisywanych słowami: „ból”, „radość”, „smutek”, „nadzieja” itd., nie można ich zrozumieć. Wbrew Kartezjuszowi umysł, aby był zdolny do nabywania i spełniania takich funkcji, jak: opanowanie języka, tworzenie i dysponowanie wiedzą itd., musi być włączony w szerszą strukturę żywego, ludzkiego organizmu funkcjonującego w kontekście społecznym, być jego częścią.

We współczesnej semantyce Kartezjusza teorię języka określa się jako referencjalną (realistyczną). Teoria referencjalna opatruje przedmioty nazwami własnymi, dlatego też określa się ją jako teorię etykietek: nazwy nazywają przedmioty, natomiast do nazywania stanów rzeczy i własności służą zdania i predykaty. Ciekawe jednak, że Kartezjusz, chwalać rozum ludzki jako „wszechstronne narzędzie”, wskazuje na jego zdolność do oznajmiania innym naszych myśli właśnie za pomocą języka, co sugeruje, jak głosi teza tej rozprawki, ściślejszy związek myślenia i języka w ramach pojęciowego modelu świata (Descartes 1970: 65).

1. Przedjęzykowe stadium uczenia się języka

Badania psycholingwistyczne dostarczają argumentów za istnieniem przedjęzykowego poziomu doświadczenia. Na przedjęzykowe tło składa się znajomość percepcyjna, performatywna i poznawcza sytuacji, która poprzedza i zarazem umożliwia wprowadzenie i ustalenie rozróżnień już czysto pojęciowych (językowych). Innymi słowy — język potrzebuje właściwego podłoża, kontekstu, który umożliwia definiowanie, określanie przez odniesienie, relacje do innych rzeczy (Bateson 1994: 32)³.

Piaget zauważył, że ze znajomością percepcyjną, spostrzeżeniem zmysłowej cechy wiąże się najczęściej moment funkcjonalny. Dziecko np. rozpoznaje piłkę jako coś, co się może toczyć, a więc coś, co pełni jakąś rolę w odniesieniu do jego motorycznej sytuacji. Piłka odcina się w polu widzenia przez to, że może być toczona. Dziecko powtarza czynność toczenia natychmiast i wiele razy, gdy piłka znajdzie się w jego polu widzenia — toczenie aktywizuje i zaspakaja jego zdolność poruszania się. Dziecko dostrzega jednocześnie określone percepcyjnie cechy piłki — że jest mniej lub bardziej okrągłym przedmiotem.

Znajomość percepcyjna umożliwia wczesną fazę przyswajania języka — przejście od gestu do słowa. W tej fazie rozwoju dziecko łączy wyrażenia językowe, pewne twory idealne, z konkretnym doświadczeniem. Zaobserwowano,

³ Bateson twierdzi, że każde spostrzeżenie jest czynnością klasyfikacji, a wzorce są wszechobecne w języku (Bateson 1994: 29, 43).

że u dzieci powiązanie znaczenia motorycznego i językowego, gestu i słowa jest poprzedzone przez powiązanie wskazania ręką z kierunkiem spojrzenia, a więc z dwiema oznakami już wcześniej przyswojonymi (Clark 1976: 94). Zdolność posługiwania się znakami symbolicznymi, jakimi są zaimki wskazujące: *to* — *tamto*, zakłada zrozumienie gestów wskazywania spełnianych zwykle przez rękę (ramię) i palec, najczęściej wraz z towarzyszącym skierowaniem spojrzenia. Doświadczeniu, że coś w polu widzenia odcina się i ściąga na siebie uwagę własną i partnera, zwykle towarzyszy zwrócenie się ciałem w tym kierunku. To zachowanie, stając się oznaką tego, co się odcina w polu widzenia, zyskuje funkcję wskazywania.

Wydaje się, że genetycznie przygotowani jesteśmy do uchwycenia owej funkcji znakowej — wskazywania. Już w piątej minucie życia dziecko szuka kontaktu wzrokowego z matką, a siedmioletni patrzą dorosłym w oczy, a nie na twarz, by wiedzieć, dokąd ma kierować się ich uwaga (Oksaar 1979: 149). Postępowi od spojrzenia i gestu wskazywania do werbalnego wskazywania (*Deixis*) towarzyszy zmiana w strukturze pola percepcji, postęp od przedmiotów rzucających się w oczy w bezpośredniej bliskości, poprzez coraz dalej położone, aż do przedmiotów znajdujących się poza polem widzenia. Dziecko przechodzi od tego, co obecne, do tego, co nieobecne, a następnie od znaków dla tego, co obecne, do znaków dla tego, co nieobecne⁴.

Własnością opisaną sytuacji jest to, że przedmiot ujęty zostaje jako pewne *to oto* i jako *coś takiego*, czyli jako indywidualny przedmiot pewnego rodzaju. Ten „dwojaki” charakter przedmiotów spostrzeżenia jest podstawą dla semiotycznego rozróżnienia oznaki i typu (*Token, Type*). Dlatego przedmiot nie jest tylko piłką, lecz jednocześnie pewnym prymitywnym typem znaku, przykładem (reprezentantem) PIŁKI. Piłką być to okazywać cechy: okrągły, elastyczny, do toczenia; podane cechy: okrągły, elastyczny, do toczenia pozwalają rozpoznać ów przedmiot jako piłkę. Percepcja, jak widać, ma także abstrakcyjny charakter, pomija bowiem nieistotne własności obiektu i jednocześnie wprowadza perspektywę podmiotu, który decyduje, jakie własności zostaną wybrane i ujęte jako aspekt, wskazówka czy oznaka. Na przykład można widzieć dom od strony dachu czy też dach jako oznakę istnienia domu, gdy widzi się tylko dach; ten sam kształt, która czyni przedmiot nożem, jest jednocześnie znakiem tego, że ów przedmiot może być użyty jako nóż. Jako wskazówki funkcjonują również rzeczy podobne i kontrastujące: biała sukienka może przypominać ośnieżony szczyt, osobę, która ją nosiła, ale też czarną sukienkę.

Przedjęzykowa faza przyswajania języka jest podstawą dla pojęć wywodzących się z percepcji. Istnieją między jednostkowymi zjawiskami podobieństwa i różnice, pokrewieństwa i preferencje wewnętrzne i one gwarantują istnienie

⁴ Zdaniem Holensteina uobecnienie tego, co nieobecne, staje się możliwe dzięki przejściu od znaków wskazujących, czy też pokazujących, do znaków jedynie przedstawiających (Holenstein 1980: 25).

spostrzeżeniowych rozróżnień. Spostrzeganie to klasyfikowanie; zawiera kodowanie między sprawozdaniem a tym, czego ono dotyczy. Podstawowe pojęcia są rozumiane przez dzieci zanim jeszcze potrafią one je wyrazić w formie symbolicznej. Dziecko np. chowa i wyjmuje zabawkę i w ten sposób utrwała pojęcie ciągłości istnienia obiektów. Nabywa się w dzieciństwie w ten sposób tysiące pojęć zanim jeszcze potrafi się wyrazić je symbolicznie. Polegamy na nich przede wszystkim dlatego, że zostały uchwycone przez różnicujące zachowanie, a dopiero następnie potwierdzone i utrwalone przez odpowiednie wyrażenia językowe.

Teza głosząca, że pojęcia percepcyjne są wynikiem procedur generalizacji i klasyfikacji przedmiotów przyswojonych percepcyjnie i funkcjonalnie, znajduje potwierdzenie w badaniach początkowego stadium uczenia się języka, czyli stadium kodowania pojęć (Pavilionis 1991: 135)⁵.

Zrozumienie językowych odróżnień (fonemów) zakłada najpierw spostrzeżenie znaków językowych, czyli umiejętność odróżnienia zmysłowych postaci, które funkcjonują jako znaki języka. Zakłada spostrzeżenie znaków językowych i uświadomienie ich sobie jako specyficznych, jak i sytuacji ich użycia. Znak opiera się na określonej postaci, zwykle są to linie, słupki, strzałki, podłużne obiekty, które ściągają na siebie uwagę. Znaki językowe są spostrzegane jako znaki czegoś i jako znaki innych znaków, co wymaga zmiany nastawienia podmiotu, przejścia od asocjacyjnego powiązania (konkretnego) do semiotycznego (abstrakcyjnego) odsyłania (*Verweisung*). Powiązania wyodrębnionych w ten sposób przedmiotów ze znakami, znaczeniami i spostrzegania ich odtąd jako znaki przedmiotów doświadczenia i innych znaków.

Odróżnienia językowe pozwalają na poprawianie i precyzowanie już użytych w percepcji i działaniu odróżnień. To, że w pełni możemy na nich polegać, jest następstwem wspólnej kontroli, wspólnego użycia języka. Dopiero w tym kontekście należy mówić o językowym normowaniu. Potwierdza to Lorenz, powołując się na badania nad funkcją kierowania, wprowadzania, wyboru i uzasadniania przez językowe odróżnienia zakresu i zasięgu zachowań. Badania te wskazują, że uczymy się pewnych zachowań i dzięki nim rozróżnień, nie uświadamiając ich sobie jednocześnie językowo (Lorenz 1970: 174). Następnie to, co kodowane przez słowo, odnoszone jest do określonej struktury pojęć, związanej z kolei z innymi takimi strukturami. Ten sam znak językowy może być więc użyty do kodowania różnych pojęć i poprzez inne pojęcia wiąże się z całym systemem konceptualnym. W ten sposób system konceptualny bierze udział w interpretowaniu znaku, w nim wyraża się jego znaczenie. System konceptualny ponadto przechowuje i interpretuje za pomocą swoich struktur kontekst językowy i sytuacyjny użycia wyrażenia, gwarantuje dostęp do dowolnej dziedziny systemu zawierającej istotną dla interpretacji informację.

⁵ W dalszej części wielokrotnie nawiązuję do tego artykułu.

Nie tylko psycholingwistyka, również fenomenologia ukazuje związek doświadczenia z odpowiadającymi mu pojęciami, mechanizmy zależności czynnościowego i treściowego aspektu doświadczenia. Chodzi o tzw. reguły noematyczne lub ich zbiory, które powstają w konkretnych sytuacjach doświadczenia i odtąd funkcjonują automatycznie. Husserl nazywa je strukturami habitualnymi, ich sieć wyznacza style poznawcze typów możliwych przedmiotów doświadczenia. Z tego powodu każdy obiekt, każdy przedmiot odsyła do pewnej określonej prawideł struktury. Dzięki niej świat i jego przedmioty doświadczane są jako uporządkowane w typy, np. jako drzewa, krzewy, zwierzęta; bliżej, jako sosna, dzika róża, pies, żmija, jaskółka itd. Każde nowe indywiduum jest już wstępnie znane, przywołuje bowiem to, co podobne i już jest uchwycone zgodnie ze swoim typem w horyzoncie możliwego doświadczenia i z odpowiednimi wskazówkami podobieństwa; posiada niejako „z góry” przewidziane typy własności jeszcze niedoświadczonych, lecz już oczekiwanych. Dlatego, kiedy widzimy kota, natychmiast antycypujemy jego możliwe sposoby zachowania, typowy sposób poruszania się, jedzenia, zabawy itd. Również wiedza *explicite*, to, co wiemy o kotach, drzewach, figurach geometrycznych itd., wchodzi do tła doświadczenia. Reguły stanowią jakby „ramę” sensu, który ma zarazem charakter semantyczny czy *quasi*-semantyczny i pragmatyczny, łączy bowiem treściowy i czynnościowy aspekt doświadczenia, wyznaczając egzemplifikację pojęcia, czyli systematyczne konteksty jego doświadczenia⁶.

„Lingwiści specjalizują się w fonetyce, fonologii, morfologii, syntaktyce, leksykografii, ontologiach, semantyce, pragmatyce i innych dziedzinach, ale język zależy od integracji wielomodalnej informacji, łączy się z percepcją i myśleniem” (Duch 2011: 462). Wyjaśnienie znaczenia językowego odsyła do całości — sprzężonych ze sobą procesów spostrzegania, myślenia i działania.

⁶ Husserl w ostatnim okresie swojej działalności wysunął program fenomenologicznych badań genetycznych. W ich świetle jednostki każdej dziedziny i kategorii, w tym także przedmiot transcendentny — rzecz oraz spójny i nacechowany prawidłowościami świat został określony jako korelat „intencjonalnej konstytucji”, „systemu intencjonalnych operacji”. System jest dynamiczną jednością o własnym stylu działania i historii. Wraz z nim konstytuują się warunki możliwości doświadczenia *resp.* poznania przedmiotów doświadczenia — każdemu przedmiotowi odpowiada pewien kompleks przeżyć połączonych przez związki motywacyjne: jednoczesność, sukcesję, styczność, ciągłość i homogeniczność. Tworzą one jedność strumienia świadomości jako jedność motywacji. Owa przestrzeń, linie sił to struktury konstytuujące sens doświadczenia. Jego moment prospektywny wyznaczony jest przez odwołanie się do przeżywanej całości, jeszcze przed-tematycznej znajomości świata, z której wydobywany jest przez wyszczególnianie (*Besondern*) aktualnie doświadczany przedmiot. Wraz z aktualnym przedmiotem współdana jest skala możliwych, różnych sposobów dania, czyli potencjalny horyzont, który może być urzeczywistniony w dalszym biegu doświadczenia. Styl owego określania jest już „ściśle z góry przypisany”. Optymalne prezentacje są strukturalnie (intencjonalnie) powiązane z odpowiadającymi im pojęciami. Na warstwie habitualnej nadbudowują się struktury kategorialne oraz charaktery etyczne, estetyczne i kulturowe. Podane typy powiązań i struktur tworzą świadomość jako spójny i jednolity system, pole doświadczenia, w którym ujmowane są rzeczy i związki. Więcej na ten temat zob.: Maciejczak 2010.

2. Znaczenie deskryptywne nazw i wyrażeń okazjonalnych

W świetle systemu konceptualnego można przekonująco interpretować znaczenie wyrażenia okazjonalnego, np. ja, ty, on, tutaj, teraz itd.

Wyróżnia się aspekt konwencjonalny — wyrażenia okazjonalne mogą być rozpatrywane jako nazwy, odpowiednio, pierwszej, drugiej, trzeciej osoby itd., są one „ustalonymi” znaczeniami odpowiadających im wyrażeń, oraz aspekt pragmatyczny (deskryptywny) — pojęcia, które mamy o konkretnym lub możliwym obiekcie odniesienia wyrażenia okazjonalnego. Taki obiekt może być też przedmiotem odniesienia wielu innych wyrażeń, które wiążą się z określoną informacją, czyli pojęciami, jakimi ktoś dysponuje o tym obiekcie. Właściwe użycie wyrażeń okazjonalnych zakłada dostrzeżenie kontekstu referencji, rozróżnienie odpowiednich obiektów, skonstruowanie określonego pojęcia danego obiektu w systemie konceptualnym. Znaczenie deskryptywne nazw lub wyrażenia okazjonalnego jest „obrazem” w jakimś systemie konceptualnym, który może być ustalony werbalnie przez różne deskrypcje: odniesienie przedmiotowe staje się wówczas rezultatem kombinacji „ustalonych” i „deskryptywnych” znaczeń (Pavilionis 1991: 136).

Deskryptywne znaczenia są sposobami wyrażania i tym samym środkiem reprezentacji i prezentacji odpowiadających im obiektów. Nie ma potrzeby, aby użytkownik języka potrafił je wyrazić, ponieważ jako „obraz” zmienia się ono zależnie od samego obiektu i naszej perspektywy. Jego nazwa zachowuje ów obiekt przy każdej jego przemianie (np. w świecie powieści), jest „sztywnym desygnatorem” zrelatywizowanym do systemu konceptualnego, nie zaś w sensie obiektywnej identyczności. Wyrażenia okazjonalnego, np. „ja”, nie da się zredukować do żadnej deskrypcji reprezentującej moją myśl o sobie, ponieważ jego niewyraźność jest własnością systemu konceptualnego, a żadne wyrażenie nie wyczerpie treści z nim skojarzonych. Użycie zaimka „ja” nie oznacza przeto zdobycia samoświadomości, pewności własnego istnienia, lecz odkrycie względności, wymienności ról, która ma miejsce w społecznej komunikacji.

W psychologii rozwojowej mówi się w tym kontekście o przewyżczeniu dziecięcego egocentryzmu. Kiedy dziecko dysponuje już strukturą zdania oznajmującego, jego wypowiedź uwalnia się od sytuacji, w której akt mowy miał miejsce, od cech niegdyś wyznaczających odniesienie — rozdzieleniu ulega związek przedmiotów i zdarzeń od sposobów dania. W tej sytuacji widać potrzebę narzędzia zdolnego powiązać daną sytuację z wyrażonym w zdaniu zdarzeniem, uwzględnionymi w nim przedmiotami, czyli identyczność czasu, przestrzeni i niekiedy także osób — spełniającej akt mowy i biorącej udział w wypowiedzianym zdarzeniu. Tym narzędziem są właśnie relacyjne wyrażenia wskazujące. Przywracają one zdaniom odniesienie do aktualnej sytuacji.

Określenia relacyjne odwracalne uwalniają między innymi od egocentryczności wczesnodziecięcego sposobu widzenia. Dopóki wymiennność wskaźników

(*ja-ty, tu-tam, lewa-prawa*) nie ustali się, pokaże w poprawnym użyciu w sytuacji dialogu, nie ma stałego kryterium na to, że zostały one użyte do określenia relacji, a nie jako wyrażenia kwalifikujące — np.: „Brat znaczy tyle, co chłopiec” (Elkind 1962). Użycia kwalifikujące związane są bardziej z postawą konkretną, relacyjne natomiast świadczą o zdobytej przestrzeni możliwych pozycji, możliwych działań, postawie kategorialnej. Znany jest przykład dziecka, które przed opanowaniem predykatywnego użycia języka nie potrafiło ani zainicjować dialogu, ani odpowiadać na czysto informacyjne pytania słowami *tak* lub *nie* (Holenstein 1980: 35). Jakobson w tym kontekście mówi o dwóch uwalniających fazach w rozwoju językowym dziecka (Holenstein 1980: 186). Dopiero chwytając wymiennność, zwrotność wskaźników, dziecko staje się zdolne do właściwego dialogu, który wymaga takiej wymienności specyficznie językowych ról społecznych. Partner w prawdziwej rozmowie to nie tylko nadawca i odbiorca sygnałów, lecz również informacji symbolicznej i potwierdzeń.

Widać jasno, dlaczego nazwy własne i wyrażenia okazjonalne są idealnymi środkami porozumiewania się. Są neutralne wobec różnych pojęć skojarzonych z nimi i dlatego nie ma potrzeby istnienia zgody w odniesieniu do związanych z nimi deskrypcji. Co więcej, one pozwalają wiązać nie tylko różne pojęcia z tym samym obiektem przez jednego i tego samego użytkownika, lecz także różne doświadczenia poznawcze, różne myśli różnych użytkowników języka jako przypuszczalnie odnoszące się do tych samych obiektów. To dlatego język nie służy tylko do zwerbalizowania przedjęzykowego i językowego „obrazu” świata. Służy również do przybliżania indywidualnych „obrazów” do „obrazów” danej wspólnoty językowej — ustalania intersubiektywnych różnic, artykulacji świata przy pomocy wspólnego kodu. Ucząc się poprawnego użycia wyrażen językowych, przyswajamy odpowiednie rozróżnienia, klasyfikacje dotyczące wspólnego świata, w którym żyją komunikujące się podmioty. Przyswojenie języka to nie tylko opanowanie środków kodowania pojęć systemu konceptualnego, lecz również opanowanie środków społecznej komunikacji, czym jest konwencjonalne zorientowanie takich systemów. Jest to drugi, obok kodowania pojęć, konieczny warunek społecznego porozumiewania się indywidualnych użytkowników języka.

Manipulując pojęciami, możemy budować nowe struktury konceptualne, tworzyć nowe „obrazy” świata, ale ta logiczna możliwość tworzenia nowych pojęć zdeterminowana jest przez system konceptualny. Struktury konceptualne odnoszą się do możliwego doświadczenia jednostki, do możliwych stanów rzeczy, przedmiotów nierealnych, wyobrażonych — do światów możliwych. Gdy dziecko już potrafi budować zdania oznajmujące, wychodząc od tego, co rzeczywiste, mówi o tym, co możliwe i co nierealne: „W możliwości przekraczania granic rzeczywistego doświadczenia zawiera się poznawcze znaczenie symbolizmu, jak i też języka w ogóle” (Pavilionis 1991: 140). Indywidualne

„obrazy” mogą daleko odbiegać od „obrazów” odzwierciedlających świat realny, niemniej nie zrywają ciągłości z pojęciami odzwierciedlającymi rzeczywiste doświadczenie jednostki; nawet w najbardziej zdeformowanych światach osób cierpiących na schizofrenię znajdują się ślady intersubiektywnego świata. Wszystkie możliwe „obrazy” rzeczywiste i nierzeczywiste, konkretne i abstrakcyjne jako struktury konceptualne są częścią jednego systemu konceptualnego, w którym są powiązane w sposób ciągły. Ciągłość tłumaczy też obecność śladów języka potocznego w językach abstrakcyjnych.

3. Sensowność jako interpretacja w określonym systemie konceptualnym

Wyrażenie ma sens, o ile znajdzie się jego interpretację w określonym systemie. Treść znaczeniowa słowa zmienia się wraz z jego werbalnym otoczeniem, nie inaczej jak wygląd przedmiotu ze zmianą punktu widzenia. Sensowność to konstruowalność określonego „obrazu” za pomocą znaczeń zawartych w systemie, czyli wiązanie pojęć w pewną całość, nadawanie formy, wypełnianie luk między pojęciami, przekraczanie różnic, dystansu, kontrastów, przeciwieństw — dowolnego rodzaju nieporównywalności.

Przekonującym świadectwem procesu konstruowania (interpretacji) są wyobrażenia powiązania, przybliżania i osiągnięcia zawarte w etymologiach wyrażeń językowych, składających się na pole semantyczne słów „znaczyć” i „rozumieć”: łacińskie *comprehendere*, angielskie *to catch, to size, to follow*, francuskie *saisir, comprendre*, niemieckie *fassen, ergreifen, begreifen*, polskie *pojmować, ujmować, chwytać* itd. Dystans pomiędzy podmiotem a przedmiotem rozumienia uwidacznia się w niemieckim *verstehen*, angielskim *understand*. Litewskie *prasme* (znaczenie) i *suprasti* (rozumieć) to przekraczanie, przybliżanie, przyswojenie przedmiotu przez podmiot. Pavilionis, podając i dyskutując podane przykłady, sądzi, że w ich polu semantycznym widoczna jest idea bezpośredniości, intencjonalności — samego w sobie wyróżniającego się probierza znaczenia (Pavilionis 1991: 143), czyli kontekstu percepcyjnego i funkcjonalnego tych pojęć, o czym była mowa powyżej.

Źródłem znaczenia (sensowności) jest interpretacja w określonym systemie konceptualnym i dlatego wyrażenie może być bezsensowne w innym systemie czy w innej strukturze tego samego systemu, może być interpretowane przez różne pojęcia, różnie ze sobą powiązane. System podaje odpowiednią interpretację, jeśli istnieje odpowiedni kontekst wskazujący fragment systemu konieczny dla interpretacji. Chodzi o pojęcia lub struktury pojęć powiązane relacjami interpretacji z pojęciami skojarzonymi bezpośrednio z przedmiotem, sytuacją, tekstem itd. Taki fragment, zwany blokiem (modułem) istotnej informacji obejmującej konstytutywne pojęcia o różnym stopniu abstrakcji i treści, jest podstawą tworzenia przez jego użytkownika sensownych (możliwych) obrazów.

Użytkownik przede wszystkim wybiera te, które uważa za prawdziwe, a zbiór tak wyróżnionych przekonań, indywidualny system przekonań wyrażają tzw. postawy przekonaniowe (*propositional attitudes*): „Jestem przekonany, że...”, „Sądzę, że...” itd. Jednostkowy system wiedzy obejmuje informację dotyczącą codziennego doświadczenia (wraz z przedwerbalnym okresem konstruowania systemu konceptualnego), osobistą historię oraz systematyczną wiedzę kodowaną w tekstach naukowych. Ta obiektywna wiedza obejmuje pojęcia konwencjonalne — uzgodnioną wiedzę o świecie, ramę porozumiewania się między nosicielami różnych systemów konceptualnych. Elementy konwencjonalne, odzwierciedlające społeczne, poznawcze doświadczenie jednostki, to intersubiektywne sensory i sądy logiczne, ale rozpatrywane w oderwaniu od indywidualnego systemu konceptualnego przekształcają się w tzw. wiedzę obiektywną, która jest abstrakcją z jednostkowych systemów wiedzy i istnieje symbolicznie w korpusie tekstów naukowych, stanowiąc historycznie i społecznie zdeterminowane „naukowe obrazy świata” (Pavilionis 1991: 146).

3.1. Przykłady roli kontekstu w interpretacji anomalii semantycznych

Uważa się, że ograniczenie łączliwości, czyli zabronionego połączenia znaczeń, stanowi część sensu wyrazu i decyduje też o jego użyciu. Na przykład wyrazów „farba” i „niema” nie wolno łączyć — o farbie nie można orzekać, że jest niema, gdyż w rezultacie otrzymujemy wypowiedź bezsensowną. Słynne zdanie Chomsky’ego: „Bezbarwne zielone myśli wściekle śpią” — jest również przykładem takiego zabronionego połączenia.

Czy całkowicie? Czy nie trzeba wziąć pod uwagę kontekstu, czyli miejsca defektywnego zdania w całości wypowiedzi? Czy sam ciąg znaków składających się na wypowiedź przesądza o jej zrozumiałości i sensowności? Czy rozumienie znaku, czyli uchwycenie, co on znaczy, co denotuje, możliwe jest tylko na gruncie reguł, jakie wyznaczają kod?

Hormann przytacza następujące historyjki: „Wczoraj za świeżo pomalowaną ścianą policja odkryła zwłoki uduszonej kobiety. Policji nie udało się, jak dotąd, znaleźć żadnych odcisków palców ani innych wskazówek pomocnych w ustaleniu tożsamości ofiary i mordercy — farba jest niema”. „Pewnego razu, gdy Chomsky miał lat szesnaście i był jeszcze uczniem liceum, do jego pokoju, późnym wieczorem, weszła matka. Chomsky już spał, przewracał się jednak niespokojnie z boku na bok, zgrzytając zębami. Widząc to, pani Chomska powiedziała: »Oho! Bezbarwne zielone myśli wściekle śpią«” (Hormann 1994: 94).

Niezrozumiałe zdania stają się sensowne przez umieszczenie w kontekście. Ich znaczenie nie wynika automatycznie z systematycznego przekładu znaków zgodnie z regułami kodu. To decyzja słuchacza określiła te zdania jako metaforę, nie zaś jako niepoprawne. Rozumienie znaczenia wyrażań nie jest wyłącznie wynikiem uchwycenia, co dany znak znaczy, ale przede wszystkim

docieraniem do zamierzeń mówiącego, do tego, co ma on na myśli. Słuchacz, starając się zrozumieć, co mówiący ma na myśli, kieruje się zasadą, że wypowiedź jest sensowna i zgodnie z tą zasadą analizuje znaki i połączenia wyrazów tak, aby osiągnąć z góry założony cel — sensowność wypowiedzi. Zmienia w miarę potrzeby kody, pomija to, co usłyszał, zignoruje niekiedy reguły gramatyki, wymyśli nowe znaczenia słowa. Nie ma ustalonego programu, są wybierane różne programy: sensowność bądź jej brak trafniej określić jako to, co może słuchacz osiągnąć. To dlatego pojęcie znaczenia wyrażen językowych pozostaje wieloznaczne. Wieloznaczność wcale nie dyskredytuje jego ważności (wbrew sceptycyzmowi Quine’a głoszącego, że skoro nie da się znaczenia określić w pojęciach behawioralnych, to w ogóle nie ma sensu staranie o ugruntowanie semantyki). O znaczeniu należy mówić w odniesieniu do celów poznawczych i komunikacji, które stawiają sobie mówca i słuchacz (co podkreślają: Strawson, Grice, Austin, Searle).

Teoria znaczenia językowego powinna brać pod uwagę to, co jego użytkownik z nim robi, to znaczy, co mógłby, powinien itd. mieć na myśli, używając danego wyrażenia w danej sytuacji. „Słuchacz kieruje się w stronę sensu — w tradycyjnym filozoficznym znaczeniu terminu *kierować się ku* [czemuś], odpowiadającemu łacińskiemu terminowi *intendere* — ponieważ czynienie świata zrozumiałym, a stąd też czynienie zrozumiałymi wypowiedzi na tym świecie, jest jedną z ważniejszych antropologicznych potrzeb człowieka, nawet czymś dla niego niezbędnym” (Hormann 1994: 98–99). Słuchacz, rozstrzygając, czy ma do czynienia z anomalią semantyczną czy metaforą, wyrażeniem idiomatycznym czy nie, aktywnie wydobywa sens wypowiedzi, dysponując niejasną wiedzą o tym, co mówiący może mieć na myśli w tej sytuacji, w tym kontekście. Rozumienie jest konkretyzacją tej wiedzy. Istnieją różne poziomy rozumienia, o których decydują nastawienie i interesy słuchacza. Wypowiedzi mówiącego odwołują się do czegoś, co słuchacz już zna; na tej podstawie spaja nowe składniki z wcześniejszymi, kolejno odrzuca nietrafne, wstępne rozumienia, modyfikuje dotychczasowe sposoby interpretacji i pojęcia do momentu, w którym uzna, że osiągnął ostateczny poziom oczekiwany w danym momencie, czyli określony sens — to, co mówiący ma na myśli.

W naszych umysłach mamy jak gdyby pewną wiedzę o danej sytuacji i o działaniach, które osoby znajdujące się w tej sytuacji mogłyby podjąć, gdyby w ogóle zdecydowały się na działanie. Oczekujemy, że będą one działały w sposób rozumiały, stosując się w dużym stopniu do tych samych reguł, motywów, konwencji i gramatyk, do których my sami stosujemy się w tej sytuacji. [...]. W tym procesie tworzenia informacji słuchacz kieruje się tym, co wie o świecie, tym, co mówiący ma na myśli i wypowiada, a także swoją dominującą tendencją, by spostrzegać zdarzenia w świecie jako sensowne. [...] Gdy rozumienie z poziomów powierzchownych przenika na głębsze, dźwięki, wyrazy i wypowiedzi stają się jak gdyby przezroczyste, a w świadomości słuchacza pojawia się to, co mówiący ma na myśli (Hormann 1994: 99–101).

Wypowiedź jest raczej instrukcją dla słuchacza: myśl w ten oto sposób, spostrzegaj te a te związki.

Podsumowanie

Powyższe uwagi skłaniają do wniosku, że znaczenie językowe należy rozpatrywać w jego związkach z systemem pojęć, procesami percepcyjnymi, doświadczeniem i działaniem w świecie. Znaczenie językowe nie jest, przywołajmy stwierdzenie Putnama, „gotowym wyrobem”. Odsyła ono do szerszego dynamicznego modelu umysłu, „modelu świata” i w jego świetle nie jest gotową treścią, ideą, składnikiem semantycznym itd., lecz system procedur, operacji ustalania sensu wypowiedzi. Z tego punktu widzenia znane w filozofii języka referencyjne koncepcje znaczenia — Davidsona, Kripkego, Montague, Hintikka — bliższe są wyjaśnienia zagadnienia znaczenia. Natomiast koncepcje niereferencyjne — Husserla z *Badań logicznych* — gdzie znaczenie określono jako istniejące niezależnie od aktów, coś poza-światowego, identycznego i niezmiennego należy raczej odrzucić. Podobnie jak stanowiska Fregego, Russella i wczesnego Wittgensteina, utrzymujące, że dziedzina myśli istnieje obiektywnie i niezależnie od uchwytowania i wyrażania. Kolejnym przykładem jest teoria Chomsky’ego, w której znaczenie wyrażen jest pewną kombinacją niezależnych od konkretnego języka wrodzonych atomów semantycznych, a idealny użytkownik języka nie potrzebuje odwoływania się do wiedzy o świecie, wystarczy mu wiedza dotycząca tych znaczeń oraz opartych na nich związków wynikania semantycznego.

Skoro zdolność do rozpoznawania konkretnej sytuacji lub przedmiotu, jako mających ogólne własności, pełni rolę podstawy dla orientacji podmiotu w świecie, a następnie formowania pojęć i znaczeń, to nie należy języka przeciwstawiać rzeczywistości, odrywać język od procesów jej poznawania. Znaczenie znaku zależy od jego interpretowania — posługiwanie się znakami wymagającymi zrozumienia wiąże się z poznawaniem rzeczywistości jako zbliżaniem się do prawdy (Pelc 1994: 13–26). Dlatego w opisie czynności nadawania, odbierania i przetwarzania znaków (semioza) należy uwzględnić podmiot tych czynności: „Nie ma znaków poza ich użyciem, nie ma użycia znaku bez interpretowania tego znaku, nie ma interpretowania znaku bez podmiotu poznającego” (Pelc 1994: 15).

Język nie jest więc kodem dla myśli. Gdyby tak było, należałoby widzieć pojęcia (treści pojęciowe) jako zakodowane w słowach, a złożenia pojęć jako zakodowane w zdaniach. Rolą słów i zdań byłoby odzwierciedlanie struktury i treści pojęć. Mielibyśmy taką sytuację: nadawca koduje swoją myśl w języku, a słuchacz następnie ją rozkodowuje. Język jako kod dla myśli nie ma z nią nic wspólnego. W konsekwencji można by żywić myśl niezależnie od języka, chwycić (rozumieć) myśli i zachodzące między nimi relacje zanim przyswoimy sobie

język, przypisywać innym posiadanie myśli, niezależnie od wiedzy, jaką mają o języku. Ale tak z pewnością nie jest — wyrażanie myśli różni się od jej identyfikowania, dlatego nie wystarczy przedstawić myśl w języku, żeby ją zrozumieć, trzeba jeszcze wiedzieć, jaką myśl zdanie wyraża. Istotną cechą języka jest więc to, że jego zwroty i zdania rzeczywiście wyrażają swe znaczenia. Na tym polega różnica między językiem a kodem, decydująca o tym, że opanowanie języka umożliwia mówiącemu chwyatanie nowych myśli w nim wyrażonych (Dummett 1993: 61–62). Język jest nośnikiem myśli (*vehicle*), zdanie nie koduje myśli, ale ją wyraża i z tego powodu nie można pozbawić myśli jej językowego kształtu⁷. Tylko wtedy, gdybyśmy potrafili odczytywać i przysyłać myśli bezpośrednio, język byłby zbędny.

Powtórzmy: pojęcie modelu świata i w jego ramach funkcjonujących systemów konceptualnych pozwala zakwestionować częste w filozofii języka naturalnego absolutyzowanie znaczenia abstrahujące od funkcji, jakie język pełni w poznawaniu i porozumiewaniu się. Absolutyzacja znaczenia zamyka podmiot wewnątrz języka, co nie pozwala na wyjaśnienie referencji wyrażen językowych, ich związku z oznaczanymi przez nie obiektami świata. Język, umożliwiając komunikowanie się między ludźmi, umożliwia tym samym bycie we wspólnym, intersubiektywnym świecie i właśnie z tego powodu badanie struktury języka rzuca światło na najogólniejsze cechy rzeczywistości ludzkiej, pojęcia, za pomocą których myślimy o świecie (Dummett 1998: 8). Można mieć nadzieję, że teoria znaczenia, dysponując pojęciem modelu świata, obejmie wreszcie w jednolitej teorii niedające się dotąd powiązać w stopniu zadowalającym aspekty: relacji do innych użytkowników języka, do innych wyrażen tego samego języka, odniesienia do rzeczywistości, czyli tego, jak jest przez wyrażenia językowe reprezentowana. Innymi słowy, wyjaśni najważniejsze funkcje języka: poznawczą i komunikacyjną, uwzględniając jednocześnie cele i konteksty użycia wyrażen językowych i rozwiąże spór między zwolennikami teorii komunikacji-intencji i semantyki formalnej. Z pewnością znaczenie słów i zdań jest w znacznym stopniu sprawą konwencji i reguł, ale w pełni można je zrozumieć dopiero w sytuacjach komunikacji, których warunkiem są intencje mówiącego ukierunkowane na odbiorcę. Reguły są niezbędne do kierowania działaniami i określają ich cele, są regułami komunikacji, mówiący, przestrzegając ich, może osiągać swój cel, spełnić komunikację-intencję. I to stanowi ich istotę, służą do tego celu.

Idea modelu świata i języka jako jego aspektu pozwala zakwestionować mocną odmianę relatywizmu językowego, uzależniającą pojęcie prawdy od schematu pojęciowego zawartego w językach etnicznych. Chociaż język skłania

⁷ Sporną kwestią jest to, czy wszystko z dziedziny myśli można wyjaśnić przez odwołanie do kategorii językowych. Szczególnie treści percepcji niektórzy autorzy uważają za niepojęciowe, dlatego twierdzą, że sądy sformułowane na ich podstawie podważają tezę o pierwotności eksplanacyjnej filozofii języka wobec filozofii myśli (Peacocke 1997: 16).

do odwołania się w trakcie mówienia do aspektów doświadczenia, które są skodyfikowane w kategoriach gramatycznych, to niemniej ów schemat pojęciowy jest nieustannie weryfikowany przez ogólniejsze związki i zależności z całością struktur poznawczych, a ostatecznie przez spójność i równowagę całego modelu świata. Choć schematy języków różnią się między sobą, to jednak model (system) narzuca im pewne granice. To dlatego, choć w języku Eskimosów istnieje ponad 100 wyrażen oznaczających barwy śniegu, wiemy, o jakim przedmiocie jest mowa i o jakich jego własnościach. Część kategoryzacji pokrywa się. Widać też wpływy względów praktycznych. Choć konceptualny obraz świata zależy od czynników społecznych, kulturowych i osobistego doświadczenia, to jednak dzięki budowaniu i manipulowaniu systemami pojęciowymi można orientować systemy indywidualne na przyjęte w danej społeczności językowej ogólne normy, co jest warunkiem społecznego porozumiewania się użytkowników zarówno tego samego języka, jak i różnych języków. Idea modelu świata umożliwia komunikację ponad lokalnymi układami pojęciowymi: lokalne układy są zakorzenione w warstwie przedwerbalnej, w strukturach postaciowych percepcji i historii dotychczasowego doświadczenia — w całości systemu, przeto różnice w pojęciowych kategoriach języków etnicznych ostatecznie nie kłócą się z ideą obiektywnej prawdy, o ile potraktujemy ją jako „ideę regulatywną”, której treścią jest niezależność od układu pojęciowego. „Świat nie jest wyrobem gotowym”, toteż systemy symboliczne kształtują percepcję i jest ona od nich nieodłączna. Jednak nie ma tu pełnej dowolności i dlatego nie jesteśmy więźniami struktur gramatycznych i pojęciowych naszego języka i możemy wyjść poza kategoryzację w nim utrwaloną. Skoro ukryte w języku wzorce nie stanowią jedynych wzorców myślenia i poznawania i co ważne nie są też ustanawiane arbitralnie, to można bronić realizmu ludzkiego poznania i przeciwstawić się próbom uzależnienia prawdy od zmiennej historii społeczności językowej.

Bibliografia

- BATESON G. (1994), *Umysł i przyroda* (przeł. A. Tanalska-Dulęba). Warszawa: PWN.
 CLARK E. (1976), *From Gesture to Word. Human Growth and Development*. Oxford.
 DĄBROWSKA E., KUBIŃSKI W. (red.) (2003), *Akwizycja języka w świetle językoznawstwa kognitywnego*. Kraków: Universitas.
 DESCARTES R. (1970), *Rozprawa o metodzie właściwego kierowania rozumem i poszukiwania prawdy w naukach* (przeł. W. Wojciechowska). Warszawa: PWN.
 DUCH W. (2011), *Jak reprezentowane są pojęcia w mózgu i co z tego wynika*, [w:] *Pojęcia. Jak reprezentujemy świat*, Bremer J. Chuderski A. (red.). Kraków: Universitas.

- DUMMETT M. (1993), *Język i komunikacja*, [w:] *Filozofia języka* (przeł. T. Szubka), Stanosz B. (red.). Warszawa: Wydawnictwo Spacja, 54–77.
- DUMMETT M. (1998), *Logiczna podstawa metafizyki* (przeł. W. Sady). Warszawa: PWN.
- ELKIND D. (1962), *Childrens Conceptions of Brother and Sister*, "Piaget Replication Study V. The Journal of Genetic Psychology" 100, 129–136.
- EY H. ([1956] 1968), *La conscience*. Paris.
- HOLENSTEIN E. (1980), *Von der Hintergebarkeit der Sprache. Kognitive Unterlagen der Sprache. Anhang: Zwei Vorträge von Roman Jakobson*. Frankfurt.
- HORMANN H. (1994), *Zagadnienia procesu rozumienia* (przeł. M. Witkowski), [w:] *Znaczenie i prawda. Pisma semiotyczne*, Pelc J. (red.). Warszawa, 93–101.
- LORENZ K. (1970), *Elemente der Sprachkritik*. Frankfurt.
- MACIEJCZAK M. (2010), *Znaczenie językowe jako fragment pojęciowego modelu świata*, [w:] tenże, *Intencjonalność i znaczenie językowe*. Warszawa: Wydawnictwo Uniwersytetu Kardynała Wyszyńskiego.
- OKSAAR E. (1979), *Spracherwerb und Kindersprache in evolutiver Sicht*, [w:] *Der Mensch und seine Sprache*, Peisl A., Mohler A. (red.). Frankfurt.
- PAVILIONIS R. I. (1991), *Język, znaczenie, rozumienie i relatywizm*, [w:] Muszyński Z. (red.), *Język, znaczenie, rozumienie i relatywizm*. Warszawa: Wydawnictwo Uniwersytetu Marii Curie-Skłodowskiej, 131–169.
- PEACOCKE C. (1997), *Concepts without words*, [w:] Heck, R. G. jr. (ed.), *Language, Thought, and Logic: Essays in Honour of Michael Dummett*. Oxford, 1–33.
- PELC J. (1994), *Znaczenie i prawda: semioza — poznanie — interpretacja*, [w:] *Znaczenie i prawda. Rozprawy semiotyczne*. Warszawa: PWN, 13–26.
- PIAGET J. (1977), *Psychologia i epistemologia* (przeł. Z. Zakrzewska). Warszawa: PWN.
- PÓŁTAWSKI A. (1996), *Problematyka doświadczenia zewnętrznego w filozofii Romana Ingardena*, część II, „Kwartalnik Filozoficzny”, t. XXIV, z. 4.
- SHUGAR G., SMOczyńska M. (red.) (1980), *Badania nad rozwojem języka dziecka*. Warszawa.
- STRAWSON P. F. (1994), *Analiza i metafizyka. Wstęp do filozofii* (przeł. A. Globler) Kraków.

Rozdział 6

Czy filozofia języka potrzebuje pojęcia sądu logicznego?*

Filip Kawczyński (Uniwersytet Warszawski)

Słowa kluczowe: sąd logiczny, zdanie, znaczenie, sens, prawda, nastawienia sądeniowe, przekonania, nośniki wartości logicznej

1. Wstęp

Pozostając w zgodzie z założeniem dotyczącym funkcji, jaką ma pełnić niniejszy tom, a które zostało wyrażone w jego podtytule *Inspiracje i kierunki rozwoju*, w pracy tej postaram się *zainspirować* do namysłu nad sądami logicznymi tych, którzy dotąd omijali je szerokim łukiem. Z założenia tekst ten nie ma być ani apologią jakiegoś wybranego stanowiska wobec sądów, ani polemiką z którymś z możliwych ujęć, lecz właśnie zachętą do bliższego przyjrzenia się jednemu z kluczowych pojęć w filozoficznych badaniach nad językiem.

Gdyby chcieć stworzyć ranking najważniejszych, a zarazem najbardziej kontrowersyjnych pojęć z zakresu filozofii języka, pojęcie sądu logicznego z pewnością zajęłoby w takim zestawieniu wysoką lokatę. Spór o sądy stanowi kwestię wciąż otwartą, niebywale szeroką i bezpośrednio powiązaną z licznymi trudnymi i fundamentalnymi zagadnieniami filozoficznymi, jak choćby istnienie abstraktów, poznanie abstraktów oraz natura prawdy. Wstępny namysł nad przyczynami problematyczności pojęcia sądu prowadzi do wniosku, iż są nimi w głównej mierze pewne zobowiązania ontologiczne, które wiążą się nieuchronnie z przyjmowaniem niemal każdej teorii angażującej sądy. Bez chwili wahania zgodzimy się, iż prawdą jest, że Napoleon wygrał bitwę pod Austerlitz i że był on doskonałym wodzem; jednakowoż, gdy każe się nam uznać, iż istnieją byty, takie jak sąd [że Napoleon zwyciężył pod Austerlitz] lub też własność *bycia doskonałym wodzem*, zaczynamy być podejrzliwi (być może ze względu na „upodobanie do krajobrazu pustynnego”, o którym wspominał

* Praca naukowa finansowana ze środków budżetowych na naukę Ministerstwa Nauki i Szkolnictwa Wyższego w roku 2012 (projekt badawczy nr 0174/B/H03/2010/39).

Quine 1948: 23) i domagamy się dodatkowych racji za uznaniem ich istnienia. Jeszcze trudniej byłoby nas nakłonić do powiedzenia, iż istnieje sąd [że Kutuzow zwyciężył pod Austerlitz], który — w przeciwieństwie do sądu [że Napoleon zwyciężył pod Austerlitz] — w jakiś sposób nie przystaje do rzeczywistości lub też niewłaściwie z nią koresponduje, wobec czego nie możemy o takim sądzie orzec, że jest prawdą. Filozofów, którzy — kierując się tego rodzaju odmianą zasady Brzytwy Ockhama (lub innymi względami) — odrzucili pojęcie sądu, określam jako *antyrealistów* w kwestii sądów.

Z drugiej strony, korzyści teoretyczne z przyjęcia sądów są na tyle duże, że wielu filozofów — których możemy określić jako *realistów*¹ w kwestii sądów — zaakceptowało pojęcie sądu kosztem uwikłania się we wspomniane niedogodności ontologiczne. Co ciekawe, w literaturze filozoficznej zdecydowanie łatwiej znaleźć argumenty przeciwko sądom niż za sądami. Przyczyny takiego stanu rzeczy upatrywać można w tym, iż moc eksplanacyjna pojęcia sądu jest tak duża, że niektórzy autorzy uznają ją za rację wystarczającą do włączenia sądów w obręb swoich teorii i nie czują się w obowiązku, by uzasadniać ten krok. Przyjrzyjmy się wobec tego temu, na czym dokładnie polega atrakcyjność pojęcia sądu logicznego, tj. jakie fenomeny z zakresu filozofii języka pozwala ono wyjaśnić i w jaki sposób się tego dokonuje. Poniżej przedstawiam charakterystykę czterech zasadniczych funkcji przypisywanych sądom, których określenie i analiza ma dawać odpowiedź na pytanie postawione w tytule niniejszej pracy. Zgodnie ze wspomnianymi wcześniej założeniami dotyczącymi jej charakteru do minimum ograniczam w niej argumentację za sądami lub przeciwko nim², większą uwagę poświęcając przy tym pobudkom, jakie kierują lub kierowały tymi, którzy sądom przyznają rację bytu.

2. Sądy jako przedmioty pewnych aktów psychicznych

Wśród realistów powszechny jest pogląd, zgodnie z którym sądy są przedmiotami pewnej grupy aktów psychicznych, do których zalicza się m.in.: sądzenie, przekonanie, wierzenie, wątplenie, przypuszczanie, uznawanie (asercja), negowanie, myślenie itp. Twierdzi się, iż w sytuacji, w której jakaś osoba jest

¹ Wspomniane dwa główne podejścia do kwestii sądów będę w dalszej części pracy nazywał krótko „realizmem” oraz „antyrealizmem”, natomiast zwolenników tych stanowisk odpowiednio „realistami” i „antyrealistami”.

² Jak już także wspominałem, nie zamierzam prezentować tutaj przeglądu różnych stanowisk w kwestii natury sądów. Zasadniczo wyróżnia się dwa główne nurty w zakresie tego zagadnienia — strukturalny oraz funkcjonalny. Najważniejszymi przedstawicielami pierwszego z nich są: Gottlob Frege (1891, 1892, 1918), Bertrand Russell (1903), Scott Soames (1985, 1987, 2010a, 2010b), Nathan Salmon (1986, 1989) i Jeffrey King (2007), natomiast drugiego — Rudolf Carnap (1947), David Kaplan (1977, 1978, 1989), David Lewis (1970) oraz (być może przede wszystkim) Robert Stalnaker (1999). Przeglądowe omówienie (w języku polskim) teorii funkcyjnych oraz strukturalnych czytelnik znajdzie w artykule Tadeusza Ciecierskiego (2003).

o czymś przekonana, wierzy w coś, wątpi w coś, coś przypuszcza, zaprzecza czemuś itd. owym czymś, co stanowi przedmiot tych wszystkich aktów mentalnych, jest sąd. Inaczej mówiąc, akty psychiczne tego rodzaju polegają na pozostawianiu pewnej osoby w określonej relacji do pewnego sądu. Choć jest to fraza niezbyt zgrabna, można powiedzieć, iż sąd jest tym, co może być wierzone, sądzone itd. (por. Soames 2010: 2). Idąc za Arthurem Priorem (1971: 4–5), można omawianą ideę wyrazić w jeszcze inny sposób i stwierdzić, że sądy są przedmiotami myśli, tj. tym, co przez nas myślane³. W bardzo jasny sposób omawiane podejście do sądów relacjonuje Paul Horwich:

Według zwolenników sądów w każdym wypadku, gdy ktoś ma jakieś wierzenie, jakieś pragnienie, jakąś nadzieję lub inne spośród tzw. nastawień sądzeniowych, wówczas stan umysłu tej osoby polega na tym, że istnieje pewna relacja łącząca ją oraz obiekt specjalnego typu: mianowicie rzecz, która jest wierzona, pożądana jako faktycznie występująca, czy wobec której ma się nadzieję, iż faktycznie wystąpi itd. Tym samym, jeśli Oskar wierzy, że psy gryzą, przyjmuje się, że jego wierzenie sprowadza się do zachodzenia pewnej relacji, wierzenia, pomiędzy nim a pewnym sądem: mianowicie sądem, że psy gryzą (Horwich 1990: 86).

Podstawą uznania sądów za przedmioty wspomnianych aktów psychicznych jest spostrzeżenie analogiczne do tego, które Bertrand Russell (1912: 15) poczynił wobec danych zmysłowych i wrażeń, konstatując, że czym innym jest (określony jako wrażenie) akt czy też doświadczenie uświadamiania sobie pewnych danych zmysłowych, a czym innym są owe dane. Podobnie jest w wypadku, na przykład, bycia przekonanym — mój akt bycia przekonanym, że Napoleon zwyciężył pod Austerlitz, jest czymś różnym od tego, co stanowi treść owego przekonania (ta dystynkcja dotyczy, rzecz jasna, nie tylko bycia przekonanym, lecz także pozostałych aktów psychicznych, o których była mowa powyżej). Mówiąc najkrócej, według zwolenników sądów racją za przyjmowaniem sądów jest to, iż w żadnym wypadku nie należy utożsamiać tego, co stanowi przedmiot danego aktu — tego, na co ów akt jest nakierowany — z samym tym aktem⁴.

³ Podana interpretacja to pierwszy z dwóch wskazanych przez Priora sensów, w jakich mówimy o przedmiotach myśli. Zgodnie z tym ujęciem, gdy myślę, że Napoleon wygrał pod Austerlitz, przedmiotem mojej myśli jest sąd [że Napoleon wygrał pod Austerlitz]. Zgodnie z drugim odczytaniem, przedmiotem mojej myśli jest m.in. Napoleon.

⁴ Dla zachowania ścisłości wyводу należy w tym miejscu wskazać raczej przemawiające za tym, że zdania — które wydają się innymi naturalnymi kandydatami do pełnienia tej roli — nie są przedmiotami omawianych aktów psychicznych. Najogólniej rzecz biorąc, przypisanie tej funkcji zdaniom prowadzi do zbytniego rozdrobnienia tych aktów. Okazuje się bowiem, iż jeśli w jednym wypadku przedmiotem mojego aktu bycia przekonanym jest zdanie „Napoleon zwyciężył pod Austerlitz”, a innym razem zdanie „Napoleon wygrał pod Austerlitz”, to należy uznać, iż są to dwa różne akty bycia przekonanym. Wobec tego może okazać się, że prawdą jest, iż „Filip jest przekonany, że Napoleon zwyciężył pod Austerlitz”, natomiast fałszem, że „Filip jest przekonany, że Napoleon wygrał pod Austerlitz” (to drugie

W ujęciach historycznych często jako tego, który wprowadził do filozofii języka pojęcie sądu (w nowoczesnym rozumieniu), wymienia się Bernarda Bolzana (por. np. Loux 1998: 133; McGrath 2007; Prior 1971: 6). To, co później zwykło określać się jako sąd, Bolzano (1837: §19) nazywał *zdaniem w sobie* (*Sätze an sich*). Zdania w sobie w ujęciu Bolzana są treściami aktów psychicznych, a przy tym czymś różnym od tych aktów i niezależnym od nich. Różnica pomiędzy zdaniami w sobie, czyli sądami, a procesami mentalnymi mającymi owe sądy za swe przedmioty jest wyraźnie widoczna w stwierdzeniu Bolzana, w którym przestrzega on, że „nie wolno nam także mylić go (sądu — przyp. F. K.) z ideą obecną w świadomości, a także z uznaniem jego prawdziwości” (Bolzano 1837: 49)⁵.

Zasługi Bolzana jako fundatora dyskusji o sądach są niekwestionowane, jednak tymi, których można uznać za odpowiedzialnych w największym stopniu za to, iż pojęcie sądu na stałe zagościło w filozofii języka, są G. E. Moore oraz Frege, w pismach których również można znaleźć wyraźne oddzielenie aktu psychicznego polegającego na pojmowaniu sądu od samego sądu.

Jak wiadomo, przełom XIX i XX wieku był okresem, w którym Moore wraz z Russellem, wzajemnie się inspirując, odeszli od zajmowanych wcześniej stanowisk idealistycznych i zdeklarowali się jako realiści. W artykule stanowiącym wyraz zmian zachodzących w poglądach Moore’a — których ważnym

zdanie mogło mi po prostu nigdy nie przyjść do głowy). Konsekwencja ta jest bez wątpienia niezgodna z intuicjami, które każą nam uznać, iż w obu powyższych wypadkach jestem przekonany o tym samym — mianowicie o tym, że Napoleon był zwycięzcą bitwy pod Austerlitz, a przekonanie to mogę wyrazić za pomocą różnych zdań (w szczególności przy pomocy dwóch powyższych). Obrońca koncepcji zdań jako podmiotów omawianych aktów psychicznych może na to odpowiedzieć, że zdanie ze „zwyciężył” oraz zdanie z „wygrał” to w zasadzie to samo zdanie, ponieważ czasowniki te mają to samo znaczenie, są synonimami. Takie postawienie sprawy jest jednak równoznaczne z uznaniem sądów za przedmioty tych aktów — przynajmniej przy założeniu poprawności zasady kompozycyjności, wedle której znaczenie zdania (za które uważa się sąd przez nie wyrażany — por. poniżej część 4) jest funkcją znaczeń poszczególnych wyrazów tworzących to zdanie. Kryterium, zgodnie z którym „Napoleon zwyciężył...” i „Napoleon wygrał...” można uznać za ten sam przedmiot aktu bycia przekonanym, jest kryterium opartym na identyfikacji sądów wyrażanych przez te zdania. Krótko mówiąc, stwierdzając, iż te zdania są tożsame, ma się w istocie na myśli to, że wyrażają one ten sam sąd. Słowem, dwa akty psychiczne — jeden angażujący zdanie „Napoleon zwyciężył...”, a drugi zdanie „Napoleon wygrał...” — odpadają pod ten sam akt-typ pojęty jako zbiór wszystkich aktów, w których jestem przekonany o tym, że Napoleon był zwycięzcą bitwy pod Austerlitz. (Można również uznać, iż omawiane zdania są tożsame, ponieważ opisują ten sam fakt — wówczas jednak niechybnie popada się w problem z łamigłówką Fregego — por. poniżej część 4).

⁵ Gwoli ścisłości warto dodać, iż Bolzano niemieckiego terminu „Urteil” — który najczęściej jest traktowany jako odpowiednik angielskiego „proposition” i polskiego „sąd” — używał w odniesieniu do wspomnianego w cytowanym sformułowaniu uznawania prawdziwości. Mówiąc inaczej, „sąd” w terminologii Bolzana odpowiada temu, co zwykle określamy jako „osąd” czy „akt sądzenia”, tj. akt psychiczny polegający na asercji lub odrzuceniu danego zdania w sobie (zob. Bolzano 1837: §34). Pomimo tych przeszkód terminologicznych nie ma wątpliwości co do tego, iż to właśnie Bolzanowskie zdanie w sobie odpowiada temu, co tradycyjnie rozumie się przez „sąd” w filozofii języka.

aspektem był postulat obiektywności⁶ sądów — o sądach pisze on, że „ich natura jest obojętna na to, czy ktoś je myśli, czy nie” (Moore 1899: 179). W późniejszej pracy kwestię odróżnienia sądu od aktu sądzenia Moore stawia jeszcze jaśniej:

Chcę tylko stwierdzić po prostu to, że każde z tych zdań logicznych [tj. sądów logicznych — przyp. F. K.]⁷ jest czymś, co można i trzeba odróżnić od aktu ujmowania, w którym ono jest ujmowane. Wszystkie akty ujmowania są podobne do siebie pod tym względem, że są aktami ujmowania i że są aktami tego samego rodzaju. Ale różnią się one pod tym względem, że jeden jest ujmowaniem jednego zdania logicznego [tj. sądu logicznego — przyp. F. K.], drugi zaś ujmowaniem innego (Moore 1953: 113).

W tym fragmencie Moore nie tylko dokonuje wyraźnego rozróżnienia na sąd i akt sądzenia, lecz daje również pewne wskazówki dotyczące tego, jaka zachodzi pomiędzy nimi relacja — jest to mianowicie relacja *ujmowania*, pod którą podpadają inne, wcześniej wspomniane związki pomiędzy podmiotem myślącym a sądem, takie jak bycie przekonanym, sądzenie, wierzenie itd. Moore stwierdza wprost, iż niemożliwe jest, by dwa różne sądy zostały ujęte w tym samym akcie. Nie do końca jasne na gruncie jego koncepcji pozostaje natomiast to, czy jeśli dwie osoby ujmują ten sam sąd lub jedna osoba ujmuje ten sam sąd w różnych momentach, to należy mówić o jednym akcie ujmowania tego właśnie sądu, egzemplifikowanym wielokrotnie, czy też raczej o wielu różnych aktach, w których ujmowany jest ten sam sąd. Wydaje się, że Moore skłania się raczej ku pierwszemu z tych rozwiązań, choć nie stwierdza tego *explicite*. Można to jednak wywnioskować z jego uwag, w których podkreśla, że sąd może być ujmowany również poza momentami, w których zdanie wyrażające ów sąd jest postrzegane lub słyszane (1953: 115–7) oraz wskazuje, iż ujmowanie przez kogoś jakiegoś sądu w danym momencie często nie wiąże się z ujmowaniem przez tę osobą w tej chwili rzeczy, których sąd dotyczy (1953: 131). Z uwag tych — oraz innych zrekapitulowanych wcześniej fragmentów koncepcji Moore’a — jasno wynika, iż Moore uznawał sądy za coś niezależnego od umysłów ujmujących te sądy. Stanowisko przyznające sądom autonomię może przybrać słabszą lub mocniejszą postać; w wersji słabej przyjmuje się, iż sąd nie istnieje, dopóki nie zostanie przez jakiś umysł pomyślany, natomiast od tego momentu sąd istnieje całkowicie niezależnie od istnienia jakichkolwiek

⁶ Moore — wbrew idealistom, takim jak Bradley — uważał, że sądy istnieją niezależnie od podmiotów pojmujących te sądy.

⁷ W polskim przekładzie cytowanej książki Moore’a używany przez niego termin „proposition” przetłumaczony jest jako „zdanie logiczne”. Sądzę, że niezastosowanie w tym wypadku zwrotu „sąd logiczny” spowodowane jest tym, że tłumaczenie owo powstawało w okresie, gdy polska terminologia filozoficznojęzykowa nie była jeszcze utarta. Obecnie, gdy ma ona już pewną tradycję i jest w wysokim stopniu ujednolicona, z pewnością w przekładzie użyto by zwrotu „sąd logiczny”. Ponadto, nie ma wątpliwości co do tego, że przedmiotem rozważań Moore’a jest właśnie to, co współcześnie rozumiane jest jako sąd.

umysłów. W wersji mocnej twierdzi się, że sądy istnieją z konieczności, co oznacza, na przykład, że istniałyby również wtedy, gdyby w ogóle nie zaistniały jakiegokolwiek umysły. W swojej wczesnej pracy poświęconej sądom Moore przyjmuje mocną wersję omawianego stanowiska. Jego zdaniem sądy są złożone z pojęć, natomiast pojęcia zostają określone jako niezmiennie i pierwotne wobec aktów ujmowania (Moore 1899: 179).

W wypadku Fregego odróżnienie sądów od psychicznych aktów ujmowania sądów było konsekwencją zajmowanego przezeń stanowiska zwanego antypsychologizmem, którego ogólna teza głosi, iż logika nie jest dyscypliną, której przedmiotem są procesy myślenia (tym zajmuje się psychologia), lecz bada ona byty będące przedmiotami myślenia (por. np. Frege 1892: 65–66; 1918: 123–126)⁸. Jednym z takich bytów są sądy, które Frege nazywa „myśłami” i które zostają wyraźnie oddzielone od aktów sądzenia:

Przez myśl rozumiem nie subiektywny akt myślenia, lecz jego obiektywną treść, która może być wspólną własnością wielu (Frege 1892: 68).

Zarówno ujęcie myśli, jak i jej osąd, są to akty poznającego podmiotu i należą do psychologii. W obu tych aktach jest jednak coś, co do psychologii nie należy, mianowicie myśl (Frege 1919: 134).

Postawienie przez Fregego granicy pomiędzy procesami myślowymi i przedstawieniami, z jednej strony, a myślami/sądami, z drugiej, jest nie tylko istotną dystynkcją w ramach jego dociekań logicznych, lecz także określeniem przedmiotu badań logiki i psychologii. Logika nie bada zrelatywizowanych do poszczególnych podmiotów przedstawień czy procesów myślenia — to pole logika oddaje psychologii — lecz obiektywne sądy, niezależne od podmiotów je pojmujących. Mówiąc obrazowo, a zatem siłą rzeczy niezbyt precyzyjnie, można stwierdzić, iż w koncepcji Fregego sąd pozostaje wobec psychicznego aktu sądzenia w takim stosunku, w jakim utwór muzyczny pozostaje wobec czynności odgrywania go na instrumencie. Podobnie jak Moore, Frege również uważał sądy za niezależne od umysłów i twierdził, iż „myśląc, nie wytwarzamy myśli, lecz je tylko ujmujemy” (Frege 1918: 125), a tym, co daje nam dostęp do myśli, jest język.

Rozróżnienie na sąd i sądzenie było obecne również w polskiej filozofii języka. Kazimierz Ajdukiewicz (1934: 146–149) przeprowadził dystynkcję na sąd w sensie psychologicznym oraz sąd w sensie logicznym. Sądy w sensie psychologicznym Ajdukiewicz określa jako procesy sądzenia, wydawania pewnego osądu; innymi słowy, są to akty psychiczne zachodzące w umyśle. Procesy te dzielą się na dwie podkategorie ze względu na możliwość ich ujęcia w języku. Jeśli z danego procesu sądzenia można zdać relację za pomocą wyrażen

⁸ W sprawie antypsychologizmu Fregego por. także Gut (2005: R. II).

jakiegoś języka, mówimy o procesach artykułowanych; jeśli zaś żadne wyrażenie językowe nie koresponduje dokładnie z sądzeniem, które odbywa się w umyśle, mamy do czynienia z procesem nieartykułowanym. Sąd w sensie logicznym jest natomiast tym, co stanowi treść artykułowanych procesów sądzenia czy też — posługując się terminologią wcześniej przeze mnie stosowaną — przedmiotem pewnych aktów psychicznych.

W powyższych — można powiedzieć: fundatorskich — koncepcjach sądów logicznych odróżnienie sądu od mentalnego procesu sądzenia jest wyraźne. Nie twierdzę, iż w każdym wypadku rozróżnienie to jest przeprowadzone w ten sam sposób. Na przykład, można zauważyć, iż w kontekście teorii Fregego o sądzie należy myśleć jako o czymś w pełni określonym, do czego akt sądzenia się odnosi i na kształt czego ów akt nie ma wpływu. Wydaje się, iż podobne intuicje stoją za ujęciem Moore'a. Inaczej jest natomiast w wypadku koncepcji Ajdukiewicza, który stwierdza, iż sądy logiczne stanowią treść sądów w sensie psychologicznym, wobec czego można domniemywać, że akt sądzenia w jakiś sposób determinuje treść sądu. Różnice te nie zaburzają jednak ogólnego obrazu, wedle którego sąd jako przedmiot psychicznego aktu sądzenia jest czymś różnym od tego aktu.

Przydzielenie sądom funkcji bycia przedmiotami pewnych aktów psychicznych ma tę korzyść teoretyczną, że pozwala, na przykład, wyjaśnić fenomen zbieżności przekonań. Jeśli osoba X jest przekonana, że Napoleon zwyciężył pod Austerlitz i jakaś inna osoba Y również uważa, że Napoleon zwyciężył pod Austerlitz, to można orzec, iż X i Y są przekonane o tym samym. Wyjaśnienie owej zbieżności przekonań X i Y nie może polegać na uznaniu, iż redukuje się ona do tożsamości procesów psychicznych X i Y, ponieważ nie istnieją dwa takie procesy mentalne, że jeden z nich jest procesem mentalnym X, drugi jest procesem mentalnym Y i procesy te są tożsame. Jeśli natomiast do głosu w tej sprawie dopuścimy sądy i za poprawny opis omawianej sytuacji uznamy:

- X jest przekonany [że Napoleon zwyciężył pod Austerlitz]
- Y jest przekonany [że Napoleon zwyciężył pod Austerlitz]

uzyskujemy wówczas przejrzyste wyjaśnienie zjawiska zbieżności przekonań w terminach relacji. Zarówno X, jak i Y pozostają w relacji bycia przekonanym do sądu [że Napoleon zwyciężył pod Austerlitz] czy też, mówiąc inaczej, ów sąd jest przedmiotem zarówno pewnych aktów psychicznych X, jak i pewnych aktów psychicznych Y.

Opisywana tu funkcja bycia przedmiotem aktów psychicznych dla antyrealistów, chcących argumentować przeciwko sądom, stanowi najprawdopodobniej, by tak rzec, najtwardszy orzech do zgryzienia. Tym, co mogą oni zarzucić adherentom sądów, jest zbędność sądów w wyjaśnianiu analizowanego feno-

menu. Można powiedzieć, iż gdy ktoś jest przekonany, że Napoleon zwyciężył pod Austerlitz, to nie dzieje się nic ponad to, że osoba ta po prostu postrzega pewne zdanie (np. „Napoleon zwyciężył pod Austerlitz”) i rozumie jego znaczenie. Nie ma potrzeby wprowadzania do tego obrazu dodatkowej czynności polegającej na ujmowaniu jakiegoś abstrakcyjnego bytu (por. Lycan 2000: 84). W odparciu tego zarzutu pomocne okazują się pewne uwagi poczynione przez Moore’a.

Po pierwsze, wskazuje on, że możemy ujmować sąd wyrażany przez dane zdanie mimo tego, że nie postrzegamy akurat tego zdania, a nawet wówczas, gdy „nie mamy przed naszym umysłem żadnych obrazów słów, które by były jego [tj. sądu — przyp. F. K.] wyrazem” (Moore 1953: 117). Mówiąc ogólnie, Moore przyjmował pogląd, wedle którego sądy można ujmować (przynajmniej w jakimś stopniu) nie tylko za pomocą aktów językowych (pojętych w szerokim sensie, tj. zawierających nie tylko wyrażanie sądów w mowie lub piśmie, lecz również przepowiadanie sobie zdań w myślach). Wydaje się, iż podobne intuicje leżały u podstaw wspomnianego wcześniej Ajdukiewiczowskiego pojęcia nieartykułowanego aktu sądzenia. Akty tego rodzaju można opisać jako te, w ramach których w umyśle pojawia się jakiś sąd, który jednak nie daje się dobrze wyrazić za pomocą języka:

Co krok spotykamy się w życiu codziennym z takimi procesami sądzenia. Widząc przy przechodzeniu przez jezdnię zbliżający się samochód, wydaję sąd, lecz chyba żadne zdanie języka nie odpowiada dokładnie memu sądowi. Tak samo rzecz się przedstawia, gdy przypominamy sobie sprawę, którą mamy załatwić. Tak też jest, gdy przychodzi pierwszy pomysł przy rozwiązywaniu zagadnienia naukowego. Znaną jest rzeczą, ile to trudu kosztuje, nim pierwszy, zrazu wcale niedający się wyrazić pomysł, tak dalece się w umyśle wyjaśni, że można go ująć w słowa (Ajdukiewicz 1934: 147).

Podane przez Ajdukiewicza przykłady sądzenia nieartykułowanego charakteryzują się tym, że sąd pojawiający się w *głowie* jest czymś — jak sam Ajdukiewicz określa to w innym miejscu — „mglistym”, natomiast każda próba ubrania tego sądu w słowa owocuje otrzymaniem zdania wyrażającego sąd daleko bardziej określony niż ten wyjściowy⁹. Moore, z drugiej strony, wskazuje na przykłady rozmiękania się języka i sądów, w których to wyrażenie językowe jest czymś „mglistym”, natomiast sąd wydaje się dobrze określony. Mowa tu o wyrażeniach takich, jak „Ogień!” albo „Deszcz”, które nie są w ogóle zdaniami, niemniej jednak wyrażają, zdaniem Moore’a, pewne sądy.

⁹ Trzeba w tym miejscu dodać, że kwestia możliwości istnienia przekonań niezależnych od języka bynajmniej nie jest niekontrowersyjna. Na jej temat toczy się debata wykraczająca poza zakres tej pracy, którą wywołał Davidson (1982), prezentując argumentację na rzecz twierdzenia, iż niemożliwe jest posiadanie przekonań przez istoty, które nie władają językiem.

Ogólnie rzecz biorąc, zarówno Moore, jak i Ajdukiewicz starają się wykazać, iż nie istnieje ściśle przyporządkowanie pomiędzy komunikatami językowymi a sądami, co stanowi odparcie naszkicowanego powyżej argumentu antyrealisty.

Ponadto Moore stawia tezę — w moim przekonaniu niekontrowersyjną — iż rozumienie zdania, do którego odwołuje się w swej krytyce przeciwnik sądów, nie jest procesem tak automatycznym, jak go widzi antyrealista:

Jest zupełnie jasne, jak myślę, że gdy rozumiemy znaczenie zdania gramatycznego, to w naszych umysłach zachodzi coś innego poza samym słyszeniem słów, z których składa się zdanie gramatyczne. Łatwo możecie się o tym przekonać, przeciwstawiając to, co zachodzi, gdy słyszycie zdanie gramatyczne, które rozumiecie, temu, co zachodzi, gdy słyszycie zdanie gramatyczne, którego nie rozumiecie. W pierwszym wypadku z pewnością poza samym słyszeniem wyrazów zachodzi inny akt świadomości, ujmowanie ich znaczenia, którego to aktu nie ma w drugim przypadku (Moore 1953: 112–3).

Antyrealista mógłby w tym miejscu replikować, twierdząc, że został błędnie zrozumiany, ponieważ wcale nie uważa rozumienia języka za tak prymitywny proces, jak mu się to przypisuje. Jego zarzut dotyczył raczej tego, iż poza rozumieniem znaczenia zdania nie jest potrzebne wprowadzanie dodatkowych procesów, takich jak ujmowanie sądów. Odpowiedź realisty może być w tym wypadku tylko jedna: rozumienie znaczenia to nic innego, jak ujmowanie sądu wyrażanego przez to zdanie. Okazuje się więc, iż antyrealista posługuje się terminem „znaczenie” w taki sposób, w jaki realista używa pojęcia „sądu”. Spór pomiędzy realistą i antyrealistą okazuje się kwestią czysto werbalną, ponieważ obaj akceptują to samo, lecz różnie to nazywają. Jedynym, co może w takiej sytuacji uczynić antyrealistę, by utrzymać swój sceptycyzm wobec sądów, jest odżegnanie się od pojęcia „znaczenia” i rzeczywiście niektórzy filozofowie przyjmują takie stanowisko. Ta kwestia zostanie nieco szerzej omówiona poniżej, w części 4.

3. Sądy jako odniesienie wyrażen o postaci „*że p*” w zdaniach o *nastawieniach sądzeniowych*

Z uznawaniem sądów za przedmioty aktów psychicznych bezpośrednio wiąże się kwestia analizy zdań o tzw. nastawieniach sądzeniowych (*propositional attitudes*), tj. takich zdań, jak: „*X* sądzi, że *p*”, „*X* jest przekonany, że *p*”, „*X* przypuszcza, że *p*” itp., a więc zdań o ogólnej postaci: „*X* ϕ że *p*”. Z syntaktycznego punktu widzenia zdania tego rodzaju składają się z wyrażen, którym można przypisać następujące pozycje syntaktyczne:

- operator główny (1,0) — predykat („ ϕ ”) oznaczający dane nastawienie sądzeniowe, np. bycie przekonanym, sądenie, przypuszczanie itp.;
- pierwszy argument (1,1) — fraza rzeczownikowa („*X*”) pełniąca niemal

zawsze funkcję podmiotu gramatycznego zdania i oznaczającą osobę, której przypisuje się w zdaniu dane nastawienie sądzeniowe;

- drugi argument (1,2) — wyrażenie o postaci „że p ”, gdzie „ p ” reprezentuje zdanie opisujące w szczegółach, jakie nastawienie sądzeniowe jest przypisywane podmiotowi. Frazę „że p ” określa się również niekiedy jako nominalizację zdania p .

W oparciu o omówioną w poprzednim punkcie tezę głoszącą, iż sądy są przedmiotami nastawień sądzeniowych, należy przyjąć, że występujące w zdaniach o nastawieniach sądzeniowych predykaty, takie jak „jest przekonany”, „sądzi”, „przypuszcza” itp., oznaczają relację zachodzącą pomiędzy osobą będącą odniesieniem wyrażenia zajmującego w rozpatrywanym zdaniu pozycję syntaktyczną pierwszego argumentu a czymś, co stanowi desygnat lub desygnaty wyrażenia występującego na pozycji drugiego argumentu. Zadanie teoretyczne, jakie wobec tego przed nami staje, polega na określeniu, co jest odniesieniem wyrażenia będącego drugim argumentem, a mówiąc inaczej, w relacji do czego pozostaje ten, kto przyjmuje pewne nastawienie sądzeniowe.

Realściści dostarczają prostego rozwiązania tej kwestii, polegającego na przyjęciu, że predykaty na pozycji operatora głównego oznaczają relacje dwuargumentowe, których pierwszym argumentem jest podmiot (osoba lub osoba-by), natomiast drugim sąd desygnowany przez wyrażenie „że p ”. Innymi słowy, na gruncie tego stanowiska przyjmuje się, iż wyrażenia o postaci „że p ” — występujące w zdaniach jako drugi argument — są z semantycznego punktu widzenia nazwami (jednostkowymi), tj. wyrażeniami, które oznaczają swój (jeden) desygnat. Twierdzi się następnie, że desygnatem tym jest sąd [że p] (por. np. Alston 2001: 44–45; Horwich 1990: 86–92; Lycan 1974: 431; Loux 1998: 134; McGrath 2007; Richard 1990). Rozwiązanie to koresponduje dobrze z intuicjami, które podpowiadają, iż wyrażenia o postaci „że p ” mają charakter nazwowy. Gdyby odwołać się do podręcznikowej definicji nazwy, zgodnie z którą wyrażenie jest nazwą, gdy można je podstawić na miejsce „A” lub „B” w zdaniu podmiotowo-orzecznikowym „A jest B” (por. Jadacki 2001: 124), to naturalnie wyrażenie „(to,) że Napoleon zwyciężył pod Austerlitz” należy uznać za nazwę, ponieważ możemy bez popadania w błąd utworzyć zdanie z tym wyrażeniem w miejscu podmiotu, np. „(To,) że Napoleon zwyciężył pod Austerlitz jest przyczyną odwołania Kutuzowa z funkcji naczelnego wodza armii rosyjskiej”.

Rozważmy zdanie:

(A) Maria wierzy, że Napoleon zwyciężył pod Austerlitz.

Zgodnie z interpretacją realistyczną, operatorem głównym w tym zdaniu jest predykat „wierzy”, który wyraża relację dwuargumentową, zachodzącą pomię-

dzy Marią a sądem [że Napoleon zwyciężył pod Austerlitz]. Analiza semantyczna tego zdania przedstawia się natomiast następująco:

- (i) „Maria” — n , „wierzy” — z/nn , „że Napoleon zwyciężył pod Austerlitz” — n , „że” — n/z , „Napoleon zwyciężył pod Austerlitz” — z

W tym ujęciu „wierzy” pełni w (A) funkcję funktora zdaniotwórczego od dwóch argumentów nazwowych (czyli typowego predykatu dwuargumentowego). Kluczowym punktem tej analizy jest uznanie słówka „że” za tzw. reifikator, tj. funktor przekształcający zdanie w jego nominalizację, która podpada pod kategorię nazw. Desygnatem powstałej w ten sposób nazwy jest sąd [że Napoleon zwyciężył pod Austerlitz].

Analiza ta pozwala dość łatwo wybrnąć z problemu tzw. kontekstów nieprzezroczystych, tj. takich, w których nie zachodzi wymienialność *salva veritate*, przez co należy rozumieć, że w kontekstach tego typu wymiana wyrażen na inne, posiadające to samo odniesienie, nie gwarantuje zachowania wartości logicznej. Na przykład jeśli w zdaniu (A) — o którym zakładamy, że jest prawdą — dokonamy wymiany nazwy „Napoleon” na inną, również desygnującą Napoleona, powiedzmy „zwycięzca spod Ulm”, otrzymamy:

(A') Maria wierzy, że zwycięzca spod Ulm zwyciężył pod Austerlitz.

które okaże się fałszem choćby w takiej sytuacji, gdy Maria nie wie, że zwycięzca spod Ulm i Napoleon to ta sama osoba. Realista może wyjaśnić brak spełnienia warunku *salva veritate* mimo koreferencyjności zamienionych wyrażen przez odwołanie się do różnicy sądów, które są desygnatami fraz występujących w (A) i (A') na pozycjach drugiego argumentu. W szczególności realista może orzec, iż sąd [że Napoleon zwyciężył pod Austerlitz] różni się od sądu [że zwycięzca spod Ulm zwyciężył pod Austerlitz], a przeto (A) zdaje sprawę z zachodzenia relacji pomiędzy Marią a innym przedmiotem niż ten, o którym mówi się w (A').

U podstaw odrzucenia powyższej, realistycznej analizy zdań o nastawieniach sądzeniowych leży fundamentalne dla stanowiska antyrealistycznego przekonanie o tym, że sądy są niepotrzebne do stworzenia dobrej teorii języka, ponieważ z zagadnieniami, do rozwiązania których sądy zostały w pewnym sensie powołane, można poradzić sobie bez odwoływania się do sądów. Antyrealiści uważają, że mówienie o sądach nie jest niczym innym, jak stawianiem — w niepotrzebnie zagmatwany sposób — zwyczajnych tez metajęzykowych dotyczących zdań.

Jeśli chodzi, w szczególności, o analizę zdań o nastawieniach sądzeniowych, to przeciwnicy sądów twierdzą, iż zdania te nie zdają sprawy z zachodzenia

jakiejs relacji wiążącej osobę X i sąd [że *p*], lecz z relacji zachodzącej albo pomiędzy X a zdaniem „*p*” (Quine 1960, 1970; Sellars 1963), albo pomiędzy X a przedmiotami, o których mowa w „*p*” (Prior 1971, Russell 1912).

Mówiąc inaczej, zasadniczą pomyłką, której popełnienie antyrealiści zarzucają realistom, jest tzw. błąd reifikacji czy też nominalizacji, polegający na uznaniu za nazwę — a zatem wyrażenie, którego funkcją semantyczną jest oznaczanie jakiegoś przedmiotu — czegoś, co w istocie nie gra roli nazwy (por. Dummett 2000: 266). Przeciwnicy sądów opowiadają się za uznawaniem słówka „że” za składnik wyrażen określających nastawienia sądzeniowe, a nie za reifikator. W związku z tym według antyrealistów głównymi funktorami/operatorami zdań o nastawieniach sądzeniowych nie są takie wyrażenia, jak „wierzy”, „sądzi” itp., lecz „wierzy, że”, „sądzi, że” itp.; innymi słowy, zdania te podpadają pod schemat o postaci „... sądzi, że ...”, a nie „... sądzi ...”. Analiza semantyczno-kategorialna zdania (A), wynikająca ze stanowiska antyrealistycznego, która nie pociąga konieczności uznawania sądów, przedstawia się następująco:

(ii) *Maria — n, wierzy, że — z/zn, Napoleon zwyciężył pod Austerlitz — z*

„Wierzy, że” gra tutaj rolę funktora zdaniotwórczego od dwóch argumentów: jednego nazwowego i jednego zdaniowego; zdanie będące drugim argumentem nie jest znominalizowane, przeto w całym (A) nie zawiera się nazwa, którą można by posądzać o desygnowanie jakiegoś sądu.

Najbardziej zażarty przeciwnik sądów, za którego można chyba uważać Willarda Van Ormana Quine’a (1960: §44) twierdził — w ramach swojego programu „ucieczki od intensji” — że w zdaniach o nastawieniach sądzeniowych nie mówi się o relacji pomiędzy osobą a sądem, lecz po prostu o tym, iż jakaś osoba odnosi się z asercją do pewnego zdania, czyli uznaje owo zdanie za prawdziwe. Quine jednakże nie mógł zgodzić się na uznanie, iż rozwiniętą i właściwą postacią zdania (A) jest:

Maria wierzy, że „Napoleon zwyciężył pod Austerlitz” jest prawdą.

z tego względu, że wyznawał pewną wersję deflacyjnej teorii prawdy (mianowicie tzw. dyskutowalną teorię prawdy), wedle której określenia „jest prawdziwy” i „jest fałszywy” są zbędne logicznie i semantycznie, ponieważ wygłoszenie „*p* jest prawdą” to ni mniej, ni więcej niż stwierdzenie „*p*”. W związku z tym Quine parafrazuje zdanie (A) jako:

Maria uznaje-za-prawdę¹⁰ „Napoleon zwyciężył pod Austerlitz”.

¹⁰ W oryginale: „believes-true”.

Zgodnie z zapowiedzią, nie będę wdawał się tutaj w szczegóły argumentacji przeciwko sądom. Warto jednak dodać, iż pomysły Quine'a dotyczące analizy zdań o nastawieniach sądzeniowych były w bardzo ciekawy sposób rozwijane przez Wilfrida Sellarsa (1963) oraz przez Priora (1971)¹¹. Wszystkie te analizy napotykają jednak na pewne przeszkody, z którymi nie są w stanie sobie poradzić i które stanowią tym samym broń w ręku zwolennika sądów. Problemy te, z jednej strony, przypominają kwestie kłopotliwe dla obrońców deflacyjnej teorii prawdy¹², z drugiej zaś strony, dotyczą tego, iż pewne treści (zwane przez realistów sądami) mogą być przedmiotem nastawień sądzeniowych, choć nie dadzą wyrazić się w języku (por. omówione wcześniej Ajdukiewiczowskie rozróżnienie na sądy artykułowane i nieartykułowane). Wnikliwe sprawozdanie z argumentacji antyrealistów i zarzutów, jakie można w ich stronę skierować, przedstawia Michael Loux (1998: 142–51).

Robiąc wyjątek od przyjętej wcześniej reguły niewdawania się w szczegóły argumentacji dotyczącej kwestii sądów, chciałbym w tym miejscu zwrócić szczególną uwagę na pewien argument, który jawi się jako wyjątkowo interesujący w świetle niniejszych rozważań. Jest to argument przeciwko analizowaniu zdań o nastawieniach sądzeniowych na sposób przedstawiony w (i), jednakowoż wnioskiem płynącym z tej argumentacji nie jest bynajmniej odrzucenie sądów.

Rozumowanie to, przedstawione przez Mieszka Tałasiewicza (2006: 124–130) odwołuje się — ogólnie rzecz biorąc — do samych zasad przeprowadzania analizy semantyczno-kategorialnej¹³. Nasze przykładowe zdanie, będące częścią (A), „Napoleon zwyciężył pod Austerlitz” można znominalizować na dwa różne sposoby, w efekcie czego uzyskamy albo nazwę „(to,) że Napoleon zwyciężył pod Austerlitz”, albo nazwę „zwycięstwo Napoleona pod Austerlitz”. Początkowo można odnieść wrażenie, że z punktu widzenia semantyki nie ma między tymi nazwami żadnej różnicy — wydaje się, iż opisują ten sam fragment świata (pewne zdarzenie) oraz nie różnią się pod względem posiadanego znaczenia. Tałasiewicz twierdzi jednak, iż identyczność własności semantycznych tych nazw jest pozorna. Parafrazując przykład Tałasiewicza, możemy powiedzieć, iż istnieje pewna istotna różnica pomiędzy takimi zdaniami, jak:

(B) Jan pamięta, że Napoleon zwyciężył pod Austerlitz.

(B') Jan pamięta zwycięstwo Napoleona pod Austerlitz.

¹¹ Stanowisko Priora w bardzo wielu szczegółach różni się od podejścia Quine'a, można jednak powiedzieć, iż Priora krytyka realizmu jest utrzymana w jak najbardziej quine'owskim duchu.

¹² Chodzi tutaj o znane przykłady zdań w rodzaju: „Maria wierzy, że wszystko, co powiedział Napoleon było prawdą/fałszem”. W takim wypadku przed antyrealistą staje trudne zadanie wskazania zdania, które jest przedmiotem wierzenia Marii.

¹³ Należy dodać, iż Tałasiewicz w swoim wywodzie powołuje się w dużym stopniu na prace Angeliki Kratzer.

Różnica ta polega przede wszystkim na tym, że do warunków prawdziwości (B') należy wymóg, by Jan był świadkiem zwycięstwa Napoleona — tylko wówczas można powiedzieć o Janie, iż pamięta to zwycięstwo — podczas gdy w warunkach prawdziwości (B) nie stawia się takich wymagań; Jan może żyć w XXI wieku i nadal będzie można o nim prawdziwie powiedzieć, iż pamięta, że Napoleon zwyciężył pod Austerlitz.

Idąc dalej tropem przykładów Tałasiewicza, możemy skonstruować zdanie, w którego skład wchodzi oba omawiane nominalizacje, np.:

(To,) Że Napoleon zwyciężył pod Austerlitz było Janowi wiadome, ale o (samym) zwycięstwie Jan nie wiedział nic.

które nie wydaje się niespójne wewnętrznie, choć niewątpliwie byłoby takie, gdyby użyte w nim nazwy-nominalizacje rzeczywiście oznaczały to samo. W związku z tym Tałasiewicz uważa, że o ile analiza taka jak (i) (tj. realistyczna), z jednym funktorem od dwóch argumentów nazwowych, jest poprawna dla (B), o tyle jest ona błędna w odniesieniu do (B') (co może świadczyć na rzecz antyrealizmu). Jego zdaniem zwycięstwo Napoleona pod Austerlitz (będące odniesieniem nazwy „zwycięstwo Napoleona pod Austerlitz”) stanowi tzw. korelat semantyczny zdania „Napoleon zwyciężył pod Austerlitz”, natomiast sąd (w terminologii Tałasiewicza, zapożyczonej od Kratzer: sytuacja-sąd): [że Napoleon zwyciężył pod Austerlitz] nie jest korelatem powyższego zdania, lecz pewnego zdania sformułowanego w metajęzyku, mianowicie: „(Zdanie) »Napoleon zwyciężył pod Austerlitz« jest prawdą”.

W opinii Tałasiewicza wspomniany błąd reifikacji, którym obarczona jest rzekomo analiza (i), bierze się z niesłusznego uznania, że gdy czasowniki takie jak „wierzy”, „sądzi”, „pamięta” itp. występują w sformułowaniach o postaci „X wierzy, że p”, „X sądzi, że p” itd., pełnią tę samą funkcję semantyczną, co w wypadku wyrażen w rodzaju „X sądzi o”, tj. funkcję funktora zdaniotwórczego od dwóch argumentów nazwowych:

Na przykład w zdaniach „X sądzi, że p” czy „X pamięta, że p” można za nazwę „(to,) że p” podstawić inną nazwę, dajmy na to „Y”, otrzymując poprawne i sensowne zdania „X sądzi Y-ka” i „X pamięta Y-ka”. Rzecz w tym, że czym innym jest pamiętanie jakiegoś Y-ka, a czym innym — pamiętanie, że p. Pamiętać Y-ka to mniej więcej tyle, co mieć reprezentację mentalną pewnej konkretnej rzeczy lub konkretnego zdarzenia Y — jak wtedy, kiedy pamięta się Jana lub przyjazd Jana. Natomiast pamiętać, że p — to tyle, co uświadamiać sobie, że p zachodzi (Tałasiewicz 2006: 129).

Tałasiewicz przyjmuje zatem, że wyrażenia takie jak „wierzy, że” itp. są syntaktycznie proste, a nie są zwrotami złożonymi z czasownika oraz słówka „że”. Wobec tego za poprawną analizę zdania (A) uznaje on (ii), która w rozwiniętej postaci przedstawia się, jego zdaniem, następująco:

- (iii) *Maria* — *n*, *wierzy, że* — *z/nz*, „*Napoleon zwyciężył pod Austerlitz*” jest prawdą — *z*, „*Napoleon zwyciężył pod Austerlitz*” — *n*, jest prawdą — *z/n*

Wyrażenie „jest prawdą”¹⁴ pełni tutaj funkcję typowego predykatu, tworzącego zdanie z nazwy, a nazwą tą jest w tym wypadku wzięte w cudzysłów zdanie dotyczące Napoleona. Jak wspomniałem, Tałasiewicz przyjmuje, iż korelatem metajęzykowych zdań o postaci „*p* jest prawdą” jest sąd [że *p*]. Okazuje się zatem, iż zgodnie z (iii) użyty w zdaniu (A) funktor „wierzy, że” wyraża relację zachodzącą pomiędzy Marią a sądem [że Napoleon zwyciężył pod Austerlitz], co jest wnioskiem pokrywającym się z efektami realistycznej analizy zdań o nastawieniach sądzeniowych.

Argumentacja Tałasiewicza jest zatem szczególnie interesująca w świetle rozważań nad sądami, ponieważ, z jednej strony, neguje się w niej poprawność typowej realistycznej analizy zdań o nastawieniach sądzeniowych, a z drugiej strony, jej rezultatem jest konkurencyjna analiza, która również angażuje pojęcie sądu i to w sposób niemal identyczny z tym, jaki miał miejsce w analizie realistycznej.

4. Sądy jako sensy/znaczenia zdania

Kolejnym zadaniem powierzonym sądom przez realistów jest bycie sensem/znaczeniem zdania oznajmującego (por. np. Loux 1998: 141–142; Lycan 1974: 130; McGrath 2007). Wiąże się to ściśle z dwiema wcześniejszymi funkcjami przypisywanymi sądom, albowiem, jak zostanie pokazane za chwilę, tym, co w istocie stanowi przedmiot bycia przekonanym, wierzenia, wątpienia itd., nie jest ani samo zdanie, ani jego odniesienie, lecz właśnie znaczenie zdania.

Pojęcie znaczenia odgrywa w filozofii języka rolę nie do przecenienia. Kłasykcyjne ujęcie dystynkcji znaczenie–odniesienie¹⁵ znajdziemy w przełomowej pracy Fregego (1892). Przypomnijmy, że powodem, dla którego Frege wprowadza owo rozróżnienie jest problem, który wszedł na stałe do słownika filozofów jako „łamiągówka Fregego” („Frege’s puzzle”), odnoszący się do różnicy w wartości poznawczej zdań dotyczących identyczności. W zdaniu „Gwiazda Wieczorna jest tożsama z Gwiazdą Wieczorną” stwierdza się, iż pewien przedmiot jest tożsamy z samym sobą; w zdaniu „Gwiazda Wieczorna jest tożsama

¹⁴ Tałasiewicz zastrzega, iż nie przyjmuje się tutaj żadnej z wersji deflacyjnej teorii prawdy.

¹⁵ Posługuję się tutaj utartą już terminologią, niezgodną z tym, w jaki sposób Wolniewicz tłumaczy na polski Fregowskie „Sinn” i „Bedeutung”. Choć w moim przekonaniu Wolniewicz postępuje słusznie w tym, że nie poprawia sformułowań Fregego i pozostawia w polskim przekładzie pewną niedoskonałość terminologiczną, występującą w oryginale, to, aby nie wprowadzać niepotrzebnego zamętu, terminy „sens” i „znaczenie” traktuję jako wyrażenia posiadające dokładnie ten sam... sens/znaczenie, czyli synonimy. Tym zaś, co odpowiada Fregowskiemu „Bedeutung” we współczesnej terminologii jest termin „odniesienie” (a także „nominat” oraz w zależności od kontekstu — „desygnat” lub „denotacja”).

z Gwiazdą Poranną” mówi się dokładnie to samo, ponieważ Gwiazda Wieczorna i Gwiazda Poranna to to samo ciało niebieskie. Jednocześnie o zdaniach tych można powiedzieć, iż są nośnikami nieco innego zasobu informacji — inna jest ich wartość poznawcza. Pierwsze zdanie z pary jest truizmem, podczas gdy drugie może stanowić znaczące rozszerzenie naszej wiedzy. Frege rozwiązuje tę kwestię poprzez oddzielenie tego, co desygnują wyrażenia „Gwiazda Wieczorna” i „Gwiazda Poranna” od tego, co te nazwy wyrażają, tj. od tego, co jest ich sensem/znaczeniem (por. Frege 1892: 60–62). Łamigłóвка Fregego została sformułowana w odniesieniu do nazw, jednak analogiczną można przedstawić również dla zdań. W parze zdań: „Gwiazda Wieczorna jest oświetlana przez Słońce”, „Gwiazda Poranna jest oświetlana przez Słońce” oba zdają się mówić o dokładnie tym samym — mianowicie, że pewna gwiazda jest oświetlana światłem słonecznym. W sytuacji jednak, w której odbiorca tych zdań nie zdaje sobie sprawy, że Gwiazda Wieczorna i Gwiazda Poranna to ten sam obiekt, zdania te będą miały nieco inną wartość poznawczą (por. Frege 1892: 68). O zdaniach również można powiedzieć, iż wyrażają różne sensy (ze względu na różne sensy zawartych w nich nazw „Gwiazda Wieczorna” i „Gwiazda Poranna”) i to właśnie ta różnica odpowiedzialna jest za ich odmienną wartość poznawczą. Określenie przez Fregego pojęcia znaczenia było równoznaczne z odejściem od prostszego, przyjmowanego dotychczas stanowiska (którego najbardziej znanym przedstawicielem był John Stuart Mill 1843: ks. I, r. 2), wedle którego — w uproszczeniu — pomiędzy znakami języka a odpowiadającymi im przedmiotami w świecie zachodzi relacja oznaczania, niezapośredniczona przez żadne dodatkowe medium. Na gruncie tego poglądu wytłumaczenie różnicy w wartości poznawczej w zdaniach omówionych powyżej nastręczało olbrzymie trudności. Funkcję wspomnianego medium, dzięki któremu można wyjaśnić fenomen odmiennej wartości poznawczej, pełni w ramach teorii Fregego pojęcie znaczenia.

Od kandydata na znaczenie zdania wymaga się zasadniczo, by spełniał dwa warunki: po pierwsze, musi być tym czymś, co jest wspólne wypowiedziom takim, jak np. „Śnieg jest biały”, „Snow is white”, „Schnee ist weiss” itd., lub też np. „Mój teść dobrze gotuje” i „Ojciec mojej żony dobrze gotuje”; po drugie, znaczenie musi determinować, co jest odniesieniem zdania. Sądom przypisuje się czynienie za dość obu tym warunkom.

Realisci twierdzą, że gdy Polak wypowiada zdanie „Śnieg jest biały”, wyraża ten sam sąd, co Anglik wygłaszający „Snow is white”. Zastosowanie pojęcia sądu pozwala zatem wyjaśnić, w jaki sposób następuje tłumaczenie zdań jakiegoś języka naturalnego na zdania innego języka naturalnego: z grubsza rzecz biorąc, polega ono na odnajdowaniu w obu językach takich zdań, które wyrażają te same sądy. Zdania spełniające ten warunek możemy określić jako synonimy, ponieważ przy użyciu pojęcia sądu synonimiczność można łatwo zdefi-

niować jako tożsamość wyrażanego sądu. Synonimiczne są więc również, należące do tego samego języka naturalnego, zdania „Mój teść...” oraz „Ojciec mojej żony...”.

Jeśli chodzi o wyznaczanie przez sądy odniesienia zdań, to w zależności od tego, jakie założenia przyjmujemy co do natury sądów, może się to odbywać dwojako.

Pierwsza możliwość to postrzeganie sądów jako obiektów złożonych, posiadających określoną strukturę. Elementami tej struktury są sensy wyrażane przez poszczególne wyrażenia proste składające się na dane zdanie¹⁶. Najczęściej przyjmuje się dodatkowo, iż struktura sądu ma charakter kompozycyjny — jest wyznaczona przez to, jakie sensy są jej elementami i w jakim porządku one występują (co, z kolei, ma związek z charakterystyką składniową zdania). W kwestii wyznaczania przez tego rodzaju strukturalne sądy odniesienia zdań możemy mówić o dwóch możliwościach. Gdy przyjmuje się, iż odniesieniem zdania jest jakiś fragment rzeczywistości, jakaś sytuacja, zdarzenie czy fakt, to ponieważ sens każdego z wyrażenń składowych wyznacza, co jest odniesieniem tego wyrażenia, możemy uznać, iż sąd pojęty jako całość wyznacza właśnie ów *wycinek* rzeczywistości, który można uznać za odniesienie zdania. Jeśli natomiast za odniesienie zdania uznamy jego wartość logiczną, jak to czynił Frege, wówczas prawdziwość lub fałszywość traktowana jest jako coś sprzężonego ze znaczeniem danego zdania (w pewnym sensie wartość logiczna jest w takim wypadku własnością znaczenia zdania). Zdaniem Fregego, gdy ujmujemy myśl wyrażaną przez zdanie (czyli jego znaczenie), to myśl ta zawsze jest prawdziwa albo fałszywa¹⁷, a więc zawsze to myśl przesądza o tym, jakie jest odniesienie zdania.

Druga opcja w podejściu do sądów polega na utożsamieniu sądu z funkcją. Dziedziną takiej funkcji jest zbiór światów możliwych, natomiast wartości, które może taka funkcja przyjmować są dwie: prawda albo fałsz. Na gruncie tego stanowiska przyjmuje się również, iż odniesieniem zdania jest wartość logiczna. Dlatego też można powiedzieć, iż sąd — jako funkcja „przeprowadzająca” światy możliwe w wartości logiczne — wyznacza odniesienie zdań¹⁸.

¹⁶ Wyjątkiem od tej reguły są sądy indywidualne (*singular propositions*), których składnikami nie są w pewnych wypadkach sensy wyrażenń prostych, lecz obiekty przez nie desygnowane.

¹⁷ Choć ujęcie myśli nie zawsze wiąże się z ujęciem jej wartości logicznej — por. Frege (1918: 125).

¹⁸ Podobnie jak w przypadku funkcji przypisywanych sądom przy analizowaniu zdań o nastawieniach sądzeniowych, również w sytuacji traktowania sądów jako sensów wyrażanych przez zdania, argumentację przeciwko sądom, która odbiła się najszerszym echem, przedstawił Quine. Zakres tej argumentacji wybiega daleko poza ramy niniejszej pracy, dlatego jedynie o niej wspominam. Quine’owska krytyka sądów jest rdzeniem jego programu „ucieczki od intensji”, którego przejawami są choćby tak doniosłe i niebywale szeroko komentowane kwestie jak: a) zanegowanie przez Quine’a istnienia podziału na prawdy analityczne i syntetyczne, a tym samym zaprzeczenie istnieniu jakichkolwiek prawd analitycznych; b) Quine’owski postulat niezdeterminowania przekładu, zgodnie z którym nie istnieją żadne

5. Sądy jako nośniki wartości logicznej

Wszystkie wyżej opisane role przypisywane sądom są niezwykle ważne z punktu widzenia dociekań filozoficznojęzykowych, jednakże chyba najbardziej *odpowiedzialnym* zadaniem, któremu sądy mają sprostać, jest bycie pierwotnym nośnikiem wartości logicznej. Sądy charakteryzuje się często jako właśnie to, o czym orzeka się prawdziwość lub fałszywość (por. Ajdukiewicz 1967: 113; Loux 1998: 135; Moore 1899: 180, 1953: 120). Podkreśla się tutaj bycie *pierwotnym* nośnikiem prawdziwości, ponieważ zdaniem zwolenników sądów nic nie stoi na przeszkodzie, by przyjąć konwencję językową, w ramach której wartość logiczną orzeka się również o zdaniach — należy przy tym jednak pamiętać, iż mówienie o prawdziwości zdań jest pewnym uproszczeniem, gdyż zdania jako takie wartości logicznej nie mają. Aby zobaczyć dlaczego to właśnie sądy zasługują na miano pierwotnych nośników prawdy, prześledzimy, jak w tej roli sprawdzają się ich główni konkurenci¹⁹.

Oprócz sądów naturalnymi kandydatami do bycia nośnikami wartości logicznej są zdania. Ponieważ zdanie może być pojmowane albo jako typ, albo jako egzemplarz, należy osobno rozważyć te dwie możliwości.

Zdanie-typ, pojęte jako zbiór wszystkich równokształtnych napisów (równobrzmiących sekwencji dźwiękowych) będących zdaniem, nie może być uznane za podstawowy nośnik prawdziwości przede wszystkim z powodów, które wskazał Peter Strawson (1950) w swej słynnej krytyce Russella teorii deskrypcji.

Na pytanie o wartość logiczną zdania-typu Strawson (1950: 385–386) odpowiada, iż to zdanie (jak wszystkie zdania-typy) nie ma wartości logicznej, wobec czego pytanie jest źle postawione. Zastanówmy się nad zdaniem-typem „Obecny premier Polski jest Kaszubem”. Wartość logiczną można orzekać

fakty dotyczące znaczeń wyrażeń języka pozwalające ocenić rzetelnie, które wyrażenia są synonimiczne; i wreszcie, będąca następstwem dwóch poprzednich: c) bodźcowa teoria znaczenia, odrzucająca wszelkie przedmioty intensionalne (w szczególności sądy), która daje dobre — zdaniem Quine’a — wyjaśnienie tych fenomenów, które realiści wyjaśniają poprzez odwołanie się do pojęcia sądu (wymienione wątki pojawiają się w wielu pracach Quine’a — podstawę powyższego streszczenia stanowiły przede wszystkim: 1951, 1960, 1970).

Ciekawe polemiki z Quine’owską krytyką sądów prowadzili Harman (1967) i Lycan (1974). W skrócie — Harman stara się wykazać, iż Quine odrzuca sądy nie dlatego, że są one obiektami podejrzanymi metafizycznie, lecz dlatego, że nie pozwalają one rozwiązać problemów, do rozwiązania których zostały powołane. Lycan w przekonujący, moim zdaniem, sposób pokazuje, że sprawy mają się dokładnie odwrotnie. Jeśli przyjmie się istnienie takich bytów, jak sądy, to faktycznie można podać rozstrzygnięcia pewnych istotnych kwestii, a na dodatek rozstrzygnięcia te charakteryzuje relatywna prostota i elegancja. Jak jednak twierdzi Lycan, sądy są niedopuszczalne ze względu na ich status metafizyczny i właśnie dlatego słusznie — zdaniem Lycana — odrzuca je Quine, który idzie o krok dalej i stara się stworzyć teorię, która będzie w stanie poradzić sobie z tymi samymi zagadnieniami bez angażowania przedmiotów intensionalnych.

¹⁹ Szerokie omówienie kwestii nośników prawdziwości, znacznie wybiegające poza cele i tematykę niniejszej pracy, znaleźć można w Rojszczak (2005).

o konkretnych egzemplarzach tego zdania²⁰ — np. egzemplarz, w którym „obecny premier Polski” odnosi się do Donalda Tuska (tj. egzemplarz użyty np. w roku 2012), jest prawdziwy, natomiast egzemplarz, w którym owa deskrypcja oznacza Tadeusza Mazowieckiego (tj. egzemplarz użyty np. w maju 1990 roku), jest fałszem. Ocena wartości logicznej egzemplarzy jest możliwa, ponieważ w ich wypadku wiadomo, kto jest odniesieniem owej deskrypcji. Strawson uważa, że dopóki zdanie-typ nie zostanie umieszczone w jakimś kontekście, co wiąże się często z ustaleniem odniesienia jego wyrażenń składowych, niemożliwa jest ocena prawdziwości. Strawson posuwa się w swoich tezach jeszcze dalej i głosi, że nonsensem jest mówienie o wartości logicznej jako cesze przysługującej zdaniu-typowi będącemu klasą zdań-egzemplarzy w wypadku *wszelkich* zdań, a nie tylko tych, które trzeba rozpatrywać ze względu na jakiś kontekst, by określić, co oznaczają poszczególne wyrażenia tworzące zdanie. Podobnym nonsensem byłoby przypisywanie własności posiadania wszystkich boków równych zbiorowi zawierającemu wszystkie kwadraty i prostokąty. O wszystkich elementach takiego zbioru można prawdziwie lub fałszywie orzekać posiadanie tej cechy, jednak ocena, czy posiada ją również ten zbiór jest zadaniem absurdalnym.

Nieco zabawnie, ale przy tym bardzo trafnie, stanowisko Strawsona streszcza Michael Dummett (2000), który każe nam wyobrazić sobie, że bierzemy udział w lekcji języka obcego, na przykład niemieckiego, podczas której nauczyciel daje nam za zadanie przetłumaczyć na niemiecki zdanie „Ten facet właśnie zbił wszystkie kieliszki do wina”. Zdanie to w takiej sytuacji pojęte jest jako typ i jak stwierdza Dummett:

Nonsensem byłoby pytać nauczyciela, o którego faceta lub o które kieliszki chodzi: żadne odniesienie się do określonego człowieka, kieliszków czy zdarzenia nie miało miejsca ani nie było nawet zamierzone; w związku z tym równie nonsensowne byłoby pytanie, czy to zdanie jest prawdziwe czy fałszywe (Dummett 2000: 264).

William Alston (2001: 43–4) podaje inny przykład przemawiający za rezygnacją z uznawania zdań-typów za nośniki wartości logicznej. Wyobraźmy sobie sytuację, w której dwie osoby odczytują wynik pomiaru z mierników — każda ze swojego. Obie osoby wygłaszają zdanie „Wskazówka jest na siódemce”; obaj mierzący posłużyli się zatem tym samym zdaniem-typem. Jednakowoż jedna z osób uległa złudzeniu optycznemu — na jej mierniku wskazówka nie była na siódemce, lecz w innym miejscu. Wówczas okazuje się, że to samo zdanie-typ ma inną wartość logiczną w zależności od tego, która z osób go używa²¹.

²⁰ Pomijam tutaj komplikację w postaci Strawsonowskiego rozróżnienia na użycia i wypowiedzenia, ponieważ nie zaburza ono argumentacji przeciwko przypisywaniu wartości logicznej zdaniom-typom.

²¹ Alston podaje również humorystyczny przykład, w którym ktoś chce ocenić wartość logiczną wypowiedzianego przeze mnie zdania-typu „Ja jestem głodny” i w razie, gdyby okazało

Podając rację za tym, iż utożsamienie sądu ze zdaniem-typem jest błędem, Richard Kirkham (1992: 57) wskazuje na pewną kwestię, która, jak sądzę, jest istotna również w kontekście poszukiwania pierwotnych nośników prawdziwości. Kirkham konstatuje, iż zasadniczą różnicą pomiędzy sądem a zdaniem-typem jest to, że sąd jako byt abstrakcyjny istnieje niezależnie od tego, czy ktoś ten sąd wyraził, natomiast zdanie-typ to ni mniej, ni więcej niż zbiór wszystkich równokształtnych i równobrzmiących zdań-egzemplarzy, w związku z czym, jeśli jakiegoś zdania nikt nigdy nie wygłosił (innymi słowy, nie zaistniał żaden jego egzemplarz), owo zdanie pojęte jako typ nie istnieje.

Przypuśćmy, że wygłaszam zdanie „Filip Kawczyński ma serce we właściwym miejscu”. Założmy dodatkowo dwie rzeczy, które najprawdopodobniej są zgodne z rzeczywistością: po pierwsze, że prawdą jest, iż moje serce jest prawidłowo ulokowane w klatce piersiowej; oraz po drugie, że moje użycie jest pierwszym użyciem takiego zdania w historii — dotąd jeszcze nikt nie wygłosił takiego komunikatu. W takiej sytuacji należy powiedzieć, iż przed moim użyciem nie istniało zdanie-typ o takim kształcie, ponieważ nie istniał żaden jego egzemplarz²². Skoro zaś zdanie-typ nie istniało, nie było ono dotąd ani prawdziwe, ani fałszywe. Jednocześnie mamy intuicję, by uznać, iż to, że Filip Kawczyński ma serce we właściwym miejscu, było prawdą również przed moim użyciem, tj. zanim istniało zdanie-typ „Filip Kawczyński ma serce we właściwym miejscu”. Dlatego też lepszym rozwiązaniem jest uznanie, iż nośnikiem wartości logicznej (w tym wypadku — prawdy) nie jest zdanie-typ, które zaistniało dopiero wraz z moim użyciem, lecz raczej sąd [że Filip Kawczyński ma serce we właściwym miejscu], który — jako przedmiot abstrakcyjny — istnieje i posiada wartość logiczną, nawet jeśli nikt go nie wyraził, wobec czego był on prawdziwy również, na przykład, pięć lat temu, czyli dokładnie pięć lat przed moim użyciem i zaistnieniem zdania-typu.

Po zdyskwalifikowaniu zdań-typów jako kandydatów na nośniki prawdziwości do rozpatrzenia pozostaje nam możliwość uznania za nie egzemplarzy zdań.

Mówiąc bardzo ogólnie, zdania-typy nie nadają się na nośniki prawdziwości, ponieważ w wypadku zdania pojętego jako typ często nie wiadomo do czego

się prawdziwe, przygotować dla mnie posiłek. Może się zdarzyć, iż zdanie to może z prawdziwego stać się fałszywym zanim ta osoba zdąży przygotować jedzenie (ponieważ w międzyczasie zaspokoje swój głód w inny sposób). Wynika stąd, że jeśli za wszelką cenę chcemy uznać, iż podstawowymi nośnikami prawdy są zdania-typy, musimy liczyć się z tą konsekwencją, że w wypadku wielu zdań, ich wartość logiczna będzie zmieniać się w czasie, a niekiedy — jak w powyższym przykładzie — w zawrotnym wręcz tempie.

²² Ktoś mógłby replikować, że takie zdanie-typ istniało przed moim użyciem i było po prostu zbiorem pustym. Aby jednak określić jakikolwiek zbiór, trzeba podać warunek, który musi spełniać przedmiot, by należeć do tego zbioru, podczas gdy sformułowanie tego warunku jako: do zbioru Z należy każdy element równokształtny/równobrzmiący z „Filip Kawczyński ma serce we właściwym miejscu”, nieuchronnie wiąże się z użyciem zdania o tym kształcie, przez co zbiór rzekomo pusty okazuje się mieć przynajmniej jeden element.

odnoszą się niektóre wyrażenia wchodzące w jego skład. Najpowszechniejszy przykład takiej sytuacji to występowanie w zdaniu wyrażen okazjonalnych, których wartości semantyczne można określić pod warunkiem, że weźmie się pod uwagę kontekst, w którym zostały użyte. Jako że zdanie-egzemplarz to napis lub dźwięk pojawiający się w pewnym kontekście, w którym odniesienie wszystkich jego elementów jest ustalone, wspomniany problem nie dotyczy egzemplarzy.

Uznanie egzemplarzy za nośniki prawdziwości nie wymyka się jednak pewnym kłopotom, którymi obarczone było obsadzenie w tej roli typów. Założyliśmy, iż przed moim użyciem wyrażenia „Filip Kawczyński ma serce we właściwym miejscu” nie zaistniał żaden egzemplarz tego zdania. Przeto przed tym użyciem nie można było o egzemplarzu tego zdania orzec prawdziwości, co pozostaje w konflikcie z intuicją, która każe nam uznać, iż to, że mam serce we właściwym miejscu, było prawdą znacznie wcześniej niż nastąpiło moje użycie. Dummett (2000: 264) ujmuje to w ten sposób, że nie należy zawężać dziedziny prawdziwości do wypowiedzi faktycznie wygłoszonych, a do tego właśnie prowadzi potraktowanie zdań-egzemplarzy jako pierwotnych nośników prawdy.

Alston (2001: 44) zauważa, że w pierwszej chwili trudno w ogóle zrozumieć, co miałyby znaczyć, że egzemplarz, który jest po prostu ciągiem dźwięków lub serią kresek i okręgów na papierze, miałyby mieć własność bycia prawdziwym lub bycia fałszywym. Dodaje jednak, iż możliwe jest wprowadzenie następującej konwencji, która pozwala to zrozumieć: egzemplarz zdania jest zawsze użyty do przekazania jakiejś treści; jeśli ta treść jest prawdziwa, wówczas prawdziwość możemy orzec również o egzemplarzu (i analogicznie z fałszywością)²³. Jak sugeruje Alston, takie rozwiązanie prowadzi wprost do uznania sądów za pierwotne nośniki prawdziwości. Jasno bowiem widać, że egzemplarze nie są pierwotnymi nośnikami wartości logicznej w tym sensie, że dziedziczą ją po odpowiedniej treści. Owa treść natomiast to nic innego, jak sąd.

6. Dodatkowe funkcje przypisywane sądom

Konsekwencją uznania sądów za nośniki wartości logicznej jest przypisanie im pełnienia co najmniej dwóch dodatkowych funkcji. Pierwsza z nich to bycie

²³ Alston stwierdza również, iż podobnie ma się sprawa z typami — jeśli upieramy się przy tym, by przypisywać wartość logiczną typom, to okazuje się, iż, aby dokonać takiego przypisania, musimy najpierw rozpoznać treść przekazywaną przez ten typ (czyli *de facto* rozpoznać treść przekazywaną przez egzemplarz lub egzemplarze tego typu umieszczone w kontekście i posiadające ustalone wartości semantyczne wszystkich wyrażen składowych; jeśli okaże się, że wszystkie egzemplarze we wszystkich kontekstach mają tę samą treść, to można przypisać ją również typowi), ocenić wartość logiczną owej treści i dopiero kierując się tą oceną, możemy orzec wartość logiczną typu (która będzie w każdym wypadku, rzecz jasna, identyczna z wartością logiczną odpowiedniej treści).

nośnikiem własności modalnych (por. King 2007: 2, Kripke 1972, Stanley 2002). Mówiąc najkrócej, sądy są tym, co stanowi właściwe argumenty operatorów konieczności oraz możliwości. Gdy mówimy, że z konieczności dwa plus dwa jest cztery lub że było możliwe, by Napoleon wygrał pod Waterloo, w gruncie rzeczy stwierdzamy, iż sąd [że dwa plus dwa jest cztery] jest konieczny (tj. jest prawdziwy we wszystkich okolicznościach ewaluacji) oraz że sąd [że Napoleon wygrał pod Waterloo] jest możliwy (tj. nie jest tak, iż jest on fałszywy we wszystkich okolicznościach ewaluacji).

Druga z funkcji sądów wynikających z bycia pierwotnym nośnikiem wartości logicznej to bycie (pierwotnym) argumentem relacji logicznych, takich jak wykluczenie, wykluczanie, sprzeczność itd. (por. Loux 1998: 139). Powiedzenie, że jeśli Napoleon urodził się na Korsyce, to Napoleon urodził się na wyspie, jest zatem w istocie stwierdzeniem, iż sąd [że Napoleon urodził się na Korsyce] implikuje sąd [że Napoleon urodził się na wyspie²⁴].

Opisane powyżej funkcje, których pełnienie przypisuje się sądom, bez wątpienia świadczą o dużej wartości teoretycznej pojęcia sądu. Jak widzieliśmy, włączenie tego pojęcia do teorii pozwala na jej gruncie podać spójne i — by posłużyć się wyrażeniem, którego logicy używają niekiedy w odniesieniu do dowodów — *eleganckie* wyjaśnienia pewnych fundamentalnych dla filozofii języka kwestii. Do powyższej listy problemów rozwiązywanych przez akceptację sądów Lycan (1974: 431–2) dodaje dwie inne kwestie — również istotne, choć mniej fundamentalne niż te wymienione powyżej. Po pierwsze, odwołanie się do pojęcia sądu stanowi dobry punkt wyjścia do konstruowania teorii prawdy. Po drugie, koncepcja angażująca sądy może dobrze zdawać sprawę z różnicy pomiędzy fałszywością z konieczności a brakiem znaczenia (czyli niewyrażaniem jakiegokolwiek sądu).

* * *

Sądzę, iż bez popadania w przesadę można powiedzieć, że pojęcie sądu jest niezwykle użytecznym narzędziem, gdy chodzi o eksplikację wielu zasadniczych fenomenów językowych. Jednocześnie, jak wspomniałem, wielu filozofów zrezygnowało z korzyści płynących z przyznania sądom racji bytu ze względu na problemy związane z naturą sądów oraz ich statusem ontycznym. Po przeszło stu latach posługiwania się pojęciem sądu te dwie ostatnie kwestie wciąż stanowią najważniejsze i najtrudniejsze wyzwanie stojące przed reali-

²⁴ Nie zamierzam upierać się przy tym, iż w wypadku sądu bycie nośnikiem własności modalnych oraz bycie argumentem relacji logicznych wynika z bycia podstawowym nośnikiem wartości logicznej; można uważać, iż kierunek tej zależności jest przeciwny. W świetle niniejszych rozważań rozstrzygnięcie tej kwestii nie jest istotne. Jakkolwiek, ponieważ zarówno pojęcia modalne, jak i stałe logiczne oznaczające relacje logiczne definiuje się poprzez odwołanie się do pojęcia prawdziwości, stanowisko, wedle którego zależność pomiędzy branymi tutaj pod uwagę funkcjami sądu przyjmuje taką postać, jak przed chwilą wspomniana, wydaje się słuszne.

stami. Przez ten czas filozofia języka, językoznawstwo, filozofia umysłu czy kognitywistyka wypracowały szereg nowych narzędzi teoretycznych — jak się jednak okazuje, sądy logiczne w pewnych sprawach okazują się niezastąpione. Jako zachętę do dalszej batalii o sądy warto więc może spytać o to, czy nie nadszedł czas, by uzasadnienia dla sądów szukać na innej płaszczyźnie niż *tradycyjna* filozofia języka, logika czy ontologia. Być może z pomocą mogą tu przyjść dziedziny dalece bardziej uwikłane w empirię, takie jak kognitywistyka czy współczesna psychologia (której nie należy mylić z *psychologizmem* z przełomu XIX i XX wieku). Mam nadzieję, że niniejszy tekst, w którym starałem się pokazać, jak wielkie korzyści płyną z akceptacji sądów logicznych, przekonał czytelnika, iż warto się w ową batalię zaangażować.

Bibliografia

- AJDUKIEWICZ K. ([1934] 1985), *Język i znaczenie*, [w:] *Język i poznanie*, t. I, Warszawa: PWN, 145–174.
- ALSTON W. (2001), *A Realist Conception of Truth*, [w:] Lynch P. (red.), *The Nature of Truth. Classic and Contemporary Perspectives*. Massachusetts: MIT Press, 41–66.
- BOLZANO B. ([1837] 2010), *Wissenschaftslehre*, [w:] tenże, *Podstawy logiki. Wybrane fragmenty I i II tomu „Teorii nauki” z uzupełniającymi streszczeniami F. Kambartela* (przeł. E. Drzazgowska). Kęty: Wydawnictwo Marek Derewiecki.
- CARNAP R. ([1947] 2007), *Znaczenie i konieczność*, [w:] tenże, *Pisma semantyczne* (przeł. T. Ciecierski, M. Poręba, M. Sala, B. Stanosz). Warszawa: Aletheia. 201–465.
- CIECIERSKI T. (2003), *O pojęciu sądu logicznego*, „Przegląd Filozoficzny” 48, 125–144.
- DAVIDSON D. ([1982] 1992), *Zwierzęta racjonalne*, [w:] tenże, *Eseje o prawdzie, języku i umyśle* (przeł. C. Cieśliński). Warszawa: PWN, 234–250.
- DUMMETT M. (2000), *Of What Kind of Thing is Truth a Property?*, [w:] Blackburn S. i Simmons K. (red.), *Truth*. New York: Oxford University Press, 264–281.
- FREGE G. ([1891] 1977), *Funkcja i pojęcie*, [w:] tenże, *Pisma semantyczne* (przeł. B. Wolniewicz). Warszawa: PWN, 18–44.
- FREGE G. ([1892] 1977), *Sens i znaczenie*, [w:] tenże, *Pisma semantyczne* (przeł. B. Wolniewicz). Warszawa: PWN, 60–88.
- FREGE G. ([1918] 1977), *Myśl — studium logiczne*, [w:] tenże, *Pisma semantyczne* (przeł. B. Wolniewicz). Warszawa: PWN, 101–129.
- FREGE G. ([1919] 1977), *Szkic dla Darmseaedtera*, [w:] tenże, *Pisma semantyczne* (przeł. B. Wolniewicz). Warszawa: PWN, 133–139.
- GUT A. (2005), *Gottlob Frege i problemy filozofii współczesnej*. Lublin: KUL.
- HARMAN G. (1967), *Quine on Meaning and Existence, I. The Death of Meaning*, „The Review of Metaphysics” 21, 124–151.
- HORWICH P. ([1990] 2005), *Truth*. Oxford University Press.
- JADACKI J. ([2001] 2002), *Spór o granice języka. Elementy semiotyki logicznej i metodologii*. Warszawa: Semper.

- KAPLAN D. ([1977] 1989), *Demonstratives*, [w:] Almog J., Perry J., Wettstein H. (red.), *Themes from Kaplan*. New York: Oxford University Press, 481–563.
- KAPLAN D. (1978), *Dthat*, [w:] Cole P., Morgan J. (red.), *Syntax and Semantics. Vol. 9*. New York: Academic Press, 221–243.
- KAPLAN D. (1989), *Afterthoughts*, [w:] Almog J., Perry J., Wettstein H. (red.), *Themes from Kaplan*. New York: Oxford University Press, 567–614.
- KING J. (2007), *The Nature and Structure of Content*. Oxford University Press.
- KIRKHAM R. ([1992] 2001), *Theories of Truth. A Critical Introduction*. Massachusetts: MIT Press.
- KRIPKE S. ([1972] 2001), *Nazywanie a konieczność* (przeł. B. Chwedeńczuk). Warszawa: Aletheia.
- LEWIS D. (1970), *General Semantics*, "Synthese" 22, 18–67.
- LOUX M. ([1998] 2003), *Metaphysics. A Contemporary Introduction*. New York: Routledge.
- LYCAN W. (1974), *Can Propositions Explain Anything?*, "Canadian Journal of Philosophy" 3, 427–434.
- LYCAN W. ([2000] 2002), *Philosophy of Language. A Contemporary Introduction*. New York: Routledge.
- MCGRATH M. (2007), *Propositions*, [w:] Zalta E. (red.), *The Stanford Encyclopedia of Philosophy*, <http://plato.stanford.edu/entries/propositions>.
- MILL J. S. ([1843] 1962), *System logiki dedukcyjnej i indukcyjnej* (przeł. Cz. Znamierowski). Kraków: PWN.
- MOORE G. E. (1899), *The Nature of Judgment*, "Mind" 8, 176–193.
- MOORE G. E. ([1953] 1967), *O głównych zagadnieniach filozofii* (przeł. Cz. Znamierowski). Warszawa: PWN.
- PRIOR A. (1971), *Objects of Thoughts*. Oxford University Press.
- QUINE W. V. O. ([1948] 2000), *O tym, co istnieje*, [w:] tenże, *Z punktu widzenia logiki* (przeł. B. Stanosz). Warszawa: Aletheia, 29–48.
- QUINE W. V. O. ([1951] 2000), *Dwa dogmaty empiryzmu*, [w:] tenże, *Z punktu widzenia logiki* (przeł. B. Stanosz). Warszawa: Aletheia, 49–76.
- QUINE W. V. O. ([1960] 1999), *Słowo i przedmiot* (przeł. C. Cieśliński). Warszawa: Aletheia.
- QUINE W. V. O. ([1970] 1977), *Filozofia logiki* (przeł. H. Mortimer). Warszawa: PWN.
- RICHARD M. ([1990] 2004), *Propositional Attitudes. An Essay on Thoughts and How We Ascribe Them*. Cambridge: Cambridge University Press.
- ROJSZCZAK A. (2005), *From the Act of Judging to the Sentence. Synthese Library 328*. Springer.
- RUSSELL B. ([1903] 2008), *The Principles of Mathematics*. Merchant Books.
- RUSSELL B. ([1912] 2004), *Problemy filozofii* (przeł. W. Sady). Warszawa: PWN.
- SALMON N. (1986), *Frege's Puzzle*. Massachusetts: MIT Press.
- SALMON N. (1989), *Tense and Singular Propositions*, [w:] Almog J., Perry J., Wettstein H. (red.), *Themes from Kaplan*. New York: Oxford University Press, 331–391.
- SELLARS W. (1963), *Abstract Entities*, "The Review of Metaphysics" 16, 627–671.
- SOAMES S. (1985), *Lost Innocence*, "Linguistic and Philosophy" 8, 59–71.
- SOAMES S. ([1987] 2009), *Direct Reference, Propositional Attitudes, and Semantic Content*, [w:] tenże, *Philosophical Essays. Volume II*. New Jersey: Princeton University Press, 33–71.

- SOAMES S. (2010a), *Propositions*, [w:] *Witryna internetowa Scotta Soamesa*, http://www.bcf.usc.edu/~soames/forthcoming_papers/Propositions.pdf, (ukaze się w: Russell G. i Graff Fara D. (red.), *The Routledge Companion to Philosophy of Language*. Routledge).
- SOAMES S. (2010b), *What is Meaning?*. New Jersey: Princeton University Press.
- STALNAKER R. (1999), *Context and Content*. New York: Oxford University Press.
- STRAWSON P. ([1950] 1967), *O odnoszeniu się użyć wyrażen do przedmiotów*, [w:] Pelc J. (red.) *Logika i język. Studia z semiotyki logicznej*. Warszawa: PWN, 277–295.
- TAŁASIEWICZ M. (2006), *Filozofia składni*. Warszawa: Semper.

Rozdział 7

Niedosłowność i interpretacja

Janusz Maciaszek (Uniwersytet Łódzki)

Słowa kluczowe: metafora, metonimia, ironia, aluzja, teoria prawdy Tarskiego, interpretacja, znaczenie, przekonanie, implikowanie konwersacyjne

Cel artykułu jest dwojaki. Po pierwsze, uzasadniam w nim i wyjaśniam Donalda Davidsona niekognitywną¹ teorię metafory, odwołując się do jego teorii interpretacji oraz teorii działania. Po drugie, rozszerzam analizę metafory na pozostałe wyrażenia niedosłowne, zwane tradycyjnie figurami retorycznymi. W szczególności przeprowadzam analizę mechanizmu rozpoznawania zdań niedosłownych w oparciu o teorię implikowania konwersacyjnego Grice’a oraz teorię interpretacji i racjonalności Davidsona, a następnie proponuję kryterium rozpoznawania metafory w wypowiedzeniach zdań interpretowanych niedosłownie. Nawiązując do poglądów Davidsona, wprowadzam pojęcie komunikowania semantycznego oraz pojęcie komunikowania pozasemantycznego, które stanowi — w moim przekonaniu — jeden z istotnych przejawów działania metafory. Proponuję również nowe podejście do ironii, traktując ją — wbrew utartemu zwyczajowi — za wypowiedzenie dosłowne.

1. Wstęp

Na temat metafory napisano bardzo wiele i z bardzo różnych punktów widzenia. Wiele spośród prac poświęconych temu zagadnieniu ma — w moim przekonaniu — jeden podstawowy mankament. Aby dokonać intuicyjnego rozróżnienia między znaczeniem dosłownym i metaforycznym, autorzy posługują się potocznym pojęciem znaczenia, nie odwołując się do jakiegokolwiek ścisłej teorii tego pojęcia. Znaczenie w sensie przedteoretycznym może być definiowane na dwa zasadnicze sposoby. Jedna z możliwych definicji ma charakter psychologiczny i została podana przez Kazimierza Ajdukiewicza. Wedle tej definicji

¹ Termin ten pojawia się w: Reiner i Camp (2006) w odniesieniu do teorii Davidsona oraz Rorty’ego.

znaczenie terminu to jego „sposób rozumienia” (zob. Ajdukiewicz 1965). Z kolei Ludwig Wittgenstein proponował określić znaczenie jako sposób użycia wyrażenia w grze językowej (zob. Wittgenstein 1953). Mamy zatem dwie pozornie wykluczające się definicje sprawozdawcze tego samego terminu. Piszę tu o pozornym wykluczaniu się, gdyż faktycznie obie definicje zdają sprawę z sensu tego samego przedteoretycznego pojęcia, lecz czynią to z dwóch różnych punktów widzenia. Z punktu widzenia pierwszej osoby słuszna jest psychologiczna definicja Ajdukiewicza — znaczenie wyrażenia dla mówiącego to sposób, w jaki je rozumie. Z kolei punktowi widzenia trzeciej osoby odpowiada definicja Wittgensteina — znaczenia przypisywane słowom mówiącego odpowiadają sposobowi, w jaki ich używa.

Obie definicje nie są zbyt przydatne w przypadku analizy wypowiedzeń niedosłownych, gdyż prowadzą do trywialnych konkluzji. Ponieważ wypowiedzenia niedosłowne rozumiemy inaczej niż dosłowne, to przypisujemy im inne znaczenia. Z kolei fakt, że nasz rozmówca używa ich inaczej niż wypowiedzeń dosłownych — zmieniając jakby grę językową — prowadzi do tej samej konkluzji. Myślę, że wiele teorii metafory milcząco zakłada przedteoretyczne pojęcie znaczenia w jednej lub drugiej wersji, co prowadzi do konstatacji, że istnieją specjalne znaczenia metaforyczne. Z intuicyjnego punktu widzenia stwierdzenie to wydaje się oczywiste. Jednak pojęcie znaczenia metaforycznego podpada pod te same definicje sprawozdawcze co pojęcie znaczenia dosłownego, stawiając pod znakiem zapytania sensowność tego rozróżnienia. Ponadto podejście to nie wyjaśnia, dlaczego podawanie znaczeń metaforycznych, tj. parafrazowanie metafor, natrafia w wielu przypadkach na zasadnicze trudności.

Aby uniknąć wspomnianych trudności, właściwa teoria metafory i innych wyrażen niedosłownych powinna być oparta na ścisłej teorii znaczenia. W przeciwnym przypadku teoria — choćby niezwykle atrakcyjna i trafiająca do wyobraźni — nie jest w stanie wyjaśnić fenomenu niedosłowności w sposób satysfakcjonujący. Jak będę się starał pokazać, niekognitywna teoria metafory Davidsona, przy pewnych uzupełnieniach, spełnia wymieniony postulat metodologiczny, stanowiąc jednocześnie punkt wyjścia do analizy pozostałych figur retorycznych.

Przed przejściem do dalszych rozważań należy poczynić oczywiste zastrzeżenie. W artykule będę pisał o niedosłownych wypowiedzeniach wyrażen, nie zaś o samych wyrażeniach. O niedosłowności można bowiem mówić jedynie w przypadku okazów wyrażen, które są wypowiedzane, słyszane, pisane lub odczytywane w konkretnych okolicznościach. Na owe okoliczności składać się będą okoliczności zewnętrzne — a dokładniej wiedza uczestników rozmowy na temat okoliczności wewnętrznych — oraz przekonania i inne postawy propozycyjalne uczestników rozmowy — w szczególności zaś wiedza faktograficzna oraz przypuszczenia na temat przekonań innych uczestników rozmowy.

Jeśli jednak przytaczanie frazy „wypowiedzenie wyrażenia” będzie niezręczne, będę pisał po prostu o wyrażeniach. Pragnę również zwrócić uwagę na zamierzoną dwuznaczność tytułu. Termin „interpretacja” bywa używany przeze mnie w szerokim sensie potocznym, który nie wymaga wyjaśnienia, lecz częściej używam go w ścisłym sensie podanym przez Davidsona. Mam nadzieję, że w konkretnych przypadkach będzie jasne, o jaki sens tego terminu chodzi.

2. Przegląd tradycyjnych figur retorycznych

Aby zilustrować niedostatki podejścia do metafory i innych figur retorycznych bez oparcia na ścisłej teorii znaczenia, odwołajmy się do kilku definicji słownikowych i encyklopedycznych. *Słownik języka polskiego PWN* (2002) — dalej (SJP) — proponuje rozumieć metaforę jako „występowanie wyrazu w nowym znaczeniu łączącym się ze znaczeniem realnym w sposób obrazowy, plastyczny i rozumianym poprzez odniesienie do znaczenia realnego”. W definicji tej nie wspomina się *explicite* o kontekście, w którym pojawiło się wyrażenie, operuje się natomiast pojęciem znaczenia metaforycznego, które pozostaje w pewnym związku ze znaczeniem dosłownym.

Z kolei *Nowa encyklopedia powszechna PWN* z (1995) — dalej (NEP) — definiuje ją jako wyrażenie, w którym „słowa obce sobie znaczeniowo, pozostając w związkach składniowej zależności, uzyskują nowe znaczenie, zw. metaforycznym”. Definicja ta, podobnie jak poprzednia, odwołuje się do pojęcia znaczenia metaforycznego. W odróżnieniu od poprzedniej zwraca jednak uwagę na kontekst składniowy wyrażenia. Przytoczone w NEP jako przykład słowo „wstęga” w wyrażeniu „wstęga rzeki” prowokuje niewątpliwie wiele skojarzeń². Gdyby jednak należało podać znaczenie terminu „wstęga” we wskazanym kontekście, tj. sparafrazować ten termin, natrafilibyśmy na spore trudności. Trudności te mają różny stopień w poszczególnych przypadkach metafor. Bywają metafory twórcze (najczęściej literackie lub poetyckie), żywe i skostniałe (np. metafory zwierzęce). Metafory skostniałe jest stosunkowo łatwo parafrazować, metafory żywe opierają się niekiedy parafrazom (np. szekspirowskie „Julia jest słońcem”)³.

Do wyrażen metaforycznych zalicza się niekiedy oksymorony, hiperbole i litoty. Przez oksymoron rozumie się „metaforyczne zestawienie wyrazów o przeciwstawnym znaczeniu, np. gorzkie szczęście, wymowne milczenie” (SJP) lub „epitet sprzeczny, w którym zestawienie przeciwstawnych znaczeniowo składników (epitetu i wyrazu określanego) prowadzi do metaforycznego

² Każde wyrażenie użyte niedosłownie może być w tym samym kontekście składniowym użyte dosłownie, np. „wstęga” byłaby użyta dosłownie jedynie w nieprawdopodobnych okolicznościach, np. w bajce, w której rzeka posiadałaby jakąś wstęgę.

³ Jest to jednocześnie przykład na to, że niektóre metafory nie poddają się skostnieniu, chociaż funkcjonują od setek lat. Sam podział metafor na żywe i skostniałe jest nieostry, gdyż nie opiera się na jasnych kryteriach.

przekształcenia ich znaczeń, daje efekt paradoksu, np. czarna błyskawica” (NEP). Wydaje się, że nie zawsze można jednoznacznie odróżnić oksymoron od zwykłej metafory, np. określenie „mroźne piekło” na okolice podbiegunowe może, lecz nie musi, być interpretowane jako oksymoron. Przez piekło można bowiem rozumieć miejsce, w którym pobyt stanowi konsekwencję własnych czynów, nie zaś jedynie jako miejsce, w którym panuje ekstremalnie wysoka temperatura. W pierwszym przypadku byłaby to najprawdopodobniej „zwykła” metafora, w drugim zaś oksymoron. Nie sądzę jednak, że efekt zaskoczenia musi być w obydwu przypadkach odmienny. Ponadto wiele typowych metafor ma charakter semantycznych bezsensów (np. „zezowate szczęście”) i może wywołać ten sam efekt paradoksu i zaskoczenia, co stanowi argument na rzecz zaliczenia oksymoronów do wypowiedzi metaforycznych. Do oksymoronów, a dokładniej zdań oksymoronicznych, zaliczyłbym ponadto zdania będące sprzecznościami przy odczytaniu dosłownym, np. słynne „Jestem za, a nawet przeciw”, „Masz rację i jej nie masz”, „Każdy plus ma swoje minusy”.

Hiperbola jest to „zwrot stylistyczny polegający na zamierzonej przesadzie” (SJP). Hiperbole mogą być bardzo złożone, gdzie efekt przesady budowany jest przez nagromadzenie wyrażen metaforycznych i porównań — mogą być również prostymi wyrażeniami typu „morze łez”, „mordercza praca” lub „pot lał się strumieniami”. Niekiedy hiperbola może być użyta do osiągnięcia celu ironicznego, np. uwaga wygłoszona do studenta, który nic nie robi: „Grozi panu śmierć z przepracowania”. Przeciwnieństwem hiperboli jest litota. W większości przypadków litoty są bardzo skonwencjonalizowane, np. „wyjdę na sekundę” lub „jeszcze po kropelce”, i mogą być interpretowane dosłownie jedynie w bardzo nietypowych lub nieprawdopodobnych kontekstach.

Warto zauważyć, że przytaczane tu definicje figur retorycznych odwołują się zazwyczaj do typowych efektów ich działania, co może sugerować, że wypowiedzenia dosłowne nie wywołują takich efektów. Z pewnością nie może to być uznane za ogólną regułę — wyobraźmy sobie choćby komiczny efekt wywołany przez użycie zakraplacza przy wypowiedzeniu ostatniego z przytoczonych przykładów, który stałby się w tym kontekście wypowiedzeniem dosłownym.

Przez metonimię rozumie się „figurę stylistyczną polegającą na oznaczaniu przedmiotu za pomocą nazwy oznaczającej inny przedmiot związany z poprzednim” (SJP). NEP proponuje bardzo podobne określenie, podając jednocześnie dwa cele, dla których stosuje się metonimię. Byłoby to uzyskanie skrótu (np. „czytać Mickiewicza” zamiast „czytać dzieła Mickiewicza”) lub uzyskanie efektu ekspresywnego („dobre pióro” zamiast „dobry pisarz”). Metonimia, nawet jeśli rozumiemy je jedynie jako skróty myślowe, są z reguły bardzo skonwencjonalizowane i zazwyczaj używane jedynie w pewnych zestawieniach słów. O pisarzach nie mówi się zazwyczaj „pióra” bez kwalifikacji „dobry”

lub „najlepszy”; o marynarzach nie mówi się „ręce” poza utartymi zwrotami, np. „wszystkie ręce na pokład”. W przypadku metonimii nie mówi się zazwyczaj o szczególnym znaczeniu metonimicznym, chociaż metonimie jako skróty myślowe zazwyczaj można bardzo łatwo parafrazować. W niektórych przypadkach metonimie maskują brak dokładnej informacji lub pozwalają pomijać zbędne szczegóły, np. „Moskwa mówi nie” — w takim przypadku parafraza ta może być długa i niezgrabna⁴. Pomimo charakterystycznego dla metonimii skonwencjonalizowania wydaje się, że dla każdego z wymienionych przykładów można sobie — być może z trudem — wyobrazić sytuację, w której metonimia zostanie zinterpretowana jak metafora⁵.

Aluzja jest to „wypowiedź formułowana z intencją przywołania skojarzeń niezbędnych do zrozumienia wyrażanej treści” (NWP). Jako przykład podane jest wyrażenie „walka z wiatrakami” traktowane jako aluzja do dzieła Cervantesa. Jeżeli jednak nie traktujemy „rozumienia” aluzji jako testu na odczytanie słuchacza, to definicja ta wydaje się chybiona. Przede wszystkim przez aluzję rozumie się zazwyczaj zdanie — „walczyć z wiatrakami” można zatem rozumieć jako wspólny trzon nieskończenie wielu aluzji, np. „Twoje starania to walka z wiatrakami”. Po drugie, „walczyć z wiatrakami” w tych zdaniach pełni funkcję typowej metafory. Samo w sobie nie stanowi to jeszcze zarzutu wobec przytoczonej definicji, gdyż nie jest wcale powiedziane, że aluzja nie może metafory zawierać. Po trzecie, trudno jest podać jasne kryteria różnicy „przywołania skojarzeń” i zwykłego odwołania się do kontekstu, który obejmuje również wiedzę uczestników rozmowy. Na przykład zdanie „Nie jestem maszyną” może wywołać cały szereg skojarzeń, które mogą być uznane za „aluzję do” sytuacji znanej rozmówcom. Typowe aluzje nie zawierają jednak metafor, np. jeżeli student wychodzący z egzaminu wypowiada zdanie „Jak dwója to dwója”, to odwołuje się do znanej słuchaczom sytuacji, w której przed chwilą się znalazł (egzamin), wiedzy słuchaczy na swój temat oraz na temat egzaminu i egzaminatora. W takim przypadku wyrażenia „dwója” nie potraktuje się jako użytego metaforycznie. Oczywiście nie jest łatwo sparafrazować, co chciał powiedzieć ów student — być może chciał powiedzieć, że się nie przejmuje, że się zgadza z egzaminatorem, że się nie podda. Możliwe jest również, że chciał po prostu dodać sobie odwagi lub zaprezentować się przed kolegami jako „beztroski luzak”. Prowadzi to do wniosku, że „przywołanie skojarzeń”, na które powołuje się podana definicja aluzji, jest cechą oczywistą, a przez to nieciekawą⁶.

Ironia jest to „ukryta drwina, utajone szyderstwo, złośliwość zawarta w wypowiedzi pozornie aprobującej” (SJP). W definicji tej następuje odwołanie do dwóch kryteriów — typowych efektów działania ironii i do przeciwstawienia

⁴ „Moskwa” oznacza tu bliżej niesprecyzowany ośrodek władzy lub jego przedstawiciela.

⁵ Koncepcja, którą przedstawiam w tym esej, uwzględnia takie przypadki mieszane.

⁶ Zdanie „Walka z biurokracją jest walką z wiatrakami” stanowi w równym stopniu aluzję do Cervantesa, jak i do powszechnie znanej sytuacji panującej w urzędach.

przekonań faktycznie żywionych przez mówiącego przekonaniom wyrażonym w wypowiedzeniu. Typowe przypadki ironii spełniają oba warunki — ironia zazwyczaj bywa złośliwa, zaś charakterystyczny ton ironiczny stanowi wskaźnik „nie wierzcie w to, co mówię”. Można podawać jednak przykłady, które nasuwają wątpliwości, np. życzliwa ironia i autoironia. Nie zawsze mamy również do czynienia z przeciwstawieniem treści lub przekonań, np. gdy nauczyciel zwraca się ironicznym tonem do bardzo zdolnego ucznia (np. wybitnego laureata olimpiady przedmiotowej), który nie uważa na zajęciach „Ty już wszystko umiesz”. W takim przypadku nie da się przypisać mówiącemu przekonania, że uczeń jeszcze wszystkiego nie umie, gdyż byłby to oczywisty fałsz⁷.

Podsumowując ten przegląd, warto zauważyć, że tradycyjne figury retoryczne wyróżniane są według niejednorodnych kryteriów, np. semantycznych (odwołanie się do znaczenia dosłownego i metaforycznego), logicznych (pojęcie sprzeczności), kontekstu wypowiedzenia, wiedzy lub do spodziewanych efektów działania. W dalszej części artykułu będę się starał pokazać, że wątpliwości sprowokowane przez podane definicje można rozwiązać, podając jasne kryteria rozpoznawania wypowiedzeń niedosłownych.

3. Grice: rozpoznawanie niedosłowności i implikowanie konwersacyjne

W klasycznym artykule *Logika a konwersacja* (1975) H. P. Grice rozwinął teorię wypowiedzeń niedosłownych opartą na pojęciu konwersacyjnego implikowania treści. Grice wyróżnia dwa rodzaje znaczenia (treści): dosłowne lub konwencjonalne (*conventional meaning*) oraz niedosłowne lub mówiącego (*speaker's meaning*), które słuchacz może implikować konwersacyjnie z wypowiedzeń. Implikowanie konwersacyjne ma miejsce, gdy nastąpi pozorne złamanie jednej z tzw. maksym konwersacyjnych. Milczącym założeniem procedury implikowania konwersacyjnego jest przekonanie, z którym prawdopodobnie zgodziłby się również Grice, że mówiący jest istotą racjonalną, tj. działającą z pewnego powodu (*with a reason*) (zob. Davidson 1963)⁸. Implikowanie konwersacyjne można bowiem traktować jako podanie powodu, dla którego mówiący złamał dana maksymę.

Maksyma ilości domaga się, aby podawać tyle informacji, ile jest na danym etapie rozmowy potrzebne. Maksymę tę można złamać faktycznie lub pozornie. Faktyczne złamanie tej maksymy to przemilczanie ważnych informacji (ukrywanie czegoś) lub zalewanie słuchaczy strumieniem informacji zbędnych (gadatliwość). Szczególnym przypadkiem pozornego złamania tej maksymy są wypowiedzenia aluzyjne. W takim przypadku mówiący wypowiada zdanie,

⁷ Przypadek ten przeanalizuję dokładnie w paragrafie 6.

⁸ Pojęcie racjonalności u Davidsona nie jest jednoznaczne. W różnych pismach określa racjonalność jako działanie z powodu, posiadanie postaw propozycjonalnych, niesprzeczność przekonań. Sądzę, że wszystkie te określenia można sprowadzić do warunków koniecznych znajomości języka.

które wyraża przekonanie oczywiste dla wszystkich uczestników rozmowy. Ponieważ wypowiedzenie takie nie niesie żadnej nowej informacji, to słuchacze traktują złamanie maksymy jako pozorne i implikują konwersacyjnie, co mówiący chciał powiedzieć. W procedurze tej bierze się pod uwagę cały kontekst rozmowy oraz wiedzę i przypuszczenia o przekonaniach jej uczestników. Słyszając deputowanego, który zwraca się do ministra finansów „Panie ministrze, dwa plus dwa równa się cztery” implikujemy konwersacyjnie, co ów deputowany faktycznie chciał powiedzieć, gdyż nikt nie przypuszcza, aby chciał poinformować ministra o zachodzeniu tej prostej równości. Co mówiący faktycznie chciał powiedzieć, zależy od całego kontekstu, na który składa się temat i przebieg dyskusji, dyskusje, które miały miejsce w przeszłości, wiedza na temat sytuacji gospodarczej, przypuszczenia na temat wzajemnej wiedzy na ten temat itd.

Teoria Grice’a nie opisuje wszystkich efektów, które mogą zostać wywołane przez wypowiedzenie niedosłowne, ograniczając się do komunikowania niedosłownych treści. Pozatreściowe efekty bywają osiągane niekiedy poprzez uświadomienie implikowanej treści, np. uświadomienie sobie treści implikowanej dla podanego przykładu może wywołać złość lub chęć rewanżu. Często jednak można osiągnąć ten efekt bez wyraźnego komunikowania owych treści. Wyobraźmy sobie, że jeden z uczestników transakcji wypowiada zdanie „jak handel to handel”. Z punktu widzenia teorii Grice’a jest to aluzja, gdyż w sensie dosłownym zdanie podpada pod tautologiczny schemat $p \Rightarrow p$ lub $A = A$. Co jednak można implikować konwersacyjnie z tego zdania? Wydaje się, że ważniejszy tu jest pozakomunikacyjny efekt jego wypowiedzenia, np. wzbudzenie atmosfery zaufania lub zniwelowanie napięcia.

Maksyma jakości głosi, że należy mówić szczerze i jedynie to, do czego mamy podstawy. Kłamstwo oraz rozpowszechnianie bezpodstawnych plotek stanowi faktyczne złamanie tej maksymy. Pozorne łamanie maksymy jakości prowadzi do uznania wypowiedzenia za metaforyczne lub ironiczne, przy czym wypowiedzenia metaforyczne obejmowałyby oksymorony, hiperbole i litoty. Wypowiedzenie jest traktowane w danym kontekście jako łamiące pozornie maksymę jakości, jeśli wyraża przekonanie, którego nie żywi nikt z uczestników konwersacji i wszyscy o sobie wzajemnie przypuszczają, że nikt go nie żywi. Tak jest w przypadku ewidentnych lub powszechnie znanych fałszów, np. „Julia jest słońcem” oraz sprzeczności — „Jestem za, a nawet przeciw”. W takim przypadku słuchacz implikuje konwersacyjnie, co mówiący chciał faktycznie powiedzieć. Warto zauważyć, że maksyma jakości nie może być jedynym wyznacznikiem metafory. O ile wypowiedzenie „Jesteś trybem w maszynie” łamie maksymę jakości i powinno być traktowane jako wypowiedzenie metaforyczne, o tyle jego negacja — „Nie jestem trybem w maszynie” — nie tracąc metaforycznego charakteru, staje się, zgodnie z kryterium Grice’a, aluzją.

W przypadku ironii sprawa wygląda nieco inaczej, gdyż zazwyczaj nie dotyczy ona powszechnie uznawanych przekonań, lecz specyficznych ocen i przekonań mówiącego. Powiedzenie o znanym polityku „On jest wybitnym mężem stanu” może być zgodne z przekonaniami słuchacza lub nie, ale o uznaniu go za ironię decydują przede wszystkim przypuszczenia na temat przekonań mówiącego. Od szczerego wyrażenia przekonania oraz od kłamstwa ironia różni się tym, że mówiący nie żywi przekonania, które wyraził, a w przeciwieństwie do kłamstwa nie stara się tego faktu ukryć. Charakterystyczny ton ironiczny stanowi wskaźnik „nie wiercie w to, co mówię”. Wskaźnik ten bywa niepotrzebny, jeśli przekonania mówiącego są tak dobrze znane słuchaczom, że nie zinterpretują wypowiedzenia ani jako szczerego wyrazu przekonania, ani jako kłamstwa. Bywa również, że ironia wyraża oczywisty fałsz. Wyobraźmy sobie, że zdenerwowany niewiedzą ucznia nauczyciel geografii powie: „Tak, tak, a Zakopane leży nad Bałtykiem”. O zakwalifikowaniu tego wypowiedzenia do ironii decyduje przede wszystkim ironiczny ton lub sam kontekst wypowiedzenia. Sprawa jest jednak nieco bardziej złożona. Ironiczny ton może bowiem — jak w przykładzie gadatliwego ucznia, który podałem w poprzednim paragrafie — towarzyszyć zdaniu w sposób oczywisty prawdziwemu, zgodnemu z przekonaniami zarówno mówiącego, jak i słuchacza. Może to sugerować, że teoria implikowania konwersacyjnego nie wystarcza do adekwatnego opisu ironii.

Teoria Grice’a miała osiągnąć dwa cele. Pierwszy z nich to podanie kryterium rozpoznawania figur retorycznych, drugi to komunikowanie treści niedosłownych (*speaker's meaning*) — na tym właśnie skupił się Searle (zob. Searle 1979) oraz (Martinich 1984). Niestety sam Grice (ani potem Searle) nie zdefiniował znaczenia, lecz podał warunki, kiedy wyrażenie językowe jest znaczące, (zob. Grice 1957). Znaczenie dosłowne było dla niego znaczeniem konwencjonalnym, zaś implikowane konwersacyjnie znaczenie niedosłowne było niekonwencjonalne. Można jednak w tym miejscu zgłosić pewne zastrzeżenia. Po pierwsze, wypowiedzenia niedosłowne charakteryzują się niekiedy dużym stopniem skonwencjonalizowania implikowanych treści, co pozwala na ich łatwe parafrazowanie, np. metafory zwierzęce. Po drugie, jak w przypadku ironii, treść wypowiedzenia stanowi negację treści dosłownej, a zatem trudno jest mówić o braku konwencji. A zatem stwierdzenie, że znaczenie dosłowne jest konwencjonalne, zaś niedosłowne jest niekonwencjonalne jest wyraźnym uproszczeniem. Przeciwno przypisywaniu nadmiernej roli konwencji językowym wystąpił Davidson, który traktował konwencje jedynie jako środki, które ułatwiają komunikację językową, nie stanowiąc o samej jej możliwości.

4. Teoria Davidsona

Przed przedstawieniem teorii metafory Davidsona należy poświęcić nieco uwagi jego teorii znaczenia, gdyż — jak już wielokrotnie to podkreślałem — dopiero z tej perspektywy teoria ta staje się jasna⁹. Teoria znaczenia stanowi dla Davidsona część teorii interpretacji. Interpretowanie słów drugiego człowieka (a nawet własnych słów interpretatora wypowiedzianych lub zapisanych w przeszłości) polega na równoczesnym przypisaniu znaczeń słowom mówiącego, a jemu samemu przekonań. Znaczenia są przypisywane przez teorię prawdy Tarskiego (zob. Tarski 1933), która jest konstruowana i weryfikowana podczas rozmowy. Owo przypisywanie znaczeń, które jest faktycznie przekładem wyrażen języka mówiącego na język interpretatora, realizowane jest przez ogół reguł semantycznych teorii Tarskiego¹⁰. Prowadzi to do holizmu znaczeniowego, który przejawia się tym, że nie można podać znaczenia tylko jednego wyrażenia, nie ekstrapolując teorii prawdy na pozostałe wyrażenia idiolektu mówiącego. Należy również zauważyć, że w przeciwieństwie do wielu współczesnych teorii, które traktują kontekst jako wyznacznik pragmatyki, semantyka Davidsona również uwzględnia kontekst wypowiedzi. Jest to — jak sądzę — jedna z konsekwencji wspomnianego wcześniej antykonwencjonalizmu. Przez pragmatykę Davidson rozumie natomiast badanie konsekwencji wypowiedzi. Rozróżnienie to ma dla przedstawionej w tym eseju koncepcji dwojakie znaczenie. Po pierwsze, komunikacja językowa może być realizowana zarówno na poziomie semantycznym (przypisywanie przekonań w procesie interpretacji), jak i na poziomie pragmatycznym (pozasemantycznym). Komunikacyjna funkcja wypowiedzi niedosłownych mieści się właśnie na tym poziomie. Po drugie, kontekst wypowiedzi odgrywa istotną rolę zarówno na poziomie semantycznym, jak i pragmatycznym.

Dwuaspektowa procedura interpretacji nie jest zdeterminowana, tj. możliwych jest wiele sposobów przekładu oraz przypisania przekonań, chociaż najczęściej decydujemy się na ten najbardziej prawdopodobny. Ponieważ dysponujemy zawsze skończoną próbką wypowiedzi rozmówcy, teoria prawdy ma charakter ekstrapolacji owej próbki na cały język, przy czym teoria ta może być weryfikowana i uzupełniana w miarę rozwoju konwersacji. Teoria prawdy może być stosowana zarówno do interpretowania wypowiedzi w innym języku, jak i do interpretowania słów rozmówcy posługującego się tym samym

⁹ Jest rzeczą zastanawiającą, że w swej teorii metafory Davidson nie odwołuje się *explicit* do własnej teorii znaczenia.

¹⁰ Teoria prawdy Tarskiego jest podatna na rozliczne interpretacje, również odmienne od prezentowanej w tym eseju. Do teorii prawdy Davidson odwoływał się bezpośrednio lub pośrednio w większości swych artykułów. Zagadnieniom teorii interpretacji i teorii prawdy Tarskiego poświęcony jest przede wszystkim tom *Inquiries into Truth and Interpretation*, Davidson (2001) oraz częściowo polski przekład wybranych artykułów (Davidson 1992). Z ostatnich opracowań tego zagadnienia można polecić monografię Ludwiga i Lepore'a (2005).

językiem co interpretator¹¹. Ograniczając się do sytuacji, w której idiolekty mówiącego i interpretatora należą do tego samego języka, interpretator zakłada zazwyczaj, że rozmówca przypisuje równokształtnym słowom te same znaczenia — może natomiast różnić się od niego niektórymi przekonaniem. Prowadzi to do wstępnego przyjęcia trywialnej teorii prawdy, której reguły semantyczne przyporządkowują słowom mówiącego ich równokształtne odpowiedniki w języku interpretatora. Jeśli jednak wstępna teoria obliuguje nas do przypisania mówiącemu nieprawdopodobnych lub absurdalnych przekonań, wówczas musi być zweryfikowana¹². W procesie weryfikacji odwołujemy się do tzw. T-zdań postaci „Zdanie Z jest prawdziwe wtedy i tylko wtedy, gdy *p*”, gdzie Z jest zdaniem idiolektu mówiącego, zaś *p* jego przekładem na idiolekt interpretatora. Oczywiście do weryfikacji teorii nie nadają się wszystkie zdania, lecz te, które wyrażają przekonania empiryczne lub pewne ogólnie akceptowane prawdy¹³. W przypadku tego typu przekonań niezgoda między mówiącym a interpretatorem skłania tego drugiego do zmiany teorii prawdy.

Aby to zilustrować, można się odwołać do autentycznej sytuacji, z którą zetknęło się prawdopodobnie wielu mieszkańców Polski centralnej podczas pobytu w Małopolsce. Będąc pierwszy raz w górach, wypowiedź „X wyszedł na pole” mieszkaniiec Polski centralnej przełoży na równobrzmiące zdanie swojego idiolektu (trywialna teoria prawdy). Gdy jednak zobaczy X-a przed domem, przypisze mówiącemu mylne przekonanie, że X jest na polu lub uzna, że miała miejsce próba wprowadzenia go w błąd. Kiedy sytuacja się powtórzy, interpretator zrozumie, że należy zmienić teorię prawdy i przypisać mówiącemu przekonanie, że X wyszedł na dwór, zaś wyrażenie „na pole” przetłumaczyć na „na dwór”. Z tak zweryfikowanej teorii prawdy będzie wynikało między innymi T-zdanie „Wypowiedzenie X wyszedł na pole jest prawdziwe wtedy i tylko wtedy, gdy X wyszedł na dwór”.

W procesie interpretacji przypisanie przekonań odbywa się na gruncie semantycznym¹⁴, gdyż uwikłane jest w proces budowania semantyki Tarskiego, zaś znaczenia przypisane słowom mówiącego są znaczeniami dosłownymi. Zauważmy jednak, że znaczeń w sensie Davidsona nie można utożsamiać z żadnymi bytami mentalnymi, fizycznymi i abstrakcyjnym. W swych wcześniejszych publikacjach Davidson argumentował, że znaczenia nie są bytami (zob.

¹¹ Jedną z konsekwencji antykonwencjonalizmu Davidsona jest krytyka tradycyjnej roli przypisywanej językowi. To, co nazywamy językiem, jest jedynie rodziną podobnych indywidualnych idiolektów. Owo konwencjonalne podobieństwo idiolektów stanowi jedynie o ułatwieniu procesu interpretacji, nie stanowi zaś jej warunku koniecznego.

¹² Jest to również konsekwencja założonej racjonalności użytkowników języka.

¹³ Gdyż są to zdania zazwyczaj wyrażające wspólne przekonania mówiącego i interpretatora.

¹⁴ Teoria interpretacji Davidsona nie wymaga ani bezbłędного przypisania przekonań mówiącemu, zgodności przekonań mówiącego i interpretatora, ani przypisania mu wszystkich możliwych przekonań. Proces interpretacji jest możliwy dzięki zachodzeniu tzw. zasady życzliwości, która w jednym z podanych przez Davidsona sformułowań głosi, że dla większości przekonań rozmówcy są zgodni, zaś przekonania te w większości są prawdziwe.

Davidson 1967). Spośród przytaczanych argumentów szczególną funkcję dla naszych rozważań pełni argument z wyuczalności języka (zob. Davidson 1965). Konsekwencją wyuczalności języka jest to, że teoria znaczenia musi być rekurencyjna, gdyż na podstawie znajomości znaczeń skończonej liczby słów powinna podawać znaczenia nieskończonej ilości zdań. Pozornie warunek ten spełnia teoria semantyczna Fregego, która odwołuje się do sensów jako bytów abstrakcyjnych, pozostając teorią rekurencyjną. Sens wyrażenia złożonego jest bowiem wartością funkcji (sensu funktora) dla określonego sensu jej argumentu. Najpoważniejszy zarzut Davidsona wobec semantyki Fregego dotyczy tzw. mowy zależnej (zob. Frege 1892). Rozwiązanie tego problemu przez Fregego miało ocalić zagrożoną zasadę zastępowania ekstensjonalnego. Aby to osiągnąć, Frege uznał, że w zdaniach podrzędnych w mowie zależnej, tj. w zdaniach postaci „X sądzi, że ...”, sens zdania podrzędnego nie jest jego sensem zwykłym, lecz sensem specjalnym, zaś sens zwykły staje się w tym kontekście jego odniesieniem¹⁵. Ponieważ zdaniem podrzędnym w mowie zależnej może być kolejne zdanie postaci „X sądzi, że ...”, generuje to dla każdego wyrażenia nieskończoną ilość wzajemnie nieredukowalnych sensów. To, w którym sensie dane wyrażenie zostało użyte w zdaniu, zależy od funkcji, jaką w tym zdaniu pełni. W rezultacie — jak argumentuje Davidson — gdyby semantyka Fregego adekwatnie opisywała język naturalny, nie byłibyśmy w stanie się nigdy go nauczyć.

Przytoczyłem tutaj argument przeciw semantyce Fregego, gdyż przypadek specjalnych sensów w mowie zależnej jest analogiczny do specjalnych sensów metaforycznych, które wyrażenia mogą mieć w pewnych specjalnych kontekstach metaforycznych. W artykule *What Metaphors Mean?* Davidson przytacza m.in. ten argument przeciw wszelkim teoriom metafory, które operują pojęciem sensu metaforycznego, a w konsekwencji zakładają, że metafora jest nośnikiem treści semantycznej (zob. Davidson 1978). Gdyby wyrażenia użyte metaforycznie miały specjalne sensory zależne od kontekstu użycia, nie byłibyśmy w stanie ich poznać w procesie nauki języka, gdyż nie da się zbudować rekurencyjnej semantyki, która mogłaby je podać. Specyficzne działanie metafory ma dla Davidsona charakter pozasemantyczny. Zgodnie ze stanowiskiem, które prezentuję w tym eseju, ów pozasemantyczny charakter ma również funkcja komunikacyjna metafory.

Aby to zilustrować, a zarazem rozszerzyć analizę Davidsona na inne figury retoryczne, odwołajmy się do pewnego przykładu, który wykorzystałem wcześniej do innych celów (zob. Maciaszek 2009). Wyobraźmy sobie, że ktoś mówi o wspólnym znajomym X, którego wszyscy uważają za osobę niezwykle uczynną, lecz niezbyt czytanną i wykształconą, że jest erudyta. Wypowiedzenie to może być interpretowane na wiele sposobów, z których przytoczę trzy:

¹⁵ Odniesieniem zwykłym jest wartość logiczna.

1. Mówiący przypisuje słowu „erudyta” to samo znaczenie co interpretator słowu „altruista” i uważa, że *X* jest altruistą¹⁶;

2. Mówiący przypisuje słowu „erudyta” to samo znaczenie co interpretator, lecz wbrew opinii przypisywanej mówiącemu uważa, że *X* erudytą nie jest¹⁷;

3. Mówiący przypisuje słowu „erudyta” to samo znaczenie co interpretator, lecz uważa, że *X* erudytą nie jest i to samo przekonanie przypisuje mówiącemu. Jest to przypadek kłamstwa, ironii lub metafory.

Aby odpowiedzieć na pytanie, z którym przypadkiem mamy do czynienia, należy odwołać się do szerszego zespołu przekonań, innych postaw propozycyjalnych oraz do szeroko rozumianego kontekstu wypowiedzenia obejmującego dotychczasowy przebieg i okoliczności rozmowy oraz wiedzę interpretatora na temat mówiącego, a w szczególności przypuszczeń na temat jego przekonań, preferencji itp.

Pomijając dość oczywiste przypadki 1 i 2, skoncentrujmy się na przypadku ostatnim. Z kłamstwem mielibyśmy do czynienia, gdy uznamy wprawdzie, że mówiący nie uważa *X*-a za erudytę, ale pragnie, aby interpretator był przekonany, że tak uważa. Z punktu widzenia semantyki znaczenie zdania będącego kłamstwem nie różni się od znaczenia zdania wypowiedzianego szczerze — różnica dotyczy przekonań i intencji przypisanych mówiącemu, np. mówiący pragnie zarekomendować *X*-a na stanowisko wymagające tej cechy. Z ironią mielibyśmy do czynienia, gdy uznamy, że mówiący nie uważa *X*-a za erudytę, ale jednocześnie uznamy, że wcale nie pragnie ukryć, że uważa coś przeciwnego. Z punktu widzenia teorii interpretacji Davidsona znaczenie wypowiedzi ironicznej nie różni się niczym od znaczenia wypowiedzi szczerzej i kłamstwa, zaś rozpoznanie ironii, np. po charakterystycznym tonie, polega na uznaniu, że mówiący żywi przekonanie przeciwne do wyrażanego przez zdanie ironiczne.

Wyobraźmy sobie jednak, że mówiący nazywa *X*-a erudytą, opisując jego rozliczne walory, które jednak z erudycją nie mają nic wspólnego. Okoliczności i przebieg rozmowy wskazują, że nie jest to kłamstwo, nic nie wskazuje również na ironię lub na „niestandardowe” rozumienie terminu „erudyta”. W takim przypadku wypowiedzenie „*X* jest erudytą” zinterpretujemy w sposób standardowy¹⁸, lecz jednocześnie zinterpretujemy je jako niedosłowne, gdyż warunkiem standardowej interpretacji tego terminu jest, wbrew werbalnym deklaracjom mówiącego, przypisanie mu przekonania, że *X* erudytą nie jest. Przypadek ten stawia pozornie pod znakiem zapytania racjonalność mówiącego. Aby uznać go za istotę racjonalną, szukamy powodów¹⁹, dla których wypo-

¹⁶ Mamy tu do czynienia z różnicą znaczeń, co stanowi przypadek niekonwencjonalnego rozumienia słowa „erudyta”, stawiający pod znakiem zapytania erudycję mówiącego.

¹⁷ Mamy tu do czynienia ze zwykłą różnicą przekonań na temat *X*-a.

¹⁸ Tj. przyjmujemy teorię prawdy, w której słowo „erudyta” rozumiane jest jako erudyta.

¹⁹ Termin „powód” (*reason*) został wyeksplikowany przez Davidsona w pracach poświęconych teorii działania. Nie wchodząc w szczegóły, powodem nazywamy racjonalną przyczynę, tj. zdarzenie opisane w specjalny sposób za pomocą terminów psychologicznych, np. przekonań i intencji (zob. Davidson 1963).

wiedział zdanie „*X* jest erudytą”. Przypuszczalne powody mogą mieć charakter pozakomunikacyjny, np. chęć zmiany nastawienia do *X*-a lub zaskoczenia słuchacza itp., lub charakter komunikacyjny, np. chęć wyrażenia jakiegoś przekonania na temat *X*-a lub chęć zwrócenia uwagi na jego szczególne cechy, które mogą być opisane w postaci zdań. W wielu teoriach metafory zdania wyrażające owe pozasemantycznie komunikowane przekonania utożsamiane są z parafrazą (podaniem znaczenia metaforycznego). Z punktu widzenia teorii interpretacji jest to nieporozumienie, gdyż do przekonań tych nie dochodzimy w procesie interpretacji, lecz w procesie racjonalizacji zachowania mówiącego. W swej teorii metafory Davidson nie odwoływał się do wprowadzonego przez niego pojęcia racjonalizacji, koncentrując się na konsekwencjach działania metafory. Sądzę jednak, że przedstawiona tu propozycja jest zgodna z duchem jego poglądów²⁰.

Na koniec przyjrzyjmy się pewnej modyfikacji podanego przykładu, która polega na tym, że *X* jest powszechnie znany ze swojej erudycji, interpretator uważa go za erudytę i to samo przekonanie przypisuje mówiącemu. Jest to typowy przykład aluzji, dla którego — podobnie jak w przypadku zdania metaforycznego — znaczenie wypowiedzenia jest jego znaczeniem dosłownym, zaś przypisanie mówiącemu pewnych przekonań i intencji, których nie bierzemy jednak pod uwagę, konstruując teorię prawdy, stanowi rekonstrukcję powodu, dla którego mówiący wypowiedział oczywistą i powszechnie znaną prawdę.

W przekazie semantycznym, tj. w ramach interpretacji obejmującej teorię prawdy, wypowiedzenia mówiącego tłumaczone są na wypowiedzenia interpretatora niezależnie od tego, czy którykolwiek z nich faktycznie żywi przekonania aktualizowane przez owe wypowiedzenia²¹. We wszystkich przypadkach poza pierwszym z odpowiedniej teorii prawdy wynika T-zdanie „Zdanie *X* jest erudytą jest prawdziwe wtedy i tylko wtedy, gdy *X* jest erudytą”. Jednocześnie przypisuje się mówiącemu — w zależności od przypadku — przekonanie, że *X* jest erudytą lub nim nie jest oraz szereg innych przekonań. Są to te przekonania, które są aktualizowane przez zdania używane w procesie weryfikacji tej teorii. Aby przypisać słowu „erudyta” standardowe znaczenie, musimy przypisać mówiącemu pewne przekonania na temat erudytów, np. „Każdy erudyta jest czytany” lub „Każdy erudyta używa bogatego słownictwa”, które można następnie zweryfikować za pomocą odpowiednich T-zdań. O takich przekonaniach powiemy, że pozostają w semantycznym związku z wypowiedzeniem „*X* jest erudytą”.

²⁰ W referacie wygłoszonym na VII Konferencji Hiszpańskiego Towarzystwa Logiki, Metodologii i Filozofii Nauki (Santiago de Compostela 18–20 lipca 2012) rozwinąłem teorię komunikowania pozasemantycznego w kategoriach teorii racjonalności Davidsona.

²¹ Sformułowanie, które głosi, że wypowiedzenie zdania aktualizuje przekonanie, pochodzi od Kazimierza Ajdukiewicza (1931).

Metafora może zwracać uwagę lub sugerować pewne podobieństwa przedmiotów, może nas irytować, prowokować do myślenia, śmieszyć, obrażać lub sugerować przekonania mówiącego, które nie pozostają w semantycznym związku z jej wypowiedzeniem²². Przekonania, których chęć wyrażenia stanowi domniemany powód wypowiedzenia zdania niedosłownego, nie pojawiają się w procesie interpretacji, tj. nie są niezbędne do skonstruowania teorii prawdy. Dlatego też komunikowanie przekonań za pomocą aluzji lub zdania metaforycznego jest komunikowaniem pozasemantycznym. Próby ustalania powodów wypowiedzeń są często skazane na niepewność i niepowodzenie. Wyjaśnia to, dlaczego metafory nie da się parafrazować w sposób jednoznaczny, a niekiedy nie poddaje się ona w ogóle prostym parafrazom.

Zgodnie z teoria Grice'a zdania interpretujemy jako metaforyczne, gdy uznamy, że doszło do pozornego złamania maksymy ilości. Podobny mechanizm rozpoznawania ma miejsce w przypadku zdania metonimicznego. Jednak jeżeli metonimię uznamy jedynie za skrót myślowy, to komunikowanie przekonań ma czysto semantyczny charakter. Na przykład dla wypowiedzenia zdania „Po ulicy maszerowały czerwone berety” konstruujemy teorię prawdy, z której wynika następujące T-zdanie: „Zdanie (mówiącego) *Po ulicy maszerowały czerwone berety* jest prawdziwe wtedy i tylko wtedy, gdy po ulicy maszerowali żołnierze w czerwonych beretach”. Problem pojawia się, gdy uznajemy, że użycie metonimii wykracza poza pełnienie funkcji skrótu myślowego. W takim przypadku mielibyśmy do czynienia z metaforycznym użyciem metonimii i obok podanej przed chwilą teorii prawdy skonstruowalibyśmy taką, z której wynika T-zdanie postaci: „Zdanie (mówiącego) *Po ulicy maszerowały czerwone berety* jest prawdziwe wtedy i tylko wtedy, gdy po ulicy maszerowały czerwone berety” (w sensie dosłownym). Możemy mieć również do czynienia z sytuacją, gdy uświadamiamy sobie obie możliwe interpretacje, co wyjaśniałoby komiczny efekt, jaki w pewnych okolicznościach wywołuje użycie niektórych metonimii.

Nieco wcześniej zwróciłem uwagę na fakt, że negacja zdania metaforycznego, np. „Nie jestem trybem w maszynie”, jest w ujęciu Grice'a aluzją, z czego wynika, że aluzje mogą zawierać metafory. Z drugiej strony, niektóre wypowiedzenia zdań metaforycznych mogą metafor nie zawierać, np. zdania oksymoroniczne „Dwa plus dwa to pięć” lub „Każdy plus ma swoje minusy”. Ponadto negacje typowych aluzji często mają wyczuwalny oksymoroniczny charakter. Mamy zatem następującą sytuację. Wedle kryterium Grice'a rozróżnienie wypowiedzeń zdań na metaforyczne i aluzje opiera się na odwołaniu się do różnych maksym — jakości i ilości. Jednak negacja zdania metaforycznego

²² Mówiący mógł chcieć wyrazić przekonanie, że *X* ma wszechstronne uzdolnienia artystyczne, które nie pozostawałoby w żadnym semantycznym związku z przekonaniem, że *X* jest erudyta. W takim przypadku metafora „erudyta” mogłaby być sparafrazowana jako „osoba o wszechstronnych uzdolnieniach artystycznych”.

może być aluzją, zaś negacja aluzji zdaniem metaforycznym (oksymoronicznym). Nie ma w tym nic dziwnego, gdyż wg teorii informacji wygłoszenie powszechnie uznanego zdania nie niesie żadnej informacji, natomiast wygłoszenie zdania powszechnie odrzucanego niesie nadmiar informacji. W szczególności wygłoszenie sprzeczności (oksymoronu) niesie nieskończenie wiele informacji²³. Być może należałoby mówić o zdaniach niedosłownych, w których może, lecz nie musi, występować metafora. Sam Davidson nie jest tutaj zbyt precyzyjny. Nie odróżnia wyraźnie metafory od zdania metaforycznego, określając jako metafory zarówno ewidentne fałsze, jak i oczywiste prawdy, czyli aluzje w proponowanej tu terminologii²⁴.

5. Metafora

W paragrafie tym proponuję pewne kryterium kwalifikowania wyrażenia występującego w zdaniu niedosłownym jako metafory²⁵. Kryterium to odwołuje się do holizmu przekonaniowego, który — oprócz Davidsona — głosiło wielu filozofów języka, np. Ajdukiewicz (1931) oraz Quine (1960). Zgodnie z tym poglądem nie można mieć tylko jednego przekonania lub (w ujęciu alternatywnym), jeżeli ktoś ma jedno przekonanie, to musi posiadać ich znacznie więcej, co stanowi warunek konieczny znajomości języka. Oczywiście nie sposób określić, ile przekonań należy posiadać, aby władać danym językiem, ale — jak twierdzi Davidson — aby znać język musimy dysponować całą „siecią” przekonań i innych postaw propozycyjalnych.

Zacznijmy od prostego przykładu zdania „*X* jest osłem” wypowiedzianego o *X*-ie, o którym wszyscy wiedzą, że jest przedstawicielem *homo sapiens*. Przekonanie, które jest wyrażane przez to zdanie, nie znajduje się w systemie przekonań żadnego uczestnika rozmowy, dzięki czemu jego wypowiedzenie ma charakter niedosłowny. Dlaczego jednak nie ulega dla nas wątpliwości, że metaforą jest właśnie wyrażenie „osioł”? Jak się zdaje, decyduje o tym zawartość systemów przekonań uczestników rozmowy, obejmujących przekonanie, że *X* jest człowiekiem. Musimy ponadto przyjąć, że użytkownicy języka są istotami racjonalnymi, wyposażonymi w minimalną kompetencją logiczną, co przejawia się zdolnością do przeprowadzania prostych inferencji logicznych. W naszym przypadku chodziłoby o inferencję ze zdania „*X* jest człowiekiem” do zdania „*X* nie jest osłem”. Uogólniając to spostrzeżenie, powiemy, że wyrażenie α w wy-

²³ Może to prowadzić do konkluzji, że maksyma jakości faktycznie da się zredukować do maksymy ilości. Wątku tego nie będę jednak tutaj rozwijał.

²⁴ „Patent falsity is the usual case with metaphor, but on occasion patent truth will do as well”, po czym podaje przykład „Business is business”, który podpada pod etykietę aluzji (zob. Davidson 1978/2001: 258).

²⁵ W literaturze przedmiotu traktuje się metaforę jako coś, czego rozpoznanie w zdaniu jest oczywiste. Nie udało mi się dotąd spotkać kryterium, które pozwalałoby wskazać na metaforę w zdaniu metaforycznym.

powiedzeniu metaforycznym W jest metaforą, gdy istnieje powszechne przekonanie V , które zawiera wyrażenie β tej samej kategorii syntaktycznej co α oraz z $V(\beta)$ wynika $\sim W(\alpha)$ ²⁶. Jeżeli nie potrafimy znaleźć takiego zdania, to mamy do czynienia z wyrażeniem metaforycznym, w którym nie potrafimy wyróżnić metafory, np. ze zdaniem oksymoronicznym.

Przyjrzyjmy się teraz przykładowi aluzji „Nie jestem trybem w maszynie”. Według powszechnego przekonania mówiący jest człowiekiem, z czego wynika, że trybem w maszynie z pewnością nie jest. A zatem wyrażenie α występujące w aluzji W jest metaforą, gdy istnieje powszechnie uznawane przekonanie V , które zawiera wyrażenie β tej samej kategorii syntaktycznej co α oraz z $V(\beta)$ wynika $W(\alpha)$. Jeżeli nie potrafimy znaleźć takiego zdania, to mamy do czynienia z aluzją, która metafory prawdopodobnie nie zawiera, np. zdanie „Kulturalni ludzie wycierają obuwie” wypowiedziane do osoby, która w zabłoconych butach wchodzi na jasny dywan.

Przedstawiona tu propozycja ma charakter szkicowy i może być zastosowana jedynie do prostych przypadków metafory. W przypadku nagromadzenia środków metaforycznych interpretator może nie móc rozpoznać metafory, gdyż ma trudność z uzmysłowieniem sobie przekonania oznaczonego jako $V(\beta)$. Wyjaśnia to, dlaczego interpretator może się mylić lub wahać się, czy dane wyrażenie faktycznie jest użyte metaforycznie. Decydują o tym jego przypuszczenia — niekiedy mylne — na temat powszechnie żywionych przekonań oraz stopień komplikacji struktury składniowej wypowiedzenia.

Na zakończenie paragrafu rozpatrzmy przypadek trudnego do zinterpretowania wypowiedzenia „Chrystus był cieślą”. Zdanie to może być w pewnych okolicznościach rozumiane dosłownie, gdyż Chrystus w młodości faktycznie wyuczył się tego zawodu. A zatem jeśli jest ono wygłoszone w gronie osób, dla których przekonanie o tym, że Chrystus był z zawodu cieślą, nie jest powszechne, jest wypowiedzeniem dosłownym. Problem zaczyna się przy zinterpretowaniu tego zdania jako niedosłownego, np. przy wypowiedzeniu go w gronie osób, które doskonale znają szczegóły biografii Chrystusa. Pierwsza możliwość polega na potraktowaniu go jako paradygmatycznej aluzji (przyp. 1), tj. aluzji niezawierającej metafory. Niemniej jednak intuicja nam podpowiada, że termin „cieśla” został użyty metaforycznie. A zatem druga możliwość — jak się okazuje o wiele bardziej kłopotliwa — polega na potraktowaniu go jako zdania niedosłownego z metaforycznie użytym terminem „cieśla”. Gdyby była to aluzja z metaforą „cieśla” (przyp. 2a), natrafiamy na problem, gdyż nie ma chyba powszechnie uznawanego przekonania, w którym termin „cieśla” zastąpiony byłby jakimś innym terminem, z którego wynikałoby to zdanie. Z drugiej strony, gdyby było to zdanie metaforyczne (przyp. 2b), to powszech-

²⁶ Przez powszechne przekonanie rozumiem przekonanie przypisywane przez rozmówców sobie nawzajem oraz osobom z ich najbliższego otoczenia. Wyrażenie α występuje w przekonaniu, gdy występuje w zdaniu, które je wyraża.

nie uznawanym przekonaniem musiałaby być jego negacja, która głosi, że Chrystus cieślą nie był.

Proponowane rozwiązanie tego dylematu odwołuje się do subtelnej niejednoznaczności terminu „cieśla”, która dotyczy również nazw innych zawodów oraz pełnionych funkcji, np. rektora lub dziekana. W jednym z możliwych znaczeń, aby być cieślą, należy ten zawód faktycznie wykonywać. W drugim znaczeniu wystarczyło się go kiedyś wyuczyć²⁷. Niejednoznaczność ta ma odbicie w dwóch konkurencyjnych interpretacjach zdań z terminem „cieśla”, które podawałyby różne znaczenia tego terminu i różniłyby się przekonaniem przypisywanymi mówiącemu oraz przekonaniem interpretatora. Dla jednej z tych interpretacji zdanie „Chrystus był cieślą” wyraża powszechnie żywione przekonanie. Dla drugiej, powszechnie żywione przekonanie głosiłoby, że Chrystus cieślą nie jest, gdyż w okresie życia, o który chodzi, zajmował się nauczaniem, czynił cuda i nie wykonywał tego rzemiosła. W pierwszym przypadku byłaby to aluzja bez metafory (przyp. 1), w drugim zdanie metaforyczne zawierające metaforę „cieśla” (przyp. 2b). Jednocześnie uświadamiana lub jedynie przeczuwana niejednoznaczność interpretacji może wywołać dodatkowy efekt wahania i niepewności.

6. Analiza parataktyczna i ironia

Nieco wcześniej, wbrew utartemu pogładowi, zaliczyłem wypowiedzenia ironiczne do wypowiedzeń dosłownych, gdyż przekonanie aktualizowane wypowiedzeniem ironicznym — w przeciwieństwie do wypowiedzenia aluzyjnego i metaforycznego — jest komunikowane w sposób semantyczny i jest po prostu przeciwne przekonaniu aktualizowanemu przez wypowiedziane zdanie. Sytuacja ta da się łatwo opisać za pomocą analizy parataktycznej Davidsona.

Analiza parataktyczna została wprowadzona przez Davidsona dla potrzeb analizy mowy zależnej²⁸, zdań performatywnych oraz wyrażen cudzysłowowych. Polega ona na potraktowaniu wypowiedzenia zdania z poziomu powierzchniowego jako pary wypowiedzeń zdań z poziomu głębszego, przy czym pierwsze wypowiedzenie wskazuje na drugie²⁹. Dla podanego w (1968) przykładu „Galileusz powiedział, że Ziemia się porusza”, Davidson proponuje: „Galileusz powiedział to. Ziemia się porusza”. Zdaniem Davidsona oba wypowiedzenia zdań mają odrębne warunki prawdziwości i nie pozostają ze sobą w żadnym związku logicznym. Dla zdania performatywnego „Czołgaj się do najbliższego

²⁷ Na przykład wśród pracowników Instytutu Filozofii UŁ niemal każdy ukończył również jakieś studia niefilozoficzne. Dlatego też w drugim ze wskazanych znaczeń są wśród nas fizycy, matematycy, filolodzy, historycy, inżynierowie, a nawet farmaceuci. Jednak odpowiedź na pytanie, czy faktycznie w Instytucie Filozofii pracują fizycy, matematycy, filolodzy, historycy, inżynierowie i farmaceuci, wcale nie jest oczywista.

²⁸ Aby uniknąć trudności, o których była mowa w paragrafie 4.

²⁹ A zatem wypowiedzenie to ma charakter zdania okazjonalnego postaci „To jest ...”.

drzewa” wypowiedzianego przez oficera do szeregowca można zaproponować następującą analizę: „Moje następne wypowiedzenie jest rozkazem. Czołgasz się do najbliższego drzewa”³⁰. Propozycja ta pozwala rozwiązać wielokrotnie dyskutowany problem wartości logicznej zdań performatywnych. Ponieważ wypowiedzenia te są faktycznie parami wypowiedzeń zdań, odpowiadają im pary wartości logicznych, np. [0,1] odpowiadałoby przypadkowi fałszywego rozkazu (np. wydanego przez przebranego za oficera oszusta), który został jednak wykonany, zaś [1,0] — prawdziwego rozkazu, który został zbojkotowany.

Analizę parataktyczną można z powodzeniem zastosować do ironii. W typowych przypadkach wypowiedzeń ironicznym pojawia się ironiczny ton, który pełni funkcję wskaźnika, aby nie wierzyć w to, co zostało powiedziane. Dlatego też wypowiedziane ironicznym tonem „On jest wybitnym mężem stanu” (o polityku), może być interpretowane w następujący sposób: „To (tj. następne wypowiedzenie) nie wyraża mojego przekonania. On jest wybitnym mężem stanu”. Z tak analizowanym wypowiedzeniem wiążą się cztery pary wartości logicznych przypisanych mu przez interpretatora:

1. [1,0] — prawdziwa ironia, gdyż interpretator uważa, że mówiący — podobnie jak on sam — nie wierzy w wypowiedziane zdanie³¹;
2. [1,1] — prawdziwa ironia, gdyż interpretator uważa, że mówiący — w przeciwieństwie do niego — nie wierzy w wypowiedziane zdanie³²;
3. [0,1] — fałszywa ironia, gdyż interpretator uważa, że mówiący — podobnie jak sam interpretator — wierzy w wypowiedziane zdanie³³;
4. [0,0] — fałszywa ironia, gdyż interpretator uważa, że mówiący — w przeciwieństwie do interpretatora — wierzy w wypowiedziane zdanie³⁴.

Z punktu widzenia interpretatora prawdziwość ironii to nic innego niż uznanie zrekonstruowanego w tej analizie okazjonalnego wypowiedzenia „To nie wyraża mojego przekonania” za prawdziwe. Oczywiście w grę wchodzi możliwość pomyłki interpretatora, niepewność lub nawet niemożliwość przypisania wypowiedzeniu pierwszej z wartości logicznych. Niekiedy interpretator nie ma podstaw do przypisania mówiącemu ani przekonania wyrażonego przez drugie wypowiedzenie, ani przekonania przeciwnego — subtelna ironia bywa niekiedy trudna do rozpoznania. Zauważmy jednak, że o ile jest sens mówienia o prawdziwej lub fałszywej ironii, o tyle nie ma sensu mówić o prawdziwej metaforze lub aluzji. Mogą one być jedynie nietrafione, niewłaściwe lub nieodosowne.

³⁰ Nawiązuję tutaj do rozważań Davidsona z (1979), lecz przykład pochodzi ode mnie.

³¹ Prawdopodobnie interpretator reprezentuje tę samą opcję polityczną co mówiący i wypowiedzenie go rozbawi.

³² W tym przypadku interpretator pewnie uzna to za przejaw niezwykle złośliwej ironii.

³³ Może to być przejaw tzw. „życzliwej” ironii lub autoironii.

³⁴ Odpowiada to nieudanemu wypowiedzeniu „życzliwej” ironii lub autoironii, co może wywołać rozbawienie interpretatora. Wyobraźmy sobie, że osoba, którą uważamy za wyjątkowo nieinteligentną mówi sama o sobie ironicznym tonem „Jaki ja jestem głupi!”.

Pewne zastrzeżenia wobec przedstawionego tu rozwiązania może wzbudzać oczywisty fakt, że ironia — podobnie jak metafora czy aluzja — działa w sposób pozakomunikacyjny, np. wzbudza emocje, złość, rozbawienie lub upokorzenie. Owe efekty nie przysługują jednak wyłącznie wypowiedziom niedosłownym. Obrazić, rozśmieszyć czy upokorzyć można za pomocą wypowiedzi dosłownych i niedosłownych. Co więcej, aby osiągnąć ten efekt, nie musimy w ogóle posługiwać się językiem — możemy go osiągnąć gestem lub miną³⁵. Davidson twierdził, że jeżeli wypowiedzenie zostało zinterpretowane, to może być użyte na nieskończenie wiele sposobów (zob. Davidson 1978), dotyczy to zarówno standardowych wypowiedzi dosłownych, ironii, metafor, jak i aluzji.

Przeciwko przedstawionej tu propozycji można wysunąć jeszcze jeden zarzut. Ironiczny ton nie jest warunkiem koniecznym zinterpretowania wypowiedzenia jako ironiczne. Niekiedy wystarczy do tego sam szeroko rozumiany kontekst³⁶, choć ironiczny ton jest niewątpliwie elementem konwencjonalnym ułatwiającym rozpoznanie ironii. Zgodnie z Davidsonowską filozofią języka konwencje nie stanowią warunku niezbędnego dla procesu komunikacji językowej. Motyw ten pojawia się w bardzo wielu pracach Davidsona, zwłaszcza zaś w (Davidson 1984). Zdaniem Davidsona konwencje stanowią jedynie środki usprawnienia komunikacji językowej, tj. procesu wzajemnej interpretacji, nie stanowią zaś jego istoty. Stwierdzenie to należy jednak rozumieć w sposób ostrożny. Z pewnością systematyczne łamanie wszystkich konwencji językowych drastycznie utrudnia proces komunikacji. Chodzi tu raczej o to, że nie można wskazać jakiegokolwiek konwencji, której incydentalne złamanie uniemożliwiłoby skuteczną interpretację³⁷. Każda konwencja może być złamana pod warunkiem, że większość konwencji jest przestrzegana. Ironizowanie bywa właściwie interpretowane nawet bez ironicznego tonu.

Na zakończenie tego paragrafu rozpatrzmy cztery przypadki mieszane, które w większości były już wspomniane wcześniej. Nauczyciel zwraca się ironicznym tonem do wybitnego, lecz gadatliwego ucznia: „Ty już wszystko umiesz”. Zgodnie z przedstawioną tu propozycją byłby to typowy przykład fałszywej ironii (przyp. [0,1])³⁸. Poczucie, że mamy do czynienia z wypowiedzeniem niedosłownym bierze się stąd, że drugie wypowiedzenie (w analizie parataktycznej) wyraża żywione powszechnie przekonanie. Pełni ono zatem

³⁵ Zauważmy również, że za pomocą gestów lub min możemy komunikować również przekonania i inne postawy propozycyjne.

³⁶ Na przykład przytoczona w tym paragrafie wypowiedź w ustach szefa partii, która niedawno zadebiutowała w polskim parlamencie, o przewodniczącym innej opozycyjnej partii prawicowej.

³⁷ Należy również zwrócić uwagę na fakt, że Davidson dopuszcza możliwość interpretacji bez jakichkolwiek konwencji (interpretacja radykalna). Komunikacja w takim przypadku byłaby jednak skrajnie nieefektywna.

³⁸ Być może byłby to przykład „życzliwej” ironii.

rolę aluzji, którą może komunikować coś w sposób pozasemantyczny. To samo wypowiedzenie kierowane do słabego ucznia byłoby prawdziwą ironią (przyp. [1,0]).

Menadżera, który nosi niezbyt świeżą białą koszulę, określamy ironicznie mianem „białego kołnierzyka” (typowa metonimia). Jeżeli wypowiedzenie „On jest białym kołnierzykiem” interpretujemy standardowo jako „On jest menadżerem”, mielibyśmy do czynienia z fałszywą ironią (przyp. [0,1]). Komiczny efekt tego wypowiedzenia jest prawdopodobnie rezultatem kontrastu między przekonaniem, które standardowo przypisuje się menadżerom (nienaganny wygląd), a przypadkiem tej konkretnej osoby. W szczególności można podejrzewać, że mówiący mógł mieć na myśli inną interpretację zdania „On jest białym kołnierzykiem”, dla której z odpowiedniej teorii prawdy wynikałoby T-zdanie: „Wypowiedzenie *On jest białym kołnierzykiem* jest prawdziwe wtedy i tylko wtedy, gdy on nosi biały kołnierzyk”. W takim przypadku byłaby to prawdziwa ironia, zaś samo uświadomienie sobie możliwości niestandardowej interpretacji wywołuje komiczny efekt związany z oksymoronicznym aspektem wypowiedzenia³⁹.

Rozpatrzmy obecnie przypadek ironii z metaforą. Przed egzaminem poprawkowym egzaminator zwraca się ironicznym tonem do poprawkowiczów „Orły już czekają”. Ponieważ słowo „orzeł” jest metaforą, mamy do czynienia z interpretacją, w której z teorii prawdy wynika następujące T-zdanie: „Wypowiedzenie *Orły już czekają* jest prawdziwe wtedy i tylko wtedy, gdy orły już czekają”, gdzie przez orły rozumiemy duże drapieżne ptaki. Komiczny efekt polega na uświadomieniu możliwości, że wypowiedzenie „To nie wyraża mojego przekonania” nie wskazuje na „Orły już czekają”, lecz na wypowiedzenie wyrażające przekonanie, że czekają na mnie wybitni studenci, będące typową parafrazą tej skonwencjonalizowanej metafory. Zauważmy również, że ilość możliwości byłaby jeszcze większa w przypadku, gdyby owi poprawkowicze doskonale zdawali egzaminy z innych przedmiotów. Doszłaby wówczas możliwość zinterpretowania ironii jako fałszywej, co dodatkowo zmodyfikowałoby efekt jej działania.

Wypowiedzenie przez nauczyciela geografii przerażonego niewiedzą ucznia słów „Zakopane leży nad Bałtykiem” jest niewątpliwie prawdziwą ironią (w zamierzeniu mówiącego przypadek [1,0]). Ponieważ jednak przypisywane mówiącemu przekonanie, że Zakopane nad Bałtykiem nie leży, jest powszechnie akceptowane, wypowiedzenie to ma wyraźnie aluzyjny charakter i pozwala się domyślać powodów, dla których zostało użyte. W szczególności możemy się

³⁹ Możliwa jest również skrajnie nieprawdopodobna interpretacja, że osoba ta faktycznie jest białym kołnierzykiem. Wówczas działanie wypowiedzenia przypominałoby działanie błędu składniowego zwanego amfibologią, np. „Kowalski wyrzucił na śmietnik kredens wraz z żoną”. Możliwość istnienia nieprawdopodobnej interpretacji wywołuje tutaj wyraźny efekt komiczny.

spodziewać, że nauczyciel chciał zakomunikować jakieś niepochlebne treści na temat ucznia.

7. Działanie metafory językowej i pozajęzykowej

Wypowiedzenia niedosłowne stosujemy ze względu na spodziewane efekty ich działania. Nie oznacza to jednak, że można je adekwatnie w tych terminach zdefiniować⁴⁰. Do owych efektów można zaliczyć ośmieszenie kogoś, rozbawienie, zdenerwowanie, obrażenie lub sprawienie przyjemności. W szczególności może to być przekazanie pewnego przekonania, co prowadzi do zwyczajnego parafrazowania wypowiedzeń niedosłownych (metafor i aluzji). Ponieważ parafraza uznawana jest za dosłowny synonim wypowiedzenia użytego niedosłownie, prowadzi to wielu teoretyków do zaakceptowania pojęcia znaczenia metaforycznego. Z punktu widzenia przedstawionej w tym eseju propozycji parafrazy, o ile jest możliwa, jest jedynie opisem domniemych przekonań mówiącego, których chęć zakomunikowania mogła być powodem wypowiedzenia.

Wypowiedzenie zdania jest — wedle Davidsona — zdarzeniem fizycznym, które posiada swoje przyczyny i skutki będące również zdarzeniami. Wypowiadając zdania, pragniemy osiągnąć pewne skutki w psychice naszych słuchaczy, w szczególności pragniemy zakomunikować im własne przekonania. Możemy to czynić na dwa sposoby — w sposób semantyczny (w ramach interpretacji) lub pozasemantyczny. W drugim przypadku pragniemy, aby interpretator szukał powodów, tj. opisanych za pomocą terminów mentalnych przyczyn naszych wypowiedzeń. Efekty naszych wypowiedzeń są zasadniczo nieprzewidywalne, gdyż zdarzenia opisane jako mentalne nie podpadają, w przeciwieństwie do zdarzeń fizycznych, pod ścisłe prawa (zob. Davidson 1970). Również interpretator nie jest w stanie w sposób pewny i jednoznaczny określać powodów wypowiedzenia wyrażenia niedosłownego. Zdaje się to dobrze charakteryzować działanie metafory poetyckiej, która jest otwarta i wymyka się jednoznacznym parafrazom, działając jednocześnie na emocje i uczucia słuchaczy. Spowodowane jest to tym, że zazwyczaj nie da się jednoznacznie podać powodów wypowiedzenia pewnych słów przez poetę, a on sam nie jest w stanie przewidzieć skutków ich użycia. Komunikacja pozasemantyczna wypowiedzeń niedosłownych jest obarczona o wiele większą niepewnością i niejednoznacznością niż komunikacja semantyczna. Jest jednak, w odróżnieniu od tej drugiej, interesująca i nietrywialna⁴¹.

⁴⁰ Próby takich definicji przytoczyłem w paragrafie 2.

⁴¹ Oczywiście nie wszyscy filozofowie zgodziliby się z tym, że metafora może coś komunikować ciekawego na drodze pozasemantycznej. Przykładem może być Paul Ricoeur, który stał na stanowisku, że metafora może być nietrywialnym środkiem komunikacji semantycznej (zob. Ricoeur 1989). Aby zakwestionować tego typu stanowisko, należy powtórzyć postulat sformułowany na początku tego artykułu: adekwatna teoria metafory powinna być nadbudowana nad ścisłą teorią znaczenia (semantyką). Teoria Davidsona z pewnością ten wymóg spełnia.

Oryginalna teoria metafory Davidsona może sprawiać wrażenie zbyt radykalnej i częściowo niezgodnej z przedstawionymi w tym eseju poglądami, szczególnie wobec stwierdzeń, że metafora nie posiada żadnej zawartości i nie pełni żadnej funkcji poznawczej (kognitywnej). Jak sądzę, problem leży w terminologii użytej przez Davidsona, który rozumiał funkcję poznawczą w wąskim sensie komunikowania treści dosłownych na drodze semantycznej. Jednak sam Davidson *implicite* uznawał możliwość komunikacji pozasemantycznej. Aby opisać przekaz charakterystyczny dla metafory, Davidson powołuje się na analogię między metaforą a wyrocznią, przywołując słowa Heraklita „*It does not say, and it does not hide, it intimates*”⁴² (zob. Davidson (1978/2001: 262).

W monografii poświęconej metaforze Samuel Guttenplan (2005) twierdzi, że metafora polega na orzekaniu rzeczy, nie zaś słowa, np. Romeo, mówiąc „Julia jest słońcem”, nie orzeka o niej wyrażenia językowego, lecz rzecz (słońce). Metafora byłaby tutaj „wyrażeniem” z poziomu „zero”, czyli poziomu rzeczy, jeśli przez język przedmiotowy rozumiemy poziom pierwszy, zaś przez metajęzyk poziom drugi. Pomijając pewne trudności związane z tym poglądem, którego nie będę tutaj dokładnie charakteryzował, można powiedzieć, że rozpoznanie wyrażenie jako metaforycznego sprawia, że nie bierzemy pod uwagę znaczenia słowa, lecz samą rzecz lub jej obraz. Pozornie wydawać się może, że da się wyprowadzić bardzo prostą analogię między tym stanowiskiem a stanowiskiem Davidsona. Każda rzecz, a nawet jej wyobrażenie, może być opisywana na nieskończenie wiele sposobów, a zatem — podobnie jak metafora — byłaby niewyczerpalną kopalnią treści i skojarzeń. Metafora działa bowiem jak sama rzecz lub jej wyobrażenie — pozwala dostrzec podobieństwa, analogie i ukryte cechy w tym, o czym się ją orzeka. Podobnie jak rzecz może się podobać lub nie podobać, rozśmieszać lub denerwować, niepokoić lub uspokajać.

Jak się zdaje, analogia ta jest zbyt prosta, choć sam Davidson był prawdopodobnie jej zwolennikiem, pisząc „Obrazu nie da się wymienić na słowa”, Davidson (1978/2001: 263). Twierdzę tak, gdyż podobnie jak słowo może być interpretowane na wiele sposobów (niezderminowanie interpretacji), tak obraz i rzecz mogą być opisane na wiele sposobów. A zatem opis rzeczy lub jej obrazu byłby raczej odpowiednikiem interpretacji dosłownej. Jednak w pomysłu Guttenplana tkwi możliwość rozszerzenia pojęcia metafory na obszary pozajęzykowe, np. sztuki plastyczne. W przypadku metafory pozajęzykowej nie da się w prosty sposób zastosować semantycznych kryteriów rozpoznawania, które zostały przedstawione w tym eseju⁴³. Można jednak wskazać na daleko idącą analogię metafory w sztuce do niedosłowności językowej. Jeżeli uznamy, że motywem namalowania obrazu nie było jedynie przedstawienie wielu rze-

⁴² „Ona niczego nie mówi i niczego nie ukrywa — ona daje do zrozumienia” (przekład mój). Trudno oprzeć się skojarzeniu z implikacją konwersacyjną Grice’a.

⁴³ W szczególności nie da się zastosować do tego celu maksym konwersacyjnych. O wiele bardziej obiecujące zdaje się odwołanie się do szeroko rozumianego pojęcia racjonalności.

czy posiadających tę samą i powszechnie znaną własność, np. zestawienie ła-two psujących się produktów, jesteśmy skłonni — pomijając oczywiście aspekt estetyczny dzieła — interpretować martwą naturę niedosłownie, np. jako aluzję do śmierci, przemijania. Dochodzimy do wniosku, że artysta, którego traktujemy jako istotę racjonalną, chciał nam coś powiedzieć, np. o przemijaniu wszystkich rzeczy, ale namalował jedynie nagromadzenia owoców i ryb, o których wszyscy doskonale wiedzą, że są bardzo nietrwałe. W tym przypadku o niedosłownej interpretacji obrazu decydowałaby narzucająca się i oczywista dla wszystkich wspólna cecha przedstawionych obiektów.

Z podobnego powodu jesteśmy skłonni interpretować metaforycznie obrazy Hieronima Boscha, w których pojawiają się dziwne i absurdalne zestawienia niepasujących do siebie rzeczy oraz hybryd ludzi i zwierząt. Ów dysonans poznawczy spowodowany ewidentną fałszywością tego, co zostało przedstawione sprawia, że doszukujemy się w tych obrazach ukrytych treści, które obok innych efektów ich działania, np. efektu zaskoczenia, zaniepokojenia, a niekiedy obawy lub odrazy, mogą być opisane słowami, chociaż na próżno szukać bezpośredniego odpowiednika tego opisu na płótnie. Ów opis efektów działania wykracza poza dosłowny opis obrazu, podobnie jak opis efektów działania metaforycznego wiersza, łącznie z przekonaniami przypisanymi jego twórcy, wykracza poza to, co dosłownie zostało w nim napisane⁴⁴. Niedosłowne elementy w malarstwie mogą wywołać u osoby wrażliwej i pomysłowej — podobnie jak wypowiedzenia niedosłowne w języku — nieprzewidziane i niespodziewane skutki.

Każda metafora, językowa czy też pozajęzykowa, pozostaje otwarta, zaś parafraza, jeśli uznamy ją za obowiązującą, prowadzi do jej uśmiercenia. Interpretacja niedosłowności jest — jak twierdził Davidson — przedsięwzięciem twórczym, w którym rola konwencji jest o wiele mniejsza niż w interpretacji dosłownej. Polega ona bowiem na szukaniu powodów wypowiedzenia szczególnych zestawień słów lub — w przypadku metafory pozajęzykowej — przedstawienia szczególnych zestawień rzeczy. Interpretując dosłownie, tworzymy teorię prawdy o ściśle określonych regułach. Interpretując niedosłowności, możemy polegać jedynie na doświadczeniu i wycuciu, które z trudem daje się ująć w postaci reguł⁴⁵.

⁴⁴ Z tego punktu widzenia symbole w malarstwie, np. lew jako symbol władzy, byłyby odpowiednikami skostniałych metafor. Lew siedzący na tronie jest ewidentnie fałszywym przedstawieniem rzeczywistości i domaga się interpretacji metaforycznej.

⁴⁵ Próby podania takich reguł można znaleźć m.in. u Searle'a (1979) i Martinicha (1984).

Bibliografia

- AJDUKIEWICZ K. (1931), *O znaczeniu wyrażań*, [w:] *Księga Pamiątkowa Polskiego Towarzystwa Filozoficznego we Lwowie*. Lwów: Polskie Towarzystwo Filozoficzne, Skł. gł. w Księgarni Książnica — Atlas, 31–77. Przedruk [w:] Ajdukiewicz K. (1960), *Język i poznanie*, t. 1. Warszawa: PWN, 102–36.
- AJDUKIEWICZ K. (1965), *Logika pragmatyczna*. Warszawa: PWN.
- DAVIDSON D. (1963), *Actions, Reason, and Causes*, „The Journal of Philosophy” 60, 685–701. Przedruk w: Davidson (2001), 3–19.
- DAVIDSON D. (1965), *Theories of Meaning and Learnable Languages*, [w:] *Logic, Methodology and Philosophy of Science: Proceedings of the 1964 International Congress*. Amsterdam: North Holland Publishing Co., 383–394. Przedruk [w:] Davidson (2001), 3–15.
- DAVIDSON D. (1967), *Truth and Meaning*, „Synthèse” 17, 304–323. Tłumaczenie polskie: *Prawda i znaczenie* (przeł. J. Gryz), w: Davidson (1992), 3–32.
- DAVIDSON D. (1968), *On Saying That*, „Synthèse” 19, 130–146. Przedruk [w:] Davidson (2001), 93–108.
- DAVIDSON D. (1970), *Mental Events*, [w:] Foster L., Swanson J. W. (red.), *Experience and Theory*. The University of Massachusetts Press i Duckworth. Tłumaczenie polskie: *Zdarzenia mentalne* (przeł. T. Baszniak), w: Davidson (1992), 163–193.
- DAVIDSON D. (1978), *What Metaphors Mean*, „Critical Inquiry” 5, 31–47. Przedruk [w:] Davidson (2001), 245–264.
- DAVIDSON D. (1979), *Moods and performance*, [w:] Margalit A. (red.), *Meaning and Use*. Dordrecht: Reidel. Przedruk [w:] Davidson (2001), 109–121.
- DAVIDSON D. (1984), *Communication and Convention*, „Synthèse” 59, 3–17. Przedruk [w:] Davidson (2001), 265–280.
- DAVIDSON D. (1992), *Eseje o prawdzie, języku i umyśle*. Warszawa: PWN.
- DAVIDSON D. (2001), *Inquiries into Truth and Interpretation*. Oxford: Clarendon Press.
- FREGE G. (1892), *Über Sinn und Bedeutung*, [w:] *Zeitschrift für Philosophie und Philosophische Kritik*, vol. 100, s. 25–50. Tłumaczenie polskie: *Sens i znaczenie*, [w:] Frege G. (1977), *Pisma semantyczne* (przeł. B. Wolniewicz). Warszawa: PWN, 60–88.
- GRICE H. P. (1957), *Meaning*, „Philosophical Review” 66, 377–388.
- GRICE H. P. (1975), *Logic and Conversation*, [w:] Cole P., Morgan J. (red.), *Syntax and Semantics*, vol. 3. Academic Press: London. Tłumaczenie polskie: *Logika a konwersacja* (przeł. B. Stanosz), w: Stanosz B. (red.) (1980), *Język w świetle nauki*. Warszawa: Czytelnik, 91–114.
- GUTTENPLAN S. (2005), *Objects of Metaphor*. Oxford: Clarendon Press.
- LEPORE E., LUDWIG K. (2007), *Donald Davidson's Truth-Theoretic Semantics*. Oxford: Clarendon Press.
- MACIASZEK J. (2008), *Znaczenie, prawda, przekonania. Problematyka znaczenia w filozofii języka*. Łódź: Wydawnictwo Uniwersytetu Łódzkiego.
- MACIASZEK J. (2009), *Filozofia języka Donalda Davidsona*, [w:] Habrajska G. (red.) *Rozmowy o komunikacji* 3, 147–170.
- MARTINICH A. P. (1984), *A Theory of Metaphor*, „Journal of Literary Semantics” 13, 35–56.
- NOWA ENCYKLOPEDIA POWSZECHNA (1995), t. 1–6. Warszawa: PWN.

- QUINE W. V. O. (1960), *Word and Object*. The Massachusetts Institute of Technology. Tłumaczenie polskie: *Słowo i rzecz* (1999) (przeł. B. Stanosz). Warszawa: Fundacja Aletheia, Wydawnictwo Spacja.
- REINER M., CAMP E. (2006), *Metaphor*, [w:] Lepore E., Smith B. (red.), *Handbook of Philosophy of Language*. Oxford University Press, 845–864.
- RICOEUR P. (1989), *Metafora i symbol*, [w:] *Język, tekst interpretacja*. Warszawa: PIW, 123–155.
- SEARLE J. R. (1979), *Metaphor*, [w:] *Expression and Meaning: Studies in the Theory of Speech Acts*. Cambridge: Cambridge University Press, 76–116.
- SZYMCZAK M. (red.) (2002), *Słownik języka polskiego*, t. 1–3. Warszawa: PWN.
- TARSKI A. (1933), *Pojęcie prawdy w językach nauk dedukcyjnych*, „Prace Towarzystwa Naukowego Warszawskiego, Wydział II Nauk Matematyczno-Fizycznych”, nr 34. Warszawa. Przedruk [w:] Tarski A. (1995), *Pisma logiczno-filozoficzne*. Warszawa: PWN, 9–172.
- WITTGENSTEIN L. (1953), *Philosophical Investigations*. Oxford: Blackwel Publishing. Wydanie polskie: *Dociekania filozoficzne* (2000) (przeł. B. Wolniewicz). Warszawa: PWN (wydanie II).

Rozdział 8

Debata między minimalizmem semantycznym a kontekstualizmem*

Joanna Odrowąż-Sypniewska (Uniwersytet Warszawski)

Słowa kluczowe: Bach, Borg, Cappelen, indeksykalizm, kontekstualizm, Lepore, minimalizm semantyczny, okazjonalność, Recanati, Searle, Salmon, Stanley, Travis, zależność od kontekstu

1. Wprowadzenie: semantyka a pragmatyka

W literaturze światowej bardzo żywo dyskutowane jest ostatnio zagadnienie podziału kompetencji między semantyką a pragmatyką. Spór dotyczący przebiegu granicy między tymi dwiema dziedzinami badań może wydawać się sporem czysto akademickim, dopóki nie uświadomimy sobie, że — przynajmniej przy pewnym rozumieniu — jest to spór o to, czemu przysługują warunki prawdziwości i — w konsekwencji — wartość logiczna.

Zgodnie z funkcjonującym od czasów Charlesa Morrisa (1938) podziałem, semantyka zajmuje się znaczeniem językowym, odniesieniem i warunkami prawdziwości wyrażen, a pragmatyka — użyciem wyrażen i tym, co nadawca miał na myśli, mówiąc to, co powiedział, czyli zamierzonym przez nadawcę znaczeniem wypowiedzi. Semantyka wyznacza warunki prawdziwości zdań, a pragmatyka zajmuje się znaczeniem konkretnych użyczeń zdań. To 'pragmatyczne' znaczenie wykracza poza znaczenie językowe i nie jest istotne dla wyznaczenia warunków prawdziwości. Będące nośnikami wartości logicznych sądy są wyrażane przez zdania oznajmujące i nie zależą od tego, co wypowiadający te zdania nadawca ma na myśli (tj. są niezależne od znaczenia mówiącego). Sądy wyrażają te zdania oznajmujące, które mają określony sens i z których usunięto wszelkie okazjonalizmy, wieloznaczności i nieostrości. Do wyrażen okazjonalnych zalicza się głównie wyrażenia z tak zwanej listy Kaplanowskiej („ja”, „ty”,

* Projekt został sfinansowany ze środków Narodowego Centrum Nauki (grant NN101 000940).

„on”, „ona”, „to”, „wczoraj”, „dziś”, „jutro”, „aktualny”, „teraźniejszy”) oraz takie rzeczowniki, jak: „wróg”, „obcy”, „obcokrajowiec”, „przyjaciół” i takie przymiotniki, jak: „lokalny”, „narodowy”, „zagraniczny”. Po usunięciu okazjonalizmów, nieostrości i wieloznaczności, otrzymujemy zdania ‘wieczne’, którym można ‘raz na zawsze’ przypisać określoną wartość logiczną. Treść semantyczna zdań jest funkcją stałych znaczeń wyrażen nieokazjonalnych, treści odniesieniowych wyrażen okazjonalnych oraz sposobu ich połączenia¹. Tradycyjnie rozumiana semantyka dostrzega udział kontekstu w ustalaniu wyrażanych sądów, ale zakłada, że rolą kontekstu jest jedynie determinowanie odniesienia wyrażen okazjonalnych i ewentualne usuwanie wieloznaczności leksykalnych (np. w zależności od kontekstu „Jan gra na rogu” może znaczyć, że Jan gra na rogu ulicy albo że Jan gra na dętym instrumencie muzycznym). Znaczenie kontekstu w pragmatyce jest natomiast bardzo szerokie i obejmuje wszystko to, co może być ważne przy ustalaniu intencji nadawcy.

Innymi słowy, zgodnie z poglądem tradycyjnym semantyka zajmuje się wyznaczaniem treści semantycznej zdań, a pragmatyka bada moc illokucyjną i perlokucyjną² wypowiedzi oraz generowane przez nadawcę implikatury konwersacyjne³. Za J. P. Grice’em przyjmuje się podział na to, co powiedziane (*what is said*) i to co dawane do zrozumienia (znaczenie mówiącego, *what is meant (but not said)*) i przyjmuje się, że to pierwsze jest domeną semantyki, a to drugie należy do pragmatyki. Załóżmy, że François Recanati na pytanie, czy umie gotować, odpowiada „Jestem Francuzem” (por. Recanati 2004: 5). Wtedy tym, co powiedziane w jego zdaniu, jest to, że Recanati jest Francuzem (inaczej mówiąc, zdanie przez niego wypowiedziane semantycznie wyraża sąd, że Recanati jest Francuzem). Oczywiście sama znajomość języka nie wystarczy do sformułowania tego sądu. Zdanie „Jestem Francuzem” zawiera podmiot domyślny „ja”, który jest wyrażeniem okazjonalnym. Reguły językowe wskazują, że do tego, aby ustalić, do czego „ja” odnosi się w danym kontekście, musimy wiedzieć, kto w tym kontekście jest nadawcą. Zatem do tego, żeby ustalić, że wypowiedziane przez Recanatiego zdanie „Jestem Francuzem” wyraża sąd, że François Recanati jest Francuzem, musimy ‘zajrzeć’ do kontekstu i zobaczyć, kto to zdanie wygłosił. Takie ‘zajrzenie’ mieści się jednak w ramach semantyki, bowiem to reguły językowe wskazują, kiedy i gdzie należy się do kontekstu odwołać. To, co powiedziane, odbiega od znaczenia konwencjonalnego tylko

¹ Bardziej radykalnym stanowiskiem będzie pogląd głoszący, że „semantyczna interpretacja wyrażenia złożonego w jest rezultatem złożenia stałych (*standing*) znaczeń elementów leksykalnych w zgodzie z semantycznymi regułami składania odpowiadającymi strukturze syntaktycznej w” (Stanley 2007: 137). W myśl takiej koncepcji semantyczną treść mają tylko zdania-typy, a zdanie „Ja jestem synem prezydenta” nie wyraża sądu, ponieważ nie można mu przypisać wartości logicznej. Co do stałych znaczeń por. przyp. 18.

² Rozróżnienie czynności lokucyjnych, illokucyjnych i perlokucyjnych zaproponował J. L. Austin (zob. *How to do things with words*, Oxford University Press 1962).

³ Pojęcie *implikatur konwersacyjnych* wprowadził J. P. Grice w słynnym artykule *Logic and Conversation* (1975).

minimalnie; tylko wtedy, gdy konieczne jest uzupełnienie znaczenia zdania i spowodowanie, żeby wyrażało ono sąd (zob. Recanati 2004: 7). Akceptacja tak rozumianej *zasady minimalizmu* jest cechą charakterystyczną tradycyjnie pojmowanej semantyki. Jedynym procesem pragmatycznym uznawanym przez tak rozumianą semantykę jest nasycanie (saturacja), która polega na kontekstowym przypisywaniu wartości semantycznych wyrażeniom okazjonalnym. Saturacja jest kontrolowana przez reguły językowe, bo to język sam wskazuje, kiedy należy się do kontekstu odwołać.

Mówiąc w powyższym kontekście, że jest Francuzem, Recanati daje oczywiście do zrozumienia, że świetnie gotuje. Pozytywna odpowiedź na pytanie „Czy umiesz gotować?” jest przez Recanatiego implikowana konwersacyjnie. Zadaniem pragmatyki jest analiza sposobu, w jaki rozmówca Recanatiego odgaduje, że odpowiedź „Jestem Francuzem” jest w tym kontekście równoznaczna z odpowiedzią „Tak, umiem gotować”.

Ten tradycyjny podział zadań między semantyką a pragmatyką został ostatnio zakwestionowany. Zwrócono bowiem uwagę, że wpływ kontekstu na znaczenie zdań jest znacznie większy niż to zakładano. W ostatnich latach często twierdzić, że zależność kontekstowa jest zjawiskiem powszechnym i wyrażenia od kontekstu zależne stanowią większość (a nie mniejszość, jak do tej pory sądzono) wyrażen języka naturalnego. Jeśli wpływ kontekstu nie zostanie należycie uwzględniony, to nie da się zdaniom (ani ich wypowiedzeniom) przypisać właściwej wartości logicznej. Twierdzi się, na przykład, że wszystkie zdania:

- (1) Wszyscy zdali egzamin;
- (2) Jan (który ma 185 cm wzrostu) jest wysoki;
- (3) To jabłko jest czerwone;
- (4) Wiem, że samolot wyląduje w Chicago;
- (5) Nina jest gotowa;
- (6) Te liście są zielone;

w sposób istotny zależą od kontekstu. Jeśli chodzi o zdanie (1), to kontekst determinuje uniwersum dyskursu⁴. Jeśli uniwersum dyskursu stanowią studenci, którzy zapisali się na kurs, to zdanie (1) wyraża sąd, że wszyscy studenci, którzy zapisali się na kurs, zdali egzamin. Jeśli zaś uniwersum dyskursu stanowią studenci, którzy chodzili na zajęcia w określonej grupie, to zdanie (1) wyraża sąd, że wszyscy studenci, którzy chodzili na zajęcia z daną grupą, zdali egzamin. Zdanie (2) — w zależności od kontekstu — może znaczyć na przykład, że Jan jest wysoki jak na piłkarza lub że Jan jest wysoki jak na koszykarza. Zdanie (3) może być prawdziwe, gdy wskazywane jabłko ma czerwoną skórę

⁴ Kontekst determinuje także to, o jaki egzamin chodzi, ale tę zależność kontekstową tutaj pomijam.

lub — w szczególnym kontekście, w którym mowa o jabłkach zaatakowanych przez grzyb, który zabarwia ich miąższ — gdy ma czerwony miąższ.

Stephen Cohen opowiada następującą historię, mającą nam uzmysłowić zależność kontekstową czasownika „wiedzieć” (1999: 59). Załóżmy, że Maria i Jan lecą z Los Angeles do Nowego Jorku. Ich samolot ma po drodze lądować w Chicago. Marii i Janowi bardzo na tym zależy, ponieważ mają tam załatwić ważną sprawę. Chcą się więc upewnić, czy rzeczywiście samolot będzie tam lądował. Stają w kolejce do informacji i czekając, pytają pana Smitha, który stoi za nimi. Smith zagląda do rozkładu lotów, który trzyma w ręku i mówi: „Wiem, że samolot wyląduje w Chicago”. Maria mówi jednak: „W rozkładzie może być błąd. Zatem Smith nie wie, czy samolot wyląduje w Chicago”. Potem okazuje się, że samolot rzeczywiście miał międzylądowanie w Chicago. Można uważać, że Smith w swoim kontekście, w którym obowiązują stosunkowo niskie standardy epistemiczne (Smithowi nie zależy na międzylądowaniu w Chicago), *wie*, że samolot wyląduje w Chicago. Jednakże w kontekście Marii i Jana, w którym obowiązują wysokie standardy epistemiczne (bo Marii i Janowi bardzo zależy na tym międzylądowaniu), Smith *nie wie*, czy samolot wyląduje w Chicago.

Zdanie (5) w jednym kontekście może znaczyć, że Nina jest gotowa do egzaminu, a w innym, że jest gotowa do wyjścia na imprezę. Przykład (6) „Te liście są zielone” pochodzi od Charlesa Trávisa (1997). Załóżmy, że klon Pii ma brązowe liście. Ponieważ Pia uważa, że liście powinny być zielone, przemalowuje wszystkie liście na zielono. Mówi „Teraz jest lepiej. Teraz te liście są zielone”. Chwilę potem przyjaciel botanik przychodzi do Pii i mówi, że poszukuje zielonych liści do badań. Pia odpowiada „Te liście są zielone, możesz je wziąć”. Botanik jednak twierdzi, że liście nie są zielone.

Przykłady te mają pokazać, że wartość logiczna zdań zmienia się wraz z kontekstem, w jakim są wypowiedziane (przykłady 1, 3, 4, 6) albo że zdanie nie posiada żadnej wartości, dopóki nie zostanie przez kontekst uzupełnione (przykłady 2, 5)⁵. Można ich również użyć do pokazania, że warunki prawdziwości przypisane przez semantykę rozmiągają się z warunkami intuicyjnymi. Na przykład, gdybyśmy chcieli się odwoływać do semantycznie określonych warunków prawdziwości dla zdania „Wszyscy zdali egzamin”, to musielibyśmy uznać, że to zdanie głosi, że wszyscy ludzie zdali egzamin⁶ i — co za tym idzie — musielibyśmy uznać je za fałszywe. Tymczasem w wielu kontekstach intuicyjnie uznajemy takie zdanie za prawdziwe, ponieważ odpowiednio zawężamy dziedzinę kwantyfikacji.

Widać wyraźnie, że zależność kontekstowa wymienionych zdań wykracza poza przypisanie wyrażeniom okazjonalnym ich odniesień i usunięcie wielo-

⁵ Niektórzy kwestionują takie postawienie sprawy. Na przykład Borg oraz Cappelen i Lepore uważają, że wszystkie wymienione wyżej zdania wyrażają sądy i uzupełnianie ich przez kontekst nie jest potrzebne (zob. Cappelen, Lepore 2005 oraz Borg 2004).

⁶ Dla uproszczenia zakładam tutaj, że tylko ludzie mogą zdawać egzaminy.

znaczności leksykalnej. Co więcej, twierdzi się, że te zdania nie są w żadnym sensie wyjątkowe. Są typowymi przykładami zdań języka naturalnego i dlatego, odwołując się do nich, można wykazać, że zależność kontekstowa jest znacznie bardziej rozpowszechniona niż dotychczas sądzono. Jak widzieliśmy, w tradycyjnie rozumianych rozważaniach semantycznych kontekst może być uwzględniony tylko wtedy, gdy legitymizują to reguły znaczeniowe (na przykład, gdy zdanie zawiera wyrażenie okazjonalne). Są zatem tutaj trzy wyjścia: albo trzeba postulować istnienie w strukturze logicznej zdań takich składników, które pozwalałyby na odwołanie się do kontekstu (tj. postulować istnienie *ukrytych* okazjonalizmów), albo trzeba uznać, że nieuwzględniająca w wystarczającej mierze kontekstu semantyka co prawda determinuje wartości logiczne zdań, ale wyznaczone przez nią wartości nie są tymi, które rozmówcy intuicyjnie zdaniom przypisują, albo wreszcie trzeba przyznać, że semantyka nie jest w stanie wykonać swego podstawowego zadania, ponieważ nie jest w stanie wyznaczyć warunków prawdziwości większości zdań. Pierwszy z tych poglądów jest głoszony przez indeksykalistów, którzy doceniają nieredukowalność zależności kontekstowej wyrażeń, ale wpływ kontekstu dopuszczają tylko tam, gdzie wyrażnie wskazują na to reguły językowe. Indeksykaliści — których najbardziej znanym przedstawicielem jest Jason Stanley — przyznają, że wpływ kontekstu jest znacznie większy niż się wydawało tradycyjnym semantynom, ale — tak jak oni — uważają, że ten wpływ jest kontrolowany przez reguły językowe. Indeksykaliści całą zależność kontekstową sprowadzają do okazjonalności. Okazjonalność jest zjawiskiem już dawno ‘oswojonym’, więc indeksykaliści muszą nas jedynie przekonać, że wyrażeń okazjonalnych jest znacznie więcej niż nam się do tej pory wydawało. Zwolennicy drugiego podejścia (minimaliści, tacy jak Emma Borg oraz Herman Cappelen i Ernie Lepore⁷) wyrażnie oddzielają dwa rodzaje treści związane ze zdaniem: treść zdania (sąd minimalny) oraz treść komunikowaną. Sądem minimalnym wyrażanym przez zdanie „Jan zamówił kawę” jest sąd *to, że Jan zamówił kawę*, natomiast sądem komunikowanym przez to zdanie wypowiedziane w kawiarni jest na przykład *to, że Jan zamówił filiżankę kawy* (a nie worek czy kroplę kawy). Zwolennikami trzeciego podejścia są kontekstualiści, wśród których można wyróżnić kontekstualistów umiarkowanych i kontekstualistów radykalnych. Kontekstualiści umiarkowani (tacy jak Recanati) twierdzą, że pojęcie sądu *minimalnego* jest teoretycznie bezwartościowe. Sądy minimalne można wyróżnić, ale one do niczego nie służą; w szczególności nie pełnią żadnej funkcji w komunikacji. Odbiorca, który słyszy wypowiedziane w kawiarni zdanie „Zamówiłem kawę”, od razu ‘chwytą’ sąd komunikowany *Zamówiłem filiżankę kawy* i nie potrzebuje pośrednictwa sądu minimalnego. Kontekstualiści radykalni (tacy jak Travis) negują samo istnienie sądów minimalnych. Według nich

⁷ Recanati nazywa takie stanowisko „synkretyzmem” (2004).

zdania jako takie w ogóle nie wyrażają sądów i nie są ani prawdziwe, ani fałszywe. Warunki prawdziwości posiadają tylko wypowiedzi — zdania wypowiedziane w konkretnych okolicznościach, bowiem tylko w kontekście konkretnego aktu mowy zdanie ma wystarczająco określoną treść, żeby można było ją ocenić pod względem wartości logicznej. Bez informacji dotyczącej aktu mowy, w którym zdanie zostało użyte, nie jesteśmy w stanie orzec, czy zdanie to jest prawdziwe, czy fałszywe. Kontekstualiści twierdzą zatem, że tylko konkretne wypowiedzenia zdań mają wartość logiczną. Żeby tę wartość określić, musimy najpierw zinterpretować zdanie pragmatycznie — bez takiej interpretacji zdanie nie wyraża żadnego sądu (nawet minimalnego sądu semantycznego)^{8,9}.

Jeśli kontekstualiści mają rację, to zadania semantyki są znacznie skromniejsze niż się do tej pory wydawało. Semantycy powinni koncentrować się na znaczeniu leksykalnym wyrażen i określaniu ich odniesienia, nie do nich należy natomiast wskazywanie, jakie sądy zostały wyrażone i jakie są ich warunki prawdziwości. Tych ostatnich zadań nie da się wykonać bez uwzględnienia badań pragmatycznych. Tym, co wyraża sądy, są konkretne akty mowy, a nie zdania (nawet nie zdania-egzemplarze). To do pragmatyków należy zatem określanie wartości logicznej określonych wypowiedzeń. Co więcej — zdaniem zwolenników kontekstualizmu — nawet na określanie znaczenia leksykalnego wyrażen mają wpływ badania pragmatyczne. Kontekstualiści, tacy jak Recanati, twierdzą bowiem, iż wyrażenie „otworzyć” znaczy co innego w kontekście „otworzyć drzwi”, a co innego w kontekście „otworzyć butelkę” czy „otworzyć książkę”. Ponieważ istotnym elementem kontekstu są intencje i zachowania użytkowników języka, znaczenie i warunki prawdziwości nie są determinowane przez reguły językowe. W szczególności, kontekstualiści twierdzą, że T-równoważności w rodzaju „Śnieg jest biały’ jest prawdziwe ztw, gdy śnieg jest biały” są fałszywe, ponieważ nie uwzględniają zależności kontekstowej. Warunki prawdziwości determinowane są w trakcie procesów pragmatycznych, a nie semantycznych. Dlatego niektórzy kontekstualiści proponują zastąpienie semantyki pragmatyką warunków prawdziwości (*truth-conditional pragmatics*).

2. Stanowiska w sporze: próba charakterystyki

Żadnej z osób zabierających obecnie głos w debacie nie można zaliczyć do zwolenników stanowiska tradycyjnego w czystej postaci. Recanati to stanowisko nazywa „literalizmem” lub „minimalizmem semantycznym”. Ja zarezerwuję dla tego podejścia nazwę „literalizm”, natomiast „minimalistycznymi” będę nazywała

⁸ Mówi się tutaj o „pragmatycznym błędnym kole” (*the pragmatic circle*): Grice twierdził, że do ustalenia zamierzonego znaczenia nadawcy trzeba najpierw ustalić to, co powiedziane. Jeśli do ustalenia tego, co powiedziane, konieczne jest odwołanie się do kontekstu, to mamy błędne koło (zob. Korta, Perry 2007). Por. tutaj jednak Bach: *The top 10 misconceptions about implicature*.

⁹ Jest jeszcze podejście pośrednie między drugim i trzecim, którego głównym reprezentantem jest Kent Bach (zob. niżej, pkt 3).

stanowiska, które wyewoluowały z literalizmu pod wpływem argumentów kontekstualistów. Debata, która jest przedmiotem niniejszego eseju, jest sporem między minimalistami a kontekstualistami.

2.1. Zasady podziału

Każdy z uczestników debaty definiuje minimalizm i kontekstualizm na własny użytek i to najczęściej nie przejmując się tym, jak robią to inni. Tymczasem widać wyraźnie, że poszczególni teoretycy stosują różne zasady podziału i w konsekwencji ich klasyfikacje też są różne. Co ciekawe, są tacy filozofowie, którzy przy każdej próbie klasyfikacji wpadają do przegródki z napisem „kontekstualizm”, natomiast po drugiej stronie sporu panuje znacznie większe zmieszanie i nie ma takiej osoby, która zawsze — niezależnie od wybranej zasady podziału — okazywałaby się minimalistą. Przyjrzyjmy się teraz różnym sposobom interpretacji debaty minimalizm — kontekstualizm i związanymi z nimi różnymi charakterystykami stanowisk w sporze.

2.1.1. Propozycjonalność treści semantycznej

Pierwszy — wskazany powyżej — sposób interpretacji sprowadza omawianą debatę do sporu między tymi, którzy sądzą, że nośnikami wartości logicznych są zdania, a tymi, którzy uważają, że prawdziwe i fałszywe mogą być tylko konkretne wypowiedzenia. Ci drudzy uważają zazwyczaj, że dla większości zdań semantyka nie jest w stanie dostarczyć takich treści, które można oceniać pod względem prawdziwości lub fałszywości. Zdania języka naturalnego są bowiem zazwyczaj niekompletne, wskutek czego ich znaczenie jest niedookreślone i nie nadaje się do oceny pod względem wartości logicznej. W tym sensie kontekstualistami są Charles Travis oraz zwolennicy teorii relewancji: Robyn Carston, Dan Sperber i Deirde Wilson. Kontekstualistą jest także Kent Bach, który twierdzi, że analiza semantyczna często dostarcza tylko nieposiadających warunków prawdziwości rdzeni sądów. Minimalistami w tym rozumieniu są m.in. Emma Borg, Stanley, Herman Cappelen i Ernest Lepore, a także postulujący sądy zwrotne John Perry i Kepa Korta¹⁰.

2.1.2. Intuicyjne warunki prawdziwości

Inaczej debatę minimalizm — kontekstualizm można interpretować jako spór dotyczący intuicyjnych warunków prawdziwości wypowiedzi: tego, co intuicyjnie powiedziane. Innymi słowy, chodzi o to, czy treść semantyczna jest tym,

¹⁰ Pojęcie sądu zwrotnego (*reflexive proposition*) wprowadził John Perry w *The Problem of the Essential Indexical and Other Essays* (1993), gdzie powołuje się na Reichenbacha *Elements of Symbolic Logic* (zob. też *Reference and Reflexivity* 2001). Sąd jest zwrotny, gdy w swojej treści zawiera swoją nazwę. Sądem zwrotnym wyrażanym przez zdanie z „Ja jestem Polką” jest: nadawca z jest Polką.

co stwierdzane przez nadawcę wypowiedzi zdaniowej. Przy takiej interpretacji linia podziału między minimalistami a kontekstualistami przebiega w innym miejscu. Minimalistami są bowiem ci, którzy uznają, że podanie intuicyjnych warunków prawdziwości jest rzeczą semantyki, a kontekstualistami ci, którzy uważają, że to zadanie należy do pragmatyki. Tak debatę minimalizm — kontekstualizm widzi Recanati w swojej najnowszej książce *Truth-Conditional Pragmatics* i tak interpretuje ją Stanley. Stanley jest w tym sensie minimalistą, ale minimalistką w tym rozumieniu nie jest już autorka książki *Minimal Semantics* Borg, bowiem jej zdaniem to, co powiedziane, jest rzeczą pragmatyki. Również przy tym rozumieniu Bach okazuje się kontekstualistą. Kontekstualistami są także autorzy książki *Insensitive Semantics. A Defense of Semantic Minimalism and Speech Acts Pluralism*, Cappelen i Lepore, którzy wyraźnie twierdzą — wbrew Grice'owi — że nie należy utożsamiać sądu minimalnego z tym, co powiedziane (zob. Cappelen, Lepore 2005: 150).

Te dwa sposoby interpretacji debaty można potraktować jako dwie niezależne osie sporu i skrzyżować ze sobą¹¹. Jedną oś stanowi pytanie, czy semantyka dostarcza treści propozycjonalnych (to jest takich, które można ocenić pod względem wartości logicznej), a drugą pytanie, czy semantyka dostarcza intuicyjnych warunków prawdziwości. Twierdząco na oba pytania odpowiada Stanley. Bach, Borg, Cappelen i Lepore oraz Recanati odpowiadają twierdząco na pytanie pierwsze, ale przecząco na drugie. Radykalni kontekstualiści, tacy jak John Searle i Travis, odpowiedzą negatywnie na oba pytania. Tych, którzy pozytywnie odpowiadają na pytanie pierwsze, można jeszcze podzielić ze względu na to, czy uważają, że wszystkie syntaktycznie poprawne zdania oznajmujące mają treść propozycjonalną: Borg, Stanley, Cappelen i Lepore uważają, że tak, a Bach i Recanati uważają, że nie. Przy czym Recanati sądzi, że zdań, które mają semantyczną treść propozycjonalną jest bardzo niewiele, a Bach uważa, że jest ich całkiem sporo.

2.1.3. Ilość wyrażenń zależnych od kontekstu

Jeszcze inaczej debatę między minimalizmem a kontekstualizmem interpretują Cappelen i Lepore, według których minimalistami są ci, którzy postulują bardzo ograniczoną liczbę wyrażenń kontekstowo zależnych, a radykalnymi kontekstualistami ci, którzy twierdzą, że wszystkie wyrażenia są zależne od kontekstu¹². Cappelen i Lepore wymieniają jeszcze umiarkowany kontekstualizm,

¹¹ Otrzymany w efekcie podział będzie podziałem zależnym, ponieważ nie może być ko- goś, kto uważa zarazem, że nie ma propozycjonalnej treści semantycznej, jak i że semantyka wyznacza intuicyjne warunki prawdziwości.

¹² Cappelen i Lepore zamiennie mówią o okazjonalności (*indexicality*) i zależności od kontekstu (*context-sensitivity*). Ja „okazjonalność” rozumiem tutaj wąsko, jako zależność od kontekstu kontrolowaną językowo. Tak rozumiana okazjonalność jest mniejszym wyzwaniem dla minimalistów niż zależność kontekstowa innego rodzaju (niekontrolowana semantycznie).

według którego bardzo wiele pozornie niezależnych od kontekstu wyrażeń w istocie od kontekstu zależy. Autorzy *Insensitive Semantics* uważają jednak, że umiarkowany kontekstualizm jest stanowiskiem niestabilnym, ponieważ każdy, kto robi krok w stronę kontekstualizmu, po równi pochyłej 'zjedzie' do kontekstualizmu radykalnego.

Cappelen i Lepore traktują jako cechę definiującą kontekstualizm radykalny to, że żadne zdanie języka naturalnego semantycznie nie wyraża sądu¹³. Skoro bowiem wszystkie wyrażenia są zależne od kontekstu, to żadne zdanie składające się z tych wyrażeń nie będzie miało niezależnych od kontekstu warunków prawdziwości¹⁴. Kontekstualizm umiarkowany nie jest aż tak radykalny i jego zwolennicy sądzą, że są takie zdania, które semantycznie wyrażają sądy. Zwolennicy kontekstualizmu umiarkowanego korzystają — zdaniem Cappelena i Lepore'a — ze strategii niespodziewanego wyrażenia kontekstowo zależnego, strategii ukrytego wyrażenia kontekstowo zależnego oraz strategii niewyartykułowanego składnika. Strategia niespodziewanego wyrażenia kontekstowo zależnego polega na twierdzeniu, że jakieś wyrażenie powszechnie do tej pory uważane za niezależne od kontekstu (np. „wie, że”) należy jednak uznać za kontekstowo zależne. Strategia ukrytego wyrażenia kontekstowo zależnego polega na postulowaniu niewidocznego w strukturze powierzchniowej, ale obecnego w formie logicznej zdania wyrażenia kontekstowo zależnego. Zwolennicy tej strategii twierdzą np., że prawdziwą strukturą zdania „Jan jest wysoki”, jest „Jan jest wysoki *jak na X*”¹⁵. Strategia niewyartykułowanego składnika polega na twierdzeniu, że do wyrażanych treści należy dodawać determinowane kontekstowo składniki, mimo że nie odpowiada im żaden element składni (jawny ani ukryty)¹⁶. Recanati na przykład uważa, że w sądzie wyrażonym przez zdanie „Pada deszcz” znajduje się miejsce, o którym się mówi, że w nim pada deszcz, mimo iż składnia tego zdania tego miejsca nie uwzględnia¹⁷.

¹³ Stanowisko takie przypisują m.in. Wittgensteinowi, Austinowi, zwolennikom teorii relewancji, a także Searle'owi i Travisowi (zob. 2005: 6).

¹⁴ Cappelen i Lepore twierdzą, że kontekstualizm radykalny, po pierwsze, nie ma racji co do powszechności zależności kontekstowej w języku (co mają wykazać ich testy na zależność kontekstową — zob. 2005; rozdz. 7), po drugie, czyni komunikację niemożliwą (ponieważ treść wszystkich wyrażeń zmienia się wraz z kontekstem, to rozmówcy nigdy nie wiedzą o czym mówią — zob. 2005; rozdz. 8), a po trzecie, jest stanowiskiem wewnętrznie sprzecznym (ponieważ implikuje, że są prawdziwe wypowiedzenia zdania „Kontekstualizm radykalny jest fałszywy” — zob. 2005; rozdz. 9).

¹⁵ Można powątpiewać, czy rzeczywiście strategia ukrytego okazjonalizmu jest strategią kontekstualistyczną. Do takich ukrytych okazjonalizmów odwołuje się bowiem i Stanley, który jest indeksykalistą, i Borg, która jest minimalistką.

¹⁶ Stanley podaje następującą definicję: „*x* jest niewyartykułowanym składnikiem wypowiedzi w ztw, gdy (1) *x* jest elementem dostarczonym przez kontekst do *w* i (2) *x* nie jest wartością semantyczną żadnego składnika formy logicznej wypowiedzianego zdania” (Stanley 2007: 47).

¹⁷ Prawidłowa analiza zdania „Pada deszcz” budzi jednak wiele kontrowersji i jest przedmiotem ożywionej debaty (zob. np. Perry 1986, Stanley 2007: 218 i Recanati 2010, rozdz. 3).

Borg (2006) zauważa, że niektórzy zabierający głos w debacie kontekstualiści odwołują się do argumentu z niewłaściwości warunków prawdziwości determinowanych przez semantykę, a to zakłada, że semantyka jakieś warunki prawdziwości jednak determinuje. Jak już wspominałam, gdybyśmy chcieli się odwoływać do semantycznie określonych warunków prawdziwości dla zdania „Wszyscy zdali egzamin”, to musielibyśmy uznać, że to zdanie jest fałszywe. Tymczasem w wielu kontekstach intuicyjnie uznajemy takie zdanie za prawdziwe (ponieważ odpowiednio zawężamy dziedzinę kwantyfikacji). Ci, którzy używają argumentów z niewłaściwości semantycznych warunków prawdziwości, muszą zatem przyznać, że nie jest tak, że żadne zdanie nie ma semantycznie determinowanych warunków prawdziwości. Są takie zdania, które te warunki mają, tyle że są to zazwyczaj warunki niewłaściwe, tj. niezgodne z intencjami rozmówców. Żaden kontekstualista, który odwołuje się do argumentu z niewłaściwości, nie jest więc kontekstualistą radykalnym w sensie Cappelena i Lepore’a.

Na dokonany przeze mnie powyżej podział można teraz jeszcze nałożyć kryterium postulowanej ilości wyrażenń zależnych od kontekstu. Do tych, którzy uważają, że jest ich mało, należą Borg, Cappelen i Lepore oraz Bach, a do tych, którzy znacznie poszerzają ich liczbę, należą Stanley i Recanati. Przy czym Stanley i Recanati zupełnie inaczej interpretują takie niejawne wyrażenia kontekstowo zależne. Stanley uważa, iż są to ukryte okazjonalizmy (bo w głębokiej strukturze syntaktycznej zdania widać, że posiadają zmienne, które muszą zostać uzupełnione przez kontekst), Recanati zaś sądzi, że ich zależność kontekstowa jest kontrolowana przez kontekst, a nie przez reguły językowe.

2.1.4. Silne i słabe wpływy pragmatyczne

King i Stanley debatę między minimalizmem a kontekstualizmem interpretują jako spór o dopuszczalność silnych i słabych wpływów pragmatycznych, gdzie:

(s)łaby wpływ pragmatyczny na to, co jest komunikowane w wypowiedzi, zachodzi w wypadku, w którym kontekst (włączając intencje nadawcy) wyznacza interpretację elementu leksykalnego w zgodzie ze stałym znaczeniem¹⁸ tego elementu leksykalnego. Silny wpływ pragmatyczny na to, co komunikowane, jest takim wpływem kontekstu na to, co komunikowane, który nie jest jedynie pragmatyczny w słabym sensie (King, Stanley 2005/2007: 140).

2.2. Proponowana klasyfikacja

Proponuję następującą klasyfikację. *Minimalizm simplicitate* to stanowisko, zgodnie z którym:

¹⁸ Stałe znaczenie to znaczenie wyrażenia-typu przypisane mu przez konwencje językowe (por. Recanati 2010: 33).

- (T1) tylko niektóre zdania oznajmujące mają propozycjonalną treść semantyczną;
- (T2) wyrażenia zależne od kontekstu tworzą dość wąski zbiór (pokrywający się mniej więcej z listą Kaplana — por. wyżej); oraz
- (T3) treść semantyczna podlega jedynie słabym wpływom semantycznym (w sensie Kinga i Stanleya).

Przedstawicielem minimalizmu *simplicite* w tym sensie jest Bach. *Minimalizm radykalny* to stanowisko akceptujące (T2), (T3) oraz twierdzenie głoszące, że

- (T4) wszystkie syntaktycznie poprawne zdania oznajmujące semantycznie wyrażają sądy.

Radykalnymi minimalistami są w tym rozumieniu Borg oraz Cappelen i Lepore. Przy czym Borg twierdzi, że przy ustalaniu semantycznej treści nie możemy odwoływać się do intencji, a Cappelen i Lepore dopuszczają takie odwołania. Natomiast *minimalista umiarkowany* (indeksykalista) to ktoś, kto — tak jak Stanley — akceptuje (T3) i (T4), ale nie akceptuje (T2). *Umiarkowany kontekstualista* akceptuje (T1), ale nie akceptuje ani (T2), ani (T3), ani (T4). Takim umiarkowanym kontekstualistą jest Recanati. *Radykalni kontekstualiści*, do których zaliczają się Travis i Searle, odrzucają wszystkie tezy (T1) — (T4). Cechą charakterystyczną kontekstualizmu jest zatem odrzucenie tez (T2) i (T3).

Moja charakterystyka powyższych stanowisk odbiega istotnie od tej zaproponowanej przez Recanatiego w *Literal Meaning* (2004). Recanati wyróżnia tam minimalistów (literalistów), którzy kierują się zasadą minimalizmu, synkretystów (którzy poza treścią semantyczną postulują treść komunikowaną), indeksykalistów, *quasi*-kontekstualistów i kontekstualistów radykalnych. Zaproponowana przeze mnie klasyfikacja zgadza się z klasyfikacją Recanatiego tylko w wypadku kontekstualistów. Problem z jego klasyfikacją polega na tym, że wśród osób obecnie biorących udział w debacie nie ma nikogo, kogo można by zaliczyć do minimalistów (jedynym kandydatem jest Stanley, ale on został już przypisany do indeksykalistów), a wszystkich tych, którzy sami siebie określają jako minimalistów (Borg, Bach, Cappelen i Lepore) trzeba zaliczyć do zwolenników stanowiska synkretycznego.

W *Truth-Conditional Pragmatics* (2010) Recanati pisze natomiast, że jedynym interesującym minimalizmem jest ten, który postuluje, że intuicyjne warunki prawdziwości są generowane semantycznie. Moim zdaniem, jedynym minimalistą spełniającym kryteria z *Truth-Conditional Pragmatics* jest znowu indeksykalista Stanley. Poza Stanleyem nikt z biorących udział w debacie bowiem nie twierdzi, że semantyka dostarcza intuicyjnych warunków prawdziwości wypowiedzi. Dlatego proponuję własną klasyfikację stanowisk. W mojej klasyfikacji ci, którzy sami uważają się za minimalistów, są przypisani do minimalizmu, a Stanley jest traktowany jako minimalista umiarkowany, co wydaje

mi się zgodne z intuicjami, bowiem z jednej strony Stanley akceptuje tezę (T3) mówiącą o kontroli języka nad wpływem kontekstu, ale z drugiej strony zależność od kontekstu widzi niemal wszędzie. Mojej klasyfikacji można natomiast postawić zarzut, że rozmija się z intuicjami Bacha, który uważa siebie za minimalistę bardziej radykalnego niż Cappelen i Lepore z tego względu, że odrzuca propozycjonalizm. Ten zarzut jest jednak łatwy do odparcia. Po pierwsze, Bach jest jedynym, który określa swoje stanowisko jako radykalny minimalizm. Recanati zalicza go do synkretystów, a Cappelen i Lepore — wręcz do kontekstualistów. Po drugie, stanowisko, które głosi, że *wszystkie* zdania wyrażają semantycznie sądy wydaje się bardziej radykalne semantycznie niż stanowisko, które mówi, że *tylko niektóre* zdania to robią. Poważniejszy zarzut dotyczy potraktowania przeze mnie założenia dotyczącego tego, czy semantyka dostarcza intuicyjnych warunków prawdziwości. Jak już wspominałam, Recanati uważa, że jedyną interesującą wersją minimalizmu jest ta, która twierdzi, że tak jest. Borg uważa natomiast, że cechą charakterystyczną minimalizmu jest odrzucenie tej tezy. U mnie Stanley, który tę tezę akceptuje, jest minimalistą umiarkowanym, a minimalistami *simplicite* są ci, którzy ją odrzucają. Przyznaję, że stanowisko, którego twierdzeniem jest, że semantyka dostarcza intuicyjnych warunków prawdziwości jest *radykalniejsze* niż stanowisko, które to odrzuca. Powodem, dla którego nie wpisałam tej tezy do cech radykalnego minimalizmu ani nawet minimalizmu *simplicite*, jest to, że nikt — poza Stanleyem — takiego poglądu nie broni. Stanleya zaś zaliczam do minimalistów umiarkowanych nie ze względu na to, że tę tezę głosi, ale ze względu na to, że nie głosi innej tezy (mianowicie tezy o ograniczonej ilości wyrażen kontekstowo zależnych). Stanley różni się od minimalistów *simplicite* i radykalnych także tym, że nie postuluje istnienia minimalnej treści semantycznej. Dla niego treść semantyczna jest bowiem bogatą treścią intuicyjną uzupełnianą przez kontekst pod kontrolą języka.

3. Minimalizm Bacha

Nie ma tu niestety miejsca na szczegółowe omówienie poszczególnych odmian minimalizmu, dlatego dokładniej przedstawię tylko stanowisko Bacha, które uważam za najbardziej obiecujące. Bach uważa za oczywiste, że to, co nadawca chce przekazać za pomocą swojej wypowiedzi, wykracza poza to, co wypowiedziane przez niego zdanie znaczy, nawet jeśli zdanie to nie jest wieloznaczne ani nieostre i nie zawiera wyrażen okazjonalnych. Bach nazywa to twierdzenie kontekstualnym banałem (*contextualist platitude*), ale podkreśla, że przyjęcie go nie prowadzi wcale do kontekstualizmu. Jego zdaniem banał kontekstualny jest do pogodzenia z minimalizmem semantycznym i twierdzeniem, że znaczenie zdania jest wyznaczane na podstawie znaczenia jego składników i sposobu ich połączenia.

Bach od kontekstualistów różni się tym, że nie uważa, żeby kontekst miał wpływ na treść semantyczną (poza oczywistymi przypadkami wyrażen okazjonalnych), a od minimalistów Borg, Cappelena i Lepore'a tym, że odrzuca propozycjonalizm, czyli założenie, że każde pozbawione okazjonalizmów zdanie oznajmujące wyraża sąd^{19, 20}. Zdaniem Bacha, należy odróżnić zdania zależne od kontekstu od zdań semantycznie niepełnych (*semantically incomplete*). Te drugie mają treść, która nie jest sądem i nie da się jej przypisać warunków prawdziwości. Bach twierdzi, że zdania takie, jak „Piotr jest gotowy” lub „Paweł ma już dosyć” wyrażają rdzenie sądów, a nie całe sądy. Założenie o istnieniu niepełności semantycznej nie prowadzi do kontekstualizmu pod warunkiem, że nie zakłada się propozycjonalizmu i dopuszcza, że niektóre zdania mogą nie wyrażać sądów.

W zależności od tego, czy wyrażają sądy, czy tylko rdzenie sądów, zdania dzielą się na pełne (*complete*) i niepełne (*incomplete*). Bach przyznaje, że wskazanie, które zdania są pełne, a które nie, nie jest sprawą łatwą, ale nie znaczy to, że nie ma takiej dystynkcji. Nie dysponujemy ogólnym kryterium odróżniania jednych zdań od drugih, ale możemy to rozstrzygać dla poszczególnych przypadków. Bach zauważa, że z punktu widzenia komunikacji to, czy zdanie wyraża sąd minimalny, czy rdzeń sądu, ma bardzo niewielkie znaczenie. Tak czy inaczej, odbiorca musi wzbogacić treść semantyczną, żeby odebrać treść, którą nadawca chciał przekazać²¹. Jego zdaniem nie ma zatem co kruszyć kopii o utrzymanie propozycjonalizmu. Jeśli zaś chodzi o przydatność treści semantycznych, to Bach twierdzi, że treść semantyczna jest zawarta w informacji przekazywanej przez zdanie. Jest dostępna dla słuchacza nawet wtedy, gdy ten z niej nie korzysta. Słuchacz mógłby tę treść 'obliczyć', nawet jeśli w danej sytuacji tego nie robi.

Żeby przejść od tego, co zostało powiedziane, do tego, co nadawca chciał zakomunikować, najczęściej trzeba dokonać uzupełnienia (*completion*) albo rozwinęcia (*expansion*). Uzupełnienie jest konieczne, jeśli wypowiedziane zdanie nie było pełne. Jeśli zaś nadawca wypowiedział zdanie pełne, ale to, co chciał przekazać, jest bardziej rozwinięte lub bardziej szczegółowe, to odbiorca powinien dokonać rozwinęcia. Uzupełnienie i rozwiniecie to dwa procesy, które Bach określa mianem implisytury (*implicature*)²². U Bacha można zatem

¹⁹ Trzeba tu wspomnieć, że Cappelen i Lepore w odpowiedzi twierdzą, że nigdzie w swojej książce nie zakładają propozycjonalizmu i żaden argument tam sformułowany nie zależy od prawdziwości tej tezy (zob. Cappelen, Lepore 2006: 469). Bach podtrzymuje swój zarzut w 2006b. Zobacz też Cappelen, Lepore (2006b) i Bach (2006c).

²⁰ Do minimalistów odrzucających propozycjonalizm należy również Soames, który zamiast o rdzeniach sądów mówi o matrycach sądów (zob. Soames 2005).

²¹ Zob. Bach (2006b: 5).

²² Tak rozumianą implisyturę należy odróżnić od Grice'owskiej implikatury (zob. np. Bach 1994). Jedną z różnic jest to, że implikatury są zazwyczaj zewnętrzne wobec tego, co powiedziane, a implisytury są konstruowane z tego, co powiedziane (zob. Bach 2001: 20).

wyróżnić następujące poziomy analizy znaczenia: 1. znaczenie językowe, 2. to, co powiedziane, 3. to, co powiedziane uzupełnione (w przypadku zdań semantycznie niepełnych) lub rozszerzone (w wypadku zdań semantycznie pełnych), 4. to, co implikowane. To, co powiedziane, Bach rozumie z grubsza po Grice'owsku, tyle że odrzuca założenie, że to, co powiedziane, musi być częścią tego, co przekazane. Zachowuje natomiast Grice'owski warunek korelacji syntaktycznej, który mówi, że to, co powiedziane, ma odpowiadać składnikom zdania, ich porządkowi i charakterowi semantycznemu²³.

4. Kontekstualizm

Jeden z głównych orędowników kontekstualizmu, Recanati, odrzuca zasadę minimalizmu (która, jak pamiętamy, głosi, że to, co powiedziane, odbiega od znaczenia dosłownego tylko wtedy, gdy jest to konieczne dla odtworzenia wyrażanego sądu) na rzecz zasady dostępności, zgodnie z którą to, co powiedziane, musi być intuicyjnie dostępne (*intuitively accessible*) dla (normalnych) uczestników konwersacji (Recanati 2004: 20). A ponieważ to, co jest dostępne, jest — jego zdaniem — wynikiem zarówno procesów semantycznych, jak i pragmatycznych procesów niezdeteminowanych językowo, to *to, co powiedziane*, staje się w jego koncepcji pojęciem pragmatycznym, a nie semantycznym²⁴. Według minimalizmu znaczenie językowe (znaczenie zdań-typów) po uzupełnieniu przez nasycanie daje to, co powiedziane, które można następnie uzupełniać dodatkowymi procesami pragmatycznymi tak, aby otrzymać to, co nadawca chciał zakomunikować. Kontekstualista widzi rzecz inaczej. Znaczenie językowe jest uzupełniane w drodze prymarnych procesów pragmatycznych, na które składa się nasycanie oraz opcjonalne (*optional*) procesy pragmatyczne²⁵.

Interpretacje pragmatyczne można podzielić na prepropozycyjalne i postpropozycyjalne. Do postpropozycyjalnych należą procesy sekundarne, ustalające m.in. implikatury konwersacyjne. Do prepropozycyjalnych należą procesy prymarne, wśród których można wyróżnić procesy obowiązkowe i opcjonalne. Te pierwsze są regulowane językowo i uruchamiane przez znaczenie językowe (oddolne), a te drugie są determinowane pragmatycznie i uruchamiane przez kontekst (odgórne). Do procesów opcjonalnych (inaczej: modulacji) należą dowolne wzbogacanie (*free enrichment*), osłabianie (*loosening*) i przeniesienie (*transfer*). Załóżmy, że A pyta: „Czy Kasia jest głodna?”, a B jej odpowiada: „Ona jadła śniadanie”. Dzięki procesom interpretacji sekundarnej A zrozumie, iż B daje do zrozumienia, że Kasia nie jest głodna (to, co B daje

²³ Zob. Bach (2001: 15).

²⁴ Pod tym względem zgadza się z nim minimalistka Borg.

²⁵ Tutaj kontekstualizm radykalny różni się od kontekstualizmu umiarkowanego, ponieważ jego zwolennicy twierdzą, że opcjonalne procesy pragmatyczne są konieczne.

do zrozumienia, będzie oczywiście zależało od kontekstu; jeśli rozmowa odbywa się wieczorem, *B* poprzez swoją wypowiedź da do zrozumienia, że Kasia *jest* głodna). W interpretacji prymarnej możemy wyróżnić dwa etapy. Proces oddolny pozwala na ustalenie, że „ona” odnosi się do Kasi. Dzięki interpretacji odgórnej ustalimy zaś, że chodzi o to, iż Kasia jadła *dziś* śniadanie. Ustalenie odniesienia „ona” jest regulowane przez znaczenie językowe. Zaimek „ona” jest ‘jawnym’ okazjonalizmem i reguły językowe mówią, że do pełnego zrozumienia zdania, w którym to wyrażenie występuje, konieczne jest ustalenie, do kogo ów zaimek się odnosi. Uzupełnienie, że chodzi o śniadanie jedzone w dniu rozmowy, następuje już w efekcie procesu regulowanego kontekstowo. Zdanie „Ona jadła śniadanie” nie zawiera bowiem ‘jawnego’ miejsca na określenie czasu i — po uzupełnieniu odniesienia przedmiotowego zaimka „ona” — wyraża pełen sąd. Zgodnie z regułami semantycznymi zdanie to należy rozumieć tak, że Kasia przynajmniej raz w swoim życiu jadła śniadanie (przed czasem, w którym odbywa się wymiana zdań). Jest jednak oczywiste, że w większości kontekstów nadawcy nie chodzi o przekazanie informacji, że osoba, o której mowa, kiedyś jadła śniadanie. Do wskazania, że chodzi o śniadanie jedzone w dniu rozmowy, potrzebny jest zatem — regulowany kontekstowo, a nie językowo — proces odgórny. Ten i podobne przykłady mają świadczyć, że przypisywanie warunków prawdziwości należy do pragmatyki, a nie do semantyki (por. Odrowąż-Sypniewska 2010).

Proces dowolnego wzbogacania ilustruje przykład:

Maria wyjęła klucz i otworzyła drzwi.

Ponieważ stosujemy dowolne wzbogacanie, zdanie to rozumiemy jako:

*Maria wyjęła klucz i otworzyła drzwi tym kluczem*²⁶.

Przeciwieństwem wzbogacania jest osłabianie. Ma ono miejsce na przykład w interpretacji zdania:

Bankomat połknął kartę.

Interpretując to zdanie, osłabiamy znaczenie wyrażenia „połknął” tak, żeby mogło się ono stosować również do bankomatów. Innym procesem opcjonalnym jest przeniesienie. Zdanie „Stoję na parkingu” w określonych okolicznościach rozumiemy jako mówiące, że samochód nadawcy stoi na parkingu, a nie, że sam nadawca tam stoi²⁷. Podobnie zdanie „Kanapka z szynką wyszła bez

²⁶ Recanati zauważa, że nie jest to jedyne możliwe wytłumaczenie, ponieważ można tę interpretację wyjaśnić, odwołując się do niewyartykułowanego składnika (Recanati 2004: 25).

²⁷ „Stoję na parkingu” dosłownie denotuje własność posiadaną przez samochód (gdy został postawiony na parkingu), a wskutek przeniesienia denotuje przysługującą właścicielowi samochodu własność wynikającą z faktu, że jego samochód posiada tę pierwszą własność (zob. Recanati 2004: 26).

płatenia” interpretujemy jako „Ten, kto zamówił kanapkę z szynką, wyszedł bez płatania”²⁸.

Warto tutaj podkreślić, że kontekstualiści odrzucają koncepcję Grice’a²⁹. Albowiem według nich do tego, aby odczytać sąd przekazywany, nie jest konieczne wcześniejsze odczytanie sądu minimalnego. U Grice’a do wyprowadzenia implikatury konieczne jest uświadomienie sobie tego, co zostało powiedziane. Według kontekstualistów zaś wzbogacanie, osłabianie i przeniesienie stosuje się bezpośrednio do znaczeń dosłownych poszczególnych wyrażen. Znaczenie dosłowne *zdania* nie jest nigdy ‘obliczane’. Na przykład ktoś, kto słyszy „Kanapka z szynką wyszła bez płatania”, od razu przejdzie do sformułowania „Ten, kto zamówił kanapkę z szynką, wyszedł bez płatania” i nie będzie formułował ‘pośredniego’ sądu minimalnego. Zderzenie wyrażenia „kanapka z szynką” ze słowem „wychodzić” spowoduje natychmiastowe przejście do interpretacji niedosłownej. Pełen sąd dosłowny, któremu można by przypisać warunki prawdziwości, nie zostanie zatem nigdy sformułowany. Proces przeniesienia zostaje uruchomiany ‘lokalnie’, gdy tylko wypowiedziana zostaje część „kanapka wyszła”, a nie globalnie, dopiero po sformułowaniu całego absurdalnego sądu dotyczącego wychodzącej kanapki.

Recanati podkreśla także, że procesy pragmatyczne są interaktywne. Gdy słyszymy zdanie „miasto śpi”, to możemy „miasto” zinterpretować dosłownie i wtedy wymusi to niedosłowną (osłabioną) interpretację wyrażenia „spanie” albo nazwę „miasto” zinterpretować metonimicznie jako „mieszkańcy miasta” i wtedy „spanie” będzie mogło być interpretowane dosłownie. Wybór interpretacji w jednym miejscu od razu wpływa na sposób interpretacji w innym. Ponownie widać, że nie dochodzi do zbudowania całego sądu dosłownego.

Według ‘pełnokrwistych’ kontekstualistów, takich jak Searle i Travis, sądów minimalnych w ogóle nie ma³⁰. Nie ma bowiem takiego poziomu, na którym otrzymywalibyśmy coś, czemu można przypisać warunki prawdziwości, a co powstało bez odwołania się do procesów odgórnych (por. Recanati 2004: 90). Searle na przykład twierdzi, że sąd można wyrazić tylko na tle określonych (niewyartykułowanych) założeń:

²⁸ „Kanapka z szynką” dosłownie denotuje kanapkę z szynką, a wskutek procesu przeniesienia kogoś, kto tę kanapkę zmkówił (zob. Recanati 2004: 26).

²⁹ Jak widzieliśmy wcześniej, koncepcję Grice’a odrzucają również minimaliści Cappelen i Lepore, natomiast akceptuje ją (z pewnymi modyfikacjami) Bach.

³⁰ Według mojej klasyfikacji do kontekstualistów radykalnych należy zaliczyć również Katarzynę Jaszczółt, mimo iż jest ona zwolenniczką *semantyki* znaczeń domyślnych (*default semantics*). Jaszczółt twierdzi, że automatycznie wzbogacamy znaczenia wyrażen prostych do postaci domyślnej (np. zdanie „Chłopiec był niegrzeczny, więc *matka* go skarciła” zostanie od razu odczytane jako „Chłopiec był niegrzeczny, więc *jego matka* go skarciła”) i formalną semantyczną analizę znaczenia należy stosować bezpośrednio do wypowiedzi intencjonalnych. Jaszczółt mówi o semantyce (a nie pragmatyce) warunków prawdziwości zastosowanej do „przedmiotu badań pragmatyki” (to jest do znaczeń zamierzonych przez nadawcę wypowiedzi, a nie do znaczeń leksykalnych) (zob. Jaszczółt 2006: 140). Jaszczółt akceptuje zatem (T4), ale odrzuca (T2) i (T3).

Założmy, że wchodzę do restauracji i zamawiam danie. Założmy, że mówię, mówiąc dosłownie, „Poproszę stek ze smażonymi ziemniakami”. [...] Zakładam, że nie dostarczą tego dania do mojego domu albo do pracy. Zakładam, że stek nie będzie zatopiony w betonie, ani skamieniały, że nie zostanie wepchnięty do mojej kieszeni, ani rozsmarowany na mojej głowie. Ale żadne z tych założeń nie zostało wyrażone (*made explicit*) w mojej dosłownej wypowiedzi (Searle 1992: 180).

Minimalista semantyczny powie, że nawet jeśli stek został dostarczony do domu zamawiającego zatopiony w betonie, to prośba „Poproszę stek z ziemniakami” została spełniona. Kontekstualista temu zaprzeczy i powie, że warunki spełnienia dla tej prośby spełnione jednak nie zostały.

Innym przykładem rozważanym przez Searle’a jest czasownik „*cut*”. Searle twierdzi, że mimo iż czasownik ten nie jest wieloznaczny, to wnosi różny wkład do warunków prawdziwości zdania „*Bill cut the grass*” i zdania „*Bill cut the cake*”. Jego zdaniem dzieje się tak, ponieważ inne założenia tła odgrywają rolę, gdy „*cut*” używane jest w połączeniu z „*grass*”, a inne — gdy jest używane w połączeniu z „*cake*”. Celowo nie przetłumaczyłam tego przykładu, ponieważ w języku polskim użyjemy innego czasownika dla trawy, a innego dla ciasta: „Bill ścinał trawę” i „Bill pokroił ciasto”. Istnienie w innym języku dwóch różnych terminów dla wyrażenia dwóch różnych sensów związanych z jakimś terminem w danym języku można potraktować jako argument za tym, że ów termin w języku wyjściowym jest wieloznaczny (por. Kripke 1979). Jeśli tak, to musimy uznać, że Searle nie ma racji, twierdząc, że „*cut*” jest wyrażeniem jednoznacznym.

Searle twierdzi, że takie zdania, jak „*Bill cut the grass*” nie mają warunków prawdziwości, dopóki nie uwzględnimy założeń tła istotnych dla danego kontekstu. W normalnym kontekście pocięcie trawnika na części nie będzie uprawdzało zdania „*Bill cut the grass*”, ale można wyobrazić sobie taki kontekst, w którym takie pocięcie będzie tym, o co chodzi nadawcy (np. gdy trawnik jest hodowany na sprzedaż i przed sprzedażą musi być pocięty na dające się transportować kawałki) (zob. Recanati 2004: 92)³¹. To, że bez założeń tła jesteśmy bezradni, Searle ilustruje przykładem rozkazu „*Cut the sun*”. Ponieważ dla takiego zdania nie ma założeń tła, nie wiemy, co byłoby spełnieniem tego rozkazu. Wydaje się jednak, że można na to odpowiedzieć, wskazując, że „*Cut the sun*” jest nonsensem semantycznym, takim jak „mąka schudła grzybowo”. Nie da się pokroić ani ścinać słońca, więc taki rozkaz jest niewykonalny, niezależnie od tego, jaki konkretnie sposób krojenia lub cięcia byśmy rozważali.

Innym przykładem jest czasownik „otwierać” (zob. Recanati 2004: 93). Otwierać można (między innymi) drzwi, okna, oczy, butelki i rany. Każdą z tych rzeczy otwiera się trochę inaczej. Rany otwiera się na przykład skalpelem. Jednakże w większości kontekstów ktoś, kto usiłowałby naciąć skalpelem drzwi,

³¹ Po polsku w tej sytuacji powiedzielibyśmy „Bill pokroił trawnik”, a nie „Bill ścinał trawnik”.

nie zostałyby potraktowany jako wypełniający prośbę „Otwórz drzwi”. Może się jednak tak zdarzyć, że proszący będzie oczekiwał właśnie takiego działania od otwierającego. Co więcej, nawet jeśli przyjmiemy, że drzwi w danej sytuacji należy otworzyć kluczem, to sposób, w jaki należy tego dokonać, jest determinowany przez kontekst. Zazwyczaj chodzi o włożenie klucza do zamka i przekręcenie go, ale może też chodzić o podważenie kluczem zamka lub całych drzwi. Przykład ten ma również ilustrować fakt, że bez znajomości założeń tła danego kontekstu nie wiemy, jakie są warunki prawdziwości wypowiedzanych w tym kontekście zdań. To kontekst determinuje, jakie działanie będzie wypełnieniem prośby, a jakie nie.

5. Minimalizm czy kontekstualizm?

Może wydawać się, że stanowisko kontekstualistyczne jest znacznie bardziej intuicyjne. Wszyscy zgodzimy się, że znaczenie wyrażen zmienia się trochę w zależności od kontekstu. Metafora porównująca znaczenia wyrażen prostych do dających się kształtować worków z piaskiem jest bardziej przekonująca niż porównanie znaczeń do cegieł³². Trzeba jednak pamiętać, że minimaliści nie twierdzą wcale, że znaczenia wyrażen prostych są takimi niepodlegającymi kształtowaniu twardymi cegłami. Znaczenia tych wyrażen bardzo często są ogólne i niedookreślone. Jeśli słyszę zdanie, „Jan zjadł supkę” i wiem, że jest to wypowiedź prawdziwa, to wiem, że Jan zjadł supkę, ale nie wiem, *jak* ją zjadł. Zakładam, że zjadł ją normalnie łyżką (tak jak się zazwyczaj je supkę), ale jeśli okaże się, że Jan zjadł tę supkę widelcem (zmęczył się bardzo, dużo czasu mu to zajęło, ale się udało), to wypowiedziane zdanie nadal będzie prawdziwe. Co więcej, moim zdaniem, znaczenie zdania „Jan zjadł supkę” będzie takie samo zarówno w kontekście, w którym zjadł supkę łyżką, jak i w kontekście, w którym zjadł supkę widelcem. Kontekstualiści ułatwiają sobie zadanie, wybierając tak ogólne czasowniki jak „ciąć” („cut”) i „otwierać” i na tych przykładach przekonując, że znaczenie zmienia się wraz z kontekstem. Jeśli jednak zastosujemy tę analizę do innych (mniej ‘ogólnych’) czasowników, to zobaczymy, że dla tych przypadków jest znacznie mniej intuicyjna.

Podobnie, na pierwszy rzut oka jest oczywiste, że ktoś, kto prosi o stek w restauracji, nie będzie usatysfakcjonowany, jeśli przywiozą mu do domu stek zatopiony w betonie albo utopiony w akwarium. Może się zatem wydawać, że kontekstualista ma rację: prośba zamawiającego nie zostałaby w ten sposób spełniona, a zatem założenia tła — o których mówi Searle — wpływają na znaczenie wypowiedzi i zmieniają ich warunki prawdziwości. Czy jednak na pewno tak jest? Recanati i Searle koncentrują się tutaj na prośbie i jej warunkach spełnienia. Ponieważ nas interesują głównie warunki prawdziwości, to wróćmy

³² Por. Recanati (2004) i (2010). Jest to metafora zaproponowana przez Jonathana Cohena.

na grunt zdań oznajmujących³³. Załóżmy, że klient dostaje w restauracji stek utopiony w akwariu. Zgodzimy się wszyscy, że nie o to mu chodziło, ale zastanówmy się, jaka w tej sytuacji będzie wartość logiczna zdania „Klient dostał stek” (po usunięciu okazjonalności czasowej). Searle i Recanati muszą twierdzić, że zdanie to jest fałszywe (ponieważ klient nie dostał steku na talerzu, ze sztućcami itp.). Jednakże moim zdaniem, intuicyjnie to zdanie jest prawdziwe. Prawdą jest bowiem, że klient dostał stek. Co prawda nie taki, jaki chciał, ale jednak dostał. Jeśli tylko to, co pływało w akwariu (czy było zatopione w betonie), można nazwać „stekiem”, to trzeba uznać, że ów stek dostał. Analogicznie w sytuacji, gdy stek został mu wciśnięty do kieszeni czy rozciśnięty na głowie. Jak wytłumaczyć to, że mimo iż klient dostał stek, to jego prośba o stek nie została spełniona? Wydaje się, że pewnym rozwiązaniem byłoby tutaj twierdzenie, że warunki spełnialności próśb są bardziej gruboziarniste niż warunki prawdziwości odpowiednich zdań twierdzących³⁴.

Podobnej koncepcji broni Delia Graff Fara. Graff Fara twierdzi, że sprawozdanie z czyjegós pragnienia może być prawdziwe, nawet jeśli treść zdania podrzędnego jest bardziej ogólna niż treść pragnienia, które powoduje, że sprawozdanie jest prawdziwe. Graff Fara pisze, że „przypisanie pragnienia [...] o postaci ‘*A* chce *C*’ jest prawdziwe zawsze i tylko, gdy *A* ma pragnienie [...], którego dokładną treścią jest sąd *Q*, dla pewnego *Q*, które pociąga sąd wyrażony przez zdanie podrzędne *C*” (2003: 159). Zatem przypisanie pragnienia „*A* chce *C*” może być prawdziwe, nawet jeśli *A* chce czegoś bardziej określonego niż to, co mówi *C*. Moim zdaniem, tę analizę można rozciągnąć również na rozkazy i prośby. Jeśli *A* mówi „Przynieś mi *c*”, to jego rozkaz będzie spełniony pod warunkiem, że *A* dostanie *c*, nawet jeśli to, czego *A* rzeczywiście pragnął, jest bardziej określone. Jeśli proszę o stek, to moja prośba będzie spełniona, jeśli dostanę stek (jakikolwiek by on nie był). *Ja* mogę nie być usatysfakcjonowana sposobem, w jaki moja prośba została spełniona, ale nie mogę zaprzeczyć, że ta *prośba* została spełniona. W myśl tej koncepcji to, czy przyniesienie skamieniałego steku można traktować jako spełnienie prośby „Poproszę stek”, zależy tylko od tego, czy skamieniały stek można uznać za stek.

Poza tym wydaje się, że stanowisko Searle’a (i Recanatiego) prowadzi do innych nieintuicyjnych konsekwencji. Załóżmy, że klient lubi tylko dobrze wysmażane steki (możemy nawet założyć, że obsługa restauracji dobrze o tym wie). Klient mówi „Poproszę stek” i dostaje stek usmażony, ale nie dobrze wysmażony. Jeśli chcemy pozostać w zgodzie z Searle’em, to musimy przyjąć, że prośba klienta nie została spełniona (bo prosił o dobrze wysmażony stek,

³³ Za dyskusję dziękuję tutaj Filipowi Kawczyńskiemu.

³⁴ Analiza minimalistyczna wydaje się jeszcze bardziej intuicyjna dla przykładu ze zdaniem „Poproszę kawę” wygłoszonym w kawiarni. Jeśli klient dostanie worek ziarnistej kawy, a nie filizankę płynnej, to pewnie nie będzie zadowolony, ale trudno zaprzeczyć, że zdanie „Klient dostał kawę” będzie w tym wypadku prawdziwe.

a takiego nie dostał). Konsekwentnie trzeba by przyjąć, że w tym kontekście zdanie „Klient dostał stek” jest fałszywe. Możemy dalej założyć, że nasz klient lubi tylko dobrze wysmażone steki polane sosem z zielonym pieprzem itd. Wydaje się, że ktoś, kto zaczyna podążać za Searle’em, nie może się zatrzymać. Nie wiadomo, jak wytyczyć granicę między tym, co wchodzi w skład znaczenia „stek” przy danej okazji, a tym, co w skład tego znaczenia nie wchodzi. Jeśli klient ma bardzo sprecyzowane oczekiwania względem swojego steku, to wszystkie te oczekiwania wchodzi w skład znaczenia słowa „stek” wypowiedzianego przy tej okazji, a w konsekwencji w sytuacji, gdy stek przez niego otrzymany będzie choćby nieznacznie odbiegał od ideału, zdanie „Klient dostał stek” będzie musiało być uznane za fałszywe. Nie da się ukryć, że takie stanowisko nie jest już tak intuicyjne, jak Searle i Recanati starali się pokazać (żeby nie powiedzieć, że jest po prostu *nieintuicyjne*).

Podobnie można analizować przykład z otwieraniem drzwi. Jeśli chodzi mi o to, żeby ktoś otworzył drzwi kluczem i podając mu ten klucz mówię „Otwórz drzwi” (bo chcę się na przykład przekonać, czy potrafi to zrobić), to jeśli poproszony otworzy drzwi po prostu naciskając klamkę, to w pewnym sensie moja prośba nie zostanie spełniona. Zatem punkt dla kontekstualisty. Ale jeśli zastanowimy się nad wartością logiczną zdania „B otworzył drzwi”, to musimy, moim zdaniem, przyznać, że to zdanie jest prawdziwe, mimo iż B zrobił to w inny niż zamierzony przez proszącego sposób. Jedynym wyjściem dla kontekstualisty byłoby postulowanie dwuznaczności czasownika „otworzyć”: Gdy proszę o otworenie drzwi, to w istocie proszę o otworenie₁ drzwi, a gdy stwierdzam, że B otworzył drzwi, to stwierdzam „B otworzył₂ drzwi”. Wydaje mi się to jednak niepotrzebnym (i *ad hoc*) mnożeniem sensów ponad potrzebę, a ponadto takie stanowisko staje się niefalsyfikowalne, bo zawsze tam, gdzie pojawią się niewygodne wartości logiczne, kontekstualista będzie utrzymywał, że zmienił się sens ocenianego zdania.

Rozważmy jeszcze przykład Trávisa z zielonymi liśćmi. Minimalista może utrzymywać, że znaczenie zdania „Te liście są zielone” jest dość ogólne i zdanie to jest prawdziwe, jeśli powierzchnia liści jest zielona. Jest oczywiście domniemany sposób, na jaki przedmioty danego rodzaju są zielone (dla liści — zielone ‘z natury’), ale ten sposób nie wchodzi w skład znaczenia frazy „zielony x”. Gdy słyszymy „Te liście są zielone”, to oczywiście mamy pewne oczekiwania co do sposobu, na jaki liście są zielone, ale ten sposób nie wchodzi w skład znaczenia. Jeśli powierzchnia liści jest zielona, to zdanie „Te liście są zielone” jest prawdziwe. Znowu może to nie być takie bycie zielonym, o jakie chodziło pytającemu, ale niemniej jednak jest to bycie zielonym³⁵.

³⁵ Oczywiście zdanie „Te liście są zielone” jest nieostre i możemy mieć wątpliwości, czy w jakimś wypadku jest prawdziwe, czy fałszywe np. gdy liście początkowo były całkiem zielone, ale teraz zaczęły już więdnąć.

Uważam, że we wszystkich powyższych przypadkach można twierdzić, że takie wyrażenia, jak „otworzył”, „pokroił”, „zielony” mają ogólny sens, który w różnych przypadkach jest różnie realizowany. W istocie większość, jeśli nie wszystkie, czasowniki i przymiotniki działają w ten sposób. Można jeść zupę, jeść lody, jeść kurczaka i jeść spaghetti³⁶. Każdą z tych rzeczy je się trochę inaczej: zupę łyżką, kurczaka zazwyczaj rękami, spaghetti łyżką i widelcem, a lody się liże. Czy to znaczy, że w każdym ze zdań „Jan zjadł zupę”, „Jan zjadł lody”, „Jan zjadł kurczaka” i „Jan zjadł spaghetti” „zjadł” ma inne znaczenie? Konsekwentny kontekstualista powinien powiedzieć, że tak, ale taka odpowiedź wydaje się co najmniej dziwaczna. Znacznie lepiej przyjąć, że jest ogólny czasownik „jeść”, który stosuje się do większości sytuacji, w których się odżywiamy. Stanley zauważa, że gdy słyszymy „Jan zjadł śniadanie”, to zakładamy, że dokonał tej czynności sam w standardowy sposób. Ale to zdanie byłoby prawdziwe również, gdyby Jan został przez kogoś nakarmiony łyżeczką. Dla Stanleya jest to dowód, że sposób, w jaki ktoś zjadł śniadanie, nie wchodzi w skład treści semantycznej i intuicyjnych warunków prawdziwości zdania „x zjadł śniadanie” (zob. Stanley 2007: 228). Recanati podkreśla wielokrotnie, że inaczej czerwone są ptaki, inaczej samochody, a inaczej arbuzy (zob. Recanati 2004: 135 i 2010: 55). Spodziewamy się, że czerwony arbuz będzie czerwony w środku, a nie na zewnątrz, że czerwony ptak będzie miał upierzenie w kolorze czerwonym, a czerwony samochód — czerwoną karoserię. Czy to jednak świadczy o tym, że „czerwony” znaczy co innego w połączeniu z różnymi rzeczownikami? Wydaje mi się, że tutaj również trzeba oddzielić wiedzę o świecie od znajomości znaczenia wrażeń. Oczekujemy, że czerwony arbuz będzie miał czerwony miąższ, ale to nasze oczekiwanie bierze się z naszej wiedzy dotyczącej arbuzów, a nie cała taka wiedza wchodzi w skład znaczenia słowa „arbuz”. Arbuz o czerwonej skórze też będzie czerwonym arbuzem.

O ile sens czasowników jest dość ogólny, to sens rzeczowników, przynajmniej tych będących nazwami rodzajów naturalnych, jest — moim zdaniem — znacznie bardziej określony. „Kaczka”³⁷ w swoim podstawowym sensie odnosi się do kaczek-ptaków i tylko rozszerzając sens tego słowa, można je zastosować do pluszowych czy drewnianych kaczek. Jeśli mówimy o jakiejś pluszowej zabawce „kaczka”, to musimy liczyć się z tym, że ktoś nas (złośliwie) poprawi i powie „To nie jest kaczka. To jest pluszowa zabawka w kształcie kaczki”. To, że większość nazw ogólnych ma taki sens podstawowy wyróżniony, można zobaczyć, zachowując się jak ‘semantyczny pedant’. Nawet najbardziej czepialska i złośliwa osoba nie będzie mogła powiedzieć o kacze (ptaku), że to nie jest kaczka. Travis twierdzi, że zdanie „To jest kaczka” nie ma wartości logicznej w oderwaniu od kontekstu, ponieważ w kontekście, w którym

³⁶ Za przykład dziękuję Katarzynie Kuś.

³⁷ Ignoruję tutaj pozostałe sensy słowa „kaczka”, takie jak np. „kaczka dziennikarska”.

nadawca wskazuje drewniany model kaczkę, który uznaje za prawdziwą kaczkę, zdanie „To jest kaczka” będzie fałszywe, ale w kontekście, w którym nadawca wskazuje drewniany model kaczkę, chcąc odróżnić ten model od modelu na przykład czajki, zdanie „To jest kaczka” będzie prawdziwe (zob. Travis 1974: 51). Jednakże możemy do obu tych sytuacji wprowadzić zewnętrznego ‘pedantycznego’ obserwatora. Taki zewnętrzny obserwator — nieczuły na subtelności kontekstu — w obu tych wypadkach będzie mógł zaprotestować: „To nie jest kaczka. To jest drewniany model kaczkę”. Przez uczestników rozmowy odbywającej się w drugim kontekście zostanie uznany najprawdopodobniej właśnie za czepiającego się pedanta, ale nikt nie będzie mógł mu zarzucić, że nie ma racji. Natomiast w żadnym kontekście taki zewnętrzny obserwator nie może powiedzieć, że wypowiedziane o kaczcze-ptaku zdanie „To jest kaczka” jest nieprawdziwe.

Podsumowując, trzeba odróżnić to, co powiedziane, od interpretacji tego, co powiedziane³⁸. Jeśli mówię „Poproszę stek”, to tym, co powiedziałam, jest „poproszę stek” i moja prośba zostanie spełniona, jeśli dostanę stek. Natomiast to, że mówię to w restauracji, powoduje, że moja wypowiedź jest interpretowana w określony sposób i zostanie zrozumiana jako prośba o gotowy stek na talerzu ze sztuczkami. Moja wypowiedź jest częścią szerszego kontekstu i przy jej interpretacji ten kontekst jest oczywiście brany pod uwagę, ale nie ma powodu, aby twierdzić, że *powiedziałam*, że chcę dostać usmażony stek na talerzu.

Uważam, że minimalistka Borg ma rację, gdy mówi, że treść semantyczna jest stosunkowo ‘cienka’ i chociaż towarzyszy jej bardzo czasem bogata wiedza o świecie, to ta druga nie jest częścią tej pierwszej (zob. Borg 2004). Inaczej jednak niż Borg uważam, że należy porzucić propozycjonalizm i przyznać, że niektóre zdania języka naturalnego semantycznie nie wyrażają sądów. My im automatycznie przypisujemy wartości logiczne, ponieważ automatycznie je uzupełniamy, odwołując się do informacji kontekstowych, ale nie zmienia to faktu, że ich treść semantyczna jest niepełna. Upieranie się przez niektórych minimalistów, że treścią semantyczną zdania „Piotr jest gotowy” jest to, że Piotr jest do czegoś (w sensie: czegokolwiek) gotowy, wydaje się bardzo naciągane. Lepszą strategią jest przyznanie, że to zdanie wyraża tylko rdzeń sądu, który do pełnego sądu musi być uzupełniony przez kontekst wypowiedzenia. Czy jednak przyjęcie takiego stanowiska nie sprowadza się do kapitulacji i oddania pola kontekstualistom? Stanley twierdzi na przykład, że różnica między odrzucającym propozycjonalizm minimalizmem Bacha a kontekstualizmem jest tylko różnicą stopnia, ponieważ jedni i drudzy zgadzają się, że semantyka nie determinuje pełnych sądów wyrażanych przez zdania i różni się

³⁸ Michael Devitt mówi tu odpowiednio o „metafizyce znaczenia” i „epistemologii interpretacji”, a Stanley o „informacjach przekazywanych w akcie mowy” i „informacjach należących do wiedzy tła” (Stanley 2007: 198).

tylko co do oceny globalności tego zjawiska. Moim zdaniem tak nie jest. Stanley przeoczył tutaj jedną zasadniczą różnicę. Mianowicie minimaliści zakładają, że wyrażenia proste mają stałe ogólne znaczenia, które są ich wkładem w sądy (lub rdzenie sądów) wyrażane przez zdania. Kontekstualiści natomiast twierdzą, że takich znaczeń albo nie ma, albo nawet jeśli są, to nie są tym, co wyrażenia proste 'wkładają' do sądów. Ich wkładem jest bowiem zawsze znaczenie modulowane, a modulacja jest sterowana pragmatycznie przez kontekst, niezależnie od reguł językowych. Posługując się terminologią Recanatiego, można powiedzieć, że minimaliści przyjmują zarówno presupozycję Fregowską głoszącą, iż konwencje językowe wyposażają wyrażenia w sensy, jak i założenie Fregowskie głoszące, że wkład wyrażenia użytego (razem z innymi wyrażeniami) do wygłoszenia pełnej wypowiedzi jest sensem, który ono posiada dzięki konwencjom językowym (Recanati 2010: 17). Natomiast kontekstualiści, nawet jeśli akceptują Fregowską presupozycję, to odrzucają założenie. Co więcej, zdecydowana większość kontekstualistów to ci, którzy odrzucają jedno i drugie. Zatem w mojej opinii różnica między minimalizmem *simplicite* a kontekstualizmem jest różnicą jakościową, a nie tylko ilościową.

Bibliografia

- BACH K. (2001), *You Don't Say?*, "Synthese" 128, 15–44.
- BACH K. (2005), *Context ex Machina*, [w:] Szabó Z. (red.), *Semantics vs. Pragmatics*. Oxford: Oxford University Press, 15–44.
- BACH K. (1994), *Conversational Implicature*, "Mind & Language" 9(2), 124–162.
- BACH K. (2006b), *Minimalism for Dummies: Reply to Cappelen and Lepore*, <http://userwww.sfsu.edu/~kbach/replytoC&L.pdf>.
- BACH K. (2006c), *From the Strange to the Bizarre: Another reply to Cappelen and Lepore*, <http://userwww.sfsu.edu/~kbach/reply2toC&L.pdf>.
- BORG E. (2004), *Minimal Semantics*. Oxford: Oxford University Press.
- BORG E. (2007), *Minimalism versus Contextualism in Semantics*, [w:] Preyer G., Peter G. (red.), *Context-Sensitivity and Semantic Minimalism. New Essays on Semantics and Pragmatics*. Oxford: Oxford University Press, 339–360.
- CAPPELEN H., LEPORE E. (2005), *Insensitive Semantics. A Defense of Semantic Minimalism and Speech Act Pluralism*. Oxford: Blackwell Publishing.
- CAPPELEN H., LEPORE E. (2006), *Replies*, "Philosophy and Phenomenological Research" 73(2), 469–492.
- COHEN S. (1999), *Contextualism, Skepticism and the Structure of Reasons*, "Philosophical Perspectives 13: Epistemology", 57–89.
- GRAFF FARA D. (2003), *Desires, Scope and Tense*, "Philosophical Perspectives" 17(1), 141–164.
- JASZCZOŁT K. (2006), *Pomiędzy semantyką a pragmatyką*, [w:] Stalmaszczyk P. (red.), *Metodologie językoznawstwa. Podstawy teoretyczne*. Łódź: Wydawnictwo Uniwersytetu Łódzkiego, 131–154.

- KING J., STANLEY J. (2005/2007), *Semantics, Pragmatics, and the Role of Semantic Content*, [w:] Stanley J., *Language in Context*, 133–181.
- KORTA K., PERRY J. (2006), *Varieties of Minimalist Semantics*, "Philosophy and Phenomenological Research" 73(2), 451–459.
- KORTA K., PERRY J. (2006), *The Pragmatic Circle*, "Synthese" 165(3), 347–357.
- ODROWĄŻ-SYPNIEWSKA J. (2010), *Sprawozdania w mowie zależnej*, „Przegląd Filozoficzny” 75, 263–276.
- RECANATI F. (2004), *Literal Meaning*. Cambridge: Cambridge University Press.
- RECANATI F. (2010), *Truth-Conditional Pragmatics*. Oxford: Oxford University Press.
- STANLEY J. (2007), *Language in Context*. Oxford: Oxford University Press.
- TRAVIS C. (1975), *Saying and Understanding*. New York: New York University Press.
- TRAVIS C. (1997), *Pragmatics* [w:] Hale B., Wright C. (red.), *A Companion to the Philosophy of Language*. Oxford: Blackwell Publishers, 87–107.

Rozdział 9

Nowe spojrzenie na referencję i kwantyfikację. Filozoficzne implikacje dynamicznego zwrotu w semantyce

Justyna Grudzińska (Uniwersytet Warszawski)

Słowa kluczowe: Dynamiczna Logika Pluralna (DPIL), Dynamiczna Logika Predykatów (DPL), Dynamiczna Logika Typów (DLT), Modele wielowymiarowe, Teoria Reprezentacji Dyskursu (DRT)

1. Wstęp

Teorie dynamiczne proponują nowe spojrzenie na znaczenie, odchodzące od paradygmatu statycznego wywodzącego się z prac Tarskiego ([1944] 1995), Davidsona (1967) oraz Montague (1973). Centralnym pojęciem paradygmatu klasycznego jest pojęcie prawdziwości w modelu. Podejście statyczne utożsamia znaczenie zdania z jego warunkami prawdziwości. Mówiąc technicznie, standardowa semantyka Tarskiego dla języka logiki predykatów pierwszego rzędu (LP) interpretuje formułę jako zbiór wartościowań (ściśle, ten zbiór wartościowań, który spełnia daną formułę w danym modelu). W podejściu dynamicznym główną rolę przestaje odgrywać pojęcie prawdziwości w modelu. Centralnym pojęciem staje się pojęcie „potencjału do zmiany kontekstu” (*context-change potential*). Teorie dynamiczne utożsamiają znaczenie zdania ze zmianą, jaką to zdanie wprowadza do kontekstu. Interpretacja zdania aktualizuje kontekst: przekształca kontekst zastany w nowy. Semantyki wpisujące się w nurt dynamiczny interpretują zdanie jako relację między kontekstami: wejściowym i wyjściowym, przy czym w zależności od systemu kontekst jest różnie modelowany jako struktura reprezentacji dyskursu, wartościowanie, stan informacyjny.

Zwrot dynamiczny włącza do semantyki pragmatyczne aspekty interpretacji, takie jak kontekst. W klasycznym paradygmacie rola kontekstu w interpretacji języka nie zostaje zupełnie pominięta. Podejście statyczne modeluje

zależność interpretacji wyrażeń języka naturalnego od kontekstu, wykorzystując zmienne wolne do reprezentowania takich wyrażeń, jak „ja” oraz wartościowania (funkcje przypisujące obiekty do zmiennych) do reprezentowania kontekstu. Interpretacja zmiennej wolnej zależy od wartościowania tak, jak interpretacja wyrażenia „ja” zależy od kontekstu. W podejściu dynamicznym zwraca się jednak uwagę na to, że interpretacja wyrażeń nie tylko zależy od kontekstu, ale ten kontekst również zmienia. Podstawowa obserwacja lingwistyczna motywująca zwrot dynamiczny w semantyce była taka, że wyrażenia rzeczownikowe wprowadzają do dyskursu obiekty i że w dalej następujących partiach dyskursu możemy do tak wprowadzonych obiektów przy pomocy wyrażeń rzeczownikowych się odnosić. Ilustrują to następujące przykłady:

- (1) Jakiś człowiek wszedł do pokoju. Powitała go cisza.
- (2) Każdy mężczyzna kocha jakąś kobietę. (Oni) przynoszą im kwiaty.
- (3) Trzy dziewczynki podniosły stół. Było im ciężko.

(1)–(3) pokazują, że wyrażenia rzeczownikowe, takie jak „jakiś człowiek”, „każdy mężczyzna”, „trzy dziewczynki” (tzw. poprzedniki relacji anaforycznej), mogą zmieniać kontekst przez dodawanie nowych obiektów. (1)–(3) pokazują także to, że takie wyrażenia rzeczownikowe, jak „go”, „oni”, „im” (tzw. wyrażenia anaforyczne) mogą odnosić się do wcześniej wprowadzonych już obiektów. Nową dynamiczną semantykę dla wyrażeń rzeczownikowych definiuje się przez uwzględnienie ich roli w dyskursie, przez uwzględnienie ich potencjału do zmiany kontekstu. Semantyki dynamiczne modelują nie tylko zależność interpretacji wyrażeń od kontekstu, ale także właśnie samą zmianę kontekstu.

W filozofii języka i — szerzej — w paradygmacie statycznym postuluje się podział wyrażeń rzeczownikowych na kwantyfikujące (np. „każdy człowiek”) i referencyjne (np. „Jan”). Przyjmuje się, że istnieje zbiór cech semantycznych (syntaktycznych), który pozwala odróżnić jedno wyrażenie od drugich, obejmujący m.in. wartość semantyczną (własności wyższego rzędu w przypadku wyrażeń kwantyfikujących versus indywidua w przypadku wyrażeń referencyjnych), wchodzenie w interakcje zasięgowe (wieloznaczności zdań z kwantyfikatorami versus jednoznaczność zdań z nazwami). Ponieważ w postulowany podział nie wpisują się równie łatwo wszystkie wyrażenia rzeczownikowe, prowadzone są spory o status bardziej problematycznych zwrotów. Najstarszy taki spór, bo sięgający połowy zeszłego wieku, toczy się o deskrypcje określone (np. „ten mężczyzna”)¹. Podobne, chociaż nieco bardziej współczesne, kontrowersje budzi również semantyka deskrypcji nieokreślonych (np. „jakiś mężczyzna”)² oraz zaimków anaforycznych (np. „on”)³. W dalszej części tekstu

¹ Polemika Strawsona ([1950] 1967) z teorią deskrypcji Russella ([1905] 1967).

² Fodor i Sag (1982), Ludlow i Neale (1991).

³ Evans (1977 a, b, 1980), Neale (1990).

omówione zostaną główne systemy dynamiczne ze szczególnym uwzględnieniem nowej, odmiennej od klasycznej roli wyrażen rzeczownikowych. W toku dyskusji będę próbowała dowodzić, że wraz ze zwrotem dynamicznym wyłania się nowa unifikująca teoria wyrażen rzeczownikowych, prowadząca do unieważnienia klasycznego podziału wyrażen na referencyjne i kwantyfikujące.

2. Teoria Reprezentacji Dyskursu (DRT)

Teoria Reprezentacji Dyskursu Kampa (1981) historycznie stanowi pierwszą wersję semantyki dynamicznej⁴. Slogan o potencjale do zmiany kontekstu przekłada się na reguły interpretacyjne dla poszczególnych wyrażen i konstrukcji syntaktycznych. Zastosowane do danego zdania *Z* w kontekście *C* reguły DRT definiują wkład, jaki zdanie *Z* wnosi do kontekstu *C*. W ten sposób kontekst *C* jest przekształcany w nowy kontekst, w którym interpretuje się kolejno następujące zdanie, które z kolei prowadzi do następnego kontekstu itd. Działanie reguł DRT zilustrujmy nieformalnie na przykładzie:

(1) Jakiś człowiek wszedł do pokoju. Powitała go cisza.

Reguła dla deskrypcji nieokreślonej „jakiś człowiek” wprowadza do dyskursu tzw. znacznik dyskursu — zmienną *x*, która reprezentuje pewien obiekt, spełniający dwa warunki *x jest człowiekiem* oraz *x wszedł do pokoju*. Zbiór znaczników dyskursu (uniwersum) wraz ze zbiorem warunków nałożonych na znaczniki tworzą strukturę reprezentacji dyskursu (DRS). Znacznik dyskursu wprowadzony do uniwersum DRS ustanawia poprzednik anaforyczny dla zaimka. Reguła dla zaimka „go” wprowadza nowy znacznik, który zostaje utożsamiony z dostępnym w kontekście poprzednikiem⁵. DRS dla zdania (1) jest prawdziwy, jeżeli istnieje obiekt (reprezentowany przez *x*), który spełnia nałożone nań warunki: *x jest człowiekiem*, *x wszedł do pokoju* oraz *x-a powitała cisza* — wszystkie znaczniki z uniwersum DRS otrzymują tym samym interpretację egzystencjalną.

DRT posługuje się niestandardowym językiem logicznym DRS-ów. Pojedynczy DRS jest parą $\langle U, Con \rangle$, składającą się z uniwersum DRS-a (zbiór znaczników) oraz zbioru warunków *Con*. Ponadto język DRT używa spójników negacji, alternatywy oraz implikacji. DRS-y dla poszczególnych zdań są generowane przez reguły DRT. DRS dla (1) tworzy się przy pomocy reguły dołączającej w prefiksie zbiór zmiennych do zbioru warunków. Reguła ta kompensuje brak koniunkcji i kwantyfikacji. Prefiksowane zmienne funkcjonują jak mechanizm kwantyfikacji, a prefiksowany zbiór warunków traktowany jest jako ko-

⁴ Opis DRT został przygotowany na podstawie czterech głównych źródeł: Kamp (1981), Kamp i Reyle (1993), Eijck i Kamp (1997), Nouwen (2003).

⁵ Podobne reguły rządzą anaforycznymi użyciami deskrypcji określonych.

niunkcja. Prosty DRS, taki jak $[x][Cx, WPx, PCx]$, ma takie same warunki prawdziwości jak formuła $\exists x(Cx \& WPx \& PCx)$ w LP. Spójniki negacji, implikacji i alternatywy przekształcają proste DRS-y w złożone warunki. Przykładowo rozważmy regułę dla konstrukcji warunkowej:

(4) Jeśli jakiś człowiek wszedł do pokoju, powitała go cisza.

Reguła dla (4) tworzy złożony warunek składający się z zagnieżdżonych DRS-ów K i K' , $K \Rightarrow K'$. Znaczenie $K \Rightarrow K'$ jest takie, że dla dowolnych znaczników z uniwersum K (które spełniają warunki z K) musi być również spełnione K' . Znaczniki z poprzednika konstrukcji warunkowej otrzymują tym samym interpretację uniwersalną. DRS dla (4): $[] [[x][Cx, WPx] \Rightarrow [] [PCx]]$ ma takie same warunki prawdziwości jak formuła $\forall x((Cx \& WPx) \rightarrow PCx)$ w LP. Na wzór konstrukcji warunkowych DRT analizuje także kwantyfikację uniwersalną. Rozważmy klasyczny przykład ‘anafory osłej’:

(5) Każdy farmer, który posiada (jakiegoś) osła, bije go.

Znaczenie konstrukcji ‘osłej’ jest takie, że dowolne znaczniki x, y z uniwersum zagnieżdżonego DRS-a, które spełniają warunki x jest farmerem, y jest osłem, x posiada y (tzw. restrykcję kwantyfikatora), muszą spełniać także warunek x bije y (tzw. ciało kwantyfikatora). DRS dla (5): $[] [[x,y][Fx, Oy, Pxy] \Rightarrow [] [Bxy]]$ ma takie same warunki prawdziwości jak formuła $\forall x \forall y((Fx \& Oy \& Pxy) \rightarrow Bxy)$ w LP.

Semantyka dla DRS-ów formułowana jest w terminach ‘funkcji osadzających’, gdzie przez funkcję osadzającą rozumie się funkcję częściową przypisującą znacznikom dyskursu (zmiennym) obiekty w danym modelu M . Modelem dla języka DRS jest para $\langle U, I \rangle$, gdzie U to dziedzina interpretacji (nieskończony zbiór obiektów), a I to funkcja interpretacji, która każdemu predykatowi przypisuje zbiór n -tek elementów U , odpowiednio do liczby argumentów predykatu. Na przykład funkcja osadzająca f weryfikuje (1) w M wtw, gdy dziedzina f zawiera x i $f(x)$ jest człowiekiem, $f(x)$ wszedł do pokoju i $f(x)$ powitała cisza w M . Funkcja osadzająca f weryfikuje (5) w M wtw, gdy dziedzina f jest pusta i każde rozszerzenie f, f' , z dziedziną x, y , takie że $f'(x)$ jest farmerem, $f'(y)$ jest osłem, $f'(x)$ posiada $f'(y)$ w M rozszerza się do funkcji f'' , takiej że $f''(x)$ bije $f''(y)$. Funkcja f'' rozszerza f , jeśli dziedzina f'' obejmuje dziedzinę f oraz funkcje f i f'' zgadzają się na przypisywanych wartościach ze względu na tę połączoną dziedzinę.

Warunki prawdziwości nie wyczerpują znaczenia zdań w DRT. Istotną częścią znaczenia są możliwości otwierane przez dany DRS dla interpretacji kolejnych zdań. Warunki określające dopuszczalne powiązania anaforyczne

przypominają ograniczenia syntaktyczne nałożone na wiązanie zmiennych w LP. Nieformalnie rzecz ujmując, tylko znaczniki z uniwersum głównego DRS-a są dostępne dla zaimków anaforycznych z następnymi zdaniami. W DRS-ie dla (5) znacznik x należy do uniwersum zagnieżdżonego DRS-a, ale uniwersum głównego DRS-a pozostaje puste []. Znacznik z uniwersum zagnieżdżonego K wprowadza poprzednik, który jest dostępny wprawdzie dla zaimków anaforycznych w K' , ale nie dla zaimków w dalej następujących zdaniach. W terminach dostępności tłumaczy się kontrast pomiędzy poprawnym (4) a niepoprawnym (6):

- (4) Jeśli jakiś człowiek wszedł do pokoju, powitała go cisza.
 (6) Jeśli jakiś człowiek wszedł do pokoju, powitała go cisza. ? (On) przyniósł kwiaty.

W przypadku kwantyfikatora uniwersalnego „każdy farmer” uniwersum głównego DRS-a również jest puste [], czym tłumaczy się niepoprawność dyskursu (7):

- (7) Każdy farmer, który posiada (jakiegoś) osła, bije go. ? (On) trzyma go za ogon.

Reguła dla nazwy własnej objaśnia z kolei poprawność (8):

- (8) Jeśli Pedro wszedł do pokoju, powitała go cisza. (On) przyniósł kwiaty.

Znacznik odpowiadający nazwie własnej wyciągany jest zawsze do uniwersum głównego DRS-a.

W LP wyrażenia referencyjne (np. „Jan”) modeluje się standardowo przy pomocy stałych indywiduowych, deskrypcje nieokreślone (np. „jakiś człowiek”) — przy pomocy kwantyfikacji egzystencjalnej. W przypadku zaimków anaforycznych („swoje”, „go”) dominować zdaje się pogląd, zgodnie z którym zaimek anaforyczny należy interpretować w trojaki sposób: (a) jako zmienną związaną przez wyrażenie kwantyfikujące w kontekstach typu „Każda mama kocha swoje dziecko”; (b) jako wyrażenie (ko)referencyjne w zdaniach „Jan wszedł do pokoju. Powitała go cisza”; (c) jako deskrypcję określoną w kontekstach typu „Jakiś człowiek wszedł do pokoju. Powitała go (tego człowieka, który wszedł do pokoju) cisza”. Deskrypcje określone budzą więcej kontrowersji i przez różnych autorów klasyfikowane są jako wyrażenia bądź referencyjne, bądź kwantyfikujące. DRT proponuje radykalnie odmienną unifikującą perspektywę na semantykę wszystkich tych wyrażań. W języku logiki DRS-ów nazwy własne, deskrypcje nieokreślone i określone oraz zaimki modeluje się

przy pomocy zwykłych zmiennych. Jak pokazują analizy (1) i (4), siła kwantyfikująca deskrypcji nieokreślonej przywiązana jest do poszczególnych typów konstrukcji: siła egzystencjalna — do prostego DRS-a, siła uniwersalna — do konstrukcji warunkowej. Jedyna różnica pomiędzy deskrypcjami nieokreślonymi (oraz nazwami własnymi) a zaimkami (także deskrypcjami określonymi) polega na tym, że te pierwsze reprezentuje się przy pomocy nowych znaczników, podczas gdy te drugie — przy pomocy już dostępnych (starych) znaczników.

DRT wraz z nową koncepcją znaczenia uwzględniającą (obok warunków prawdziwości) efekt dynamiczny wyrażen rzeczownikowych wykopuje natomiast nowy podział semantyczny. Chodzi o różnicę w potencjale dynamicznym deskrypcji nieokreślonych z jednej strony i zwrotów kwantyfikujących, takich jak „każdy człowiek”, z drugiej. Deskrypcje nieokreślone (także nazwy własne) zmieniają kontekst przez dodanie nowych obiektów, do których można się odnieść w dalej następujących zdaniach. Wyrażenia kwantyfikujące, takie jak „każdy farmer”, nie zmieniają globalnego kontekstu — są jedynie wewnętrznie dynamiczne i zewnętrznie statyczne. Wprowadzają znaczniki, które są wprawdzie dostępne (lokalnie) dla zaimków anaforycznych w ciele kwantifikatora, ale nie dla zaimków w dalej następujących zdaniach. Rozwój semantyk dynamicznych wiąże się z dalszą dynamizacją najpierw samego pojęcia znaczenia w Dynamicznej Logice Predykatów Groenendijka i Stokhofa (1991), potem również zachowania wszelkich wyrażen rzeczownikowych w Dynamicznej Logice Pluralnej van den Berga (1996).

3. Dynamiczna Logika Predykatów (DPL)

DPL definiuje ‘dynamiczne’ znaczenie bez odwoływania się do pośredniej reprezentacji znaczenia (w postaci DRS-ów)⁶. W DPL dynamika znaczenia zostaje bezpośrednio wbudowana w semantykę wyrażen. Język DPL to standardowy język LP. Slogan o potencjale do zmiany kontekstu przekłada się na minimalną modyfikację standardowej semantyki Tarskiego dla LP: zamiast interpretować formułę jako zbiór wartościowań, interpretuje się ją jako zbiór par wartościowań (zbiór par kontekstów: wejściowego i wyjściowego). Podobnie jak w DRT, dynamika interpretacji w DPL wiąże się z informacją o obiektach, które wprowadza się do dyskursu i do których można się następnie odnosić przy pomocy zaimków anaforycznych. Z wyrażeniami rzeczownikowymi łączy się zmienne, a informacja o ich możliwych wartościach zostaje zakodowana w wartościowaniach. Ponieważ wartościowania są funkcjami przypisującymi obiekty do zmiennych, różnica pomiędzy wartościowaniami (kontekstami) wejściowym

⁶ Opis DPL został przygotowany w oparciu o trzy podstawowe źródła: Groenendijk i Stokhof (1991), Dekker (2008b) oraz notatki do kursu Profesor Bittner *Introduction to Dynamic Semantics* (wiosna 2010, dostępne okresowo na stronie: <http://www.rci.rutgers.edu/~mbittner>).

i wyjściowym może być tylko taka, że jakiś nowy obiekt został przypisany do jednej czy więcej zmiennych. Przypomnijmy przykład ze wstępu:

(1) Jakiś człowiek wszedł do pokoju. Powitała go cisza.

Pierwsze zdanie (1) „Jakiś człowiek wszedł do pokoju” wyraża relację pomiędzy dwoma wartościowaniami: g (kontekstem wejściowym) oraz h (kontekstem wyjściowym), gdzie h różni się od g co najwyżej przypisaniem do x pewnego obiektu, który jest człowiekiem i wszedł do pokoju. Ujmując to samo bardziej formalnie, (1) denotuje zbiór par wartościowań $\langle g, h \rangle$, takich że $g[x]h$ i $h(x)$ jest człowiekiem i $h(x)$ wszedł do pokoju. Dalej następujące zdanie „Powitała go cisza” jest interpretowane ze względu na wartościowanie h (zaktualizowany kontekst). Przez powiązanie zaimka „go” ze zmienną x możemy odnieść się do tego człowieka, który wszedł do pokoju: $h(x)$ powitała cisza.

Język DPL jest standardowym językiem LP, składającym się ze zbioru stałych indywiduowych ($a, b, c \dots$) i predykatowych ($R^1, R^2, R^3 \dots$) oraz zbioru zmiennych ($x, y, z \dots$). Formuły buduje się z formuł atomowych przy pomocy negacji, koniunkcji, alternatywy, implikacji oraz egzystencjalnej i uniwersalnej kwantyfikacji. Modelem dla języka DPL jest standardowo para $\langle D, I \rangle$, gdzie D to dziedzina interpretacji (nieskończony zbiór obiektów), a I to funkcja interpretacji dla stałych indywiduowych i predykatowych, która każdej stałej indywiduowej przypisuje elementy dziedziny, $I(c) \in D$ oraz każdemu predykatowi przypisuje zbiór n -tek elementów D , odpowiednio do liczby argumentów predykatu, $I(R) \in D^n$. Wartościowania $g(h, k, l \dots)$ to funkcje przypisujące elementy dziedziny do zmiennych, $g(x) \in D$. Mówimy, że h jest wariantem wartościowania g , $g[x]h$ wtw, gdy h jest takie jak g za wyjątkiem wartości przypisanej do x .

Charakterystycznymi dynamicznymi elementami DPL są definicje kwantyfikatora egzystencjalnego oraz koniunkcji. Zmiana w wartościowaniach (zmiana kontekstu) dokonuje się za sprawą interpretacji kwantyfikacji egzystencjalnej:

$$\models \exists x \varphi \models^M = \{ \langle g, h \rangle : \exists k [g[x]k \ \& \ \langle k, h \rangle \in \varphi] \}.$$

Zgodnie z podaną definicją, interpretacja $\exists x \varphi$ denotuje zbiór takich par wartościowań $\langle g, h \rangle$, że istnieje wartościowanie k , takie że k różni się od g co najwyżej na wartości przypisanej do x oraz razem z h tworzy możliwą parę wartościowań (wejściowego i wyjściowego) $\langle k, h \rangle$ dla φ . Kwantyfikator egzystencjalny zmienia kontekst przez ustalenie wartości przypisanej do określonej zmiennej (która dodatkowo spełnia zasięg φ). Mówiąc metaforycznie, kwantyfikator egzystencjalny zmienia kontekst przez dodanie nowego „indywiduum dyskursowego”. Z kolei koniunkcja sama w sobie nie zmienia kontekstu, ale zachowuje (składa) potencjalne zmiany związane z każdym z członów koniunkcji:

$$\|\varphi \wedge \psi\|^M = \{ \langle g, h \rangle : \exists k (\langle g, k \rangle \in \varphi \ \& \ \langle k, h \rangle \in \psi) \}.$$

Zgodnie z podaną definicją, interpretacja $\varphi \wedge \psi$ względem wartościowania wejściowego g może dać na wyjściu wartościowanie h wtw, gdy istnieje wartościowanie k , takie że interpretacja φ w g prowadzi do wartościowania k , a interpretacja ψ w k umożliwia z kolei przejście do h . W ten sposób uchwycona zostaje ‘dynamiczna’ intuicja, zgodnie z którą kolejność członów koniunkcji ma znaczenie dla jej interpretacji — koniunkcja w DPL nie jest przemienne. Pozostałe definicje DPL są statyczne. Na przykład predykaty przybierają postać testów, które nie zmieniając samych wartościowań, pozostawiają jedynie te, które spełniają predykat:

$$\|Rt_1 \dots t_n\|^M = \{ \langle g, h \rangle : g = h \ \& \ \langle \|t_1\|^M, g \dots \|t_n\|^M, g \rangle \in I(R) \}.$$

Statyczną interpretację otrzymują również operatory negacji, alternatywy oraz implikacji, z tym że w przypadku implikacji $\varphi \rightarrow \psi$ dynamiczne efekty φ mogą mieć wpływ na interpretację ψ . Podobnie jak w DRT implikacja i kwantyfikacja uniwersalna są wewnętrznie dynamiczne i zewnętrznie statyczne. Definicja dla kwantyfikacji uniwersalnej wygląda następująco:

$$\|\forall x \varphi\|^M = \{ \langle g, h \rangle : g = h \ \& \ \forall k (g[x]k \Rightarrow \exists j \langle k, j \rangle \in \|\varphi\|^M) \}.$$

Warunek $g = h$ mówi, że interpretacja $\forall x \varphi$ nie zmienia kontekstu (podobnie jak interpretacja wszelkich ‘testów’) — zwraca wartościowania wyjściowe g wtw, gdy każde wartościowanie k różniące się od g co najwyżej na x razem z j tworzy możliwą parę wartościowań (wejściowego i wyjściowego) dla φ .

DPL modeluje deskrypcje nieokreślone i nazwy własne przy pomocy kwantyfikacji egzystencjalnej; zaimki anaforyczne — przy pomocy zmiennych związanych. W DPL nieadekwatne staje się klasyczne pojęcie wiązania zmiennej, gdzie wiązanie semantyczne zawsze idzie w parze z wiązaniem syntaktycznym. W myśl reguł nowej semantyki syntaktycznie wolna zmienna może zostać semantycznie związana przez pojawiający się wcześniej w tekście kwantyfikator egzystencjalny. Jeden mechanizm semantycznego wiązania zmiennej zastępuje co najmniej trzy mechanizmy przywoływane w LP do objaśnienia anafory: (a) mechanizm syntaktycznego (semantycznego) wiązania zmiennej w kontekstach typu „Każda mama kocha swoje dziecko”; (b) mechanizm koreferencji w zdaniach „Jan wszedł do pokoju. Powitała go cisza”; (c) mechanizm rekonstruowania odpowiedniej deskrypcji określonej w kontekstach typu „Jakiś człowiek wszedł do pokoju. Powitała go (tego człowieka, który wszedł do pokoju) cisza”. Podobnie jak DRT, w DPL semantyczna relacja pomiędzy zaimkami anaforycznymi a ich poprzednikami ustanawiana jest

przez przywiązanie do nich ('koindeksowanych') zmiennych. W przypadkach typu (a) zmienna interpretowana jest ze względu na iterowane konteksty (iteracja związana jest z kwantyfikatorem kontrolującym daną zmienną). W przypadkach typu (b) i (c) kwantyfikator egzystencjalny działający na danej zmiennej ustala dla niej nową wartość, która zachowuje się aż do momentu wystąpienia kolejnego kwantyfikatora działającego na tej samej zmiennej.

Podobnie jak DRT DPL utrzymuje także podział wyrażeń rzeczownikowych na dynamiczne (deskrypcje nieokreślone, nazwy własne) i statyczne (takie wyrażenia, jak „każdy mężczyzna”). Oba podejścia analizują kwantyfikację uniwersalną jako wewnętrznie dynamiczną i zewnętrznie statyczną. I oba stają wobec tzw. problemu anafory pluralnej. Dyskurs traktuje nie tylko o pojedynczych indywiduach, ale także o zbiorach indywiduów (mnogościach, grupach), do których można poprawnie odnosić się przy pomocy zaimków anaforycznych w liczbie mnogiej:

(9) Każdy senator podziwia prezydenta. (Oni) przynoszą mu kwiaty.

Przykłady w rodzaju (9) pokazują, że kwantyfikacja uniwersalna idzie w parze z efektami dynamicznymi, które nie mają wyłącznie lokalnego charakteru. Kwantyfikacja uniwersalna może, jak zdaje się sugerować (9), zmieniać kontekst globalny, ustanawiając poprzednik anaforyczny dla zaimków z dalszych zdań.

Zaproponowane w ramach DRT rozwiązanie problemu anafory pluralnej wiązało się z wprowadzeniem dodatkowego operatora tworzącego złożony warunek, którego interpretacja oparta została na pojęciu kwantyfikatora uogólnionego. W teorii kwantyfikatorów uogólnionych zdanie postaci $Q(A)(B)$ interpretowane jest jako relacja na zbiorach A i B . Na przykład zdanie „Każde A jest B ” interpretowane jest jako zbiór par zbiorów $\langle A, B \rangle$, takich że $A \subseteq B$. W DRT kwantyfikacja uogólniona otrzymuje interpretację *quasi*-dynamiczną. Na przykład (9) po pierwsze testuje, czy świat jest taki, że wszyscy senatorzy mają własność podziwiania prezydenta (to jest, czy zbiór senatorów zawiera się w zbiorze fanów prezydenta), ale po drugie zmienia kontekst przez dodanie do kontekstu globalnego znacznika X — maksymalnego zbioru x -ów, takich że x jest senatorem i x podziwia prezydenta. Taki poprzednik staje się odniesieniem dla zaimka anaforycznego w liczbie mnogiej (nie pojedynczej), czym tłumaczony jest kontrast pomiędzy poprawną kontynuacją „(Oni) przynoszą mu kwiaty” a niepoprawnym dyskursem

(10) Każdy senator podziwia prezydenta. ? (On) przynosi mu kwiaty.

Poprzedniki dla anafory pluralnej tworzone są przy pomocy dodatkowej operacji abstrakcji rekonstruuującej poprzednik anaforyczny z materiału struktury

kwantyfikującej $Q(A)(B)$. Na przykład abstrakcja w (9) tworzy maksymalny zbiór senatorów podziwiających prezydenta. Kamp i Reyle sugerują, że abstrakcja stanowi rodzaj „reguły wnioskowania na DRS-ach” (1993: 344). Jak zwraca na to uwagę Nouwen (2003: 31), nie istnieje naturalny sposób na włączenie tak pomyślanej reguły konstrukcji znacznika X do dynamicznej semantyki DPL. Pełna dynamizacja kwantyfikacji uogólnionej dokonała się dopiero w Dynamicznej Logice Pluralnej van den Berga (1996) i wiązała z kolejną, dalej idącą modyfikacją semantyki standardowej.

4. Dynamiczna Logika Pluralna (DPIL)

Przypomnijmy przykład ze wstępu:

(2) Każdy mężczyzna kocha jakąś kobietę. (Oni) przynoszą im kwiaty.

Przykład (2) pokazuje, że wyrażenie kwantyfikujące wprowadza do dyskursu zbiór obiektów i że przy pomocy zaimka „oni” można się odnieść do tak wprowadzonego zbioru. Semantyka DPIL modyfikuje DPL w sposób umożliwiający włączenie zjawiska anafory pluralnej w zakres teorii dynamicznych⁷. Język DPIL to język LP/DPL (wzbogacony o operatory dystrybutywności Δ oraz maksymalności $\mathbf{M}$). Zaproponowana modyfikacja semantyki DPL związana jest z nowym pomysłem na modelowanie kontekstu: zamiast interpretować formułę jako relację na wartościowaniach, interpretuje się ją jako relację na zbiorach wartościowań, nazwanych stanami informacyjnymi F , $(G, H, I, \dots)$. Dla przykładu — stan informacyjny może składać się z czterech wartościowań f_1, f_2, f_3, f_4 :

$$\{\langle x, b \rangle, \langle y, a \rangle\}, \{\langle x, k \rangle, \langle y, m \rangle\}, \{\langle x, j \rangle, \langle y, m \rangle\}, \{\langle x, j \rangle, \langle y, j \rangle\}.$$

Stany informacyjne kodują informację o zbiorach obiektów przypisanych do zmiennych. W naszym przykładowym stanie informacyjnym do x przypisany jest zbiór $\{b, k, j\}$ (powiedzmy: Bogdan, Karol, Jan), do y — zbiór $\{a, m, j\}$ (powiedzmy: Anna, Maria, Joanna). Van den Berg zwraca uwagę na to, że zdanie „Każdy mężczyzna kocha jakąś kobietę” wprowadza relację zależności (pomiedzy każdym z mężczyzn a kochaną przez niego kobietą), do której również można się odnieść anaforycznie. W zdaniu „(Oni) przynoszą im kwiaty” mówi się w sposób oczywisty o tym, że każdy mężczyzna przynosi kwiaty kochanej przez niego kobiecie, nie zaś cudzej wybrance: „im” odnosi się do odpowiednich kobiet kochanych przez odpowiednich mężczyzn. Stany informacyjne pozwalają kodować także tę informację o zależnościach pomiędzy elementami zbiorów,

⁷ Opis DPIL przygotowałam na podstawie trzech głównych źródeł: van den Berg (1996), Nouwen (2003) oraz notatek do kursu Profesor Bittner *Introduction to Dynamic Semantics* (wiosna 2010, dostępne okresowo na stronie: <http://www.rci.rutgers.edu/~mbittner/>).

w naszym przykładzie pomiędzy poszczególnymi mężczyznami a kobietami: Bogdanem i Anną (przy pomocy wartościowania f_1), Karolem i Marią (przy pomocy wartościowania f_2) itd.

W DPIL dwa operatory odgrywają podstawową rolę w modelowaniu zależności anaforycznych: operator ε oraz operator dystrybutywności Δ . Operator ε spełnia funkcję analogiczną do kwantyfikatora egzystencjalnego z DPL:

$$\| \varepsilon_x \|^\mathbf{M} = \{ \langle G, H \rangle : G[x]H \}$$

ε_x zmienia kontekst przez dodanie nowego „obiekту dyskursowego” — wprowadza zbiór wszystkich wartości przypisanych do x . Posłużmy się przykładem ze wstępu:

(3) Trzy dziewczynki podniosły stół. Było im ciężko.

Pierwsze zdanie (3) wyraża relację pomiędzy dwoma stanami informacyjnymi: G (kontekstem wejściowym) oraz H (kontekstem wyjściowym), gdzie H różni się od G co najwyżej przypisaniem do zmiennej x zbioru dziewczynek (trójelementowego), które to dziewczynki zbiorowo podniosły stół. Zdanie „Trzy dziewczynki podniosły stół” otrzymuje automatycznie interpretację kolektywną: dziewczynki wspólnie (jako zbiór/grupa) wzięły udział w podnoszeniu stołu. Dodanie operatora dystrybutywności Δ umożliwia dystrybucję predykatu „podniosły stół” po elementach zbioru dziewczynek. Ponieważ definicja operatora Δ jest raczej skomplikowana, poprzestaniemy tutaj na przedstawieniu głównej jej idei: formuła $\Delta_x \varphi$ jest prawdziwa o zbiorze przypisanym do x wtw, gdy φ jest prawdziwa o każdym elemencie x . Opatrzono operatorem Δ zdanie „Trzy dziewczynki podniosły stół” mówi o tym, że każda z trzech dziewczynek podniosła stół. Mówiąc dokładniej, zdanie to dopuszcza dalszą wieloznaczność związaną z zasięgiem operatora Δ : w zależności od tego, gdzie w zdaniu umieszczony zostaje Δ , znaczy ono bądź to, że każda z dziewczynek podniosła ten sam stół, bądź to, że każda podniosła jakiś stół (możliwe, że innym).

W DPIL operator dystrybutywności jest także źródłem wspomnianych zależności pomiędzy elementami dwóch zbiorów w takich zdaniach, jak „Każdy mężczyzna kocha jakąś kobietę”. Mówiąc nieformalnie: jeśli jeden ze zbiorów (w tym przypadku zbiór mężczyzn) wprowadzony zostaje w kontekście dystrybutywnym, drugi zbiór (w tym przypadku zbiór kobiet) staje się funkcjonalnie zależny od zbioru, po którym się dystrybuuje: do każdego z mężczyzn przypisany zostaje zbiór kochanych przez niego kobiet (w naszym przykładzie do Bogdana przypisana zostaje Anna, do Karola Maria, do Jana Maria i Joanna). Mówiąc językiem systemu van den Berga, Δ_x umożliwia dostęp do każdego elementu d zbioru x — Δ_x umożliwia zejście do tzw. podstawów informacyjnych

stanu F , $F|_{x=d}$, gdzie $F|_{x=\text{Bogdan}}(y) = \{\text{Anna}\}$, $F|_{x=\text{Karol}}(y) = \{\text{Maria}\}$, $F|_{x=\text{Jan}}(y) = \{\text{Maria}, \text{Joanna}\}$. Interpretacja kolejnego zdania „Oni przynoszą im kwiaty” wymusza nową dystrybucję po x . Podstany informacyjne, które teraz bierzemy pod uwagę, mają zakodowaną informację o zależnościach pomiędzy elementami x i y — możemy tym samym przy pomocy zaimka „im” anaforycznie odnieść się do odpowiednich kobiet kochanych przez odpowiednich mężczyzn.

Interpretując zdania jako relacje na stanach informacyjnych, zmieniamy kontekst przez dodawanie indywiduów, zbiorów indywiduów oraz zależności pomiędzy elementami zbiorów (w przypadku zdań z dwoma czy więcej kwantyfikatorami). W porównaniu z DRT i DPL semantyka dynamiczna van den Berga pozwala modelować potencjał dynamiczny wyrażeń rzeczownikowych w dużo bogatszym zakresie. Jak była o tym mowa, DRT i DPL wprowadziły podział semantyczny wyrażeń rzeczownikowych na dynamiczne i statyczne. W pierwszych systemach dynamicznych deskrypcje nieokreślone otrzymują interpretację dynamiczną, tj. zmieniają kontekst przez dodawanie nowych obiektów. Wyrażenia kwantyfikujące, takie jak „każdy mężczyzna”, mają zaś interpretację statyczną w DPL i jedynie *quasi*-dynamiczną w DRT (za sprawą wprowadzonej nieco *ad hoc* operacji abstrakcji). Dopiero DPL van den Berga proponuje pełną unifikację wszystkich wyrażeń rzeczownikowych. Dowolny zwrot rzeczownikowy zmienia kontekst przez dodanie nowego obiektu (indywiduum, zbioru indywiduów). Wyrażenia kwantyfikujące dołączają do tego dodatkowo niezależny 'kwantyfikujący' (ilościowy) warunek nałożony na rzeczywistość.

Van den Berg podaje nową dynamiczną definicję dla kwantyfikatora. Za dynamizacją kwantyfikacji stoi obserwacja, że wyrażenia kwantyfikujące umożliwiają odniesienie anaforyczne do zbiorów, zasadniczo dwojakiego typu. Mówiąc ogólnie, w przypadku struktury kwantyfikującej $Q(A)(B)$ są to maksymalny zbiór A (zbiór wszystkich rzeczy, które są A) oraz maksymalny zbiór $A \cap B$ (zbiór wszystkich rzeczy, które są A i B). Przykładowo:

- (11) Niewielu pracowników przyszło na zebranie. Ale oni (wszyscy pracownicy) pojawili się na bankiecie po zebraniu.
- (12) Niewielu pracowników przyszło na zebranie. Oni (wszyscy pracownicy, którzy przyszli na zebranie) postanowili nie podejmować żadnych ważnych decyzji.

W standardowej teorii kwantyfikatorów uogólnionych, jak była o tym mowa, struktura kwantyfikująca postaci $Q(A)(B)$ interpretowana jest jako relacja na zbiorach A i B : bierzemy zbiory odpowiadające restrykcji i ciału kwantyfikatora (odpowiednio A i B) i stosujemy do nich kwantyfikator $Q(A)(B)$. W tradycji dynamicznej (Chierchia 1992, 1995; Kamp, Reyle 1993) zdania kwantyfikujące

zaczęto interpretować jako relację na zbiorach A oraz $A \cap B$ ze względu na możliwość pojawienia się powiązań anaforycznych pomiędzy restryktorem kwantyfikatora A a ciałem B . Na przykład w zdaniu „Większość farmerów, która posiada jakiegoś ośła, bije go” zaimbek anaforyczny w B odsyła do materiału w A . W myśl zaproponowanej analizy $Q(A)(A \cap B)$ takie zdanie mówi, że „Większość farmerów, która posiada jakiegoś ośła, posiada ośła i go bije”. Konserwatywność zostaje wprost wbudowana w definicję kwantyfikatora⁸. Definicja van den Berga dla struktury kwantyfikującej $Q_x(\varphi, \psi)$ sięga do powyższych ustaleń, odpowiednio je ‘dynamizując’. Najpierw znajdujemy maksymalne zbiory, które odpowiadają formułom φ oraz ψ . Przypisujemy je do zmiennych x' i x , przy czym $x \subseteq x'$ — konserwatywność wbudowana zostaje w definicję w postaci warunku, że zbiór wprowadzony przez ciało kwantyfikatora jest zawsze podzbiorem zbioru wprowadzonego przez restrykcję. Maksymalność zbiorów zawarowana zostaje przy pomocy operatora maksymalności $\mathbf{M}$, który znowu scharakteryzujemy jedynie nieformalnie: formuła $\mathbf{M}_x\varphi$ jest prawdziwa o zbiorze przypisanym do x wtw, gdy nie istnieje nadzbiór x , który także uprawdziwiłby φ (x jest maksymalnym zbiorem spełniającym φ). Kolejny składnik definicji van den Berga wiąże się z obserwacją, że kwantyfikacja zawsze ma miejsce w odniesieniu do jakiegoś kontekstu, ‘zbioru kontekstowego’, który ustala dziedzinę kwantyfikacji, tj. jeśli $Q(A)(B)$ jest interpretowane w odniesieniu do kontekstu C , to jest interpretowane jako $Q(A \cap C)(B)$. Wreszcie na koniec dodajemy warunek ‘kwantyfikujący’: $Q(x', x)$. Pełna definicja dla kwantyfikacji wygląda następująco (gdzie y to zbiór kontekstowy):

$$Q_y^x(\varphi, \psi) := \varepsilon_{x' \subseteq y} \& \varepsilon_{x \subseteq x'} \& \mathbf{M}_{x' \subseteq y}(\varphi[x/x']) \& \mathbf{M}_{x \subseteq x'}(\psi) \& Q(x', x)$$

W systemie van den Berga wyrażenia kwantyfikujące są zawsze dynamiczne: zmieniają kontekst przez dodanie maksymalnych zbiorów x' oraz x . W rozwinęciach DPIL (Brasoveanu 2008) niefortunności dyskursu typu:

(10) Każdy senator podziwia prezydenta. ? (On) przynosi mu kwiaty.

tłumaczy się konfliktem w semantycznej (nie morfologicznej) liczbie ‘mnogiego’ poprzednika (zbiorów) oraz ‘pojedynczego’ zaimka.

Najnowszy rozwój semantyk dynamicznych przyniósł w pełni kompozycyjne teorie w stylu gramatyk Montague (1973). W największym tylko skrócie przedstawimy ideę stojącą za pożenieniem semantyk dynamicznych z wymogiem kompozycyjności (zob. Muskens 1996, Bittner 2012). W myśl wymogu

⁸ Barwise i Cooper (81) zaobserwowali, że kwantyfikatory języka naturalnego spełniają warunek konserwatywności: dla dowolnego M i dowolnych $A, B \subseteq M$, $Q_M(A, B) \leftrightarrow Q_M(A, A \cap B)$. Warunek konserwatywności mówi o szczególnej roli pierwszego argumentu w strukturze kwantyfikującej $Q(A)(B)$ — A ustala uniwersum, w którym interpretuje się kwantyfikację.

kompozycyjności znaczenie zdania jest funkcją znaczeń jego prostszych części. Zastąpienie dynamicznych wariantów LP (DPL oraz DPIL) dynamicznym wariantem logiki typów wykorzystanej w gramatykach Montague pozwoliło dynamicznym teoriom interpretacji zejść na poziom składowych zdania. W logice typów do każdego wyrażenia języka przypisany jest jednoznacznie jego typ, gdzie typ klasyfikuje wyrażenia języka pod względem własności składniowych i semantycznych. W ‘statycznej’ logice typów (LT) zbiór typów obejmuje (a) dwa typy podstawowe: typ e (indywiduów) oraz typ t (wartości logicznych) oraz (b) typy złożone (funkcji) konstruowane z typów podstawowych. Na przykład nazwy mogą stanowić wyrażenia typu e, których interpretacją są indywidua; zdania są wyrażeniami typu t, których interpretacją są wartości logiczne; predykaty jednoargumentowe to wyrażenia typu et, których interpretacją są funkcje przypisujące każdemu indywiduum jedną z wartości logicznych. Główną innowacją Dynamicznej Logiki Typów (DLT) polega na wprowadzeniu dodatkowego typu podstawowego s (stanów), którego interpretacją jest obiekt semantyczny naśladujący wartościowania — w systemach Muskensa (1996) i Bittner (2012) są to ciągi indywiduów. W DLT do zdania przypisany zostaje typ sst — zdanie interpretuje się jako relację na stanach/kontekstach (obiektach typu s). Do „indywiduum dyskursowego” przypisuje się typ se, którego interpretacją jest obiekt semantyczny naśladujący zmienne — funkcje ze stanów do indywiduów. Podobnie jak w DPIL, w DLT aktualizacja stanu (kontekstu) wiąże się z dodaniem tzw. świadka: indywiduum dyskursowego, zbioru indywiduów dyskursowych, także zależności pomiędzy indywiduami dyskursowymi. DLT pozwala modelować potencjał dynamiczny wyrażen rzeczownikowych w tak samo bogatym zakresie.

5. Wielowymiarowy system cech — nowy model znaczenia dla referencji i kwantyfikacji w języku naturalnym

Rozwój semantyk dynamicznych podaje w wątpliwość istnienie zbioru cech, który sztywno oddzielałyby takie wyrażenia, jak „Jan” od zwrotów kwantyfikujących, takich jak „każdy człowiek”. We wstępie wspomniano o dwóch takich cechach, które miałyby uzasadniać klasyczny podział wyrażen na referencyjne i kwantyfikujące. Pierwsza różnica dotyczy tzw. wartości semantycznej. W klasycznym paradygmacie wyrażenia referencyjne oznaczają indywidua — zdanie z wyrażeniem referencyjnym wyraża sąd jednostkowy, tj. sąd dotyczący konkretnego indywiduum. Jak zwraca uwagę Glanzberg (2008: 208–209), określenie ‘jednostkowy’ jest tu nieco niefortunne, gdyż takie myśli mogą dotyczyć więcej niż jednego indywiduum — chodzi raczej o sąd mówiący o konkretnych — tych, a nie innych — indywiduach. Zdanie „Jan i Maria podnieśli stół” także wyraża sąd ‘jednostkowy’ o Janie i Marii — tych, a nie innych osobach. W standardowym ujęciu wyrażenia kwantyfikujące denotują z kolei własności dru-

giego rzędu (zbiory zbiorów indywiduów); zdania kwantyfikujące $Q(A)(B)$ mówią o relacjach na własnościach (zbiorach A i B) — wyrażają sądy ogólne, tj. sądy o jakichkolwiek indywiduach spełniających dane własności. W perspektywie dynamicznej kryterium ‘jednostkowości/ogólności’ staje się wątpliwe: wszystkie wyrażenia rzeczownikowe wprowadzają do dyskursu obiekty. Nazwy własne i deskrypcje nieokreślone — indywidua, nazwy zbiorowe („Jan i Maria”) oraz wyrażenia kwantyfikujące („każdy człowiek”) — zbiory indywiduów. Wyrażenia kwantyfikujące dodatkowo wyrażają niezależny ‘kwantyfikujący’ (ilościowy) warunek nałożony na rzeczywistość. Gdy zakres badania rozszerzy się na wypowiedzi wielozdaniowe i związki anaforyczne, intuicje dotyczące jednostkowości/ogólności sądu zaczynają zawodzić. Jak ilustruje to zjawisko anafory pluralnej, zaimki odnoszą się do obiektów (konkretnych zbiorów) wprowadzanych przez kwantyfikatory. Co więcej, same wyrażenia kwantyfikujące mogą odnosić się anaforycznie do wcześniej wprowadzonych obiektów (konkretnych indywiduów):

(13) Janek, Piotr, Michał i Ania weszli do pokoju. Każdy chłopiec miał czapkę.

Wydaje się, że w kontekście wypowiedzi (13) zdanie „Każdy chłopiec miał czapkę” wyraża sąd jednostkowy o zbiorze trzech chłopców: Janku, Piotrze i Michale, tych właśnie konkretnych osobach. Glanzberg (2008: 223–224) powołuje się na taką samą intuicję przy okazji omawiania problematycznego statusu kwantifikatora „obydwa” (*both*) i na jej wsparcie przywołuje test Kiplana, który można również zastosować do naszego przykładu. Nazwijmy sąd wyrażony przez drugie zdanie w wypowiedzi (13) S . Przypuśćmy, że w jakichś innych okolicznościach chłopcy wchodzący do pokoju mieliby gołe głowy. Sąd wygłoszony w takim kontekście przy użyciu zdania „Każdy chłopiec miał czapkę” byłby fałszywy, ale S intuicyjnie pozostaje prawdziwy. Intuicje dotyczące prawdziwości i fałszywości S zdają się wiązać z losami tego zbioru indywiduów wprowadzonego do dyskursu, do którego anaforycznie odnosi się zwrot „każdy chłopiec”, nie zaś z jakimkolwiek zbiorem spełniającym dane własności. S intuicyjnie jest sądem jednostkowym. Van den Berg (1996: 6, 9, 36) sugeruje, że wszystkie wyrażenia rzeczownikowe mogą odnosić się anaforycznie do obiektów dyskursowych. Intuicje dotyczące jednostkowości/ogólności sądu nie mogą zatem uzasadniać w sposób dostateczny podziału na wyrażenia referencyjne i kwantyfikujące.

Druga standardowo przywoływana różnica między wyrażeniami referencyjnymi a kwantyfikującymi dotyczy tzw. ‘zachowania zasięgowego’. Przykładowo, porównajmy zdania:

(14) Jan podniósł jakiś stół.

(15) Każdy człowiek podniósł jakiś stół.

(14) jest jednoznaczne, a (15) wieloznaczne. Wieloznaczność (15) przypisuje się interakcji dwóch wyrażeń kwantyfikujących, dwóch operatorów z zasięgiem: bądź „każdy człowiek” ma szeroki zasięg i stoły mogą być różne dla każdego człowieka, bądź „jakiś stół” ma szeroki zasięg i to ten sam stół został podniesiony przez każdego człowieka. Jednoznaczność (14) tłumaczy się zwykle w ten sposób, że wyrażenia referencyjne (takie jak „Jan”) nie mają zasięgu, a zatem nie generują wieloznaczności zasięgowych. Ale różnice związane z zachowaniem zasięgowym można także tłumaczyć w inny sposób. Przede wszystkim podobny typ ‘wieloznaczności zasięgowych’, jak w przykładzie (14), można zaobserwować w przypadku nazw zbiorowych (pluralnych wyrażeń referencyjnych):

(16) Jan i Maria podnieśli jakiś stół.

Przykład (16) sugeruje, że za powstawanie ‘wieloznaczności zasięgowych’ odpowiada nie tyle kwantyfikująca natura wyrażenia, co jego pluralny charakter. Jak pisze Szabolcsi (2010: 104), jednostkowe (semantycznie) wyrażenia, jak w przykładzie (14), nie mogą w sposób oczywisty generować zmienności w referencji drugiego wyrażenia. Źródła ‘wieloznaczności zasięgowych’ nie należy zatem doszukiwać się w kwantyfikującej naturze wyrażenia, ale raczej w działaniu operatora dystrybutywności w rodzaju Δ van den Berga. Nie znaczy to, że wyrażenia „Jan i Maria” oraz „każdy mężczyzna” niczym się nie różnią. Jak pisało o tym wielu autorów (m.in. Kamp i Reyle 1993, Nouwen 2003), nazwy zbiorowe (także wyrażenia liczbowe, takie jak „trzy dziewczynki”) dopuszczają interpretację kolektywną: Jan i Maria (jako grupa) podnieśli stół. Wyrażenia kwantyfikujące zdają się zaś z natury być dystrybutywne. Zdanie (15) nie może znaczyć, że grupa osób uczestniczyła w podnoszeniu stołu i że ta grupa obejmowała wszystkich ludzi. Ale, jak pokazuje przykład (16), cecha dystrybutywności znowu nie może uzasadniać w sposób dostateczny podziału wyrażeń na referencyjne i kwantyfikujące.

Glanzberg (2008) i Szabolcsi (2010) sugerują, że pojęcie kwantyfikacji w języku naturalnym (czy szerzej, wyrażenia rzeczownikowego w ujęciu Szabolcsi) stanowi faktycznie przecięcie pewnej liczby niezależnych cech, które nie zawsze współwystępują w grupach odpowiadających klasycznemu podziałowi wyrażeń na referencyjne i kwantyfikujące. Jak pisze Szabolcsi (2010: 87), ten klasyczny podział zdaje się raczej „objawem lingwistycznej naiwności”. Bardzo podobna filozofia wyrażeń rzeczownikowych zdaje się stać za teoriami znaczenia w stylu modeli wielowymiarowych (Asher i Wang 2003, Dekker 2008,

Piasecki 2000). Główna idea modelu wielowymiarowego jest taka, że znaczenie wyrażen rzeczownikowych powinno być opisywane przy pomocy kombinacji pewnej liczby niezależnych aspektów znaczenia (cech), lokujących się na kilku poziomach. W modelu Dekkera są to poziom treści stwierdzanej (poziom warunków prawdziwości), treści presuponowanej oraz efektu dynamicznego. Na przykład zdanie (14) wprowadza do dyskursu indywiduum *d* (efekt dynamiczny), presuponuje, że *d* jest Janem i stwierdza, że *d* podniosło stół; zdanie (15) wprowadza do dyskursu maksymalny zbiór ludzi, którzy podnieśli stół (efekt dynamiczny), presuponuje istnienie maksymalnego zbioru ludzi (w ujęciu van den Berga ten zbiór będzie zawsze podzbiorem zbioru kontekstowego), wreszcie w postaci treści stwierdzanej wnosi warunek kwantyfikujący. Lista cech składająca się na efekt dynamiczny wyrażen rzeczownikowych może być oczywiście znacznie bogatsza i obejmować także m.in.: (a) relacje zależności pomiędzy elementami (mnogich) obiektów dyskursowych, (b) strukturę obiektów dyskursowych, tj. to, czy mówi się o grupach obiektów (interpretacja kolektywna), czy o pojedynczych indywiduach (odczytanie dystrybutywne), (c) cechę maksymalności⁹. Zbiór wszystkich cech (składających się na efekt dynamiczny, poziom treści stwierdzanej i presuponowanej) wyznacza przestrzeń możliwych znaczeń dla wyrażen rzeczownikowych. W naszym przekonaniu ustalenie pełnej listy cech wyrażen rzeczownikowych i scharakteryzowanie klasy ich dopuszczalnych konfiguracji obiecuje znacznie subtelniejszy od klasycznego opis zachowania wyrażen rzeczownikowych, uwzględniający dużo więcej niż jedynie dwie możliwe kombinacje cech (odpowiadające tradycyjnemu podziałowi na wyrażenia referencyjne i kwantyfikujące). Paradoksalnie przyjęcie unifikującej perspektywy na wyrażenia rzeczownikowe może prowadzić do teorii semantycznej dużo wrażliwszej na subtelności rzeczywistości lingwistycznej.

Bibliografia

- ASHER N., WANG L. (2003), *Ambiguity and anaphora with plurals in discourse*, "Proceedings of SALT" 13, <http://www.utexas.edu/cola/depts/philosophy/faculty/asher/papers.html>.
- BARWISE J. R. COOPER (1981), *Generalized quantifiers and natural language*, "Linguistics and Philosophy" 4 (2), 159–219.

⁹ Wyrażenia anaforyczne na dystrybutywnie interpretowanych wyrażeniach muszą odnosić się do maksymalnych zbiorów. Niepoprawny jest dyskurs: „Większość studentów napisało artykuł. *Być może, inni studenci też to zrobili”. Wyrażenia anaforyczne na kolektywnie interpretowanych wyrażeniach nie muszą odnosić się do maksymalnych zbiorów, o czym świadczy poprawność dyskursu: „Dwóch studentów napisało artykuł. Być może, inni studenci zrobili to samo”.

- BELLERT I. (1989), *Feature System for Quantification Structures in Natural Language*. Dordrecht: Foris Pub.
- BENTHEM J. VAN, MEULEN A. TER (red.) (1997), *Handbook of Logic and Language*. Amsterdam: Elsevier Science.
- BITTNER M. (2012), *Temporality: Universals and Variation*, książka w przygotowaniu, dostępna na stronie <http://www.rci.rutgers.edu/~mbittner/>.
- BRASOVEANU A. (2008), *Donkey Pluralities: Plural Information States Versus Non-Atomic Individuals*, "Linguistics and Philosophy" 31 (2), 129–209.
- CHIERCHIA G. (1992), *Anaphora and dynamic binding*, "Linguistics and Philosophy" 15 (2), 111–83.
- CHIERCHIA G. (1995), *Dynamics of Meaning. Anaphora, Presupposition, and the Theory of Grammar*. Chicago: The University of Chicago Press.
- DAVIDSON D. (1967), *Truth and Meaning*, "Synthese" 17, 304–323.
- DEKKER P. (2008a), *A multi-dimensional treatment of quantification in extraordinary English*, "Linguistics and Philosophy" 31 (1), 101–127.
- DEKKER P. (2008b), *A guide to dynamic semantics*, "ILLC Prepublications" PP-2008–42.
- EIJCK J. VAN, KAMP H. (1997), *Representing discourse in context*, [w:] Benthem J. van, Meulen A. ter (red.), 179–238.
- EVANS G. (1977a), *Pronouns, Quantifiers, and Relative Clauses (I)*, "Canadian Journal of Philosophy" 7(3), 467–536.
- EVANS G. (1977b), *Pronouns, Quantifiers, and Relative Clauses (II): Appendix*, "Canadian Journal of Philosophy" 7 (4), 777–797.
- EVANS G. (1980), *Pronouns*, "Linguistic Inquiry" 11 (2), 337–362.
- FODOR J. D., SAG. I. (1982), *Referential and Quantificational Indefinites*, "Linguistics and Philosophy" 5 (3), 355–398.
- GLANZBERG M. (2008), *Quantification and Contributing Objects to Thoughts*, "Philosophical Perspectives" 22, 207–231.
- GROENENDIJK J., STOKHOF. M. (1991), *Dynamic Predicate Logic*, "Linguistics and Philosophy" 14 (1), 39–100.
- HEIM I. (1990), *E-type pronouns and donkey anaphora*, "Linguistics and Philosophy" 13 (2), 137–78.
- KAMP H. (1981), *A theory of truth and semantic representation*, [w:] Groenendijk J., Janssen T., Stokhof M. (red.), *Truth, Interpretation and Information*. Dordrecht: Foris, 1–41.
- KAMP H., REYLE U. (1993), *From Discourse to Logic*. Dordrecht: Kluwer.
- KARTTUNEN L. (1971), *Discourse Referents*. Technical report, RAND Corporation, Reproduced by the Indiana University Linguistics Club.
- LUDLOW P., S. NEALE. (1991), *Indefinite Descriptions: In Defense of Russell*, "Linguistics and Philosophy" 14 (2), 171–202.
- MONTAGUE R. (1973), *The proper treatment of quantification in ordinary English*, [w:] Hintikka J., Moravcsik J., Suppes P. (red.), *Approaches to Natural Language*. Dordrecht: Reidel, 221–242.
- MUSKENS R. (1996), *Combining Montague semantics and discourse representation*, "Linguistics and Philosophy" 19 (2), 143–186.
- NEALE S. (1990) *Descriptions*. Cambridge: MIT Press Books.
- NOUWEN R. (2003), *Plural pronominal anaphora in context: dynamic aspects of quantification*, Ph.D. thesis, UiL-OTS, Utrecht, LOT dissertation series, No. 84.

- NOUWEN R. (2007), *On dependent pronouns and dynamic semantics*, "Journal of Philosophical Logic" 36 (2), 123–154.
- PIASECKI M. (2000), *The Multidimensional Approach to Quantification and Reference in the Noun Phrase*, [w:] Dorfmueller-Karpusa K., Vretta-Panidou E. (red.), *Thessaloniker interkulturelle Analysen. Akten des 33. Linguistischen Kolloquiums*. Frankfurt/M: Peter Lang, 239–248.
- PIASECKI M. (2005), *Compositional, Variable Free, Binding Free and Structure Oriented Discourse Interpretation*, [w:] Vetulani Z. (red.), *Human Language Technologies as a Challenge for Computer Science and Linguistics, Proceedings of the 2nd Language & Technology Conference Poznań*. Poznań: Wydawnictwo Poznańskie, 395–399.
- RUSSELL B. (1905), *On Denoting*, „Mind” 14, 479–493.
- RUSSELL B. (1967), *Denotowanie*, [w:] Pelc J. (oprac.), *Logika i język: studia z semiotyki logicznej*. Warszawa: PWN, 253–275.
- STRAWSON P. F. (1950), *On Referring*, „Mind” 59, 320–334.
- STRAWSON P. F. (1967), *O odnoszeniu się użycia wyrażen do przedmiotów*, [w:] Pelc J. (oprac.), *Logika i język: studia z semiotyki logicznej*. Warszawa: PWN, 377–413.
- SZABOLCSI A. (2010), *Quantification*. Cambridge: Cambridge University Press.
- BERG M. H. VAN DEN (1996), *The Internal Structure of Discourse*, Ph.D. thesis. Universiteit van Amsterdam, Amsterdam.
- TARSKI A. (1944), *The Semantic Conception of Truth and the Foundations of Semantics*, „Philosophy and Phenomenological Research” 4, 341–376.
- TARSKI A. (1995), *Semantyczna koncepcja prawdy i podstawy semantyki*, [w:] tenże, *Pisma logiczno-filozoficzne*, t. I. Prawda (przeł. J. Zygmunt). Warszawa: PWN, 228–283.
- ZAWADOWSKI M. (1989), *Formalization of the feature system in terms of pre-orders*, [w:] Bellert I., 155–175.

Rozdział 10

Co to jest „termin jednostkowy”?

Mieszko Tałasiewicz (Uniwersytet Warszawski)

Słowa kluczowe: termin jednostkowy, termin ogólny, podmiot, predykat, gramatyka kategorialna, kategoria semantyczna nazwy, wyrażenie nienasycone, Geach Peter T., Dummett Michael, Strawson Peter F., wartość semantyczna wyrażenia

Terminy jednostkowe (*singular terms*) są przedmiotem długiej, ale wciąż ożywionej debaty¹. Wiele poglądów na ich naturę ogniskuje się wokół dwóch bazowych intuicji.

W myśl pierwszej intuicji pojęcie terminu jednostkowego obejmuje m.in. takie własności, które sprawiają, że terminy jednostkowe skontrastowane są z predykatami. W szczególności terminy takie uważa się za kompletne/nasycone, podczas gdy predykaty są niekompletne/nienasycone (w Fregowskim sensie):

[Singular term] is an expression that can function as the subject of the sentence but not as the predicate (Stirton 2000: 196).

Singular terms form one of the two categories of *complete* expressions, the other being that of *sentences*. In devising criteria for singular terms, we may presuppose a capacity to recognize well-formed sentences as such (Hale 2001a: 33).

Predicates can be defined in terms of singular terms and sentences as follows: a predicate is what remains when one or more singular terms has been removed from a sentence. (We assume that we have some means for determining what a sentence is) (Wetzel 1990: 239)².

¹ 'A fair amount of ink has been spilt by recent philosophers of language in the attempt to achieve a precise definition of a class of English expressions which they call 'singular terms'(Stirton 2000: 191).

² Cytaty z dzieł angielskojęzycznych zdecydowałem się podawać w oryginale. Niniejszy tekst jest adresowany do osób zorientowanych w problematyce współczesnej filozofii języka, a zatem bez wątpienia znających język angielski przynajmniej biernie w stopniu wystarczającym do zrozumienia tych cytatów. Przekład jest już interpretacją; nie ma powodu, bym

W myśl drugiej intuicji pojęcie terminu jednostkowego obejmuje m.in. takie własności, które sprawiają, że terminy jednostkowe skontrastowane są z terminami ogólnymi. W szczególności wartością semantyczną takich terminów, jak się uważa, są pojedyncze indywidua:

Semantically, a singular term is any expression the function of which is to convey a reference to a particular object (Hale 2001a: 31). [Genuine singular terms] take individuals as their semantic values (Justice 2007: 363), [or are engaged in certain] relationship [...] to just one object (Devitt 1974: 183).

Powstaje pytanie, jak te intuicje mają się do siebie.

Po pierwsze (I), można uznać — co czyni większość autorów — że pomimo tego dwoistego zaplecza intuicyjnego pojęcie terminu jednostkowego nie jest dwuznaczne. Autorzy ci uważają, że terminy jednostkowe spełniają obie intuicje łącznie; więcej nawet: że obie te intuicje nie są niezależne. Można sądzić na przykład — ich zdaniem — iż to, że terminy jednostkowe są nasycone, zależy jakoś od tego, że odnoszą się do pojedynczych indywiduów. Skoro zaś te intuicje są systematycznie powiązane, to powiązanie to może być w jakiś sposób teoretycznie wykorzystane.

Jednym ze sposobów takiego wykorzystania (A) jest znajdowanie jakichś semantyczno-ontologicznych wyróżników terminów jednostkowych w opozycji do terminów ogólnych (na podstawie domniemanej w drugiej intuicji ich specjalnej relacji do indywiduów) i następnie używanie tych wyróżników jako kryteriów wyodrębniania podmiotu i predykatu (czy szerzej — jako narzędzi badania składni logicznej). Drugim — poniekąd przeciwnym sposobem takiego wykorzystania (B) — jest znajdowanie jakichś syntaktyczno-logicznych wyróżników terminów jednostkowych w opozycji do predykatów (na podstawie domniemanej w pierwszej intuicji ich specjalnej roli składniowej) i następnie używanie tych wyróżników jako narzędzi badania ontologii, na przykład do poszukiwania rodzajów indywiduów.

Na sposób (A) próbował wykorzystać połączenie obu tych intuicji P. F. Strawson. Szczegółowo opisałem jego przedsięwzięcie w (Tałasiewicz 2011).

w tym miejscu narzucał czytelnikowi — świadomie lub co gorsza nieświadomie — jakąś interpretację. Artykuł zawiera pewną oryginalną propozycję badawczą, rzuconą na pewne tło literaturowe, zarysowane w tych cytatach, i do samego tego tła nie chciałbym niczego „domalowywać” — aby wyraźniej widać było, na czym polega moja propozycja. Rozumiem oczywiście wagę zadania polegającego na spolszczaniu terminologii filozoficznej; sądzą jednak, że zadanie to powinno się podejmować nie mimochodem, w tekstach, w których co innego jest głównym zamysłem, lecz w specjalnie do tego przeznaczonych pracach: przekładach najcenniejszych dzieł obcych, opracowanych przez specjalizujących się w danej tematyce tłumaczy, oraz w specjalnych artykułach sprawozdawczych (nb. z tego powodu nie podzielałam pojawiającej się niekiedy opinii, że prace „czysto” sprawozdawcze są bezwartościowe: nie są — a ich wielka wartość polega właśnie na opisanu po polsku jakiegoś fragmentu mapy nowo odkrywanego w filozofii terytorium). Niniejsza praca nie stawia sobie jednak takiego celu, a podawanie obu wersji językowych wydaje mi się zupełnie nieekonomiczne.

W tym miejscu wystarczy wspomnieć, że zaprojektował on całą ontologię — przedstawioną w słynnej książce *Individuals* — w tym celu, by poprzez ontologiczne wyodrębnienie indywiduów i pojęć (*individuals and concepts*) ugruntować logiczne rozróżnienie podmiotu i predykatu; przez pewien czas używał nawet wyrażenia ‘termin jednostkowy’ zamiennie z wyrażeniem ‘podmiot’. Przedsięwzięcie to okazało się jednak porażką (a nawet, jak argumentuję we wspomnianej pracy, było skazane na porażkę od samego początku).

Sposobu (B) próbowało wielu autorów. Stirton (2000) dzieli próby tego rodzaju na dwie „tradycje”: Arystotelesowską (i) i Dummettowską (ii), w zależności od tego, jakiego rodzaju logiczno-syntaktyczne kryterium identyfikacji terminów jednostkowych jest w nich przyjęte.

Kryterium Arystotelesowskie (I.B.i) mówi, że terminy jednostkowe nie mogą być negowane, podczas gdy predykaty mogą. Filozofem, który próbował podać wyjaśnienie rozróżnienia podmiotu i predykatu na bazie różnicy w działaniu negacji, był na przykład P. T. Geach (1980). Próba ta, jak uzasadniał Strawson — w (Tałasiewicz 2011) omówiłem z aprobatą argumenty Strawsona w tej materii — również nie może się powieść, bowiem różnice w negowalności nie wyjaśniają rozróżnienia podmiotu i predykatu: przeciwnie, to one same domagają się wyjaśnienia i mogą je znaleźć dopiero wtedy, kiedy uprzednio dowiemy się, na jakiej podstawie wyróżnia się podmiot i predykat.

Kryterium Dummettowskie (I.B.ii) składa się z kilku testów inferencyjnych, które — mówiąc z grubsza — polegają na tym, że jeśli jakiś termin *t* jest terminem jednostkowym w pewnych zdaniach, to pewne inne zdania mogą (lub właśnie nie mogą) być z tych zdań wyprowadzone³. Głównym zwolennikiem tego kryterium — stąd jego nazwa — jest sam Michael Dummett, a prócz niego także np. Bob Hale i Crispin Wright.

Niestety, perspektywy kryterium Dummettowskiego również nie wyglądają różowo, ponieważ również do niego stosują się te zarzuty, które Strawson podnosił przeciwko kryterium Arystotelesowskiemu (zob. wyżej), a mianowicie przede wszystkim zarzut *ignotum per ignotum*. Kryterium Dummettowskie — pewien wzorzec inferencyjny stowarzyszony z występowaniem w zdaniach terminów jednostkowych — nie może wyjaśnić rozróżnienia podmiotu i predykatu z tego samego powodu, z którego w kryterium Arystotelesowskim nie może tego zrobić wzorzec stosowania negacji. To sam ten wzorzec wymaga wyjaśnienia. Zanim się go zastosuje do wyjaśnienia terminów jednostkowych, należy odpowiedzieć na pytanie: co takiego jest w terminach jednostkowych, co sprawia, że zdania z tymi terminami zachowują się tak, a nie inaczej w inferencjach?

³ Dokładne sformułowanie kryterium Dummettowskiego jest dość skomplikowane, ma wiele wariantów, a każdy z nich w dodatku był niezliczoną ilość razy cytowany w literaturze; powstrzymam się więc od podania kolejnego cytatu — rozmaite wersje można znaleźć m.in. w (Dummett 1973), (Hale 1979), (Wetzel 1990), (Stirton 2000), (Hale 2001a) czy (Hale 2001b).

Niezależnie od tego zarzutu można postawić inne. Zacytujmy manifest zwolenników kryterium Dummetta:

[W]e set about the task, armed with some firm intuitions (about which expressions ought, and which ought not, to be classified as singular terms) about a limited range of cases; and are then looking for a set of general tests which deliver the right verdict in the cases about which we are confident and additionally rule on further cases, where intuitions can't be relied upon (Hale 2001b: 69).

Jak widać, ich punktem wyjścia są ‘solidne intuicje co do tego, które wyrażenia powinny, a które nie powinny liczyć się za terminy jednostkowe’. Jednakże krytyka przedstawiona w (Wetzel 1990) czy (Stirton 2000) pokazuje, że w tej materii intuicje filozofów są jak najdalsze od najsłabiej nawet rozumianej ‘solidności’. Kryterium Dummetta w szczególności zalicza do grona terminów jednostkowych pewne wyrażenia, które powinny być — intuicyjnie (także w myśl intuicji samego Dummetta i jego zwolenników) — wykluczone, a wyklucza takie, które powinny być zaliczone. Wszystkie próby naprawienia tego kryterium skutkują dodawaniem do niego kolejnych epicykli i komplikowania go ponad wszelką rozsądną miarę.

Nie ma sensu przepisywać tu krytyki Wetzel czy Stirtona, warto jednak dodać od siebie jeszcze jedną ilustrację trudności, w jakie popada kryterium Dummettowskie. Oto Bob Hale (2001a: 34) obwieszcza, że częścią zadania rozjaśnienia pojęcia terminu jednostkowego jest zadanie odróżnienia terminów jednostkowych od szerszej klasy wyrażen rzeczownikowych (*substantial expressions*), która zawiera m.in. charakterystyczne dla języka naturalnego słowa i zwroty kwantyfikatorowe. Klasę tę, według Hale’a, możemy wyróżnić metodami czysto syntaktycznymi, jej elementami bowiem są dokładnie te wyrażenia, które mogą występować *salva congruitate* tam, gdzie mogą wystąpić nazwy własne w zwyczajnym sensie.

Tymczasem jednak pojęcie nazwy własnej nie ma w filozofii zwyczajnego sensu: jest to jedno z najbardziej „naładowanych” teoretycznie pojęć. Najbliższym kandydatem na taki standardowy sens jest sens przyjmowany w ujęciu Saula Kripkego. Ujęcie to jednak szczególnie wyraźnie eksponuje kauzalny związek pomiędzy nazwami a nazywanymi obiektami, co pasuje do kierunku wyjaśniania A, ale nie do kierunku B. Zwłaszcza Hale nie może z tego związku korzystać: jego teoria terminów jednostkowych ma na celu wypracowanie narzędzi poszukiwawczych w ontologii matematyki (Hale 2001a: 32). Uzależnianie tego narzędzia od związków kauzalnych między bytami matematycznymi a praktyką językową nie ma najmniejszego sensu.

Co więcej, nawet gdybyśmy przypuścili, że nazwy własne są jakoś po prostu dane (bez zakładania związków kauzalnych z nazywanymi obiektami), nie moglibyśmy uzyskać klasy wyrażen rzeczownikowych — jeśliby miała ona obejmować frazy kwantyfikatorowe — na podstawie relacji wymienialności

salva congruitate z nazwami własnymi. A to dlatego, że nazwy własne po prostu nie są wymienne z frazami kwantyfikatorowymi *salva congruitate*, czyli — jak to ujmuje Hale — „bez niszczącego wpływu na strukturę zdania” (*‘without destructive effect upon the grammar of a sentence’* (Hale 2001a: 34)).

W tej sprawie należy się kilka słów wyjaśnienia. Istotnie, od publikacji Geacha (1970) często utrzymywano, że w zdaniach typu *‘Raleigh smokes’* słowo *‘smokes’* jest funktorem, którego argumentem jest *‘Raleigh’*, a w zdaniach typu *‘Every man smokes’* fraza *‘Every man’* jest funktorem, natomiast *‘smokes’* — jego argumentem⁴. Analiza semantyczno-kategorialna⁵ wygląda więc odpowiednio:

- | | |
|-----------------|-------------|
| (1) Raleigh | smokes |
| (n) (1,1) | (s/n) (1,0) |
| (2) Every man | smokes |
| (s/(s/n)) (1,0) | (s/n) (1,1) |

Na pierwszy rzut oka może się wydawać, że skoro zarówno nazwa własna (n), jak i fraza kwantyfikatorowa (s/(s/n)) łącznie z predykatem (s/n) mogą stworzyć zdanie (s), to nazwa i ta fraza są wymienne. W rzeczywistości nie są — w każdym razie nie „bez niszczącego wpływu na strukturę zdania”. Struktura zdania to nie tylko wybór kategorii semantycznych, ale także ich układ, ich *pozycje syntaktyczne*. Jeżeli wymienimy nazwę *‘Raleigh’* na frazę kwantyfikatorową *‘Every man’*, musimy odpowiednio zmienić pozycję syntaktyczną predykatu *‘smokes’*: w (1) jest to operator, w (2) — nie jest. A zatem wymiana nazwy własnej na frazę kwantyfikatorową zachowuje jedynie „zdaniowość” zdania, ale nie zachowuje struktury gramatycznej tego zdania. Zarówno (1), jak i (2), są zdaniami, ale mają istotnie różną składnię⁶.

W rezultacie nadal nie wiemy, czym są *‘wyrażenia rzeczownikowe’*, a tym samym nie wiemy, czym są terminy jednostkowe.

Problem tkwi w tym, że połączenie wyjściowych dwóch intuicji leżących u podstaw koncepcji terminów jednostkowych, tak często zakładane bezdyskusyjnie, w istocie bynajmniej nie jest nieproblematyczne. Istnieje mianowicie poważny powód, by sądzić, że wyjściowe intuicje wcale nie są zależne, a co za

⁴ Są powody, by sądzić, że to jest błędna analiza (zob. Tałasiewicz 2010). To nie ma jednak znaczenia dla niniejszych wywodów.

⁵ Stosuję notację kategorii semantycznych z artykułu (Ajdukiewicza 1935), a pozycji syntaktycznych z artykułu (Ajdukiewicz 1960).

⁶ Z tego powodu należy bardzo podejrzliwie podchodzić do wszelkiego rodzaju reguł podnoszenia typu. Geach często się na nie skarżył. Pisał m.in.: *‘[distinguishing a name and a quantified phrase] is a profound insight, ignored by those who would lump together proper names and phrases like ‘every man’ as Noun Phrases [albo ‘substantial expressions’, moglibyśmy dodać...]; [in fact] we have two different syntactical categories’* (Geach 1970: 4). Smutną ironią może wydawać się fakt, że to ci właśnie, którzy ignorowali ten „głęboki wgląd”, nazwali reguły podnoszenia typów „regułami Geacha”. Więcej o tym nieporozumieniu można przeczytać w (Humberstone 2005) i (Tałasiewicz 2009).

tym idzie, pojęcie terminu jednostkowego ugruntowane w jednej z nich jest inne niż pojęcie ugruntowane w drugiej: ‘termin jednostkowy’ jest wtedy wyrażeniem dwuznacznym, a rozumowania podstawiające pod to wyrażenie oba te znaczenia są ekwiwokacjami.

Powodem tym jest to, że połączenie obu intuicji narusza bardzo ważne rozróżnienie pomiędzy **rodzajami** wyrażen a **rolami składniowymi** pełnionymi przez te wyrażenia. Rozróżnienie to jest przedstawione *in extenso* wraz z uzasadnieniem tezy, że jest to rozróżnienie ważne, w (Tałasiewicz 2012). W tym miejscu przytoczę tylko dwa najbardziej charakterystyczne fragmenty:

Nazwy, zdania i funktry to nie są wyrażenia o pewnych własnościach, lecz **role** składniowe, jakie mogą być pełnione przez rozmaite wyrażenia. Niektóre wyrażenia lepiej się nadają do pełnienia pewnych ról niż inne, to jasne. Ale są to tylko różnice stopnia naturalności pewnych konstrukcji językowych; na takich różnicach nie można budować logiki języka. Generalnie to samo wyrażenie może w rozmaitych okolicznościach pełnić rozmaite role. Na przykład wyrażenie ‘żółty’ może pełnić rolę nazwy, jak w zdaniu ‘żółty jest jasnym kolorem’, lub funktrora, jak w zdaniu ‘żółty samochód uderzył w drzewo’. Dla zachowania przejrzystości pozostaniemy przy tradycyjnym sposobie wysłowienia i dalej będziemy nazywać zdania, nazwy i funktry *kategoriami*, trzeba mieć jednak zawsze w pamięci, że jest to w gruncie rzeczy nadużycie terminologiczne. Kategorie to nie są zbiory wyrażen, tylko role składniowe. Różni aktorzy mogą grać tę samą rolę, ten sam aktor może grać różne role (Tałasiewicz 2012: 155).

Być może dla ostatecznego wyjaśnienia, o co chodzi w odróżnianiu z jednej strony ról składniowych od wyrażen pełniących te role, a z drugiej ról składniowych jako takich od układów tych ról w danym wyrażeniu złożonym (pozycji syntaktycznych), warto będzie posłużyć się poglądową metaforą. Rozważmy zadanie doprowadzenia do porządku trawnika pokrytego opadłymi liśćmi. Wypełnienie tego zadania rozkłada się na wypełnienie następujących ról: grabienie liści, zbieranie liści do pojemnika, transportowanie pojemnika z liśćmi do kompostowego dołu, opróżnianie pojemnika, transportowanie pustego pojemnika z powrotem do trawnika, dalsze grabienie liści, zbieranie itd. aż trawnik jest oczyszczony. Te role mogą być pełnione przez różne obiekty, spośród których niektóre lepiej się do danej roli nadają, inne gorzej. Grabienie na przykład wygodnie robić grabiami. Ale miotła domowa, z braku lepszych instrumentów, też się nada, albo parę patyków w garści. Nawet gołymi rękami można zbierać liście, choć nie jest to ani miłe, ani wygodne, ani efektywne. Tak więc słowo ‘grabie’ może oznaczać przedmiot zaprojektowany do grabienia, ale może też oznaczać cokolwiek, co pełni funkcję grabienia (czyli *de facto* samą tę rolę — wypełnić ją bowiem może praktycznie każdy przedmiot, który jest fizycznie zdolny przesunąć liść). To jest pierwsze odróżnienie. Jeżeli teraz będziemy mówić już o samych rolach, nie przedmiotach, które je pełnią, zobaczymy, że potrzebujemy jeszcze dalszej charakterystyki procesu porządkowania trawnika: aby ten proces zakończył się sukcesem, poszczególne role muszą zostać odegrane w odpowiednim *porządku* (zupełnie niezależnie od tego, jakie przedmioty je pełnią) — i to jest drugie odróżnienie. Jeżeli najpierw grabilibyśmy liście, potem zbierali je do pojemnika, potem opróżniali pojemnik (z powrotem na trawnik), potem transportowali pusty pojemnik do kompostu, potem

transportowali ów pusty pojemnik z powrotem do trawnika, potem znowu grabili itd. — nigdy nie doprowadzilibyśmy trawnika do porządku. Tak więc rozróżnienie pomiędzy wyborem ról a ich układem w konkretnym zastosowaniu nie powinno zamazywać rozróżnienia pomiędzy rolami a obiektami, które mogą je pełnić (Tałasiewicz 2012: 158–9).

W naszym wypadku chodzi o to, że bycie podmiotem (a nie predykatem) i bycie w relacji desygnowania z pojedynczym przedmiotem (a nie z wieloma przedmiotami) to są dwie zupełnie różne, niezależne rzeczy. W rzeczywistości z jednej strony istnieją podmioty desygnujące pojedyncze obiekty ('**chłopiec** kopnął piłkę') i podmioty desygnujące wiele przedmiotów ('**chłopaki** nie płaczą'), a z drugiej strony wyrażenia desygnujące pojedynczy przedmiot stoją niekiedy w pozycji podmiotów ('**Najwyższa góra świata** znajduje się w Himalajach'), a niekiedy predykatów ('Mount Everest jest **najwyższą górą świata**'). Dwie podstawowe intuicje dotyczące pojęcia terminu jednostkowego *de facto* nie mają ze sobą nic wspólnego i oczekiwanie jakiejś systematycznej koincydencji pomiędzy nimi jest nierozsądne.

Właściwe podejście do zagadnienia terminów jednostkowych powinno zatem wyrażać się w tezie (II): pojęcie terminu jednostkowego jest dwuznaczne (a nawet mocniej: są dwa niezależne pojęcia terminu jednostkowego). W jednym znaczeniu wyrażenie 'termin jednostkowy' jest dokładnie synonimiczne z wyrażeniem 'podmiot'; w drugim znaczeniu 'termin jednostkowy' oznacza wyrażenie, którego wartością semantyczną jest pojedynczy obiekt.

Opracowanie pierwszego znaczenia znajduje się w (Tałasiewicz 2012): z grubsza chodzi tu o nazwę w sensie gramatyki kategorialnej. Drugie znaczenie musi być opracowywane dalej. Krótki rzut oka wystarczy by dostrzec, że jest tu sporo niebanalnych zagadnień do rozstrzygnięcia, gdyż wyrażenia mogą wskazywać swoją wartość semantyczną na wiele sposobów. Porównajmy następujące przykłady (przykłady są angielskie, ponieważ ze względu na występowanie rodzajników angielszczyzna pozwala łatwo uwidocznić zróżnicowania, które są w polszczyźnie zatarte):

- (3) **Anna** left the house in hurry.
- (4) **A girl** left the house in hurry.
- (5) **She** left the house in hurry.
- (6) **This girl** left the house in hurry.
- (7) **Someone** left the house in hurry.
- (8) **The prime minister** left the house in hurry
- (9) **The prime minister** is entitled to declare war.

Wszystkie pogrubione wyrażenia mają w jakimś sensie pojedynczy obiekt jako wartość semantyczną, w każdym wypadku jednak jest to jakiś uchwytne inny

sens. (3) to nazwa własna, (4) to deskrypcja nieokreślona, (5) zaimek, (6) demonstratyw, (7) fraza kwantyfikatorowa, (8) deskrypcja określona użyta referencyjnie, (9) deskrypcja określona użyta atrybutywnie.

Które z nich zasługują na miano terminów jednostkowych? Jakie dodatkowe kryteria powinniśmy przyjąć?

Można to rozstrzygać rozmaicie. Jedną z opcji jest przyjęcie poglądu Michaela Devitta, który twierdzi, że relacja pomiędzy wyrażeniem a jego wartością semantyczną to ta, która ma kluczowe znaczenie dla wartości logicznej zdań zawierających to wyrażenie (*'the relationship in question is picked out by its crucial bearing on the truth value of sentences containing the [expression]'* (Devitt 1974: 183)). Takie kryterium uznaje za terminy jednostkowe tylko tzw. sztywne desygnatory. Tego rodzaju wyrażenia występują w (3), (5), (6) i (8)⁷. Inną opcję wybiera John Justice, który za terminy jednostkowe uważa wszystkie wyrażenia denotujące zbiory jednostkowe, do których należą obiekty, które jako jedyne spełniają warunki wyznaczone przez znaczenie tych wyrażen (*'[expressions that denote singletons containing the individuals that uniquely satisfy their sense-conditions]'* (Justice 2007: 368)) — sztywność w tym wypadku jest sprawą regulowaną pozytywnie lub negatywnie w ramach tych warunków określonych w znaczeniu. To kryterium pozwala zaliczyć do terminów jednostkowych, przynajmniej *prima facie*, wszystkie pogrubione wyrażenia od (3) do (9).

Jest również wiele innych opcji, jednak szczegółowa dyskusja nad tym, co warto, a czego nie warto uważać za termin jednostkowy w sensie drugim, wykracza poza ramy niniejszego artykułu (i raczej ma niewiele wspólnego z zagadnieniami logiki języka — dla tych zagadnień kluczowy jest sens pierwszy: termin jednostkowy jako podmiot).

Bibliografia

- AJDUKIEWICZ K. ([1935] 1985), *O spójności syntaktycznej*, [w:] *Język i poznanie*, t. I. Warszawa: PWN, 222–242.
- AJDUKIEWICZ K. ([1960] 1985), *Związki składniowe między członami zdań oznajmujących*, [w:] *Język i poznanie*, t. II. Warszawa: PWN, 344–355.
- DEVITT M. (1974), *Singular Terms*, "The Journal of Philosophy", Vol. 71, No. 7, 183–205.
- DUMMETT M. (1973), *Frege: Philosophy of Language*. London: Duckworth.
- DUMMETT M. (1993), *Origins of Analytical Philosophy*, Cambridge, MA: Harvard University Press.
- GEACH P. T. (1970), *The Program of Syntax*, "Synthese" 22, 3–17.
- GEACH P. T. (1980), *Reference and Generality*. Ithaca, NY: Cornell University Press.

⁷ Devitt wylicza tu kategorie nazw własnych, demonstratywów, zaimków osobowych i deskrypcji określonych w użyciu referencyjnym w sensie Donellana (Devitt 1974: 184).

- HALE B. (1979), *Strawson, Geach and Dummett on Singular Terms and Predicates*, "Synthese" 42, 275–295.
- HALE B. (2001a), *Singular Terms (I)*, [w:] HALE B., WRIGHT C. (2001), 31–47.
- HALE B. (2001b), *Singular Terms (II)*, [w:] HALE B., WRIGHT C. (2001), 48–71.
- HALE B., WRIGHT C. (2001), *The Reason's Proper Study: Essays towards a Neo-Fregean Philosophy of Mathematics*. Oxford: Oxford University Press.
- HUMBERSTONE L. (2005), *Geach's Categorical Grammar*, "Linguistics and Philosophy" 28, 281–317.
- JUSTICE J. (2007), *Unified Semantics of Singular Terms*, "The Philosophical Quarterly" 57, No. 228, 363–373.
- STIRTON W. R. (2000), *Singular Term, Subject and Predicate*, "The Philosophical Quarterly" 50, No. 199, 191–207.
- STRAWSON P. F. (1959), *Individuals*. London: Routledge.
- TALAŚIEWICZ M. (2009), *Geach's Program and the Foundations of Categorical Grammar*, "Acta Semiotica Fennica" XXXIV, 1731–1742.
- TALAŚIEWICZ M. (2010), *Philosophy of Syntax. Foundational Topics*. Dordrecht: Springer.
- TALAŚIEWICZ M. (2011), *Podmiot i predykat: krytyka koncepcji Petera F. Strawsona*, „Przegląd Filozoficzny” 3 (79), 209–226.
- TALAŚIEWICZ M. (2012), *Podmiot i predykat: filozoficzne podstawy gramatyki kategoryalnej*, „Przegląd Filozoficzny” 1 (81), 147–165.
- WETZEL L. (1990), *Dummett's Criteria for Singular Terms*, "Mind" 99, No. 394, 239–254.

Rozdział 11

Czy „kot” jest racją? Spór o normatywność znaczenia

Joanna Klimczyk (Instytut Filozofii i Socjologii PAN)

Słowa kluczowe: znaczenie, warunki poprawnego użycia, normatywność znaczenia, standard, preskryptywny, racja

1. Wstęp

Teza *semantycznego normatywizmu* (dalej SN) głosi, że znaczenie językowe jest wewnętrznie normatywne. Daniel Whiting, główny orędownik tego stanowiska, tak oto precyzuje jego istotę:

Fakty dotyczące znaczenia [...] są wewnętrznie kierujące działaniem [*action-guiding*] albo preskryptywne; w szczególności mają implikacje dla tego, co podmiot *może* lub *powinien* (nie może, nie powinien) zrobić (Whiting 2009a: 536)*.

Inaczej rzecz ujmując, SN można scharakteryzować jako pogląd, wedle którego „to, co wyrażenie znaczy, (samo) dostarcza *racji* do używania (lub nieużywania) go w określony sposób” (Whiting 2009a: 537).

Na pierwszy rzut oka mogłoby się wydawać, że tak jasno sformułowana teza pozwala uczestnikom debaty dojść do równie klarownej konkluzji co do tego, czy znaczenie jest czy też nie jest normatywne w przedstawionym wyżej sensie. A przynajmniej przedstawić argumenty i kontrargumenty dotyczące *konkretnej* tezy, która jest zarzewiem sporu. Nic bardziej mylnego. Dyskusja wokół tego, czy znaczenie *samo w sobie* w rozumieniu realizmu semantycznego jest normatywne, rozpoczęła się po ukazaniu się książki Saula Kripkego *Wittgenstein o regułach i języku prywatnym* pod koniec lat 80. XX wieku, apogeum osiągnęła w latach 90., a ponieważ nie zakończyła się definitywnym rozstrzygnięciem, wciąż powraca, choć obecnie już nie jako samodzielny temat, lecz

* Wszystkie przekłady zamieszczone w niniejszym tekście, o ile nie zaznaczono inaczej, są tłumaczeniem własnym.

raczej jako część argumentu za lub przeciw tezie dotyczącej normatywności treści umysłowej (zob. Glüer, Wikforss 2009a, b; Boghossian 2003, 2005a; Gibbard 2003; Steglich-Petersen 2010). Jak zauważyły pionierki semantycznego antynormatywizmu, Kathrin Glüer i Åsa Wikforss, obecnie slogan głoszący, że znaczenie jest normatywne, stracił na atrakcyjności, a niektórzy z jego obrońców przeszli do obozu krytyków (najbardziej spektakularnym przykładem jest Paul Boghossian, który w swoim artykule (Boghossian 1989) utrzymywał, że ze względu na istotny związek między znaczeniem a warunkami poprawności znaczenie jest normatywne, a w ostatnich latach zmienił zdanie i dołączył do grona krytyków semantycznego normatywizmu (Boghossian 2003: 39)).

Dlaczego SN wzbudza dziś mniej emocji i nie jest już w centrum debat toczonych przez filozofów języka i umysłu, skoro dyskusja nie przyniosła zadowalającej odpowiedzi na pytanie, czy znaczenie jest wewnętrznie normatywne? Co sprawiło, że filozofowie zgodnie uznali, że debata utknęła w martwym punkcie, sprowadzając się już wyłącznie do zderzenia intuicji? Intuicje zwolenników SN nakazują im wierzyć, że za przyjęciem SN przemawiają dobre racje, podczas gdy intuicje przeciwników podpowiadają im, że SN jest poglądem nie do utrzymania.

W niniejszym artykule zamierzam: 1) pokazać, że dyskusja między semantycznymi normatywistami i antynormatywistami utknęła w martwym punkcie, ponieważ obie strony sporu odmiennie interpretują tezę SN, czego konsekwencją jest częste mijanie się w argumentach; 2) krytycznie odnieść się do popularnego w literaturze przedmiotu przekonania, że obie strony dysponują równie silnymi argumentami, których źródło i zarazem siła tkwią w intuicjach; wreszcie, 3) wskazać główne wady i niekonsekwencje w argumentacji obu stron, które potęgują chaos pojęciowy, stanowiący charakterystyczną cechę tej dyskusji.

2. Geneza sporu

Jedna rzecz wydaje się nie budzić kontrowersji zarówno wśród normatywistów, jak i antynormatywistów, a mianowicie geneza sporu. Dyskusję, jak się zgodnie przyjmuje, wywołała książka Kripkego poświęcona interpretacji Wittgensteinowskiego podejścia do reguł i języka prywatnego. We wspomnianej publikacji Kripke posłużył się argumentem z normatywności jako bronią przeciwko dyspozycjonalistycznym teoriom znaczenia. Argument Kripkego z normatywności można przedstawić następująco: jeśli mam na myśli *kota*, gdy używam słowa „kot”, wówczas to, że mam na myśli kota, zobowiązuje mnie do używania słowa „kot” zgodnie z jego ekstensją, w odniesieniu do kotów, a nie do nie-kotów. Jeśli użyłam słowa „kot” niezgodnie z jego znaczeniem, to użyłam go niepoprawnie. Jeśli przez „kot” mam na myśli kota, to w odpowiedzi na

pytanie o to, jakie zwierzę wyleguje się na moim fotelu (o ile jest nim kot), *powinnam* odpowiedzieć: „Kot”.

„Związek między znaczeniem (*meaning*) i zamiarem a przyszłymi działaniami jest *normatywny*, a nie *opisowy*” — twierdzi Kripke (2005: 66) i uważa, że teoria znaczenia, która nie jest w stanie uwzględnić tego wymogu normatywności, jest z góry skazana na porażkę. Jak łatwo spostrzec, głównym obiektem ataku Kripkego, stały się tak zwane teorie „czystego użycia” (*pure use theories*) (por. Wikforss 2001: 204). Niezależnie od zasadności zarzutu Kripkego i tego, czy głównym problemem teorii „czystego użycia” jest rzeczywiście normatywność (a nie raczej to, czy pozwalają one wyjaśnić błąd językowy), wielu filozofów potraktowało argument z normatywności jako intuicyjnie przekonujący test wiarygodności teorii znaczenia. W głośnym artykule z 1989 roku Paul Boghossian pisał:

Fakt, że określone wyrażenie coś znaczy, implikuje, że istnieje cały zbiór prawd normatywnych dotyczących mojego zachowania względem tego wyrażenia (Boghossian 1989: 513).

I mimo że Boghossian w późniejszych artykułach wycofał się z poparcia dla semantycznego normatywizmu, intuicje dotyczące tego, że istnieje silny związek między znaczeniem wyrażenia a tym, jak powinno się go używać, podziela wielu uczestników tego sporu¹. Kluczowym problemem pozostaje *charakter* tego związku; czy znaczenie samo w sobie generuje bądź implikuje normatywne konsekwencje dotyczące jego użycia, czy raczej znaczenie *nabiera* normatywnego charakteru, o ile dodane zostaną dodatkowe założenia dotyczące np. intencji komunikacyjnej mówiącego, tego, co leży w jego interesie, lub jaki związek ma użycie słowa w zgodzie z jego warunkami poprawności z osiągnięciem przez niego jakichś pozajęzykowych celów. To właśnie kwestia tego, czy *samo* znaczenie jest normatywne, a więc jest wewnętrznie normatywne, czy raczej jest normatywne zewnętrznie, z uwagi na cele lub intencje mówiącego, jest od lat 90. punktem spornym. Przyjrzymy się teraz głównym argumentom formułowanym w tej debacie.

3. Normatywność a warunki poprawnego użycia

Teza SN wydaje się jasna. Jego zwolennicy przekonują, że *fakt*, iż słowo „kot” w języku polskim ma określone warunki poprawności, *sam w sobie* jest *racją* do tego, żeby używać go *poprawnie*, czyli zgodnie z jego znaczeniem. Słowa „fakt”, „samo w sobie”, „racja” oraz „poprawnie” celowo zostały napisane kursywą, ponieważ to, jak się je interpretuje, ma istotny wpływ na siłę argumentów

¹ Do grona semantycznych normatywistów należą m.in. S. Blackburn (1984); D. Bloor (1997); P. Boghossian (1989); R. Brandom (1994); A. Gibbard (2003, MS); H.-J. Glock (1996); J. McDowell (1994); P. Pettit (1990); D. Whiting (2007, 2009a, b); C. Wright (1984).

formułowanych przez poszczególnych autorów. Zaczniemy więc od tego, jakie znaczenie przypisują tym słowom uczestnicy debaty.

Znalezienie odpowiedzi na to pytanie nie jest trudne w przypadku semantycznych antynormatywistów², ponieważ wprost wskazują oni, przeciwko jakiej tezie i jakiemu jej rozumieniu występują. W ich ujęciu stanowisko SN głosi, że fakt, iż słowo ma określone znaczenie, ma tę konsekwencję, że pewne jego użycia są poprawne, a inne nie. Inaczej rzecz ujmując, w ich przekonaniu, to, że wyrażenia posiadające znaczenie posiadają warunki poprawnego użycia, nie mówi nam jeszcze nic na temat tego, *jak* powinniśmy używać owych wyrażzeń. Weźmy słowo „kot”. W języku polskim słowo „kot” ma określoną denotację, odnosi się do określonego typu zwierząt, które posiadają genom charakterystyczny dla kotów. Powiedzieć, że słowo „kot” posiada warunki poprawnego użycia, to w związku z tym stwierdzić tylko, że słowo to *poprawnie* jest zastosować w odniesieniu do kotów, natomiast niepoprawnie w odniesieniu do zwierząt, które nie posiadają cech wymienionych w definicji słowa „kot” lub — prościej rzecz ujmując — nie są kotami.

Twierdzenie, że wyrażenia posiadające znaczenie mają warunki poprawnego użycia, można w sposób bardziej formalny, przedstawić następująco:

(P) s znaczy $F \rightarrow \forall x$ (s stosuje się poprawnie względem $x \leftrightarrow x$ jest f)³

Jeśli zgodzimy się, że powyższa formuła poprawnie wyraża warunek poprawnego użycia, to, z tego, że prawdą jest, iż słowo „kot” ma poprawne zastosowanie wyłącznie w odniesieniu do kotów, jeszcze nic *normatywnego* nie wynika. Powiedzieć, że słowo „kot” poprawnie jest używać w odniesieniu do kotów to tyle, co odnotować pewien *fakt* dotyczący znaczenia. Fakt ten sam w sobie pozbawiony jest normatywnej siły, ponieważ warunki poprawnego użycia słowa kot nie *instruuje* nas, co powinniśmy zrobić. Norma semantyczna jest bezosobowa i jako taka sama w sobie nie daje *nikomu* żadnej racji do działania⁴. Jeśli (P) coś implikuje, to tylko tyle, że mamy tu do czynienia z *kategoryzacją* czy *uporządkowaniem* użyć danego słowa „w dwa podstawowe semantyczne rodzaje; na przykład: prawdziwe i nieprawdziwe” (Glüer, Wikforss 2008: 2, 2009a: 7). Jeśli znaczenie słowa mówi nam tylko, do jakich obiektów poprawnie jest je stosować, to jakim sposobem można bezpośrednio wnosić z tego, że poprawne, tj. zgodne z ekstensją użycie wyrażenia jest tym, które *powinno się* zastosować lub do zastosowania którego *ma się rację*? (Glüer, Wikforss 2008: 2). To, co powinno się lub czego nie powinno się zrobić z danym wyrażeniem,

² Wśród semantycznych antynormatywistów są m.in. A. Bilgrami (1993); D. Davidson (1984); F. Dretske (2000); K. Glüer (1999, 2001); K. Glüer i P. Pagin (1999); K. Glüer i Å. Wikforss (2008, 2009a, b, 2010a, b); A. Hattiangadi (2006, 2007, 2009); A. Miller (2010a, b); D. Papineau (1999).

³ Por. Glüer, Wikforss (2009a: 35); Hattiangadi (2009: 55); Whiting (2009a: 537).

⁴ Na temat tego, kiedy norma jest normatywna, zob. von Wright (1963); Humberstone (1971).

albo to, czy ma się rację, żeby użyć go poprawnie lub niepoprawnie, nie jest kwestią *semantyczną*, lecz *pragmatyczną* lub *moralną*. A zatem *jeśli* chcę, żebyś mnie zrozumiał, gdy wypowiadam zdanie: „Kot leży na mojej sofie” jako mówiące coś o *kocie*, wówczas mam *rację*, utrzymują semantyczni antynormatywiści, żeby użyć słowa „kot” zgodnie z jego warunkami poprawności. W takiej sytuacji racją do użycia słowa „kot” poprawnie nie jest dla mnie to, że słowo „kot” *oznacza* kota, lecz to, że chcę wywołać odpowiedni skutek swoją wypowiedzią, a mianowicie poinformować cię, że na mojej sofie leży kot. Jednak równie dobrze moja motywacja do tego, żeby użyć „kot” w odniesieniu do kota, może mieć charakter nie pragmatyczny, lecz moralny. Oto jestem przekonana, że należy mówić prawdę i jeśli to *kot* leży na mojej sofie, to mam *obowiązek* powiedzieć: „Kot leży na sofie”, a nie na przykład: „Łasica leży na sofie”. Mimo że zgadzają się z tezą, iż względy moralne mogą odgrywać istotną rolę w tym, jak użytkownicy danego języka posługują się wyrażeniami tego języka (zgodnie z ich znaczeniem), semantyczni antynormatywiści zaprzeczają, że *moralną* powinność należy utożsamiać z powinnością *semantyczną*. W ich ocenie, choć nie dają temu wyrazu *expressis verbis*, samo wyrażenie „powinność semantyczna” brzmi co najmniej dziwacznie, ponieważ sugeruje, że istnieją obowiązki natury semantycznej, których niewypełnienie prowadzi do krytyki lub ją przynajmniej uzasadnia. To z kolei sugeruje, że język, na który składają się wyrażenia posiadające znaczenie, jest wartością samą w sobie tak jak moralność. *Nieoszukiwanie* może być wartością samą w sobie, ale czy wartością samą w sobie jest *używanie słowa „kot” wyłącznie w odniesieniu do kota*? Innymi słowy, według semantycznych antynormatywistów, kiedy wskazując na kota, mówię „łasica”, o moim zachowaniu językowym można powiedzieć jedynie, że złe, czyli niepoprawnie zastosowałam słowo „kot”. W żadnym razie nie implikuje to twierdzenia, że zrobiłam coś, czego *nie powinnam* była zrobić. Być może jesteś krótkowzrocznym mordercą kotów i przy okazji zagorzałym fanem łasic, a celowe wprowadzenie cię w błąd, które od strony semantycznej należy zakwalifikować jako niepoprawne, od strony moralnej jest jak najbardziej słuszne i uzasadnione. W tym miejscu dochodzimy do istotnej różnicy w interpretacji głównych pojęć, którymi posługują się uczestnicy dyskusji. Kiedy semantyczny antynormatywista argumentuje przeciwko tezie głoszącej, że znaczenie jest „wewnętrznie”, „samo” lub „z istoty” *normatywne*, to obiektem jego krytyki jest następująca teza: „samo znaczenie rozumiane jako wyrażające warunki poprawnego użycia jest preskryptywne lub kieruje działaniem”. Termin „preskryptywny” antynormatywiści interpretują dość ściśle, jako pochodzący od preskrypcji, a preskrypcja jest typem normy, która składa się z sześciu „elementów” (por. von Wright 1963): 1) *charakteru* (czy norma mówi, że coś *powinno się, można* lub *musi się* zrobić lub nie zrobić); 2) *treści* (co norma nakazuje, zakazuje, lub na co zezwala); 3) *warunków zastosowania* (czy norma jest

kategoryczna, a zatem warunki jej zastosowania są ujęte w treści normy, czy też hipotetyczna, tzn. warunków jej zastosowania nie można wyprowadzić z samej normy); 4) *autorytetu* (kto jest autorem preskrypcji); 5) *podmiot(ów)* (do kogo preskrypcja jest adresowana), wreszcie 6) *okoliczności* (najczęściej określenia miejsca lub czasu, w którym norma ma zastosowanie, np. „dzisiejszego popołudnia”, „raz w roku”, „zawsze”). Poza tymi sześcioma „komponentami” preskrypcji von Wright wyróżnił dwa dodatkowe elementy preskrypcji: *promulgację* i *sankcję*. Jeśli charakterystykę preskrypcji przedstawioną przez von Wrighta uznamy za poprawną, wówczas, jak sugerują semantyczni antynormatywiści, o tym, czy jakaś norma jest czy też nie jest preskryptywna, decyduje to, czy spełnia ona powyższe warunki czy też nie. Ponadto, jeśli uwzględnimy argument Kripkego z normatywności, okaże się, że norma, która jest preskryptywna, to taka, która *kieruje* działaniem (*action-guiding*), a zatem w sensie dosłownym dostarcza *nam*, działającym *konkretnych* wskazówek, *co* i *kiedy* powinniśmy zrobić (Glüer, Wikforss 2010b: 758). Preskrypcja w rodzaju „Myj zęby po każdym posiłku” jest czytelna i sama w sobie może kierować naszym działaniem, ponieważ informuje nas, co i kiedy powinniśmy lub mamy silną rację zrobić, natomiast norma semantyczna wyrażająca warunki poprawnego użycia *sama w sobie* preskrypcją w ogóle nie jest.

Przy założeniu, że preskrypcje są normatywne w sensie dostarczania racji do działania⁵, dostarczają one również istotnych wskazówek dotyczących działania, ponieważ mówią wyraźnie, do kogo i w jakich okolicznościach preskrypcja się stosuje. Zatem *realnie* kierują działaniem (Glüer, Wikforss 2009a: 32–33).

Opór semantycznych antynormatywistów przeciwko tezie, że znaczenia, czyli warunki poprawnego użycia wyrażenia *same* są *racjami* do działania bądź je implikują, wydaje się zrozumiały. Nie twierdzą oni bowiem, że znaczenia słów nie *mogą* stanowić racji do ich używania w określony sposób, a tylko, że sam *fakt* dotyczący tego, jakie obiekty denotuje słowo „kot”, nie ma jeszcze siły normatywnej. Jeśli (P) wyraża jedynie warunki poprawnego użycia *s*, a zatem poprawna interpretacja (P) to taka, która mówi, jakie użycia *s* są poprawne, to bezpośrednią implikacją (P) jest co najwyżej kategoryzacja na poprawne i nie-

⁵ Jest to kwestia nader kontrowersyjna i będąca wciąż przedmiotem szerszego sporu o naturę normatywności. Filozofowie tacy jak Derek Parfit uważają, że preskrypcje w sensie von Wrighta nie są normatywne, ponieważ normatywność ma silny związek z racjami, leżącymi u podstaw danej normy, a nie z jej formą logiczną. Prościej rzecz ujmując, Parfit trafnie zauważa, że w przypadku wyraźnych preskrypcji w rodzaju „Przynieś mi książkę, która leży w sali konferencyjnej”, moc nakazująca przyniesienie książki nie ma nic wspólnego z racjami przemawiającymi za wykonaniem tego, co jest treścią preskrypcji. „Normatywny” charakter preskrypcji, zdaniem Parfita, jest w tym przypadku pozorny, ponieważ w preskrypcji o charakterze nakazu nie o racje chodzi, lecz o zrealizowanie nakazu. Nakaz nakazuje, ale nie dostarcza racji poza racją, która wynika z faktu, że nakaz jest fortunny; „racja” stojąca za nakazem ma, by tak rzec, charakter definicyjny, a niekoniecznie materialny. Jeśli jestem twoim zwierzchnikiem, fakt, że mówię ci, co masz zrobić (przynieść mi książkę ze stołu w sali konferencyjnej) *wystarcza*, żebyś zrobił to, co ci nakazuję, niezależnie od tego, czy istnieje jakiś dobry powód, by zaprzętać sobie głowę ową książką (por. Boghossian 2008: 475).

poprawne użycia *s*. Kategoryzacja sama w sobie, semantyczna czy pozasemantyczna, nie wywołuje bezpośrednich skutków normatywnych. Możemy podzielić przedmioty w moim pokoju na stoły i nie-stoły, ale z samego tego podziału nic jeszcze nie wynika odnośnie do tego, co ty, ja lub ktokolwiek inny *powinien* bądź *ma rację* zrobić.

Stwierdzenie, że sama kategoryzacja nie niesie bezpośrednich normatywnych konsekwencji, nie jest jednak tym samym, co stwierdzenie, że daną kategoryzacją nie można się *posłużyć* do wyciągnięcia normatywnych konsekwencji. W istocie truizmem wydaje się stwierdzenie, że *wszystko* — każdy fakt (rozumiany jako prawdziwy sąd) — *może* w określonym kontekście dla określonej osoby *stać się racją* do określonego zachowania, o ile dodane zostaną dodatkowe warunki. *Fakt*, że „kot” w języku polskim oznacza kota, z punktu widzenia użytkownika języka polskiego jest potencjalnie „racjogenny”, nawet silnie „racjogenny”, tyle że nie *sam w sobie*, lecz po uwzględnieniu psychologii użytkowników języka polskiego lub konkretnego użytkownika języka polskiego. Semantyczni antynormatywiści zwracają również uwagę na to, że jeśli zwolennicy semantycznego normatywizmu mają rację i znaczenie, czyli warunki poprawnego użycia konkretnego wyrażenia, *samo* generuje racje do działania, to na ich oponentach spoczywa ciężar dowiedzenia, że poprawność *semantyczna* jest preskryptywna.

Wyzwanie rzucone zwolennikom normatywności znaczenia jest jak najbardziej uczciwe, ponieważ to oni bronią tezy, że warunki poprawnego użycia danego wyrażenia same rodzą racje do jego używania (nieużywania) w określony sposób. Jeśli sam fakt, że „kot” oznacza kota *jest* dla mnie racją odnośnie do tego, jak powinnam używać tego słowa, to pytanie, *jak to jest możliwe*, jest jak najbardziej na miejscu, zważywszy że (P) nie dostarcza informacji o tym, kto *powinien* lub *ma silną* rację używać słowa zgodnie z jego ekstensją.

4. Normatywność pojęcia „poprawny”

Jak na to wyzwanie odpowiada semantyczny normatywista?⁶

Po pierwsze, zaczyna od przypomnienia, że obie strony zgadzają się co do tego, iż fakt, że dane wyrażenie ma określone warunki poprawności, *implikuje*, że pewne użycia tego wyrażenia są *poprawne*, a inne *niepoprawne* (Whiting 2009a: 537). Następnie stara się przekonać, że *czymś oczywistym, intuicyjnie*

⁶ W dalszej części tego artykułu będę odwoływać się wyłącznie do argumentów przedstawionych przez Daniela Whitinga, ponieważ inni zwolennicy semantycznego normatywizmu, wymienieni w przypisie 5, nie odwołują się bezpośrednio do zarzutów sformułowanych przez semantycznych antynormatywistów. Zamiast polemiki podejmują próbę autorskiego przedstawienia pozytywnego ujęcia normatywności znaczenia. Por. zwł. Brandom (1994); Gibbard (MS). W jakim stopniu ich propozycje teoretyczne stanowią odpowiedź na zarzuty antynormatywistów i czy w ogóle są próbą obrony tezy SN w interesującym nas sensie jest osobną i mocno dyskusyjną kwestią.

wiarygodnym jest uznać, że słowo „poprawny” dopuszcza jedną prawidłową interpretację, *normatywną* lub *preskryptywną* (oba terminy w argumentacji głównego orędownika semantycznego normatywizmu są używane zamiennie i sugerują relację między sprawcą a działaniem: racja jest racją dla kogoś do zrobienia czegoś)⁷. Zdaniem Whitinga słowo „poprawny” jest pojęciem normatywnym *sui generis* tak jak paradygmatycznie normatywny termin „powinien”. Prostą konsekwencją takiego poglądu jest to, że ilekroć w zdaniu, które jest prawdziwe, pojawia się termin *sui generis* normatywny, niech to będzie „powinien”, „poprawnie jest”, „słusznie jest” itp., ma ono bezpośrednie normatywne następstwa dotyczące działania. Zgodnie z powyższym, jeśli prawdziwe jest zdanie głoszące, że „kot” w języku polskim oznacza kota, co jest innym sposobem powiedzenia, że *poprawnie jest stosować* słowo „kot” w odniesieniu do kota (w języku polskim), wówczas każdy (domyślnie: każdy użytkownik języka polskiego) *ma* rację, żeby używać słowa „kot” w odniesieniu do kota, a nie na przykład łasicy. Co więcej, pisze Whiting, skoro semantyczni antynormatywiści zachowują ostrożność co do tego, czy „poprawny” jest czy też nie jest terminem *sui generis* normatywnym, dopuszczając, iż „poprawny” można zinterpretować na dwa sposoby: normatywnie i nienormatywnie (Glüer, Wikforss 2008: 2, 2009a: 36–37; Hattiangadi 2009: 59–61), to sami powinni najpierw wyjaśnić, dlaczego są przekonani, że w kontekstach moralnych słowu „poprawny” odpowiada interpretacja raczej normatywna, a w semantycznym nienormatywna. Czy w tej sprawie dysponują jakąkolwiek racją poza własną intuicją? Jeśli semantyczni antynormatywiści, kontynuując swój wywód Whiting, uważają, że „poprawny” jest słowem wieloznacznym, to podatni są na błąd ekwiwokacji, widoczny chociażby w poniższym rozumowaniu (Whiting 2009a: 538):

- (1) Zofia zachowała się poprawnie, kiedy oddała właścicielowi portfel, który znalazła na ulicy.
- (2) Zofia w sposób poprawny użyła słowa „czerwony” w odniesieniu do czerwonego przedmiotu.
- (3) Zatem Zofia zachowała się poprawnie dwa razy.

Zarzut Whitinga pod adresem antynormatywisty brzmi następująco: jeśli w zdaniu (1) słowu „poprawnie” nadamy (zdroworozsądkową) interpretację normatywną, a w zdaniu (2) rzekomo oczywistą nienormatywną, to musimy uznać wniosek wyrażony w zdaniu (3). Whiting przyznaje, że argument, nazwijmy go przyczynkowym (ponieważ w żadnym razie nie rozstrzyga, czy „poprawny” jest terminem jednoznacznie normatywnym czy też nie jest) argumentem z *zagrożenia ekwiwokacji*, nie ma charakteru decydującego, niemniej

⁷ Takie rozumienie racji we współczesnej debacie poświęconej normatywności określa się jako *deliberatywne*, ponieważ jest silnie związane z perspektywą pierwszoosobową; por. Schroeder (2011: 141); Finlay (2010a); Wedgwood (2007).

pozwala przerzucić ciężar dowodu na antynormatywistów, których pogląd na wieloznaczność terminu „poprawny” ma wyraźnie kłócić się z naszymi intuicjami.

Ponieważ Whiting sprowadza dyskusję do poziomu intuicji, sugerując przez to, że antynormatywiście, jeśli chce dowieść słuszności swojego stanowiska, nie pozostaje nic innego, jak zaoferować nam jakiś intuicyjnie rozstrzygający przykład, warto zadać pytanie, czy rzeczywiście antynormatywista ma jakikolwiek powód, by szukać wsparcia w przywoływanych *ad hoc* przykładach, jak tego oczekuje od niego Whiting. Sądzę, że nie ma, ponieważ po pierwszej rundzie (nazwijmy ją rundą „normatywność poprawności”) i tak zachowuje niewielką przewagę. A jest tak z dwóch dość łatwo zauważalnych powodów: po pierwsze, rozumowanie, które Whiting uważa za jednoznacznie przemawiające za normatywnością terminu „poprawnie”, wcale nie *pokazuje* tego, co ma rzekomo pokazać, czyli że to rzekoma normatywność słowa „poprawnie” w zdaniu (1) przemawia na rzecz normatywnej interpretacji tego zdania, a nie reszta zdania, kontekst, w którym słowo „poprawnie” występuje, jak również kontekst, w którym zdanie (1) zostaje wypowiedziane, skłania nas do przyjęcia interpretacji normatywnej jako właściwej. Po drugie, nie jest wcale oczywiste, dlaczego *właściwą* i oczywistą samą przez się interpretacją słowa „poprawnie” w zdaniu (1) miałyby być akurat interpretacja preskryptywna. Równie poprawną i dopuszczalną interpretacją zdania (1) jest taka, w której *to, że Zofia zachowała się poprawnie, zwracając portfel właścicielowi, znaczy tylko tyle, że Zofia postąpiła zgodnie z normą (moralną) nakazującą zwracać znalezione rzeczy prawowitemu właścicielowi*. W tym wypadku mamy do czynienia z opisem zachowania Zofii z punktu widzenia konkretnej normy moralnej. Wydaje się, że Whiting uważa niepreskryptywną interpretację słowa „poprawnie” w zdaniu (1) za nienaturalną, ponieważ sądzi, że słowo „poprawnie” zachowuje się w tym zdaniu jak niekontrowersyjne pojęcia, takie jak „słusznie”, „właściwie”, które z istoty swojej mają charakter normatywny i informują nas, co mamy *rację* zrobić lub nie zrobić. Kłopot z sugestią Whitinga polega na tym, że nawet przy założeniu, że normatywne odczytanie słowa „poprawnie” w zdaniu (1) bardziej zgadza się z intuicjami niektórych osób i „poprawnie” jest normatywne w sensie, w jakim normatywne są inne wymienione wyżej pojęcia, to aby argument przeciwko antynormatywiście był skuteczny, musi on wykazać, że wspomniane terminy mają *jedno* niekontrowersyjne i normatywne (preskryptywne) znaczenie. Jest to ważne z tego względu, że gdyby się okazało, iż paradygmatycznie normatywne pojęcia, takie jak „powinien”, „słusznie” itd., wcale nie są jednoznacznie normatywne, semantyczny normatywista znalazły się w bardzo trudnej sytuacji. W takim przypadku musiałby dysponować wyjątkowo silnym argumentem za jednoznacznością normatywności terminu, który jest pod tym względem całkowicie *wyjątkowy*, ponieważ nie ma innego terminu, który byłby również bezwyjątkowo normatywny. Innymi słowy, zawsze ilekroć

„słusznie”, „niesłusznie” lub inne paradygmatycznie normatywne pojęcia, takie jak „powinien”, wystąpią w zdaniu, jedyną prawidłową interpretacją tych terminów byłaby interpretacja preskryptywna. Tymczasem jest to kwestia wysoce kontrowersyjna⁸ i równie mocno sprzeczna z naszymi intuicjami. Weźmy zdanie „Poprawnie jest wypełnić formularz urzędowy, uzupełniając wszystkie puste miejsca wymaganymi i prawdziwymi informacjami”. Dlaczego „poprawnie” w zacytowanym zdaniu miałoby implikować cokolwiek na temat tego, co powinnam zrobić? Co ważniejsze, gdyby Whiting miał rację i takim pojęciem, jak „poprawny”, przysługiwałaby jednoznacznie preskryptywna interpretacja, wówczas zadanie wykazania *niemożliwości* nienormatywnego odczytania tak zwanych paradygmatycznie normatywnych terminów nie nastęrczałoby większych problemów. Tymczasem nie sposób dowieść, że słowu „poprawny” przysługuje tylko jedna, normatywna interpretacja. Jak można bowiem dowieść niemożliwości czegoś, co ma charakter empiryczny? Zwróćmy uwagę, że nawet jeśli wykluczmy warunek dowiedzenia niemożliwości niepreskryptywnej interpretacji słowa „poprawnie” jako zbyt silny i skoncentrujemy się na słabszym, głoszącym, że słowo „poprawnie” w *większości standardowych* użyciach zachowuje się preskryptywnie, nie posunie nas to wcale naprzód z tego względu, że nie bardzo wiemy, które sposoby jego użycia uznać za standardowe, oraz dlaczego ewentualne standardowe użycia preskryptywne w jednej dziedzinie (przykładowo w moralnej) miałyby automatycznie sprawiać, że pozamoralne użycia słowa „poprawnie”, na przykład semantyczne, również są preskryptywne. Być może, jak sugerują antynormatywiści, jest tak, że *moralna* poprawność ma wymiar wewnętrznie preskryptywny, ale dlaczego mielibyśmy przyjąć, że tak samo dzieje się w przypadku poprawności *semantycznej*?

5. Normatywność i standard

Jeszcze innym sposobem na udzielenie krytycznej odpowiedzi na sugestię Whitinga dotyczącą normatywności pojęć i takich wyrażań, jak „poprawny” czy „jest poprawnie”, jest zwrócenie uwagi na to, jak *działa* standard albo norma. Odwołanie się do standardu czy normy wydaje się uzasadnione, ponieważ pojęcia i wyrażenia normatywne często stanowią ich element. Na przykład norma etykietalna poucza nas, że nie wolno trzymać łokci na stole, gdy spożywa się obiad w towarzystwie, a norma gramatyczna języka polskiego instruuje nas, że poprawnie jest tworzyć nazwy osobowe żeńskie od męskich przez dodanie przyrostka *-ka* (choć oczywiście są tu wyjątki). Czy dlatego, że istnieją pewne normy lub standardy sformułowane przy użyciu rzekomych jednoznacznie

⁸ Obecnie trwa debata, czy pojęcie „powinien” ma jedno ewaluatywne znaczenie (powinno być tak, że...), z którego dopiero można wywieść deliberatywną powinność (X powinien zrobić coś), czy też jest pojęciem wieloznacznym, któremu w niektórych kontekstach przysługuje interpretacja ewaluatywna (w mojej terminologii: nienormatywna), z kolei w innych deliberatywna (preskryptywna); zob. przypis 7.

normatywnych pojęć, uzasadnione jest stwierdzenie, że owe normy albo standardy *automatycznie* generują normatywne konsekwencje na temat dopuszczalnych zachowań?

Anandi Hattiangadi, należąca do grona antynormatywistów, uważa, że normatywność (preskryptywność) norm nie zawsze jest oczywista.

Niekiedy — pisze Hattiangadi — powiedzieć, że coś jest słuszne (lub poprawne) *nie* implikuje preskrypcji; jest raczej stwierdzeniem, że spełnia się określony *standard* [...]. Powiedzieć, że jakieś użycie jest „poprawne”, to zaledwie opisać je w określony sposób — w świetle normy lub standardu ustanowionego przez znaczenie owego terminu (Hattiangadi 2006: 223–225; por. Kusch 2006: 60; Wikforss 2001: 205 i n.).

Żeby zilustrować nienormatywny charakter *standardu*, Hattiangadi podaje przykład ze swoim synem Vikramem. Wyobraźmy sobie, że istnieje oficjalny standard mówiący o tym, jaki wzrost muszą mieć dzieci, aby mogły bawić się na określonym placu zabaw. Załóżmy, że wymóg ten jako dolną granicę wzrostu ustanawia jeden metr wysokości oraz że Vikram (szczęśliwie) spełnia ten warunek (ma więcej niż jeden metr wzrostu). Czy to, że Vikram spełnia ów standard *samo w sobie* rodzi jakiegokolwiek normatywne konsekwencje? Nie — to, że Vikram spełnia lub nie spełnia określonego wymogu zawartego w standardzie jest kwestią nienormatywną, lecz jednym z faktów na temat Vikrama. Standard ów nabrałby mocy normatywnej, mógłby być racją do działania, o ile Vikram chciałby pobawić się na owym placu zabaw. Jeśli Vikram nie wykazywałby zainteresowania zabawą na tym placu, standard mówiący, jakim dzieciom wolno korzystać z placu zabaw, pozostałby wyłącznie opisem warunków koniecznych do tego, aby móc się na nim bawić. Podobnie rzecz wygląda w przypadku standardu poprawnego użycia danego wyrażenia. To, że istnieje norma semantyczna informująca, jaka jest ekstensja słowa „kot”, należy rozumieć jako stwierdzenie na temat tego, jakie jest poprawne użycie słowa „kot” w języku polskim. Z samej normy stwierdzającej poprawność jakiegoś zachowania nie sposób jeszcze wyprowadzić *racji do działania*⁹. Whiting przyznaje, że przykład podany przez Hattiangadi pokazuje, że *sam* standard nie jest jeszcze normatywny, ale gdybyśmy dołożyli warunek, że ów standard *obowiązuje (in force)*, wówczas niewątpliwie rodziłby skutki normatywne, z których najważniejszym byłaby uzasadniona krytyka w przypadku naruszenia owego standardu.

⁹ W tym kontekście przydatne może być zaproponowane przez Parfita rozróżnienie „faktów i własności niosących normatywne znaczenie” (*facts and properties with normative importance*) oraz „faktów i własności o normatywnym znaczeniu” (*facts and properties about normative importance*). To, że „kot” oznacza kota, jest przypadkiem faktu niosącego normatywne znaczenie, ale nie o normatywnym znaczeniu, ponieważ fakt ten sam w sobie nie jest racją dla nikogo, żeby użyć słowa „kot” w odniesieniu do kota. Przykładem faktu o normatywnym znaczeniu byłoby to, że to, iż „kot” oznacza kota, jest racją dla Marka, żeby powiedzieć „kot”, a nie „kuna” w sytuacji S; por. Parfit (2011: 279). Zob. też Finlay (2010b: 334).

Kontrargument Whitinga, nazwijmy go „argumentem z obowiązywania”, wydaje się chybiony z punktu widzenia celu, któremu ma się przysłużyć, a mianowicie pokazania, że *samo* znaczenie jest preskryptywne. Po pierwsze, to, że standard *obowiązuje* (przez co rozumieć będziemy to, że został nałożony przez uprawniony do tego podmiot, w tym przypadku zarządcę placu zabaw), samo w sobie nie *jest racją* do działania, jest nią dopiero wówczas, gdy jakieś dziecko chce bawić się na tym placu. Niemniej standard, który jest *ważny* (tzn. spełnia warunki normy instytucjonalnej rodzącej określone konsekwencje prawne), jest *racją instytucjonalną* do tego, żeby go przestrzegać, a zatem rodzi *pewne* konsekwencje normatywne. Jak należy w tym przypadku rozumieć „konsekwencje normatywne”? Otóż wydaje się, że jako *opis* następstw związanych z naruszeniem określonej normy. Jeśli standard instytucjonalny mówi, co jest dozwolone, to zyskujemy wiedzę o tym, jakie zachowanie jest niedozwolone (ewentualnie jakie sankcje grożą w razie jego naruszenia), a nie wskazówkę praktyczną dotyczącą tego, co *powinniśmy* (nie *powinnismy*) zrobić. To, co powinnam zrobić albo co mam rację zrobić, nie zależy wyłącznie od tego, czy istnieje jakiś standard, który stosuje się do mojej sytuacji, lecz od nienormatywnych cech sytuacji, w której się aktualnie znajduję. Przypuśćmy, że standard mówiący o tym, jaki wzrost muszą mieć dzieci bawiące się na placu zabaw, *obowiązuje*, a zatem jeśli Maja nie spełnia tego standardu, to *nie powinna* wchodzić na plac zabaw. Czy fakt, że ów standard obowiązuje jest *bezwzględną racją* do tego, żeby Maja nie postawiła stopy na owym placu? Wydaje się, że nie; gdyby bowiem okazało się, że Maja, która nie ma jeszcze 1 metra wzrostu, *dzięki temu faktowi*, może prześlizgnąć się i, na przykład, pomóc bratu otworzyć furtkę, która się zatrzasnęła, powiedzielibyśmy, że w takiej sytuacji Maja ma rację (lub że istnieje racja, żeby Maja...), pomimo tego, co standard głosi, wejść na plac zabaw.

Niepreskryptywny charakter standardu instytucjonalnego ilustruje przykład sformułowany przez Alexandra Millera. Wyobraźmy sobie, że w pewnej wspólnocie obowiązuje standard zezwalający dzieciom na korzystanie z placu zabaw, o ile spełniają one następujący warunek: w przynajmniej jeden wtorek w ciągu ubiegłego roku zjadły one na śniadanie płatki (zob. Miller 2010a: 9). Czy uznamy, że standard ten jest racją do tego, żeby go spełniać? Odpowiedź Millera jest negatywna. Zanim uznamy, że istnieją racje do tego, by trzymać się danej normy, musimy wpięrow ustalić, czy racje przemawiające za ową normą, należą do *właściwego* rodzaju racji, a zatem dostarczają powodu, by zaakceptować ów standard i go wypełniać. Jeśli, tak jak w tym przykładzie, warunek nałożony na dzieci, które mogą wejść na plac zabaw, nie ma nic wspólnego z *rozsądnymi* racjami przemawiającymi za dopuszczeniem dzieci na plac zabaw (do rozsądnych racji zaliczmy wzrost, wagę lub inne istotne ze względów bezpieczeństwa cechy), norma ta, *mimo że obowiązuje*, nie implikuje żadnych racji.

Zwróćmy ponadto uwagę, że sugestia Whitinga, iż *obowiązywanie* standardu „uwalnia” jego preskryptywną moc (w przypadku norm instytucjonalnych), nawet jeśli jest trafna, nie wydaje się szczególnie pomocna jako argument na rzecz normatywności normy semantycznej. Jest tak dlatego, że nie wiemy, na czym miałyby polegać analogia między obowiązywaniem standardu instytucjonalnego, głoszącego, jakie dzieci mogą korzystać z placu zabaw (A), a obowiązywaniem standardu semantycznego, mówiącego, jakie użycie danego wyrażenia jest poprawne w świetle jakiejś normy (B). O ile nie mamy problemu ze zrozumieniem, co to znaczy, że standard w sensie (A) *obowiązuje*, tzn. wiemy, kto jest władny decydować o zasadach dostępu do placu zabaw, o tyle mamy kłopot z interpretacją tego, co to znaczy, że norma semantyczna *obowiązuje*. W jakim sensie to, że „kot” odnosi się do kota, *obowiązuje* i *zobowiązuje* kogokolwiek do czegoś? Z pewnością nie w takim samym, w jakim obowiązuje norma zezwalająca niektórym dzieciom bawić się na placu zabaw. Nie ma bowiem autorytetu instytucjonalnego, który „narzuca” znaczenia wyrażen i ustala sankcje za ich niewłaściwe stosowanie. Zatem nawet jeśli zgodzimy się, że dopiero „skorygowanie” przykładu Hattiangadi zaproponowane przez Whitinga, które polega na dodaniu do normy mówiącej, jakie dzieci mogą bawić się na placu zabaw, warunku *bycia obowiązującą*, sprawia, że norma ta rodzi normatywne konsekwencje lub nabiera *ważności*, to wciąż nic istotnego nam to nie mówi o tym, czy norma semantyczna jest wewnętrznie normatywna czy nie. Brakuje nam bowiem informacji na temat tego, jak rozumieć *obowiązywanie* normy wyrażającej warunki poprawnego użycia danego wyrażenia.

Jeśli zatem Whiting używa słowa „obowiązuje” (*in force*) w odniesieniu do normy semantycznej, to musi mieć na myśli coś innego niż jej *legalny* charakter. Spróbujmy teraz zastanowić się, o jak rozumiane *obowiązywanie* tu chodzi, oraz czy faktycznie w istotny sposób przyczynia się ono do wzmocnienia stanowiska semantycznego normatywisty?

W tym celu rozważmy następującą formułę, wyrażającą stanowisko SN:

(P') *s* znaczy $F \rightarrow \forall x (S \text{ nie powinien (stosować } s \text{ w odniesieniu do } x) \leftrightarrow x \text{ nie jest } f)$

Formuła (P') jest logicznie równoważna formule (P''):

(P'') *s* znaczy $F \rightarrow \forall x (S \text{ może (stosować } s \text{ w odniesieniu do } x) \leftrightarrow x \text{ jest } f)$

Whiting uważa, że *logiczna* poprawność owych formuł ma wymiar preskryptywny (por. Whiting 2009a: 544). A zatem, skoro (P') oraz (P'') poprawnie wyrażają warunki poprawnego użycia słowa *s* ((P'') jest prostym przekształceniem (P')), to *same* te formuły implikują racje dotyczące tego, jak *można* (nie

można) używać s. Whiting zaznacza, że nie chodzi mu o to, że (P') czy (P'') informują, co *można* (w sensie zbioru dostępnych dla konkretnej jednostki opcji działań) zrobić, bo można zrobić wiele rzeczy, na przykład założyć nogę na nogę, kiedy się siedzi, którego to zachowania w zwykłych okolicznościach raczej nie ujęłoby się w kategoriach poprawności, a jedynie, że *stwierdzenie (sąd) na temat niepoprawności* jakiegoś działania implikuje bez żadnych dodatkowych założeń, że *ktos nie powinien* go zrealizować (Whiting 2009a: 545, podkr. J. K.).

Nazwijmy formuły (P') i (P''), wyrażające stanowisko semantycznego normatywizmu, albo — w terminologii Whitinga — po prostu normatywizmu, odpowiednio, *poprawność „mocno” preskrybująca* oraz *poprawność „słabo” preskrybująca*, ponieważ pierwsza wyrażona jest przy użyciu silnie normatywnego pojęcia „powinien”, a druga za pomocą słabszego pojęcia normatywnego „może”. (Cudzystów ma zwrócić uwagę, że różnica między poprawnością „mocno” preskrybującą a poprawnością „słabo” preskrybującą jest pozorna, ponieważ formuły te są równoważne). Jak zauważają zgodnie semantyczni antynormatywiści, przyjmując normatywną interpretację SN, uzyskujemy bardzo kontrintuicyjny rezultat. Okazuje się bowiem, że za każdym razem, kiedy wypowiadam fałszywe zdanie — niezależnie od tego, czy chcę okłamać mojego rozmówcę, żartuję czy zamierzam być ironiczna albo po prostu żywię fałszywe przekonanie — postępuję *wbrew* obowiązкови *semantycznemu*; robię coś, czego *nie powinnam robić* (por. Boghossian 2005a: 207; Hattiangadi 2009: 59; Glüer, Wikforss 2008: 4; Whiting 2009a: 546). Nie mówiąc już o tym, że semantyczny normatywista nie przedstawia żadnego argumentu, dlaczego mamy uznać, że częścią znaczenia słowa „kot” jest obowiązek lub racja do jego poprawnego stosowania — innymi słowy, dlaczego, żeby mieć na myśli *kota*, kiedy wypowiadam słowo „kot”, niezbędne jest *zobowiązanie* do poprawnego używania tego słowa? Dlaczego mamy przyjąć, że kiedy mam na myśli właśnie *kota*, moje myśli nie ograniczają się do kota, lecz rozciągają się na to, co mogę lub nie mogę, powinnam lub nie powinnam zrobić ze słowem „kot”? Normatywista, taki jak Whiting, nie wyjaśnia nam, dlaczego częścią znaczenia miałoby być użycie; dlaczego znajomość znaczenia jakiegoś słowa miałaby iść w parze ze zobowiązaniem do jego poprawnego użycia. Na zarzut ten Whiting odpowiada, osłabiając wymowę swoich wcześniejszych stwierdzeń i zarazem precyzując, co ma na myśli. Otóż jeśli „poprawność semantyczną”, a raczej zdanie o tym, co można lub powinno się zrobić, zinterpretuje się jako rację *prima facie*, stanowisko normatywisty stanie się mniej problematyczne (por. Whiting 2009a: 536). Racja *prima facie* to bowiem taka racja, która *niekoniecznie* jest racją decydującą. Wyrażenie idei normatywizmu semantycznego w kategoriach racji *prima facie* ma tę zaletę, że zgadza się z naszymi potocznymi intuicjami dotyczącymi racji. Czy nie sądzimy, że jeśli słowo „kot” oznacza kota, a nie jakieś inne zwierzę, to mamy rację (być może nierozstrzygającą) używać

słowa „kot” w odniesieniu do kota? Wreszcie, czy semantyczny antynormatywista, który mocno oponuje przeciwko *kategorycznej* interpretacji normatywności znaczenia, nie twierdzi czegoś podobnego, to znaczy, że na ogół użytkownicy danego języka mają rację, żeby używać wyrażen tego języka zgodnie z ich znaczeniem, a zatem dostrzegają w normie semantycznej rację *prima facie*?

Wprowadzenie do dyskusji poświęconej wiarygodności tezy na temat normatywności znaczenia, jednego z kluczowych pojęć w teorii działania i moralności, mocno ją gmatwa. Dlatego warto w tym miejscu powrócić do zasygnalizowanego w punkcie 3. problemu z interpretacją kluczowych terminów używanych w tej debacie.

6. Normatywność i racje praktyczne

Zacznijmy od przypomnienia, że kiedy semantyczny antynormatywista używa terminu „racja”, używa go w zgodzie ze standardowym rozumieniem przyjętym w filozofii praktycznej. *Racja* jest relacją wiążącą sprawcę z działaniem: *ktoś* ma¹⁰ rację zrobić *coś* w warunkach *W*. Racja opiera się na faktach, lecz się do nich nie sprowadza. Żeby coś stanowiło rację dla kogoś do zrobienia czegoś, niezbędne wydaje się istnienie jakiegoś związku owego „czegoś” z układem motywacyjnym działającego. Owo „coś” stanowi albo nieodzowną część celu, którego się pragnie, albo po prostu nim jest.

W związku z tym nie istnieje racja w ścisłym sensie „semantyczna”, to znaczy taka, która wypływa z samego faktu, że wyrażenia danego języka mają określone warunki poprawności i która zarazem jest niezależna od pragnień, celów czy intencji danego użytkownika konkretnego języka. Semantyczni antynormatywiści w ogóle zachowują sceptycyzm względem realizmu normatywnego w sensie Mackie’go (tzn. co do istnienia własności, które charakteryzowałyby się wbudowaną preskryptywnością) w odniesieniu do wszelkich rodzajów racji, przede wszystkim moralnych, ale też poznawczych. Na przykład to, że dwa plus dwa daje cztery, nie jest jeszcze w ich opinii racją do tego, żeby mieć przekonanie, że dwa plus dwa daje cztery. Niezależnie jednak od

¹⁰ Niektórzy filozofowie, np. Mark Schroeder, nie bez słuszności zauważają, że ten potoczny sposób wyrażania się o racjach w kategoriach posesywnych (*mam* rację) jest mylący, ponieważ sugeruje, że racja jest czymś, co się *nabywa* i *posiada*, zatem ma charakter subiektywny, gdy tymczasem bardziej właściwym ujęciem racji jest to, w którym akcentuje się komponent niezależności racji od podmiotu. „Racja” jest raczej terminem relacyjnym: coś znajduje się w relacji bycia racją dla kogoś. To, czy coś jest racją, w sensie posiadania własności *bycia* racją, przemawiająca za zrealizowaniem określonego działania, jest w istotnym sensie kwestią zachodzących stanów rzeczy czy faktów, np. tego, czy dziś wieczór jest przyjęcie u Juliana, oraz tego, że ja lubię chodzić na przyjęcia, zwłaszcza te organizowane przez Juliana. To, że Julian urządza dziś wieczorem przyjęcie, oraz to, że lubię przyjęcia organizowane przez Juliana, *staje się* dla mnie racją, żeby dziś wieczór iść do niego, o ile pójście na przyjęcie jest tym, co chcę robić dzisiejszego wieczora i zarazem *jest* dla mnie racją, żeby tam iść, jeśli prawdą jest, że szczególnie lubię przyjęcia organizowane przez Juliana. Więcej zob. Schroeder (2008).

tego, czy antynormatywiści mają rację, że nie ma racji niesubiektywnych, tzn. niezrelatywizowanych do pragnień, celów i intencji konkretnych podmiotów, wydaje się, iż mają oni rację, utrzymując, iż fakt, że „kot” oznacza kota, sam nie rodzi jeszcze dla nikogo racji do używania słowa „kot” w taki czy inny sposób. Jeśli semantyczny normatywista jest odmiennego zdania, to powinien wykazać, że fakty dotyczące znaczenia wyrażeń są *przynajmniej* tak szczególnym rodzajem motywacyjnie skutecznych faktów, jak fakty moralne w ujęciu realistów moralnych. Dopóki zwolennik normatywności znaczenia przekonująco nie uargumentuje normatywnej „samowystarczalności” faktów dotyczących poprawnego użycia wyrażeń, dopóty jego twierdzenie nie jest niczym innym niż wyrazem intuicji dotyczących funkcjonowania pewnych pojęć — intuicji, których jego oponent, nie podziela.

Innymi słowy, semantyczny antynormatywista nie twierdzi, iż *to, że „kot” oznacza kota*, nie może być dla kogoś racją do użycia słowa kot w zgodzie z jego ekstensją, jest tylko zdania, że fakty dotyczące znaczenia wyrażeń są preskryptywnie nieaktywne. W związku z czym odrzuca sugestię Whitinga, że znaczenia mogą być *prima facie* racjami, ponieważ koncepcja racji *prima facie* zakłada realizm racji, tzn. pogląd, że racje biorą się z obiektywnego charakteru konkretnej sytuacji, są faktami, które błędnie interpretuje się jako przemawiające za działaniem w wyniku nieuwzględnienia wszystkich istotnych dla konkretnej sytuacji czynników (por. Ross 2007: 19). W terminologii wprowadzonej przez Rossa to, co Whiting określa mianem *prima facie* racji, jest *obowiązkiem prima facie*. Ta zmiana terminologiczna niesie ze sobą istotne konsekwencje teoretyczne i interpretacyjne. Zaczniemy od tych pierwszych: kiedy Whiting twierdzi, że samo znaczenie jakiegoś wyrażenia jest *prima facie* racją do jego stosowania zgodnie z warunkami poprawności tego wyrażenia, to zdaje się mieć na myśli tylko tyle: to, że poprawnie jest używać słowa „kot” w odniesieniu do kota, *przemawia za tym*, żeby słowa „kot” używać, gdy chce się powiedzieć coś o kocie, a nie o łasicy, ale nie znaczy to jeszcze, że użycie słowa „kot” w odniesieniu do kota *zawsze* jest tym, co *powinno się* lub *ma się* rację zrobić. Whiting, prawdopodobnie nieświadomie, uznaje koncepcję racji podobną do przedstawionej przez T. M. Scanlona w książce *What We Owe to Each Other* (1998; por. Cuneo 2007; Dancy 2000; Parfit 2006; Shafer-Landau 2003). Scanlon przedstawia w niej ujęcie racji jako *względu, który przemawia za czymś* (*consideration that counts for, speaks in favour*). W tym sensie racją może być cokolwiek, co nadaje się na przesłankę (rozsądnego) sylogizmu praktycznego. W takim sensie racjami są fakty interpretowane jako *prawdziwe sądy na temat stanów świata*. Realistyczna teoria racji pasuje do stanowiska Rossa w tym sensie, że obaj filozofowie są realistami w odniesieniu do racji. W tym miejscu podobieństwo się jednak kończy, ponieważ Ross pisze o *obowiązkach prima facie*. Racje *prima facie* rozumiane jako fakty przemawiające za jakimś działa-

niem stosunkowo łatwo jest zastąpić innymi racjami, np. jeśli preferencje ulegną zmianie. Jest tak dlatego, że ich siła normatywna jest stosunkowo słaba. To, że lubię lody waniliowe, oraz fakt, że znajdują się one w mojej zamrażarce, składa się na *prima facie* rację, żeby je zjeść, o ile tylko mam na nie właśnie ochotę. Jeśli jednak w drodze do kuchni uznam, że większą przyjemność sprawi mi wypicie kakao, to *to, że lody waniliowe znajdują się w moim zamrażalniku*, przestanie być dla mnie racją, żeby je zjeść lub pozostanie racją pozorną, czyli racją *prima facie*. Sprawa nie jest tak prosta w przypadku *obowiązków prima facie*. Jeśli obiecałam ci, że pójdziemy do kina, a w międzyczasie zadzwonił mój brat z pilną prośbą o odebranie jego dziecka z przedszkola, to wydaje się, że mój obowiązek dotrzymania słowa został przeważony przez mój obowiązek względem brata, a precyzyjnie przez jego treść: odebranie dziecka z przedszkola wydaje się rzeczą ważniejszą niż pójście do kina. Niezależnie jednak od tego, co okazuje się moim prawdziwym obowiązkiem w tej konkretnej sytuacji, dotrzymanie obietnicy, że pójdę z tobą do kina, nie staje się mniej ważne, nie przestaje być tym, co *powinnam była zrobić*, gdyby nie zaszły nadzwyczajne okoliczności (telefon od mojego brata). Ponieważ zaistniały nowe okoliczności, moja obietnica przestała być obowiązkiem właściwym, lecz pozostała obowiązkiem *prima facie*. I tu pojawia się druga wątpliwość, interpretacyjna: czy Whiting utożsamia semantyczne *racje prima facie* z obowiązkiem *prima facie* w rozumieniu Rossa? Wreszcie: czy takie utożsamienie jest zasadne? Wydaje się, że na pierwsze pytanie odpowiedź jest pozytywna. W istocie Whiting uważa, że *racje prima facie*, podobnie jak obowiązki *prima facie* w ujęciu Rossa, są niezależne od pragnień i innych stanów psychologicznych podmiotu. W rozumieniu obydwu autorów *to, co należy zrobić*, zależy od tego, *jak się rzeczy mają*, a zatem od tego, jaki jest *obiektywny charakter sytuacji*. Ale jeśli podobieństwo racji *prima facie* w ujęciu Whitinga i obowiązku *prima facie* w rozumieniu Rossa ma być uzasadnione tym, że w obu przypadkach źródłem racji jest to, jak się rzeczy mają, a nie to, czego podmiot pragnie, ani to, co sądzi na jakiś temat itd., to moja *prima facie* racja, żeby zjeść lody waniliowe, jest zarazem obowiązkiem *prima facie*! A to dlatego, że *to, że lubię lody waniliowe*, jest faktem obiektywnym na temat tego, jakie mam preferencje; faktem, wobec którego mogę zajmować dowolną postawę (np. być rozczarowana) i który w połączeniu z innymi faktami (lody są w mojej zamrażarce, jestem głodna, nie mam do jedzenia nic poza lodami i kakao itd.) składa się na *obiektywny charakter sytuacji*, który jest źródłem zarówno racji, jak i obowiązku *prima facie*. Niemniej wniosek, że mam *prima facie* obowiązek zjeść lody waniliowe podobnie jak mam *prima facie* obowiązek dotrzymać słowa, wydaje się absurdalny. Jeśli Whiting chce uniknąć takiej konkluzji, musi wyjaśnić nam, dlaczego *prima facie* racja *semantyczna* miałaby zachowywać się jak moralny obowiązek *prima facie*, a *prima facie* racja odwołująca się do czyichś preferencji już nie. Whiting uważa za

rzecz bezdyskusyjną, że *jeśli* prawdą jest, iż powinno się dotrzymywać słowa, to prawdziwość tego zdania implikuje *prima facie* obowiązek dotrzymywania słowa (por. Wedgwood 2007). Analogicznie, jeśli prawdą jest, że powinno się używać słów zgodnie z tym, co oznaczają, to prawdziwość tego zdania implikuje *prima facie* rację, aby używać słów zgodnie z ich ekstensją. Pogląd znajdujący wyrażenie w powyższych sformułowaniach stanowi przykład *mocnego realizmu w odniesieniu do racji*. Trzy charakterystyczne cechy tego stanowiska to: 1) racje mają postać prawdziwych sądów, którym odpowiadają stany rzeczy w świecie (są *faktami*); 2) *fakty* są niekontrowersyjnie normatywne w tym sensie, że są potencjalnymi racjami działań (to, jaki *fakt* (jakie *fakty*) stanie się aktualną racją, zależy od kontekstu); oraz 3) *fakty* są racjami działań, niezależnymi od pragnień czy przekonań podmiotu. Największy problem z mocnym realizmem racji zastosowanym do znaczenia wyrażen polega na tym, że nie za bardzo wiadomo, dlaczego norma wyrażająca warunki poprawnego stosowania określonego wyrażenia miałaby być racją *niezależną* od stanów psychologicznych konkretnego użytkownika danego języka, miałaby być semantycznym obowiązkiem *prima facie*? Dlaczego, niezależnie od tego, *co chcę powiedzieć* (jaka jest moja intencja komunikacyjna), *powinam* używać słów zgodnie z ich znaczeniem? To, jak używam słów, jest kwestią pragmatyczną¹¹; jeśli słowa „bogaty” umyślnie użyję w odniesieniu do osoby niezamożnej, to niewątpliwie użyję go w sposób niepoprawny, niezgodny z jego znaczeniem. Dlaczego jednak, jak utrzymują semantyczni antynormatywiści, miałoby to oznaczać, że zrobiłam coś, czego nie powinnam była zrobić, oraz że w związku z tym moje zachowanie słowne zasługuje na krytykę? Whiting uważa, że antynormatywiści są w błędzie, bowiem nawet jeśli w konkretnych okolicznościach mogłoby się okazać, że ze względów pozajęzykowych niepoprawne zastosowanie słowa „kot” jest uzasadnione, niemniej nie zmienia to faktu, że używając słowa „kot” w odniesieniu do nie-kota, użytkownik języka polskiego zachował się *niewłaściwie* z semantycznego punktu widzenia, a ponieważ pojęcie „niewłaściwy” ma intuicyjnie charakter normatywny, krytyka — niechby była ona słaba —

¹¹ Przypomnijmy, że cała debata dotycząca normatywności znaczenia toczy się w ścisłe zdefiniowanych ramach: spór nie dotyczy tego, czy jakaś *konkretna* teoria znaczenia może uwzględnić warunek normatywności sformułowany przez Kripkego (na przykład jeśli uznamy pragmatykę za nierozdzielnie związaną z semantyką, pominiemy różnicę między *znaczeniem a użyciem*, lecz wówczas pytanie o to, czy *samo* znaczenie jest normatywne okaże się źle postawione), lecz o to, czy można obronić sugestię Kripkego, jakoby normatywność była rozstrzygającym, przedteoretycznym testem, który *każda* teoria znaczenia musi spełnić, aby zasługiwać na miano *teorii znaczenia*; por. Wikforss (2001: 207). Semantyczni antynormatywiści zaprzeczają, że (P) *s* znaczy $F \rightarrow \forall x$ (*s* stosuje się poprawnie względem $x \leftrightarrow x$ jest *f*), w którym „poprawnie” odnosi się do *poprawności semantycznej*, uzasadnia interpretację preskryptywną. Ich argument jest następujący: jeśli w miejsce wyrażenia „stosuje się poprawnie”, w (P) podstawimy dowolną inną relację na określenie relacji słowo — świat, niech to będzie „odnosi się poprawnie do...”, „jest prawdą o *s*, że...” czy „denotuje”, to nie wynikają *żadne bezpośrednio* normatywne konsekwencje dotyczące użycia *s*; por. Hattiangadi (2009: 55); Glüer, Wikforss (2008: 2).

jest uzasadniona¹². Przekonanie Whitinga wydaje się trafne i nietrafne zarazem. Jednak w obu przypadkach nie wiadomo, w jaki sposób miałyby to pomóc normatywiście. Jest ono trafne w tym sensie, w jakim trafny jest opis czyjegoś zachowania w świetle danej normy (semantycznej czy jakiegokolwiek innej), nietrafne natomiast z tego względu, że nie jest wcale jasne, czemu termin „właściwy” w tym kontekście należy interpretować w sposób normatywny, a nie jedynie ewaluatywny. O ile niekontrowersyjne wydaje się to, że normatywność jest związana z ewaluacją, o tyle wydaje się, że relacja nie zachodzi w drugą stronę: nie wszystko, co jest ewaluatywne, jest *zarazem* normatywne. Jeśli powiesz, że mój strój jest właściwy, nie musi to zaraz oznaczać, że ubrałam się tak, *jak powinnam*, może być również wyrazem oceny, iż biorąc pod uwagę sytuację i pewne normy obyczajowe, mój strój jest po prostu stosowny. Jeśli ubrałabym się niezgodnie z *dress code*m, za to zgodnie z zasadą „noś ubranie, w którym czujesz się komfortowo”, pogwałcenie normy obyczajowej samo przez się nie oznacza jeszcze, że zachowałam się nie tak, jak powinnam była się zachować. To, jak powinnam była się zachować, nie jest kwestią zastosowania się do normy, która pasuje do okoliczności, w których się znajduję, lecz racji, dla których postąpiłam tak, jak postąpiłam.

O kimś, kto w sposób rozmyślny systematycznie używałby słów niezgodnie z ich znaczeniem, można by powiedzieć, że jest nierozsądny, ale nie, że robi coś, czego *nie powinien* robić. W tym sensie norma znaczeniowa, czyli fakt, że wyrażenia danego języka mają określoną denotację, jest niejako *propozycją* dotyczącą używania słów danego języka. Propozycja z definicji ma warunkowy charakter: *jeśli* chcesz użyć słowa „kot” poprawnie, masz rację, aby użyć go w odniesieniu do kota, a nie łasicy. Co prawda racja, żeby użyć słowa „kot” w określony sposób *zależy* od tego, jakie obiekty słowo „kot” w języku polskim w rzeczywistości denotuje (i jeśli normatywista twierdzi, że warunki poprawności określonego wyrażenia w konkretnym języku są *obiektywne* w tym sensie, że są niezależne od stanów psychologicznych użytkownika tego języka, antynormatywista bez trudu się z nim zgadza), lecz nie jest *bezpośrednią konsekwencją* faktu, że „kot” oznacza kota.

7. Podsumowanie

Spór między semantycznym normatywistą a antynormatywistą wydaje się mieć charakter fundamentalny a zarazem nierozstrzygalny. Pierwszy jako realista w odniesieniu do racji broni silnej tezy, wedle której prawdziwość zdań dotyczących norm implikuje racje do działania (niedziałania) zgodnego z treścią danej normy. Z kolei jego przeciwnik, semantyczny antynormatywista, zajmuje stanowisko antyrealistyczne w odniesieniu do racji, tzn. uważa, iż z samego

¹² Whiting buduje swój argument w oparciu o niewłaściwe użycie słowa „bogaty”. Zob. Whiting (2009a: 548).

faktu, że dana norma istnieje lub jest ważna (*in force*), żadne bezpośrednio normatywne konsekwencje nie wynikają. *Racje* nie są po prostu pochodną tego, jak się rzeczy mają. Czy można zasadnie powiedzieć, że któraś ze stron ma przewagę? Jak zapowiedziałam we *Wprowadzeniu*, sądzę, że można. W tym celu należy przypomnieć sobie, czego *dokładnie* dotyczy dyskusja. Otóż debata dotyczy tego, czy znaczenie językowe *jako* *znaczenie językowe* jest z *istoty* (wewnętrznie, bez dodatkowych warunków) normatywne w sensie preskryptywnym, tzn. dostarcza racji do działania. Racja w tym kontekście nie jest luźnym terminem, który znaczy tyle co „wzgląd przemawiający za czymś” — w tak luźnym sensie racją jest każdy fakt. Termin „racja” używany jest jako pojęcie techniczne, zapożyczone z filozofii praktycznej, w którym oznacza relację między sprawcą a działaniem. Jeśli zatem problem przedstawimy w formie pytania: czy to, że „kot” oznacza **kota**, jest lub implikuje *rację* dla *kogokolwiek* do *używania* słowa „kot” zgodnie z tym, co słowo to oznacza, odpowiedź brzmi: nie. W *tych* sensie terminu „racja” „kot” nie jest racją, choć z pewnością w swobodniejszym, potocznym sensie może nią być. Zwróćmy uwagę, że nasza odpowiedź nie jest kwestią intuicji, lecz odwołuje się do znaczenia takich terminów, jak „fakt”, „racja”, „norma”, „preskryptywny”, oraz tego, w jaki sposób dyskusja została sprofilowana (zob. przypis 11). Przewaga semantycznego antynormatywisty w tak *konkretnie* zdefiniowanym sporze nie oznacza ostatecznej dyskwalifikacji semantycznego normatywizmu. Stawia jednak zwolennika normatywności znaczenia przed następującym dylematem. Aby przeważać szalę na swoją stronę, ma on dwa wyjścia: albo przekona nas, że tezę dotyczącą normatywności znaczenia należy interpretować w sposób luźniejszy, potoczny, ale wówczas przedmiot jego sporu z antynormatywistą zniknie. Albo, żeby uwiarygodnić swoje stanowisko, pokaże, że norma znaczeniowa jest *przynajmniej* tak samo normatywna jak norma moralna. Jest to strategia bardzo nieatrakcyjna: trafność argumentu w jednej dziedzinie zależna jest od jego trafności w innej dziedzinie. W dodatku jest bardzo ryzykowna: jeśli okazałoby się, że z powodzeniem można zakwestionować wewnętrzną preskryptywność normy moralnej, normatywista w odniesieniu do znaczenia straciłby punkt zaczepienia dla swojej tezy (Klimczyk MS). Jeśli bowiem paradygmatyczna wewnętrzna normatywność normy moralnej jest fikcją, jeszcze większą fikcją musi wydać się przypuszczenie, że „kot” komukolwiek cokolwiek nakazuje.

Bibliografia

- BILGRAMI A. (1993), *Norms and Meaning*, [w:] Stoecker R. (red.), *Reflecting Davidson*. Berlin: New York, 121–144.
- BLACKBURN S. (1984), *The Individual Strikes Back*, „Synthese” 58, 281–301.

- BLOOR D. (1997), *Wittgenstein, Rules and Institutions*. London: Routledge.
- BOGHOSSIAN P. (1989), *The Rule-Following Considerations*, "Mind" 98, 507–549.
- BOGHOSSIAN P. (2003), *The Normativity of Content*, "Philosophical Issues" 13, 31–45.
- BOGHOSSIAN P. (2005a), *Is Meaning Normative?*, [w:] Beckerman A., Nitz C. (red.), *Philosophy — Science — Scientific Philosophy*. Paderborn: Mentis.
- BOGHOSSIAN P. (2005b), *Rules, Meaning and Intention-Discussion*, "Philosophical Studies" 124, 185–197.
- BOGHOSSIAN P. (2008), *Epistemic Rules*, "The Journal of Philosophy" 105 (9), 472–500.
- BRANDON R. (1995), *Making It Explicit*. Cambridge, Mass.: Harvard University Press.
- BROOME J. (2004), *Reasons*, [w:] Wallace J., Smith M., Scheffler S., Pettit P. (red.), *Reason and Value: Themes from the Moral Philosophy of Joseph Raz*. Oxford: Oxford University Press.
- DAVIDSON D. (1984), *Inquiries Into Truth and Interpretation*. Oxford: Oxford University Press.
- DRETSKE F. (2000), *Perception, Knowledge and Belief: Selected Essays*. Cambridge: Cambridge University Press.
- FINLAY S. (2010a), *What Ought Probably Means, and Why You Can't Detach It*, "Synthese" 177 (1), 67–89.
- FINLAY S. (2010b), *Recent Work on Normativity*, "Analysis" 70 (2), 331–346.
- GIBBARD A. (2003), *Thoughts and Norms*, "Philosophical Issues" 13, 83–98.
- GIBBARD A. MS. *Oughts and Thoughts*, <http://www-personal.umich.edu/~gibbard/bk10/Book-draft-Contents.pdf>.
- GLOCK H.-J. (1996), *Necessity and Normativity*, [w:] Sluga H., Stern D. G. (red.), *The Cambridge Companion to Wittgenstein*. Cambridge: Cambridge University Press, 198–225.
- GLÜER K. (1999), *Sense and Prescriptivity*, "Acta Analytica" 14, 111–128.
- GLÜER K. (2001), *Dreams and Nightmares: Conventions, Norms and Meaning in Davidson's Philosophy of Language*, [w:] Kotatko P., Pagin P., Segal G. (red.), *Interpreting Davidson*. Stanford: CSLI Publications, 53–74.
- GLÜER K., WIKFORSS Å. (2008), *Against Normativity Again: Reply to Whiting*, MS, <http://people.su.se/%7Ekgl/Reply%20to%20Whiting.pdf>
- GLÜER K., WIKFORSS Å. (2009a), *Against Content Normativity*, "Mind" 118, 31–70.
- GLÜER K., WIKFORSS Å. (2009b), *The Normativity of Meaning and Content*, [w:] Zalta E. (red.), *Stanford Encyclopedia of Philosophy*, <http://plato.stanford.edu/entries/meaning-normativity/>.
- GLÜER K., WIKFORSS Å. (2010a), *The Rule Is Not Needed: Rule-Following and Meaning in Later Wittgenstein*, [w:] Whiting D. (red.), *The Later Wittgenstein on Meaning*. Basingstoke: Palgrave Macmillan, 148–166.
- GLÜER K., WIKFORSS Å. (2010b), *The Truth Norm and Guidance: A Reply to Steglich-Petersen*, "Mind" 119, 757–761.
- GLÜER K., PAGIN P. (1999), *Rules of Meaning and Practical Reasoning*, "Synthese" 117, 207–227.
- HATTIANGADI A. (2006), *Is Meaning Normative?*, "Mind and Language" 21, 220–240.
- HATTIANGADI A. (2007), *Oughts and Thoughts: Rule-Following and the Normativity of Content*. Oxford: Oxford University Press.
- HATTIANGADI A. (2009), *Some More Thoughts on Semantic Oughts: A Reply to Daniel Whiting*, "Analysis" 69, 54–63.
- HUMBERSTONE I. L. (1971), *Two Sorts of 'Ought's'*, "Analysis" 32, 8–11.

- KLIMCZYK J. MS. *The Myth of Intrinsic Normativity*.
- KRIPKE S. (1982), *Wittgenstein on Rules and Private Language*. Oxford: Blackwell.
- KRIPKE S. (2005), *Wittgenstein o regułach i języku prywatnym* (przeł. K. Pośłajko, L. Wroński). Warszawa: Fundacja Aletheia.
- KUSCH M. (2006), *A Sceptical Guide to Meaning and Rules: Defending Kripke's Wittgenstein*. Chesham: Acumen.
- MCDOWELL J. (1994), *Mind and World*. Cambridge: Cambridge University Press.
- MILLER A. (2006), *Meaning Scepticism*, [w:] Devitt M., Hanley R. (red.), *The Blackwell Guide to the Philosophy of Language*. Oxford: Blackwell.
- MILLER A. (2010a), *The Argument From Queerness and the Normativity of Meaning*, [w:] Grajner M., Rami A. (red.), *Truth, Existence and Realism*. Paderborn: Mentis.
- MILLER A. (2010b), *Semantic Realism and the Argument from Motivational Internalism*, [w:] Schantz R. (red.), *What Is Meaning?* Berlin–New York: Mouton de Gruyter.
- PAPINEAU D. (1999), *Normativity and Judgment*, "Aristotelian Society Supplementary Volume" 73, 17–43.
- PARFIT D. (2011), *On What Matters*. New York: Oxford University Press, vol. 2.
- PETTIT P. (1990), *The Reality of Rule-Following*, "Mind" 99, 1–21.
- ROSS D. (2007), *The Right and the Good*, Stratton-Lake Ph. (red.). Oxford: Oxford University Press.
- SCHROEDER M. (2008), *Having Reasons*, "Philosophical Studies" 139, 57–71.
- SCHROEDER M. (2011), *Ought, Agents, and Actions*, "Philosophical Review" 120 (1), 1–41.
- STEGLICH-PETERSEN A. (2010), *The Truth Norm and Guidance: A Reply to Glüer and Wikforss*, "Mind" 119, 1–7.
- WRIGHT G. H. VON (1963), *Norm and Action* (1958–1960 Gifford Lectures, St. Andrews), <http://www.giffordlectures.org/Browse.asp?PubID=TPNORM&Cover=TRUE>.
- WEDGWOOD R. (2007), *The Nature of Normativity*. Oxford: Oxford University Press.
- WHITING D. (2007), *The Normativity of Meaning Defended*, "Analysis" 67, 133–140.
- WHITING D. (2009a), *Is Meaning Fraught With Ought?*, "Pacific Philosophical Quarterly" 90, 535–555.
- WHITING D. (2009b), *On Epistemic Conceptions of Meaning: Use, Meaning and Normativity*, "European Journal of Philosophy" 17, 416–434.
- WHITING D. (w druku), *Should I Believe the Truth?*, "Dialectica".
- WIKFORSS Å. (2001), *Semantic Normativity*, "Philosophical Studies" 102, 203–226.
- WRIGHT C. (1984), *Kripke's Account of the Argument Against Private Language*, "The Journal of Philosophy" 81, 759–777.

Rozdział 12

Przyczynowe teorie nazw a argument Stephena Sticha¹

Tadeusz Ciecierski, Katarzyna Kuś (Uniwersytet Warszawski)

Słowa kluczowe: argument Sticha, metasemantyka, nazwa własna, przyczynowa teoria nazw, teoria bezpośredniego odniesienia/bezpośrednie odniesienie, teoria hybrydowa nazw własnych, teoria łańcucha przyczynowego, teoria okazjonalna nazw własnych

Pod hasłem przyczynowej teorii nazw kryje się we współczesnej filozofii języka pogląd (a właściwie zespół poglądów) na temat semantycznych oraz metasemantycznych właściwości wyrażeń użytych imiennie². Rozważywszy pytania semantyczne „Co jest wartością semantyczną wyrażenia użytego imiennie?” oraz „Jaki jest wkład użytego imiennie wyrażenia w warunki prawdziwości zdań, w których wyrażenie to występuje?”, zwolennik przyczynowej teorii nazw udziela odpowiedzi: „Indywidualum, do którego wyrażenie użyte imiennie się odnosi”. Twierdzenie to ma również implikować stanowisko negatywne względem konotacyjnych lub deskrypcyjnych teorii nazw własnych — wkładem w warunki prawdziwości zdania nie są warunki opisowe, które mają w jakikolwiek sposób identyfikować przedmiot. Nawiązując do znanej koncepcji semantycznej Davida Kaplana, można semantyczną część przyczynowej teorii nazw określić mianem *teorii bezpośredniego odniesienia*. Ta decyzja terminologiczna jest uzasadniona także przez niebanalne związki i podobieństwa w funkcjonowaniu nazw własnych (wyrażeń użytych imiennie) oraz wyrażeń

¹ Tekst ukazał się po raz pierwszy [w:] *Przegląd Filozoficzny — Nowa Seria*, 2010, nr 3 (75), s. 35–52. Bardzo dziękujemy Joannie Odrowąż-Sypniewskiej oraz Mieszkowi Tałasiewiczowi za uwagi do wcześniejszej wersji tekstu oraz redaktorom tomu — Joannie Odrowąż-Sypniewskiej i Justynie Grudzińskiej za zgodę na przedruk tekstu.

² Rozróżnienie pojęciowe na *semantykę* i *metasemantykę* — teorię wartości semantycznych oraz teorię objaśniającą przyczyny, dla których wyrażenia mają określone wartości semantyczne — pochodzi od Davida Kaplana (por. Kaplan 1989, zwłaszcza s. 573–576). Jego ciekawe zastosowania filozoficzne może odnaleźć Czytelnik w pracach Roberta Stalnakera (zob. zwłaszcza *Ways a World Might Be*) oraz Agustina Rayo (zob. zwłaszcza *A Metasemantic Account of Vagueness*).

okazjonalnych³. Odpowiadając na pytanie metasemantyczne „Na mocy jakich faktów wyrażenie użyte imiennie ma taką, a nie inną wartość semantyczną?”, zwolennik przyczynowej teorii nazw udziela odpowiedzi: „fakty te obejmują historię użyc danego wyrażenia od momentu (wzorcowego) pierwszego chrztu przez kolejne jego użycia aż do aktualnego jego użycia”. Metasemantyczną część przyczynowej teorii nazw można nazwać *teorią łańcucha przyczynowego*.

Z dwóch wspomnianych składników teorii przyczynowej pierwszy — *teoria bezpośredniego odniesienia* — jest, w opinii autorów niniejszej pracy, stanowiskiem słusznym. Argumenty, które przytacza się na jego korzyść, są bardzo mocne, te zaś, które przytacza się przeciw, albo okazują się dotyczyć także konkurencyjnego stanowiska deskryptywistyczno-konotacyjnego⁴, albo opierają się na rozlicznych nieporozumieniach, wśród których najbardziej rzucającym się w oczy jest nieodróżnianie semantycznego i metasemantycznego aspektu teorii przyczynowej⁵. Wiele słusznym wątpliwości — także wśród autorów uważanych za zwolenników przyczynowej teorii nazw — wzbudza natomiast metasemantyczna część tej koncepcji — *teoria łańcucha przyczynowego*. Wokół niej właśnie skupiona jest na przykład subtelna i szczegółowa krytyka Garetha Evansa (1973). Temu też aspektowi teorii przyczynowej chcielibyśmy poświęcić niniejszy artykuł. W szczególności zamierzamy rozważyć implikacje, które ma dla *teorii łańcucha przyczynowego* pewien następujący kłopotliwy przykład podany przez Stephena Sticha (1986: 96–97).

W przykładzie mamy do czynienia z dwiema osobami: Borysem A. Nogo-odnikiem i Borysem B. Nogoodnikiem. Borys A. dokonuje morderstwa, lecz jego ofiara przed śmiercią go rozpoznaje i zdąża ostatkiem sił powiedzieć policjantowi (powiedzmy, że jest nim Jan Nowak): „To Borys Nogoodnik!”. Policja prowadzi śledztwo i na podstawie śladów wnioskuje o właściwościach zabójcy (dowiaduje się na przykład, że pali papierosy marki Popularne, jest kibicem hokeja na trawie, ma numer buta 45). Na podstawie przedśmiertnej wypowiedzi zamordowanego przyjmuje także, że zabójcą jest ktoś im bliżej nieznanym, do którego ofiara odnosiła się, lub próbowała się odnosić, za pomocą nazwy „Borys Nogoodnik”. Założmy teraz, że zbierane przez policję dane stopniowo wskazują na Borysa B. Nogoodnika. Według Sticha istnieje silna intuicja przemawiająca za tym, że gdy policjant Jan Nowak mówi: „Zabójcą jest Borys Nogo-odnik”, odnosi się raczej do Borysa B., a nie do Borysa A. Dzieje się tak, mimo że domniemany związek przyczynowy odnaleźć możemy jedynie między wypowiedziami Jana Nowaka a Borysem A., do którego odnosiła się ofiara, gdy

³ Często zwraca się uwagę na to, że już u Russella tzw. *nazwy własne w sensie logicznym* były po prostu wyrażeniami okazjonalnymi (przede wszystkim chodzi tu o zaimek wskazujący „to”), które to z kolei otrzymywały u Russella analizę w kategoriach *teorii bezpośredniego odniesienia*. Russell łączył zatem deskrypcjonizm w kwestii zwykłych nazw własnych z teorią bezpośredniego odniesienia w kwestii nazw własnych w logicznym sensie tego słowa.

⁴ Zob. Kripke (1996: 386–387).

⁵ Zob. np. Stalnaker (2003: 171–181), a także Hess (2009).

w obecności policjanta wypowiadała zdanie „To Borys Nogoodnik!”. Policja weszła w posiadanie nazwy „Borys Nogoodnik” użytej w odniesieniu do Borysa A. Nazwa ta powinna więc zachować to odniesienie również przy jej dalszych użyciach przez policję i powinna za każdym razem niezależnie od tego, czy osoby jej używające są w stanie zidentyfikować go jednoznacznie, odnosić się do Borysa A., lecz nie do Borysa B. Problem dla zwolennika *teorii łańcucha przyczynowego* — przynajmniej wtedy, gdy rozważamy ją na tym ogólnym i abstrakcyjnym poziomie — powstaje, jeśli zgodzi się on z intuicją, że mimo uczestnictwa policjantów w odpowiednim łańcuchu przyczynowym, w którym wszyscy uczestnicy należą do tej samej wspólnoty językowej i dzielą intencję stosowania nazwy do tej samej osoby, policjanci po odnalezieniu Borysa B., twierdząc, że zabójcą jest Borys Nogoodnik, odnoszą się do Borysa B. Przykład swój podsumowuje Stich stanowczym stwierdzeniem: „Wypadek ten [...] stawia [...] pod znakiem zapytania ogólną adekwatność przyczynowej teorii nazwy własnych” (Stich 1983: 97).

Opisaną wyżej sytuację można, obok znanych w literaturze przedmiotu wypadków, w których dochodzi do przerwania łańcucha lub do poplątania dwóch łańcuchów, dołączyć do listy zagadek, z którymi musi poradzić sobie każda dobra *teoria łańcucha przyczynowego*.

I. Warianty przykładu Sticha

Dla analizy powyższego przykładu istotne wydają się trzy następujące kwestie: po pierwsze, charakter związku między śladami a Borysem B., po drugie, to, czy zebrane przez policję dowody są jednoznacznie identyfikujące, a po trzecie, znaczenie, jakie mają określone warianty przykładu Sticha dla tych wersji *teorii łańcucha przyczynowego*, które (w przeciwieństwie na przykład do oryginalnej teorii Kripkego) mówią nam coś więcej o funkcji relacji przyczynowości w historii użyc danego wyrażenia (nazwy własnej) typu. Rozpocznijmy od pierwszych dwóch kwestii.

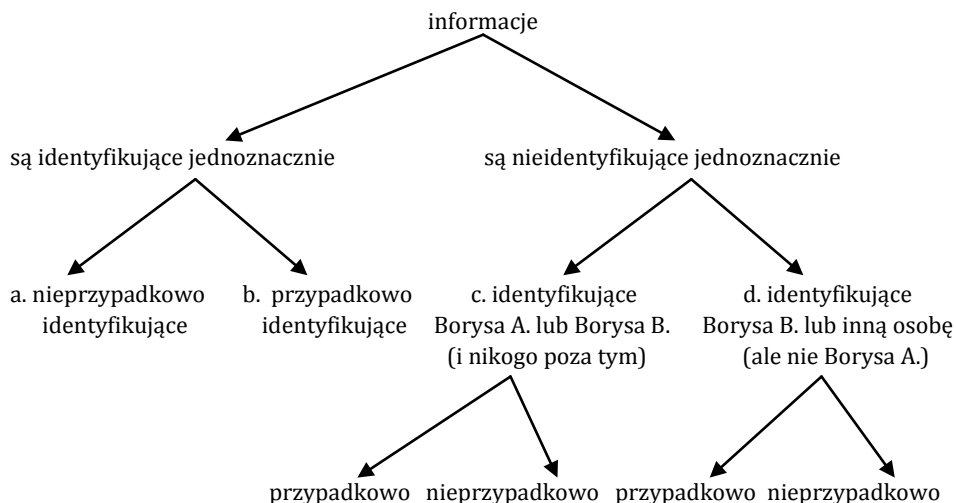
Przykład Sticha można interpretować na kilka sposobów — zależnie od tego, co dokładnie rozumiemy przez dowody wskazujące na konkretną osobę (tu Borysa B.). Możemy przede wszystkim przyjąć, że dowody znalezione na miejscu zbrodni jednoznacznie wskazują na pewną osobę, to znaczy istnieje jedna i tylko jedna osoba, która posiada własności związane z dowodami. To jednak może w kontekście naszej historii znaczyć jedną z dwóch rzeczy. Po pierwsze, możemy zakładać, iż prawdziwy zabójca jest bardzo przebiegłym szpiegiem i umyślnie tak spreparował dowody, że pozwalają jednoznacznie zidentyfikować jakąś konkretną znaną mu osobę, w naszym wypadku Borysa B.⁶ Druga możliwość jest natomiast następująca: dowody znalezione na miej-

⁶ Wydaje się, że ten właśnie wariant miał na myśli Stich, który twierdzi przy tym, że jeżeli ślady zostały spreparowane w taki sposób, aby wskazywały na Borysa B., to nie ma związku przyczynowego między nim a tymi śladami. Jest to dyskusyjna teza, którą zwolennik *Teorii łańcucha przyczynowego* z powodzeniem może kwestionować (por. poniżej).

scu zbrodni jednoznacznie identyfikują pewną osobę, robią to jednak w sposób przypadkowy oraz niezwiązany z przebiegłymi intencjami Borysa A. Możemy sobie na przykład wyobrazić, że kwadrans przed morderstwem w miejscu zbrodni przebywały trzy przypadkowe osoby, z których każda pozostawiła po sobie jeden ślad: *a* (odcisk buta o numerze 45), *b* (pudełko po papierosach Popularne) oraz *c* (przedarty bilet wstępu na mecz hokeja na trawie, który odbył się dzień przed morderstwem). Każdy z tych dowodów odrębnie jest nieidentyfikujący, ale wszystkie razem mogą identyfikować pewną konkretną osobę, którą może być nasz Borys B. Możemy także założyć, że Borys A. spreprował zbiór jednoznacznie identyfikujących dowodów, choć nie ma zielonego pojęcia, którą osobę one jednoznacznie identyfikują.

Jeszcze inna możliwość jest następująca: tak, jak to się często zdarza w powieściach i filmach detektywistycznych, dowody odnalezione na miejscu zbrodni jedynie z pewnym prawdopodobieństwem wskazują określoną osobę, a więc faktycznie nie są w żadnym razie identyfikujące. Taki z kolei wypadek ma dwa interesujące warianty. Po pierwsze, dowody mogą niejednoznacznie wskazywać jednocześnie na Borysa A. oraz Borysa B. (zauważmy, że fakt, czy osoby te faktycznie mają na imię „Borys Nogoodnik”, nie jest tu istotny, wystarczy, że policja zakłada, iż tak jest). Po drugie, dowody mogą wskazywać na Borysa B. oraz jakąś jeszcze inną osobę. W efekcie otrzymujemy następującą typologię interpretacji przykładu Sticha:

Typologia interpretacji przykładu Sticha



Należy tu poczynić jedną jeszcze ważną uwagę. Ofiara zabójstwa może użyć nazwy „Borys Nogoodnik” z dwójaką intencją. Po pierwsze, może chcieć po prostu przekazać w swojej wypowiedzi informację, która w jej przekonaniu

może najlepiej pomóc policji w identyfikacji zabójcy: że zabójca nazywa się „Borys Nogoodnik”. W takim wypadku mamy do czynienia z *deskrypcyjnym użyciem imienia własnego* (lub — równoważnie — z nieimiennym użyciem wyrażenia „Borys Nogoodnik”). Jeśli jednak ofiara wypowiada to zdanie nie z intencją podania informacji identyfikującej (są to wypadki mniej typowe, ale z całą pewnością możliwe), lecz ze zwykłą intencją charakterystyczną np. dla sytuacji, w których ktoś opowiada o znanej lub nieznaney nam bliżej osobie, używając jej imienia, to mamy do czynienia z imiennym użyciem wyrażenia „Borys Nogoodnik”. W sytuacji, gdy nie znaliśmy tej nazwy, zapoznanie się z odpowiednią wypowiedzią sprawia, że zaczynamy zbieranie pewnych wiadomości pod wspólnym szyldem i zakładamy, iż kryje się pod nim konkretna osoba. W sytuacji zaś, gdy znaliśmy już tę nazwę, to, o ile nie żywimy zarazem przekonania o jej przynależności do dwóch różnych praktyk językowych (wieloznaczności), jesteśmy po prostu gotowi uzupełniać o nowe informacje nasz — istniejący już — zbiór wiadomości na temat domniemanego jedyne go nosiciela tej nazwy. Jeśli którykolwiek z dwóch wyróżnionych typów użycia ma być problematyczny dla *teorii łańcucha przyczynowego*, to jest to jedynie drugi — imienny — sposób użycia. W pierwszym wypadku ofiara mogłaby bowiem po prostu powiedzieć „To osoba o imieniu *Borys Noogodnik*”.

II. Dwie wersje teorii łańcucha przyczynowego

Przykład Sticha dobitnie wskazuje, że rację miał Evans, pisząc, że:

Teoria przyczynowa ignoruje [...] znaczenie towarzyszącego kontekstu i traktuje zdolność denotowania czegoś jako magiczną sztukę, która została w jakiś sposób raz wykonana i której raz wykonanej nie można utracić (Evans 1993: 231).

Jednym z podstawowych zarzutów, które postawić można teorii przyczynowej, jest niejasna rola, którą w jej strukturze odgrywa pojęcie *przyczynowości*. Czy zakłada ono nieintencjonalny charakter determinant wartości semantycznych nazw? Co dokładnie jest w tej koncepcji przedmiotem związku przyczynowego? Drugie z pytań — pytanie o dziedzinę i przeciwdziedzinę relacji przyczynowej (lub — jeśli trzymać się metafory łańcucha przyczynowego — pytanie o naturę tego, co spaja ogniwa łańcucha) — jest tu szczególnie istotne. Rysują się tu przynajmniej dwie możliwe odpowiedzi: jedna inspirowana okazjonalną teorią Pelczara i Rainsbury’ego, druga inspirowana tzw. hybrydową teorią Evansa. Obie starają się powiedzieć coś więcej na temat roli pojęcia przyczynowości w przyczynowych teoriach nazw, odnoszą się także bezpośrednio do pytania o rolę czynników intencjonalnych przy odpowiedzi na pytania meta-semantyczne.

a) Okazjonalna teoria łańcucha przyczynowego

Zgodnie z pierwszą z wyżej wspomnianych teorii relacja przyczynowa łączy ze sobą pojedyncze wypowiedzi lub akty mowy, w których użyta jest dana nazwa własna. Modelowo na początku takiego łańcucha znajduje się akt mowy, w którym ktoś nadaje komuś lub czemuś pewne imię własne. Oczywiście w bardzo wielu wypadkach nie możemy mówić o żadnym wąsko pojętym akcie chrztu, możemy wtedy jednak zawsze odnaleźć jakiś proces, który jest ekwiwalentem takiego aktu. Na końcu łańcucha znajduje się natomiast określona wypowiedź, w której użyte jest to samo wyrażenie-typ. Ponieważ wiemy, że te same nazwy własne-typy mogą mieć różne odniesienia (*vide* „Katon”) oraz iż odniesienia nazw mogą zmieniać się w czasie (*vide* przykład Evansa nazwy „Madagaskar”), musimy uwzględniać w naszym opisie możliwość wystąpienia kilku aktów mowy będących aktami chrztu lub ich ekwiwalentami. W związku z tym w naszej konstrukcji powinniśmy nie tyle mówić o konkretnym wypadku użycia nazwy, co raczej o użyciu nazwy w kontekście, w którym ma moc pewien określony akt mowy, będący aktem chrztu lub jego ekwiwalentem. Tak więc jeśli w pewnej określonej rozmowie ktoś wygłasza zdanie „Katon domagał się zniszczenia Kartaginy”, to Grice’owska maksyma jakości decyduje o tym, że aktem mowy, który jest aktem pierwotnego chrztu, a który ma moc w tym kontekście, jest akt wiążący nazwę „Katon” z Katonem Starszym. Taka wersja teorii przyczynowej zaproponowana między innymi w pracy Pelczara i Rainsbury’ego (1998) określona zostanie mianem teorii okazjonalnej, ponieważ nazwy własne są w niej traktowane jak szczególnego typu wyrażania okazjonalne — takie, których odniesienie zależne jest od związku między aktualną wypowiedzią a mającym moc w kontekście tej wypowiedzi aktem mowy.

Warto zwrócić też uwagę na to, że na pytanie, jakie czynniki decydują o tym, iż w danych okolicznościach moc ma konkretny akt chrztu pierwotnego, nie ma prostej odpowiedzi. Jak widzieliśmy powyżej, znaczącą rolę mogą tu odgrywać czynniki pragmatyczne, takie jak założenia na temat intencji i przekonania uczestników sytuacji komunikacyjnej, racjonalności konwersacji itp. W żadnym razie nie możemy tu jednak twierdzić, że determinacja wartości semantycznej nazwy własnej ma charakter nieintencjonalny. Jak wskazuje przykład z Katonem, sprawy mają się wprost przeciwnie. To, że w tym kontekście moc ma akt chrztu związany z Katonem Starszym, jest wynikiem tego, iż zakładamy, że nadawca wypowiedzi posiada pewną wiedzę historyczną oraz że jego intencją jest powiedzieć coś prawdziwego.

Można się słusznie zapytać, gdzie w ogóle w tej koncepcji jest miejsce dla przyczynowości. Związek przyczynowy między pewną aktualną wypowiedzią a mającym miejsce w przeszłości aktem mowy w przytłaczającej liczbie wypadków po prostu nie zachodzi. Możemy co najwyżej powiedzieć, że w indywidualnej historii nadawcy wypowiedzi musi istnieć zdarzenie będące aktem

mowy obejmującym użycie imienne danego wyrażenia. Z tym z kolei wiąże się inne zdarzenie-akt mowy. O każdym nadawcy wypowiedzi w takim łańcuchu możemy powiedzieć, że w jego indywidualnej historii istnieje akt mowy, w którym miał on okazję zapoznać się z daną nazwą. W całej zaś klasie takich inicjujących aktów mowy istnieje jeden lub kilka aktów, które możemy utożsamić z chrztem pierwotnym. Z jednej strony istnienie odpowiednich aktów inicjujących jest tu warunkiem koniecznym późniejszych sensownych użyć danej nazwy, z drugiej zaś strony pewną rolę gra tu przyczynowość, która związana jest z faktem nabycia umiejętności posługiwania się daną nazwą w danej społeczności językowej. Zaznaczmy przy tym, że nie jest to rola z gatunku wielkich — nazwa „teorii przyczynowej” lub „teorii łańcucha przyczynowego” jest tu nadana koncepcji okazjonalnej nieco na wyrost.

Z teorią okazjonalną związane są także pewne problemy innego rodzaju. Pierwszym, który się narzuca, a który być może można uznać za zaletę tej teorii, a nie wadę, jest to, że teoria okazjonalna w dość nieprecyzyjny sposób opisuje czynniki mające wpływ na fakt, iż w wypadku konkretnej wypowiedzi moc ma ten, a nie inny akt mowy będący ekwiwalentem chrztu pierwotnego. Istnieje duża asymetria między tym wypadkiem a wypadkami wąsko pojętych wyrażań okazjonalnych, na przykład zaimków osobowych, w których wypadku możemy powiedzieć po prostu, iż zachodzi pewien prosty fakt na temat kontekstu wypowiedzi — na przykład taki, że Katarzyna Kuś jest nadawcą, a Tadeusz Ciecierski odbiorcą wypowiedzi. W tym aspekcie teoria okazjonalna pozostaje w wyraźnym kontraście do omawianej przez nas poniżej *hybrydowej teorii łańcucha przyczynowego*, w której *explicite* podejmuje się kwestie determinant intencji referencyjnych.

Inne, poważniejsze chyba trudności, związane są z postulowaną analogią między nazwami własnymi a nazwami okazjonalnymi. Po pierwsze, warto zwrócić uwagę na to, że w koncepcji nazw własnych jako wyrażań okazjonalnych to, co w teoriach przyczynowych uchodziłoby za determinantę metasemantyczną, staje się tu po prostu determinantą semantyczną. Przypomnijmy, że pytanie metasemantyczne „Na mocy jakich faktów wyrażenie ma wartość semantyczną (lub znacznie) taką, jaką ma?” jest tu rozumiane w sensie przyczynowym lub genetycznym. W takim rozumieniu czynniki kontekstowe, które są w pewnym sensie determinantami odniesienia nazw okazjonalnych w określonych kontekstach, nie zasługują w żadnej mierze na miano „czynników metasemantycznych”. Są one raczej czynnikami semantycznymi, na których operują związane z danym wyrażeniem okazjonalnym odpowiednie reguły językowe (na przykład reguła stwierdzająca, że odniesieniem wyrażenia „ja” jest nadawca wypowiedzi). Ponieważ nazwy własne uzyskują na gruncie omawianej koncepcji interpretację okazjonalną, to dana nazwa własna ma takie, a nie inne odniesienie na mocy faktu semantycznego będącego wynikiem zastosowania odpo-

wiedniej reguły językowej (wyrażenie NN odnosi się w kontekście C do przedmiotu wskazanego przez ekwiwalent aktu chrztu pierwotnego, który ma moc w kontekście C). Możemy w pewnym sensie powiedzieć, że teoria okazjonalna w ogóle nie udziela odpowiedzi na nasze oryginalne pytanie, a jej zwolennik pomieszał dwa różne jego znaczenia — semantyczne i metasemantyczne, za czym notabene kryje się słuszny przytyk pod adresem wieloznaczności pytania metasemantycznego.

Dodatkowo teoria okazjonalna, traktując nazwy własne jak nazwy okazjonalne, powinna wyjaśniać następujący brak analogii między wypadkami wąsko pojętych nazw okazjonalnych oraz wypadkami nazw własnych. Gdy do czynienia mamy ze zwykłymi wyrażeniami okazjonalnymi, z jednym wyrażeniem-typem związana jest jedna reguła językowa, która determinuje wartość semantyczną tego wyrażenia w każdym konkretnym kontekście. Nie zdarza się właściwie w językach naturalnych (z nielicznymi wyjątkami, por. polskie „dziś” i „dzisiaj”), by z dwoma wyrażeniami okazjonalnymi-typami związana była jedna reguła językowa wiążąca ich wartości semantyczne z kontekstem. W wypadku nazw własnych jest zupełnie inaczej — tutaj z *każdym* różnokształtnym wyrażeniem użytym imiennie związana jest jedna i ta sama reguła językowa, zgodnie z którą odniesieniem jest przedmiot wskazany przez akt mowy będący ekwiwalentem chrztu pierwotnego, a który ma moc w tym kontekście. *Teoria okazjonalna* w żadnym razie nie wyjaśnia tej asymetrii między zwykłymi wyrażeniami okazjonalnymi a wyrażeniami w użyciu imiennym. Z drugiej strony trudność tę można być może obrócić na jej korzyść, jeśli zwróci się uwagę na fakt, iż bardzo dobrze pasuje ona do tezy o etykietowym charakterze nazw, która jest podstawą *teorii bezpośredniego odniesienia*. Gdyby podążyć tym tokiem rozumowania, konsekwentnie trzeba by było powiedzieć, że wszystkie nazwy własne-egzemplarze podpadają w gruncie rzeczy pod jedną nazwę własną-typ (kształt etykiety lub miejsce, w którym ją umieściliśmy, nie mają znaczenia).

b) Hybrydowa teoria łańcucha przyczynowego

Mianem „hybrydowej teorii nazw własnych” określa się w literaturze przedmiotu koncepcje, które inspirowane są pracami Garetha Evansa⁷. Rozwiązania te — jak wskazuje ich nazwa — łączą ze sobą elementy *teorii łańcucha przyczynowego* oraz tezę o ważnej roli deskryptywnych i intencjonalnych warunków w wyznaczaniu odniesienia wyrażenia użytego imiennie. Poniżej przedstawimy w zarysie teorię najbliższą oryginalnym pomysłom Evansa.

Jedno z najważniejszych założeń tej koncepcji zawarte jest w następującym cytacie:

⁷ (Evans 1993), a także (Evans [1982] 2009); omówienie teorii Evansa znaleźć może Czytelnik także w: (Kawczyński 2010).

Jest coś absurdalnego w przypuszczeniu, że zamierzonym przedmiotem odniesienia jakiegoś całkowicie potocznego użycia nazwy przez mówiącego mógłby być przedmiot kompletnie izolowany (przyczynowo) od społeczności i kultury użytkownika, pełniący tę rolę na mocy faktu, że odpowiada lepiej niż coś innego wiążące deskrypcji, którą mówiący wiąże z nazwą. Zgodziłbym się z Kripkem, że absurdalność tkwi w braku relacji przyczynowej między wchodzącym w grę przedmiotem a mówiącym. Wydaje się jednak, że błędnie umiejscowił on tę relację przyczynową; ważna relacja przyczynowa zachodzi między stanami i działaniami tego przedmiotu a zasobem informacji mówiącego, nie zaś między nadawaniem nazwy temu przedmiotowi a dzisiejszym jej użyciem przez mówiącego (Evans 1993: 235).

Jak wyjaśnia dalej Evans, przez „zasób informacji mówiącego” rozumie się tu informację związaną przez mówiącego z określoną nazwą. Evans rozróżnia wyraźnie w swojej koncepcji dwie kwestie: po pierwsze, zagadnienie (przyczynowego) związku między przedmiotem (potencjalnym odniesieniem) a pewnymi zbiorami informacji, które podmiot kojarzy z nazwą (za Evansem możemy nazwać takie zbiory „dossier informacji”), a po drugie, problem czynników, które determinować mogą intencję odnoszenia się (przy użyciu określonego wyrażenia) do pewnego przedmiotu. Rozważając pierwszą z tych kwestii, wprowadza pojęcia dominującego źródła informacji oraz zapożycza od Kaplana pojęcie *bycia o czymś* (Kaplan 1996: 357–361). Może się na przykład zdarzyć, że w przekazach historycznych przypisuje się określony zbiór czynów pewnej osobie, chociaż faktycznie za czynami tymi kryje się więcej osób. Historycy uważają na przykład, iż *Iliada* i *Odyseja* są dziełami dwóch różnych autorów, ale prawdopodobnie dominującym źródłem informacji kojarzonych z nazwą „Homer” jest twórca pierwszego z tych eposów⁸. Jest tak prawdopodobnie dlatego, że *Iliada* uchodzi za utwór ważniejszy oraz lepiej znany. Z kolei *bycie o czymś* jest własnością zbiorów (dossier) informacji — dany zbiór informacji jest o przedmiocie, który jest w danym wypadku dominującym źródłem informacji⁹. Pojęcie to zakłada istnienie pewnego związku przyczynowego między przedmiotem a faktem pojawienia się określonej informacji. Przedmiot wraz ze swoimi właściwościami i zachowaniami uczestniczy w istotnym stopniu w genezie danego zasobu informacji. Kaplan w swojej oryginalnej koncepcji odwoływał się tu do analogii z obrazami i zdjęciami — pewne zdjęcie lub obraz (w naszym wypadku pewna informacja) jest o tym raczej niż o innym przedmiocie, ponieważ przedmiot ten odgrywał istotną przyczynową rolę w fakcie powstania obrazu lub fotografii (w naszym wypadku zbioru informacji). Evans zapożycza tę ideę, lecz modyfikuje ją odrobinę, stosując pojęcie *bycia o* do zbiorów

⁸ Choć między autorami tego tekstu istnieje rozbieżność zdań w tej sprawie, jeden z uczestników sporu — ceniący bardziej *Odyseję* — okazał się słabszy psychicznie i ustąpił.

⁹ Nieco mylący może się tu okazać polski przekład *The causal theory of names*. Evans, objaśniając tę kwestię, pisze jednak: „Consequently a cluster of dossier of information can be dominantly of an item though it contains elements whose source is different”. A więc wyraźnie podkreśla, że tylko niektóre elementy dossier mogą mieć źródło inne niż przedmiot, o którym jest to dossier.

informacji, a nie tylko do pojedynczych informacji. Kaplan prawdopodobnie powiedziałby, że opisany powyżej wypadek Homera to sytuacja, w której opisane przez nas dossier informacji jest faktycznie o dwóch osobach jednocześnie, podczas gdy Evans dzięki dopuszczeniu możliwości hierarchizacji i liczenia informacji w zbiorze może powiedzieć, że dossier to jest o jednej osobie — o tej, która jest dominującym źródłem informacji.

Co zaś tyczy się kwestii determinant intencji odniesienia się do danego przedmiotu, zgodnie z hipotezą Evansa, w większości wypadków — choć podkreśla on, że nie zawsze — podmiot, używając imiennie danego wyrażenia, ma zamiar odnieść się do przedmiotu, który jest dominującym źródłem informacji kojarzonych z tym wyrażeniem. Pomysł ten zostaje wykorzystany przez Evansa w roboczej definicji użycia imiennego pewnego wyrażenia w pewnej społeczności w odniesieniu do pewnego konkretnego przedmiotu. Proponuje on, by uznać, że pewne wyrażenie występuje w pewnej społeczności w użyciu imiennym, po pierwsze, gdy istnieje będąca przedmiotem wiedzy podzielanej przez członków tej społeczności praktyka używania imiennie tego wyrażenia w odniesieniu do danego przedmiotu. Po drugie zaś, gdy powodzenie w odniesieniu się do danego przedmiotu za pomocą tej nazwy zależy od tego, czy uczestniczący w rozmowie dzielą wiedzę, że w ich społeczności używa się tej nazwy z zamiarem odniesienia się do tego, a nie innego przedmiotu. Co istotne, nie jest to wspólna wiedza na temat własności jednoznacznie identyfikujących pewien przedmiot¹⁰.

Evansowska hipoteza, zgodnie z którą zamierzonym przedmiotem odniesienia jest dominujące źródło informacji, nie obowiązuje dla sytuacji określanych mianem „uległego użycia wyrażenia imiennego”. W takich wypadkach używa się danego wyrażenia z intencją odniesienia się do tego samego przedmiotu, do którego odnieśli się przedmówca lub grupa przedmówców. Jeśli mamy do czynienia z wypadkiem uległego użycia wyrażenia, to intencja ta jest nadrzędna względem intencji odniesienia się do dominującego źródła informacji. Nieuległe użycie wyrażenia imiennego ma natomiast miejsce, gdy intencją mówiącego jest odniesienie się do dominującego źródła informacji. Przedstawiona wyżej w zarysie *hybrydowa teoria łańcucha przyczynowego* cechuje się dużym stopniem wyrafinowania, między innymi z tej przyczyny, że zagadnienia referencji łączy z problematyką z zakresu teorii działania i filozofii umysłu. Jest to jednak cecha, której trzeba najprawdopodobniej wymagać od adekwatnej teorii nazw własnych¹¹.

¹⁰ Evans mówi w tym wypadku o opozycji między a) używaniem pewnego wyrażenia z intencją odniesienia się do pewnego przedmiotu, ponieważ wiemy, że używamy tego wyrażenia z intencją odniesienia się do niego, a b) używaniem danego wyrażenia z intencją odniesienia się do danego przedmiotu, ponieważ na przykład wiemy, iż przedmiot ten posiada pewne cechy (zob. Evans 1993: 239).

¹¹ Swoją wersję *Teorii łańcucha przyczynowego* rozwijał Evans w pracach, które złożyły się na wydaną pośmiertnie książkę *The Varieties of Reference*. W pracy tej między innymi

c) Dwie teorie łańcucha przyczynowego: okazjonalna versus hybrydowa

Jak zwracaliśmy już uwagę, teoria okazjonalna jest teorią bardziej ogólną i mniej precyzyjną od teorii hybrydowej. W szczególności ta pierwsza nie mówi nam wiele o faktach, które decydują o tym, że w określonym kontekście obowiązuje ten raczej, a nie inny akt mowy będący ekwiwalentem chrztu pierwotnego. Ta ogólność jest przyczyną, dla której ktoś mógłby słusznie zauważyć, że można prawdopodobnie pogodzić ze sobą obie teorie i zaproponować okazjonalny wariant teorii hybrydowej. W największym zarysie teoria taka wyglądałaby w sposób następujący. Twierdzilibyśmy, że na to, który akt mowy będący ekwiwalentem aktu chrztu pierwotnego ma moc w danym kontekście, wpływ mają intencje odniesienia się do tego, a nie innego przedmiotu. Intencje te z kolei otrzymywałyby analizę analogiczną do omówionej wyżej analizy Evansowskiej. Taka hybryda *teorii okazjonalnej* i *teorii hybrydowej* warta jest zapewne bliższego rozważenia, które jednak na mocy aktu decyzji autorów niniejszego artykułu wypada poza jego ramy.

III. Przykład Sticha a teoria łańcucha przyczynowego

Czy — tak, jak chciałby Stich — jego przykład stanowi problem dla *teorii łańcucha przyczynowego*? Po pierwsze, należy odpowiedzieć tu na pytanie, czy rację ma Stich, że nasze intuicje związane ze sformułowanymi przez policję w toku śledztwa wypowiedziami zawierającymi nazwę „Borys Nogoodnik” przemawiają za tym, że odpowiednie wypowiedzi nie są o Borysie A. Nogoodniku. To znaczyć może jedną z dwóch rzeczy — albo że są o Borysie B. Nogoodniku (lub jakiejś innej niż Borys A. Nogoodnik osobie), albo że w ogóle nie potrafimy stwierdzić, o jakim są przedmiocie. Oba wypadki są — przynajmniej na pierwszy rzut oka — wyzwaniem dla zwolennika *teorii łańcucha przyczynowego*. Jeśli nasze intuicje nie wskazują na jeden z tych wypadków, problem dla *teorii łańcucha przyczynowego* w ogóle nie powstaje. Drugą kwestią, którą należy tu rozważyć, jest pytanie, czy dana wersja teorii przyczynowej jest w stanie wytłumaczyć nasze intuicyjne sądy wskazujące, że odniesieniem nie jest Borys A. Nogoodnik. Przykład Sticha (a właściwie pewien konkretny jego wariant) jest zatem problemem dla danej wersji *teorii łańcucha przyczynowego*, o ile pewnym intuicjom semantycznym (dotyczącym odniesienia danej nazwy) towarzyszy brak wyjaśnienia teoretycznego, który pozwoliłby uzgodnić te intuicje z konkretną wersją *teorii łańcucha przyczynowego*.

wprowadzone zostają pojęcia *praktyki językowej* i *przynależności do praktyki językowej*. Wedle Evansa warunkiem skutecznego odniesienia się przy użyciu wieloznacznej nazwy własnej jest to, iż użytkownik języka jest w stanie, choć nie musi faktycznie tego robić, posługując się nieuległe tą nazwą, wskazać, do której z alternatywnych praktyk językowych związanych z nią przynależy. Pytanie: „Czy policjant wypowiadający zdanie jest w stanie zaimplementować, do której praktyki używania imienia Borys Nogoodnik należy?”, odpowiedź: „Nie”.

a) Interpretacje, w których informacje przypadkowo jednoznacznie identyfikują Borysa B.

Jak pisaliśmy wyżej, może się zdarzyć, że dowody odnalezione na miejscu zbrodni lub zgromadzone w toku śledztwa identyfikują w sposób przypadkowy jakąś bliżej nieznaną zabójcy osobę, która może, choć nie musi, nazywać się „Borys Nogoodnik”. Najbardziej jasny jest wypadek, w którym mamy do czynienia z jednoznaczną identyfikacją — tylko jeden przedmiot we wszechświecie posiada właściwości, na które wskazują dowody (pali Popularne, ma numer buta 45, był w danym dniu na meczu hokeja na trawie). Czy w takim wypadku, podkreślmy raz jeszcze — najbardziej wyrazistym z klasy wszystkich, w których identyfikacja przedmiotu ma charakter przypadkowy — prawdą jest to, co twierdzi Stich, a mianowicie, że nazwa „Borys Nogoodnik” użyta przez policjantów w śledztwie z całą pewnością nie odnosi się do Borysa A. Nogoodnika? Wydaje się, iż tak nie jest. Aby pokazać, do jak absurdalnej sytuacji prowadziła-by odpowiedź twierdząca na to pytanie, rozważmy następujący przykład. Wyobraźmy sobie, że Borys A. Nogoodnik pozostawił na miejscu zbrodni szereg śladów, które wskazują na jego osobę. Zarazem jednak do śladów tych doszedł jeden, niezwiązany z nimi w żaden sposób, lecz na przykład pozostawiony przez przechodzącą przez miejsce zbrodni na godzinę przed zabójstwem Monikę Bellucci, która zatarła też przypadkowo jeden ze śladów zostawionych przez Borysa A. Nogoodnika. Okazuje się następnie, że wszystkie te ślady zebrane razem nie wskazują ani Moniki Bellucci, ani Borysa Nogoodnika, tylko jeszcze jakąś inną osobę — na przykład najstarszą mieszkankę wsi Baniocha pod Warszawą. W takiej sytuacji, zgodnie z intuicją Sticha, policjant, używając nazwy „Borys Nogoodnik”, odnosi się do tejże sędziwej mieszkanki Baniochy. Jest to sytuacja dziwna, w której musielibyśmy założyć, że odnosimy się przy użyciu nazwy własnej do jakiegoś zupełnie odseparowanego przyczynowo od użytkowników języka przedmiotu. Samo w sobie nie musi to dyskwalifikować wniosku Sticha, ale prawdziwie dyskwalifikujące zdaje się być tu rozważenie pytania, co powiedziałby policjant, gdyby poznał prawdę na temat całej historii. Nic nie przemawia na rzecz wniosku, że uznałby nazwę „Borys Nogoodnik” za niezmiennie odnoszącą się do wspomnianej mieszkanki wsi Baniocha. Twierdziłby raczej, że w przeszłości mylił się, co do odniesienia nazwy. Ponieważ wszystkie pozostałe wypadki, w których możemy mówić o identyfikacji przez przypadek, są mniej wyraziste od tego, który tutaj rozważamy, we wszystkich sytuacjach będzie wyglądała identycznie: nic nie wspiera w nich intuicji, że odnosimy się do osoby różnej od Borysa A. Nogoodnika. Analiza tej klasy wypadków nie wymaga zatem rozpatrywania ich stosunku do okazjonalnej lub hybrydowej koncepcji łańcucha przyczynowego¹².

¹² Jak się wydaje, przyczyną, dla której Stich ocenił, że nasze intuicje w takich wypadkach wskazują na Borysa B. Nogoodnika jako odniesienie nazwy „Borys Nogoodnik”, jest to,

b) Interpretacje, w których informacje nieprzypadkowo i niejednoznacznie identyfikują Borysa B.

Z kolei zastanówmy się nad tymi wypadkami, w których identyfikacja nie ma charakteru przypadkowego, lecz jest zarazem niejednoznaczna. W takich wypadkach zebrane dowody wskazują na jakąś mniej lub bardziej liczną klasę osób, wśród których, jak zakładamy, znalazła się osoba, którą chciał wskazać zabójca (ponownie możemy zauważyć, że osoba ta może, ale nie musi, nazywać się „Borys B. Nogoodnik”). Rozróżniliśmy w obrębie takich sytuacji te, w których pośród osób niejednoznacznie wskazanych przez dowody znalazł się Borys A. Nogoodnik, oraz takie, w których nie ma go w klasie osób podejrzanych. Zapytajmy ponownie, czy rację ma tu Stich, że odpowiednia wypowiedź policji nie jest o Borysie A. Nogoodniku lub jest o osobie od niego różnej. Wydaje się, iż wbrew zapewnieniom Sticha nie zachodzi żadna z tych sytuacji. Z dwóch wyróżnionych wyżej wypadków zupełnie oczywisty jest ten, w którym zakładamy, że w klasie wskazywanych osób jest Borys A. W takiej sytuacji nasze przekonanie, iż policja odnosi się w rzeczywistości do Borysa A., ulega zrozumiałemu wzmocnieniu. Nie dość bowiem, że policjanci weszli w posiadanie nazwy własnej z intencją odniesienia się do Borysa A., to także zebrane przez nich informacje identyfikujące wyróżniają klasę osób, wśród których jest nasz zabójca.

W drugim wypadku, gdy niejednoznacznie oraz w sposób nieprzypadkowy ślady wskazują na pewną klasę osób, wśród których nie ma Borysa A., nasza ocena sytuacji nie jest w żadnym razie jednoznaczna. Można twierdzić, że tak naprawdę policjant nadal odnosi się do Borysa A., można też zajmować stanowisko, zgodnie z którym policjant po analizie dowodów, wypowiadając zdanie „Borys Nogoodnik jest zabójcą”, nie odnosi się do Borysa A., ponieważ nie odnosi się do żadnej określonej osoby (powiedzielibyśmy wówczas prawdopodobnie, że jego wypowiedź nie posiada wartości logicznej). Ponieważ nasze intuicje w żadnym razie nie opowiadają się za którymkolwiek z tych wypadków, powinniśmy rozważyć je oba i zastanowić się, czy i w jakiej mierze stanowią one kłopot dla konkretnych wariantów teorii przyczynowej.

Jeśli twierdzimy, że w opisanej sytuacji policjant mówi coś o Borysie A. Nogoodniku, to oczywiście, tak jak wcześniej, intuicja Sticha zostaje oddalona. Jednak powstaje tu pewna subtelna trudność, z którą musi się zmierzyć zwolennik *hybrydowej teorii łańcucha przyczynowego*. Przypomnijmy, iż zgodnie z jednym z jej założeń użytkownik nazwy zazwyczaj ma zamiar odnieść się do związanego z nią przedmiotu będącego dominującym źródłem informacji. Jeśli

że zamiast rozważać osobę o dowolnym imieniu i nazwisku, wziął on pod uwagę osobę, która przypadkowo nazywa się „Borys Nogoodnik”, tym samym przeniósł on prawdopodobnie naszą własną, konstruktorów przykładu, wiedzę na temat możliwych odniesień nazwy „Borys Nogoodnik” na jej właściwości, gdy użyta jest ona przez policjanta.

dowody zostały, tak jak to założyliśmy, sporządzone umyślnie, to nie możemy wykluczyć możliwości, że w klasie potencjalnych odniesień danej nazwy jest też jakaś osoba, która jest dominującym źródłem informacji związanym z tą nazwą, a nie jest Borysem A. Nogoodnikiem. Zwolennik teorii hybrydowej musi wyjaśnić w jakiś sposób to, że w takim wypadku odnosimy się nadal do Borysa A. Nogoodnika, a nie do domniemanego faktycznego dominującego źródła informacji. Prawdopodobnie odpowiedź zwolennika teorii hybrydowej powinna tu przybrać postać następującą: w posiadanym przez policjanta dossier informacji związanych z nazwą „Borys Nogoodnik” musi znajdować się informacja, że osoba ta jest zabójcą naszej ofiary. Jeśli założymy, że informacja taka faktycznie znajduje się w dossier i żadne inne informacje nie są od niej istotniejsze (bo nie wskazują na nikogo jednoznacznie), to oczywiście policjant odnosi się do Borysa A. To on jest bowiem dominującym źródłem informacji.

Teoria hybrydowa pozwala również na wyjaśnienie twierdzenia, że Stich ma rację i policjant nie odnosi się do Borysa A. W takiej sytuacji musimy uznać, iż faktyczny zabójca przestaje być dominującym źródłem informacji związanych z nazwą „Borys Nogoodnik”. Zgodnie bowiem z tą teorią nazwa niezwiązana z żadnym jednym dominującym źródłem informacji nie odnosi się do niczego¹³. Z kolei teoria okazjonalna również dysponuje w tym wypadku bardzo eleganckim wyjaśnieniem całej sytuacji. Wystarczy, że jej zwolennik zauważy, iż w zaistniałej sytuacji kontekst nie dostarcza nam żadnej informacji na temat tego, który akt z aktów chrztu pierwotnego ma moc w tym określonym kontekście. Jeśli wśród osób wskazanych przez dowody nikt nie nazywa się „Borys Nogoodnik”, to w pewnym sensie po prostu nie ma takiego aktu, jeśli natomiast pewna osoba się tak nazywa, to nic w okolicznościach wypowiedzi nie wyróżnia aktu chrztu lub jego ekwiwalentu związanego z tą osobą. Aby posłużyć się analogią, pierwszy wypadek przypominałby sytuację, w której ktoś w pewnych okolicznościach używałby wyrażenia wskazującego, mając halucynacje, w związku z tym nie istnieje obiekt wskazywany, a drugi sytuację, w której ktoś wskazuje jakiś przedmiot w sposób nieostry lub niejednoznaczny (czy to królik czy noga królika?).

c) Interpretacje, w których informacje nieprzypadkowo i jednoznacznie identyfikują Borysa B.

Ostatni wypadek, który musimy rozważyć, to sytuacja prawdopodobnie najbardziej klarowna i jednoznaczna, a zatem wypadek, w którym przebiegły Borys A. Nogoodnik tak spreparował dowody, że jednoznacznie wskazują na Borysa B.

¹³ Por. przykład z *Przyczynowej teorii nazw* dwóch dziewcząt nazywanych w pewnej społeczności imieniem „Goldilocks”. Evans twierdzi tu, że jeżeli członkowie społeczności nie zdają sobie sprawy, że w rzeczywistości mają do czynienia z dwiema osobami, to nie jest tak, że nazwa „Goldilocks” odnosi się do obu dziewcząt, lecz że w ogóle nie ma odniesienia.

Nogoodnika. Czy mamy w takiej sytuacji rzeczywiście intuicje, że policjant po przeprowadzeniu śledztwa, używając nazwy „Borys Nogoodnik”, może od pewnego momentu mówić coś o Borysie B. Nogoodniku? Wydaje się, że Stich ma tutaj słuszność i w pewnych sytuacjach faktycznie bylibyśmy skłonni uznać odnośne wypowiedzi za będące o osobie, na którą wskazują dowody. Jeśli Stich ma tu rację, jak dwie opisane wyżej teorie łańcucha przyczynowego mogą sobie poradzić z tym faktem?

Zacznijmy od *hybrydowej teorii łańcucha przyczynowego*. Bez wątplenia, choćby przez sam fakt, że dowody zostały świadomie spreparowane przez Borysa A. Nogoodnika w taki sposób, aby wskazywały na Borysa B. Nogoodnika, powiemy po prostu, iż Borys B. może stać się w pewnym momencie dominującym źródłem informacji związanych z nazwą „Borys Nogoodnik”. Dzieje się tak, ponieważ powstanie dowodów oraz ich charakter są ściśle związane z intencjami i przekonaniami Borysa A., które to z kolei są w szerokim sensie tego słowa informacyjnie i przyczynowo związane z cechami i zachowaniami Borysa B. Nogoodnika. W takiej sytuacji oczywiście użytkownik lub grupa użytkowników nazwy „Borys Nogoodnik” będą używali jej z zamiarem odniesienia się do dominującego źródła informacji, którym w tym wypadku jest Borys B. Nogoodnik. Jeśli mamy tu rację, to przykład Sticha nie jest problemem dla *hybrydowej teorii łańcucha przyczynowego*. Nie da się oczywiście wykluczyć, że w intencji autora tego przykładu nie był on pomyślany jako źródło kłopotu dla tej koncepcji, nie możemy tu jednak przy braku odpowiednich mocy poznawczych wczuwać się w jego intencje.

Co natomiast z wyjaśnieniem, które może zaoferować nam teoria okazjonalna? Wydaje się, że przykład Sticha rzeczywiście stanowi dla niej bardzo poważne wyzwanie. Najpoważniejszym kłopotem jest dla niej to, iż — jak podkreślaliśmy to wcześniej kilkakrotnie — osoba, na którą wskazują dowody, w żadnym momencie swojego istnienia nie musi być nazywana Borysem Nogoodnikiem. W takiej sytuacji nie istnieje ani żaden oficjalny akt chrztu pierwotnego, ani żaden jego ekwiwalent, który może mieć moc w kontekście wypowiedzi policji zawierającej na przykład zdanie: „Borys Nogoodnik jest kibicem hokeja na trawie”. Nie ma zatem tutaj żadnego parametru kontekstowego, który mógłby sprawiać, iż nasze wypowiedzi są o osobie, na którą wskazują ślady. Jedyny zarys odpowiedzi na ten problem kryłby się w powiedzeniu, że podczas śledztwa zaczyna ucierać się nowa praktyka używania imienia „Borys Nogoodnik”, która w pewnym momencie składa się na ponowny akt chrztu wiążący rzeczoną nazwę z osobą wskazywaną przez ślady. Jednakże nic nie zmienia faktu, iż w swojej ogólnej wersji teoria okazjonalna nie tylko nic nie wspomina o takich możliwościach, lecz także nie wiadomo nawet, jak i w jakich kategoriach opisywałyby narodziny takiej nowej praktyki (w tym miejscu teoria okazjonalna pozostaje w znaczącym kontraście do Evansowskiej teorii hybrydowej, zwłaszcza jej wersji przedstawionej w *Varieties of Reference*).

IV. Morały filozoficzne

Argument Sticha, jak staraliśmy się pokazać wyżej, można interpretować na kilka sposobów. Większość z nich może na pierwszy rzut oka wydawać się wyzwaniem dla teorii łańcucha przyczynowego w jej najogólniejszej, dość nieprecyzyjnej wersji. Jednakże bliższa analiza wariantów przykładu Sticha oraz bardziej precyzyjne przedstawienie *teorii łańcucha przyczynowego* sprawia, że nie sposób uważać go za problematyczny: albo intuicje, które mają być źródłem problemu, nie są takie, jak sugeruje Stich, albo choć są takie, bardziej precyzyjne wersje koncepcji łańcucha przyczynowego potrafią wyjaśnić fenomen zmiany odniesienia wyrażenia użytego imiennie.

W jednym wypadku wszakże okazuje się, że okazjonalna wersja łańcucha przyczynowego nie daje sobie dobrze rady z problemem. Jest on zatem argumentem krzyżowym, który wskazuje, iż hybrydowa teoria łańcucha przyczynowego w przeciwieństwie do teorii okazjonalnej stanowi lepszą realizację pomysłów analizy nazw własnych, które wywodzą się z prac takich autorów, jak Barcan, Kripke, Geach lub Donnellan. Omówiony przez nas problem można interpretować także jako wyzwanie dla teorii okazjonalnej, od której zwolenników wymaga powiedzenia czegoś więcej o praktyce imiennego użycia wyrażen.

Przy okazji jednak analiza argumentu Sticha wskazuje na szereg subtelnych i interesujących problemów mieszczących się na styku filozofii języka, umysłu i działania. Jako taki w opinii autorów zasługuje na uwagę filozoficzną nie mniejszą od głośniejszych przykładów dyskutowanych przez Kripkego, Donnellana oraz Evansa.

Bibliografia

- EVANS G. (1993), *Przyczynowa teoria nazw*, [w:] Stanosz B. (red.), *Fragmenty filozofii analitycznej*. Warszawa: SPACJA, 226–245.
- EVANS G. ([1982] 2009), *The Varieties of Reference*. Oxford: Clarendon Press.
- HESS L. (2009), *Nazwy własne — fakty i mity*, „Filozofia Nauki” 2 (66), 123–128.
- KAPLAN D. (1989), *Afterthoughts*, [w:] Almog J., Wettstein H., Perry J., (red.), *Themes from Kaplan*, Oxford: Oxford University Press, 568–614.
- KAPLAN D. (1996), *Quantifying in*, [w:] Martinich A. P. (red.), *The Philosophy of Language*. Oxford: Oxford University Press, 347–368.
- KAWCZYŃSKI F. (2010), *The Hybrid Theory of Reference for Proper Names*, [w:] Stalmaszczyk P. (red.), *Philosophy of Language and Linguistics*, Frankfurt am Main: Ontos Verlag, 137–150.
- KRIPKE S. (1996), *A Puzzle about Belief*, [w:] Martinich A. P. (red.), *The Philosophy of Language*. Oxford: Oxford University Press, 382–410.
- PELCZAR M., RAINSBURY J. (1998), *The Indexical Character of Names*, „Synthese” 114 (2), 293–319.

- RAYO A. (2010), *A Metasemantic Account of Vagueness*, [w:] Dietz R., Moruzzi S. (red.), *Cuts and Clouds. Vagueness, Its Nature and Its Logic*. Oxford: Oxford University Press, 23–45.
- STALNAKER R. C. (2003), *Reference and Necessity*, [w:] Stalnaker R., *Ways a World Might Be*. Oxford: Oxford University Press, 165–187.
- STICH S. (1983), *From Folk Psychology to Cognitive Science: The Case Against Belief*. Cambridge and London: MIT Press.

Rozdział 13

Filozoficzne podstawy teorii gier — najnowsze tendencje

Elżbieta Chrzanowska-Kluczevska (Uniwersytet Jagielloński)

Słowa kluczowe: gry agonalne, gry językowe, gry kooperacji, paradygmat strategiczny, semantyka teoriogrowa, A. K. Dixit, B. J. Nalebuff, J. Hintikka, L. Wittgenstein

1. Krótka historia teorii gier

Teoria gier stanowi stosunkowo młodą dyscypliną naukową, której początki sięgają lat trzydziestych ubiegłego wieku. Jest też w dużej mierze kreacją filozofów, których myśl biegła równolegle do myśli matematyków, co w rezultacie zaowocowało wypracowaniem formalnego modelu strategicznego, ukazującego sposób podejmowania decyzji racjonalnych. Dziś paradygmat teoriogrowy święci największe sukcesy w zastosowaniu do działań ekonomicznych, ale obecny jest również w szeroko pojętym modelu nauk społecznych. Warto jednak pamiętać, że rozważania na temat przydatności pojęcia *gry* w badaniach naukowych i praktyce społeczno-ekonomicznej rozpoczęły się w dużej mierze od rozważań nad miejscem gry i zabawy w języku naturalnym. Dlatego też poniższy artykuł odnosić się będzie nie tylko do podstawowych założeń teorii gier *in abstracto*, lecz również, a może szczególnie, do ich wagi w filozofii języka i w językoznawstwie.

Naukową analizę koncepcji gier otwiera we współczesnej humanistyce klasyczna pozycja Johana Huizingi, holenderskiego historyka i filozofa kultury, zatytułowana *Homo ludens. Zabawa jako źródło kultury* ([1938] 1967). Huizinga dostrzega w grach i zabawach społecznych konieczny element rozwoju cywilizacyjnego rodzaju ludzkiego. Podkreśla w nich szczególnie ważne w jego mniemaniu aspekty: 1) *ludyczny*, zwany też *eskapistycznym*, który daje graczowi radość ucieczki od rzeczywistości w autonomiczny, tymczasowy świat gry/zabawy oraz 2) *agonalny* (*agonistyczny*), a więc wymiar współzawodnictwa. Ze względu na element *konfliktu* obecny w tej drugiej kategorii gier niektórzy badacze upatrują

historycznej genezy teorii gier już u Arystotelesa, a dokładnie rzecz ujmując w *Polityce*. Prototypem gry agonicznej pozostają niezmiennie szachy, historycznie wywodzące się (podobnie jak warcaby) z repliki pola bitewnego.

W eseju Huizingi opis zjawisk typowo językowych obejmuje w pierwszej kolejności gry na poziomie mikrotekstu (frazę, zdanie), w szczególności figury stylistyczne (metafora personifikująca, alegoria), gry słowne (kalambury) oraz „igraszki etymologiczne” (Huizinga 1967: 286), które autor klasyfikuje jako „treści ludyczne” współczesnej mu nauki¹.

Jeśli chodzi natomiast o dłuższe formy językowe, a więc na poziomie funkcjonalnym tekstu/dyskursu, Huizinga omawia takie gatunki, jak poezja (szczególnie liryczna), dramat, mity, jak również turnieje zagadek i dysputy różnego rodzaju, a więc dyskurs filozoficzny, religijny i akademicki.

Na ogół jednak właściwy ingres terminu *gra* do współczesnego językoznawstwa kojarzymy z nazwiskiem Ludwiga Wittgensteina, który w podobnym okresie, lecz niezależnie od Huizingi, oddawał się rozważaniom na temat roli gier w języku naturalnym. Początki koncepcji *gier językowych* sięgają czasów *Niebieskiego i Brązowego zeszytu* (szkice do późniejszych *Dociekań filozoficznych* datujące się na lata 1933–1936), a jej dojrzała wersja znajduje się w *Dociekaniach* ([1953] 2000). Stąd pochodzi znana definicja:

„Grą językową” nazywać też będę całość złożoną z języka i czynności, w które jest on wpleciony (Wittgenstein 2000: 12).

Od razu więc zarysowuje się zgoła odmienna koncepcja gry językowej niż ta proponowana przez Huizingę. O ile dla holenderskiego badacza gry jako zjawisko są zawsze wyjątkowe, usytuowane we własnej przestrzeni, wyodrębnionej czasowo z rzeczywistości, o tyle pragmatycznie zorientowana teoria Wittgensteina nadaje grom charakter aktywności wszechobecnej w języku naturalnym, bynajmniej nie niezwyklej, a wprost przeciwnie, odgrywającej ważną rolę pomostu pomiędzy językiem a rzeczywistością. W zasadzie więc wszystkie koncepcje gier w naukach humanistycznych, społecznych i ekonomicznych będą mieścić się albo w paradygmacie Huizingi², albo też w paradygmacie Wittgensteina.

Z funkcjonalnego punktu widzenia, gry cytowane przez Wittgensteina mają miejsce na mikropoziomie zdania (np. „rozkazywanie”, „prośenie”, „dziękowanie”, „przeklinanie” itd., por. Wittgenstein 2000: 21) lub też na makropoziomie struktur tekstowych („modlenie się”, „opisywanie przedmiotu”, „zda-

¹ Co ciekawe, Huizinga (tamże) traktuje również nazbyt abstrakcyjne dywagacje współczesnych mu syntaktyków jako zajęcia bardziej ludyczne niż praktyczne. Zresztą według niego „cała działalność naukowa zostaje wciągnięta na tory ludyzmu przez dążenie do rywalizacji”.

² Krytykę Wittgensteinowskiej koncepcji gry jako wszechobecnej w języku przeprowadził M. Black (1986, przywoływany również w Marion 2006: 261–262). W podobnym duchu wypowiadam się w mojej monografii poświęconej grom językowym (Chrzanowska-Kluczeńska 2004). Natomiast A. K. Dixit i B. J. Nalebuff ([2008] 2010) traktują całe życie społeczne człowieka jako skomplikowaną sieć powiązanych ze sobą gier.

wanie sprawy z zajścia”, „wysuwanie hipotez i ich sprawdzanie”, „wymyślanie historyjki”, „odgrywanie czegoś jak w teatrze”, „opowiadanie”, „opowiadanie dowcipu”, a nawet typowa gra tłumaczy, czyli „przekładanie z jednego języka na inny”). Teorię gier językowych rozwinęli następnie, jako *teorię aktów mowy*, John L. Austin na gruncie brytyjskim (1962) oraz John R. Searle na gruncie amerykańskim (1969), aplikując ją w szczególności do użycia zdań ze szczególną intencją i w szczególnym kontekście społecznym (*mikroakty mowy*). Z czasem dynamiczny rozwój pragmatyki językoznawczej doprowadził do rozszerzenia pojęcia *akt mowy* na struktury tekstowe (*makroakty mowy*, realizowane albo globalnie, albo pod postacią łańcuchów mikroaktów).

Cofając się nieco w czasie, musimy powrócić do okresu II wojny światowej, kiedy to badania nad koncepcją *gry* w obrębie nauk ścisłych, zintensyfikowane naoczną egzemplifikacją konfliktu społecznego w postaci działań militarnych, doczekały się swego zwieńczenia w klasycznej pozycji autorstwa matematyka Johna von Neumanna i ekonomisty Oscara Morgensterna *Theory of Games and Economic Behavior* (1944, 1947). Od niej rozpoczyna się, trwający nieprzerwanie do dziś, okres bardzo intensywnej aplikacji matematycznej teorii gier do ekonomii, jak również sprzężenia zwrotnego między tymi naukami, w wyniku czego teoriogrowa idealizacja matematyków stopniowo podlega modyfikacjom w kierunku przystosowania jej do zachowań ludzkich, często zaskakujących i nieprzewidywalnych w praktyce.

Następnym etapem rozwoju tej teorii stały się badania matematyka Johna Forbesa Nasha, Jr. (przeprowadzone w latach 50., a nagrodzone w roku 1994 pierwszą Nagrodą Nobla honorującą wykorzystanie matematycznej teorii gier w ekonomii). Wypracowane wówczas pojęcie *równowagi Nasha*, jako wypadkowej optymalnych strategii uczestników w danej rozgrywce, używane jest w badaniach struktur wszystkich gier, a szczególnie gier symultanicznych, o większej ilości graczy i strategii niż w klasycznej grze dwuosobowej o sumie zerowej i ruchach naprzemiennych (reprezentowanej typowo przez szachy).

Kolejnym klasykiem w rozwoju tej teorii stała się pozycja *Games and Decisions* R. Duncana Luce i Howarda Raiffy ([1958]1964), która nie tylko podkreśla rolę *strategii* jako kluczowego narzędzia metodologicznego w badaniu gier od czasów von Neumanna, ale również wykazuje przynależność teorii gier do struktury nadrzędnej, zwanej *teorią decyzji*. Ta ostatnia ma swój głęboki wymiar filozoficzny, stanowi bowiem część składową *teorii racjonalnego zachowania*³. Avinash K. Dixit i Barry J. Nalebuff to z kolei autorzy najbardziej chyba

³ Zapewne z tej przyczyny badacze historii teorii gier wymieniają niekiedy nazwisko Barucha Spinozy, który jako racjonalista wypowiadał się na temat rozumności czynów ludzkich, czyli ich zgodności z naturą ludzką i słuszności w danej sytuacji (por. Tatarkiewicz [1931] 1947 II: 100). Ograniczone ramy moich rozważań nie pozwalają mi na dalsze odwołania się do innych teorii filozoficznych zajmujących się rozumnym zachowaniem istot ludzkich. Warto natomiast uświadomić sobie, że to właśnie działania irracjonalne człowieka są największym wyzwaniem dla współczesnej teorii gier czy teorii AI (Artificial Intelligence).

poczytnej pozycji przybliżającej rolę strategii w działaniach ekonomicznych, *Thinking Strategically* (1991). Kontynuacją ich zabiegów popularyzujących teorię gier wśród czytelników, nawet tych, którzy nie posiadają głębszych podstaw matematycznych, jest ich kolejna książka — *The Art of Strategy: A Game Theorist's Guide to Success in Business and Life* ([2008] 2010), która powinna wzbudzić także zainteresowania językoznawców i uświadomić im, że koncepcja *strategii*, bardzo mocno obecnie zdywersyfikowana w ekonomii, to znakomity kandydat na *kategorię integracyjną*, która śmiało może stać się pomostem pomiędzy naukami ścisłymi i społecznymi⁴.

Powracając do pierwszego nurtu naszych rozważań, czyli do rozwoju teorii gier na polu nauk humanistycznych i społecznych, należy odwołać się do ważnego studium socjologicznego, rozwijającego pomysły Huizingi na drodze konstruktywnej polemiki, autorstwa Rogera Caillois, pt. *Les jeux et les hommes* ([1958] 1997). Wprowadziło ono stosowaną do dziś taksonomię gier i zabaw społecznych obejmującą cztery kategorie: 1) *gry typu agon* — oparte na rywalizacji pomiędzy graczami, która może przybierać różne formy na skali od ostrego konfliktu interesów do w miarę pokojowego współzawodnictwa, 2) *gry aleatoryczne (Glücksspiele)* — odwołujące się do czynników losowych, pozostających poza kontrolą graczy, z kośćmi czy ruletką jako prototypem, gdzie miejsce nieistniejącej strategii zajmuje rachunek prawdopodobieństwa, 3) *gry mimetyczne* — oparte na naśladownictwie (teatr, taniec, bale maskowe, a w języku naturalnym gry ikoniczne, często fonetyczne, jak np. zabawy onomatopieczne, figura zwana paronomazją itd.), oraz 4) *zabawy typu ilinx*, których cechą charakterystyczną jest brak jasnych reguł na rzecz doświadczenia mocnych wrażeń (rytuały orgiastyczne, sporty ekstremalne).

Wkrótce potem socjologiczna teoria gier uzyskała kolejne pogłębienie poprzez wskazanie na dodatkowy, psychologiczny wymiar zachowań interpersonalnych. W ramach tzw. analizy transakcyjnej wprowadził ją amerykański psychoterapeuta Eric Berne w książce *Games People Play* ([1964] 2000). Co ciekawe, w tym ujęciu — które można przyjąć za absolutne przeciwieństwo teorii Wittgensteina — gry są z natury niezwykle i zawsze nieuczciwe, z ukrytą motywacją i zasadzką zastawioną na współgracza. W tym znaczeniu są one pokrewne całej gamie *gier podstępnych i nieszczerých*, które w odniesieniu do języka omawiam bardziej szczegółowo w innym miejscu (Chrzanowska-Kluczeńska 2004, w rozdziale poświęconym patologiom i perwersjom gier). Do problemu gier podstępnych powrócimy poniżej (2.4.).

⁴ Dixit i Nalebuff (2010: 443–446) podają bardzo przydatny przegląd najbardziej znanych pozycji bibliograficznych z dziedziny teorii gier. Nieco inną listę „lektur obowiązkowych” zamieszczają w swej książce M. Maławski, A. Wieczorek i H. Sosnowska ([1997] 2004: 198–199). Na temat *konsiliencji*, czyli zbieżności pomiędzy naukami, por. E. O. Wilson ([1998] 2002) i E. Chrzanowska-Kluczeńska (2007).

Obecnie psychologiczny wymiar zachowania „growego”, choć niełatwo poddający się obiektywnemu opisowi, zostaje stopniowo włączany do ekonomicznej teorii gier, w której od dawna zauważono już wpływ elementu subiektywnego na podejmowanie decyzji. Świadczy o tym nowe podejście do ekonomii oparte na *grach behawioralnych*, zwane ekonomią behawioralną, czy też takie dziedziny badań, jak neuroekonomia lub emocjonika (por. Hill [2008] 2010), w których bada się emocjonalne reakcje osób, związane np. z poczuciem złości lub strachu, i ich wpływ na podejmowanie decyzji (por. też Dixit i Nalebuff 2010: 53–54). Powyższy trend w badaniach biznesowo-marketingowych posiada swoje odzwierciedlenie w językoznawstwie kognitywnym, szczególnie w tzw. *third generation cognitive science*, gdzie myśl „ucieleśniona” jest także często myślą naznaczoną uczuciami i jako taka znajduje swój wyraz w ekspresji językowej⁵.

Ponieważ nasze rozważania toczą się dwutorowo, powracamy znów do matematycznej teorii gier, która w połączeniu z Wittgensteinowską koncepcją wszechobecną i instrumentalną gry językowej stała się podstawą dla tzw. *game-theoretical semantics* (GTS), czyli *semantyki teoriogrowej*, stworzonej i rozwijanej nieprzerwanie od schyłku lat sześćdziesiątych ubiegłego stulecia przez fińskiego logika i filozofa języka Jaakko Hintikę oraz jego szkołę (podstawowe pozycje teoriiotwórcze to: Hintikka 1973, Hintikka i Kulas 1983, Hintikka i Kulas 1985). W tym modelu gry rozgrywane są na poziomie zdania, a nagrodą jest ustalenie prawdziwości lub fałszywości zawartego w nim sądu logicznego (*proposition*). Są to typowe gry dwuosobowe o sumie zerowej, gdzie gracz nr 1 (weryfikator) staje naprzeciw rywalowi, czyli Natury (gracz nr 2) i zadaje jej pytania pomagające w ustaleniu wartości prawdziwościowej zdania. GTS jest semantyką teoriomodelową, aksjomatycznie opartą na koncepcji *światów możliwych*, zaczerpniętej z amerykańskiej szkoły filozofii analitycznej i logiki modalnej (Saul A. Kripke [1963] 1971, Nicholas Rescher 1975, David Lewis 1979). Jest też semantyką teorioprawdziwościową, w stylu Alfreda Tarskiego czy Donalda Davidsona (więcej na ten temat piszę w Chrzanowska-Kluczevska 2004, 2012).

W swej zasadniczej postaci semantyka teoriogrowa zajmuje się grami na poziomie funkcjonalnym zdania, podobnie jak w zrębie koncepcji Wittgensteina czy podstawowej wersji teorii aktów illokucyjnych. Jednak już w roku 1983, w ważnej pozycji *Dialogue Games. An Approach to Discourse Analysis*, Lauri Carlson poszerzył pojęcie gry, tak aby swym zasięgiem obejmowała ona (jako łańcuch *podgier*) cały tekst/dyskurs, słusznie wycofując się z wymogu poszukiwania wartości prawdziwościowej jako cechy przynależnej w języku naturalnym tylko zdaniom twierdzącym, a której przyporządkowanie tekstowi

⁵ Margaret H. Freeman (2009: 11) za patronów tego nurtu w kognitywizmie uznaje Maurice’a Merleau-Ponty’ego i Susanne K. Langer.

(zwłaszcza fikcyjnemu) jest bezzasadne. Carlson poddał analizie nie tylko język mówiony (dwu- lub wieloosobowe konwersacje), ale również teksty pisane, którym pragnął narzucić zunifikowaną formę dialogiczną⁶.

W roku 1990 Hintikka, w metodologicznie ważnym artykule *Paradigms for Language Theory*, przeciwstawiał GTS modelowi generatywnemu Noama Chomsky'ego, argumentując na rzecz większej przydatności swojego semantyczno-pragmatycznego systemu do całościowego opisu tekstów/dyskursów, a więc do największych jednostek językowych, w których dopiero w pełni ukazuje się kreatywna moc języka ludzkiego. GTS od swego początku pomyślana była jako sformalizowana kontynuacja teorii gier językowych Wittgensteina, dlatego też gry postrzegane są tu jako narzędzia, których celem jest ugruntowanie języka w praktyce. Hintikka nazywa czasem swój system „semantyką dynamiczną”, służącą eksploracji świata aktualnego i światów alternatywnych w poszukiwaniu nie tylko odniesienia i wartości prawdziwościowej, ale globalnego sensu wypowiedzi. Należy ubolewać, że semantyka teoriogrowa, ze względu na wysoki stopień formalizacji, nie znalazła właściwego odzewu wśród językoznawców. Natomiast współgra ona znakomicie z różnymi formami logiki dialogowej i systemami informatycznymi jako „popularny system semantyki formalnej”, jak określa ją Mathieu Marion (2006: 255).

W dalszej części naszych rozważań zastanowimy się nad różnymi kwestiami natury filozoficznej i metodologicznej, które nasuwają się w odniesieniu do opisu rzeczywistości językowej i pozajęzykowej za pomocą instrumentarium zapożyczonego z teorii gier.

2. Czy każda gra posiada formę dialogową?

Powyższe zapytanie wynika bezpośrednio z faktu, iż najbardziej typowa gra posiada formę dwuosobową (lub dwuzespołową, gdzie zespół/drużynę/koalicję można traktować jako pojedynczego gracza, por. Bacharach 2006). Zakładając więc model dwuosobowy, winniśmy zadać kolejne pytanie o naturę takiego bardzo ogólnie potraktowanego „dialogu”.

2.1. Gry agonalne

Formuła gry opartej na współzawodnictwie czy konflikcie jest prawdopodobnie najstarszą analizowaną w piśmiennictwie, jeśli uznać Arystotelesa za inicjatora takich rozważań. Pamiętać jednak warto, że Arystoteles podkreślał ostatecznie naturalną skłonność człowieka do bycia istotą społeczną. W bliższej nam czasowo tradycji anglosaskiej przywołuje się w tym kontekście Thomasa Hobbesa (którego poglądy polityczne znane są na ogół z jego jedyne

⁶ Polemizuję z tym pomysłem jako niefortunnym, szczególnie w odniesieniu do tekstów literackich (Chrzanowska-Kluczeńska 2004: 4.1.2.)

pisanego po angielsku traktatu *Leviathan*, 1651). Hobbes, w opozycji do Arystotelesa, podkreślał często aspołeczną i egoistyczną naturę istoty ludzkiej. Skoro człowiek za jedyne dobro uznaje własny interes, nie będzie liczył się z pożytkiem bliźniego. Tatariewicz (1947 II: 94) reasumuje stanowisko Hobbesa w sprawie postępowania jednostki w ten sposób:

Więc też każdy z natury rzeczy usiłuje wszystko, czego domaga się jego egoizm, osiągnąć dla siebie i jeśli trzeba, wywalczyć. Nie ma praw ani obowiązków, które by go ograniczały; ile wywalczy, zależy tylko od jego sił. Stąd *status naturalis* człowieka: walka wszystkich przeciw wszystkim. *Homo homini lupus*.

Oto opis gry opartej na konflikcie w najczystszej, skrajnej postaci. Jak pamiętamy, gry współzawodnictwa opisywał Huizinga i Caillois (ten drugi parę ruchów przyległych, wykonanych przez przeciwników w takiej grze nazywał zespołem *provokacja — reakcja*). Również agonistyczna pozostaje wizja świata w kategoriach makrogrzy biochemicznej o kosmicznych wymiarach (w której rozgrywającymi są Życie kontra Śmierć, a biernymi pionkami organizmy żywe), przedstawiona w znanej pozycji *Das Spiel* autorstwa Manfreda Eigena i Ruth Winkler ([1975] 1983).

Na pozytywny aspekt gier współzawodnictwa, zarówno w życiu społecznym, jak i w odniesieniu do gatunków dyskursywnych, wskazuje natomiast Jean-François Lyotard w eseju *La condition postmoderne* (1979). W jego przekonaniu „mówić to walczyć, w sensie współzawodniczyć”, a „akty językowe należą do pewnej ogólnej agonistyki” (Lyotard [1979] 1997: 45)⁷. Problem więzi społecznych sprowadza się również do gry stawiania pytań (tamże 61), która jest z natury konfliktowa, ale też twórcza i innowacyjna:

Żeby w ten sposób zrozumieć stosunki społeczne, niezależnie od poziomu, na jakim się je rozpatruje, potrzebna jest nie tylko teoria komunikacji, ale teoria gier, zawierająca pośród swoich przesłanek agonistykę (Lyotard 1997: 63).

Warto przywołać fakt, że podstawowa wersja GTS Hintikki i jego następców posiada również formę „pytanie — odpowiedź”. Metafory graczy w kolejnych wersjach tej teorii to: *proponujący — oponent, weryfikator — falsyfikator, ja — Natura, Eloise — Abelard, atakujący — obrońca* (Rahman i Dégrement 2006), *budowniczy (systemu) — krytyk* (van Benthem 2006) czy wreszcie *system — środowisko* (ten zdepersonifikowany schemat proponuje dla programów informatycznych Abramsky 2006).

Na polu nauk ekonomicznych, jak twierdzą Dixit i Nalebuff (2010), aż do lat dziewięćdziesiątych XX wieku współzawodnictwo, skoncentrowane egoistycznie

⁷ Lyotard postrzega grę językową nie tylko w kategoriach współzawodnictwa oponentów, jak na przykład w dyskursie akademickim czy politycznym. „Potężnym adwersarzem” w takich grach może stać się sam język, przeciwko którego ustalonym i często skostniałym zasadom gra kreatywnie człowiek (Lyotard 1997: 45).

na własnej osobie gracza i na jego partykularnym interesie, było trendem dominującym, czego najlepszym dowodem jest obowiązująca wówczas definicja *strategii*:

Myślenie strategiczne to sztuka przechytrzenia przeciwnika, przy świadomości, że przeciwnik będzie starał się zrobić to samo z nami (Dixit i Nalebuff 2010: x, tłum. E. C.-K.).

Jednakże w wypadku języka ludzkiego *agon*, choć niewątpliwie obecny w wielu gatunkach mowy i pisma (debata polityczna, dysputa naukowa czy religijna, kłótnia; polemika polityczno-społeczna, artykuł argumentacyjny itp.), nie jest bynajmniej strategią dominującą. Należy więc przypatrzeć się kolejnej niezwykle ważnej kategorii gier.

2.2. Gry kooperacji

Zdecydowanie przeważające w języku naturalnym, którego podstawową funkcją pozostaje nieodmiennie komunikacja oparta na porozumieniu, gry współdziałania — jako przeciwstawienie gier agonalnych — są też niezwykle istotne w każdej działalności społecznej. Dyskurs filozoficzny na ich temat sięga poglądów Arystotelesa dotyczących istoty ludzkiej jako istoty społecznej. Również wspomniana powyżej teoria spinozańska (przyp. 3), zogniskowana na opisie człowieka doskonałego, czyli rozumnego, może być uznana za prekursorkę współczesnej *teorii decyzji*. Posłuchajmy zatem niderlandzkiego myśliciela (pierwszy cytat pochodzi z jego *Traktatu politycznego*, drugi natomiast z *Etyki*):

Gdy dwie osoby umawiają się ze sobą i łączą swe siły, to mogą wspólnie więcej dokonać, a skutkiem tego posiadają pospołu więcej prawa do natury, aniżeli każdy oddzielnie. A im więcej osób łączy w ten sposób swe potrzeby, tym więcej prawa będą posiadali wszyscy razem (Spinoza [1670] 2003: 342 (§ 13)).

Dla ludzi jest nade wszystko pożyteczne łączenie swych obyczajów i wiązanie się wzajemne takimi węzłami, które najlepiej czynią z nich wszystkich jedność [...].

Ale do tego potrzeba umiejętności i czujności. Ludzie bowiem są zmienni (gdyż rzadko zdarzają się tacy, którzy żyją według przepisu rozumu), a jednak są poważnie zawistni i bardziej skłonni do zemsty niż do sympatii (Spinoza [1662–5] 2010: 298, rozdz. XII i XIII).

Dziś, w świetle teorii gier, pierwszy cytat odnieść można do gier wieloosobowych (tzw. *koalicji*, por. Dixit i Nalebuff 2010), w których ludzie stowarzyszają się dla uzyskania większej wypłaty (*payoff*, czyli wygrana lub przegrana w grze), porozumiewając się także co do wspólnej strategii⁸. Natomiast drugi fragment, choć podkreśla raz jeszcze pożytek płynący z kooperacji, zwraca też uwagę na te afektywne cechy charakteru człowieka, które łatwo przechylają

⁸ Niektóre z uwag Spinozy zbliżają się nawet do opisu *gry o sumie zerowej*, gdzie zysk jednej strony w grze równoważy się ze stratą drugiej.

szale na stronę wrogości i konfliktu czy też gier nieszczerých i zwodniczych. Tak więc Spinoza, ukazując słabość ludzkiego charakteru, w istocie odnosi się pośrednio do trzech kategorii gier: kooperacji, *agonu* i gier podstępnych.

Na dwa wieki przed Huizingą na temat ludzkiej potrzeby angażowania się w gry wypowiadał się David Hume w *A Treatise of Human Nature* ([1740] 2009, Book II, Section X 'Of curiosity'). Hume twierdzi ponadto (Book III, Section V 'Of the Obligation of Promises') — zresztą w sposób do złudzenia przypominający współczesne rozważania na temat gier kooperacji, a szczególnie przyczyn, dla których się w nie angażujemy — iż ludzie, którzy z natury są samolubni, a ich szczodrość ograniczona, skłaniają się do działania na korzyść drugiego tylko wówczas, gdy liczą na korzyść z odwzajemnienia. W tym kontekście Hume podaje dwa opisy „growych” zachowań człowieka, które niekiedy są cytowane przez teoretyków gier jako przykład gier dwuosobowych symultanicznych, poprzedzonych jednak spekulacją graczy na temat możliwych korzyści. Jedna sytuacja odnosi się do dwu sąsiadów-gospodarzy, dla których korzystnie byłoby pomóc sobie wzajemnie przy żniwach. Obaj nie czują do siebie sympatii. Maja dwie opcje do wyboru: 1) pomóc sobie dla obopólnego pożytku lub też 2) zrezygnować z pomocy drugiemu, istnieje bowiem ryzyko, że sąsiad nie odwzajemni usługi, a nawet nie okaże wdzięczności. Niestety, ludzie często dokonują drugiego wyboru (okazując brak zaufania do współgracza), co w wypadku dwu nieufnych chłopów oznacza stratę w plonach. W tejże samej księdze (Section VII, 'Of the Origin of Government') Hume twierdzi, że zawarcie porozumienia między dwoma współwłaścicielami łąki w celu jej osuszenia, korzystne dla obu, jest łatwiejsze niż porozumienie się co do podobnej wspólnej akcji przez np. tysiąc osób. Ten kontrast między grą dwuosobową a zawiązaniem szerokiej koalicji prowadzi do argumentacji na rzecz utworzenia „społeczeństwa politycznego”.

Dylemat gry wieloosobowej symultanicznej wymienia też bardzo krótko Jean Jacques Rousseau w rozprawie *Discours sur l'origine et les fondements de l'inégalité parmi les hommes* (1750), w części poświęconej wspólnym interesom i przedsięwzięciom⁹. Polowanie na jelenia, któremu oddaje się grupa pierwotnych myśliwych, jest wystawieniem na próbę ich zdolności do kooperacji: albo pozostaną wszyscy w ukryciu i z zaskoczenia upolują jelenia (zysk dla wszystkich) lub też na widok przebiegającego zająca wyskoczą z kryjówki, a interes partykularny tego, który zająca pochwyli, w żadnej mierze nie zrekompensuje straty wszystkich myśliwych, jeleni bowiem spłoszy się i ucieknie¹⁰.

⁹ To jedno zdanie z Rousseau, które stało się ogólnie znane w teorii gier jako *stag hunt*, brzmi w wersji angielskiej następująco: „If a deer was to be taken, everyone saw that, in order to succeed, he must abide faithfully by his post: but if a hare happened to come within the reach of any one of them, it is not to be doubted that he pursued it without scruple, and, having seized his prey, cared very little, if by doing so he caused his companions to miss theirs (zob. też formalny opis tej gry na http://en.wikipedia.org/wiki/Stag_hunt).

¹⁰ Bardziej skomplikowaną wersję tej gry symultanicznej, gdy myśliwi mogą upolować trzy rodzaje zwierzyny (bizona, jelenia bądź zająca) omawiają Dixit i Nalebuff (2010: 102–103).

W kategoriach współczesnej teorii gier zarówno zachowanie wieśniaków u Hume'a, jak i pierwotnych myśliwych u Rousseau, pokrewne jest jednej z najtrudniejszych do wyjaśnienia w kategoriach strategicznych gier, zwanej *dylematem więźnia*. Jest to historia dobrze znana z powieści (i jej ekranizacji) *Z zimną krwią* autorstwa Trumana Capote. Dwóch więźniów, oskarżonych o morderstwo (na które brak jest jednak dowodów), pozostaje uwięzionych w dwu celach. Nie mogą więc porozumiewać się ze sobą. Policja namawia obu na złożenie szybkich zeznań obciążających drugą stronę w zamian za umorzenie kary jednemu z nich. Gra ta dyskutowana jest *in extenso* w każdym podręczniku teorii gier (por. Malawski et al. 2004; Dixit i Nalebuff 2010). Pomimo, iż posiada ona dwa punkty równowagi gry (*równowaga Nasha*¹¹), tylko jedna z opcji jest naprawdę korzystna dla obu graczy: obaj powinni milczeć. Tymczasem z reguły dochodzi do sytuacji odwrotnej: obydwaj obciążają siebie nawzajem, w ten sposób doprowadzając do wykazania własnej winy. Podłoże takiej decyzji jest psychologiczne, a Dixit i Nalebuff (2010: 89) wiążą podobne zachowania w innych grach z *imperatywem kategorycznym* Kanta, który mówi:

Postępuj tylko według takiej maksymy, dzięki której możesz zarazem chcieć, żeby stała się powszechnym prawem (Kant [1785] 1984: 50)¹².

To oznacza, iż gracze często mylnie zakładają, że druga osoba będzie powodowana takimi samymi względami i podobną motywacją jak oni sami. Niestety, taka strategia może okazać się zawodna, ponieważ gry ludzkie dalekie są od idealnych założeń racjonalnego i etycznego postępowania w każdej sytuacji.

Potrzeba kooperacji jest w obecnym świecie ogromna, musi się tu jednak brać pod uwagę cały splot czynników kontekstowych, nie tylko związanych z psychologią uczestników. Michael Bacharach (2006: xviii, 112) zasłużył się formalnym opisem strategii dla gier zespołowych opartych na czynnikach społecznych, takich jak identyfikacja z grupą. Wprowadził on pojęcie *ramy gry*, która jest niczym innym jak przyjęciem odpowiedniej perspektywy, indywidualnej (*I' frame*) lub grupowej (*'we' frame*). Według Bacharacha człowiek ewolucyjnie został przygotowany do różnych form kooperacji, łatwych (sytuacje wspólnego interesu) lub trudniejszych (typu *dylemat więźnia*), lecz zawsze odbywających się w ściśle określonym kontekście, a nie w przestrzeni mate-

¹¹ Dixit i Nalebuff (2009: 130) komentują tzw. *punkt stabilizacji rozważań*, czyli *równowagę gry* w taki sposób: „Taki wynik, gdzie działanie każdego z uczestników jest dla niego najlepsze w świetle przekonania o wyborze drugiego gracza i działanie każdego uczestnika zgadza się z przekonaniem o tym działaniu drugiego gracza, bardzo łatwo rozwiązuje problem kwadratury koła, czyli myślenia o myśleniu”. Autorzy wykazują więc, że strategia stanowi *epistemiczną superstrukturę* gry na zasadzie: „gracz nr 1 myśli, że gracz nr 2 myśli, że gracz nr 1...”.

¹² Przy innej okazji Dixit i Nalebuff (2010: 74) przywołują bliskie duchowi imperatywu polecenie Jezusa pochodzące z „Kazania na górze” (Łk. 6: 31): „A jak chcecie, aby wam ludzie czynili, tak i wy im czyńcie”.

matycznej idealizacji, zwanej *przestrzenią gry* (por. van Benthem 1990). Matematyczno-ekonomiczno-socjologiczne teoretyzowanie Bacharacha, postrzegające gry jako nieodmiennie rozgrywane w ściśle określonej rzeczywistości społecznej, uzupełniają zgrabnie rozważania Ilkki Niiniluoto (2006) na gruncie semantyki teoriogrowej, a dotyczące „gry z punktu widzenia x” oraz takich pojęć, jak „zbiorowa wiedza” czy „kolektywne wierzenia”, które autor traktuje jako operatory epistemiczne rozszerzone na agensa grupowego, wzbogacające logikę modalną.

Przechodząc na grunt rozważań nad naturą języka, można zauważyć, że Huizinga nie omawiał w ogóle gier kooperacyjnych językowych (a tylko społeczne, np. gry zespołowe), natomiast Caillois całkowicie pominął tę ogólną kategorię gier w swojej taksonomii. To Wittgensteinowi zawdzięczamy dopiero spojrzenie na gry językowe jako typowo skierowane na współdziałanie z otoczeniem. Ten wątek kooperacji podejmowali potem wszyscy badacze aktywni na polu teorii aktów mowy oraz H. Paul Grice ([1975] 1977) poprzez swą Zasadę Kooperacji (*Cooperation Principle*¹³). To ona w znacznej mierze odpowiedzialna jest za koherentną i skuteczną komunikację interpersonalną w mowie i piśmie (choć niekoniecznie w dyskursie literackim, gdzie „udziwnienie” jako zasada artystyczna często kładzie kres łatwej współpracy między twórcą a interpretatorem).

W samej semantyce teoriogrowej można zauważyć pewną dwoistość w podejściu do gier. Podstawowe gry na poziomie zdania odbywają się raczej na zasadzie *agonu*, jeśli chodzi o wartościowanie logiczne. Z drugiej strony, *gry dyskursywne* Carlsona oparte są na *dialogach*, których celem jest negocjacja wspólnego stanowiska. Za takim podejściem stoi dosyć mocna filozoficzna tradycja niemieckiej szkoły dialogu, do której zaliczyć trzeba Hansa-Georga Gadamera ([1960] 1993) i Jürgena Habermasa ([1985] 2000). Ten pierwszy posługiwał się również metaforą stawiania pytań Matce Naturze (Pani Fortunie) w naszym współdziałaniu ze światem i próbie jego interpretacji, ten drugi z kolei odwoływał się do koncepcji gry (dwu- lub trzyosobowej) w dochodzeniu do konsensusu społecznego na drodze porozumienia i poszukiwania wspólnych informacji. Dla obu więc forma dialogiczna stanowi o największej wartości gry, czy to pomiędzy graczami indywidualnymi czy też zespołowymi. Język naturalny w tym procesie odgrywa zasadniczą rolę, będąc niekiedy aktywnym uczestnikiem gry. Również fenomenologiczne rozważania Merleau-Ponty’ego, zatytułowane *Proza świata* ([1969] 1976), a szczególnie niezwykle ekspresyjny esej *Obecni w słowie*, zawierają pochwałę bytowania w świecie poprzez dialog z „innymi”:

¹³ Van Benthem (1990: 250) traktuje globalne konwencje językowe typu postulatów Grice’a jako posiadające wyraźny charakter teoriogrowy. Co ciekawe, Hintikka (1990: 185) uważa maksymy konwersacyjne Grice’a nie za strategie, ale raczej podrzędne im reguły definicyjne gry dyskursywnej.

W słowie urzeczywistnia się niemożliwa zgoda dwóch rywalizujących ze sobą całości, ale nie dlatego, żebyśmy dzięki niemu mieli powrócić do siebie i odnaleźć jakiegoś jedyne go ducha, w którym mielibyśmy uczestniczyć, lecz dlatego, że nas dotyczy, dotyka nas z ukosa, pociąga nas, uwodzi, przekształca nas w innych i *vice versa*, jako że niweczy granice pomiędzy tym, co moje, a tym, co cudze [...] (Merleau-Ponty 1976: 78).

Oto poetycki opis gry dwuosobowej, tym ciekawszy z formalnego punktu widzenia, iż nie jest to bynajmniej klasyczna gra dwuosobowa o sumie zerowej (*agon*), ale niezwykle połączenie strategii prowadzącej do *zgody z rywalizacją*. O takich właśnie grach od lat dziewięćdziesiątych XX wieku coraz częściej mówi teoria gier ekonomicznych.

2.3. Gry typu *co-opetition*

W ostatnich dwu dekadach teoretycy gier uświadomili sobie na podstawie obserwacji empirycznych, że dobra strategia powinna w wielu sytuacjach być połączeniem *kooperacji* i *współzawodnictwa*, a więc *zgody* i *rywalizacji*, co w języku angielskim Adam Brandenburger i Barry J. Nalebuff (1996) ochrzczili kontaminacją słów *co-operation* i *competition*, czyli *co-opetition* (*kooperacja-konflikt?*). Do tej nowej formuły dopasowano też nową definicję strategii dla gry o podobnym, mieszanym charakterze:

[Strategia] to także sztuka odnajdywania sposobów na współpracę, nawet jeśli innymi [graczami] kieruje interes własny, a nie uczynność. [...] To sztuka interpretacji i odkrywania informacji. To sztuka wstawiania się w położenie innych [„wchodzenie w ich buty”], aby przewidzieć i mieć wpływ na ich postępowanie (Dixit i Nalebuff 2010: x, tłum. E. C.-K.).

Niewątpliwie gry tego typu posiadają bardziej złożoną strukturę niż proste strategie gier agonistycznych i kooperacyjnych. Często będą to gry wieloosobowe o strategiach mieszanych lub *zrandomizowanych* (czyli często odmienianych przez gracza w obrębie tej samej gry). Przywołany przez van Benthema (2006: 124) casus procesu sądowego, czyli gry o skomplikowanej strukturze, możemy uznać za dyskurs, w którym wyróżniają się różne typy gier: *agon* na linii prokurator — adwokat, prokurator — oskarżony, być może też na linii prokurator/adwokat — sędzia; ława przysięgłych może (choć nie musi) działać jako koalicja; adwokat i oskarżony powinni współpracować, ale wiemy z praktyki sądowej, że często pozostają w konflikcie; świadkowie często zmieniają swe sympatie itp. W opisach skomplikowanych gier tego typu, społecznych i językowych, zawsze trzeba wziąć pod uwagę element emocjonalny i irracjonalny w zachowaniu graczy, słowem zakładają one potrzebę humanizacji i psychologizacji matematycznej teorii gier.

2.4. Gry o niedoskonałej informacji i manipulacja

Od samego początku rozwoju semantyki teoriogrowej Hintikka i jego następcy (np. van Benthem 2006) zwracali uwagę na fakt, że gry odgrywane przez ludzi w bardzo zasadniczy sposób odbiegają od zakładanej przez matematyczną teorię gier idealizacji, w której gracze obdarzeni są absolutną pamięcią co do ruchów swoich i współgraczy oraz pełną świadomością własnej i cudzej strategii. Gry takie zwane są *grami o pełnej informacji*. W rzeczywistości gry ludzkie należy w większości wypadków traktować jako *gry o niepełnej/niedoskonałej informacji*. Piszą o tym interesująco Shahid Rahman i Cédric Dégrement (2006), wskazując na fakt, że w grach społecznych występuje nagminnie zjawisko *luk informacyjnych*, a często informacja dociera do graczy z opóźnieniem. Ten niedoskonały przekaz informacji, związany często z celowym *zatajeniem informacji* (jak np. w pokerze), powoduje, iż pojęcie idealnego agensa w grze matematycznej w rzeczywistości powinno zostać zastąpione pojęciem gracza *ograniczonego kognitywnie*. Fakt, że informacja lub jej brak wpływa znacząco na przebieg gry, wiąże się ze zjawiskiem *manipulacji informacją*, czego pochodną jest również *manipulacja językowa*, znana od niepamiętnych czasów filozofom i retorom. Dixit i Nalebuff (2010: 238) wymieniają następujące proste sytuacje strategiczne, związane z manipulacją informacją: 1) niektórzy gracze posiadają więcej danych niż ich współgracze na temat, od którego zależy wypłata gry; 2) gracze, którzy posiadają dodatkową informację, chętnie ją ukrywają, 3) z kolei niektórzy gracze ciesząc się dostępem do dodatkowej informacji chętnie udostępniają ją innym, wreszcie 4) gracze posiadający niewystarczającą informację będą chcieli wydobyć więcej danych od współgraczy (np. w sposób uczciwy, jak Król Salomon).

Łatwo więc zauważyć, że kłopoty z informacją w grze mogą (choć nie muszą) prowadzić do zachowań ocenianych pejoratywnie, a więc do całej gamy gier oszukańczych.

Teoretycy gier czasami odwołują się tu do doktryny pisarza politycznego Niccoló Machiavellego, historyka i dyplomaty florenckiego okresu Odrodzenia, który w traktacie *Il principe (Księżę)* wyłożył naganną etycznie doktrynę „celu uświęcającego środki”. Tak więc, jeżeli okoliczności tego wymagają, trzeba czasem stosować w rozgrywce z przeciwnikiem metody niemoralne. W kategorii gier będą to posunięcia podstępne. Machiavelli zastrzegał, że środki takie powinny być stosowane tylko w koniecznym zakresie, szczególnie wtedy, gdy cel sam w sobie jest szczytny. Wiemy jednak z praktyki, że podobne przyzwolenie prowadzi często do różnego rodzaju nadużyć, w słowie lub czynie.

Warto przypomnieć, że definicja *gier transakcyjnych* (codziennych zachowań interpersonalnych) według Berne’a w sposób eksplicytny odwołuje się do postępowania obłudnego i podstępnego:

Grą nazywamy serię komplementarnych transakcji ukrytych, prowadzących do dobrze określonego, dającego się przewidzieć wyniku. Mówiąc bardziej opisowo, jest to okresowy, często powtarzający się zestaw transakcji, pozornie bez zarzutu, o utajonej motywacji, czy też bardziej potocznie, seria posunięć z pułapką albo „sztuczką”. Gry wyraźnie różnią się od procedur, rytuałów i rozrywek dzięki swym dwóm podstawowym właściwościom: (1) swojej ukrytej jakości i (2) wyplacie. Procedury mogą być udane, rytuały skuteczne, a rozrywki korzystne, lecz wszystkie one są z definicji szczerze; mogą pociągać za sobą rywalizację, ale nie konflikt, a ich zakończenie może być sensacyjne, lecz nie dramatyczne. *Każda gra jest natomiast w swoim założeniu nieuczciwa, a wynik ma wydźwięk dramatyczny, a nie jedynie ekscytujący* (Berne [1964] 2000: 37, podkr. moje).

Podejście reprezentowane przez Berne’a, choć mieści się w paradygmacie Huizingi (gra to działalność szczególna, a nie powszechna), kładzie nacisk na cechy rozgrywki negatywne, choć w perwersyjny sposób ludyczne. Teoria ta bardzo znacząco odbiega natomiast od paradygmatu gier Wittgensteina-Hintikki, w którym emfaza położona została na pozytywny wymiar współdziałania między ludźmi.

Blisko stąd do pokrewnego problemu manipulacji informacją (wiedzą) czy uczuciami w tekście/dyskursie. *Manipulacja językowa* rozumiana jest tu zazwyczaj jako odmiana perswazji z ukrytym celem lub też perswazja nie do końca lub w ogóle nierozpoznana przez odbiorcę (por. Chrzanowska-Kluczevska, w druku 2013). Podobnie jak w grach transakcyjnych Berne’a, jako szczególne działanie strategiczne może być ona źródłem zadowolenia jednej ze stron biorących udział w takiej rozgrywce, natomiast odbiorca szczerzy lub naiwny może nawet nie rozpoznać tej sytuacji jako „growej”. Ograniczone ramy naszej dyskusji nie pozwalają mi na bardziej szczegółowe rozważanie manipulacji w języku, jednak — co warto podkreślić — istnieje w tej materii duża zbieżność między rozważaniami o ogólnej naturze manipulacji w życiu społecznym i ekonomicznym, a jej językową realizacją.

3. Paradygmat strategiczny w semantyce języka naturalnego

Van Benthem (1990) plasuje semantykę teoriogrową Hintikki w obrębie lingwistyki zorientowanej kognitywnie, ponieważ gry językowe w tym systemie istnieją nie tylko „w głowie użytkownika”, ale przede wszystkim ułatwiają kontakty ze światem zewnętrznym. Klasyfikuje on gry jako *kategorię unifikacyjną*, ważną we wszystkich naukach kognitywnych, które zajmują się interakcją pomiędzy istotami inteligentnymi i ich środowiskiem, powtarzając za Hintikką następującą tezę: „Każda racjonalna aktywność kognitywna może być grą” (van Benthem 1990: 236).

Dodatkowo autor odwołuje się tu do rozróżnienia, wprowadzonego przez Hintikkę, na tzw. *I-games* (*internal games*, gry wewnątrz systemu językowego) oraz *E-games* (*external games/environmental games*, gry zewnętrzne na styku

języka i świata/środowiska). Podział ten odpowiada z grubsza innej wprowadzonej przez Hintikka (1990) dystynkcji na dwa typy reguł w grze: *reguły niższego rzędu*, tzw. *reguły definicyjne gry*, które konstytuują daną grę i określają właściwe dla niej ruchy, oraz *reguły wyższego rzędu*, tzw. *reguły strategiczne*, które pomagają nam rozgrywać grę w sposób prowadzący do zwycięstwa oraz do najkorzystniejszej w danej sytuacji wypłaty. W odniesieniu do języka te pierwsze to rekursywne reguły składni i semantyki, działające najczęściej na poziomie zdania i odpowiedzialne za jego gramatyczność i sensowność. Natomiast reguły strategiczne, zwane często *superstrukturą gry*, pełnią dwojaką funkcję: po pierwsze, przyczyniają się do budowy tekstu/dyskursu i jego *koherencji* (w znaczeniu globalnej spójności semantycznej), po drugie, odnoszą tekst/dyskurs do jego pragmatycznej obudowy, a więc stają się podstawą dla interpretacji danego tekstu w ściśle określonym kontekście.

Podobnie Maciej Grochowski (1993: 100) w analizie poświęconej eksplikacji semantycznej leksemu *gra* w języku naturalnym wyróżnia dwa typy gier społecznych, zależnie od celu, jaki przyświeca zaangażowanym w nie graczom: gry *nastawione na przebieg* oraz gry *nastawione na rezultat*. Odwołuje się tu on do propozycji Jerzego Jarzębskiego (1977), aby wyodrębnić *cel wewnętrzny* i *cel zewnętrzny gry*. W sensie GTS gry pierwszego typu w odniesieniu do języka naturalnego będą gramami o gramatyczność i sens (oraz wartość prawdziwościową dla ograniczonej grupy zdań deklaratywnych), natomiast gry drugiego typu, odgrywane w przestrzeni dyskursu, będą gramami o skuteczność języka w działaniu i o sprawną komunikację. Te pierwsze Hintikka i jego szkoła określają często jako *połączenia horyzontalne* (wewnątrzjęzykowe), te drugie jako *relacje wertykalne* (zewnętrzne, łączące język ze światem)¹⁴.

Hintikka wielokrotnie podkreślał wagę *paradygmatu strategicznego* dla badania języka ludzkiego (por. Hintikka 1990), uważając, że systemy generatywne Noama Chomsky'ego nigdy nie przekroczyły poziomu reguł definicyjnych. Dopiero przyjęcie paradygmatu strategicznego umożliwia holistyczne podejście do znaczenia tekstu i konieczne dla jego właściwej interpretacji odwołanie się do świata aktualnego oraz innych światów możliwych (modalność, fikcja). Podobnie van Benthem (1990) ubolewa, iż do tej pory *paradygmat teoriogrowy* nie wygrał z *paradygmatem obliczeniowym* (*computational paradigm*), a więc z paradygmatem proceduralnym w znaczeniu ściśle zaprogramowanego przechodzenia od jednego do drugiego stanu w reprezentacji języka ludzkiego lub sztucznego. Chodzi tu o etapy konstruowania informacji, którym jednak miano „stanów kognitywnych” przysługiwać może wyłącznie w odniesieniu do inteligentnego współgrania z otoczeniem.

¹⁴ Niezwykle ważnym problemem związanym z podziałem gier językowych na dwa typy funkcjonalne pozostaje sprawa kompozycyjności i kontekstualizmu w odniesieniu do GTS, której warto by było poświęcić odrębne rozważania (por. Haaparanta 2006, Chrzanowska-Kluczevska 2012).

Dwie dekady, które upłynęły od czasu przytoczonych powyżej dyskusji warsztatowych prowadzonych przez semantyków formalnych na temat natury gier językowych, przyniosły ze sobą dynamiczny rozwój pragmatyki językoznawczej, teorii dyskursu, a również samej teorii gier, która od matematycznej abstrakcji przesunęła się ku kompleksowym badaniom zachowań ludzkich w dziedzinach ekonomii, psychologii i socjologii. Również sama semantyka teoriogrowa zmierza sukcesywnie od logiki formalnej ku naukom kognitywnym.

Mam nadzieję, że nasze krótkie rozważania nad podstawami teorii gier społecznych i językowych ukazują nie tylko ich historyczne zakotwiczenie w filozofii, ale uświadamiają Czytelnikowi potrzebę interdyscyplinarnego spojrzenia na powszechny w naszym życiu wymiar gry i zabawy. Wydaje się więc, że twórca pojęcia *konsiliencji*, Edward O. Wilson, słusznie utrzymuje, iż:

Teorie, których podstawowe pojęcia i procesy są spójne z solidnie ugruntowaną wiedzą z innych dyscyplin, są zdecydowanie skuteczniejsze w teorii i praktyce od teorii, które takiej spójności nie wykazują (Wilson [1998] 2002: 301).

Bibliografia

- ABRAMSKY S. (2006), *Socially Responsive, Environmentally Friendly Logic*, [w:] Aho T., Pietarinen A.-V. (red.), 17–45.
- AHO T., PIETARINEN A.-V. (red.) (2006), *Truth and Games. Essays in Honour of Gabriel Sandu*, "Acta Philosophica Fennica" Vol. 78. Helsinki: University of Helsinki.
- AUSTIN J. L. (1962), *How to Do Things with Words*, Urmson J. O. (red.). Oxford: Oxford University Press (1993). *Jak działać słowami*, [w:] Austin, J. L., *Mówienie i poznawanie: rozprawy i wykłady filozoficzne*, przeł. B. Chwedeńczuk. Warszawa: Wydawnictwo Naukowe PWN).
- BACHARACH M. (2006), *Beyond Individual Choice. Teams and Frames in Game Theory*, Gold N., Sugden R. (red.). Princeton-Oxford: Princeton University Press.
- BENTHEM J. VAN. (1990), *Computation versus Play as a Paradigm for Cognition*, [w:] Haaparanta L. et al. (red.), 236–251.
- BENTHEM J. VAN. (2006), *Logical Construction Games*, [w:] Aho T., Pietarinen A.-V. (red.), 123–137.
- BERNE E. (1964), *Games People Play. The Psychology of Human Relationships*. New York: Random House (2000). *W co grają ludzie. Psychologia stosunków międzyludzkich*. Warszawa: Wydawnictwo Naukowe PWN).
- BLACK M. (1986), *Wittgenstein's Language-Games*, [w:] Shanker S. (red.), 74–88.
- BRANDENBURGER A., NALEBUFF B. J. (1996), *Co-opetition*. New York: Doubleday.
- BUJNICKI T., SŁAWIŃSKI J. (red.), (1977), *Problemy odbioru i odbiorcy*. Wrocław-Warszawa: Ossolineum.
- CAILLOIS R. (1958), *Les jeux et les hommes*. Paris: Éditions Gallimard (1997). *Gry i ludzie*. Warszawa: Volumen).
- CARLSON L. (1983), *Dialogue Games. An Approach to Discourse Analysis*. Dordrecht: Reidel.

- CHRZANOWSKA-KLUCZEWSKA E. (2004), *Language-Games: Pro and Against*. Kraków: Universitas.
- CHRZANOWSKA-KLUCZEWSKA E. (2007), *Konsiliencja, czyli o porozumieniu między naukami w trzecim tysiącleciu*, [w:] Szpila G. (red.), 15–23.
- CHRZANOWSKA-KLUCZEWSKA E. (2012), *Modal Games in Natural Language*, [w:] Stalmaszczyk P. (red.), 57–79.
- CHRZANOWSKA-KLUCZEWSKA E., *Sztuka strategii w dyskursie w świetle rozwoju teorii gier* (w druku, *Prace Komisji Neofilologicznej PAU*, t. XI, Kraków 2012, 33–48).
- CHRZANOWSKA-KLUCZEWSKA E., SZPILA G. (red.) (2009), *In Search of (Non)Sense*. New-castle u. Tyne: Cambridge Scholars Publishing.
- DIXIT A. K., NALEBUFF B. J. (1991), *Thinking Strategically: The Competitive Edge in Business, Politics, and Everyday Life*. New York: Norton (2009. *Myślenie strategiczne*. Gliwice: Onepress).
- DIXIT A. K., NALEBUFF B. J. ([2008] 2010), *The Art of Strategy. A Game Theorist's Guide to Success in Business and Life*. New York-London: W. W. Norton & Co. (2009. *Sztuka strategii. Teoria gier w biznesie i życiu prywatnym*. Warszawa: MT Biznes).
- EIGEN M., WINKLER R. (1975), *Das Spiel*. München: R. Piper Co. Verlag (1983. *Gra. Prawa natury sterują przypadkiem*. Warszawa: PIW).
- FREEMAN M. H. (2009), *Making Sense of (Non)Sense: Why Literature Counts*, [w:] Chrzanowska-Kluczevska E., Szpila G. (red.), 1–19.
- GADAMER H.-G. ([1960] 1989), *Truth and Method*. New York: Crossroad (1993. *Prawda i metoda Zarys hermeneutyki filozoficznej*. Kraków: Inter Esse).
- GRICE H. P. ([1975] 1977), *Logika a konwersacja*, „Przegląd Humanistyczny” 6, 85–100.
- GROCHOWSKI M. (1993), *Konwencje semantyczne a definiowanie wyrażen językowych*. Warszawa: Zakład Semiotyki Logicznej UW.
- HAAPARANTA L. (2006), *Compositionality, Contextuality and the Philosophical Method*, [w:] Aho T., Pietarinen A.-V. (red.), 289–301.
- HAAPARANTA L., KUSCH M., NIINILUOTO I. (red.) (1990), *Language, Knowledge and Intentionality. Perspectives on the Philosophy of Jaakko Hintikka*, „Acta Philosophica Fennica” Vol. 49. Helsinki: University of Helsinki.
- HABERMAS J. ([1985] 2000), *Filozoficzny dyskurs nowoczesności*. Kraków: Universitas.
- HILL D. (2008), *Emotionomics: Leveraging Emotions for Business Success*. London-Philadelphia: Kogan Page (2010 *Emocjonika. Wykorzystanie emocji a sukces w biznesie*. Poznań: Rebis).
- HINTIKKA J. (1973), *Logic, Language-Games and Information. Kantian Themes in the Philosophy of Logic*. Oxford: Clarendon Press.
- HINTIKKA J. (1990), *Paradigms for Language Theory*, [w:] Haaparanta L. et al. (red.), 181–209.
- HINTIKKA J., KULAS J. (1983), *The Game of Language. Studies in Game-Theoretical Semantics and Its Applications*. Dordrecht: Reidel.
- HINTIKKA J., KULAS J. (1985), *Anaphora and Definite Descriptions. Two Applications of Game-Theoretical Semantics*. Dordrecht: Reidel.
- HUIZINGA J. ([1938] 1967), *Homo ludens. Zabawa jako źródło kultury*. Warszawa: Czytelnik.

- HUME D. ([1740] 2009), *A Treatise of Human Nature. Being an Attempt to Introduce the Experimental Method of Reasoning into Moral Subjects*. The Floating Press (wydawnictwo internetowe).
- JARZĘBSKI J. (1977), *O zastosowaniu pojęcia gra w badaniach literackich*, [w:] Bujnicki T., Sławiński J. (red.), 23–46.
- KANT I. ([1785] 1984), *Uzasadnienie metafizyki moralności*. Warszawa: PWN.
- KRIPKE S. A. ([1963] 1971), *Semantical Considerations on Modal Logic*, [w:] Linsky L. (red.), 63–72.
- LEWIS D. (1979), *Possible Worlds*, [w:] Loux M. J. (red.), 182–189.
- LINSKY L. (red.) (1971), *Reference and Modality*. Oxford: Oxford University Press.
- LOUX M. J. (red.) (1979), *The Possible and the Actual*. Ithaca–London: Cornell University Press.
- LUCE R. D., RAIFFA H. (1958), *Games and Decisions. Introduction and Critical Survey*. New York: John Wiley and Sons (1964. *Gry i decyzje*. Warszawa: PWN).
- LYOTARD J.-F. (1979), *La condition postmoderne. Rapport sur le savoir*. Paris: Minuit (1997. *Kondycja ponowoczesna. Raport o stanie wiedzy*. Warszawa: Fundacja Aletheia).
- MALAWSKI M., WIECZOREK A., SOSNOWSKA H. ([1997] 2004), *Konkurencja i kooperacja. Teoria gier w ekonomii i naukach społecznych*. Warszawa: Wydawnictwo Naukowe PWN.
- MARION M. (2006), *Hintikka on Wittgenstein: From Language-Games to Game Semantics*, [w:] Aho T., Pietarinen A.-V. (red.), 255–274.
- MERLEAU-PONTY M. ([1969] 1976), *Proza świata. Eseje o mowie*. Warszawa: Czytelnik.
- NEUMANN J. VON, MORGENSTERN O. ([1944] 1947), *Theory of Games and Economic Behavior*. Princeton: Yale University Press.
- NIINILUOTO I. (2006), *The Poverty of Relative Truth*, [w:] Aho T., Pietarinen A.-V. (red.), 165–174.
- RAHMAN S., DÉGREMONT C. (2006), *The Beetle in the Box: Exploring IF Dialogues*, [w:] Aho T., Pietarinen A.-V. (red.), 91–122.
- RESCHER N. (1975), *A Theory of Possibility*. Pittsburgh: Pittsburgh University Press.
- ROUSSEAU J.-J. (1750), *A Discourse on the Origin of Inequality*. Digireads (wydawnictwo internetowe).
- SEARLE J. R. (1969), *Speech Acts. An Essay in the Philosophy of Language*. Cambridge: Cambridge University Press (1987. *Czynności mowy: rozważania z filozofii języka*, przeł. B. Chwedeńczuk. Warszawa: Instytut Wydawniczy PAX).
- SHANKER S. (red.) (1986), *Ludwig Wittgenstein. Critical Assessments 2*. New York: Croom Helm.
- SPINOZA B. ([1670] 2003), *Traktaty*. Kęty: Wydawnictwo Antyk.
- SPINOZA B. ([1662–5] 2010), *Etyka w porządku geometrycznym dowiedziona*. Warszawa: Wydawnictwo Naukowe PWN.
- STALMASZCZYK P. (red.) (2012), *Philosophical and Formal Approaches to Linguistic Analysis*. Frankfurt: Ontos Verlag.
- SZPILA G. (red.) (2007), *Język polski XXI wieku: analizy, oceny, perspektywy*, monografia z cyklu *Język trzeciego tysiąclecia*. Kraków: Tertium.
- TATARKIEWICZ W. ([1931] 1947), *Historia filozofii*, t. II (wyd. III). Warszawa: Czytelnik.
- WILSON E. O. ([1998] 2002), *Konsiliencja. Jedność wiedzy*. Poznań: Wydawnictwo Zysk i S-ka.

WITTGENSTEIN L. ([1958] 1998), *Niebieski i Brązowy zeszyt. Szkice do Dociekań filozoficznych*. Warszawa: Wydawnictwo Spacja.

WITTGENSTEIN L. (1953), *Philosophical Investigations* (transl. G. E. M. Anscombe). Oxford: Basil Blackwell (2000. *Dociekania filozoficzne*, przeł. B. Wolniewicz. Warszawa: Wydawnictwo Naukowe PWN).

Rozdział 14

Semantyka formalna i Program Minimalistyczny: od logicznej analizy języka do biolingwistycznego naturalizmu

Jarosław Jakielaszek (Uniwersytet Warszawski)

Słowa kluczowe: biolingwistyka, Chomsky, Frege, generatywizm, gramatyka Montague, Leibniz, Montague, Program Minimalistyczny, semantyka formalna, Tarski

1. Autonomia składni i autonomia zdolności językowej

Program Minimalistyczny (zob. Chomsky 1993, 1995, 2000a, 2001, 2005, 2007, 2008) jako kontynuacja myśli generatywnej okresu wcześniejszego — bezpośrednio teorii rządu i wiązania — przejmując z etapów poprzednich szeregi założeń natury ogólnej; mimo wielu szczegółowych innowacji wprowadzonych w pierwszym okresie minimalizmu, poczynając od radykalnej redukcji liczby operacji składniowych i rezygnacji ze znaczników frazowych jako formalnie poprawnej reprezentacji struktury składniowej, po uproszczenie wewnętrznej struktury komponentu składniowego i jego relacji z komponentami nieskładniowymi, minimalizm zachowuje w istotnych punktach ciągłość z tradycją sięgającą pierwszych propozycji generatywnych (zob. Stalmaszczyk 2011). Do takich podstawowych założeń należy uznanie zdolności językowej człowieka, przedstawianej jako „organ” umysłu, za podstawowy przedmiot badania językoznawstwa, którego celem jest przedstawienie wyjaśniająco adekwatnej teorii zdolności do nabywania języka charakteryzowanego jako system cechujący się dyskretnością, nieogranicznością i hierarchicznością. Niektóre z takich fundamentalnych założeń są dyskutowane i formułowane w niemal wszystkich dyskusjach podstaw minimalizmu; inne, choć o nie mniejszym znaczeniu dla kształtu teorii językoznawczej, zostały w okresie minimalistycznym zepchnięte na plan dalszy. Wśród tych ostatnich wskazać można założenie autonomii składni.

Hipoteza autonomii składni (w generatywnym rozumieniu tego pojęcia), zakładająca niezależność mechanizmów i obiektów składniowych od własności

i mechanizmów semantycznych, pragmatycznych oraz związanych z powierzchnią realizacją struktur składniowych, czyniła ze składni — centralnego przedmiotu zainteresowania analizy generatywnej — komponent poddający się badaniu niezależnemu nie tylko od związanych z komunikacją werbalną funkcjonalnych aspektów języka (zgodnie z rozróżnieniem E-języka i I-języka i wskazaniem I-języka jako właściwego przedmiotu badań językoznawczych), ale także od wszelkich własności języka *ex hypothesi* nieskładniowych, choć i nienależących wprost do cech motywowanych użyciem języka jako narzędzia komunikacji. Autonomia składni, mocno akcentowana w konsekwencji tzw. wojen językowych z nurtem Semantyki Generatywnej, pojmowana jest tradycyjnie jako, z jednej strony, hipoteza motywująca oparcie wyjaśnień istotnych dla generatywnej analizy zjawisk językowych na własnościach wyłącznie i specyficznie składniowych, z drugiej zaś strony jako uzasadnienie dla takiej architektury zdolności językowej, w której komponenty inne niż repozytorium jednostek leksykalnych mają charakter wyłącznie interpretacyjny w stosunku do struktur tworzonych przez składniowy mechanizm generatywny — w szczególności komponent składniowy jest całkowicie uniezależniony od komponentu semantycznego, którego rolą jest interpretacja (nie zaś współtworzenie) struktur składniowych. Model taki, nazywany modelem T lub modelem Y architektury zdolności językowej, zakłada jednokierunkowość przepływu informacji: wyłącznie ze składni do komponentów interpretacyjnych. Leksykon, dostarczający elementarnych obiektów manipulowanych przez mechanizm składniowy, znajduje się w istocie poza polem badania analizy generatywnej i obejmuje elementy charakteryzowane za pomocą tzw. cech należących do wzajemnie rozłącznych zbiorów własności dźwiękowych, semantycznych i składniowych (formalnych). Podobnie wyłączone są z analizy generatywnej własności i mechanizmy komponentów interpretacyjnych, brane pod uwagę o tyle tylko, o ile wydaje się to niezbędne dla odsłonięcia uniwersalnej struktury składniowej. Wypada od razu zauważyć, że, ponieważ to realizacja struktur składniowych przez komponent morfologiczno-fonologiczny odpowiada za powierzchniarne zróżnicowanie języków i zaciemnienie struktury składniowej, cechy morfologiczne jednostek elementarnych — wbrew oczekiwaniom wynikającym z założenia ścisłej separacji komponentu składniowego — mogą w istocie odgrywać rolę w budowaniu struktur syntaktycznych (i tak się dzieje bardzo wyraźnie we wczesnym modelu minimalistycznym (zob. Chomsky 1995). Gdy zaś własności morfologiczne zostaną rozdzielone między różne etapy generowania i interpretacji wyrażań, jak w modelu Morfologii Rozdzielonej, konsekwencją może być odrzucenie jednokierunkowości przepływu informacji między składnią i komponentem morfologiczno-fonologicznym (zob. Starke 2011), który postuluje cykliczność przepływu informacji między składnią i komponentem morfologicznym; w odmiennym kontekście

zbliżoną zmianę perspektywy sugeruje Uriagereka (2012) w ramach proponowanego przez siebie modelu CLASH). O ile jednak relacje między mechanizmem składniowym i mechanizmami dokonującymi realizacji powierzchniowej nie są tak jednoznaczne, jak model Y mógłby sugerować, o tyle komponent semantyczny traktowany jest jako całkowicie odrębny moduł, odseparowany od składni przez komponent LF — miejsce, w którym ostatecznie dokonuje się przekształcenia struktur składniowych w sposób umożliwiający ich interpretację przez niezależne od składni mechanizmy.

Hipoteza autonomii składni, wraz z towarzyszącym jej modelem ludzkiej zdolności językowej, nie ma w generatywizmie statusu jedynie metodologicznego (o metodologicznym aspekcie autonomii składni w generatywizmie zob. Chomsky 1980): ujęcie języka jako „organu” poznawczego, a językoznawstwa jako należącego w istocie do psychologii, po zwrocie biolingwistycznym zaś szerzej — do biologii, oznacza, że hipotetyczna niezależność składni od pozostałych składników zdolności językowej i zdolności językowej od pozostałych władz poznawczych jest hipotezą dotyczącą rzeczywistej architektury umysłu, która powinna znaleźć ostatecznie umocowanie w empirycznych badaniach nad ludzkimi zdolnościami poznawczymi. Dotyczy to również daleko posuniętej modularności samej zdolności językowej zakładanej w teorii rządu i wiązania, w której odrębnym i niezależnym modułem przypisuje się odpowiedzialność za realizację wymogów teorii kontroli, teorii wiązania itp.: ich obecność winna być zgodna nie tylko z czysto wewnątrzteoretycznymi postulatami (jak było już na etapie GB), ale także z ogólnymi postulatami metodologicznymi minimalizmu i z rezultatami badań dyscyplin analizujących psychologiczne/biologiczne uwarunkowania ukształtowania zdolności językowej. Pierwszy aspekt — zgodność z postulatami wewnątrzteoretycznymi — prowadził do daleko idącego uproszczenia architektury zdolności językowej, w której nie mogło być miejsca dla wysoko wyspecjalizowanych modułów znanych z teorii rządu i wiązania: efekty przypisywane modułowi kontroli czy modułowi wiązania redukowane są do ogólnych mechanizmów składniowych (zatem do interakcji operacji przeniesienia i operacji zgody; o redukcji zjawiska kontroli zob. już Hornstein 2001; o redukcji — całkowitej lub częściowej — zjawisk należących pierwotnie do teorii wiązania zob. Hornstein 2006, Kayne 2002, Reuland 2011). Co więcej, autonomia składni jako taka w istocie została poddana w minimalizmie daleko idącemu ograniczeniu, jedną z głównych tez minimalistycznych jest bowiem tzw. Mocna Teza Minimalistyczna:

[Mocna Teza Minimalistyczna:] Język jest najlepszym rozwiązaniem zapewniającym czytelność obiektu na łącznikach (Chomsky 2000a: 96)¹.

¹ Cytowane w tekście przekłady nietłumaczonych tekstów źródłowych pochodzą od autora.

Podobnie jak hipoteza autonomii składni, Mocna Teza Minimalistyczna nie ma charakteru wyłącznie metodologicznego, ale odzwierciedlać ma naturę zdolności językowej i jej architektury (o roli metodologicznej Mocnej Tezy Minimalistycznej i jej interpretacji jako odnoszącej się do istotnych twierdzeń teorii zob. Martin i Uriagereka (2000)). Mocna Teza Minimalistyczna, wyznaczając kierunek poszukiwań adekwatnej analizy „najlepszego rozwiązania”, nadaje generatywnemu mechanizmowi składniowemu w istocie rolę służebną wobec wymogów łączników. Choć bowiem mechanizmy składniowe oraz specyfikacja cech jednostek elementarnych i powstających z nich struktur nadal formułowane są w kategoriach specyficznych dla składni — operacje przeniesienia i zgody są charakteryzowane jako relewantne dla tzw. cech formalnych — to jednak ostatecznie komponenty zewnętrzne wobec składni określają wymogi, jakie struktury składniowe winny spełnić. Autonomia składni ma zatem bardzo wąskie granice. Warto przy tym zauważyć, że Mocna Teza Minimalistyczna, wprowadzając pojęcie optymalności (charakteryzowane m.in. przy pomocy takich pojęć, jak prostota) jest równocześnie nawiązaniem i reinterpretacją stanowiska wczesnego generatywizmu, który postulował porównywanie pod analogicznymi względami *teorii* (gramatyk) — minimalizm natomiast formułuje takie oczekiwania nie wobec konstruktów teoretycznych, ale wobec samego badanego przedmiotu (zob. Hinzen 2012c: 110).

2. Miejsce semantyki w minimalizmie

Jedną z konsekwencji przyjęcia Mocnej Tezy Minimalistycznej jest faktyczne uznanie nieskładniowych właściwości wyrażeń/struktur językowych za istotne dla procesów składniowych, a w rezultacie włączenie własności interpretacyjnych — semantycznych i pragmatycznych — do minimalistycznej analizy zdolności językowej. Chociaż tendencja taka nie jest wiązana z otwartym kwestionowaniem autonomii składni, przejawia się wyraźnie w obecności wśród cech (rozumianych technicznie) oraz w katalogu ośrodków składniowych elementów, które, niezależnie od nazwy, faktycznie są cechami semantycznymi lub semantyczno-pragmatycznymi: negacja, topik, fokus, ośrodki czasu czy aspektu postulowane są jako należące do uniwersalnego wyposażenia składniowego komponentu zdolności językowej niezależnie od realizacji powierzchniowej. *Nomina non mutant rerum naturas*: analiza formalnosemantyczna w istotnej części dyktuje analizie składniowej, jakimi elementami powinna się posługiwać, jak winny być charakteryzowane i które z nich łączyć winny bliskie relacje; cechy formalne w coraz większym stopniu stają się jedynie odpowiednikami cech semantycznych. Wpływ analiz semantycznych na praktykę minimalistyczną widać wyraźnie w przyjmowanych za punkt wyjścia dalszych analiz strukturach składniowych oraz stosowanej niekiedy równoczesnej, z założenia składnikowej, analizie semantycznej (zob. np. Pylkkänen

(2008), gdzie węzły znacznika frazowego są wprost dekorowane wyrażeniami rachunku lambda, lub Ramchand (2008), z jednoznacznie interpretacyjnie motywowaną strukturą składniową) oraz poszukiwaniu semantycznych wyjaśnień dla zjawisk, które tradycyjnie generatywizm analizował w kategoriach czysto składniowych, jak ograniczenia wyspowe i bariery dla ekstrakcji (zob. np. Truswell (2011) oraz ogólnie o tendencji do wyjaśniania tego rodzaju zjawisk przez odwołanie do mechanizmów nieskładniowych Hornstein i in. (2007)). Charakterystyczne jednak, że sam Chomsky nie włącza takich rozwiązań do swoich propozycji, co najwyżej — jak w przypadku rozbudowanych struktur tzw. tradycji kartograficznej (Cinque 1999) — uznając, że być może przyjmowane przezeń struktury składniowe mogą być uznane za ekwiwalentne pod istotnymi dla analizy względami.

Włączenie zjawisk semantycznych do analizy minimalistycznej nie ma bowiem charakteru programowego, a nawet jest do pewnego stopnia procedurą „niekanoniczną” wobec oficjalnie głoszonej przez Chomsky’ego nie tylko niewłaściwości łączenia analizy składniowej/formalnej z semantyczną (choćby i „formalną”, to jednak nie w rozumieniu składni generatywnej), ale też przekonania o niemożliwości zbudowania ściśle naukowej teorii znaczenia. Już Chomsky (1955), choć nie odmawia semantyce miejsca w językoznawstwie, jednoznacznie separuje składnię i semantykę:

Teoria językoznawcza ma dwa główne działy, składnię i semantykę. Składnia jest badaniem formy językowej. [...] Semantyka [...] zajmuje się znaczeniem i odniesieniem wyrażen językowych. Jest zatem badaniem tego, jak narzędzie, którego struktura formalna i możliwości wyrażenia są przedmiotem badania składniowego, jest rzeczywiście używane we wspólnocie językowej (Chomsky 1955: 57).

Bezpośrednio po przytoczonym wyżej fragmencie formułuje Chomsky opinię, że teoria znaczenia nie spełnia elementarnych wymogów pozwalających na wykorzystanie jej w budowie teorii języka; podobne opinie wyrazi Chomsky (1957) i będą się one przewijać w dyskusji przedmiotu badania językoznawstwa i relacji składnia — semantyka niezmiennie, z okresem minimalistycznym włącznie, w którym ciągłość spojrzenia na semantykę jest także *explicite* podkreślana:

Jest możliwe, że język naturalny ma tylko składnię i pragmatykę; „semantykę” ma tylko w znaczeniu „badania tego, jak narzędzie, którego struktura formalna i możliwości wyrażenia są przedmiotem badania składniowego, jest rzeczywiście używane we wspólnocie językowej”, by przytoczyć najwcześniejsze sformułowanie w gramatyce generatywnej, sprzed czterdziestu lat, na które wpłynął Wittgenstein, Austin i inni (Chomsky 2000b: 132).

Koncentracja na formalnym aspekcie języka motywowana jest po części względami metodologicznymi — założeniem, że w badaniu aspektu formalnego można

i należy abstrahować od interpretacji — po części zaś wynika z przyjmowanego przez Chomsky'ego rozumienia teorii znaczenia, rozumienia, którego źródła wskazuje częściowo cytowany wyżej fragment.

3. Język przedmiotowy w semantyce formalnej: problem Fregego

Pierwszy ze wskazanych czynników, tj. uznanie badania formy wyrażeń/struktur językowych za poddające się adekwatnej analizie niezależnie od ich interpretacji, jest w istocie blisko związany z kierunkiem rozwoju logiki współczesnej i powinien raczej sprzyjać włączeniu badań formalnosemantycznych do nurtu generatywnego niż stać im na drodze. Zakładana bowiem w generatywizmie niezależność składni i interpretacji, charakteryzowanych przy pomocy odrębnego aparatu i polegających na całkowicie odrębnych mechanizmach, znajduje paralełę w przejętym z dwudziestowiecznej logiki aparacie semantyki formalnej. Jej główny nurt wywodzi się z propozycji Montague, którego wielokrotnie cytowane stwierdzenie:

Nie ma w mojej opinii żadnej istotnej różnicy teoretycznej między językami naturalnymi i sztucznymi językami logików; uważam za możliwe objęcie składni i semantyki języków obu rodzajów jedną, naturalną i matematycznie precyzyjną teorią (Montague 1970a: 373),

jednoznacznie wskazuje stanowisko teoretyczne leżące u podstaw współczesnej semantyki formalnej: język naturalny traktowany jest jako zinterpretowany system formalny, którego specyfikacja, wraz ze wskazaniem alfabetu, obejmuje reguły składniowe, określające sposób budowania wyrażeń języka, oraz reguły semantyczne, określające sposób interpretacji wyrażeń. Oba komponenty pozostają zasadniczo oddzielone — Chomsky (1955), proponując badanie formalnego aspektu wyrażeń językowych, w istocie proponuje zatem badanie składni jako *niezinterpretowanego* systemu formalnego, odrzucając komponent semantyczny (zob. Bach 1989). Teoria Montague ogranicza zasadniczo możliwe relacje między obu komponentami, definiując komponenty składniowy i semantyczny algebraicznie i postulując homomorficzność przekształcenia ze składni na semantykę. Istotne jednak, że wybór interpretacji oraz szczegóły specyfikacji obu komponentów dopuszczają rozwiązania z punktu widzenia przedmiotu analizy całkowicie arbitralne. W postępowaniu takim nie idzie semantyka formalna (i wczesny generatywizm, akceptując możliwość niezależnego badania komponentu składniowego) śladami Fregego, mimo oficjalnego uznania go za ojca logiki współczesnej, ale podąża w kierunku wskazanym przez Hilberta (zob. już Hilbert 1899). Wymiana listów między Fregem i Hilbertem świadczy o istotnie odmiennym pojmowaniu właściwego sposobu budowania systemu formalnego, który dla Hilberta jest „jedynie rusztowaniem” (*nur ein Fachwerk*), a elementy, których relacje określa, można

wybrać „w dowolny sposób” — przez punkty geometryczne można rozumieć np. „miłość, prawo, kominiarz”, a mimo to

jeśli [...] się przyjmie wszystkie moje aksjomaty jako [określające] relacje między tymi rzeczami, to moje twierdzenia, np. Twierdzenie Pitagorasa, są ważne także w odniesieniu do nich. Innymi słowy: każda teoria może być zastosowana do nieskończonej liczby systemów elementów podstawowych (Hilbert do Fregego *apud* Frege 1976: 67).

Frege widzi w dopuszczalności wielorakiej interpretacji błąd:

Jeśli sformułować jakiś aksjomat tak samo, łatwo można uwierzyć, że ma się do czynienia z tym samym aksjomatem. Zależy to jednak od znaczenia; a ono jest różne zależnie od tego, czy słowa „punkt”, „prosta” itd. są rozumiane w sensie geometrii euklidesowej czy w sensie szerszym (Frege do Liebmann; Frege 1976: 149).

Dla Fregego specyfikacja strony formalnej nierozłącznie związana jest z interpretacją systemu, którego kształt zależy przeto od poddawanego formalizacji zagadnienia, co wpływa w zasadniczy sposób także na rozumienie pojęcia wynikania logicznego (zob. dalej Blanchette 1996, 2007). Dalszy rozwój logiki formalnej idzie jednak drogą wytyczoną przez Hilberta. Niezależność interpretacji od komponentu składniowego jest charakterystyczną cechą dwudziestowiecznych formalizmów i podnoszona bywa jako ich wielka zaleta — jak podsumowuje swoje omówienie modeli Kripkego Burgess (2011: 138), wprawdzie nie są one w stanie same przez się rozwikłać zagadnień filozoficznych, jednak jest zaletą teorii, że „nie będąc związaną z żadnym konkretnym pojmowaniem natury modalności, może być zastosowana do wielu”. Główny nurt semantyki formalnej w językoznawstwie — teoriomodelowe podejście Montague — cechuje neutralność co do wyboru wszystkich szczegółów specyfikacji obu komponentów, od alfabetu poczynając, na interpretacji kończąc, choć pod wieloma innymi względami podejście to realizuje szereg zasadniczych założeń Fregego, łącznie z fundamentalnym znaczeniem rozróżnienia między składnikami nasyconymi i nienasyconymi (formalizowanego w rachunku lambda), które sam Frege uważa za podstawowe dla wyjaśnienia możliwości łączenia elementów w jedną, znaczącą całość: „W logice wszelkie łączenie zdaje się polegać na uzupełnianiu czegoś niedopełnionego” (Frege 1977: 135). O ile jednak takie cechy gramatyki Montague bywają modyfikowane lub kwestionowane *in toto* (jak w propozycjach rezygnacji z redukcji beta jako sposobu formalizacji łączenia wyrażeń w interpretacji poskładniowej (zob. np. Pietroski 2005, 2006, 2010, 2011; Pietroski i Hornstein 2009), o tyle idea formalizacji w stylu Hilberta zdaje się powszechnie akceptowana, zarówno jako część tradycyjnego już podejścia do określania systemu formalnego, jak i jako gwarancja jak najdalej posuniętej neutralności aparatu — niezależność od wyboru interpretacji zdaje się najlepszym dowodem takiej cechy.

Z generatywnego punktu widzenia takie postępowanie nie jest jednak jednoznacznie pożądane. Wprawdzie możliwość przedstawienia bardzo różnych sposobów formalizacji, zmieniających specyfikację bądź komponentu składniowego, bądź komponentu semantycznego, otwiera pole do modelowania badanych zjawisk w sposób zależny jedynie od potrzeb badania, jednak brak zewnętrznych ograniczeń dla takiej swobody nie pozwala rozstrzygnąć pytań istotnych dla dążącej do adekwatności wyjaśniającej teorii generatywnej — pytań w rodzaju tych, jakie formułuje Chomsky (1980):

Żałóśmy, że zgadzamy się, że gramatyka powinna określać na poziomie reprezentacji znaczenia pewien rodzaj reprezentacji struktury kwantyfikacyjnej. Czy zapis ma znaczenie? Czy reprezentacja powinna zawierać kwantyfikatory lub zmienne, czy powinna być w notacji bez zmiennych, czy też pytanie nie ma znaczenia empirycznego? (Chomsky 1980: 63).

Przed semantyką formalną, która chce stać się częścią minimalistycznej teorii języka, stoi przeto wyzwanie, które można nazwać problemem Fregego: zbudowanie systemu formalnego, który — jak ideografia Fregego była zaprojektowana tak, by ująć wyłącznie cechy istotne dla wnioskowania — poddawałby się interpretacji jako adekwatna formalizacja poskładniowego komponentu interpretacyjnego. Język przedmiotowy powinien zatem zostać zdefiniowany w sposób odzwierciedlający realne własności semantycznego komponentu zdolności językowej, a równocześnie wolny od takich cech formalizmu, które byłyby tylko przygodnymi własnościami modelu. Pytanie o to, czy jest miejsce dla zmiennych w modelowaniu kwantyfikacji w języku naturalnym, dla generatywnego badania języka *nie może* być bez znaczenia empirycznego.

4. Interpretacja metajęzyka: problem Tarskiego

Drugi czynnik sprzyjający separacji składni i semantyki oraz przesądzający o negatywnym stanowisku Chomsky'ego wobec możliwości zbudowania naukowej teorii znaczenia to sformułowane przezeń w przytoczonych wyżej fragmentach pojmowanie semantyki w kategoriach, których źródła są jasno wskazane — obok siebie, mimo wszystkich dzielących ich różnic, wymienieni są Wittgenstein i Austin. Charakterystyczne bowiem, że już Chomsky (1955) widzi semantykę jako badającą język jako „narzędzie“, a jej zadaniem jest analiza, jak takie narzędzie „jest rzeczywiście używane we wspólnocie językowej“:

Pomyśl o narzędziach w skrzynce z narzędziami: jest tam młotek, są obcęgi, piła, śrubokręt, całówka, garnek do kleju, klej, gwoździe i śruby. — Jak różne są funkcje tych przedmiotów, tak różne są też funkcje wyrazów. (A podobieństwa znajduje się tu i tam). (Wittgenstein 1953 § 11).

Sceptycyzm wobec możliwości zbudowania naukowej teorii tak widzianego znaczenia jest przeto zrozumiały, wymagałaby przecież stworzenia teorii wszelkich możliwych gier językowych. Nie można oczekiwać spełnienia standardowych postulatów metodologicznych przez badanie zjawisk językowych, na które „można patrzeć jak na stare miasto: plątanina uliczek i placów, starych i nowych domów, domów z dobudówkami z różnych czasów” (Wittgenstein 1953 § 18). Co więcej — i co zasadniczo sprzeczne z orientacją metodologiczną generatywizmu — próba taka wymagałaby również koncentracji na badaniu języka jako „formy życia” i badaniu znaczenia jako sposobów użycia wyrażen językowych: „Tego, jak słowo funkcjonuje, nie da się zgadnąć. Trzeba *przyjrzeć się* jego zastosowaniu i z tego się uczyć” (Wittgenstein 1953 § 340). Sfera ta należy zaś do E-języka, który metodologia generatywna z zasady wyłącza z zakresu badania. Trzeba jednak przy tym od razu zauważyć, że semantyka formalna w postaci uprawianej na gruncie generatywnym nie stawia sobie podobnych celów: choć część badań semantycznych istotnie zmierza w kierunku analizy funkcjonowania języka (badania interakcji językowych, podobnie jak niektóre nurty logiki filozoficznej), główny nurt semantyki formalnej uwzględnia wpływ użycia języka na interpretację jedynie przez wprowadzanie dodatkowych elementów aparatu formalnego — takich jak zmienne kontekstowe, indeksy kontekstowe (skądinąd rozważane już przez Montague, zob. np. Montague 1970b) oraz formalizacja presupozycji lub znaczeń ekspresywnych. Wprowadzenie takich parametrów — jak i uwzględnienie niedeskrptywnych składników znaczenia — nie oznacza wszakże, by celem było stworzenie teorii obejmującej wszelkie możliwe użycia języka, choć otwarte pozostaje pytanie, czy elementy takie powinny być obecne już w interpretacji następującej bezpośrednio po operacjach składniowych, czy też raczej należą do odrębnego komponentu interpretacyjnego.

Obok zastrzeżeń dotyczących zakresu badania semantycznego, które może zostać wyznaczone przecież w sposób odmienny od rozumienia, jakie prezentuje Chomsky (1955), i którego aparat formalny nie przesądza o zaangażowaniu w badania realizacji językowej, zachowując przydatność jako narzędzie badania kompetencji językowej w generatywnym ujęciu, Chomsky wysuwa jednak także zastrzeżenia dotyczące możliwości zastosowania aparatu pojęciowego standardowo wiążanego z interpretacją teoriomodelową, w szczególności korespondencyjnego rozumienia prawdy oraz roli pojęcia odniesienia. Najogólniej rzecz ujmując, generatywizm — w tym także minimalizm — nie widzi miejsca dla relacji świat — język jako relacji wyjaśniających w teorii znaczenia wyrażen językowych, jeśli ma ona, zgodnie z założeniami całego nurtu, zajmować się I-językiem. Kwestionując czysto odniesieniową koncepcję języka, zgadza się Chomsky znów z Wittgensteinem, choć z odmiennych powodów i dochodząc do odmiennych konkluzji; obaj jednak za nietrafny uważają

pogląd, wedle którego „wyrazy języka nazywają przedmioty — zdania są splo-tami takich nazwań” (Wittgenstein 1953 § 1). Dla pojęcia odniesienia takiego, jakie występuje w większości rozważań filozofii języka dwudziestego wieku, nie ma zdaniem Chomsky’ego miejsca w naukowej teorii języka:

Nie jest w ogóle jasne, że teoria języka naturalnego i jego użycia pociąga za sobą re-lacje „denotacji” i „prawdziwości w odniesieniu do” w znaczeniu jakkolwiek zbliżo-nym do technicznej teorii znaczenia (Chomsky 2000b: 130).

Zastosowanie pojęcia prawdy w ujęciu Tarskiego, wraz z semantycznym poję-ciem spełniania, wydaje się więc na gruncie generatywnym wątpliwe; w jego koncepcji semantyki kluczową rolę odgrywają przecież pojęcia, którym gene-ratywizm odmawia roli wyjaśniającej w teorii znaczenia:

Semantyka jest dyscypliną, która — luźno mówiąc — *zajmuje się pewnymi relacjami między wyrażeniami języka a przedmiotami* (bądź „stanami rzeczy”), *do których te wyrażenia „się odnoszą”*. Jako typowe przykłady pojęć semantycznych możemy wy-mienić pojęcia: *oznaczania, spełniania i definiowania...* (Tarski 1944: 236).

Sam Tarski, jak wiadomo, również nie przewidywał objęcia języka naturalnego aparatem ufundowanym na wysuniętych przez siebie propozycjach — jego propozycje są przecież propozycjami odnoszącymi się do rygorystycznie zdefi-niowanych języków przedmiotowych. Nie wynika z tego jednak, by semantyka teoriomodelowa musiała być co do zasady niezgodna z celami i metodami ge-neratywnej analizy języka. Charakteryzacja podstawowych pojęć semantycz-nych, jakiej dokonuje Tarski, może skłaniać do uznania, że relacja język — świat jest konstytutywna dla posługującej się tym aparatem pojęciowym se-mantyki. Tak jednak nie jest i sam Chomsky niejednokrotnie szkicuje alterna-tywne ujęcie; rozważając odpowiedzi na pytanie, czy w teorii znaczenia po-winno być miejsce dla „relacji między pewnymi wyrażeniami i przedmiotami zewnętrznymi”, widzi więcej niż jedną możliwość:

Tak się zwykle zakłada, ale musimy starannie rozróżnić dwa warianty: (1) rzeczy w świecie, lub (2) rzeczy w pewnego rodzaju modelu mentalnym... (Chomsky 2000b: 129).

Stanowisko przyjmujące opcję (2) byłoby, wedle Chomsky’ego, jako internali-styczne badanie znaczenia w pełni akceptowalne jako element generatywnej teorii języka. Mogłoby się więc wydawać, że w istocie konieczna jest jedynie odmienna niż standardowo przyjmowana koncepcja dziedziny kwantyfikacji, obejmująca „przedmioty” w „modelu mentalnym” zamiast przedmiotów w świecie zewnętrznym — i semantyka teoriomodelowa może być bez trudu przyjęta jako część generatywnej teorii języka. Byłoby to jednak nieuzasadnio-ne uproszczenie. Przytoczony wyżej fragment Chomsky’ego trzeba widzieć

w szerszym kontekście generatywnych rozważań nad modelowaniem znaczenia w języku naturalnym i nad miejscem zdolności językowej wśród zdolności poznawczych; wskazuje jedynie kierunek, ale — by semantyka istotnie pod tym względem spełniła postulaty generatywne — nie wystarczy samo przecięcie więzi język/słowa — świat. Problemem, który można nazwać problemem Tarskiego dla generatywnej semantyki formalnej, jest określenie empirycznie adekwatnej interpretacji całego aparatu metajęzykowego.

5. Leksykon i rachunek pojęć: problem Leibniza

Dojrzała faza minimalizmu wyostreza problemy wskazane wyżej, otwierając zarazem na nowo szereg pytań dotyczących miejsca zdolności językowej. Hauser i in. (2002) przynoszą rozważania dotyczące miejsca zdolności językowej człowieka wśród ogółu ludzkich zdolności poznawczych nie tylko z punktu widzenia zdolności do nabywania języka, ale także z punktu widzenia ewolucji zdolności językowej. Podział na zdolność językową *sensu stricto*, ograniczoną do specyficznie językowych cech i/lub mechanizmów, oraz zdolność językową *sensu largo*, obejmującą w szczególności komponenty odpowiedzialne za interpretację semantyczną i motoryczną, choć pozornie może się wydawać bliski tradycyjnej koncepcji autonomicznej składni, w istocie stawia pytania o miejsce zdolności językowej i jej właściwą analizę w dużym stopniu na nowo. Rozwiązanie problemu Platona, tj. odpowiedź na pytanie o możliwość nabywania języka, wraz z dobrze znanym argumentem za wrodzonym charakterem Gramatyki Uniwersalnej z niewystarczalności bodźca (zob. ostatnio Berwick i in. 2011), powinno być niesprzeczne z rozwiązaniem problemu Darwina — pytaniem o naturę ewolucyjnie bardzo młodej zdolności językowej, zbyt młodej, by można było dla niej postulować wysoce specyficzne mechanizmy zakładane w teorii rządu i wiązania. Minimalistyczne propozycje wynikające z ogólnych założeń metodologicznych dotyczących budowy teorii spotykają się w ten sposób z wymogami adekwatności empirycznej wobec biolingwistycznej analizy języka. Otwarcie pytania o stopień niezależności mechanizmów językowych od mechanizmów postulowanych dla pozostałych władz poznawczych, o precyzyjną charakterystykę zdolności językowej w obu znaczeniach i o ich wzajemne relacje, choć pozornie następuje w tradycyjnych ramach tezy o autonomii, wskazuje na konieczność ponownego przemyślenia podstawowych założeń teorii (zob. obszernie Hinzen (2006) o konsekwencjach minimalizmu dla teoretyzowania relacji między zdolnością językową i pozostałymi władzami poznawczymi). Dotyczy to również Mocnej Tezy Minimalistycznej — oprócz takiego ujęcia relacji między składnią i komponentami nieskładniowymi, które zakłada, że składnia odpowiada na ich wymogi, proponować można odmienne rozwiązanie, w którym to raczej składnia nadaje faktyczny kształt operacjom komponentów nieskładniowych i wyznacza, w znacznym stopniu lub nawet

całkowicie, granice ich autonomii (zob. Hinzen 2011, 2012a, 2012c). Tradycyjna charakterystyka generatywnej adekwatności opisowej i adekwatności wyjaśniającej musi ulec zmianie: uwzględniać one powinny, stosownie do wstępnego zakreszenia granic zdolności językowych w obu znaczeniach, wymogi teorii wyjaśniających dla innych władz poznawczych. Oznacza to, z jednej strony, konieczność budowania modeli, które mogą stanowić pomost z badaniami nad biologicznymi podstawami języka i innych władz poznawczych — taki kierunek postulują Poeppel i Embick (2005), zwracający uwagę na zasadnicze niezgodności między aparatem teoretycznym minimalizmu i aparatem pojęciowym nauk empirycznych badających biologiczne podstawy procesów poznawczych; z drugiej strony, poszczególne komponenty teorii języka stają wobec inaczej sformułowanych zadań. Problem Fregego i problem Tarskiego winny w świetle rozważań biolingwistycznych znaleźć rozwiązanie nie tylko oświetlające relacje między zdolnością językową i pozostałymi zdolnościami poznawczymi, ale także poddające się interpretacji w kategoriach badań empirycznych.

Wymagające nowej analizy relacje składnia — komponenty nieskładniowe obejmują nie tylko relację składnia — poskładniowy komponent interpretacyjny, ale także relacje między leksykonem i składnią. Niedostatkiem formalnosemantycznej analizy języka jest sposób, w jaki opisywane są elementarne jednostki poddawane operacjom składniowym, tj. elementy „słownika”. Analiza minimalistyczna, podejmując do pewnego stopnia i modyfikując rozwiązania Semantyki Generatywnej, postuluje niejednokrotnie dekompozycję składniową jednostek leksykalnych (w tradycji Montague już Dowty (1979), rozwiązania zaś minimalistyczne częściowo inspirowane są analizami semantycznymi w tym duchu, częściowo propozycjami składniowymi, jakie wysuwają Hale i Keyser (2002); z minimalistycznego punktu widzenia krytykę dotychczasowych rozwiązań przedstawia Hinzen (2012b)). Problemem jest jednak semantyczna analiza jednostek elementarnych, a podstawowym źródłem trudności dla semantyki formalnej jest tu ścisła separacja stałych pozalogicznych od stałych logicznych i przyjęcie niezmienniczości za kryterium przynależności do tej ostatniej grupy:

W skrajnym przypadku rozważalibyśmy klasę *wszystkich* jednojednoznacznych przekształceń całej przestrzeni czy uniwersum dyskursu, czy też „świata” na siebie. Jakaż by to była nauka, która zajmuje się pojęciami niezmienniczymi wobec tej najobszerniejszej klasy przekształceń? W tym wypadku będziemy mieli jedynie kilka pojęć, wszystkie o charakterze bardzo ogólnym. Sugeruję, że są to pojęcia logiczne i proponuję, abyśmy nazywali pojęcie *logicznym*, jeśli jest ono niezmiennicze wobec wszystkich jednojednoznacznych odwzorowań świata na siebie (Tarski 1986: 459).

Ceną za charakteryzację pojęć logicznych w kategoriach niezmienniczości jest wprowadzanie stałych pozalogicznych — a jako takie semantyka formalna formalizuje najbardziej elementarne elementy alfabetu, które mogą odpowiadać elementom „leksykonu” — arbitralną decyzją podejmowaną przy specyfikacji systemu formalnego, oraz daleko posunięty brak determinacji ich interpretacji (zob. dalej Sher (2000) oraz Sher (2001) o relacji semantyka — logika). Daleko odbiega takie ujęcie od zamierzonej przez Leibniza *lingua characteristica*, umożliwiającej adekwatną analizę pojęć. Frege widział różnicę między własnym projektem i projektem Leibniza przede wszystkim w kierunku analizy:

Nie wychodzę zatem od pojęć, składając z nich myśli czy sądy, lecz na odwrót: rozszczepiając myśl dochodzę do jej składników. Pod tym względem moja ideografia logiczna różni się od podobnych pomysłów Leibniza i jego następców, mimo nie najszcześliwiej dobranej przeze mnie nazwy „*Begriffsschrift*” (Frege 1919: 134).

Fregego ideografia sądów nie rozdziela jednak jeszcze specyfikacji formalnej od interpretacji, nie inspirował się też programem erlangenkim Kleina w charakteryzacji stałych. Współczesna semantyka formalna, korzystając z aparatu dwudziestowiecznej logiki, staje tu wobec kolejnego wyzwania. Problemem Leibniza, bo tak można go nazwać ze względu na jego prace nad algebrą pojęć, jest dla minimalistycznej semantyki formalnej proponowanie adekwatnej teorii interakcji składni z komponentem/komponentami *przedskładniowymi* — tym/tymi, w których dokonuje się operacji na pojęciach i z których zdolność językowa czerpie zasoby tradycyjnie zamykane w niepoddającej się dalszym wyjaśnieniu kategorii „leksykonu”. Sformułowane wyżej problemy nie dowodzą jednak nieprzydatności semantyki formalnej w minimalistycznym projekcie budowy teorii zdolności językowej; wskazują natomiast potrzebę podjęcia analiz podstawowych narzędzi semantycznych z czysto minimalistycznego punktu widzenia. Potrzeba to tym większa, że znaczenie teorii semantycznej w minimalizmie jest — wobec zmian w ujmowaniu roli i miejsca składni w teorii zdolności językowej — coraz większe.

Bibliografia

- BACH E. (1989), *Informal Lectures on Formal Semantics*. New York: State University of New York Press.
- BERWICK R., PIETROSKI P., YANKAMA B., CHOMSKY N. (2011), *Poverty of the Stimulus Revisited*, „Cognitive Science” 35, 1207–1242.
- BLANCHETTE P. (1996), *Frege and Hilbert on Consistency*, „The Journal of Philosophy” 93(7), 317–336.

- BLANCHETTE P. (2007), *Frege on Consistency and Conceptual Analysis*, "Philosophia Mathematica" 15, 321–346.
- BURGESS J. (2011), *Kripke Models*, [w:] Berger A. (red.), *Saul Kripke*. Oxford: Oxford University Press, 119–140.
- CHOMSKY N. ([1955] 1975), *The Logical Structure of Linguistic Theory*. Chicago: Chicago University Press.
- CHOMSKY N. (1957), *Syntactic structures*. The Hague/Paris: Mouton.
- CHOMSKY N. (1980), *Rules and representations*. New York: Columbia University Press.
- CHOMSKY N. (1993), *A Minimalist Program for Linguistic Theory*, [w:] Hale K., Keyser S. J. (red.), *The View from Building 20: Essays in Linguistics in Honor of Sylvain Bromberger*. Cambridge, Massachusetts: MIT Press, 1–52.
- CHOMSKY N. (1995), *The Minimalist Program*. Cambridge, Massachusetts: MIT Press.
- CHOMSKY N. (2000a), *Minimalist Inquiries: the framework*, [w:] Martin R., Michaels D., Uriagereka J. (red.), *Step by Step: Minimalist Papers in Honor of Howard Lasnik*. Cambridge, Massachusetts: MIT Press, 89–155.
- CHOMSKY N. (2000b), *New Horizons in the Study of Language and Mind*. Cambridge: Cambridge University Press.
- CHOMSKY N. (2001), *Derivation by Phase*, [w:] Kenstowicz M. (red.), *Ken Hale: A life in language*. Cambridge, Massachusetts: MIT Press, 1–52.
- CHOMSKY N. (2005), *Three factors in language design*, "Linguistic Inquiry" 36, 1–22.
- CHOMSKY N. (2007), *Approaching UG from Below*, [w:] Riemsdijk H. van, Hulst H. van der (red.), *Interfaces + Recursion = Language?* Berlin: Mouton de Gruyter, 1–29.
- CHOMSKY N. (2008), *On Phases*, [w:] Freidin R., Otero C., Zubizarreta M. L. (red.), *Foundational Issues in Linguistic Theory. Essays in Honor of Jean-Roger Vergnaud*. Cambridge, Massachusetts: MIT Press, 133–167.
- CINQUE G. (1999), *Adverbs and functional heads: a cross-linguistic perspective*, Oxford: Oxford University Press.
- DOWTY D. (1979), *Word meaning and Montague grammar. The semantics of verbs and times in Generative Semantics and in Montague's PTQ*. Dordrecht: Reidel.
- FREGE G. ([1919] 1977), *Szkic dla Darmstaedtera*, [w:] Frege G., *Pisma semantyczne* (przeł. B. Wolniewicz). Warszawa: PWN, 133–139.
- FREGE G. (1976), *Wissenschaftlicher Briefwechsel. Herausgegeben, bearbeitet, eingeleitet und mit Anmerkungen versehen von Gottfried Gabriel, Hans Hermes, Friedrich Kambartel, Christian Thiel, Albert Veraart*. Hamburg: Felix Meiner Verlag.
- HALE K., KEYSER S. J. (2002), *Prolegomenon to a Theory of Argument Structure*. Cambridge, Massachusetts: MIT Press.
- HAUSER M. D., CHOMSKY N., FITCH W. T. (2002), *The Faculty of Language: What Is It, Who Has It, and How Did It Evolve?*, "Science" 298, 1569–1579.
- HILBERT D. (1899), *Grundlagen der Geometrie*. Leipzig: Teubner.
- HINZEN W., POEPEL D. (2011), *Semantics between cognitive neuroscience and linguistic theory: Guest editors' introduction*, "Language and Cognitive Processes" 26, 1297–1316.
- HINZEN W. (2006), *Mind Design and Minimal Syntax*. Oxford: Oxford University Press.

- HINZEN W. (2011), *Emergence of a Systemic Semantics through Minimal and Under-specified Codes*, [w:] Di Sciullo A. M., Boeckx C. (red.), *The Biolinguistic Enterprise. New Perspectives on the Evolution and Nature of the Human Language Faculty*. Oxford: Oxford University Press, 417–439.
- HINZEN W. (2012a), *Narrow syntax and the language of thought*, “Philosophical Psychology”, doi: 10.1080/09515089.2011.627537.
- HINZEN W. (2012b), *Syntax in the atom*, [w:] Werning W. M., Hinzen W., Machery E. (red.), *Handbook of Compositionality*. Oxford: Oxford University Press, 351–370.
- HINZEN W. (2012c), *Minimalism*, [w:] Kempson R., Fernando T., Asher N. (red.), *Philosophy of linguistics*. Amsterdam: Elsevier, 93–141.
- HORNSTEIN N. (2001), *Move! A Minimalist theory of construal*. Oxford: Blackwell.
- HORNSTEIN N. (2006), *Pronouns in a Minimalist Setting*, [w:] Kazanina N., Minai U., Monahan P. J., Taylor H. L. (red.), “University of Maryland Working Papers in Linguistics” 14, University of Maryland, 47–80.
- HORNSTEIN N., LASNIK H., URIAGEREKA J. ([2003] 2007), *The dynamics of islands: Speculations on the locality of movement*, “Linguistic Analysis” 33, 149–175.
- KAYNE R. (2002), *Pronouns and Their Antecedents*, [w:] Epstein S. D., Seely T. D. (red.), *Derivation and explanation in the minimalist program*. Oxford: Blackwell, 133–166.
- MARTIN R., URIAGEREKA J. (2000), *Introduction: Some Possible Foundations of the Minimalist Program*, [w:] Martin R., Uriagereka J. (red.), *Step by Step: Essays on minimalist syntax in honor of Howard Lasnik*. Cambridge, Massachusetts: MIT Press.
- MONTAGUE R. (1970a), *Universal grammar*, “Theoria” 36, 373–398.
- MONTAGUE R. (1970b), *Pragmatics and Intensional Logic*, “Synthese” 22, 68–94.
- PIETROSKI P. (2005), *Events and Semantic Architecture*. Oxford: Oxford University Press.
- PIETROSKI P. (2006), *Interpreting concatenation and concatenates*, “Philosophical Issues” 16(1), 221–245.
- PIETROSKI P. (2010), *Concepts, meanings and truth: First nature, second nature and hard work*, “Mind and Language” 25(3), 247–278.
- PIETROSKI P. (2011), *Minimal Semantic Instructions*, [w:] Boeckx C. (red.), *Oxford Handbook of Linguistic Minimalism*. Oxford: Oxford University Press, 472–497.
- PIETROSKI P., HORNSTEIN N. (2009), *Basic Operations*, “Catalan Journal of Linguistics” 8, 113–139.
- POEPEL D., EMBICK D. (2005), *Defining the relation between linguistics and neuroscience*, [w:] Cutler A. (red.), *Twenty-first century psycholinguistics: Four cornerstones*. Routledge, 103–118.
- PYLKKÄNEN L. (2008), *Introducing Arguments*. Cambridge, Massachusetts: MIT Press.
- RAMCHAND G. (2008), *Verb Meaning and the Lexicon: A First Phase Syntax*. Oxford: Oxford University Press.
- REULAND E. (2011), *Anaphora and Language Design*. Cambridge, Massachusetts: MIT Press.
- SHER G. (1996), *Semantics and Logic*, [w:] Lappin S. (red.), *The Handbook of Contemporary Semantic Theory*. Oxford: Blackwell, 509–535.
- SHER G. (2000), *The Logical Roots of Indeterminacy*, [w:] Sher G., Tieszen R. (red.), *Between Logic and Intuition*. Cambridge: Cambridge University Press.

- STALMASZCZYK P. (2011), *Historyczne i metodologiczne uwarunkowania współczesnego generatywizmu*, „Linguistica Copernicana” 1, 15–32.
- STARKE M. (2011), *Towards an elegant solution to language variation: Variation reduces to the size of lexically stored trees*, niepublikowany maszynopis, Center for Advanced Study in Theoretical Linguistics, University of Tromsø.
- TARSKI A. ([1944] 1995), *Semantyczna koncepcja prawdy i podstawy semantyki*, [w:] *Pisma filozoficzno-logiczne*, t. 1: *Prawda* (przeł. J. Zygmunt). Warszawa: PWN, 228–282.
- TARSKI A. ([1986] 2001), *Czym są pojęcia logiczne*, [w:] *Pisma filozoficzno-logiczne*, t. 2: *Metalogika* (przeł. J. Zygmunt). Warszawa: PWN, 446–466.
- TRUSWELL R. (2011), *Events, Phrases and Questions*. Oxford: Oxford University Press.
- URIAGEREKA J. (2012), *Spell-Out and the Minimalist Program*. Oxford: Oxford University Press.
- WITTGENSTEIN L. ([1953] 2000), *Dociekania filozoficzne* (przeł. B. Wolniewicz). Warszawa: Wydawnictwo Naukowe PWN.

Indeks

- Ajdukiewicz Kazimierz, 8, 11, 71–73, 76, 78–91, 94–95, 116–119, 128, 133, 136–137, 148, 150, 158–159, 208, 211
- Akt illokucyjny, 50–51, 57–64, 70, 74–75
- Akt społeczny, 11, 50–52, 54–60, 64, 68, 70–71, 75
- Aluzja, 136, 140, 142, 148–154, 156, 158
- Anomalie semantyczne, 96, 105
- Argument Sticha, 5, 235, 250
- Austin John, 11, 50–51, 58–61, 63, 74–76, 106, 162, 169, 254, 267, 275, 278
- Behawioryzm, 78
- Biolingwistyka, 14, 271
- Cassirer Ernst, 8, 29–35, 47–49
- Chomsky Noam, 12, 14, 105, 107, 257, 266, 271–273, 275–276, 278–280, 283–284
- Dummett Michael, 108, 110, 122, 129, 131, 133, 204, 206–207, 211–212
- Dynamiczna Logika Pluralna (DPL), 185, 190, 194
- Dynamiczna Logika Predykatów (DPL), 185, 190
- Dynamiczna Logika Typów (DLT), 185, 198
- Dyrektywa znaczeniowa, 78, 80
- Fenomenologia, 15, 18–19, 26, 101
- Formy symboliczne, 11, 25, 29–31, 47–48
- Frege Gottlob, 10–12, 61, 76, 91, 107, 112, 114, 116–117, 125, 127, 133–134, 146, 159, 211, 271, 276–278, 282–284
- Gadamer Hans-Georg, 8, 10, 15, 19, 22–24, 26–28, 262, 268
- Geach Peter T., 204, 206, 208, 211–212, 250
- Generatywizm, 271, 273–276, 279–280, 286
- Gramatyka kategoryalna, 91, 204, 210, 212
- Gramatyka Montague, 198, 271, 277
- Grice, 9, 11, 106, 136, 141–143, 149, 157, 159, 162, 166, 168, 173–174, 176, 240, 262, 268
- Gry agonalne, 252, 259
- Gry językowe, 9, 26, 252–254, 257, 262, 265–267, 279
- gry kooperacji, 252, 259–260, 262
- Heidegger Martin, 15–18, 22, 28
- hermeneutyka, 10, 15, 18, 25–26, 28, 268
- Hintikka Jaakko, 11, 107, 202, 252, 256–258, 262, 264–266, 268–269
- Husserl Edmund, 8, 11, 15, 18, 22, 24, 50, 101, 107
- Implikowanie konwersacyjn, 11, 136, 141, 143
- Indeksyzm, 161
- Interpretacja, 5, 11, 18, 24–25, 28, 71, 78, 85, 97, 100, 104–106, 110, 113, 120, 136, 138, 144–145, 147–149, 152, 154–159, 162, 166–168, 174–176, 182, 185, 186, 188, 190–193, 198, 201, 204, 214, 217–218, 220–222, 225, 227, 238, 262–263, 266, 272, 274, 276–279, 281–283
- Ironia, 11, 86, 136, 140–141, 143, 147, 152–155, 208
- Kategoria semantyczna nazwy, 204
- konstruktywizm komunikacyjny, 29–31, 49
- Kontekst, 9, 11–12, 52, 98, 100, 102, 104–105, 125, 129, 131, 138, 140–144, 146–147, 154, 162, 163–166, 168–175, 177–178, 181–182, 185–187, 190–198, 221, 230, 239, 241–242, 248
- Kontekstualizm, 5, 11, 161, 166–174, 178, 182, 266
- Kontynentalna filozofia języka, 10, 15–16, 19, 27

- Langacker Ronald, 11, 29–30, 37–46, 48
- Malinowski Bronisław, 29, 33–37, 49
- Merleau-Ponty Maurice, 8, 10, 15, 19–21, 24, 26, 28, 256, 263, 269
- Metafora, 11, 27, 87, 105–106, 136–140, 142–144, 146–153, 155–158, 178, 209, 239, 253, 262
- Metasemantyka, 235
- Metonimia, 136, 139–140, 149, 155
- Minimalizm semantyczny, 161, 166, 172
- Modele wielowymiarowe, 185, 200
- Montague Richard, 11–12, 107, 185, 197, 202, 271, 276–277, 279, 282, 284–285
- Nastawienia sądzeniowe, 111, 113, 120, 122–123
- Nazwa własna, 82, 90–91, 98, 103, 189–190, 192–193, 207–208, 211, 235–237, 240–242, 245–247
- Normatywność znaczenia, 5, 12, 213, 219, 227–228, 230, 232
- Nośniki wartości logicznej, 111, 128–132
- Obiecywanie, 50–51, 60, 65, 67–71, 73
- Okazjonalność, 161, 165, 168, 179
- Paradygmat strategiczny, 12, 252, 265, 266
- Podmiot, 20–23, 26, 30–31, 35–39, 52–54, 97, 99–100, 104, 107–108, 115–116, 120, 162, 204–206, 210–211, 213, 218, 224, 227, 229, 243–244
- Poznanie oderwane, 11, 29–30, 38, 40, 44
- Poznanie zaangażowane, 29, 37–40
- Prawda, 9, 17, 27, 33, 43, 62, 70–71, 78, 87, 96–97, 107, 109–111, 113, 121–125, 127–128, 130–132, 136, 144–145, 147–150, 155, 158–160, 165, 179, 203, 216–217, 227, 230–231, 246, 279–280, 286
- Predykat, 12, 82, 90–92, 98, 119–121, 125, 185, 188, 191–192, 195, 198, 204–206, 208, 210, 212
- Preskryptywność, 223, 227, 232
- Program Minimalistyczny, 5, 12, 271
- Przedjęzykowy poziom doświadczenia, 96, 98
- Przekonanie, 16, 27, 55–56, 85, 91, 111–114, 121, 136–137, 141–143, 145, 147–156, 159, 214, 226–227, 239, 247, 261, 275
- Przyczynowa teoria nazw, 235–236
- Quine Willard Van Orman, 11, 78, 83–84, 89, 95, 106, 112, 122–123, 127–128, 134, 150, 159
- Racja, 5, 21, 112–113, 133, 169, 177, 182, 213, 215–220, 223–224, 226–229, 231–232
- Radykalna subiektyfikacja, 29, 40, 45
- Referencja, 5, 12, 96, 102, 108, 185, 198, 200, 244
- Reinach Adolf, 5, 8, 11, 50–60, 64–76
- Ricouer Paul, 15
- Sąd logiczny, 5, 11, 111–112, 115, 133, 256
- Searle John, 10–11, 13, 50–51, 56–65, 70, 73–77, 106, 143, 158, 160–161, 168–169, 171, 176, 177–179, 254, 269
- Sellars Wilfrid, 11, 78, 86–87, 95, 122–123, 135
- Semantyka formalna, 12, 108, 257, 271, 276–279, 281–283
- Semantyka teoriogrowa, 252, 256–257, 262, 264–265, 267
- Sens, 16, 18–22, 24–28, 31–33, 35, 56, 72, 92, 95–96, 101, 104–107, 111, 125–127, 137–138, 146, 153, 161, 180–181, 183, 207, 211, 257, 266, 281
- Sensowność, 78, 88–89, 104–106, 137, 266
- Standard, 213, 222–225
- Stich Stephen, 12, 236–239, 245–251
- Strawson Peter F., 97, 106, 110, 128–129, 135, 186, 203–206, 212
- Subiektyfikacja, 29, 37, 40–42, 45, 48
- Tarski Alfred, 81, 144–145, 160, 185, 190, 203, 256, 271, 278, 280–282, 286
- Teoria bezpośredniego odniesienia / bezpośrednie odniesienie 235–236, 242
- Teoria hybrydowa nazw własnych, 235, 242

Teoria łańcucha przyczynowego, 235–237, 240, 242, 244–245, 250
 Teoria okazjonalna nazw własnych, 235, 240, 245, 250
 Teoria prawdy Tarskiego, 136, 144
 Teoria Reprezentacji Dyskursu (DRT), 185, 187
 Termin jednostkowy, 5, 12, 204–211
 Termin ogólny, 12, 204–205
 Ucieleśnienie, 29, 37–38
 Warunek poprawnego użycia, 213, 215–216, 218–219, 225
 Wirtualny dokument, 29, 45–48
 Wittgenstein Ludwig, 8–11, 14–16, 28, 70, 73, 77, 86, 88, 107, 137, 160, 213, 233–234, 252–253, 255, 257, 265, 267, 269–270, 275, 278–280, 286
 Wprowadzanie nowych znaczeń językowych, 96
 Wymaganie, 50–51, 57, 65–70, 74
 Wyrażenie nienasycone, 204
 Zależność od kontekstu, 161–164, 168, 170, 172
 Zdanie, 12, 19–20, 24, 38–39, 42, 45, 47, 50–52, 55, 59–63, 65, 67, 70–75, 79–80, 82–88, 91–94, 96, 98, 102–103, 105, 108, 111, 113–115, 118–131, 139–142, 145–153, 155–156, 161–182, 185, 187, 191, 193–199, 201, 206, 208–209, 214, 217, 220–221, 226, 228, 235, 237, 239–240, 245, 247, 249, 253, 256, 260, 262, 266, 280
 Zdanie analityczne, 50, 70–74
 Zdanie syntetyczne *a priori*, 50–51, 76
 Znaczenie, 5, 7–9, 11–12, 19–20, 23, 25, 30, 31–35, 41, 61, 63, 71–72, 78, 80–83, 85, 87–108, 110–111, 114, 118–119, 123, 125–128, 132–133, 136–138, 140–141, 143–145, 147–148, 152, 156–159, 161–163, 166–167, 170, 172–176, 178, 180–183, 185, 188, 190, 192, 198, 200–201, 208–211, 213–219, 221–228, 230–232, 237, 239, 242, 255, 266, 271, 275–280, 283
 Znaczenie językowe, 11–12, 71, 96, 107, 110, 161, 174, 213, 232
 Zobowiązanie, 50–51, 65–71, 73–75, 111, 226

