

KRZYSZTOF POSŁAJKO



TO WSZYSTKO NIC NIE ZNACZY

FAKTUALIZM I NONFAKTUALIZM
W KWESTII ZNACZENIA

To wszystko nic nie znaczy

Faktualizm i nonfaktualizm w kwestii znaczenia

Krzysztof Pośłajko



OŚRODEK
BADAŃ
FILOZOFICZNYCH

BIBLIOTEKA OŚRODKA BADAŃ FILOZOFICZNYCH

Komitety redakcyjne serii: Marta Bucholc, Tadeusz Ciecierski,
Paweł Grabarczyk, Witold Hensel, Sebastian Szymański

Copyright©2012

Krzysztof Pośłajko and *Ośrodek Badań Filozoficznych*

ISBN: 978-83-931671-4-2

Recenzja wydawnicza: prof. dr hab. Jan Woleński

Redakcja i korekta: Marcelli Sommer

Projekt okładki: Karolina Pluta

Skład komputerowy: Juliusz Doboszewski

Publikacja finansowana przez Instytut Filozofii Uniwersytetu Jagiellońskiego

Ośrodek Badań Filozoficznych

Warszawa 2012

Spis treści

Wstęp	3
1 Określenie problemu	3
2 Plan pracy	11
 I Paradoks Kripkego – Wittgensteina	 15
I.1 <i>Dramatis personae</i>	15
I.2 Wyzwanie sceptyczne	16
I.3 Odrzucenie koncepcji znaczenia i rozwiązanie sceptyczne	24
I.4 Inspiracje i zbliżone poglądy	29
I.5 W poszukiwaniu ogólnego argumentu	36
 II Normatywność znaczenia	 44
II.1 Podstawowe intuicje dotyczące normatywności znaczenia	44
II.2 Precyzacja Boghossiana	46
II.3 Krytyka tezy o normatywności znaczenia w wersji Boghossiana	49
II.4 Odrzucenie tezy o normatywności znaczenia <i>sensu stricto</i>	57
II.5 Próba zakwestionowania potocznej koncepcji normatywności zna- czenia	70
 III Normatywność a przyczynowość	 79
III.1 Słaba normatywność i jej miejsce w argumentacji sceptycznej	79

III.2	Słaba normatywność a teorie dyspozycyjne	83
III.3	Argument z przyczynowego wykluczenia	89
III.4	Przyczynowe wykluczenie a nonfaktualizm	94
IV	Rozwiązania nieprzyczynowe	97
IV.1	Dyspozycje społeczne	97
IV.2	Rodzaje naturalne	107
IV.3	Koncepcja teleologiczna	118
IV.4	<i>Qualia</i>	123
IV.5	Platonizm i nonredukcjonizm semantyczny	129
V	Nonfaktualizm i minimalizm	137
V.1	Co przemawia za nonfaktualizmem?	137
V.2	Niespójność nonfaktualizmu I: Argument z dwuznaczności	140
V.3	Niespójność nonfaktualizmu II: Argument z uogólnienia	146
V.4	Konsekwencje argumentu z uogólnienia	152
V.5	Minimalizm jako interpretacja rozwiązania sceptycznego	156
V.6	Zdania o znaczeniu a warunek dyscypliny	161
V.7	Minimalizm a nonredukcjonizm semantyczny	168
VI	Zakończenie: Konsekwencje minimalizmu	181
	Bibliografia	188
	Indeks nazwisk	202

Wstęp

1 Określenie problemu

Podtytuł niniejszej książki zawiera dwa kluczowe pojęcia: *znaczenie* i *non-faktualizm*. W języku potocznym bardzo często używamy tego pierwszego: formułujemy wypowiedzi o tym, co znaczą słowa w obcych językach („Angielskie *'turnip'* znaczy to samo co polskie *'rzepa'*”); stwierdzamy tożsamość i brak tożsamości wyrażań („słowo *'jednakowoż'* znaczy co innego niż *'bynajmniej'*”); przypisujemy znaczenia wypowiedziom innych osób („gdy mówiła, że cię kocha, nie znaczyło to wcale, że ma zamiar za ciebie wyjść”). Ogół zdań tego typu będzie w dalszej części książki nazywany dyskursem semantycznym bądź dyskursem o znaczeniu.

O ile pojęcie znaczenia jest bezpośrednio uchwytnie, to pojęcie nonfaktualizmu należy do technicznego słownika filozofii. Jest to określenie na tyle żargonowe, że zapewne nie zna go także większość wykształconych filozofów. Jednak leżąca u jego podstaw idea jest bardzo prosta. Za faktualne uznaje się te dziedziny dyskursu, w których mówi się o faktach, za nonfaktualne zaś — te, w których o faktach się nie mówi. Intuicyjnie uchwytny przykład tego rozróżnienia mogą stanowić sądy z dziedziny estetyki, o czym świadczy zdroworozsądkowa maksyma „O gustach się nie dyskutuje”. Wyznawca nonfaktualizmu estetycznego, skonfrontowany ze sporem dotyczącym oceny jakiegoś dzieła sztuki, będzie zdania, że tak

naprawdę nie mamy do czynienia z autentycznym sporem. O „Bitwie pod Grunwaldem” Matejki możemy, zgodnie z prawdą powiedzieć, że ma takie to a takie wymiary czy że znajduje się w tym a tym miejscu. Nie ma jednak sensu kłócić się o to, czy jest to dzieło wybitne, czy też po prostu kicz. W myśl — dość często spotykanej w zwykłej rozmowie — postawy pospolitego estetycznego relatywizmu, spór o ocenę dzieła jest bezprzedmiotowy. Oceny takiej nie da się sprowadzić do faktów, a zatem nie można powiedzieć, by jeden z uczestników sporu mógł mieć rację.

Nonfaktualizm jest próbą usystematyzowania intuicji stojącej za tym przekonaniem i przedstawienia jej w taki sposób, by mogła objąć także inne dziedziny niż estetyka. Na tym ogólnym poziomie nonfaktualizm można scharakteryzować następująco: jest to przekonanie, iż zdaniom z danej dziedziny dyskursu nie odpowiadają żadne fakty, a co za tym idzie, nie mogą one być ani prawdziwe, ani fałszywe.

Pojęcie dziedziny dyskursu wynika się jednoznacznej definicji. Można je spróbować wyjaśnić za pomocą pewnych przypadków wzorcowych: w naturalny sposób potrafimy odróżnić na przykład dyskurs naukowy od dyskursu estetycznego czy religijnego. Moglibyśmy powiedzieć, że tym, co odróżnia od siebie dyskursy, jest ich przedmiot. Już krótka refleksja pokazuje jednak, że nie jest to fortune ujęcie problemu. Ateista będzie przeczył temu, że dyskurs religijny ma jakikolwiek przedmiot, nie będzie się jednak wzdrygał przed stwierdzeniem, że jest coś takiego jak „dyskurs religijny” i że różni się on od innych rodzajów dyskursów.

Za wzorcowy przykład nonfaktualizmu w dziejach filozofii można uznać emotywizm w dziedzinie etyki. Emotywiści uznawali, że celem wypowiedzi etycznych nie jest opisywanie faktów, lecz wyrażanie odczuć, jakie budzą w nas pewne czyny. Jednym z pierwszych przykładów emotywizmu była teoria Alfreda J. Ayera, który tak scharakteryzował status zdań moralnych:

Jeśli (...) mówię „Kradzież jest zła”, to wypowiadam stwierdzenie, które nie posiada żadnej faktualnej treści — to jest nie wyraża sądu, który może być prawdziwy lub fałszywy. Jest tak, jakbym napisał „Kradzież!!” — gdzie kształt wykrzykników pokazuje, na mocy odpowiedniej konwencji, że mamy tu do czynienia z pewnym specjalnym uczuciem moralnego oburzenia. Jasne jest, że nie zostało tu powiedziane nic, co może być prawdziwe lub fałszywe. (...) mówiąc, iż pewien sposób postępowania jest słuszny bądź nie, nie wypowiadam żadnego faktualnego twierdzenia, nawet twierdzenia o moim stanie umysłu. Jedynie wyrażam własne uczucia moralne (Ayer 1936/1952: 107).

Ayerowska charakterystyka emotywizmu wskazuje wyraźnie na dwa kluczowe elementy każdego stanowiska nonfaktualistycznego: przekonanie, iż zdania z określonej dziedziny nie opisują żadnych faktów oraz że nie mogą one posiadać wartości logicznej. Emotywizm jest o tyle specyficzną wersją nonfaktualizmu, że uznając, iż dane zdania nie odnoszą się do faktów, twierdzi, że są one wyrazem emocji. Nie jest to jednak istotową cechą nonfaktualizmu jako takiego. Nonfaktualizm to stanowisko *stricte* negatywne, które mówi, że pewne zdania czegoś nie robią — tj. nie opisują faktów. Odpowiedź na pytanie, jaka w takim razie jest rola tych zdań w języku, będzie różna w poszczególnych wariantach stanowiska nonfaktualistycznego, a może się zdarzyć, że nie będzie jej wcale.

Stanowisko przeciwne, czyli faktualizm, jest pewną odmianą realizmu filozoficznego — jego zwolennik będzie uznawał, że wypowiedzi z danej dziedziny dyskursu odnoszą się do czegoś rzeczywistego. „Realizm” jest jednak w filozofii pojęciem notorycznie nieostrym, podobnie jak jego przeciwieństwo, czyli „antyrealizm”. Za definicję realizmu jako postawy filozoficznej może uchodzić np. przyjmowanie, że pewnego rodzaju byty egzystują „niezależnie od umysłu” czy nieograniczona akceptacja zasady wyłączonego środka, jak to postulował Michael Dummett (1976/1992: 25–31).

Wydawać by się mogło, iż sprzeciw wobec realizmu w odniesieniu do znaczenia polegać będzie na zaprzeczeniu by istniały takie szczególne byty jak znaczenia. Za przykład teorii postulujących istnienie znaczeń jako pewnych realnie istniejących, oddzielnych bytów posłużyć może koncepcja sensu stworzona przez Gottloba Fregego (Frege 1892/1977: 62). Kategoria sensu w wydaniu tego filozofa ma ewidentnie platoński posmak — Frege postulował włączenia do ontologii odrębnej dziedziny pozaczasowych, obiektywnych bytów.

Nonfaktualizm w odniesieniu do dyskursu o znaczeniu prowadzi oczywiście do zaprzeczenia istnienia tak rozumianych sensów, ale nie musi się do niego ograniczać. Za przykład filozofa, który je odrzucał, ale nie był nonfaktualistą, można uznać Quine’a (przynajmniej w pewnym okresie jego twórczości):

(...) odrzucając znaczenia, wcale nie przeczę tym samym, że słowa i zdania coś znaczą. Dzieląc formy językowe na te, które coś znaczą, i te, które nic nie znaczą, mogę zgadzać się z Iksińskim [hipotetycznym zwolennikiem istnienia znaczeń — przyp. K.P.] co do joty, mimo że Iksiński w przeciwieństwie do mnie, uważa za kryterium tego podziału *posiadanie* (w pewnym szczególnym sensie tego terminu) pewnego abstrakcyjnego bytu, który nazywa znaczeniem. Wolno mi twierdzić, że fakt, iż dane wyrażenie językowe coś znaczy (wolałbym mówić jest sensowne, aby nie otwierać drogi do hipostazowania znaczeń jako bytów), jest rzeczą ostateczną i nieredukowalną. Mogę też podjąć próbę zanalizowania tego zwrotu w terminach zachowań ludzkich, występujących w sytuacjach pojawiania się danego wyrażenia językowego (Quine 1953/1969a: 23).

Mówienie o tym, że pewne wyrażenie posiada znaczenie, nie musi być interpretowane jako stwierdzenie istnienia relacji między tym wyrażeniem a abstrakcyjnym bytem zwanym znaczeniem. Sam Quine w późniejszym okresie swojej twórczości opisał kilka — omówionych w dalszej części pracy — problemów dotyczących

statusu dyskursu semantycznego, nie zmienia to jednak trafności przytoczonej tutaj diagnozy.

Przyjmowanie istnienia rozmaitych bytów jest jedną z możliwych, lecz nie jedyną drogą, pozwalającą na uznanie danego dyskursu za opisujący rzeczywistość - na fakt ten zwraca uwagę Alex Miller w haśle *Realism* w Stanfordzkiej Encyklopedii Filozofii (Miller 2008). Można np. uznawać sensowność dyskursu psychologicznego, mówiącego o stanach mentalnych, nie przyjmując tym samym istnienia substancjalnej duszy. Do utrzymania faktualizmu w kwestii psychologii potocznej wystarczy, jeśli powiemy, że pewne byty fizyczne, takie jak ludzie, mogą mieć pewnego rodzaju przeżycia psychiczne, niesprowadzalne do wydarzeń rozgrywających się w ich mózgach. Możemy tak twierdzić, nie zobowiązując się przy tym do uznawania jakiegoś niematerialnego podłoża tych przeżyć. Po prostu stwierdzamy, że pewne byty fizyczne mogą mieć psychiczne własności — na tym polega sedno tzw. nieredukcyjnego fizykalizmu, stanowiska dość popularnego we współczesnej filozofii umysłu (stanowisko to w przystępny sposób omawia Paprzycka 2005: 25-52).

Podobnie, można uznawać za sensowny i opisujący rzeczywistość dyskurs moralny, nie uznając jednocześnie za konieczne włączenia do swojej ontologii obiektywnych, quasi-platońskich wartości. Jeśli będziemy tego typu realistami moralnymi, to powiemy, że pewne zdarzenie, np. cios siekierą zadany przez studenta prawa lichwiarce, miało tę cechę, że było czynem złym, nie mając przy tym na myśli, że istnieje jakieś abstrakcyjne „zło”, a tylko to, że ludzkie działania można obiektywnie klasyfikować jako „dobre” bądź „złe”.

Z podanych powyżej przykładów wynika, że kluczowe znaczenie dla przyjęcia bądź odrzucenia stanowiska faktualistycznego ma stosunek do istnienia pewnych szczególnych własności, nie zaś bytów. Jak pisze Paul Boghossian:

Irrealistyczne podejście do danej dziedziny dyskursu to pogląd, wedle którego żadne realne własności nie odpowiadają centralnym predykatom omawianej dziedziny (Boghossian 1990: 157).

Pojęcie irrealizmu, którym posługuje się cytowany autor, jest nieco szersze niż pojęcie nonfaktualizmu. Najogólniej określony irrealizm jest koniunkcją trzech, wzajemnie powiązanych ze sobą tez:

1. (Centralnym) predykatom z danej dziedziny dyskursu nie odpowiadają własności.
2. Zdaniom należącym do tej dziedziny dyskursu nie odpowiadają fakty.
3. Zdania należące do tej dziedziny dyskursu nie mogą być ani prawdziwe, ani fałszywe.

Tak ogólnie scharakteryzowane irrealistyczne podejście do danej dziedziny dyskursu można uściślić na dwa sposoby. W pierwszej wersji irrealizm przybiera formę teorii, wedle której predykaty należące do danej dziedziny mają puste zakresy, zaś wszystkie należące do niej zdania elementarne są fałszywe (Boghossian 1990: 159). Za wzorcową dla teorii tego typu można uznać teorię błędu Mackiego, którą autor scharakteryzował następująco:

(...) jakkolwiek większość osób, wygłaszając stwierdzenia z zakresu moralności, *implicite* twierdzi, między innymi, że wskazuje na coś obiektywnie preskryptywnego, to wszystkie te twierdzenia są fałszywe (Mackie 1977: 35).

Warto zauważyć, iż teoria błędu zachowuje pewien element realizmu. Jakkolwiek przedstawiciel tego stanowiska zakłada, iż predykaty dla danej dziedziny dyskursu mają puste zakresy, to przyznaje, że predykaty te wyznaczają one własności, które mogłyby być zrealizowane. Podobnie, nawet jeśli przyjmuje, iż wszystkie zdania elementarne są fałszywe, to uznaje zarazem, iż mają one warunki prawdziwości, które, choć niespełnione w świecie rzeczywistym, mogłyby być spełnione w świecie możliwym. Mówiąc metaforycznie, tylko przypadkiem zdarzyło się tak, że nic nie jest dobre ani złe; nie ma jednak nic nonsensownego w przypuszczeniu,

że pewne rzeczy dobrymi czy złymi być by mogły. Gdy mówiliśmy, że pewne czyny są dobre lub złe, to popełnialiśmy błąd — nie mówiliśmy jednak od rzeczy.

Druga, bardziej radykalna wersja antyrealizmu to właśnie omawiany w tej pracy nonfaktualizm (Boghossian 1990: 159). Jest to stanowisko, wedle którego zdania należące do danej dziedziny dyskursu nie mają w ogóle warunków prawdziwości, a co za tym idzie, nie odnoszą się do żadnych możliwych faktów; predykaty zaś nie wyznaczają żadnych własności. Wydaje się, że to właśnie nonfaktualista w najbardziej jednoznaczny sposób przeczy temu, by dany dyskurs dotyczył czegoś w jakimkolwiek sensie realnego. Zwolennik tej koncepcji powie np., że zdania należące do dyskursu etycznego, nie służą do mówienia o jakichkolwiek faktach. I nie będzie miał na myśli tego, że szukając w świecie faktów moralnych ze smutkiem odkryliśmy ich brak. Sama wyprawa w ich poszukiwaniu była z góry skazana na porażkę. Ze względu na ten radykalizm, stanowisko nonfaktualistyczne w odniesieniu do dowolnej dziedziny dyskursu wydaje się stanowiskiem niezmiernie ciekawym filozoficznie.

W przypadku dyskursu estetycznego czy moralnego jesteśmy w stanie łatwo zrozumieć motywację do przyjęcia nonfaktualizmu, nawet jeśli jej nie podzielamy. Gdy ktoś twierdzi, że żadne czyny nie są „tak naprawdę” dobre ani złe, albo że o żadnym dziele sztuki nie można w należyście uzasadniony sposób powiedzieć, że jest kiczowate lub wzniosłe, to doskonale rozumiemy, o co temu komuś chodzi.

W odniesieniu do dyskursu semantycznego nonfaktualizm wydaje się natomiast na pierwszy rzut oka poglądem niezbyt sensownym. Gdy dowiadujemy się, że angielskie słowo „*eventually*” nie oznacza tego samego co polskie „ewentualnie”, lub że „dziwny” w staropolszczyźnie znaczyło co innego niż w polszczyźnie współczesnej, mamy wrażenie, że w ten sposób uzyskujemy wiedzę o faktach, że zdania, które nam tę wiedzę przekazują, mogą być prawdziwe lub fałszywe. Aby uznać nonfaktualizm w odniesieniu do semantyki za pogląd zasadny, konieczna jest zatem rzetelna filozoficzna argumentacja na jego rzecz.

Na gruncie współczesnej filozofii najbardziej znana obrona tego stanowiska przedstawiona została przez Saula Kripkego w książce *Wittgenstein o regulach i języku prywatnym*. Od czasu jej wydania w 1982 roku, nonfaktualizm semantyczny stał się jedną z najszerzej dyskutowanych kwestii w analitycznej filozofii języka. I o ile niewielu autorów dało się przekonać sformułowanym przez Kripkego argumentom, o tyle wielu dostrzegło konieczność teoretycznego zmierzenia się z tezą o nonfaktualnym charakterze zdań o znaczeniu.

Zadaniem niniejszej książki jest próba racjonalnej rekonstrukcji tego sporu. Nie będę się jednak wcielał w rolę kronikarza, dokumentującego jego historyczny przebieg. Byłoby to zadanie mało sensowne, ze względu na niewielki dystans czasowy dzielący nas od początku sporu oraz fakt, że dyskusja nadal trwa i ciągle pojawiają się w niej nowe argumenty. Można natomiast pokusić się o narzucenie pewnej logicznej struktury na przedstawione argumenty oraz dokonanie oceny ich trafności. Za usprawiedliwienie tego typu historiofilozoficznej dezynwoltury może (częściowo) służyć fakt, iż czytając teksty poświęcone omawianej tu problematyce, niejednokrotnie można odnieść wrażenie, że autorzy polemizują ze stanowiskami, które sami na użytek polemik wytworzyli. W takiej sytuacji ustrukturyzowanie ich sporu, choćby nawet w nieco arbitralny sposób, może się okazać strategią owocną, dającą nadzieję na rzetelne rozpatrzenie padających w dyskusji argumentów i przedstawienie zadowalającego rozwiązania problemu.

Uprawiając filozofię języka, możemy bardziej skupić się na wymiarze filozoficznym albo na językowym. W tym drugim podejściu język to swoisty fenomen, który może być ujmowany z różnych stron przez różnych badaczy — językoznawców, psychologów, socjologów, antropologów czy filozofów wreszcie. Swoistą cechą tego podejścia jest zacieranie granic między dyscyplinami — wiele z publikacji ukazujących się pod szyldem „pragmatyki” trudno jednoznacznie zakwalifikować jako „językoznawcze” czy „filozoficzne”. W pierwszym z wymienionych paradygmatów chodzi z kolei przede wszystkim o rozstrzygnięcie problemów filozoficznych zwią-

zanych z tym, że do rozwiązywania problemów filozoficznych stosujemy metodę analizy językowej — to rozróżnienie na „filozofię języka” i „filozofię lingwistyczną” prezentował m. in. Searle (1971: 1).

W niniejszej pracy postaram się oprzeć na jeszcze innych założeniach. Problemy filozoficzne związane z językiem będą rozpatrywane w kontekście szeroko pojętego realizmu. Jedno z najbardziej istotnych pytań filozoficznych brzmi bowiem: kiedy mamy prawo stwierdzić, że nasze wypowiedzi odnoszą nas do czegoś rzeczywistego, a kiedy musimy przyznać, że nasze słowa trafiają w próżnię? Techniczne pojęcia faktualizmu i nonfaktualizmu mają precyzować ten problem, zaś język jest pewnym szczególnym obszarem tego sporu. Ilekroć toczy się spór o faktualny charakter jakiejś dziedziny dyskursu, pojawia się, przynajmniej *implicite*, założenie o sensowności wprowadzenia podziału na dyskursy faktualne i nonfaktualne, oraz o istnieniu kryterium takiego podziału. Szczegółowy problem języka jest, jak się okaże, użytecznym punktem wyjścia do rozważenia tych kwestii na poziomie ogólnym.

2 Plan pracy

Rozdział I poświęcony będzie szczegółowej analizie argumentów za nonfaktualizmem, które zostały przedstawione w pracy Kripkego. Wniosek, jaki wyłania się z tych rozważań jest zaskakujący: w rzeczzonej książce nie da się odnaleźć dobrych racji na rzecz odrzucenia faktów semantycznych. Autor prezentuje serię argumentów przeciwko konkretnym propozycjom określenia, czym są fakty semantyczne, nie zaś rozumowanie wykluczające istnienie takich faktów. Analiza prowadzi do wyróżnienia u Kripkego trzech linii argumentacji, z których najważniejsza okaże się ta odwołująca się do koncepcji normatywności znaczenia.

Rozdział II poświęcony będzie twierdzeniu o normatywnym charakterze znaczenia. Teza ta stała się przedmiotem licznych polemik, spowodowanych najpewniej niejasnościami narosłymi wokół pojęcia normatywności, jakie stosowano w istnieją-

cych dotąd ujęciach problemu znaczenia. Wniosek z tych rozważań jest następujący: jakkolwiek można mówić o normatywności znaczenia, to nie jest to normatywność taka, z jaką mamy do czynienia w przypadku dyskursu moralnego. Uniemożliwia to proste przeniesienie argumentów z dziedziny metaetyki na problem znaczenia.

Okazuje się zatem, iż argumentację za nonfaktualizmem semantycznym trzeba uzupełnić. W *Rozdziale III* zostanie podjęta próba sformułowania argumentu pokazującego, że domniemane fakty semantyczne nie są koniecznym elementem wyjaśnień przyczynowych, co z kolei ma prowadzić do uznania ich za nieistniejące. Argument ten jest próbą przeniesienia argumentu z przyczynowego wykluczenia, oryginalnie sformułowanego na gruncie filozofii umysłu, na problem znaczenia. Argument ten, w moim odczuciu, zmienia dynamikę sporu między faktualistą a nonfaktualistą, sprawiając, że *onus probandi* znajduje się po stronie zwolennika faktualizmu.

W *Rozdziale IV* zostaną omówione koncepcje, które starają się pogodzić istnienie faktów semantycznych z uznaniem ich za nieefektywne przyczynowo. Żadna z nich wszakże — jak będę starał się dowieść — nie może zostać uznana za właściwą. Faktualista zmuszony jest tym samym do obrony tezy o istnieniu nieredukowalnych faktów i relacji semantycznych. Takie rozwiązanie łamie jednak z kolei zasadę oszczędności ontologicznej. Faktualista musi zatem przytoczyć racje, z uwagi na które mielibyśmy wprowadzać tego typu fakty do naszej ontologii.

O ile *Rozdziały I — IV* przedstawiają argumenty na rzecz nonfaktualizmu, to w *Rozdziale V* zostanie zaprezentowane rozumowanie, które dowodzi, iż jest to stanowisko niemożliwe do utrzymania. Główną przesłanką na rzecz tej tezy jest odkrycie, iż nonfaktualizm prowadzi do niespójności. Aby jej uniknąć konieczne okazuje się uznanie, że przynajmniej niektóre zdania przypisujące warunki deflacyjnej prawdziwości mogą być deflacyjnie prawdziwe.

Rozumowanie sceptyka semantycznego wydaje się poprawne: nie możemy wskazać na żadne fakty, które odpowiadałyby zdaniom o znaczeniu. Z drugiej

jednak strony, nonfaktualistyczna konkluzja jest niemożliwa do zaakceptowania. Wyjściem z tej paradoksalnej sytuacji okazuje się przyjęcie minimalizmu w sprawie zdolności do prawdy. Minimalizm pozwala na twierdzenie, że zdania o znaczeniu mogą być prawdziwe lub fałszywe, jednak wyjaśnia ten fakt poprzez odwołanie się nie do cechy prawdziwości, ale do tego, jak funkcjonuje gra językowa w przypisywanie wartości logicznej. Filozoficzne konsekwencje konieczności przyjęcia minimalizmu w sprawie zdolności do prawdy zostaną omówione w Zakończeniu.

Książka ta jest poprawioną, przynajmniej w mniemaniu jej autora, wersją rozprawy doktorskiej pod tytułem *Spór faktualizmu z non-faktualizmem w kwestii znaczenia* obronionej w 2010 roku. Jakkolwiek podstawowe tezy i struktura argumentu nie uległy znaczącym modyfikacjom, to zmienił się w dużej mierze sposób ich prezentacji.

Artykuł *Spór o normatywność znaczenia w analitycznej filozofii języka*, który ukazał się w „Kwartalniku Filozoficznym”, t. XXXVII, z. 1, w roku 2009, zawiera nieznacznie zmodyfikowaną treść *Rozdziału II* niniejszej książki. Argumentacja przedstawiona w podrozdziałach 2 i 3 *Rozdziału V* stanowi główną część artykułu *O niespójności antyrealizmu w kwestii znaczenia*, który ukazał się w „Przeglądzie Filozoficznym” 2010, R. 19: 2010 Nr 3 (75).

Chciałbym podziękować promotorowi wspomnianej pracy doktorskiej, dr hab. Jerzemu Szymurze, za cenne uwagi i inspirujące dyskusje, a zwłaszcza za pozostawienie szerokiego pola swobody intelektualnej przy pisaniu pracy. Podziękowania należą się także recenzentom — profesorom Tadeuszowi Szubce i Janowi Woleńskiemu — za wnikliwą krytykę. Duży wpływ na kształt książki miały dyskusje, jakie toczyłem z dr hab. Andrzejem Nowakiem oraz Doktorantami i Studentami Instytutu Filozofii UJ. Pracę nad rozprawą ułatwiło znacząco stypendium SYLFF, które wykorzystałem w czasie pobytu na University of Birmingham w Wielkiej Brytanii w 2007 roku. Wdzięczny jestem zwłaszcza prof. Alexowi Millerowi

za cenne sugestie bibliograficzne. Osobne słowa podziękowania należą się mojej
Żonie za nieustanne wsparcie.

Rozdział I

Paradoks Kripkego – Wittgensteina

I.1 *Dramatis personae*

Kripke (1982/2007: 16–17) przedstawia rozważania zawarte w książce *Wittgenstein o regułach i języku prywatnym* jako interpretację późnej myśli Ludwiga Wittgensteina, a zwłaszcza jego „argumentu z języka prywatnego” oraz rozważań dotyczących „reguły do interpretacji reguły” (Wittgenstein 1953/2000 § 86). Wartość interpretacyjna pracy Kripkego była często podważana — charakterystyczna w tej kwestii jest opinia Collina McGinna:

Tym, co naprawdę zrobił Kripke, jest stworzenie imponującego i stanowiącego olbrzymie wyzwanie problemu, który nie ma wiele wspólnego z problemami i twierdzeniami samego Wittgensteina: w pewnym sensie Kripke i prawdziwy Wittgenstein nawet nie zajmują się tymi samymi *kwestiami* (McGinn 1987: 61).

Kripke świadomy był problemów związanych z egzegezą myśli Wittgensteina — dlatego twierdził, iż jego „praca powinna być traktowana nie jako argument Wittgensteina, ani też Kripkego, lecz jako argument Wittgensteina, tak, jak uderzył on Kripkego — tak, jak przedstawił mu się jako problem” (Kripke 1982/2007: 17–18). Podmiot odpowiedzialny za twierdzenia zawarte w książce *Wittgenstein o regulach i języku prywatnym* nazwano później, na poły żartobliwie, Kripkensteinem — ta konwencja będzie przyjęta również i w niniejszej pracy. Zastrzeżenia dotyczące relacji między historycznym Wittgensteinem, a tym, jak przedstawił go Kripke, są w tym miejscu o tyle niezbędne, iż w dalszej części pracy będą występować nawiązania do myśli autora *Dociekań filozoficznych*; nawiązania te nie są jednak czynione z myślą o podaniu jakiejś spójnej i całościowej myśli tego filozofa.

Co istotne, Kripkenstein nie jest jedyną postacią stworzoną na potrzeby podejmowanej tu analizy tekstu Kripkego. Martin Kusch (2006: 13) przestrzega przed utożsamianiem Kripkensteina ze sceptykiem Kripkego–Wittgensteina. Sceptyk jest tym, który przedstawia wątpliwości i przekonuje, że znaczenie jest złudzeniem, że kategoria ta do niczego nas w świecie nie odnosi. Sam Kripkenstein proponuje natomiast, w odpowiedzi na to wyzwanie, „sceptyczne rozwiązanie tych wątpliwości”. Rozróżnienie pomiędzy tymi dwiema postaciami jest o tyle trudne, że również propozycja Kripkensteina bywa często interpretowana jako nonfaktualistyczna.

Ani poglądy sceptyka, ani poglądy „samego” Kripkensteina nie są poglądami, które można by przypisać jakiejś realnej, historycznej postaci. Nie można ich przypisać ani historycznemu Wittgensteinowi, ani też Kripkemu, który niejednokrotnie na marginesie swojej książki dystansuje się od poglądów w niej głoszonych.

I.2 Wyzwanie sceptyczne

Swoją argumentację sceptyk Kripkego rozpoczyna od podania przykładu. Wyobraźmy sobie, iż mam przeprowadzić operację dodawania na liczbach, z których każda jest na tyle duża, iż do tej pory nie używałem ich w procesie dodawania.

Przykład taki z pewnością da się znaleźć, zważywszy na to, iż, jak każdy, wykonałem do tej pory skończoną ilość operacji matematycznych. Dla uproszczenia wyводу przyjmijmy — kontrfaktycznie — iż tymi liczbami są „68” i „57”. Zatem wykonuję operację dodawania i otrzymuję wynik „125”. Moją wypowiedź, że 68 plus 57 daje 125, uważam za poprawną, i to w dwojaki sposób. Po pierwsze, sądzę, iż poprawnie wykonałem to działanie arytmetyczne. Po drugie zaś, uważam moją wypowiedź za poprawną w sensie, jak to ujmuje Kripke, „metajęzykowym” — to znaczy, używam słowa „plus” w takim sensie, w jakim kiedyś postanowiłem go używać — jako oznaczającego pewną wybraną funkcję arytmetyczną.

Załóżmy teraz, że spotykam dziwnego sceptyka. Kwestionuje on moją pewność co do właściwej odpowiedzi w sensie, który dopiero co nazwałem „metajęzykowym”. Być może, twierdzi, w przeszłości używałem słowa „plus” tak, że zgodny z moimi intencjami wynik obliczenia „ $68 + 57$ ” wynosi 5! Sugestia sceptyka jest w oczywisty sposób szalona. (...) Pozwólmy mu jednak kontynuować. W końcu — mówi — jeśli jestem teraz tak pewny tego, że — zgodnie z tym, jak do tej pory używałem symbolu „+” — moją intencją było, by „ $68 + 57$ ” oznaczało liczbę 125, to nie może być tak dlatego, iż sam dokładnie siebie poinstruowałem, że 125 jest wynikiem przeprowadzenia operacji dodawania w tym konkretnym przypadku. Z założenia wynika, że niczego takiego nie zrobiłem. Ale oczywiście pomysł polegał na tym, by w tym nowym przypadku zastosować tę samą funkcję — czy też regułę — której w przeszłości używałem tak wiele razy. Lecz któż może stwierdzić, co to była za funkcja? W przeszłości podałem sobie jedynie skończenie wiele przykładów jej działania. Wszystkie przykłady, jak wcześniej założyliśmy, brały pod uwagę liczby mniejsze od 57. Tak więc, być może w przeszłości używałem terminów „plus” i „+” jako oznaczających funkcję, którą nazwę tu „kwus” i przyporządkuję jej symbol „*”.

$$x * y = x + y \text{ jeśli } x, y < 57$$

$$x * y = 5 \text{ w innym przypadku}$$

Kto może stwierdzić, że to nie tę funkcję miałem uprzednio na myśli, gdy używałem symbolu “+”?

Sceptyk twierdzi (albo udaje, że twierdzi), że obecnie źle interpretuję własne przeszłe użycia słów. Przez “plus”, mówi, *zawsze miałem na myśli (meant)* “kwus”; a teraz, pod wpływem jakiegoś obłąkańczego szalu albo LSD źle rozumiem to, co robiłem wcześniej (Kripke 1982/2007: 22–23).

Wprowadzony przez Kripkensteina sceptyk nie kwestionuje umiejętności liczenia swojego antagonisty — nie jest to sceptycyzm dotyczący zdolności matematycznych. Poddawana w wątpliwość jest natomiast zgodność danego użycia słowa plus z intencją dotyczącą tego, co owo słowo miało w moich ustach znaczyć.

Na pierwszy rzut oka odparcie sugestii sceptyka wydaje się sprawą prostą. W rozważanym przykładzie słowa „plus” wystarczy odnieść się do ogólnych reguł określających, w jaki sposób powinienem „dodawać”. To, że do tej pory nie wykonałem jeszcze obliczenia na tak dużych liczbach, nie ma żadnego znaczenia — opanowanie użycia danego wyrażenia nie polega przecież na wyliczaniu sobie na głos albo w myśli wszystkich możliwych sytuacji, w których mógłbym je zastosować (Kripke 1982/2007: 33–34). Polega ono raczej na opanowaniu reguły używania tego słowa — w tym przypadku dodawania — która to reguła będzie mną „kierować” we wszystkich tych przypadkach. Wydaje się, iż zwłaszcza w przypadku słowa „plus” łatwo jest podać takowe reguły — np. prawo przemienności czy łączności.

Odpowiedź sceptyka na takie *dictum* jest jednak łatwa do przewidzenia. Założmy, że opanowanie reguły rządzącej słowem „plus” polega na podaniu sobie pewnych instrukcji. W instrukcjach tych będą się jednak pojawiać pewne wy-

rażenia. Przyjmijmy, że będą one zawierały słowo „liczyć”. Kripke odpowiada tak:

Jednakże „liczę”, tak jak „plus”, zastosowałem tylko w skończenie wielu przypadkach. Sceptyk może więc kwestionować moją obecną interpretację przeszłego sposobu używania słowa „liczyć” tak, jak to czynił w przypadku słowa „plus”. Może w szczególności twierdzić, że przez „liczyć” miałem dawniej na myśli *litszyć*, gdzie „*politszyć*” ilość elementów w kupce to policzyć ją w zwykły sposób, chyba, że owa kupka została stworzona przez połączenie dwóch innych, z których jedna miała 57 lub więcej elementów — w którym to przypadku należy automatycznie odpowiedzieć “5” (Kripke 1982/2007: 16).

Jeśli, chcąc uzasadnić to, że w danym przypadku postępuję zgodnie z pewną regułą, będę się odwoływał do reguł „bardziej podstawowych”, to sceptyk po prostu zakwestionuje to, czy rozumiem wyrażenia, których używam w owych bardziej podstawowych regułach, i tak *ad infinitum*. Pojawia się tu problem nowej reguły, za pomocą której interpretujemy poprzednią regułę. Wittgenstein wyraził tę wątpliwość tak:

W jaki jednak sposób reguła może mnie pouczyć, co mam zrobić w *tym* miejscu? Cokolwiek zrobię da się przecież poprzez jakąś interpretację pogodzić z regułą (Wittgenstein 1953/2000 § 198).

Oczywiście prezentacja problemu tak, jakby dotyczył on zgodności mojego teraźniejszego użycia z moją przeszłą intencją, ma jedynie charakter dydaktyczny (Kripke 1982/2007: 42). Jeśli bowiem okazałoby się, iż nie można wskazać, w jaki sposób to, co — i jak — mówię teraz, jest realizacją mojego wcześniejszego postanowienia, iż będę używał słowa tak a nie inaczej, to nie da się również pokazać, w jaki sposób moje obecne postanowienie mogłoby określać, jak mam teraz postępować. Problem reguły do interpretacji reguły pokazuje jasno, iż

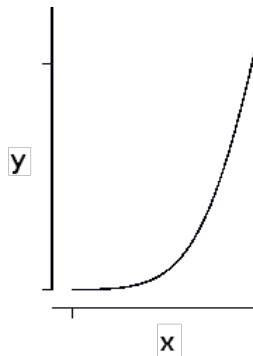
nawet obecne intencje mogą podlegać różnym interpretacjom. Gdybym w tym momencie zdefiniował jakiś symbol, wypisał reguły jego użycia i w tej samej chwili chciał go użyć, to byłbym bezradny wobec sugestii sceptyka, który mógłby przedstawić alternatywną interpretację kierującą użyciem tego symbolu.

W pierwszej chwili mogłoby się wydawać, że przedstawiony przez sceptyka Kripkego problem jest jedynie powtórzeniem problemu indukcji, a dokładniej mówiąc pewnej jego wariacji, jaką jest zagadnienie podciągania danych pod funkcję. Wyobraźmy sobie, że przebadaliśmy pewną ilość przypadków jakiegoś zjawiska. Analizując te dane, dochodzimy do „odkrycia” regularności, co uzasadnia z kolei przeświadczenie, że nasze zjawisko jest opisywane za pomocą pewnej funkcji, pod którą zebrane przez nas dane zdają się podpadać. Z tą procedurą wiążą się dwojakiego rodzaju problemy. Po pierwsze, funkcja, o której sądzimy, że poszczególne przypadki są jej realizacjami, nie jest jedyną, pod którą nasze dane mogą podpadać. Po drugie, jeśli dokonamy wyboru takiej funkcji, to czy uzasadnione będzie dokonywanie predykcji na jej podstawie?

Anegdotyczną ilustracją problemów z tego typu operacjami intelektualnymi są wszelkiego typu quasi-maltuzjańskie predykcje dotyczące np. londyńskich ulic pokrytych niewyobrażalną ilością końskich odchodów związanych z przyrostem użycia dorożek. To, że wydaje nam się, że dostrzegamy jakiś trend czy regularność, nie dowodzi jeszcze, że taka regularność istnieje i że możemy racjonalnie spodziewać się, że przyszłe wydarzenia będą zachodziły w określony przez tę domniemaną regularność sposób.

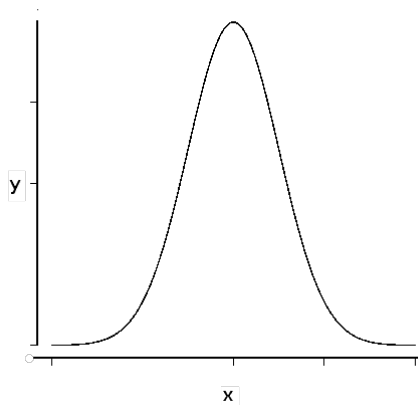
Podawane przez sceptyka Kripkego — jak i samego Wittgensteina — przykłady rażą swoją sztucznością (co zapewne jest celowym chwytem stylistycznym, mającym na celu pobudzenie uwagi czytelnika). Nikt nie będzie na serio rozważał opcji, aby na pytanie „ile jest $68 + 57$ ” odpowiadać inaczej niż „125”. Jednak w pewnych zwykłych przypadkach „kontynuowanie funkcji” nie będzie sprawą prostą. Rozważmy przykład — mamy dorysować „dalszą część” wykresu funkcji

znajdującego się poniżej.



(Rys. 1)

Na pierwszy rzut oka wydawałoby się, że w naturalny sposób powinniśmy kontynuować wykres „w górę”, to znaczy uznać, że funkcja ta jest rosnąca. Tymczasem, oryginalny rysunek wygląda tak:



(Rys. 2)

Jest to po prostu wykres krzywej Gaussa, która jest jak najbardziej naturalną funkcją i która opisuje realnie zachodzące w świecie zjawiska. Mimo to łatwo jest, obserwując tylko część jej wykresu, nabrać przekonania, że mamy do czynienia z funkcją zgoła odmienną.

Problem ekstrapolacji na podstawie ograniczonych danych jest, jako wariacja na temat zagadnienia indukcji, problemem typowo epistemicznym. Pytanie brzmi:

skąd wiemy, że przyszłe zjawiska będą opisywalne założoną przez nas funkcją? Jest to szczególnie przypadek pytania: skąd wiemy, że przyszłość będzie podobna do przeszłości? To drugie pytanie Popper uznał zaś za sedno tradycyjnego problemu indukcji (por. Popper 1972/1992: 10 -11).

Zadając pytania o zasadność stosowania danej metody poznawczej, zakładamy, iż jest jakaś rzeczywistość, którą można za pomocą tejże metody poznać. Jeśli więc mamy wątpliwości, czy posługiwanie się metodą ekstrapolacji z obecnych trendów jest dobrym sposobem przewidywania na przykład przyszłej sytuacji demograficznej, to presuponujemy tym samym, że w przyszłości będą miały miejsce pewne fakty, które uczynią naszą hipotezę prawdziwą lub fałszywą. Kripke twierdzi tymczasem, iż pytanie, przed którym stawia nas sceptyk, nie ma, wbrew pozorom, charakteru epistemologicznego, lecz ontologiczny.

(...) problem wydaje się natury epistemologicznej — w jaki sposób ktokolwiek może wiedzieć, którą z owych rzeczy miałem na myśli? Jeśli jednak weźmiemy pod uwagę to, iż cała moja historia mentalna jest zgodna z obydwojema wnioskami co do tego, co miałem na myśli, stanie się jasne, że zarzut sceptyka w rzeczywistości nie ma charakteru epistemologicznego. Chce on pokazać, że do historii mojego zachowania nie należy nic, co mogłoby określić, czy miałem na myśli plus, czy kwus — i nawet wszechwiedzący Bóg nie znajdzie tam potrzebnej informacji. Wydaje się jednak z tego wynikać, że nie było żadnego dotyczącego mnie *faktu* konstytuującego to, że miałem na myśli plus, a nie kwus (Kripke 1982/2007: 41–42).

Przykłady, w których pokazujemy, że ekstrapolacja funkcji może być sprawą trudną, nie stanowią jeszcze — same w sobie — żadnego argumentu na rzecz tezy, iż nie istnieją fakty określające, który sposób kontynuowania danej funkcji jest właściwy. Stanowią one raczej sposób sformułowania problemu. Mają w nas wzbudzić poczucie, że mamy do czynienia z pytaniem, które domaga się odpowiedzi.

Poprzez podanie przykładu ekstrapolacji funkcji z wykresu, sceptyk Kripkego chce zachwiać naszym przekonaniem, iż cokolwiek decyduje o tym, że to właśnie funkcja przedstawiona na rys. 2. jest „właściwym” (cokolwiek by to miało oznaczać) rozwinięciem funkcji przedstawionej na rys 1. Dlaczego „naturalna” metoda jej rozwinięcia — czyli uznanie, że wartości będą liniowo szły „w górę” miałyby być zła?

Oczywiście nie ma żadnego sensu, jak było to już zaznaczane, próba odpowiedzi, że taką funkcję „miałem na myśli”. Pytanie brzmi bowiem: jaki fakt określa to, co „miałem na myśli”? Należy tu wskazać na ważne założenie, które czyni Kripke. Przyjmuje on, że jest rzeczą niemożliwą, abym — ustalając znaczenie, czy też właściwy sposób postępowania — pomyślał o wszystkich możliwych przypadkach użycia danego symbolu czy wyrażenia. Trafność tego założenia najlepiej widać na przykładzie matematycznym — nikt z nas nie pomyślał przecież *explicite* o nieskończonej tabeli przypadków dodawania.

Problem nie odnosi się jednak tylko do wyrażen z języka matematyki — da się go w prosty sposób przenieść na inne wyrażenia. Jeśli zechcę na przykład nazwać dane zwierzę należące do rodzaju *elephas maximus*, „słoniem”, to sceptyk może zapytać, skąd wiem, iż tak naprawdę, używając słowa „słoń”, nie miałem na myśli „słowcy”, że słowo „słoń” do dnia 1.01.2010 odnosiło się do słoni, od tego dnia jednak odnosi się już do owiec. Kripke chętnie posługuje się zaczerpniętym od Nelsona Goodmana (1955: 74–75) przykładem słowa „ziebieski” — rzeczy są ziebieskie, jeśli do teraz były zielone, natomiast od teraz są niebieskie. Wedle Goodmana, możliwość wprowadzenia predykatu typu „ziebieski” wskazywała na problem dotyczący indukcji — jeśli do tej pory trawa była zielona, to czy mam z tego indukcyjnie wnosić, iż w przyszłości będzie zielona, czy raczej ziebieska. Sceptyk Kripkego pyta natomiast: jaki fakt decydował o tym, że gdy używałem słowa „zielony” to miałem na myśli zielony, a nie „ziebieski”? Podobnie jak w przypadku dodawania, również i tym razem nie pomyślałem przecież z góry

o wszelkich możliwych przypadkach, w których miałbym użyć słowa zielony. Nie pomyślałem przecież o tym konkretnym liściu tego kwiatu, który właśnie tym słowem chcę określić.

Podawane przez Kripkego — jak również przez samego Wittgensteina — przykłady wskazujące na możliwość alternatywnego sposobu postępowania i alternatywnych reguł, z którymi te sposoby mogą zostać uzgodnione, pokazują, że tak naprawdę zawsze jesteśmy postawieni w nowej sytuacji. Za każdym razem, kiedy mamy użyć jakiegoś wyrażenia czy też symbolu, możemy znaleźć jakiś szczegół sytuacji, w jakiej się znajdujemy, który możemy uznać za istotny dla sposobu, w jaki powinniśmy użyć danego słowa.

I.3 Odrzucenie koncepcji znaczenia i rozwiązanie sceptyczne

Jeśli chcemy odeprzeć sugestię sceptyka, to musimy wskazać na coś, co miałoby decydować o tym, że to właśnie preferowany przez nas sposób użycia jest poprawny, a nie dziwaczna sceptyczna alternatywa. Jak na razie, w dyskusji z hipotetycznym sceptykiem wskazywano tylko na dwa rodzaje faktów, które miałyby decydować o tym, w jaki sposób powinienem używać danego słowa — były to werbalne instrukcje oraz moje wcześniejsze lub teraźniejsze intencje. To, iż odwołując się do nich nie da się uzasadnić, czy na pytanie „ile jest $68 + 57$?” powinienem odpowiedzieć „125” czy „5”, jest z pewnością pouczające, jednak nie uprawnia do wniosku, iż żaden tego typu fakt nie istnieje. Kripke wymienia zatem inne propozycje teoretyczne, wskazujące na różnego typu fakty, które miałyby determinować znaczenia wyrażań i sposoby ich poprawnego użycia.

Najwięcej uwagi poświęca teorii dyspozycjonalnej. Jej zrąb został stworzony przez Gilberta Ryle’a w pracy „Czym jest umysł?”. Ryle przeciwstawia się pogładowi, jakoby najogólniej pojmowana „działalność intelektualna” miała polegać

na tym, iż „we wnętrzu podmiotu” dokonują się jakieś tajemnicze operacje. Zdaniem tego autora (Ryle 1949/1970: 71–75), gdy mówimy, iż ktoś działa inteligentnie, to przypisujemy tej osobie pewną umiejętność, posiadanie zaś tej umiejętności jest pewnego rodzaju dyspozycją do działania. Sedno całego dyskursu „mentalnego” tkwić miałoby zatem nie w przypisywaniu innym jakichś niedostępnych poznawczo stanów wewnętrznych, lecz w opisywaniu tego, jakie dyspozycje do działania posiadają.

Prosta dyspozycjonalna teoria znaczenia miałaby, zdaniem Kripkego, następującą postać: „Używając symbolu „+” mieć na myśli dodawanie to posiadać dyspozycję do podawania w odpowiedzi na pytanie o dowolną sumę „ $x + y$ ” sumy x i y ” (Kripke 1982/2007: 44). O tym, jakie jest znaczenie symboli, których używam, nie decydują, w myśl tej koncepcji, ani instrukcje, jakich sobie udzieliłem, ani podane przeze mnie wyjaśnienia tychże symboli, lecz to, w jaki sposób będę postępował — czy też, mówiąc dokładnie, jak postąpiłbym, gdybym znalazł się w odpowiedniej sytuacji. Mogę bowiem posiadać dyspozycję, nawet jeśli nie wystąpią okoliczności, w której mógłbym ją ujawnić.

Kripke przedstawia dwa argumenty przeciwko dyspozycjonalistycznemu ujęciu reguł. Pierwszy z nich odwołuje się do skończoności moich dyspozycji. Funkcja dodawania ma nieskończenie wiele zastosowań dla nieskończonej ilości argumentów. Tymczasem moje dyspozycje są skończone. Niektóre liczby są po prostu za duże, bym mógł je ogarnąć swoim skończonym i ograniczonym umysłem. Z tego powodu nie ma sensu mówić, że moje dyspozycje określają znaczenie słowa „dodawać” (Kripke 1982/2007: 50).

Drugim, o wiele bardziej przekonującym zarzutem jest to, że tego typu wyjaśnienie nie uwzględnia normatywnego charakteru pojęcia znaczenia (Kripke 1982/2007: 47–48). Analiza dyspozycyjna wyjaśnia to, co w danej sytuacji uczynię, nie podaje jednak odpowiedzi na pytanie, co uczynić powinienem. Mogę mieć przecież dyspozycję do mylenia się: na przykład do tego, by w każdym zadaniu

z zakresu dodawania pisać w wyniku 6 zamiast 9. Nikt jednak nie powie, że odpowiedź „16” jest poprawną, zgodną ze znaczeniem symbolu „+” odpowiedzią na pytanie „ile jest $11 + 8$?”.

Kolejną koncepcja znaczenia, z którą polemizuje Kripke, głosi, że nadawanie takiego a takiego znaczenia danemu wyrażeniu cechuje się pewną unikalną, introspekcyjnie dostępną jakością (Kripke 1982/2007: 71–87). Wedle Kripkego, takie wyjaśnienie również upada pod wpływem zarzutów sceptyka. W najprostszym przypadku polega ono na tym, że danemu wyrażeniu przyporządkowane jest pewne wyobrażenie, pewien mentalny obraz. Intuicyjnie uchwytym przykładem takiego wyjaśnienia może być słowo „kostka”. Poprawne użycie tego słowa będzie rezultatem porównania przedmiotu, do którego mam odnieść to słowo, z „mentalnym wzorcem” kostki, znajdującym się w moim umyśle. Tutaj pojawia się jednak od razu problem reguły do interpretacji reguły, w tym przypadku przyjmując postać problemu metody projekcji, czyli tego, co łączy obraz mentalny z odpowiadającym mu przedmiotem. Przecież można sobie wyobrazić taką metodę rzutowania, która przyporządkuje sześciannowi w moim umyśle znajdujący się przede mną ostrosłup! Nawet gdyby mój mentalny obraz przedstawiał także schemat tego, jak należy ową kostkę rzutować, to czy również tego schematu nie można by zinterpretować w jakiś niestandardowy sposób? I skąd mam wiedzieć, który schemat projekcji jest schematem poprawnym?

Jeszcze bardziej problematyczne są przypadki użycia symboli, które odnoszą się do „obiektów” niemających zmysłowej reprezentacji. Załóżmy, że domniemanym, introspekcyjnie danym stanem wewnętrznym, świadczącym o tym, że gdy używam słowa „plus”, to mam na myśli dodawanie, jest pewnego rodzaju szczególny ból głowy. W jaki sposób taki stan miałby mi wskazywać, w opisanym powyżej normatywnym sensie, jak mam dodawać? A nawet jeśli udałoby się wskazać takową relację, czyż nie byłoby łatwo zreinterpretować ją w zgodzie z jakąś „kwusopodobną” regułą?

Rozwiązaniem nie jest też potraktowanie znaczeń jako bytów „platońskich” (Kripke 1982/2007: 90–91). Rozumiana po platońsku, jako pewien wieczny byt, funkcja dodawania zawiera w sobie wszystkie nieskończenie liczne, możliwe zastosowania. Lecz sceptyczna odpowiedź jest w tym przypadku aż nadto oczywista. Choć sceptyk może się zgodzić, iż istniejąca w świecie idei funkcja dodawania wyznacza wszystkie przypadki, nawet w sensie normatywnym, to może zapytać, skąd wiem, w jakiej funkcji dotąd „partycypowałem”? Przecież „kwus” to również byt platoński, a chodzi przecież o to, by wskazać, którą z tych funkcji miałem na myśli. Platonizm zatem nie oferuje żadnego rozwiązania problemu. (Rozwinięcie tych skrótowych uwag na temat qualiów i platonizmu znajdzie czytelnik w *Rozdziale IV*).

Z rozważań tych Kripke — czy też Kripkenstein — wyciąga następujący wniosek:

Niemożliwe jest, by jakiekolwiek słowo miało jakieś znaczenie.
Każde nowe użycie to krok w nieznaną; wszelkie aktualne intencje
da się zinterpretować tak, by pozostały w zgodzie z czymkolwiek,
co postanowimy uczynić (Kripke 1982/2007: 93)

Kripkenstein, rzecz jasna, nie akceptuje tej konkluzji — twierdzi, iż wniosek taki jest całkowicie niewiarygodny i najpewniej wewnętrznie sprzeczny (Kripke 1982/2007: 117). Proponuje on *sceptyczne* rozwiązanie swojego problemu. W przeciwieństwie do rozwiązania „zwykłego” przyznaje ono rację sceptykowi, że nie da się znaleźć faktu odpowiadającego naszym wypowiedziom o znaczeniu. Nie prowadzi to jednak koniecznie do uznania tego typu zdań za zupełnie nieprawomocne. Wedle Kripkensteina, aby uzasadnić prawomocność dyskursu o znaczeniu, przy założeniu o słuszności sceptycznej argumentacji, należy odrzucić obraz języka, w którym kluczową rolę odgrywają warunki prawdziwości. Zamiast tego proponuje on przyjęcie takiego obrazu języka, w którym kluczową rolę odgrywałyby warunki stwierdzalności lub warunki uzasadnienia (Kripke 1982/2007: 119–121).

Zdaniem Kripkensteina, nie powinniśmy pytać o to, jakie „byty”, czy też „fakty”, miałyby odpowiadać wypowiedziom o znaczeniu, lecz o to, jakie jest zastosowanie tego typu wypowiedzi w naszym języku. Jeśli przyjmiemy taką optykę, to okaże się, iż mówienie o znaczeniu odgrywa niezwykle istotną rolę — stwierdzenie przez nas jako społeczność, iż dana osoba rozumie przez pewien symbol to samo, co my, jest tożsame z przyjęciem jej do wspólnoty i uznaniem za kompetentnego użytkownika języka. Jeśli na przykład stwierdzimy, że *X* posługuje się symbolem „+” tak, że znaczy to dodawanie, to możemy mu zaufać, iż gdy wejdziemy z nim w interakcję, w której będziemy używać tego symbolu, jego sposób użycia tego znaku nie będzie się znacząco różnił od naszego. Owszem, może on popełniać sporadyczne pomyłki, jednak spodziewamy się, iż jego zachowanie nie będzie przesadnie dewiacyjne (Kripke 1982/2007: 151–153).

Kripkenstein stara się nas zatem przekonać, iż na kwestię statusu dyskursu semantycznego powinniśmy patrzeć pragmatycznie — to znaczy zadowolić się konstatacją, iż spełnia on niezwykle istotną funkcję w naszym życiu. Jednocześnie, według tego rozumowania, zmuszeni jesteśmy przyjąć, że poszukiwanie faktów, które miałyby ten pragmatyczny sukces wyjaśniać, pozbawione jest sensu (Kripke 1982/2007: 155). Podejście Kripkensteina do problemu funkcji dyskursu o znaczeniu zdaje się więc sugerować, iż reprezentuje on pewnego rodzaju nonfaktualizm w tej kwestii. (W sprawie alternatywnych interpretacji Kripkensteina — jako faktualisty — patrz *Rozdział V*).

Rozwijając tę myśl, Soames sugeruje, iż nonfaktualista w dziedzinie dyskursu o znaczeniu mogłoby ujmować wypowiedzi, w których mówimy, iż ktoś przypisuje jakieś znaczenie danemu wyrażeniu, jako wypowiedzi p e r f o r m a t y w n e (Soames 1998a: 322). Pojęcie performatywu, wprowadzone przez Austina (1962/1993: 554–555), oznacza taki rodzaj wypowiedzi, w którym nie stwierdzamy faktu, lecz dokonujemy pewnej czynności. Wzorcowymi przypadkami performatywów są akty nadawania nazwy, składanie obietnic itp. W myśl propozycji Soamesa, gdy

mówimy, iż ktoś przypisuje symbolowi „+” takie samo znaczenie jak my, nie stwierdzamy więc żadnego faktu, lecz performatywnie konstytuujemy pewien stan rzeczy (jest tak przy założeniu, iż performatywy nie stwierdzają faktów). Na tej podstawie przyjmujemy następnie ludzi do naszej wspólnoty językowej i uznajemy ich za kompetentnych użytkowników języka. Ten fakt — przynależności danej osoby do wspólnoty językowej — jest przez nasze stwierdzenie konstytuowany, a nie „odkrywany”. Ostatecznie rzecz biorąc, Soames nie jest zwolennikiem tezy, aby analiza performatywna była właściwym ujęciem roli wypowiedzi o znaczeniu, nie jest to też jedyna możliwa precyzacja stanowiska Kripkensteina. W propozycji performatywnej ważne jest jednak to, że pokazuje, iż zwolennik nofaktualnego ujęcia dyskursu semantycznego może mieć do opowiedzenia ciekawą historię o tym, dlaczego posługujemy się tym dyskursem i jak wyjaśnić jego powszechność, mimo jego niezdolności do stwierdzania faktów.

I.4 Inspiracje i zbliżone poglądy

Warto w tym miejscu przyrzeć się zagadnieniom, które kierują nasze myślenie w stronę wskazywaną przez sceptyczny paradoks Kripkego-Wittgensteina. Przykładowo, jednoznaczny atak na status dyskursu semantycznego przeprowadził Nelson Goodman w swoim artykule *On Likeness of Meaning*. Uważa on, że nie można stwierdzić, iż jakieś dwa wyrazy mają to samo znaczenie (Goodman 1949: 6). Odrzuca kolejno koncepcje, które miały uzasadnić orzekanie tożsamości znaczeń. Nie może o niej, według niego, decydować ani tożsamość skojarzeń (bo skojarzenia są subiektywne), ani tożsamość odniesienia — nie rozwiązywałaby ona bowiem problemu nazw pustych („jednorożec” i „chimera” miałyby w jej myśl to samo znaczenie). Proponuje natomiast następujący test: dwa wyrażenia są tożsame, jeśli można jedno zastąpić drugim w dowolnym kontekście jego występowania, tak aby nowo powstała fraza zachowywała wartość logiczną frazy wyjściowej. Rozwiązuje to problem nazw pustych — nie sposób przecież zastąpić „chimerą” „jednorożca”

we frazie „oto jest obraz jednorożca”. Test ten jest jednakowoż tak mocny, iż żadna para wyrazów go nie zdaje — dla dowolnych dwóch wyrazów istnieje bowiem kontekst, w którym wymiana taka nie będzie możliwa. Jeśli weźmiemy wyrazy, które chcielibyśmy uznać, za znaczące to samo, np. „trójkąt” i „trójbok” to zdaniem Goodmana, zawsze będziemy mogli skonstruować przykładowy kontekst, w którym jednego wyrazu nie będzie można zastąpić drugim.

Goodman nie odmawia sensowności całemu dyskursowi o znaczeniu — stwierdza jedynie, iż nie da się znaleźć żadnego przykładu zdania, które prawdziwie orzekałoby o tożsamości znaczeń. Można natomiast według niego całkowicie poprawnie orzekać o podobieństwie znaczenia dwóch wyrażań. Istnieją bowiem prawdziwe zdania stwierdzające, iż jedno wyrażenie ma prawie to samo znaczenie, co inne (Goodman 1949: 7). Nasuwają się jednak oczywiste wątpliwości co do takiego rozwiązania — czy rzeczywiście sensowne jest stwierdzanie podobieństwa w sytuacji, w której nie można stwierdzić tożsamości? Choć Goodman nie zadaje sobie tego typu pytań, jego rozważania mogą prowadzić do postawienia bardziej radykalnych pytań na temat znaczenia.

Goodman (1955: 74–75) jest ponadto autorem cytowanego przez Kripkego i wspomnianego powyżej paradoksu dotyczącego słowa „ziebieski”. Paradoks ten ma, w zamierzeniu, głównie wymiar epistemologiczny — dotyczy tego, na jakiej podstawie możemy powiedzieć, że szmaragdy, które do tej pory były zielone będą zielone nadal — skoro możemy równie dobrze przypuszczać, że w przyszłości będą ziebieskie (czyli zielone do dowolnie wybranej chwili t , i niebieskie później), skoro teraz (przed chwilą t) można je opisać zarówno jako zielone i niebieskie. Taki sposób wysłowienia problemu Goodmana pokazuje, iż nie jest on czysto epistemologiczny. Kluczową rolę odgrywa w nim pytanie o to, w jaki sposób — za pomocą jakich predykatów — powinniśmy opisywać rzeczywistość. Mamy więc do czynienia z zagadnieniem tzw. predykatów rzutowalnych (Goodman 1955: 81–85). „Nowa zagadka indukcji” brzmi zatem następująco: za pomocą jakich

hipotez powinniśmy opisywać przypadki, do których mamy poznawczy dostęp? Czy jeśli do tej pory widzieliśmy jedynie białe łabędzie, to czy powinniśmy rezultat naszych obserwacji opisywać za pomocą zdania „wszystkie do tej pory widziane łabędzie są białe”? Może lepszy opis brzmiałby „Wszystkie łabędzie są wbiałe”, gdzie „wbiałe” znaczy „białe, jeśli znajduje się na półkuli północnej a czarne, jeśli na południowej”. Na jakiej podstawie możemy stwierdzić, że ten pierwszy opis jest właściwy? Kwestia ta zbliża problem Goodmana do problemu Kripkensteina, ale mimo zbieżności, ten pierwszy autor konsekwentnie przedstawia „nową zagadkę indukcji” jako pytanie epistemologiczne.

Kwestię statusu dyskursu semantycznego problematyzuje również Quine’owski atak na rozróżnienie analityczne — syntetyczne. Jak wiadomo, pojęcia sądu analitycznego i syntetycznego zostały wprowadzone do filozofii przez Kanta, który zdefiniował je w następujący sposób:

We wszystkich sądach, w których pomyślany jest stosunek podmiotu do orzeczenia (...), stosunek ten jest możliwy w sposób dwojaki. Albo orzeczenie B należy do podmiotu A jako coś, co jest (w sposób ukryty) zawarte w pojęciu A, albo B leży całkiem poza pojęciem A, choć pozostaje z nim w związku. W pierwszym przypadku nazywam sąd analitycznym, w drugim syntetycznym. (...) sądy analityczne są więc sądami, w których powiązanie orzeczenia z podmiotem jest pomyślane przez utożsamienie, sądy natomiast, w których to powiązanie jest pomyślane bez utożsamienia, nazywać się winny syntetycznymi. Pierwsze można nazwać sądami wyjaśniającymi, drugie zaś sądami rozszerzającymi, tamte bowiem nie dorzucają swym określeniom nic do pojęcia podmiotu, lecz jedynie przez rozbiór rozbijają je na pojęcia składowe, które już były w nim, choć mętnie, pomyślane (Kant 1787/2001: B10–B11)

Mimo tego, iż definicja Kanta jest dość niejasna, pojęcie sądu analitycznego, jako takiego, którego treść jest określona przez „zawartość” pojęcia, stało się powszechnie akceptowanym elementem filozoficznego instrumentarium konceptualnego. Pojęcie to wyjaśnia między innymi fakt, że definicje pojęciowe zdają się nie rozszerzać naszej wiedzy.

Idąc w poprzek tej tradycji, Quine (1953/1969b: 58) twierdzi, iż niemożliwe jest wprowadzenie rozróżnienia na sądy analityczne i syntetyczne. Zdanie analityczne roboczo definiuje jako takie, które „jest prawdziwe na mocy swojego znaczenia, niezależnie od faktów” (Quine 1953/1969b: 36). Podobnie jak Goodman, Quine przestrzega przed mieszaniem pojęć znaczenia i ekstensji — mogą istnieć np. predykaty o tych samych ekstensjach, lecz o różnych znaczeniach. Podaje klasyczny już przykład „zwierząt posiadających nerki” i „zwierząt posiadających serce”; prawdopodobnie wyrażenia te odnoszą się do tych samych przedmiotów, jednak ich znaczenie jest inne. W wyjaśnieniu „prawdziwości na mocy znaczenia” nie pomaga też odwołanie się do pojęcia synonimiczności — pojęcie to bowiem zakłada pojęcie znaczenia; podobnie bezowocne jest odwoływanie się do definicji — aby stwierdzić ich poprawność, również musimy dysponować pojęciem znaczenia. Kiedy Quine omawia teorie posługujące się pojęciem wymienialności, dochodzi do tych samych wniosków, co Goodman. Odwołanie się do reguł semantycznych bynajmniej nie rozwiązuje problemu. Zdaniem Quine’a są one bowiem czymś równie, jeśli nie bardziej, tajemniczym od „znanzeń”, które miały być przez owe reguły wyjaśniane.

Ostatecznie Quine odrzuca ideę, jakoby miały istnieć zdania, których prawdziwość zależałaby jedynie od znaczenia występujących w nich terminów. Dowolne zdanie może zostać zrewidowane pod wpływem empirii — o tym, że z pewnych zdań trudno nam zrezygnować, decydują względy wyłącznie pragmatyczne (Quine 1953/1969b: 66–7). Zdaniem Harmana (1967: 124–27), odrzucenie przez Quine’a istnienia zdań analitycznych prowadzi do przeświadczenia, że nie ma czegoś takiego

jak znaczenie, jeśli pojęcie to rozumieć tak, jak było ono wcześniej przedstawiane w filozofii.

Kolejną tezę Quine’a — dyskutowaną wprost przez Kripkego — problematyzującą pojęcie znaczenia była koncepcja niezdeterminowania przekładu. Quine (1960/1999: 39–67) opisuje hipotetyczną sytuację „radikalnego przekładu”, tj. taką, w której mamy dokonać przekładu z języka, który jest nam zupełnie nieznan i który prezentuje się nam bez słownika, jedynie w kontekście zachowań jego użytkowników. Załóżmy, iż osoba mająca dokonać takiego przekładu zauważa, iż tubylcy wypowiadają słowo „gavagai”, ilekroć w polu widzenia pojawi się królik. Naturalną rzeczą byłoby więc przełożenie tego wyrazu jako „królik”. Jednak zdaniem Quine’a, zebrane w ten sposób dane nie pozwalają rozstrzygnąć, czy słowo to należy tłumaczyć jako „królik”, czy raczej jako „faza królika”. Dane świadczące na korzyść jednej hipotezy będą bowiem świadczyć też na korzyść drugiej. Quine wyciąga nawet wniosek, iż „pytanie, czy w opisanej sytuacji tubyliec *naprawdę* sądzi, że A [że ma przed sobą królika — przyp. KP], czy sądzi, że B [że ma przed sobą fazę królika — przyp. KP], jest pytaniem, którego sensowność należałoby podważyć” (Quine 1970: 180–81).

Wniosek Quine’a wydaje się jednoznacznym atakiem na status dyskursu znaczeniowego. Wadą tej argumentacji jest to, iż Quine wychodzi od zdecydowanie behawiorystycznych przesłanek, na co zwraca uwagę Kripke (1982/2007: 93–97). Nawet jeśli zgodzilibyśmy się, iż jego argumenty są poprawne, to oparte są one na dość nieprzekonującej wizji, wedle której dyskurs o znaczeniu powinniśmy wyjaśniać w terminach „znaczenia bodźcowego” (por. Quine 1960/1999: 67). Założeniem rozważań nad „przekładem radykalnym” jest to, iż w takiej sytuacji dysponujemy jedynie behawioralnymi świadectwami i obserwacjami fizykalnymi i to na ich podstawie musimy dokonywać przekładu z nieznanego nam języka. Można zatem po prostu przyjąć, iż aby określić znaczenie danego terminu, potrzebne są inne niż behawioralne przesłanki, i tym samym odrzucić wniosek o niezdeterminowaniu

przekładu. Searle (1987/1993: 158) twierdzi nawet, iż rozumowanie Quine’a można odczytywać jako *reductio ad absurdum* behawiorystycznego podejścia do języka.

Inną motywację do odrzucenia dyskursu semantycznego prezentowali eliminatywiści w dziedzinie filozofii umysłu. Ideą przewodnią eliminatywizmu jest przekonanie, iż terminy tzw. „psychologii potocznej” zostaną, czy też mogą zostać, wyeliminowane z naszego dyskursu w wyniku rozwoju nauki o mózgu (por. Paprzycka 2005: 26–27). Stanowisko to zostało po raz pierwszy sformułowane przez Richarda Rorty’ego (1965/1995: 221–23), przy czym, według niego, eliminacji miał zostać poddany język dotyczący doznań. Zdaniem Rorty’ego, wypowiedzi, w których mówimy o doznaniach, nie da się zredukować do wypowiedzi o neuronach, lecz nie znaczy to, iż musimy przyjąć istnienie nieredukowalnych doznań — możemy po prostu uznać, iż nie ma czegoś takiego jak doznania. Porównuje on język doznań z językiem, w którym mowa o demonach: tego drugiego po prostu przestaliśmy używać.

Paul Churchland zasugerował (Churchland 1981: 67–68), aby w podobny sposób podejść do problemu tzw. *nastawień sądzeniowych* (*propositional attitudes*). Mianem tym określa się w filozofii analitycznej takie stany mentalne podmiotu, w których opisie występuje odniesienie do jakiegoś sądu. Przykładami zdań opisujących nastawienia sądzeniowe są: „Jan jest przekonany, że Uganda nie ma dostępu do morza”, „Zofia obawia się, że inflacja w przyszłym roku wzrośnie”, „Teresa zastanawia się, czy warto inwestować w obligacje”. We wszystkich tych przykładach mamy do czynienia z sytuacją, w której dana osoba żywi jakieś nastawienie psychiczne względem pewnego sądu — z tego powodu polski termin „nastawienia sądzeniowe” wydaje się najbardziej odpowiednim tłumaczeniem angielskiego *propositional attitudes*.

Zdaniem Churchlanda, psychologia potoczna, w której posługujemy się nastawieniami sądzeniowymi, stanowi teorię mającą na celu wyjaśnianie i przewidywanie zachowań innych ludzi. Teoria ta, pozwalając nam na koordynację interakcji z in-

nyimi ludźmi, ma dla nas sporą wartość pragmatyczną. Zdaniem omawianego autora, jest jednak obciążona poważnymi wadami: nie wyjaśnia chociażby procesu uczenia się, z trudem też daje się wkomponować w inne, dobrze działające i uznane teorie naukowe (Churchland 1981: 73–75). Nie są to oczywiście argumenty jednoznacznie wskazujące na fałszywość psychologii potocznej jako teorii, sygnalizują jednak powody, by jej eliminacja pozostała kwestią otwartą.

W jaki sposób uznanie psychologii potocznej za fałszywą teorię umysłu i w konsekwencji jej eliminacja miałyby wpływać na status dyskursu semantycznego? Wydaje się, iż gdybyśmy rzeczywiście zrezygnowali z posługiwania się zdaniem odnoszącymi się do nastawień sądzeniowych, to należałoby również zrezygnować z posługiwania się dyskursem znaczeniowym. Mówienie o tym, co znaczy dana wypowiedź w ustach innej osoby, może wydawać się czynnością nieuzasadnioną, jeśli przyjmiemy, iż osoba ta nie posiada żadnych przekonań. Czy można by uznać za zasadne stwierdzenie, że jakieś wyrażenie znaczy to samo co inne, w sytuacji, w której zrezygnowalibyśmy z mówienia o przekonaniach? Byłoby to trudne, gdyż pojęcie znaczenia wydaje się być ściśle związane z pojęciem przekonania. Jeśli założymy, iż dwa stwierdzenia znaczą to samo, to tym samym przyjmujemy, że przekonania, których wyrazami mają być te stwierdzenia, są identyczne.

Podobnie jak Quine'owi, eliminatywistom zarzucić można przyjmowanie zbyt mocnych założeń. W obu przypadkach mamy do czynienia z dość radykalną postacią scjentyzmu — czyli przekonania, że dyskurs naukowy jest jedynym dyskursem, który możemy uznać jako odnoszący się do rzeczywistości. Zgodnie z tym założeniem, jeżeli dyskurs o znaczeniu — tudzież o nastawieniach sądzeniowych — nie daje się w łatwy sposób sprowadzić do dyskursu naukowego, to należy go odrzucić. Można oczywiście scjentyzm traktować jako uprawnione stanowisko filozoficzne, jednak w tej sytuacji będziemy mieli do czynienia z argumentem zewnętrznym przeciwko faktualnemu podejściu do dyskursu semantycznego. Podejście Kripkensteinowskie wydaje się o tyle ciekawsze, że na pierwszy rzut oka

mamy tu do czynienia z zarzutem, który nie opiera się na przyjętej *a priori* koncepcji rzeczywistości, a który pokazuje, iż nasze zwyczajowe pojęcie znaczenia jest w nieusuwalny sposób problematyczne. Z tego właśnie względu w niniejszej pracy to problemy i argumenty przedstawione przez Wittgensteina w interpretacji Kripkego zostaną poddane analizie.

I.5 W poszukiwaniu ogólnego argumentu

W literaturze komentującej koncepcję Kripkensteina często zwraca się uwagę na to, iż posiadają one zasadniczą wadę — brak w nich argumentu, który na poziomie ogólnym uzasadniałby niemożność przyjęcia poglądu faktualistycznego w odniesieniu do znaczenia. Struktura tekstu Kripkego jest taka, iż po przedstawieniu problemu odrzuca on kolejne propozycje mające wyjaśnić pytanie o to, co konstituuje fakt, że dany symbol coś znaczy. Anandi Hattiangadi pisze tak:

Bez argumentu *a priori* przeciwko wszelkim możliwym faktom, które mogą konstituować to, co mam na myśli (*what I mean*), sceptyk musi zadowolić się słabszą konkluzją, że nie wiadomo nam nic o faktach, które to konstituuja. Jednakże, z tego, że nie wiadomo nam nic o faktach, które miałyby rozstrzygnąć pomiędzy hipotezą, iż mówiąc „śnieg” mam na myśli śnieg, a hipotezą, że mówiąc „śnieg” mam na myśli szneg, nie wynika, że taki fakt nie istnieje (Hattiangadi 2007: 210).

Autorka zwraca tu uwagę na ważną kwestię: bez argumentu, który wykluczyłby z góry *każdą* propozycję istnienia faktu determinującego znaczenie, nonfaktualistyczna konkluzja wydaje się nieuzasadniona. Możliwe jest bowiem, że przy omawianiu różnych możliwości jakaś — właściwa — koncepcja nam umknęła. W takiej sytuacji, jeśli chcemy na poważnie potraktować stanowisko sceptyka Kripkego, to musimy poszukać ogólnego argumentu, który uzasadniłby konkluzję mówiącą, iż fakt określający znaczenie danego słowa nie może istnieć.

Barbara Głód, w swoim studium problemu Kripkensteina trafnie wyróżniła trzy argumenty, którymi posługuje się sceptyk: „Są to (1) argument z regresu interpretacji; (2) argument ze skończoności; i (3) argument z normatywności” (Głód 2003: 86). Pytanie brzmi więc, czy któryś z tych argumentów można uznać za na tyle mocny, by eliminował każdą koncepcję znaczenia?

Argument z regresu interpretacji zdaje się nie spełniać tego wymogu. Uderza on bowiem tylko w te koncepcje, które wyjaśniają znaczenie przez odwołanie się do czegoś, co już znaczenie posiada — np. do ogólnych reguł semantycznych bądź do intencjonalnych przeżyć psychicznych użytkowników języka.

Warto w tym miejscu poczynić pewną dygresję. W analitycznej filozofii języka istnieje silny nurt, który utożsamia znaczenie językowe z intencjami użytkownika. P.F. Strawson, w swoim esej *Znaczenie i prawda*, opisuje konflikt pomiędzy zwolennikami teorii, według której znaczenie powinno się eksplikować w kategoriach komunikacji–intencji użytkowników języka, a tymi, którzy sądzą, iż znaczenie językowe powinno być interpretowane w kategoriach reguł językowych, bez odniesienia do intencji komunikacyjnych użytkowników (Strawson 1970/1993: 12–13). Ci teoretycy, którzy sądzą, iż znaczenie należy analizować odwołując się do intencji użytkownika — wzorcowy przypadek takiego podejścia można odnaleźć u Paula Grice’a (1957: 381) — zdają się mieć gotową odpowiedź na pytanie Kripkego: dane słowo w danym kontekście znaczy to a to, ponieważ wypowiadając je, mamy taką a nie inną intencję komunikacyjną. Dokonując wypowiedzi, zamierzamy przekazać coś naszemu audytorium i właśnie to zamierzenie określa znaczenie naszych słów. Przekonanie, iż odwołanie się do stanów psychicznych użytkownika pozwala rozwiązać paradoks Kripkego, wyraził, między innymi, Colin McGinn, stwierdzając, iż Kripke nie pokazuje, w jaki sposób jego problem miałby odnosić się do pojęć (McGinn 1987: 144–146).

Czy odwołanie się od intencji użytkownika pozwala istotnie w łatwy sposób rozwiązać problem Kripkego? Pozytywna odpowiedź na to pytanie wymaga przy-

jęcia, iż możliwe jest dysponowanie intencjami o określonej treści bez znajomości jakiegokolwiek języka. Samo to założenie, owszem, jest kontrowersyjne, ale nie ma powodu, by je *a priori* odrzucić.

Rozważmy tę propozycję nieco dokładniej. Stany mentalne, za pomocą których mielibyśmy wyjaśniać znaczenie wyrażań, muszą być stanami mentalnymi o określonej treści. Czy przekonanie o posiadaniu przez podmioty takich stanów jest uzasadnione w świetle sceptycznej argumentacji Kripkego? Paul Boghossian dowodzi, iż aby można było przypisywać stanom mentalnym określoną treść, należy założyć, iż mamy do czynienia z czymś na kształt języka myśli — to znaczy, iż elementy mentalne, którym przypisujemy treść, musiałyby być „syntaktycznie identyfikowalne” (Boghossian 1989: 511). Jeśli zaś przyjmujemy, iż mamy do czynienia z czymś na kształt języka myśli — to znaczy wyodrębniamy z naszego uposażenia mentalnego takie elementy, które miałyby być nośnikiem treści, to jasne staje się, że sceptyczna argumentacja Kripkego może być z równym powodzeniem zastosowana do owych nośników. Pojawia się bowiem pytanie: co sprawia, że te elementy naszego „uposażenia mentalnego” posiadają taką a nie inną treść? Wyjaśnienie znaczenia przez odniesienie do stanów mentalnych nie jest zatem właściwe, o ile nie potrafimy wyjaśnić, w jaki sposób elementy języka mentalnego mogą mieć znaczenie.

Wracając do głównego wątku, trzeba stwierdzić, że argument z regresu nie stanowi dobrej racji do odrzucenia teorii wyjaśniających status dyskursu semantycznego przez odwołanie do faktów, które same nie mają charakteru znaczeniowego. Przykładem takiej teorii jest dyskutowana przez Kripkensteina teoria dyspozycyjna: jeśli zgodzilibyśmy się, że tym, co określa znaczenie słów, są dyspozycje użytkowników języka, to wówczas problem regresu interpretacji nie pojawiłby się. Dyspozycje są bowiem czymś, co nie podlega interpretacji — posiadanie lub nieposiadanie dyspozycji jest „nagim faktem”, który sam w sobie nie ma charakteru znaczeniowego. Zatem wyjaśnienie odwołujące się do faktów, które nie wymagają

interpretacji — a do tej kategorii zaliczają się wszystkie fakty „naturalne” — jest odporne na argument z regresu.

Argument ze skończoności także rodzi wątpliwości. Przede wszystkim jest on dość niejasny. Według Kripkensteina, każde słowo ma (potencjalnie) nieskończenie wiele zastosowań, na przykład symbol „+” daje się zastosować do nieskończenie wielu argumentów, z których wiele przekracza nasze zdolności poznawcze (Kripke 198/2007: 50). Domniemany fakt determinujący znaczenie musiałby zatem określać nieskończenie wiele zastosowań użycia danego wyrazu. Jak jednak pisze Kripke „Stan taki musiałby być skończonym obiektem zawartym w naszych skończonych umysłach” (Kripke 1982/2007: 87), co jego zdaniem wyklucza możliwość określania przez niego nieskończenie wielu przypadków stosowania symbolu.

Problemy związane z tym argumentem są dwójakiego rodzaju. Po pierwsze, nie jest wcale jasne, iż stosuje się on do wszystkich wyrażeń — a argument Kripkensteina miał mieć właśnie taki charakter. O ile jest on łatwo uchwytany w przypadku wyrażeń należących do języka matematyki — w ich przypadku bowiem mamy silne przeświadczenie, iż mają one nieskończony zakres zastosowań i mogą się odnosić do elementów, których nie jesteśmy w stanie „objąć” z racji ich „wielkości”, np. bardzo dużych liczb. Czy jednak podobne rozważania mają zastosowanie w przypadku słów odnoszących się do makroskopowych obiektów fizycznych, będących częścią naszego potocznego obrazu świata? Czy przy określaniu znaczenia słowa „kaczka” muszę brać pod uwagę możliwość istnienia kaczek tak dużych, że będą one niemożliwe do pojęcia skończonym umysłem? Takie *reductio ad absurdum* może wydawać się zabiegiem nieuczciwym, lecz pokazuje ono, iż nie jest łatwo zrozumieć, jak zastosować argument z nieskończoności do dyskursów innych niż matematyczny.

Paul Boghossian pisze, iż problem nieskończoności polega na tym, że: „Jeśli mówiąc 'koń' mam na myśli *konia*, to istnieje dosłownie nieskończenie wiele prawd

na temat tego, do czego mam poprawnie stosować ten termin — do koni na Alfa Centauri czy do koni w Wielkiej Armenii i tak dalej — lecz nie do koni czy kotów” (Boghossian 1989: 509). Problem nieskończoności miałby więc polegać na tym, iż dla każdego słowa istnieje nieskończenie wiele możliwych sytuacji, w których mielibyśmy aplikować to słowo.

Tego typu nieskończoność nie wydaje się jednak przesadnie problematyczna. Istnieje wiele własności skończonych przedmiotów, którym moglibyśmy przypisać tak rozumianą „nieskończoność”. Jak zauważa Simon Blackburn, jeśli przyjrzymy się prostej dyspozycjonalnej własności — np. kruchości szklanki — to okaże się, że „istnieje nieskończona ilość miejsc i czasów, powierzchni i uderzeń, które mogą [tę dyspozycję — przyp. KP] ujawnić” (Blackburn 1984: 289). Nawet jeśli pominiemy własności dyspozycyjne — taka prosta własność przedmiotu jak masa jest w pewnym sensie nieskończona: jeśli przedmiot ma daną masę, to prawa fizyki określają, jakie siły grawitacyjne będą na nią działać, gdy znajdzie się on w danym miejscu pola grawitacyjnego innego obiektu. Jeśli zatem mój długopis posiada określoną masę, to istnieje nieskończenie wiele sił grawitacyjnych, które mogłyby na niego działać, gdyby znalazł się w określonym punkcie czasoprzestrzeni. Reasumując: albo argument z nieskończoności odnosi się jedynie do symboli o nieskończonych ekstensjach — wówczas nie jest to argument ogólny, lecz lokalny; albo odnosi się do wszystkich wyrażań — w takiej sytuacji trudno jednak zrozumieć istotę tego zarzutu. Kluczowe w rozumowaniu sceptyka było to, że mój umysł nie może wziąć pod uwagę nieskończonej ilości przypadków. Taki argument zdaje się wszakże zakładać wizję umysłu, która uwzględnia jedynie świadome przeżycia. W obrazie dyspozycjonalnym nie muszę tymczasem świadomie rozważać możliwych przypadków użycia symbolu. W ujęciu tym, nabywanie rozumienia słowa polega na tym, że wykształca się pewien mechanizm, który następnie może aktywować się w nieskończonej liczbie przypadków.

Na placu boju zostaje zatem argument z normatywności znaczenia. Z wielu względów wydaje się on bowiem najbardziej oczywistym kandydatem do uzasadnienia *a priori* i w sposób ogólny, iż zdania o znaczeniu nie odnoszą się do faktów. Kwestia relacji między zdaniami wyrażającymi normy a opisującymi fakty jest uznawana w filozofii za problematyczną przynajmniej od czasów Davida Hume'a, który pisał:

W każdym systemie moralności, z jakim dotychczas się spotykałem, stwierdzałem zawsze, że autor przez pewien czas idzie zwykłą drogą rozumowania, ustala istnienie Boga, albo robi spostrzeżenia dotyczące spraw ludzkich; aż nagle nieoczekiwanie i ze zdziwieniem znajduję, iż zamiast zwykłych spójek, jakie znajduje się w zdaniach, a mianowicie *jest* i *nie jest*, nie spotykam żadnego zdania, które by nie było powiązane słowem *powinien* albo *nie powinien*. Ta zmiana jest niedostrzegalna, lecz niemniej ma wielką doniosłość. Wobec tego bowiem, że to *powinien* albo *nie powinien* jest wyrazem pewnego nowego stosunku czy twierdzenia, przeto jest rzeczą konieczną te zwroty zauważyć i wyjaśnić; a jednocześnie konieczne jest, iżby wskazana została racja tego, co wydaje się całkiem niezrozumiałe, a mianowicie, jak ten nowy stosunek może być wydedukowany z innych stosunków, które są całkiem różne od niego (Hume 1739/1963: 227)

Jakkolwiek nie jest jasne, czy Hume był zwolennikiem jakiejś formy antyrealizmu w stosunku do dyskursu moralnego, to z przytoczonego cytatu wynika, że dostrzegał on zasadniczą różnicę między wyrażeniami, w których stwierdzamy to, co jest, a tymi, w których mówimy, co być powinno. Jeśli zakwalifikowalibyśmy wypowiedzi o znaczeniu do kategorii tych wyrażających powinność, pojawia się problem ich relacji do wypowiedzi o faktach.

Wzorcowymi przypadkami obu omawianych w Rozdziale I koncepcji antyrealistycznych — zarówno teorii błędu, jak i nonfaktualizmu — były stanowiska

z zakresu metaetyki. Etyczny realizm — czyli stanowisko, wedle którego zdania z zakresu etyki opisują „rzeczywistość zewnętrzną wobec podmiotu” — nie jest bynajmniej stanowiskiem jednoznacznie odrzucanym we współczesnej metaetyce (przegląd stanowisk realistycznych podają np. Miller 2003: 138–298, van Roojen 2008), jednak przekonanie, iż istnieje zasadnicza różnica pomiędzy zdaniami normatywnymi a zdaniami o faktach, jest dość powszechne.

Przekonanie, iż normatywny charakter dyskursu o znaczeniu stanowi podstawową trudność w wyjaśnieniu jego ontologicznego statusu często podnoszono w debatach, które omawiały sceptyczną argumentację. Hattiangadi posuwa się nawet do stwierdzenia, iż „argument Kripkego przeciwko realizmowi semantycznemu jest strukturalnie analogiczny do argumentu Ayera przeciw realizmowi moralnemu” (Hattiangadi 2007: 39). Simon Blackburn opisując problem Kripkensteina stwierdza z kolei, iż „zagadnienie brzmi: czy istnieją poprawne i niepoprawne użycia wyrażen” (Blackburn 1984: 281).

Prosty argument za antyrealizmem w odniesieniu do dyskursu semantycznego miałby zatem następującą postać:

ZDANIA O ZNACZENIU WYRAŻAJĄ NORMY.

ZDANIA WYRAŻAJĄCE NORMY NIE OPISUJĄ FAKTÓW.

ZDANIA O ZNACZENIU NIE OPISUJĄ FAKTÓW.

Obie przesłanki tego rozumowania można uznać za kontrowersyjne. Jednak w dalszej części pracy teza o niefaktualnym charakterze dyskursu etycznego zostanie przyjęta bez dowodu. Założenie to ma charakter roboczy — służyć ma ono sprawdzeniu, czy uznanie tej przesłanki pozwoliłoby uzasadnić tezę nonfaktualistyczną.

W filozofii analitycznej ostatnich lat pojawiło się wiele argumentów przemawiających za odrzuceniem pierwszej przesłanki, czyli tezy o normatywności znaczenia. Z tego względu konieczne będzie ich omówienie w tym miejscu — jeśli okazałoby

się, iż istnieją dobre racje do twierdzenia, iż zdania z dyskursu semantycznego nie wyrażają norm, to musielibyśmy odrzucić prosty argument za antyrealizmem. Teza o normatywności znaczenia wydaje się zatem kluczowa dla rozstrzygnięcia o trafności argumentacji Kripkego — z tego względu Rozdział II będzie poświęcony jej analizie.

Rozdział II

Normatywność znaczenia

II.1 Podstawowe intuicje dotyczące normatywności znaczenia

Teza o normatywności znaczenia jest przez jej zwolenników uznawana za „intuicyjną” w sensie, jaki temu słowu nadaje techniczny żargon filozofii analitycznej. Aby uznać jakąś prawdę za intuicyjną, nie potrzeba specjalnej władzy poznawczej — prawdy intuicyjne to te, które bezrefleksyjnie, nawet nieświadomie, akceptują wszyscy, a przynajmniej wszyscy kompetentni użytkownicy języka. W ten właśnie sposób wszyscy mieliby uznawać, że zdania o znaczeniu są zdaniami wyrażającymi normy. Normatywny charakter tego typu zdań polegać miałby zaś na tym, iż określają one, jak powinno się danego wyrażenia używać, a nie to, jak jest ono faktycznie używane. Mimo, iż na poziomie powierzchniowym zdania te mają formę zdań opisowych, wyrażają one *de facto normy* — mówią o tym, które sposoby użycia danego wyrażenia są poprawne, a które nie.

Teza o normatywności znaczenia, spopularyzowana we współczesnej filozofii języka przez Kripkego, została następnie podchwycona przez Crispina Wrighta (por. 1989/2002: 112–113) oraz Johna McDowella. Kripke o normatywności znaczenia pisze tak:

Założmy, że używając symbolu „+”, mam na myśli dodawanie. Jaki jest związek tego założenia z pytaniem o moją odpowiedź na problem, ile jest „68 + 57”? [...] Właściwe jest (...) podejście *normatywne*. *Nie chodzi* o to, że, jeśli używając „+”, miałem na myśli dodawanie, odpowiem „125”, ale o to, iż, jeśli chcę działać w zgodzie z tym, co w przeszłości miałem na myśli, używając symbolu „+”, *powinienem* odpowiedzieć „125”. Błędy obliczeniowe, skończoność mojego pojmowania i inne czynniki zakłócające mogą doprowadzić do tego, że nie będę posiadał *dyspozycji* do odpowiadania tak, jak *powinienem*, ale w takim przypadku nie zadziałam zgodnie z moimi intencjami. Związek między znaczeniem [*meaning*] i zamiarem a przyszłymi działaniami jest *normatywny*, a nie *opisowy* (Kripke 1982/2007: 65–66).

Zdaniem Kripkego, istnieje różnica pomiędzy faktycznymi sposobami użycia danego wyrażenia lub dyspozycjami do jego używania przez użytkownika języka a tym, jak powinien go używać. Znaczenie wyrażenia jest czymś, co wyznacza poprawne sposoby użycia. Kiedy mówimy, iż dany sposób użycia wyrażenia jest zgodny z jego znaczeniem, mamy na myśli to, iż jest to właściwy sposób użycia, nie zaś to, iż dany użytkownik będzie tego wyrażenia używał właśnie tak. Aby rozbudzić nasze intuicje w tej kwestii, Kripke zauważa, iż często określamy użycia słów jako niepoprawne. Co więcej, niektórzy użytkownicy języka mają dyspozycję do regularnego popełniania pomyłek. Z tego też powodu nieuzasadnione jest utożsamienie znaczenia danego słowa z dyspozycją czy też tendencją do używania go tak, a nie inaczej (Kripke 1982/2007: 53–54).

Fakt istnienia pomyłek nie może być oczywiście traktowany jako dowód, w ścisłym sensie tego słowa, na rzecz tezy o normatywności znaczenia; twierdzenie, że używając języka, można się mylić, zakłada bowiem poprawność bądź niepoprawność użycia słów. Odwołanie się do praktyki określania wypowiedzi jako poprawnych bądź niepoprawnych ma służyć raczej wzbudzeniu „intuicji” (w

przytoczonym wcześniej znaczeniu) przemawiających za wyżej wymienioną tezę. Przekonanie o tym, iż teza o normatywności znaczenia ma charakter bezdyskusyjny i nie wymaga uzasadnienia, zostało wyrażone w najbardziej dobitny sposób przez Johna McDowella:

W naturalny sposób myślimy o znaczeniu i rozumieniu w terminach w pewnym sensie kontraktowych. Przyjmujemy, iż nauczenie się znaczenia pewnego słowa jest tym samym co uzyskanie rozumienia, które następnie zobowiązuje nas — jeśli mamy okazję do używania rozważanego pojęcia — byśmy sądzili i mówili w pewien określony sposób (...) (McDowell 1984: 325).

Można w tym miejscu dokonać podsumowania podstawowych tez przedstawionej tu, potocznej koncepcji normatywności znaczenia:

- a. Sposoby użycia języka dzielą się na poprawne i niepoprawne.
- b. Nie każde możliwe użycie języka jest poprawne; użytkownik może się mylić, nawet w sposób systematyczny.
- c. Znaczenie słowa określa, które sposoby jego użycia są poprawne — nie zaś to, w jaki sposób mówiący będą tego słowa używać.
- d. Tezy a, b, c są wyrazem naszych potocznych przeświadczeń na temat znaczenia; nie jest potrzebna szczegółowa argumentacja na ich rzecz.
- e. Uznanie tez a, b, c jest częścią każdej właściwej filozoficznej analizy pojęcia znaczenia; jeśli jakaś teoria tych tez nie zawiera, to oparta na niej analiza jest błędna.

II.2 Precyzacja Boghossiana

Przedstawiona przez Kripkego i McDowella koncepcja normatywności znaczenia, nazwana wyżej potoczną, choć przekonująca, obciążona jest pewną wadą —

brak w niej charakterystyki tego, które sposoby użycia języka są poprawne, a które nie. Kripke i McDowell wskazują, iż dokonujemy podziału na poprawne i niepoprawne użycia, nie mówią jednak nic o kryterium — nawet słabo rozumianym — wedle którego miałby dokonywać się ten podział.

Wydaje się, iż to zauważenie tego braku skłoniło Paula Boghossiana do bardziej ścisłego wyrażenia omawianej tu intuicji. Precyzacja zaproponowana przez tego autora w znanym tekście *The rule-following considerations* stała się niemal powszechnym sposobem rozumienia tezy o normatywności znaczenia, przyjmowanym za pewnik w wielu dyskusjach, zatem warto przytoczyć ją w całej rozciągłości:

Przypuśćmy, iż wyrażenie „zielony” znaczy *zielony*. Wynika z tego bezpośrednio, iż wyrażenie „zielony” stosuje się *poprawnie* tylko do *tych* rzeczy (zielonych) nie zaś do *tamtych* (nie-zielonych). Fakt, iż dane wyrażenie coś znaczy, pociąga za sobą cały zbiór *normatywnych* prawd na temat mojego postępowania z tym wyrażeniem: mianowicie to, że moje użycie jest poprawne w odniesieniu do pewnych przedmiotów, zaś w odniesieniu do innych nie (...).

Normatywność znaczenia okazuje się zatem, innymi słowy, po prostu nowym określeniem dobrze znanej prawdy, iż — niezależnie od tego, czy ujmujemy znaczenie w terminach warunków prawdziwości czy warunków stwierdzalności — wyrażenia, które mają znaczenie, posiadają również warunki poprawnego użycia. (W pierwszej z wyżej wymienionych teorii są to warunki *prawdziwości*, w drugiej — warunki *stwierdzalności*). Odkrycie Kripkego polega na zauważeniu, że konstatacja ta może być przekształcona w warunek adekwatności dla teorii znaczenia: każda własność, która miałaby konstituować znaczenie danego słowa, musi ugruntowywać normatywność znaczenia — z każdej własności, która rzekomo konstituuje znaczenie danego słowa, powinno się dać odczytać, jakie jest poprawne użycie tego słowa (Boghossian 1989: 513).

Widać wyraźnie, iż Boghossian akceptuje tezy a - e, stanowiące rdzeń „potocznej” koncepcji normatywności znaczenia, sformułowanej przez Kripkego, McDowell’a i innych. Dodaje jednak do nich istotnie nową treść: wyrażenie jest używane poprawnie tylko wówczas, gdy odnosimy się za jego pomocą do bytów, które ono rzeczywiście oznacza.

Interpretacja Boghossiana pozwoliła Danielowi Whitingowi na ujęcie tezy o normatywności znaczenia w sposób quasi-formalny, za pomocą następującej formuły (Whiting 2007: 134):

$$(N_1) \ (w \text{ znaczy } F) \Rightarrow \forall_x \ (w \text{ stosuje się poprawnie do } x \leftrightarrow x \text{ jest } f)$$

gdzie $\forall$ oznacza kwantyfikator ogólny, w jest wyrażeniem, F oznacza pojęcie, zaś f cechę posiadaną przez przedmioty, do których odnosi się owo pojęcie.

Zdaniem Boghossiana, tak rozumiana normatywność znaczenia znajduje zastosowanie zarówno do wyrażeń języka „publicznego”, jak i do pojęć, z których miałby się składać język mentalny. Jeśli w „języku myśli” mamy „wyrażenie”, które oznacza „zielony”, to stosuje się ono poprawnie tylko i wyłącznie do przedmiotów zielonych.

Powyższe sformułowanie tezy o normatywności znaczenia prowadzi jednak, według samego Boghossiana, do pewnych problemów. Jeśli przyjrzymy się równoważności w nawiasie tezy (N_1) i skupimy uwagę na implikacji od prawej do lewej, to okaże się, iż teza o normatywności znaczenia zobowiązuje nas do stosowania danego wyrażenia w odniesieniu do wszystkich rzeczy należących do zakresu danego pojęcia. W podanym przykładzie oznaczałoby to osobliwe wymaganie nazywania wszystkich zielonych rzeczy „zielonymi”. Podobnie, chcąc użyć słowa „kot”, musielibyśmy je zastosować do każdego żyjącego na planecie egzemplarza tego gatunku, co byłoby działaniem tyleż niemożliwym do wykonania, co bezcelowym.

Wymóg odnoszenia się do wszystkich rzeczy wchodzących w zakres danego pojęcia należałoby zatem, w ślad za Boghossianem i Whitingiem, zastąpić słabszym — regułą, nakazującą odnoszenie się tylko do takich przedmiotów, które się

w owym zakresie mieszczą (Boghossian 2005: 210-211). Oznaczałoby to, że używamy słowa „zielony” poprawnie tylko wówczas, jeśli nazywamy nim przedmioty zielone. Zaś „kotami” tylko przedstawiciele odpowiedniego gatunku. Nie byłoby jednak wspomnianej wyżej konieczności nazywania wszystkich przedmiotów zielonych „zielonymi” i poszukiwania wszystkich kotów na świecie. Tę słabszą wersję precyzacji Boghossiana można sformalizować następująco (Whiting 2007:137):

$(N_2) (w \text{ znaczy } F) \Rightarrow \forall_x (w \text{ stosuje się poprawnie do } x \rightarrow x \text{ jest } f)$

Charakterystyczną cechą przedstawionej przez Boghossiana interpretacji tezy o normatywności znaczenia jest to, że względną precyzję osiąga się tu kosztem znacznego uproszczenia. Przy tym rozumieniu, jakie proponuje, rozstrzygnięcie, czy ktoś używa danego wyrażenia w danym kontekście poprawnie, okazuje się dość proste — w przypadku prostej wypowiedzi przypisującej danemu obiektowi pewną cechę, wypowiedź ta okazuje się poprawna, jeśli obiekt ten cechę tę posiada — poprawnie używam słowa „ogar”, jeśli określam tym mianem tylko psy, które do tej rasy należą, a nie np. jamniki. Uproszczenie to budzi jednak łatwe do przewidzenia obiekcje.

II.3 Krytyka tezy o normatywności znaczenia w wersji Boghossiana

Można wskazać na dwa argumenty przeciwko tezie o normatywności znaczenia, które wprost odnoszą się do zaproponowanego przez Boghossiana rozumienia tej tezy. Pierwszy z tych argumentów — „z obowiązku prawdy” — został przedstawiony m.in. przez Anandi Hattiangadi (2006: 229–230) oraz Åsę Wikforss (2001: 203–205). Jeśli tezę o normatywności znaczenia należy rozumieć tak, jak zostało to przedstawione w (N_2) , to normy semantyczne zdają się „wymagać” od użytkowników języka, aby zawsze mówili prawdę. Tymczasem nic nie stoi

na przeszkodzie, abyśmy mogli — nadal posługując się poprawnie językiem — kłamać, ironizować bądź mylić się.

Oto ktoś pyta nas o sytuację finansową Marka. Chcemy wprowadzić w błąd pytającego i mówimy, niezgodnie z prawdą: „Marek jest bogaczem”. Jeśli Marek w rzeczywistości jest człowiekiem niezbyt zamożnym, to naszą wypowiedzią złamaliśmy normę językową rządzącą użyciem słowa „bogacz”, którą wyrażać ma następujące podstawienie tezy (N_2):

Jeśli „bogacz” oznacza człowieka zamożnego, to słowo „bogacz” stosuje się poprawnie tylko do ludzi zamożnych.

Wypowiadając zdanie „Marek jest bogaczem”, popełnialibyśmy, według teorii Boghossiana, błąd językowy i zdradzalibyśmy swoją nieznajomość znaczenia słowa „bogacz”. Trudno pogodzić tę konkluzję z naszymi intuicjami dotyczącymi funkcjonowania języka. Jeśli ktoś posługuje się danym słowem, aby wprowadzić innego użytkownika języka w błąd, to najwyraźniej słowo to rozumie. Podobnie, użycie jakiegoś wyrażenia w kontekście ironicznym zdaje się raczej wskazywać na wprawę w posługiwaniu się nim, niż na nieznajomość rządzących nim reguł.

Oczywiście, można wywodzić, że istnieje moralny obowiązek mówienia prawdy, że prawda sama w sobie jest wartością. „Należy mówić prawdę” to interesująca, ale filozoficznie dość kontrowersyjna teza, której rozstrzygnięcie, jak się wydaje, powinno się dokonać na terenie etyki, nie zaś filozofii języka. Byłoby dość zaskakujące, gdyby spory etyczne można było rozstrzygać poprzez odwołanie się do potocznych intuicji na temat funkcjonowania języka. Sugestia Boghossiana zostaje tym samym sprowadzona do absurdu.

Kolejny argument przeciwko tezie o normatywności znaczenia wysunął sam Boghossian (2005: 213–216), który z wyznawcy tej tezy stał się jej zagorzałym krytykiem. Zarzut ten został pierwotnie sformułowany w odniesieniu do normatywności treści pojęć konstytuujących „język mentalny”. Jako że, zdaniem naszego autora, normatywność wyrażań języka publicznego i normatywność pojęć „języka

myśli” to jeden i ten sam problem, możliwe jest jednak odniesienie tych wątpliwości do wyrażen języka mówionego lub pisanego.

Wedle Boghossiana, niesłuszne jest mówienie o normatywności znaczenia, gdyż normatywność ową można odnosić jedynie do zdań twierdzących. Skoro tak, to zamiast o normatywności znaczenia należałoby mówić o „normatywności sądu”. Można formułować normy mówiące jedynie o tym, co powinniśmy sądzić: powinniśmy mianowicie sądzić prawdziwie. Błędem natomiast jest mówienie o normatywności znaczenia — aby pojęcie to miało sens, powinno odnosić się do wszelkiego rodzaju użyc języka, nie tylko do asercji.

Rozważmy następujący przykład: „Jan sądzi, iż następnym prezydentem Polski zostanie Janusz Korwin-Mikke”. Być może istnieje jakaś pozasemantyczna norma, która określa, że powinniśmy dążyć do tego, aby sądy, które wydajemy, były prawdziwe, i w związku z tym fraza „następny prezydent Polski” powinna być stosowana jedynie w odniesieniu do osoby, która faktycznie zostanie następnym prezydentem Polski. Tego typu frazy można jednak stosować również w innych kontekstach, na przykład: „Jan chciałby, aby następnym prezydentem Polski został Janusz Korwin-Mikke”, „Jan obawia się, iż następnym prezydentem Polski zostanie Janusz Korwin-Mikke”, „Jan wątpi, czy następnym prezydentem Polski zostanie Janusz Korwin-Mikke” itd. Jeśli w takich przypadkach również działałyby normy semantyczne typu (N_2), oznaczałoby to, iż wypowiedź postaci „Chciałbym, aby x został następnym prezydentem Polski” jest poprawna tylko wtedy, jeśli pod x podstawiona zostanie osoba, która faktycznie zostanie następnym prezydentem Polski — co wydaje się jawnie sprzeczne z naszymi potocznymi przekonaniem na temat działania języka. Dlaczego normy dotyczące użycia wyrażen same z siebie miałyby określać, które wyrażenia preferencji są poprawne?

Boghossian zauważa, iż szczególnie problematyczne dla zwolenników tezy o normatywności znaczenia jest zjawisko „rozważania sądu” (*entertaining a proposition*), to znaczy takiego nastawienia propozycjonalnego, w którym ani się sądu

nie akceptuje, ani nie odrzuca. Za językowy korelat takiego nastawienia można uznać frazę „Zastanawiam się, czy...”. Boghossian twierdzi, iż mówienie w takich kontekstach o normatywności jest zupełnie nieuzasadnione. Nie ma żadnego powodu, dla którego pewne podstawienia frazy „*x* zastanawia się, czy *y* zostanie następnym prezydentem Polski” miałyby być uznane za niepoprawne tylko dlatego, że faktycznie *y* nie zostanie następnym prezydentem Polski. Czy rozważanie scenariuszy, które się nie zrealizują, musi być zajęciem nonsensownym? Co więcej, jeśli uznamy, iż fraza „*y* zostanie następnym prezydentem Polski” odnosi się tylko i wyłącznie do osoby, która rzeczywiście zostanie następnym prezydentem Polski, to pojęcie normatywności znaczenia zawęży zakres wypowiedzi, które możemy uważać za poprawne, i w rezultacie, zamiast uchwycić nasze intuicje dotyczące posługiwania się językiem, zdaje się prowadzić do wniosków zupełnie z nimi niezgodnych. Teza o normatywności znaczenia w rozumieniu, jakie zaproponował Boghossian, jest zatem nie do utrzymania — wypływają z niej konsekwencje niemożliwe do uzgodnienia z naszymi intuicjami dotyczącymi języka, a koncepcja normatywności znaczenia była przecież pomyślana jako uchwytyująca owe intuicje.

Wydaje się jednakowoż, iż nie trzeba wylewać dziecka z kąpielą i odrzucać przekonania, że zdania o znaczeniu wyrażają normy. Rozwiązanie, które eliminuje kłopotliwe konsekwencje podejścia Boghossiana, utrzymując jednocześnie w mocy potoczną koncepcję normatywności znaczenia, zostało przedstawione przez Alana Millara i rozwinięte przez Hansa-Johana Glocka. Millar (2002: 60–61) odróżnia dwa rodzaje poprawności wypowiedzi i odpowiednio do tego dwa rodzaje pomyłek: poprawne użycie (*correct use*) oraz poprawne zastosowanie (*correct application*) i odpowiadające im błędne użycie (*misuse*) oraz błędne zastosowanie (*misapplication*).

Różnicę między tymi dwoma rodzajami błędów pokazuje Millar w przekonująco sposób na następującym przykładzie (stanowiącym modyfikację przykładu podanego w innym kontekście przez Tylera Burge’a). Zestawmy ze sobą dwa

przypadki — pacjenta i lekarza. Pacjent nie zna społecznie akceptowanego znaczenia słowa „artretyzm” i sądzi, iż słowo „artretyzm” odnosi się do każdego schorzenia powodującego ból w kończynach. Jest skłonny wypowiedzieć zdanie „Cierpię na artretyzm”, gdyż jest prawdziwie przekonany, iż doskwiera mu ból w kończynach. Lekarz, dla odmiany, zna doskonale znaczenie tego słowa, stawia jednak błędną diagnozę — na skutek błędu w rozpoznaniu przypisuje omyłkowo naszemu pacjentowi artretyzm.

Zarówno pacjent, jak i lekarz, popełniają w jakimś sensie błąd — określają mianem „artretyzmu” coś, co artretyzmem nie jest. Gdyby analizować te przypadki wedle interpretacji norm semantycznych zaproponowanej przez Boghossiana, to okazałoby się, iż obaj użyli tego słowa niezgodnie z jego znaczeniem. Różnica między tymi pomyłkami jest jednak doskonale widoczna. Pacjent rzeczywiście myli się co do znaczenia słowa artretyzm — Millar określa to błędnym użyciem (*misuse*) wyrażenia. W przypadku lekarza natomiast nie można mówić o błędnym użyciu *tout court*, lecz o błędnym zastosowaniu (*misapplication*). Znajomość znaczenia danego słowa zobowiązuje użytkownika języka do prawidłowego użycia danego wyrażenia, nie zobowiązuje go jednak do jego prawidłowego stosowania. Millar twierdzi, iż zaproponowana przez Boghossiana interpretacja tezy o normatywności znaczenia jest niepoprawna, albowiem utożsamia (*conflates*) kwestię prawdziwości oraz właściwego zastosowania z kwestią właściwego użycia słów.

Hans-Johan Glock (2005: 233–234) podaje analogię, która dobrze naświetla różnicę pomiędzy omawianymi rodzajami normatywności i błędów. Zauważa on, iż w grach typu szachy możemy wyróżnić dwa rodzaje „reguł”: po pierwsze są to zasady gry *sensu stricto*, to znaczy przepisy określające, które posunięcia są poprawne, a które nie; drugą grupę reguł stanowią zasady „dobrej gry w szachy”, określające, jakie posunięcia w grze prowadzą do zwycięstwa. Oczywiście jest, że można grać w szachy poprawnie w pierwszym sensie, to znaczy przestrzegając reguł, i popełniać jednocześnie błędy — celowo bądź wskutek ignorancji —

w drugim sensie. Według Glocka, reguły semantyczne są analogiczne do reguł pierwszego rodzaju, wyznaczających zasady gry; analogonem reguł określających, które zagrania prowadzą do zwycięstwa, są z kolei reguły wskazujące sposoby poprawnego zastosowania danego słowa. Podobnie jak można grać w szachy zgodnie z regułami pierwszego typu, łamiąc zalecenia dotyczące wygrywania, tak samo można mówić poprawnie po polsku, popełniając błędy w stosowaniu konkretnych terminów.

Warto zwrócić w tym miejscu uwagę na istotną różnicę między podejściem Millara i Glocka a podejściem Boghossiana. W tym drugim ujęciu, dla każdego kontekstu istnieje dokładnie jedna wypowiedź, którą można określić jako poprawną. Konsekwencją utożsamienia reguł poprawnego stosowania i reguł poprawnego użycia jest to, iż reguły semantyczne w pewnym sensie nie pozostawiają użytkownikom języka wyboru — każde użycie języka, które łamie normy poprawnego zastosowania, jest niepoprawne językowo. Taka konsekwencja nie wypływa natomiast z interpretacji Millara i Glocka. Jeśli istnieją poprawne użycia niebędące poprawnymi zastosowaniami, nie ma powodu, by nie przyznać, iż więcej niż jedna wypowiedź może być poprawna w danym kontekście.

Wersja Millara i Glocka pozwala również rozwiązać problem normatywności wypowiedzi innych niż asercje — np. tych wyrażających preferencje. Intuicyjnie rzecz biorąc nie każde wyrażenie tego typu jest językowo poprawne, jakkolwiek pewna doza ekscentryczności, zwłaszcza w przypadku wypowiedzi ironicznych i quasi-ironicznych, jest dopuszczalna. Pewne wypowiedzi zdają się wszakże świadczyć o tym, iż użytkownik języka albo nie zna reguł rządzących użyciem danych wyrażen albo celowo je łamie. Istnieją wypowiedzi, których nie można zakwalifikować inaczej niż jako pomyłki językowe. „Chciałbym, aby karczoch z jabłkami był następnym prezydentem Polski” może być tutaj przykładem — o ile oczywiście nie potraktujemy tego zdania jako nadmiernie wysublimowanej ironii, co byłoby jednakże sporym nadużyciem zasady życzliwości semantycznej.

Chociaż podejście Millera-Glocka unika obydwu problemów, jakie niosła za sobą interpretacja Boghossiana, płaci ono za to pewną cenę. Koncepcja Boghossiana pozwalała określić, które wypowiedzi są poprawne, a które nie, i w nietrywialny sposób opisywała warunki poprawności użycia wyrażań. Tego samego nie da się powiedzieć o ujęciu Millara-Glocka. W ramach przeprowadzonej przez nich analizy nie jest łatwo sformułować jednoznacznych reguł wyznaczających poprawne użycie wyrażań. Ich ujęcie ma walor głównie negatywny — skupia się na wykazaniu, iż nie możemy utożsamiać poprawnego użycia z poprawnym zastosowaniem. Może się tu pojawić zarzut, iż mamy do czynienia z tego typu rozwiązaniem problemu filozoficznego, które polega na unikaniu trudności poprzez rezygnację ze ścisłego wysłowienia tezy — jeśli nie wiadomo, co dokładnie dany filozof chce powiedzieć, to trudno sformułować wobec jego poglądów jakiś zarzut.

Interpretacja Boghossiana zdaje się mieć początek w takiej filozoficznej wizji języka, wedle której podstawową formą wypowiedzi jest asercja, a znaczeniami wypowiedzi stwierdzających są ich warunki prawdziwości. Kripke pisze, iż tego typu koncepcja języka jest obecna (miedzy innymi) w „pierwszej” filozofii Wittgensteina, tej z *Traktatu logiczno-filozoficznego* (por. Kripke 1982/2007: 117–120). W *Dociekaniach filozoficznych* wizja języka jest już jednak zasadniczo odmienna — asercje tracą swój uprzywilejowany charakter. Wittgenstein wzbrania się przed uznaniem, jakoby istniejącą mnogość gier językowych dało się sprowadzić do jednego wzorca:

Lecz ile istnieje rodzajów zdań? (...) Istnieje *niezliczona* ich ilość: niezliczona ilość sposobów użycia tego wszystkiego, co zwiemy „znakiem”, „słowem”, „zdaniem”. I mnogość ta nie jest czymś stałym, raz na zawsze danym (...) (Wittgenstein 1953/2000 § 23).

Według „późnego” Wittgensteina, istnieją różne gry językowe, nieredukowalne do jednej podstawowej, połączone ze sobą jedynie relacją podobieństwa rodzinnego. Jeżeli zaakceptujemy tę diagnozę, wówczas nie sposób traktować stanowiska Millera-Glocka — odrzucającego tezę o tożsamości poprawnego stosowania i po-

prawnego użycia — jako skonstruowanego *ad hoc* narzędzia do obejścia problemów związanych z tezą o normatywności znaczenia. Przyjęcie stanowiska, iż nie da się odkryć ogólnej natury gier językowych, pociąga za sobą również sceptycyzm wobec nadziei na odszukanie ogólnego wzorca poprawności językowej.

Nie sposób w tym miejscu „udowodnić” słuszności podejścia odrzucającego prymat asercji oraz warunków prawdziwości. Warto jednak poczynić następującą obserwację: gdyby, jak chcą zwolennicy „warunkowo-prawdziwościowej” teorii znaczenia, poprawnymi były tylko te użycia słów i wyrażań, które są zarazem poprawnymi aplikacjami, to do znajomości języka wystarczyłaby znajomość interpretacji — w sensie odniesienia przedmiotowego — jednostek języka oraz jego wewnętrznej struktury. Innymi słowy, aby poprawnie posługiwać się językiem, wystarczyłoby znać słownik oraz gramatykę. Tymczasem praktyka pokazuje, iż, przynajmniej w przypadku języków naturalnych, do sprawnego posługiwania się nimi potrzebne są również inne umiejętności. Nie jest to argument rozstrzygający. Zwolennik ujęcia w kategoriach warunków prawdziwości zawsze może postulować istnienie „struktury głębokiej” opierającej się na warunkach prawdziwości. Co jednak istotne, przeświadczenie, że wzorce poprawności językowej są zbyt złożone, aby można je było uchwycić jedną prostą formułą, zdaje się zawierać w zbiorze naszych intuicji dotyczących funkcjonowania języka.

Warto zauważyć, że propozycja Boghossiana nie jest całkowicie chybiona. Istnieją konteksty, w których można zasadnie utożsamić poprawne użycie z poprawnym zastosowaniem. Ogólnie rzecz biorąc, są to te konteksty, w których mówiący chce dokonać szczerzej asercji, dysponując przy tym rzetelną wiedzą. Są to jednak szczególne przypadki, które nie mogą stać się podstawą uogólnienia.

II.4 Odrzucenie tezy o normatywności znaczenia

sensu stricto

Kolejną grupę krytyków idei normatywności znaczenia stanowią autorzy, którzy do pewnego stopnia akceptują wyjściową intuicję, że użycia języka dzielą się na poprawne i niepoprawne. Twierdzą oni jednak, że intuicja ta nie wystarcza do wyciągnięcia wniosku, jakoby zdania o znaczeniu były normatywne w pełnym znaczeniu tego słowa. Krytyki te nie ograniczają się do interpretacji pojęcia normatywności znaczenia dokonanej przez Boghossiana — w równym stopniu mogą one dotyczyć bardziej liberalnego podejścia, które zaproponowali Millar i Glock.

Pierwszym argumentem tego typu jest argument z imperatywów hipotetycznych, sformułowany m.in. przez Anandi Hattiangadi (2006: 228). Twierdzi ona, iż normy semantyczne nie są „prawdziwymi normami”, lecz imperatywami hipotetycznymi. Jak wiadomo, imperatywy hipotetyczne — w odróżnieniu od kategorycznych — określają, co „należy” zrobić w danej sytuacji, aby osiągnąć zamierzony cel. Typowymi przykładami imperatywów hipotetycznych są zdania w rodzaju: „Aby dojechać na dziesiątą rano do Warszawy, należy wsiąść w pociąg odjeżdżający z Krakowa o szóstej” czy „Jeśli chcesz bezpiecznie ulokować swoje oszczędności, powinieneś kupić obligacje”. Jakkolwiek zdania te zawierają terminy „normatywne”, takie jak „należy” czy „powinieneś”, nie mamy tu do czynienia ze zdaniami wyrażającymi rzeczywiste normy. Określają one jedynie środki służące do osiągnięcia jakiegoś celu; i tak, z pierwszego przykładu nie wynika norma, która miałaby komukolwiek mówić, że powinien jechać do Warszawy.

Uznanie, iż normy semantyczne są imperatywami hipotetycznymi, pozwala na wyjaśnienie, dlaczego używamy w kontekście znaczenia wyrażeń normatywnych, przy jednoczesnym zaprzeczeniu, iż mamy tu do czynienia z rzetelnymi normami. Czy jednak takie rozumienie jest trafne? Aby odpowiedzieć na to pytanie, warto

sobie uświadomić, jakie warunki musi spełniać zdanie, aby mogło być nazwane imperatywem hipotetycznym. Nie każde zdanie normatywne mające postać okresu warunkowego należy bowiem uznać za imperatyw hipotetyczny. Zdanie „rzetelnie normatywne” może mieć formę warunkową, jeśli określa normę, obowiązek lub powinność, które zachodzą w szczególnej sytuacji (Mackie 1977: 28–29). Rozważmy to na przykładzie prawniczym: zdania typu „Jeśli uczestnik ruchu drogowego spowoduje wypadek, powinien pozostać na miejscu zdarzenia do czasu przyjazdu policji” nie są imperatywami hipotetycznymi, wskazują bowiem na bezwzględny obowiązek, który spoczywa na osobie znajdującej się w danej sytuacji. Nawet jeśli zdanie warunkowe zawiera w poprzedniku określenie intencji podmiotu, nie oznacza to jeszcze, iż mamy do czynienia z imperatywem hipotetycznym; możliwe bowiem są takie normy, które określają, co podmiot powinien zrobić, jeśli posiada daną intencję, na przykład: „Jeśli ktoś zamierza nabyć broń palną, powinien zwrócić się o oficjalne pozwolenie”. Taka norma nie jest jednak praktyczną poradą — wystąpienie o pozwolenie może nawet utrudnić zdobycie broni — lecz obowiązkiem, który należy w danych okolicznościach wypełnić.

Jakie zatem są warunki uznania danego zdania za imperatyw hipotetyczny, jeśli nie decyduje o tym sama forma zdania warunkowego? Istotną cechą imperatywu hipotetycznego jest to, iż określa on zewnętrzny cel, którego osiągnięcie jest wynikiem stosowania się do zalecenia w nim zawartego. Zostało to wyraźnie podkreślone już przez Kanta, który wprowadził to pojęcie:

Hipotetyczny imperatyw powiada więc tylko, że czyn nadaje się do osiągnięcia jakiegokolwiek *możliwego* lub *rzeczywistego* celu (Kant 1785/2001: 32).

Aby uzasadnić twierdzenie, iż normy semantyczne są imperatywami hipotetycznymi, konieczne jest zatem wskazanie zewnętrznego celu, do osiągnięcia którego prowadziłoby przestrzeganie reguł poprawnego użycia języka. Hattiangadi uważa, że takim zewnętrznym celem jest mówienie prawdy (Hattiangadi 2006:

231). Jej zdaniem, normy semantyczne mają następującą postać: „Jeśli p chce mówić prawdę, to p powinien używać wyrażeń językowych w taki-a-taki sposób”. Propozycja Hattiangadi podobna jest do interpretacji pojęcia normatywności znaczenia zaproponowanej przez Boghossiana. Z chwilą, gdy odrzucimy założenie, że prawda jest celem aktów mowy, i zgodzimy się, iż istnieją poprawne użycia języka niebędące jednocześnie poprawnymi zastosowaniami, musimy także przyznać, iż przestrzeganie reguł poprawnego użycia nie jest środkiem służącym do mówienia prawdy. Jeśli bowiem stosowanie się do jakiejś reguły może, ale nie musi prowadzić do osiągnięcia zamierzonego celu, to nie ma podstaw, by twierdzić, iż mamy do czynienia z imperatywem hipotetycznym.

Niepowodzenie tej konkretnej koncepcji starającej się określić zewnętrzny cel poprawnych wypowiedzi, nie dyskwalifikuje, rzecz jasna, teorii mówiącej, iż normy semantyczne są imperatywami hipotetycznymi. Alternatywną propozycję odpowiedzi na pytanie, do czego miałyby służyć kierowanie się regułami semantycznymi, przedstawił Alex Miller. Według niego, naturalnym celem, do którego dążymy, kierując się regułami semantycznymi, jest chęć porozumiewania się z innymi ludźmi (Miller 2006: 109). Zaletą propozycji Millera jest to, iż nie wymaga ona utożsamienia warunków poprawnego użycia z warunkami poprawnego zastosowania. To, że można „poprawnie” używać wyrażeń, dokonując jednocześnie błędnego zastosowania, nie wyklucza możliwości rozpatrywania „poprawności” użycia w kategoriach efektywności komunikacji.

Aby właściwie ocenić propozycję Millera, trzeba zdać sobie sprawę z jednej, dość oczywistej własności imperatywów hipotetycznych. Otóż gdyby sposób postępowania zalecany w następniku imperatywu okazał się nieskuteczny w realizacji celu określonego w poprzedniku, to cały imperatyw należałoby odrzucić. Jeśli zaś można wykazać, iż, pomimo swej nieskuteczności, norma ta nadal pozostaje w mocy, to analiza wyjaśniająca obowiązywanie owej normy poprzez potraktowanie jej jako następnika imperatywu hipotetycznego, byłaby niewłaściwa. Jeżeli zdanie

„Jeśli ktoś zamierza nabyć broń palną, powinien zwrócić się o oficjalne pozwolenie” mielibyśmy interpretować jako imperatyw hipotetyczny, to należałoby odczytywać je w następujący sposób: wystąpienie o oficjalne pozwolenie na posiadanie broni palnej stanowi środek do celu, jakim jest nabycie tejże. Jeśli natomiast wystąpienie o oficjalne pozwolenie byłoby w praktyce wysoce nieskutecznym środkiem do osiągnięcia tego celu — w porównaniu na przykład z nabyciem broni na czarnym rynku — to wówczas taki domniemany imperatyw hipotetyczny można by zakwestionować. W tym przypadku norma nakazująca wystąpienie o pozwolenie na broń pozostaje w mocy — fakt, iż nie prowadzi ona do zamierzonego rezultatu, nie wpływa na jej obowiązywanie.

Podobny przykład można podać dla norm semantycznych. Załóżmy, że normy te dyktują, iż w sytuacji, kiedy chcemy, aby ktoś podał nam stół, zamierzamy otwarcie wyrazić takie polecenie oraz posiadamy rzetelną wiedzę na temat tego, że przed nami znajduje się stół, to powinniśmy powiedzieć „Przynieś mi stół”. Jednak osoba, której chcemy wydać to polecenie, słabo zna język polski i z tego względu wypowiedzenie zdania „Przynieś mi stół” będzie nieskuteczne. Być może w danej sytuacji wypowiedzenie zdania „Przynieś mi krzesło” okaże się efektywniejszym środkiem komunikacji. Wydaje się jednak jasne, iż w omawianym przypadku tym, co „należało” powiedzieć — z punktu widzenia poprawności językowej — była fraza „Przynieś mi stół”. Tego typu przykłady można mnożyć — w wielu sytuacjach użycie frazy, którą uznalibyśmy za niepoprawną w danym kontekście, może być najskuteczniejszą metodą osiągnięcia porozumienia.

Jeśli obserwacja ta jest trafna, to nie można uznać norm semantycznych za imperatywy hipotetyczne mówiące nam, jak skutecznie się komunikować z innymi. Prawdą jest, że najczęściej stosowanie się do tych reguł prowadzi nas do porozumienia. Nie jest tak jednak zawsze — i to wystarczy, aby wykazać, że normy semantyczne nie są imperatywami hipotetycznymi. Oczywiście, obalenie dwóch koncepcji wskazujących zewnętrzny cel kierowania się regułami semantycznymi

nie stanowi niepodważalnego dowodu, iż takiego zewnętrznego celu nie da się podać. Jednakże *onus probandi* leży po stronie proponentów tezy o regułach semantycznych jako imperatywach hipotetycznych. Wydaje się, iż problemy, przed jakimi stoją omówione wyżej teorie, sugerują, że kierowanie się regułami semantycznymi, choć może prowadzić do osiągnięcia zewnętrznych celów, jest wartością autoteliczną. Wartość poprawnego używania języka nie sprowadza się, w myśl naszych intuicji, do praktycznych efektów.

Wedle innej koncepcji, przyjmującej do pewnego stopnia intuicje związane z tezą o normatywności znaczenia, lecz odrzucającej wniosek, że zdania z dyskursu semantycznego są „rzetelnie normatywne”, powinno się te zdania interpretować jako reguły konstytutywne. Teorię tę sformułowali Katrin Glüer i Peter Pagin (1999: 214–225, Glüer 2001: 65–72) oraz Åsa Wikforss (2001: 215). Ma ona strukturę podobną do argumentu z imperatywów hipotetycznych: proponuje interpretację zdań o znaczeniu, która ma wyjaśnić naszą skłonność do używania w stosunku do nich określeń normatywnych, nie prowadząc do wniosku, iż są to rzeczywiście zdania wyrażające normy.

Koncepcja ta zasadza się na znanym rozróżnieniu Johna Searle’a na reguły regulatywne i konstytutywne (Searle 1965/1996: 265–266). Te pierwsze mówią, jak powinniśmy się zachowywać w danych, już istniejących kontekstach. Dobrym przykładem są reguły grzeczności, wskazujące, jak powinniśmy się zachowywać w pewnych kontekstach interpersonalnych. Reguły konstytutywne określają z kolei, co jest daną czynnością w danym kontekście; nie regulują istniejących czynności, ale tworzą nowe. Przykładem może być zasada określająca, co kwalifikuje się jako „roszada” w szachach. Nie mówi ona, czy i kiedy wykonywać dane posunięcia na szachownicy, nie reguluje też uprzednio i niezależnie istniejących czynności — coś takiego jak „roszada” w ogóle nie mogłoby istnieć, gdyby nie istniała odpowiednia reguła konstytutywna.

Zgodnie z tezą o konstytutywnym charakterze reguł rządzących użyciem wyrażen, używanie na przykład słowa „zielony” w określony sposób w określonych okolicznościach kwalifikuje się jako wyrażanie pojęcia „zielony” (Glüer i Pagin 1999: 216). W tym sensie można mówić o istnieniu reguł „poprawnego” użycia wyrażen. Przyjąwszy jednak taką interpretację, trudno argumentować na rzecz „normatywności” zdań o znaczeniu. Reguły konstytutywne nie mówią bowiem nikomu, jak powinien postępować. Wskazują tylko, iż pewnego rodzaju czynności w pewnym kontekście kwalifikują się w pewien sposób; nie uprawniają do wniosku, iż ktokolwiek kiedykolwiek powinien je wykonywać.

Porównanie zasad języka do zasad gry w szachy jest uproszczeniem, z którego *notabene* jego autorzy zdają sobie sprawę. Analogia zawodzi pod jednym istotnym względem: w szachach niemożliwe jest „złamanie zasad”. Jeśli ktoś wykona ruch niezgodny z zasadami, to ten ruch po prostu nie należy do gry. Wydaje się, że w języku rzecz ma się inaczej — nie każda pomyłka językowa świadczy o tym, że nie używam danego pojęcia. Jeśli w sytuacji, w której chcę określić rasę danego psa, wygłoszę szczere stwierdzenie „to jest wyżeł”, mając przed sobą charta — by użyć przykładu, w którym można utożsamić warunki poprawnego użycia z warunkami poprawnego zastosowania — nie oznacza to jeszcze, że przestałem się posługiwać pojęciem „wyżeł”, a zacząłem posługiwać się jakimś innym, na przykład pojęciem „wyżłarta”. Intuicyjna interpretacja tego zdarzenia jest taka, iż zwyczajnie pomyliły mi się znaczenia słów „wyżeł” i „chart”. Język pod tym względem, bardziej niż szachy, przypomina gry takie jak koszykówka czy piłka nożna, w których możliwe jest złamanie reguł „w obrębie gry”. Gdy ktoś popełnia błąd kroków w koszykówce, to łamie reguły gry, ale nie przestaje przez to grać w koszykówkę. Co więcej, same reguły gry przewidują możliwość „nieprzepisowego zagrania”. Podobnie jest z językiem: gramatyka zdań określających poprawność użycia wydaje się, co zauważają również Glüer i Pagin (1999: 221), dopuszczać popełnienie pomyłki w stosowaniu danego wyrażenia.

Alternatywna interpretacja reguł konstytutywnych rozważana przez wyżej wymienionych autorów głosi, że pozostawanie w zgodzie z większością lub z odpowiednio dużą liczbą reguł danej gry konstytuuje fakt, iż się w nią gra. Podobnie w przypadku języka — pozostawanie w zgodzie z większością lub wystarczająco dużą liczbą reguł językowych miałoby konstytuować fakt, iż mówimy po polsku. Nie wydaje się jednak, aby taka koncepcja stała w sprzeczności z przekonaniem, że reguły językowe są normatywne — w obrębie gry nadal można mówić o poprawności bądź niepoprawności danego wyrażenia. Ta sama reguła może przecież pełnić dwojaką rolę: konstytutywną — decydującą wraz z pozostałymi regułami należącymi do danego zbioru, czy ogół pewnych zachowań *en bloc* spełnia wyznaczone mu standardy, oraz regulatywną — określającą, czy pewne zachowanie jest poprawne w danym kontekście. Ostatecznie jednak Glüer i Pagin odrzucają i tę interpretację reguł konstytutywnych. Ich zdaniem, z czym trudno się nie zgodzić, takie podejście prowadzi do problemów związanych z tożsamością gier. Jeśli granie w daną grę polega na kierowaniu się częścią konstytuujących ją reguł, to trudno *de facto* stwierdzić, w którą grę gramy (Glüer i Pagin 1999: 221).

Ostatecznie Glüer i Pagin proponują teorię, wedle której charakterystyczną cechą reguł konstytutywnych jest to, iż aby dana czynność mogła się odbywać, pewne reguły muszą pozostawać w mocy. Na przykład, aby dana gra klasyfikowała się jako gra w szachy, w trakcie rozgrywania partii reguły gry w szachy muszą pozostawać w mocy. Warto zauważyć, iż rozwiązanie to nie jest sprzeczne z „potoczną” koncepcją normatywności znaczenia. Taka wizja „konstytutywnego” charakteru norm semantycznych dopuszcza bowiem istnienie poprawnych oraz niepoprawnych użyć wyrażen, nawet jeśli pomyłki popełniane są systematycznie. Glüer i Pagin twierdzą, iż podejście to uniemożliwia natomiast uznanie reguł znaczeniowych za „rzetelnie normatywne”, gdyż nie stanowią one — a przynajmniej nie muszą stanowić — elementu tych rozumowań przeprowadzanych przez jednostkę, które motywują jej działania. Argument ten wydaje się jednak niezależny od ujmowania

reguł semantycznych jako konstytutywnych. Można zatem stwierdzić, iż argument, w którym uznaje się reguły językowe za reguły konstytutywne, nie wytrzymuje krytyki.

Warto w tym miejscu zauważyć, iż tezę o konstytutywnym charakterze reguł znaczeniowych postawił o wiele wcześniej na polskim gruncie Kazimierz Ajdukiewicz (1934/1985: 151–153). W przedstawionej przez niego „dyrektywnej koncepcji znaczenia” dyrektywy językowe były *explicite* określone mianem „konstytutywnych”. Zdaniem Ajdukiewicza, jeśli stykamy się z kimś, kto łamie dyrektywy językowe, to możemy stwierdzić, iż używa on słów w innym znaczeniu niż my. Wydaje się, iż przytoczone powyżej argumenty przeciwko koncepcji Pagina i Glüer uderzają również w tę koncepcję. Kwestia tego, jaka interpretacja Ajdukiewicza jest poprawna, zostanie jednak pozostawiona na boku — słusznie czy nie, koncepcja ta nie odgrywa poważniejszej roli we współczesnych debatach.

Kolejny argument kwestionujący zasadność mówienia o normatywności w przypadku zdań o znaczeniu — argument z internalizmu motywacyjnego — ma strukturę zasadniczo odmienną od tych omawianych wcześniej. Argumenty z imperatywów hipotetycznych i reguł konstytutywnych miały pokazać, dlaczego zdania o znaczeniu uznajemy za normatywne, choć *de facto* nie są one „rzetelnie normatywne”. Argument z internalizmu motywacyjnego nie podaje alternatywnego odczytania „normatywności znaczenia”. Wskazuje się tu natomiast na fakt, iż zdaniom o znaczeniu brakuje pewnej istotnej cechy, przysługującej stwierdzeniom „rzetelnie normatywnym” — przekonania dotyczące tych stwierdzeń nie są wewnętrznie motywujące.

Załączki argumentu z internalizmu motywacyjnego pojawiają się u wielu autorów. Hattiangadi (2006: 231–232) uważa na przykład, że nad przekonaniami dotyczącymi norm prawidłowego użycia wyrażen mogą wziąć górę inne pragnienia żywione przez podmiot. Innymi słowy, wystarczy, by użytkownik języka miał inne intencje niż ta, by wyrażać się poprawnie, a nie będzie ciążył na nim obowiąz-

zek stosowania się do reguł semantycznych. Zarzut ten wydaje się dość niejasny. Można go interpretować na co najmniej dwa sposoby. W pierwszej wykładni, użytkownik języka nie musi czuć się zobligowanym do używania języka poprawnie — w jego systemie motywacyjnym mogą przeważać inne pragnienia, na przykład chęć komunikacji. Jest to obserwacja bezsporna, aczkolwiek nierелеwantna. Fakt, iż ktoś nie czuje się zobligowany do przestrzegania danej reguły, nie stanowi jeszcze argumentu za tym, że owa reguła nie obowiązuje. Nawet jeśli większość kierowców nie „czuje”, iż powinna przestrzegać kodeksu drogowego, nie oznacza to jeszcze, iż kodeks ów nie wyznacza obowiązujących norm.

Przy bardziej życzliwej interpretacji argument Hattiangadi przedstawia się następująco: intencja, która motywuje nas do posługiwania się językiem w sposób odmienny od „poprawnego”, wystarcza, aby zwolnić nas z obowiązku posługiwania się językiem poprawnie. Jeśli rzeczywiście by tak było, to teza o normatywnym charakterze zdań o znaczeniu stawałaby się wątpliwa; trudno bowiem mówić, że mamy do czynienia z rzeczywistą normą w sytuacji, gdy intencja podmiotu wystarczy, by ową normę uchylić. Ta linia rozumowania wydaje się jednak niedostatecznie uzasadniona; autorce udaje się co prawda wykazać psychologiczny fakt, iż odmienne motywacje mogą przeważać nad chęcią mówienia poprawnie, ale nie wynika z tego, iż normy językowe nie obowiązują. Podkreślił to, polemizujący z Hattiangadi, Daniel Whiting:

Prawdą jest, iż mogę nie kierować się regułą z tego tylko powodu, iż nie mam na to ochoty. Lecz nie dowodzi to jeszcze, że nie obowiązuje żadna norma; moje użycie wyrażenia wciąż powinno być uznane za *niepoprawne*. Oczywiście, wykroczenie nie jest zbyt poważne (pomyłka ma charakter semantyczny, nie etyczny), jednak nie prowadzi to do wniosku, iż nie posiada ono normatywnego charakteru (Whiting 2007: 139).

Whiting twierdzi, iż normy semantyczne są normami *prima facie*, co znaczy, że mogą zostać unieważnione przez inne normy, na przykład etyczne, dobrego wychowania etc. Jeśli norma etyczna wymaga ode mnie, bym użył słowa niepoprawnie, to wówczas moje zachowanie jest, *summa summarum*, słuszne. Nie oznacza to jednak braku norm semantycznych. Można argumentować, iż wiele norm etycznych jest również normami *prima facie*, na przykład norma nakazująca poszanowanie własności może zostać unieważniona przez normę nakazującą ratowanie życia. Wydaje się zatem, iż argumentu Hattiangadi nie sposób traktować jako przekonującego argumentu przeciw normatywności znaczenia.

Argumentację przeciwko normatywności znaczenia opierającą się na kategorii motywacji podali również Glüer i Pagin (1999: 223–224). Jak zostało już powiedziane, przyjmują oni, iż reguły znaczeniowe są regułami konstytutywnymi: aby dane słowo miało znaczenie, reguły jego obowiązywania muszą pozostawać w mocy. Pozornie nie ma powodu, dla którego teza ta miałaby stać w sprzeczności z przekonaniem o normatywnym charakterze zdań o znaczeniu. Glüer i Pagin twierdzą jednak, iż tak rozumiane reguły znaczeniowe nie spełniają funkcji „kierowania” użytkownikami języka. Dana osoba może wiedzieć, jaki sposób posługiwania się słowem jest tym właściwym, ale nie uznawać tego za przesłankę podejmowanych przez siebie decyzji — tak samo, jak zawodnik grający w piłkę nożną może doskonale zdawać sobie sprawę, że uderzenie piłki ręką jest niedozwolone, a jednocześnie zdecydować się zignorować tę zasadę.

Obserwacje poczynione przez Glüer i Pagina wydają się słuszne. Pytanie o wnioski z nich płynące pozostaje jednak otwarte. Można się zgodzić z wnioskiem, że kategoria „kierowania się regułami” jest nieszczególnie użyteczna z punktu widzenia analizy zdań o znaczeniu, gdyż sugeruje, jakoby istotną ich cechą była zdolność do wpływania na użytkowników języka w zakresie ich decyzji. Tymczasem, dla zdań o znaczeniu istotne jest to, że „pozostają w mocy” niezależnie od tego, czy użytkownik bierze je pod uwagę przy podejmowaniu decyzji, czy też

tego nie robi. Nie ma jednak powodu, by to spostrzeżenie uznać za przesłankę do odrzucenia tezy o normatywności znaczenia.

Brakujące ogniwo tej argumentacji można odnaleźć u Alexa Millera, który połączył rozważania na temat reguł semantycznych z istotnym w dziedzinie metaetyki problemem internalizmu motywacyjnego (Miller 2009). Internalizm motywacyjny, zwany również internalizmem etycznym (por. np. van Roojen 2008, Saja 2007: 180), jest stanowiskiem głoszącym, iż istnieje konieczny związek między sądami moralnymi a motywacją. Ujmując rzecz w dużym uproszczeniu, internalista twierdzi, iż osoba, która uznaje jakiś sąd etyczny, jest równocześnie zmotywowana do postępowania w zgodzie z nim. Jeśli na przykład uważam, że należy przekazywać pieniądze organizacjom dobroczynnym, to będę miał motywację, aby tak właśnie czynić. Różne odmiany internalizmu różnie interpretowały charakter tej zależności — jako prawidła psychologicznego, zasady ontologicznej lub faktu dotyczącego „gramatyki” zdań etycznych (Saja 2007: 181). Często mówi się też o ograniczeniach zakresu stosowalności tej teorii, na przykład tylko do podmiotów „praktycznie racjonalnych” (Miller 2009), czyli takich, które są w stanie kierować własnym postępowaniem — w przeciwieństwie na przykład do chorych na poważną depresję, którzy nie przejawiają standardowych zachowań motywacyjnych. Abstrahując od różnic między poszczególnymi odmianami tej koncepcji, można stwierdzić, iż tezą istotną dla internalizmu motywacyjnego jest istnienie związku między uznawaniem sądów moralnych a posiadaniem motywacji.

Gdyby zastosować tezę internalizmu motywacyjnego do konkretnego przypadku zdania o znaczeniu, to brzmiałaby ona następująco:

(MI) Jest prawdą konieczną i analityczną (*conceptual*), że jeśli podmiot działający S sądzi, iż „sroka” znaczy *sroka*, to, wyjąwszy przypadek, w którym S nie jest osobą praktycznie racjonalną, S będzie posiadał motywację do używania słowa „sroka” zgodnie ze wzorami poprawnego użycia (Miller 2009).

Miller odrzuca tę tezę, lecz zauważa, iż jeśli odrzucimy interpretację Boghosiana (zrównującą poprawne użycie z poprawnym zastosowaniem), to nie będzie łatwe skonstruowanie dla niej kontrprzykładu. Z tego powodu przyjęta przez niego strategia argumentacyjna opiera się na próbie wykazania, iż argumenty tradycyjnie przytaczane na rzecz internalizmu motywacyjnego nie zachowują swej mocy, gdy zastosujemy je do zdań o znaczeniu.

Internalizm kontrastowany jest zwykle z eksternalizmem, co w odniesieniu do dyskursu moralnego ma służyć pokazaniu, iż ten pierwszy lepiej odpowiada naszym intuicjom dotyczącym funkcjonowania przekonań moralnych. Wedle eksternalistów, poza akceptacją pewnych sądów moralnych, niezbędne do wystąpienia motywacji jest również istnienie w podmiocie uogólnionego pragnienia, aby czynić dobrze. Konsekwencją eksternalizmu jest zatem dopuszczenie możliwości, że podmiot będzie szczerze przekonany o tym, iż dany sposób postępowania jest właściwy, lecz nie będzie miał motywacji, by tak postępować. Podmiot ten jest w literaturze określany mianem „amoralisty”, spór zaś dotyczy tego, czy taka postawa jest możliwa. Według Millera, szczególnie przekonującą strategią internalisty jest pokazanie, iż eksternalizm, postulując konieczność owej uogólnionej motywacji, nieprzekonująco tłumaczy zachowanie podmiotu moralnego (*morally virtuous person*).

To, czy internaliści mają w tej sprawie rację, jest osobną kwestią. Ważne jest, iż eksternalizm wydaje się zupełnie przekonujący w przypadku zagadnienia dotyczącego poprawnego użycia słów. Sama wiedza na temat tego, jak poprawnie używać słów, nie stanowi dla normalnego podmiotu wystarczającej motywacji, by posługiwać się nimi we właściwy sposób. Konieczna jest dodatkowa przesłanka motywacyjna — pragnienie posługiwania się językiem poprawnie. Wbrew Millerowi, dla zdań o znaczeniu można skonstruować rzetelny kontrprzykład dla tezy internalisty; przywoływani powyżej Glüer i Pagin wskazywali na możliwość świadomego występowania przeciw regułom językowym. Wyobraźmy sobie ko-

gość, kto doskonale wie, iż prawidłowym sposobem używania słowa „bynajmniej” w języku polskim jest stosowanie go do wzmocnienia przeczenia. Jednocześnie osoba ta może zupełnie świadomie twierdzić, iż reguła ta wcale jej nie interesuje i z upodobaniem używać tego słowa w wypowiedziach typu „Ja bynajmniej jestem kawalerem”. Werdykt jest jednoznaczny: teza internalizmu motywacyjnego w odniesieniu do zdań o znaczeniu jest fałszywa. W przypadku zdań semantycznych nie ma związku między uznawaniem ich a motywacją, by posługiwać się słowami we właściwy sposób. Czy uprawnia nas to jednak do twierdzenia, iż zdania te nie wyrażają norm?

Wydaje się, że argument ten opiera się na wąskim rozumieniu pojęcia normatywności, wedle którego jedynie zdania wyrażające oceny moralne mają status „normatywnych”; jeśli zaś jakimś zdaniom nie przysługują cechy zdań wyrażających oceny moralne, to nie mamy prawa nazywać ich „normatywnymi”. Takie ograniczenie pojęcia normatywności, którego *notabene* Miller nie dokonuje jednoznacznie, wydaje się zbyt restrykcyjne. Istnieje bowiem wiele rodzajów zdań nie mieszczących się w sferze dyskursu moralnego, do których nie stosuje się teza internalizmu, a określanych mianem normatywnych. Można tu zaliczyć normy prawne, nakazy grzeczności, tabu kulturowe i wiele innych. Glock (2005: 232–233) jako przykład zdania normatywnego przytacza normę określającą, iż profesoria uniwersyteccy powinni w dniu inauguracji roku akademickiego występować w togach. Wydaje się, że można niekontrowersyjnie nadać takiemu zdaniu miano „normatywnego”, chociaż nie stosuje się do niego teza internalizmu motywacyjnego.

Powyższe rozważania pokazują, iż pojęcie normatywności cechuje się pewną dwuznacznością: mamy do czynienia z dwoma rodzajami zdań nazywanych mianem normatywnych. Do pierwszej grupy należą zdania z dyskursu moralnego, które można nazwać mocno normatywnymi, do drugiej zaś zdania słabo normatywne, w odniesieniu do których teza internalizmu motywacyjnego nie stosuje się. Zali-

czenie zdań o znaczeniu do drugiej kategorii nie stoi w sprzeczności z „potoczną” koncepcją normatywności znaczenia, zarysowaną w pierwszej części niniejszej pracy. Wprowadzenie takiego rozróżnienia ma jednak istotne konsekwencje dla zaprezentowanego wcześniej ogólnego argumentu, który miał pokazywać, że z normatywności znaczenia wynika Nonfaktualizm — wątki te zostaną omówione dokładniej w *Rozdziale III*.

II.5 Próba zakwestionowania potocznej koncepcji normatywności znaczenia

Do ostatniej kategorii zarzutów wobec tezy o normatywności znaczenia należą te, które kwestionują posiadanie przez użytkownika języka intuicji dotyczących poprawnych i niepoprawnych użyć wyrażen. Najbardziej znana i zarazem wyrazista argumentacja tego typu została przedstawiona przez Donalda Davidsona w artykułach *Nice derangement of epitaphs* oraz *Social aspect of language*. Davidson opiera swoją krytykę na rozbudowanej teorii komunikacji językowej, której podstawowe założenie mówi, iż celem jej uczestnika jest interpretacja wypowiedzi rozmówcy. Interpretacja jest tu rozumiana podobnie jak w semantyce formalnej: podać interpretację jakiejś wypowiedzi to podać warunki prawdziwości dla zdań oraz przyporządkować zakresy predykatom i przedmioty indywidualne nazwom własnym (Davidson 1986/2005: 100–101).

Użytkownik języka dysponuje teorią uprzednią (*prior theory*) — zbiorem oczekiwań dotyczących tego, w jaki sposób jego interlokutor będzie się posługiwał językiem i jak będzie można zinterpretować jego wypowiedź. Założeniem teorii uprzedniej jest zazwyczaj to, że rozmówca będzie posługiwał się „standardowym” językiem naturalnym, na przykład językiem polskim. To oczekiwanie nie jest spełnione w każdym przypadku; może się zdarzyć — wedle Davidsona zdarza się tak bardzo często — iż mój rozmówca będzie się posługiwał językiem inaczej, niż się

spodziewałem. W takiej sytuacji porozumienie jest nadal możliwe, pod warunkiem, że zmienię swoją interpretację. W procesie dostosowania jej do nowych danych muszę kierować się zasadą życzliwości semantycznej — to znaczy przyjmować, że przekonania wyrażane przez mojego interlokutora są w większości prawdziwe oraz wewnętrznie spójne. Dzięki temu mogę stworzyć tymczasową teorię (*passing theory*) znaczenia dla języka, którym posługuje się mój partner. Jak podkreśla Davidson (1986/2005: 102-104), to zgodność tymczasowych teorii rozmówców decyduje o sukcesie komunikacji. W punkcie wyjścia mogą oni dysponować zupełnie odmiennymi teoriami uprzednimi, nie wyklucza to jednak możliwości późniejszego uzgodnienia interpretacji. Komunikacja polega zatem na nieustannym uzgadnianiu teorii tymczasowych.

Zarysowana tu w uproszczeniu teoria języka Davidsona nie pozostawia de facto miejsca na rozróżnienie na poprawne i niepoprawne użycia języka. Wyobraźmy sobie osobę, która używa słowa „bynajmniej” jako wzmocnienia twierdzenia, na przykład „Ja bynajmniej jestem kawalerem”. Wedle normatywnej koncepcji języka, tego typu użycie należy zakwalifikować jako „niepoprawne”. Zdaniem Davidsona, natomiast, to, że ktoś używa słowa w nieoczekiwany sposób, jest dla nas wskazówką do zrewidowania naszej teorii jego języka. Tak długo, jak jesteśmy w stanie rozmówcę zinterpretować, czyli dostosować naszą teorię do nowych danych, nie powinniśmy uważać jego sposobu wyrażania się za niewłaściwy. Ujmując rzecz lapidarnie, w teorii Davidsona nie ma niepoprawnych użyć języka, pewne z nich są tylko zaskakujące.

Davidson przywołuje dwa fenomeny na poparcie swojej koncepcji: malapropizmy oraz artystyczną inwencję językową (Davidson 1986/2005: 89–94). Przytacza on przykłady użyc malapropizmów w amerykańskich komediowych słuchowiskach radiowych. Za ich polski odpowiednik można uznać słuchowisko „Z pamiętnika młodej lekarki”, pełne wypowiedzi w rodzaju: „Dolega mię, pani doktor, prostota.” „Rozumię.” „Je pan kuraka rękami albo mówi pan ościennemu premierowi,

że tańczy czoczołkie na grobie socjalizmu”. Davidson zauważa, że przeciętny użytkownik języka nie ma problemu z uchwyceniem sensu takich wypowiedzi, chociaż są one czasem całkowicie niepoprawne z punktu widzenia reguł językowych. Jego zdaniem, kierowanie się regułami językowymi nie jest niezbędne, aby zostać zrozumianym. Co więcej, ujęcie języka w kategoriach reguł semantycznych ma tę słabość, iż nie wyjaśnia ani powstawania, ani rozumienia malapropizmów.

Drugim zjawiskiem, wobec którego — wedle Davidsona — zwolennicy ujmowania języka w kategoriach reguł semantycznych stają się bezradni, jest fenomen literackiej inwencji językowej. Davidson przytacza w tym kontekście twórczość Joyce’a, ale w polskiej literaturze również można odnaleźć wiele przykładów łamania konwencji językowych, choćby wiersze Leśmiana. Odrzuca on pogląd, iż tego typu użycia stanowią margines komunikacji językowej. Uważa, że tak malapropizmy, jak i inwencja literacka, są zjawiskiem dosyć częstym i istotnym.

Zdaniem Davidsona, interpretowanie języka w terminach reguł semantycznych nie jest warunkiem ani koniecznym, ani wystarczającym do opisywania komunikacji. Na tej podstawie wygłasza on niezwykle mocną na pierwszy rzut oka tezę o „nieistnieniu języka”:

(...) nie istnieje taka rzecz jak język — jeśli język miałby być tym, czego istnienie postulowało wielu filozofów i językoznawców. Nie istnieje nic, czego można by się uczyć, co można by opanowywać lub co może być wrodzone. Musimy pozbyć się idei ściśle określonej, wspólnej struktury, którą użytkownicy języka nabywają, by później stosować ją do konkretnych przypadków (Davidson 1986/2005: 107).

Czy jednak deklaracje te nie stoją w sprzeczności ze znanymi uwagami Davidsona (1970/1992: 187–188) z zakresu filozofii umysłu? Twierdzi on przecież, między innymi, że fundamentalną rolę w przypisywaniu innym podmiotom stanów mentalnych, a więc i zdolności do posługiwania się językiem, odgrywają normatywne przesłanki wynikające z ideału racjonalności. Według Davidsona, przypisywanie

komuś innemu posiadania przekonań jest możliwe tylko wówczas, gdy spełnia on pewne kryteria, na przykład gdy jego przekonania można zrekonstruować w niesprzeczny sposób w zgodzie z regułami klasycznej logiki.

Jak słusznie zauważyła Åsa Wikforss (2001: 215), powyższe rozumienie racjonalności nie jest sprzeczne z odrzuceniem istnienia reguł semantycznych i sprowadzeniem znaczenia do użycia. Podejście Davidsona do reguł językowych jest wręcz konsekwencją przyjęcia takiego poglądu. Ideał racjonalności, na którym zasadza się, według jego teorii, rozpoznawanie innego jako posiadającego intencje i posługującego się językiem racjonalnego podmiotu, ma charakter konstytutywny w ścisłym tego słowa znaczeniu. Jeśli podmiot spełnia warunki opisywane przez ów ideał, to jest tym samym podmiotem racjonalnym, jeśli zaś ich nie spełnia, to *ex definitione* podmiotem racjonalnym nie jest. Trudno, przyjąwszy taką koncepcję, utrzymywać, że ktokolwiek w jakimkolwiek sensie „powinien” być racjonalny. Przeciwnie, można argumentować, iż jeśli „coś” nie spełnia warunków „konstytutywnego ideału racjonalności”, to nie można w odniesieniu do tego czegoś mówić o żadnej powinności. Co więcej, to „coś” nie może dokonywać wyborów; jeśli „coś” nie jest racjonalnym podmiotem, to nie można przypisywać mu „poprawnego” bądź „niepoprawnego” zachowania.

Jeśli zaś podmiot spełnia wyznaczone przez ideał racjonalności warunki, to trzeba przyznać, że jest racjonalny i posiada określone intencje. Jego wypowiedzi należy interpretować kierując się zasadą życzliwości semantycznej. Celem interpretacji jest odkrycie intencji wyrażanych przez mówiącego — niezależnie od tego, jak bardzo ekstrawaganckie byłyby środki, którymi się posługuje. Klasyfikowanie wypowiedzi jako poprawnych bądź niepoprawnych jest zatem zbędne.

Intuicja poprawności językowej jest fikcją, odzwierciedlającą nasze oczekiwania względem zachowania językowego innych. Złudzenie normatywności bierze się w myśl tej koncepcji stąd, że w sytuacjach społecznych jesteśmy często korygowani przez innych użytkowników języka (Davidson 1994/2005: 117). W takich

sytuacjach mamy skłonność do zmiany własnego zachowania tak, by było zgodne z oczekiwaniami innych. Nie świadczy to jednak, iż mamy tu do czynienia z „rzetelnymi normami” — mówienie w sposób inny niż społecznie akceptowany jest pewnego rodzaju towarzyskim *faux pas*, tak samo jak w pewnych kręgach *faux pas* byłoby jedzenie ryby nożem lub mówienie „smacznego”. Złudzenie normatywności znaczenia powstaje zatem w wyniku presji społecznej i naszej tendencji do konformizmu.

Argumentacja Davidsona jest najbardziej konsekwentną i najdalej idącą krytyką tezy o normatywności znaczenia. Czy jego zarzuty są jednak słuszne? Wskazane przez Davidsona fenomeny malapropizmów i literackiej inwencji językowej nie falsyfikują bynajmniej tezy o istnieniu reguł językowych — zwraca na to uwagę m. in. Karen Green (2001: 250–252). Możliwe jest wyjaśnienie faktu, że rozumiemy tego typu wypowiedzi, uwzględniające istnienie reguł semantycznych. Może być ono oparte chociażby na teorii implikatur konwersacyjnych Grice’a. Pozostawiając na boku psychologistyczną tezę tego autora, utożsamiającą znaczenie z intencjami (Grice 1957: 381–383), należy zwrócić uwagę, że jego teoria miała wyjaśnić, jak możliwe jest użycie słów wyrażające inną treść niż ta, którą obejmuje ich konwencjonalne znaczenie (Grice 1975/2000: 272–273). Zdaniem Grice’a, kiedy orientujemy się, iż nasz interlokutor używa słów w sposób niedający pogodzić się z ich konwencjonalnym znaczeniem, musimy — kierując się zasadą życzliwości i maksymami konwersacyjnymi — starać się odczytać jego intencje. Jeśli na przykład ktoś powie: „Artykuł Quine’a był granatem wrzuconym w światek analitycznej filozofii języka.”, to będę się starał życzliwie odczytać jego wypowiedź, mimo że słowo „granat” zostało tu użyte, na pierwszy rzut oka, niepoprawnie. Uznam jednak, że mój interlokutor chce ze mną współpracować; dojdę zapewne do wniosku, iż użył tego słowa jako metafory czegoś, co powoduje wstrząs — w tym wypadku poznawczy.

Davidson twierdzi, iż w obliczu powszechności opisywanych przez niego fenomenów teoria języka opierająca się na założeniu o istnieniu konwencjonalnych znaczeń i reguł językowych jest po prostu zbędna. Jeśli dysponujemy dobrą teorią opisującą proces interpretacji naszych rozmówców — w domyśle: jego teorią — to nie potrzebujemy dodatkowych założeń o istnieniu reguł semantycznych. Teorię przyjmującą istnienie konwencjonalnych znaczeń można odrzucić kierując się zasadą brzytwy Ockhama.

Davidson odrzuca także ideę autonomicznych zobowiązań użytkownika do posługiwania się językiem poprawnie; jego zdaniem, absurdem jest postulowanie istnienia „zobowiązań wobec języka” (Davidson 1994/2005: 118). Mamy zobowiązania jedynie wobec innych podmiotów — w kwestii języka sprowadzają się one do tego, iż powinniśmy tak się nim posługiwać, aby ułatwić komunikację.

Aby odpowiedzieć Davidsonowi, trzeba zauważyć, iż teoria niepostulująca istnienia konwencjonalnych znaczeń pomija bardzo ważny aspekt rozumienia malapropizmów i wyrażeń nowatorskich literacko. W procesie ich interpretacji nie tylko dochodzimy intencji mówiącego — zauważamy także niepoprawność wypowiedzi. Większość osób rozumiejących słuchowiska radiowe oparte na malapropizmach od razu rozpoznaje je jako niepoprawne — na tym w dużej mierze polega ich wartość rozrywkowa. Chociaż Davidson twierdzi, iż fakt, że malapropizmy są zabawne, nie ma istotnego znaczenia dla ich zrozumienia (Davidson 1986/2005: 90), odrzucenie tezy o istnieniu reguł semantycznych w istotny sposób zubaża opis tego, jak wypowiedzi tego typu są przez nas odbierane.

Kwestia interpretacji tekstu w terminach poprawności bądź niepoprawności ma jeszcze większe znaczenie w przypadku dzieł bazujących na literackiej inwencji językowej. Davidson w swoim szkicu o Joysie zwraca uwagę na jego chęć „wyzwolenia się z więzów języka” (Davidson 1989/2005: 146–147) i na to, że udało mu się przekazać mnogość treści w sposób łamiący konwencje lingwistyczne. Nasuwa się refleksja: czy rzeczywiście na gruncie teorii Davidsona można w ten sposób

Joyce'a odczytywać? Skoro bowiem nie istnieją reguły językowe, to czy można dorzecznie twierdzić, że Joyce je twórczo łamał?

Davidson mógłby odpowiedzieć, iż wartość dzieł Joyce'a polega na tym, że używał on znanych i nieznanymi słów w zaskakujący sposób. Jednak użyć słowa w sposób zaskakujący to nie to samo, co użyć go niepoprawnie. W tym sensie Joyce staje się nowatorem, a nie kimś, kto łamie konwencje. Trudno wchodzić tu w polemikę na temat właściwego sposobu interpretacji dzieł literackich w stylu Joyce'a; warto jednak odnotować, iż samo przywołanie tego fenomenu nie rozstrzyga jeszcze sporu. Jeśli zgodzimy się bowiem z cytowaną powyżej tezą Whitinga, iż normy semantyczne są normami *prima facie* i, co za tym idzie, gotowi będziemy przyznać, iż inne normy mogą unieważniać reguły poprawnego użycia słów, to możemy również przyjąć, iż względy estetyczne i literackie mogą stanowić dobrą rację dla odrzucenia standardów semantycznych.

Tego typu argumenty nie obalają oczywiście teorii Davidsona; pokazują jedynie, iż teoria reguł semantycznych jest w stanie uwzględnić wskazane przez niego fenomeny. Zwolennik koncepcji Davidsona może twierdzić, iż jest ona oszczędniejsza — wyjaśnia wszystko, co jest do wyjaśnienia, bazując na mniejszej liczbie przesłanek. Jakkolwiek argument z oszczędności założeń jest niewątpliwie poważny, wydaje się, iż przewaga podejścia autora *Social aspect of language* wynika głównie z tego, jakie fakty ma ono w założeniu wyjaśniać. Podejście to koncentruje się na wyjaśnieniu fenomenu komunikacji, czyli tego, że osoby posługujące się nieraz bardzo odmiennymi językami są w stanie uzgodnić swoje interpretacje i w rezultacie koordynować działania. Teoria Davidsona opiera się — co przyznaje on *explicite* — na założeniu, iż w punkcie wyjścia mamy do czynienia z podmiotami, które potrafią posługiwać się jakimś językiem (Davidson 1986/2005: 100). Jeśli uznamy to założenie metodologiczne za słuszne, wówczas skłonni będziemy zgodzić się z Davidsonem, iż inklinacja użytkowników języka do dzielenia użyc na poprawne i niepoprawne jest mało interesująca. Tym, co ważne dla teoretyka komunikacji,

są bowiem założenia, jakie należy poczynić, aby teoria mogła funkcjonować, a nie potoczne wyobrażenia na temat języka.

Można zatem przyznać Davidsonowi rację, iż do wyjaśnienia tak rozumianej komunikacji założenie istnienia reguł semantycznych nie jest konieczne. Nie wydaje się to, nawiasem mówiąc, wnioskiem szczególnie zaskakującym w świetle rozważań zawartych w poprzedniej części pracy. Wynika z nich jasno, iż nie można utożsamiać kierowania się regułami semantycznymi z realizacją zaleceń mających prowadzić do skutecznej komunikacji. Normy językowe mają charakter przynajmniej do pewnego stopnia autonomiczny wobec celu, jakim jest komunikacja.

Problem, na wyjaśnieniu którego zależało Davidsonowi — jak dwa podmioty dysponujące językiem dochodzą do porozumienia — nie jest zatem tożsamy z tym, co rozważał Kripke, gdy wprowadzał pojęcie normatywności znaczenia. Kripke koncentruje się na sytuacji, w której użytkownik języka podejmuje decyzję co do użycia słów, na przykład zastanawia się, jak zastosować słowo „+” do nowego przypadku (Kripke 1982/2007: 22–23). Aby przypisywanie znaczenia słowu — tak jak jest ono używane przez podmiot — miało sens, musimy założyć istnienie kryteriów poprawności użycia, innych niż sama skłonność podmiotu do używania tego słowa w określony sposób. Innymi słowy, aby pojęcie znaczenia miało sens, niezbędne jest odróżnienie tego, co prawidłowe, od tego, co się prawidłowym jedynie wydaje (Kripke 1982/2007: 46).

Davidson, odnosząc się do podobnych zarzutów stawianych mu przez Dummetta, stara się odpowiedzieć na to, co nazywa „pytaniem Wittgensteina” (Davidson 1994/2005: 123–125): co odróżnia sytuację, w której ktoś kieruje się regułą, od tej, w której komuś jedynie wydaje się, że kieruje się regułą? Odpowiedź, której udziela Davidson, trudno jednak uznać za satysfakcjonującą. Twierdzi on, iż, aby możliwe było posługiwanie się językiem, nieodzowne jest uchwycenie różnicy pomiędzy prawdziwością a fałszywością własnych wypowiedzi. Inaczej mówiąc, jeśli zaczynam stosować słowo „nos”, muszę rozumieć różnicę między prawdzi-

wymi a fałszywymi wypowiedziami typu „To jest mój nos”. Różnicę tę jednostka może pojąć jedynie, gdy wchodzi w interakcje z innymi; dla osoby pozostającej w izolacji różnica ta nie istnieje. Może się więc wydawać, iż Davidson zachowuje pewne rozróżnienie na poprawne i niepoprawne użycia języka — występuje jedynie przeciwko tezie, że reguły semantyczne stanowią element wspólny dla wszystkich jego użytkowników (Davidson 1994/2005: 125).

Po uwzględnieniu tych rozważań, wizja Davidsona przedstawia się następująco: istnieje mnogość idiolektów, których esencjalną cechą jest wzajemna przekładalność — język nieprzekładalny na inny nie jest bowiem językiem. Użytkownicy idiolektów porozumiewają się ze sobą poprzez tworzenie teorii interpretacji dla języków swoich interlokutorów. Na tym poziomie analizy nie powstaje kwestia normatywności znaczenia — w trakcie procesu interpretacji muszą traktować wypowiedzi mojego interlokutora jako „dane”. Nie mam żadnych podstaw, aby uznać je za poprawne lub niepoprawne. Pytanie o normatywność znaczenia pojawia się natomiast w odniesieniu do idiolektu, na poziomie którego analizuje się znaczenie; skoro nie istnieje język wspólny, to przypisywanie znaczenia wyrażeniom musi być relatywizowane do idiolektów. W tym sensie Davidson akceptuje tezę o normatywności znaczenia: twierdzi, iż istnieją poprawne oraz niepoprawne sposoby użycia, a użytkownik języka może się mylić, gdy stosuje dane pojęcie. Wydaje się zatem, iż ostatecznie Davidson podpisałby się pod „potoczną” koncepcją normatywności znaczenia, omówioną w pierwszym punkcie niniejszego rozdziału.

Rozdział III

Normatywność a przyczynowość

III.1 Słaba normatywność i jej miejsce w argumentacji sceptycznej

W poprzednim rozdziale wykazano, że argumenty przeciwko tezie o normatywności znaczenia są nieskuteczne. Konkluzja ta okupiona jest jednak poważnym kosztem teoretycznym — wyjściową tezę należało bardzo złagodzić. Uznanie, iż w przypadku semantyki mamy do czynienia z normatywnością w słabym sensie, pociąga za sobą przede wszystkim odrzucenie kilku omawianych wcześniej przekonań. Należy przede wszystkim odrzucić precyzację dokonaną przez Boghossiana, co wiąże się z koniecznością rozróżnienia poprawnego użycia i poprawnego zastosowania wyrażen. Błędna okazuje się również terminologia „kierowania się regułami” — istotny jest bowiem fakt, iż pewne normy poprawności pozostają w mocy, a nie to, że się do nich stosujemy. Przyjęcie tezy o normatywności znaczenia nie prowadzi też, samo w sobie, do uznania istnienia konwencjonalnych reguł. Musimy

jednak przyjąć, że istnieją poprawne i niepoprawne sposoby użycia wyrażen — to rozróżnienie można zastosować również na poziomie idiolektu.

Najważniejszym ograniczeniem tezy o normatywności znaczenia, jest odrzucenie przekonania o wewnętrznie motywującym charakterze tego ostatniego. Zmuszeni jesteśmy przyjąć, że normatywność, z jaką mamy do czynienia w przypadku zdań z dyskursu semantycznego, jest inną normatywnością, niż ta, z którą mamy do czynienia w przypadku zdań etycznych. Ten wniosek uniemożliwia z kolei uznanie poprawności zarysowanego w *Rozdziale I* argumentu:

ZDANIA O ZNACZENIU WYRAŻAJĄ NORMY.

ZDANIA WYRAŻAJĄCE NORMY NIE OPISUJĄ FAKTÓW.

ZDANIA O ZNACZENIU NIE OPISUJĄ FAKTÓW.

Argument okazuje się być obciążony błędem ekwiwokacji: słowo „norma” użyte w pierwszej przesłance ma inny, słabszy sens niż w przesłance drugiej. Okazuje się, że pojęcie normatywności nie pozwala na stworzenie argumentu przeciwko faktualizmowi, który łączyłby w sobie prostotę, ogólność i aprioryczny charakter.

W tej sytuacji można by dojść do wniosku, że uznanie zdań o znaczeniu za słabo normatywne nie ma żadnych interesujących konsekwencji ontologicznych. Takiego zdania jest Hattiangadi, która wprowadza zbliżony do pojęcia słabej normatywności termin zależności od normy (*norm-relativity*). Wedle tej autorki, można o zależności od normy mówić wtedy, gdy istnieje pewien standard, z którym coś — np. zachowanie danego podmiotu — może być zgodne lub nie. W przypadku zdań o znaczeniu można mówić jedynie o normatywności rozumianej jako zależność od normy, co nie stanowi kontrargumentu wobec twierdzenia, iż dyskurs ten ma charakter opisowy.

Hattiangadi ilustruje tę tezę zgrabnym przykładem: w parku rozrywki znajduje się karuzela, na którą wstęp mają tylko dzieci wyższe niż, powiedzmy, 140 cm.

Ta norma może obowiązywać w parku, w tym sensie, iż cieszy się akceptacją i jej przestrzeganie jest egzekwowane za pomocą sankcji (...). Dziecko może chcieć spełnić ów standard, jednak to, czy go spełnia, czy nie, jest kwestią całkowicie nie-normatywnego faktu — tego, iż ma ono tyle a tyle centymetrów wzrostu. Możemy powiedzieć, iż jego wzrost jest „właściwy”, ale nie oznacza to, iż to jest ten wzrost, który powinno ono osiągnąć — jeśli pominiemy jego chęć, aby przejechać się na karuzeli (Hattiangadi 2006: 63).

Podobnie ma się rzecz ze zdaniami o znaczeniu. Można mówić, iż jakaś wypowiedź była poprawna bądź nie, ale oznacza to jedynie, iż jest ona — bądź też nie jest — zgodna z pewnym standardem. Zaś zgodność ze standardem jest, w myśl tej koncepcji, w sposób nieproblematyczny kwestią faktu.

Jednak mamy tutaj bowiem do czynienia nie tyle z rozwiązaniem problemu normatywności znaczenia, co z jego przeformulowaniem. Można zgodzić się z tezą, że to właśnie możliwość mówienia o zgodności z pewnym standardem — standardem poprawnego użycia — jest tym, co decyduje o zasadności stosowania słownictwa normatywnego do wypowiedzi o znaczeniu. Czy jednak zdania, w których stwierdzamy zgodność ze standardem, są zdaniami, w których stwierdzamy fakty?

W podanym przez Hattiangadi przykładzie niewątpliwym stwierdzeniem faktu jest powiedzenie, iż dziecko chcące wejść na karuzelę ma 140 cm wzrostu. Jednak stwierdzenie, iż taki właśnie wzrost jest zgodny z normą, może już budzić pewne wątpliwości — można bowiem zapytać, co konstytuuje odpowiadający mu fakt? Kwestia dotyczy faktualnego statusu zdań, w których stwierdzamy poprawność rozumianą jako zgodność z normą. Czyż paradoksu Kripkego nie da się zastosować do podanego przez Hattiangadi przykładu? Sceptyk mógłby stwierdzić, iż zgodnie z normą na karuzelę mogą wsiadać dzieci mające powyżej 140 cm wzrostu, przy czym wyrażenie mieć „140 cm wzrostu” oznacza mieć „140 cm wzrostu” do 01. 01. 2010 r., po tej dacie zaś znaczy „mieć 140 funtów wagi”.

Stwierdzenia wyrażające normy w słabym sensie są powszechne w języku i określają bardzo wiele fenomenów: różnego rodzaju wzorce kulturowe, przepisy grzeczności, normy prawne, reguły określające zasady gry i temu podobne. Wszystkie te zjawiska można podciągnąć pod ogólną kategorię „kierowania się regułami” — mamy w ich przypadku do czynienia z podmiotem, który, świadomie bądź nie, postępuje w zgodzie lub wbrew pewnym *implicite* bądź *explicite* przyjmowanym normom. Te normy jednak, co istotne, nie mają charakteru moralnego. Warto przy okazji zauważyć, że sam Wittgenstein omawiając problem kierowania się regułami ani razu nie podał przykładu z dziedziny etyki. Omawiane konteksty można, kierując się zarówno socjologicznym uzusem, jak i sugestią samego Wittgensteina, nazwać instytucjonalnymi. W 199. paragrafie *Dociekań* Wittgenstein pisze:

Kierowanie się regułą, komunikowanie czegoś, wydanie rozkazu, rozegranie partii szachów są to pewne zwyczaje (nawyki, instytucje) (Wittgenstein 1951/2000 §199).

Widać wyraźnie, że trop metaetyczny — poszukiwanie uzasadnienia dla non-faktualizmu w rozważaniach dotyczących statusu metafizycznego zdań etycznych — okazał się błędny. Normatywność instytucjonalna — do której zalicza się normatywność semantyczna — jest kategorią innego typu niż normatywność etyczna (o tym rozróżnieniu w kontekście metaetycznym pisze Smith 1994: 80). Nie znaczy to jednak, że ta pierwsza nie stanowi żadnego problemu filozoficznego.

Normatywność znaczenia jest szczególnym przypadkiem normatywności instytucjonalnej. W tym konkretnym przypadku wykluczone jest bowiem pewne narzucające się wyjaśnienie warunków prawdziwości wyrażen normatywnych, a mianowicie odwołanie się do słownego ujęcia normy — tego, iż ona jest napisana lub wypowiedziana. Obowiązywanie pewnych zasad w boksie można wyjaśnić przez odwołanie do reguł spisanych przez markiza Queensbury. Jednak w przypadku języka posunięcie tego typu jest nieuprawnione — pojawia się bowiem problem regresu, jako że te spisane bądź wypowiedziane reguły będą również

obdarzone słabą normatywnością. Z podobnym problemem zetkniemy się, jeśli będziemy chcieli wyjaśnić istnienie reguł językowych przez odwołanie się do intencji użytkowników. Fakty instytucjonalne nie są zatem, wbrew temu, co sugerują Hattiangadi, Glüer i Pagin, czymś, co możemy bez problemów włączyć do naszej ontologii. Z drugiej jednak strony stajemy znów przed pytaniem o *onus probandi* — nie wiemy co prawda, jakie fakty odpowiadają stwierdzeniom wyrażającym normatywność instytucjonalną, ale nie mamy żadnego argumentu, który pokazywałby, że fakty takie nie mogą istnieć.

W kolejnych paragrafach postaram się przedstawić argument, który wypełni tę lukę i wskaże na pewien istotny powód, dla którego nie można uznać stwierdzeń słabo normatywnych za opisujące fakty. Rozumowanie to będzie próbą przeniesienia tzw. argumentu z przyczynowego wykluczenia z terenu filozofii umysłu na problem znaczenia (rzecz jasna, z pewnymi koniecznymi modyfikacjami). W tym celu najpierw przeprowadzę rozumowanie pokazujące, w jaki sposób teza o słabej normatywności obala teorię dyspozycjonalną — pozwoli to uzasadnić konieczność odróżnienia dyspozycji od znaczeń. Następnie zostanie przedstawiony oryginalny argument z przyczynowego wykluczenia. Ostatnim krokiem rozumowania będzie zastosowanie go do problemu znaczeń i dyspozycji.

III.2 Słaba normatywność a teorie dyspozycjonalne

Jakkolwiek celem niniejszej rozprawy nie jest stworzenie egzegezy tekstu Kripkego, to warto zauważyć, iż nigdzie nie przedstawia on ogólnego argumentu przeciw istnieniu faktów semantycznych, który miałby taką postać, jak ten przedstawiony w części pierwszej tego rozdziału — czyli dowodzący nonfaktualizmu semantycznego na podstawie ogólnej przesłanki, mówiącej, że zdania normatywne nie oznaczają faktów. Odwołanie się do normatywności znaczenia pojawia się

u niego w kontekście polemiki z dyspozycjonalną teorią znaczenia. W skrócie ów argument został przedstawiony w *Rozdziale I*.

Zdefiniujmy teorię dyspozycjonalną jako taką, wedle której przypisywanie danemu symbolowi znaczenia jest tym samym, co posiadanie przez podmiot własności, która sprawia, że używa on tego symbolu w określony sposób. Własność ta będzie określana mianem D-własności. Takie wyrażenie teorii dyspozycjonalnej ma tę zaletę, iż nie przesądza o charakterze owej D-własności. Jedyną rzeczą, której wymaga się od owej własności, jest to, aby w jakiś sposób determinowała ona działanie podmiotu. Naturalne jest więc przypuszczenie, iż jako D-własności będą się kwalifikować przede wszystkim cechy należące do szeroko rozumianego katalogu behawiorystycznego (to znaczy opisującego działania jednostek i dyspozycje do owych działań) albo do katalogu neurofizjologicznego (opisującego wewnętrzne, biologiczne mechanizmy działania człowieka). To, jakie jednak konkretnie cechy miałyby odgrywać rolę D-własności zostanie pozostawione bez odpowiedzi. Jakkolwiek badania nad neurofizjologicznymi korelatami ludzkiej zdolności do werbalnego porozumiewania się zdają się rozwijać dość dynamicznie, to wydaje się, iż nie ma jeszcze podstaw do twierdzenia, że potrafimy dobrze opisać konkretną cechę decydującą o tym, jak używamy danego wyrażenia. Samo słowo „własność” musi zatem być w tym kontekście rozumiane bardzo liberalnie — jako coś, co możemy przypisać pewnemu podmiotowi. Jest oczywiście mało prawdopodobne (choć niewykluczone), że będzie można wyodrębnić jedną własność — identyfikowaną wedle kryteriów „naukowych” — która odpowiada za użycie danego słowa.

Odwołanie się do normatywności znaczenia nie dowodzi nieistnienia D-własności. Naturalne wydaje się dość spostrzeżenie, że większość ludzi używa symboli we względnie systematyczny sposób. Postulowanie pewnych własności dyspozycjonalnych (w szerokim sensie tego słowa) jest jednym z możliwych wyjaśnień tego zjawiska. Przedmiotem krytyki jest jednak utożsamienie posiadania D-własności ze znaczeniem.

Aby tę krytykę uczynić jaśniejszą, wprowadźmy pojęcie M-własności. M-własność jest to hipotetyczna cecha, która przysługuje wszystkim i tylko tym osobom, które posługują się danym pojęciem w określonym znaczeniu. Oczywiście, M-własności są własnościami jedynie w trywialnym sensie tego słowa — jako coś, co przysługuje podmiotom. Jeśli zatem można o mnie powiedzieć, zgodnie z prawdą, że posługuję się słowem „ogórek” w tym znaczeniu, jakie ma ono w standardowej polszczyźnie, to przysługuje mi pewna M-własność (czyli własność posługiwania się słowem „ogórek”, w tym znaczeniu, jakie ma ono w standardowej polszczyźnie).

Prosta teoria dyspozycjonalna jest zatem koncepcją, która utożsamia posiadanie M-własności z D-własnością. Tego typu tezę łatwo obalić poprzez odwołanie się do faktu normatywności znaczenia. Istnienie własności, która sprawia, że używam danego wyrażenia w pewien określony sposób, nie jest jeszcze żadną przesłanką na rzecz stwierdzenia, iż używam tego słowa poprawnie bądź nie. Przy czym poprawność, o którą tu chodzi, nie musi wcale być interpretowana w kategoriach mocnej normatywności — normy „instytucjonalne” również nie mogą być utożsamiane z dyspozycjami. Rozważmy typowy przykład normy w sensie słabym — przepisy kodeksu drogowego. W tym przypadku również możemy z dość dużą dozą pewności przyjąć, iż każdy kierowca posiada pewną własność dyspozycyjną, która określa jego zachowanie za kierownicą. Z tego, że ktoś posiada ową dyspozycję, nie wynika jednak to, że jeździ on zgodnie z przepisami.

Na poziomie ogólnym argument ów brzmiałby następująco: za każdym razem, gdy mamy do czynienia z kontekstem, w którym możemy mówić o poprawnych bądź niepoprawnych zachowaniach podmiotów — przy czym tę poprawność należy rozumieć w słabym sensie — należałoby przypuszczać, iż podmiotom, których zachowania oceniamy, możemy przypisać jakieś własności dyspozycyjne, które determinują ich zachowania. Jednak niektóre z tych dyspozycji będą prowadziły do zachowań, które należy uznać za niepoprawne z punktu widzenia obowiązują-

cych standardów. Gdyby w danym kontekście wszyscy z konieczności zachowywali się zawsze „poprawnie”, to nie byłoby w ogóle sensu mówić o tym, iż obowiązują w tym kontekście jakieś normy. Taki też jest sens słynnego *dictum* Wittgensteina z paragrafu 258. *Dociekań*, w którym stwierdza on, iż jeśli prawidłowe jest to, co się komuś prawidłowe wyda, to „o prawidłowości nie ma co mówić”.

Interesującą strategię obrony dyspozycjonalizmu rozwijają Martin i Heil. Zdaniem tych autorów (Martin i Heil 1998: 288 - 293), błędem jest „empirystyczne” traktowanie dyspozycji, czyli analizowanie ich w kategoriach okresów warunkowych. Zwolennik empirystycznego rozumienia dyspozycji będzie twierdził, że każde bez wyjątku zachowanie, jakie przejawia dany podmiot, gdy np. stosuje dany symbol, należy uznać właśnie za wynik działania dyspozycji tego podmiotu. Tymczasem, wedle omawianych filozofów, dyspozycje należy traktować realistycznie, to znaczy jako własności, które przysługują podmiotom niezależnie od tego, czy mają się okazać ujawniać czy też nie. To pozwala wyjaśnić możliwość popełniania pomyłki: osoba może posiadać odpowiednią dyspozycję, ale jej nie ujawnić — dyspozycja ta może zostać „zablokowana”. Na przykład, nie każda operacja, którą wykonam, używając symbolu „+” będzie mogła zostać uznana za realizację mojej dyspozycji; jeśli, powiedzmy, pod wpływem LSD zacznę używać tego symbolu w niestandardowy sposób, to moje działanie nie będzie efektem mojej dyspozycji — narkotyk zakłóci bowiem standardowe procesy umysłowe, które zazwyczaj determinują moje zachowanie.

Można zgodzić się z Martinem i Heilem, iż empirystyczne pojęcie dyspozycji jest błędne — nasze dyspozycje mogą pozostać nieujawnione bądź zakłócone, co nie zmienia faktu, iż je posiadamy. Jednak taka konstatacja nie prowadzi jeszcze do rozwiązania problemu normatywności znaczenia. Czyż nie jest bowiem możliwe, iż postąpię w sposób niewłaściwy nawet wówczas, gdy moja dyspozycja nie zostanie zablokowana? Kripke wskazywał (1982/2007: 57 - 58) na fakt, że możemy popełniać systematyczne pomyłki, czyli mieć dyspozycje do mylenia się. Podobnie, wielu

jest kierowców, w przypadku których to jazda zgodnie z przepisami wydaje się przykładem zaburzenia realizacji ich naturalnej dyspozycji. Martin i Heil muszą zatem założyć, że *zawsze*, gdy posługuje się symbolem poprawnie, jest to wynik działania dyspozycji, zaś każda pomyłka — efektem jej zablokowania.

Taka wersja teorii dyspozycjonalnej wydaje się opierać w dużej mierze na spekulacji. Nie wiemy bowiem, czy rzeczywiście da się wyodrębnić taką własność dyspozycjonalną, której realizacja odpowiada za wszystkie działania, które uznaliśmy za poprawne użycia danego symbolu i za żadne inne. Z drugiej jednak strony nie możemy tego wykluczyć — kwestia ta wydaje się mieć charakter empiryczny. Tylko drogą badań naukowych można uzyskać odpowiedź na pytanie, co powoduje ludzkie zachowania językowe w konkretnych przypadkach. Nie można *a priori* wykluczyć ewentualności, iż znajdzie się jakiś czynnik, który odpowiada za wszystkie poprawne zachowania językowe i tylko za nie. Tymczasem, aby uzasadnić nonfaktualizm, potrzebujemy apriorycznego argumentu przeciwko utożsamieniu poprawności z dyspozycjami.

Martin i Heil oraz, jak się zdaje, Hattiangadi próbują dokonać *naturalizacji* normatywności instytucjonalnej. Ich zdaniem fakty dotyczące tego, iż ktoś zachowuje się w jakimś kontekście poprawnie bądź nie, mogą zostać wyjaśnione w całkowicie naturalnych, czyli *eo ipso* faktualnych kategoriach. Poprawność, w myśl tej wizji, może zostać utożsamiona z posiadaniem pewnej naturalnej własności. Bardzo dobrze widać to na podanym przez Hattiangadi przykładzie dziecka „nie spełniającego normy” umożliwiającej przejażdżkę na karuzeli. W końcu określony wzrost jest całkowicie naturalną i „faktualną” własnością. Można domniemywać, iż w myśl tej koncepcji, tylko nasze (tymczasowe) ograniczenia poznawcze uniemożliwiają nam wskazanie takiej naturalnej własności w przypadku znaczenia.

Ta linia rozumowania wydaje się z zasadniczych powodów błędna. Rozważmy przypadek, w którym jednoznacznie można wskazać naturalną własność odpowiadającą niepoprawnemu zachowaniu: możemy powiedzieć, iż pojazd przekracza

dopuszczalną prędkość na autostradzie, jeśli pojazd ten porusza się z prędkością większą niż 140 km/h. Można by więc dojść do wniosku, iż mamy tutaj do czynienia z identyfikacją faktu normatywnego instytucjonalnie — łamaniem pewnego przepisu — z zupełnie naturalnym faktem poruszania się przez pewien obiekt z określoną prędkością.

Tożsamość tych dwóch faktów jest jednak pozorna — własność „przekraczania prędkości na autostradzie” jest inną własnością niż własność „poruszania się z prędkością większą niż 140 km/h”. Owszem te własności są koekstensjonalne — to znaczy, każdy przedmiot posiadający jedną własność, posiada również i drugą (jeśli ograniczymy zbiór przedmiotów do samochodów), jednak nie znaczy to, że są one tą samą własnością. Są one bowiem koekstensjonalne przygodnie — to znaczy, nie jest z konieczności tak, że przedmiot posiadający jedną z tych własności posiada automatycznie drugą. Bardzo łatwo można wyobrazić sobie hipotetyczną sytuację, w której w Polsce zniesiono ograniczenia prędkości w poruszaniu się na autostradach i samochód jadący z prędkością 150 km/h nie łamie żadnych przepisów. Podobnie, w podanym przez Hattiangadi przykładzie, posiadanie 135 cm wzrostu nie oznacza z konieczności niespełniania standardu — z łatwością można pomyśleć inny „możliwy świat”, w którym norma została zmieniona tak, by i nieco niższe dzieci mogły cieszyć się przejażdżką.

Ogólnie rzecz biorąc, za każdym razem, gdy mamy do czynienia z kontekstem „instytucjonalnym”, dokonujemy podziału zachowań czy własności ludzi na poprawne bądź nie. Wyjąwszy przypadki nieprzestrzegania regulaminu towarzystwa telepatów, klasyfikacji tej dokonujemy biorąc pod uwagę jakieś naturalne własności lub widoczne zachowania interesujących nas jednostek — innymi słowy, za poprawne uznajemy to, co cechuje się pewną naturalną własnością. Nie znaczy to jednak, iż „bycie poprawnym” można utożsamić z posiadaniem tej naturalnej własności. Fakty instytucjonalne zależą bowiem w istotny sposób od kontekstu — ta sama własność, która w jednym kontekście będzie znacznikiem poprawności

zachowania w innym może być czymś, co sprawi, że dane działanie zakwalifikowane będzie niepoprawne.

Aby uniknąć zarzutu, iż przyjmuje się tutaj jakieś mocne założenia dotyczące nie-indywidualnego charakteru języka, warto dodać, iż kontekst, o którym mowa nie musi mieć charakteru społecznego. Jeśli przyjmujemy wizję, w której o znaczeniu decydują czyjeś indywidualne intencje znaczeniowe, to możemy stwierdzić, iż to właśnie owe intencje decydują o ocenie poprawności mojego zachowania językowego. Mogę mieć dyspozycję do nazywania „ogórkami” wszystkich podłużnych zielonych warzyw, jednak ta dyspozycja może, w różnych kontekstach, bądź prowadzić do zachowań poprawnych, bądź sprawiać, że będę popełniał błędy — jeśli poprzednio świadomie zdecydowałem, iż nie będę nazywał „ogórkami” cukinii. Nawet zwolennik indywidualistycznej wizji języka musi przyznać, iż możliwa jest sytuacja, w której nabędę dyspozycji do działania niezgodnego z moją pierwotną intencją znaczeniową.

Dyspozycjonalizm okazuje się zatem błędną teorią znaczenia. Musimy odróżnić normy i własności semantyczne, które służą do określania, czy moje zachowanie językowe jest zgodne z pewnymi standardami, od własności dyspozycyjnych, które przyczynowo determinują moje zachowanie. W jaki jednak sposób ta konstatacja prowadzi nas do nonfaktualizmu semantycznego? Czyż nie możemy powiedzieć, że podmiot posługujący się językiem da się prawdziwie opisać zarówno w terminach norm, jak i dyspozycji?

III.3 Argument z przyczynowego wykluczenia

Aby wykazać, iż nie można przyjmować obu tych opisów jednocześnie, warto odwołać się do argumentu z przyczynowego wykluczenia sformułowanego przez Jaegwona Kima. Został on pierwotnie wprowadzony jako próba refutacji nieredukcyjnie fizykalnych ujęć zagadnienia relacji umysł - ciało.

Sednem nieredukcyjnego fizykalizmu (por. np. Kim 1989/2008: 77 - 78) jest połączenie ze sobą dwóch tez: przekonania, iż cała rzeczywistość ma fizykalny (materialny) charakter, z przeświadczeniem, iż nie da się opisać rzeczywistości umysłowej w języku nauk fizykalnych. Innymi słowy, nieredukcyjni fizykaliści (tacy jak Davidson 1970/1992 czy Fodor 1974/2008) sądzą, iż można mający zasadniczo fizyczną naturę świat opisywać również za pomocą słownictwa mentalistycznego i będzie to opis równie zasadny, jak ten operujący wyłącznie kategoriami fizykalnymi. Przenosząc tę tezę na poziom ontologiczny można powiedzieć, iż zwolennicy nieredukcyjnego fizykalizmu uznają, iż wszystkie byty mają charakter fizyczny, aczkolwiek niektóre z owych fizycznych bytów mogą posiadać mentalne, czyli niefizyczne, własności.

W potocznym obrazie relacji między umysłem a ciałem przyjmujemy, iż zdarzenia mentalne wpływają przyczynowo na zdarzenia fizyczne — moja chęć zjedzenia kolacji powoduje otwierający lodówkę ruch mojej ręki. Wydawać by się mogło — taka myśl przyświecała koncepcji Davidsona — iż, nieredukcyjny fizykalizm pozwala na utrzymanie tego przekonania. Jeśli utrzymujemy w zgodzie z tą koncepcją, iż wszystkie byty — w tym wydarzenia - mają charakter fizyczny, to również wydarzenie, które opisujemy jako „chęć zjedzenia przeze mnie kolacji” będzie pewnym wydarzeniem fizycznym (prawdopodobnie stanem mojego mózgu). Jako takie będzie ono mogło być uznane za przyczynę mojego zachowania, które jest niewątpliwie zjawiskiem fizycznym. Z drugiej strony zdarzenie to może być opisywalne w języku mentalistycznym — można więc orzec, iż zdarzenia mentalne powodują zdarzenia fizyczne.

Zdaniem Kima (1989/2008: 81), taka koncepcja nie daje się utrzymać, gdyż przyjęcie nieredukcyjnego fizykalizmu zmusza do uznania, że „umysłowość nie odgrywa faktycznie żadnej roli przyczynowej” (Kim 1989/2008: 81). Wniosek ten opiera się na dwóch przesłankach, wyodrębnionych przez Jespera Kallestrupa w jego studium argumentu Kima. Pierwszą przesłanką jest teza przyczyno-

wego domknięcia, głosząca, iż każdy skutek fizyczny ma dostateczną przyczynę fizyczną (Kallestrup 2006: 463). Założenie to jest akceptowane przez każdą koncepcję nieredukcyjnie fizykalistyczną, bowiem w tych teoriach wszystkie zdarzenia mają charakter fizyczny.

Kluczowe znaczenie ma druga przesłanka — teza przyczynowego wykluczenia — która głosi, iż „jeśli własność E ma wystarczającą przyczynę C , to żadna inna własność C^* , różna od C , nie może być przyczyną E ” (Kallestrup 2006: 464). Jeśli chcemy zapytać, dlaczego jakieś zdarzenie zaszło i do wyjaśnienia tego faktu wystarczy fakt, że zaszło jakieś inne zdarzenie i miało ono cechę C , to nie możemy już twierdzić, iż jakieś inne zdarzenie albo inna cecha tego samego zdarzenia spowodowały interesujący nas fakt. Jeśli mówimy, że przyczyną pożaru był fakt, iż w pomieszczeniu znajdowała się zapalona zapalka, to nie możemy powiedzieć, iż tą przyczyną było również to, że znajdowała się tam zapalka wyprodukowana w Chinach. Innymi słowy, jeśli do przyczynowego wyjaśnienia jakiegoś zjawiska wystarcza odwołanie się do pewnych własności danego przedmiotu, czy też zdarzenia, to fakt, że ten przedmiot czy zdarzenie ma również inne własności, nie jest relewantny.

W omawianym przypadku przyczynowania mentalnego, fizyczne cechy mojego ruchu ręką są w zupełności zdeterminowane — co wynika z założenia przyczynowego domknięcia — przez fizyczne cechy wydarzeń prowadzących do niego. Wśród tych zdarzeń możemy wyodrębnić pewien stan mojego mózgu, który, wedle zwolennika nieredukcyjnego fizykalizmu, może być również opisywany za pomocą terminologii mentalnej i, *eo ipso*, posiada mentalne cechy. Jednak tylko cechy fizyczne owego stanu są przyczynowo relewantne w wyjaśnianiu mojego ruchu ręką. Fakt, że to zdarzenie jest identyczne z jakimś zdarzeniem mentalnym, nie ma żadnego znaczenia dla jego fizycznych skutków. Opisując relację przyczynową zachodzącą między moim stanem neuronalnym a moim działaniem, nie musimy w ogóle brać pod uwagę istnienia jakichkolwiek stanów, które mogą być opisywane jako mentalne.

Otrzymujemy zatem tezę epifenomenalizmu typicznego, głoszącą, iż własności mentalne nie odgrywają żadnej roli przyczynowej. Kim nie poprzestaje jednak na tym i dochodzi do tezy irrealizmu (Kim 1989/2008: 81), mówiącej, iż żadnych zdarzeń i własności mentalnych po prostu nie ma. Podobne przekonanie zostało wyrażone, być może po raz pierwszy, przez brytyjskiego emergentystę Samuela Alexandra. Według tego ostatniego, „epifenomenalizm uznaje istnienie w naturze czegoś, co nie ma nic do zrobienia, nie służy żadnemu celowi, swoistego szlachetnego gatunku, który opiera się na pracy swych poddanych, lecz sam jest trzymany na pokaz, a równie dobrze mógłby być i bez wątpienia z czasem zostanie wyeliminowany” (Alexander 1920: 8 [za:] Kallestrup 2006: 468).

Kim przenosi rozumowanie Alexandra na poziom własności — twierdzi, iż jeśli jakieś własności nie są niezbędne w wyjaśnieniach przyczynowych, to ich po prostu nie ma. Związek między pojęciami własności i przyczynowości podkreślał Sidney Shoemaker (1980/1999: 256-257), według którego predykat wskazujący na realną własność wówczas, kiedy odniesienie tego predykatu do danego przedmiotu pociąga za sobą określenie jego „mocy przyczynowej”. Inaczej mówiąc, przypisując danej rzeczy jakąś własność, mówimy tym samym, w jakie relacje przyczynowe weszłaby ta rzecz, gdyby znalazła się w określonej sytuacji. Odwołanie się do pojęcia przyczynowości służy Shoemakerowi do wprowadzenia rozróżnienia między predykatami, którym odpowiadają „rzetelne” czy też „prawdziwe własności” a tymi, którym takie własności nie odpowiadają. Jako przykłady „rzetelnych” własności Shoemaker podaje „posiadanie kształtu noża” lub „bycie zrobionym ze stali”. Nazywając przedmiot stalowym, stwierdzamy tym samym, jak zachowywałby się on w określonych sytuacjach, np. że topiłby się w określonej temperaturze. Na przeciwnym biegunie lokują się predykaty, które odnoszą się do „własności z Cambridge” — przykładem takowej może być posiadana przez Sokratesa „własność” bycia podziwianym przez Platona.

Teza mówiąca, iż własności przyczynowo nieefektywne nie są „rzetelnymi” własnościami może być uznana za szczególny przypadek zasady oszczędności ontologicznej. Zasada ta, popularnie zwana „brzytwą Ockhama”, głosi, że „zbytecznie jest przyjmować wiele elementów tam, gdzie wystarczy najmniejsza ich ilość” (Ockham 1971 : 56). W omawianym przypadku owymi elementami nie są jednak byty jednostkowe ani zdarzenia, lecz własności. Testem na konieczność przyjmowania danej własności do naszej ontologii jest jej przyczynowa relewancja. Jeśli jakaś własność nie spełnia tego warunku, to możemy ją uznać za nieistniejącą.

Rozumowanie tu przedstawione można w łatwy sposób zastosować do problemu nonfaktualizmu semantycznego. Relacja własności semantycznych do własności dyspozycyjalnych wydaje się bowiem bardzo podobna do relacji własności mentalnych z fizycznymi. Mamy bowiem do czynienia z wydarzeniami, które można uznać za jednoznacznie „naturalne” — mianowicie zachowaniami językowymi. Wydarzenia te można opisywać w dwóch różnych językach — semantycznym, w którym mówimy o znaczeniu wypowiedzi i o normach, które ta wypowiedź spełnia bądź nie, oraz dyspozycyjalnym, w którym opisujemy ich przyczyny.

Z punktu widzenia celu, jakim jest wyjaśnienie zachowania, postulowanie istnienia własności semantycznych jest zbędne. Jeśli naszym celem odkrycie, co sprawia, że użytkownicy języka zachowują się werbalnie tak a nie inaczej, to podział ich zachowań na poprawne i błędne jest dla nas zupełnie nieinteresujący. Branie pod uwagę takiego podziału może wręcz zaciemniać sytuację — odróżniać od siebie zachowania niczym nie różniące się na poziomie behawioralnym. Przyczynowa historia wypowiedzenia przeze mnie jakiegoś słowa w sposób niepoprawny może być niemal identyczna z historią wypowiedzenia przeze mnie słowa w sposób poprawny. W obu przypadkach odpowiedni kontekst konwersacyjny uruchamia pewien schemat zachowania, prawdopodobnie warunkowany odpowiednim stanem mojej aparatury neuronalnej. Argument z przyczynowego wykluczenia pozwala więc stwierdzić, że skoro nie muszę się w wyjaśnieniu do niej odwoływać, to fakt

poprawności danego działania nie wpływa przyczynowo na zachowanie. Wynika to wprost z przyjętego wyżej założenia, iż normatywność zdań o znaczeniu nie implikuje ich wewnętrznie motywacyjnego charakteru.

Można wykazać, iż argument z przyczynowego wykluczenia zastosowany do własności semantycznych jest nawet bardziej skuteczny, niż ten oryginalnie zastosowany przez Kima w odniesieniu do tego, co mentalne. Nie trzeba tu bowiem zakładać, że wszelkie przyczyny mają charakter fizyczny. Nawet gdyby okazało się, iż wśród czynników determinujących zachowania językowe znajdują się jakieś nieredukowalnie mentalne własności, to i tak nie zmieni to zasadniczego schematu argumentacji. Jakiegokolwiek cechy (fizyczne, mentalne czy jeszcze inne) nie powodowałyby naszych zachowań językowych, nie mogą być one utożsamiane z własnościami semantycznymi. Zastosowanie tego argumentu nie wymaga zatem przyjęcia żadnych metafizycznych przesłanek, określających, jakie własności mogą wchodzić w relacje przyczynowe. Wniosek, iż własności semantyczne są przyczynowo nieefektywne, wynika wprost z analizy pojęcia znaczenia, *via* twierdzenie o normatywności semantycznej.

III.4 Przyczynowe wykluczenie a nonfaktualizm

Powyższe rozumowanie stanowi ogólny i aprioryczny argument na rzecz nonfaktualizmu semantycznego. Jest to argument ogólny, ponieważ nie jest on wymierzony w żadną konkretną teorię znaczenia. Na jakiegokolwiek domniemane fakty nie wskazywałyby dany teoretyk znaczenia, to, o ile jego koncepcja miałaby być uznana za poprawną analizę pojęcia znaczenia, musiałaby ona ujmować fenomen słabej normatywności. W takim zaś przypadku, można do jego teorii zastosować zasadę przyczynowego wykluczenia. Aprioryczny charakter tego argumentu polega na tym, iż wnioski z niego płynące są niezależne od empirycznych ustaleń na temat rzeczywistych przyczyn zachowań językowych.

Czy uznanie własności semantycznych za przyczynowo nieefektywne dowodzi rzeczywiście, że dyskurs semantyczny jest niefaktualny? Można tu argumentować, że własności semantyczne istnieją nawet jeśli nie są niezbędne w wyjaśnieniach przyczynowych. Wielu filozofów, w tym sam Wittgenstein (1958/1998: 41), było przekonanych, że gdy odwołujemy się do znaczenia, to nie czynimy tego w celu podania przyczyny, lecz raczej motywu czy też racji działania. Mamy więc do czynienia z dwoma różnymi sposobami wyjaśniania zachowania werbalnego. Możemy, po pierwsze, odwołać się do jego przyczyn i wówczas interesować nas będzie działanie umysłu rozumianego jako swego rodzaju mechanizm generujący wypowiedzi. Możemy też podać motywy jego działania, wśród których może znaleźć się fakt, iż właśnie taki sposób użycia jest poprawny.

Ten drugi sposób wyjaśniania zachowania jest jednak nadmiarowy. Jako iż w takim schemacie pojęciowym domniemane fakty semantyczne uznawane są z góry za nieefektywne przyczynowo, rodzi się pytanie, dlaczego w ogóle mielibyśmy opisywać ludzkie zachowania w kategoriach motywów? Nie zwiększy to przecież naszej zdolności przewidywania ich działania. Jeśli zaś odpowiemy, że chodzi nam o zrozumienie tych ludzi, to pojawi się pytanie, na czym owo rozumienie miałoby polegać? Jeśli stwierdzimy, iż rozumieć kogoś, to znać znaczenie jego wypowiedzi, to będziemy kręcić się w kółko, bowiem to właśnie kategoria znaczenia jest przez nonfaktualistę kwestionowana.

Wobec przedstawionej argumentacji można, rzecz jasna, wysunąć zastrzeżenie, że teza głosząca, iż tylko własności przyczynowo efektywne mogą znaleźć miejsce w ontologii, jest przykładem dogmatycznej metafizyki. Zwolennik nonfaktualizmu może jednak twierdzić, iż kategoria znaczenia jest po prostu zbędna. Mając do dyspozycji pełną wiedzę na temat dyspozycji, bylibyśmy w stanie przewidywać wszelkie zachowania werbalne użytkowników języka; jeśli zaś działanie dyspozycji okazałoby się indeterministyczne, to i tak będziemy w stanie przewidywać zachowania podmiotów w takim stopniu, w jakim w ogóle można je przewidzieć.

W takiej sytuacji dyskurs semantyczny okaże się po prostu zbędny, podobnie jak zbędny okazał się dyskurs magiczny czy astrologiczny.

Argument z przyczynowego wykluczenia przenosi zatem onus probandi na zwolennika faktualizmu. Jeśli mamy włączyć przyczynowo nieefektywne fakty i własności do naszej ontologii, to potrzebna jest do tego jakaś pozytywna racja. Zwolennik faktualizmu musi zatem sformułować taką koncepcję, która będzie zawierać analizę potocznego pojęcia znaczenia w kategoriach faktów czy też własności, które można włączyć do naszej ontologii pomimo, że są nieefektywne przyczynowo. Analizie tego typu koncepcji będzie poświęcony *Rozdział IV*.

Rozdział IV

Rozwiązania nieprzyczynowe

IV.1 Dyspozycje społeczne

Zdaniem Scotta Soamesa (1998b: 229), jedną z możliwych odpowiedzi na upadek teorii dyspozycjonalnej jest rozszerzenie pola faktów, wśród których chcemy odnaleźć te odpowiadające stwierdzeniom z dyskursu semantycznego. Według niego, jakkolwiek nie możemy stwierdzić, iż fakty semantyczne superwenują na faktach dotyczących dyspozycji, to nie wynika jeszcze z tego wniosek, iż nie superwenują one w ogóle na faktach nie-intencjonalnych, czy też nie-semantycznych.

Warto przypomnieć w tym miejscu znane rozróżnienie na superweniencję lokalną i superweniencję globalną. Superweniencję lokalną dotyczącą własności wybranych przedmiotów rozumie zwykle się w następujący sposób:

Własności typu A superwenują na własnościach typu B wtedy i tylko wtedy, gdy jeśli dla każdego możliwego świata, obiekty p_1 i p_2 są nieodróżnialne pod względem własności typu A , to są nieodróżnialne pod względem własności typu B (por. Paprzycka 2005: 51 i McLaughlin, Bennett 2008).

Relację superweniencji można rozumieć także globalnie. Superweniencja globalna jest relacją między możliwymi światami: jeśli mamy do czynienia z dwoma światami, które jako osobne całości są identyczne pod względem rozkładu własności bazowych, to również całościowy rozkład własności superweniujących będzie w tych światach identyczny. To, które własności uznajemy za bazowe, ma charakter względny — możliwe są różne poziomy superweniencji. Własności chemiczne, na przykład, możemy traktować jako superweniujące na fizycznych, lecz bazowe dla biologicznych. Różnica między tymi dwoma rodzajami superweniencji polega więc na tym, iż w przypadku superweniencji globalnej mówimy o rozkładzie własności na poziomie całych światów, natomiast w przypadku superweniencji lokalnej mówimy o konkretnym przedmiocie lub przedmiotach i jego / ich własnościach.

Relacja superweniencji globalnej jest, na co zwracają uwagę Kim (1989/2008: 89) i Paprzycka (2005: 51), niezwykle słaba — gdy stwierdzamy, że zachodzi, nie dowiadujemy się zbyt wiele. Jeśli weźmiemy pod uwagę nasz świat i świat możliwy, różniący się od naszego jakimś zupełnie nieistotnym szczegółem z poziomu własności bazowych — np. położeniem jednego atomu wodoru — to nie możemy powiedzieć nic na temat tego, jak w tym, niezwykle podobnym do naszego, świecie, będą się rozkładać własności wyższego rzędu. Może być tak, iż w tamtym świecie tych własności w ogóle nie będzie albo będą rozłożone w zgoła odmienny sposób.

Wydaje się zupełnie naturalnym przeświadczenie, iż własności semantyczne superweniują globalnie na jakichś — choć nie potrafimy określić jakich — własnościach podstawowych. Soames (1998b: 229) zauważa jednak, że takie przekonanie jest bardziej wyrazem naszej wiary niż uzasadnionym przeświadczeniem. Sama zaś wiara w to, że domniemane własności semantyczne pozostają w — niezwykle słabej — relacji superweniencji globalnej do bliżej nieokreślonych własności bazowych, nie uzasadnia tezy, iż zdania należące do dyskursu semantycznego mają charakter faktualny. Z drugiej jednak strony, rozważania zawarte w poprzednim rozdziale pozwalają stwierdzić, iż własności semantyczne nie superweniują lokalnie

na własnościach przyczynowo determinujących moje sposoby zachowania, a więc na moim uposażeniu neurofizjologicznym oraz dyspozycjach zachowaniowych.

Zwolennik faktualizmu może w tej sytuacji wywodzić, iż w dyskursie semantycznym mówimy o własnościach i faktach złożonych, które same w sobie nie są przyczynowo efektywne, ale ich elementy składowe mogą cechować się taką efektywnością. Przyjmowanie własności złożonych jako elementów ontologii jest rozwiązaniem dość kontrowersyjnym (odrzuca je choćby Armstrong 1997: 26), jednak odrzucenie *a priori* pomysłu, iż własności semantyczne mają właśnie taki charakter, byłoby stawianiem proponentowi realizmu w odniesieniu do dyskursu semantycznego zbyt wysokich wymagań.

Przyjęcie tego założenia oznaczałoby, iż faktów odpowiadających za prawdziwość zdań o znaczeniu można szukać także poza podmiotem i jego własnościami wewnętrznymi. Być może Kripke pomylił się przyjmując, iż odpowiedzią na pytanie o źródła znaczenia musi być wskazanie na jakiś „dotyczący mnie fakt” (Kripke 2007: 42).

Pierwszym rodzajem „zewnętrznych” faktów, do których można się odwołać poszukując odniesień dla zdań o znaczeniu, są fakty dotyczące społeczności językowej, której jestem członkiem. Tego typu podejście pozwala na rozwiązanie problemu normatywności znaczenia w łatwy sposób. Poprawne, w myśl tej koncepcji, byłyby te wypowiedzi, które zgadzają się ze społecznie akceptowanym sposobem użycia wyrażenia; natomiast niepoprawne byłyby te, które są względem niego dewiacyjne.

Odniesienie do społeczności jest *explicite* zawarte w zaproponowanym przez Kripkensteina „rozwiązaniu sceptycznym”. Według Kripkego, słynny argument Wittgensteina przeciwko językowi prywatnemu jest w istocie konsekwencją jego paradoksu (Kripke 2007: 129).

Jeśli nasze rozważania do tej pory były poprawne, to (...) o ile rozważamy jedną osobę w izolacji, pojęcie reguły jako prowadzącej

osobę ją przyjmującą nie może mieć żadnej istotnej treści. Nie ma, jak widzimy, żadnych warunków prawdziwości ani faktów, dzięki którym ta osoba pozostaje, bądź nie, w zgodzie ze swoimi przeszłymi intencjami (...).

Sytuacja staje się zasadniczo odmienna, gdy rozszerzymy nasze pole widzenia tak, by nie obejmowało ono jedynie samej osoby kierującej się regułą i pozwolimy sobie przyjrzeć się jej jako wchodzącej w interakcje z szerszą wspólnotą (Kripke 1982/2007: 144).

Kripke rozważa w tym miejscu przykład dziecka nabywającego pojęcie dodawania. Jasne jest, że nauczyciel kontrolujący postępy ucznia nie uzna każdej jego wypowiedzi za zgodną z regułą. Ponieważ każdy z nas, użytkowników języka polskiego, ma podobne doświadczenia z procesem nabywania pojęć takich jak „dodawanie”, aby uwzględnić wymiar normatywności, analiza znaczenia nie może się zatrzymywać na tym, co określa sposób mojego zachowania. Musi również brać pod uwagę to, jak szersza społeczność reaguje na moje postępowanie.

Tyler Burge w swoim znanym tekście *Individualism and the mental* opisuje przywoływaną już w *Rozdziale II* sytuację, w której pacjent zgłasza się do lekarza, uskarżając się na artretyzm (Burge 1979/1996: 129–131). Pacjent ten używa słowa „artretyzm” inaczej niż cała wspólnota językowa — sądzi, iż odnosi się ono do każdego bólu kończyn. W rzeczywistości zgłasza się więc do lekarza z zupełnie inną przypadłością. Możemy jednak wyobrazić sobie przeciwną faktom sytuację, w której modalny odpowiednik tego człowieka udałby się do lekarza z podobnym problemem, przy czym miałoby to miejsce w możliwym świecie, w którym reszta ludzkości używa słowa „artretyzm” podobnie jak on (czyli określa nim każde schorzenie kończyn). Zdaniem Burge’a (1979/1996: 131) podobne przykłady można by mnożyć, używając najrozmaitszych terminów.

Wnioski z przykładu Burge’a da się przedstawić w terminologii superweniencji. Może bowiem zdarzyć się, że osoba z naszego świata i jej modalny odpowiednik

będą doskonale identyczni pod względem fizycznym i behawioralnym, jednak jedna z nich będzie używać pewnego wyrażenia poprawnie, a druga nie. Dowodzi to, iż treść wyrażen i fakty określające wzorce poprawnego ich stosowania nie superwenują na fizycznych i behawioralnych własnościach użytkowników języka.

Na zależność treści naszych przekonań i znaczeń słów od kontekstu społecznego zwracał również uwagę Hilary Putnam w swojej teorii społecznego podziału pracy językowej. Zdaniem tego autora, często używamy pewnych wyrażen nie potrafiąc ani podać, ani zastosować kryteriów ich stosowności. Weźmy przykład słów „wiąz” i „buk” (Putnam 1975/1998: 111). Zarówno Putnam, jak i piszący te słowa, nie potrafią rozróżniać desygnatów tych wyrażen — czyli nie potrafimy odróżnić wiązków od buków. Jednak mimo to, ze słowami „wiąz” i „buk”, tak jak ich używamy, związane są różne znaczenia. Gdy mówię „wiąz”, to odnoszę się do wiązków, nie zaś do buków, nawet jeśli to, co odpowiada tym słowom „w mojej głowie” niczym się nie różni.

Aby wyjaśnić fakt, iż dwa wyrażenia w moim języku mogą mieć różne znaczenie, nawet jeśli nie potrafię powiedzieć, na czym ta różnica miałaby polegać, konieczne jest przyjęcie, iż o znaczeniu słów, których używam nie decydują moje stany wewnętrzne, lecz otoczenie społeczne, w którym się znajduję. W owym otoczeniu znajdują się z pewnością ludzie, którzy potrafią wiązki od buków odróżniać — są to tzw. eksperci, czyli ludzie, którzy opanowali metodę rozpoznawania czy dany przedmiot kwalifikuje się np. jako wiązek. To właśnie dzięki temu, że metody rozpoznawania wiązków i buków znajdują się w posiadaniu „zbiorowego ciała językowego” (Putnam 1975/1998: 113), mogę stwierdzić, że odnoszące się do nich wyrażenia mają różne znaczenie.

Obserwacja poczyniona przez Putnama jest w dużej mierze zdroworozsądkowa. Każdy kompetentny użytkownik języka zdaje sobie sprawę z tego, iż używa wielu terminów, których nie potrafi zdefiniować ani zastosować w odpowiedni sposób.

Gdy przeciętny użytkownik, niemający zbyt rozległej wiedzy na temat współczesnej fizyki, używa słowa „proton”, robi to z niewyrażonym wprost przeświadczeniem, iż jakkolwiek on sam nie potrafiłby zdefiniować tego wyrazu, to w społeczeństwie są ludzie, którzy to potrafią. Putnam (1975/1998: 114) twierdzi, iż zjawisko podziału pracy językowej ma charakter powszechny — każda społeczność posiada terminy, których kryteria stosowalności są znane jedynie wybranej grupie jej członków. Fakt, że inni mogą się nimi skutecznie posługiwać, musi być wyjaśniony poprzez odwołanie się do współpracy.

Przykłady skonstruowane przez Putnama odnoszą się jedynie do sytuacji, w których użytkownik języka nie zna warunków poprawnego stosowania danego terminu — czyli, innymi słowy, nie potrafi określić, do jakich przedmiotów ma się on odnosić. Podobne obserwacje można poczynić w odniesieniu do poprawnego użycia. Wielu użytkowników języka naturalnego — często w sposób całkowicie nieświadomy — łamie warunki poprawnego użycia. Można np. z dużą dozą pewności przyjąć, iż tylko niewielka część rodzimych użytkowników języka polskiego opanowała (i potrafi wyartykułować) zasady rządzące użyciem imiesłowu uprzedniego „przyszły” czy wyrażenia „w międzyczasie”. Wielu z nich *implicite* przyjmuje zasadę, iż nawet jeśli my sami nie umiemy poprawnie stosować pewnych wyrażeń, to w naszej społeczności istnieją ludzie, którzy warunki takiej poprawności potrafią sformułować.

Czy jednak odwołanie do społeczności pozwala odpowiedzieć na pytanie Kripkensteina? Czy znajdziemy fakty odpowiadające zdaniom z dyskursu semantycznego, jeśli rozszerzymy dziedzinę poszukiwań o fakty społeczne?

Sam Kripke jest przeświadczony, że nie jest to dobre rozwiązanie. Pisze tak:

Gdyby Wittgenstein próbował podać konieczne i wystarczające warunki pokazujące, iż to „125”, a nie „5” jest „prawidłową” odpowiedzią na „68 + 57”, mógłby zostać posądzony o popadnięcie w błędne koło. Można bowiem uważać, iż twierdzi on, że moja odpowiedź jest

poprawna wtedy i tylko wtedy, jeśli zgadza się z odpowiedziami innych. Lecz nawet jeśli zarówno sceptyk, jak i ja zaakceptowalibyśmy z góry to kryterium, to czyż sceptyk nie mógłby utrzymywać, iż podobnie jak myliłem się co do tego, co w przeszłości znaczyło „+”, tak też myliłem się co do „zgadzać się”? Tak naprawdę, próba redukcji reguły dodawania do innej reguły: „Odpowiadaj na problemy z zakresu dodawania dokładnie tak, jak robią to inni!” - podpada pod zarzuty Wittgensteina co do „reguły interpretacji reguły” w taki sam sposób, jak wszystkie inne proponowane redukcje. Wittgenstein kładłby nacisk na to, że taka reguła niepoprawnie opisuje to, co robię: gdy dodaje, nie konsultuję się z innymi (...).

Wiele z rzeczy, które można powiedzieć o jednostce w „prywatnym” modelu języka, można podobnie powiedzieć o całej wspólnocie w modelu Wittgensteina. W szczególności, jeśli cała wspólnota zgadza się co do pewnej odpowiedzi i obstaje przy niej, to nikt nie może jej poprawić (Kripke 1982/2007: 177-178).

Cytat ten pokazuje, że Kripkenstein nie przyjmuje teorii, w której znaczenie redukowane jest do opinii większości. Dostrzega on problemy, jakie by się z przyjęciem takiego rozwiązania nieuchronnie wiązały.

Podstawowym problemem jest to, że sytuacja wspólnoty nie jest wcale lepsza niż sytuacja jednostki — na co zwracał uwagę również Simon Blackburn (1984: 293–295). Według Blackburna, jeśli przyjrzymy się wspólnocie jako całości, nie znajdziemy tam żadnego faktu, który umożliwiłby nam wykluczenie tezy, iż tak naprawdę powinna ona się kierować jakąś niestandardową regułą. Być może wszyscy mylimy się, stosując dane pojęcie. Jeśli zaś chcemy wykluczyć taką hipotezę, to musimy wskazać na jakiś fakt — inny niż konsensus między użytkownikami języka. Tym, co musimy wykluczyć, jest więc hipoteza zbiorowego błędu.

Można powiedzieć, że uwagi Blackburna przenoszą problem normatywności znaczenia na poziom społeczeństwa. Jeśli zgodzimy się, że podmiotem, który nadaje znaczenia symbolom, nie jest jednostka, lecz społeczność jako całość, to w stosunku do owej społeczności powstaje problem tego, czy może ona się mylić — także w sposób systematyczny. Jeśli tak, to znaczy, że odniesienie do społeczności nie wystarcza do wskazania faktów odpowiadających zdaniom dyskursu semantycznego. W tym celu musielibyśmy wskazać na standard użycia zewnętrzny wobec społeczności.

Założenie, że każda wypowiedź, co do poprawności której panuje społeczny konsensus, jest z definicji poprawna, nie byłoby zresztą zgodne z naszym zwykłym pojęciem kierowania się regułą i koncepcją znaczenia. Jeśli bowiem rozważania z *Rozdziałów II i III* są poprawne, to możliwość popełnienia pomyłki przez użytkowników języka jest esencjalną cechą pojęcia znaczenia. Jeśli zatem wspólnota nie może się pomylić, to mówienie o znaczeniu, jakie nadaje ona stosowanym przez siebie wyrażeniom, nie ma większego sensu. Intuicyjnie wyczuwamy przecież, że w pewnych sytuacjach to mniejszość — nie zaś większość — ma słuszność co do zasad stosowania danego terminu.

Zdarza się to na przykład wtedy, gdy dokonuje się autentycznych aktów klasyfikacji. Jeśli np. przy słabej widoczności w górach spieram się z większością grupy o to, czy widoczna w oddali biała plama to owca czy owczarek podhalański, to większość kompetentnych użytkowników języka dopuszcza możliwość, że to ja mogę mieć rację. Sam fakt, iż większość jest innego zdania, nie stanowi argumentu rozstrzygającego. Co więcej, nawet jeśli wszyscy zgodzilibyśmy się, iż widzimy psa, to *implicite* przyjmujemy, iż, nawet mimo naszej zgody, możliwe jest, że wszyscy się mylimy.

Podobne przykłady pokazują również problem, przed którym stanęłaby teoria definiująca znaczenie przez odwołanie się do tego, co mówią eksperci. Załóżmy, iż pod wpływem lektury Putnama zgodzilibyśmy się na następującą formułę:

„symbol s (np. słowo „wiąz”) ma znaczenie z wtedy i tylko wtedy gdy odpowiedni ekspert sądzi, że tak jest”. Również w tym przypadku pojawia się pytanie: czy nasz ekspert może się pomylić? Jeśli tak, to nasza formuła okazuje się nieprawdziwa i jesteśmy zmuszeni odwołać się do jakiegoś faktu zewnętrznego w stosunku do zachowań i stanów mentalnych naszego eksperta. Czyli nadal nie potrafimy wskazać faktu, który odpowiada zdaniom o znaczeniu. Założenie o nieomyślności eksperta będzie miało w odniesieniu do pojęcia znaczenia identyczne konsekwencje, jak to o nieomyślności społeczności.

Kolejnym problemem omawianej koncepcji, jest to, na jakiej podstawie mielibyśmy wybierać ekspertów. Jeśli uznalibyśmy, że ekspertem od danego wyrażenia jest ta osoba, która (zazwyczaj) postępuje zgodnie z wzorcami poprawnego użycia, to nasza koncepcja nabierze charakteru błędnego koła — odwołujemy się do pojęcia eksperta chcąc wyjaśnić poprawność, lecz zarazem chcąc zdefiniować „eksperta” musimy odwołać się do pojęcia „poprawności”. Z drugiej strony możemy eksperta wybierać na zasadzie konwencjonalnej, np. głosowania; wówczas jednak kwestia wyjaśnienia poprawności pozostaje otwarta. Okazuje się, że nasze działania językowe ostatecznie opierają się na arbitralnej decyzji, co uniemożliwia rekonstrukcję normatywnego wymiaru języka.

Inna kwestia, której, jak się czasem twierdzi, teoria społeczna nie potrafi wyjaśnić została zarysowana przez Kripkego w przytoczonym powyżej cytacie. Jaka relacja powinna zachodzić między społecznie akceptowanymi standardami a moim własnym użyciem? Jak wskazywał Kripke, słowo „zgodny” można interpretować na niestandardowe sposoby. Przyjęcie, iż słowo „zgodny” ma standardowe znaczenie, naraża nas zatem, ponownie, na problem regresu interpretacji.

Zarzut regresu interpretacji w odniesieniu do tej akurat teorii znaczenia wydaje się jednak nietrafny. Zasadnie można go postawić wobec tych koncepcji, które wskazują na — mający odpowiadać zdaniom o znaczeniu — fakt, który sam w sobie ma charakter semantyczny. Z taką sytuacją mieliśmy do czynienia choćby

w przypadku koncepcji, która wyjaśniać miała znaczenia wyrażen językowych odwołując się do stanów intencjonalnych użytkowników języka. Zgodność lub niezgodność mojego zachowania z wzorcem przyjętym przez większość członków wspólnoty lub ekspertów może tymczasem być uznane za fakt czysto behawioralny. By stwierdzić tego rodzaju zgodność, nie jest konieczne odwołanie do faktów semantycznych.

Odparcie tego zarzutu nie zmienia jednak ogólnej, negatywnej oceny teorii społecznej. Jej główną wadę można ująć następująco: albo przyjmujemy, iż społeczność, czy też eksperci, na których cedujemy kompetencje semantyczne, mogą popełnić pomyłkę albo taką możliwość wykluczamy. W pierwszym przypadku pozostajemy z zadaniem wyjaśnienia tego, co konstytuuje standard, z którym sposób, w jaki wspólnota używa języka jest zgodny lub nie. Powołanie się na opinię wspólnoty pozwala wyjaśnić normatywny aspekt znaczenia, który pojawia się, gdy analizujemy posługiwanie się językiem przez jednostkę. Jednak problem normatywności zostaje w tej koncepcji jedynie przeniesiony, nie zaś rozwiązany.

Alternatywne rozwiązanie — w którym przypisujemy społeczności semantyczną nieomyślność, nie musi kłopotać się koniecznością wskazania na coś, do czego wspólnota w swoich decyzjach musiałaby się odnosić. Problemem z tym rozwiązaniem jest jednak taki, że jest ono całkowicie nieintuicyjne. Jako wspólnota zakładamy bowiem *implicite*, iż możemy się mylić. Podobny zarzut sformułował onegdaj Putnam w odniesieniu do relatywizmu pojęciowego: nie można twierdzić, iż „prawdziwy” znaczy to samo, co „zasadnie stwierdzalny w zgodzie normami kultury zachodnioeuropejskiej”, bo zdanie, w którym stwierdzalibyśmy taką równoznaczność, nie jest zasadnie stwierdzalne w kulturze zachodnioeuropejskiej (Putnam 1982/1998: 280–281). Innymi słowy — relatywizm upada, ponieważ my nie jesteśmy relatywistami. Podobnie upada społeczna teoria znaczenia, gdyż jako społeczność dopuszczamy możliwość, że to my się mylimy, a jednostka ma rację.

IV.2 Rodzaje naturalne

Upadek teorii wyjaśniających znaczenie poprzez odwołanie się do faktów społecznych można wyjaśnić w dość prosty sposób. Próbowano w nich wyjaśnić znaczenie przez odwołanie się do działań innych podmiotów. Jednak te same problemy, z którymi mierzyliśmy się w związku z pojedynczym podmiotem, dotyczyć będą także podmiotu zbiorowego.

Wyjściem z takiej sytuacji mogłoby być zwrócenie uwagi na fakty, które nie dotyczą podmiotów posługujących się językiem. Jeśli szukamy czegoś, co miałyby spełniać funkcję standardu, z którym użycie języka miałyby być zgodne bądź nie, należałoby się przede wszystkim przyjrzeć światu zewnętrznemu.

Teorią semantyczną, która bierze pod uwagę rzeczywistość zewnętrzną wobec użytkowników języka, jest przyczynowa teoria odniesienia, rozwinięta przez Kripkego i Putnama oraz — ściśle z nią powiązana — koncepcja rodzajów naturalnych. Przekonanie, iż to właśnie ta teoria pozwala rozwiązać problem postawiony przez Kripkensteina, wyraziła Penelope Maddy w tekście *How the causal theorist follows a rule*. Wielu innych komentatorów (por. np. Kusch 2006: 133), nawet jeśli nie było przekonanych do zalet tego ujęcia, podkreślało paradoksalność sytuacji, w której Kripke, omawiając problem sceptyczny w sprawie znaczenia, nie zajmuje się przyczynową teorią odniesienia, mimo iż sam był jednym z twórców tej teorii. Wydaje się, że nieuwzględnienie tego rozwiązania stanowi lukę w argumentacie sceptycznym.

Przyczynowa teoria odniesienia została rozwinięta w dwóch wersjach. Hilary Putnam przedstawił ją dla tzw. rodzajów naturalnych, zaś Saul Kripke dla nazw własnych i przedmiotów indywidualnych. Koncepcja nazw własnych Kripkego stanowiła kontrpropozycję dla tak zwanej deskrypcyjnej teorii nazw, której podwaliny stworzył Bertrand Russell w tekście *Denotowanie*. Russell (1905/1967: 253–257) zaproponował metodę analizy wyrażeń, w których występują tak zwane zwroty denotujące, takie jak „obecny król Francji”. Wedle Russella zdanie typu

„Obecny król Francji jest łysy” należy rozumieć w następujący sposób: „istnieje takie jedyne x , że x jest obecnym królem Francji i x jest łyse”. W ten sposób możemy stwierdzić, że zdania, w którym używamy tego typu zwrotów, są fałszywe, jeśli takie zwroty do niczego się w aktualnej sytuacji się nie odnoszą.

Wedle Quine’a (1953/1969a: 15–19), teorię Russella można zastosować do problemu nazw własnych — zwłaszcza do nazw pustych takich jak „Pegaz”. Jeśli nie chcemy przyjmować nieintuicyjnej tezy, że przedmiot odpowiadający takiej nazwie w pewnym sensie istnieje, to musimy znaleźć metodę analizy takich zdań, która pozwoli wyeliminować nazwę własną. Wedle Quine’a najlepszy sposób polega na potraktowaniu nazwy jako skrótu deskrypcji — np. nazwy „Pegaz”, jako skrótu dla „skrzydlaty koń”. Wówczas możemy przykładowe zdanie „Pegaz został zrodzony z krwi Meduzy” zanalizować jako „istnieje takie jedyne x , że x jest skrzydlatym koniem i x zostało zrodzone z krwi Meduzy”. Jako, że nie istnieje taka rzecz jak skrzydlaty koń, możemy ocenić całe to zdanie jako fałszywe. Konsekwencją tej teorii jest uznanie nazw własnych za skróty deskrypcji. Gdy mówimy „Pegaz”, to w niejawnym sposób wypowiadamy pewną deskrypcję określoną — chcemy odnieść się do obiektu posiadającego pewną własność.

Kripke (1972/1998: 75 - 90) wysuwa dwa zarzuty wobec teorii deskryptywnej. Pierwszy z nich to argument epistemiczny — otóż mogłoby się okazać, iż przedmiot, do którego chcemy się odnieść, nie spełnia deskrypcji, które standardowo wiążemy z jego nazwą, gdyby np. wyszło na jaw, że Sienkiewicz nie napisał *Pana Wołodyjowskiego*. Częściowym rozwiązaniem tego problemu jest teoria wiązek deskrypcji Searle’a (1958: 171–173). Według tej koncepcji, znaczenie nazwy własnej nie jest identyczne z pojedynczą deskrypcją, lecz z wiązką różnych deskrypcji różnych użytkowników języka. Ułatwia to rozważanie sytuacji, w której przedmiotowi, do którego chcemy się odnieść, nie przysługuje akurat własność, o której mówimy w naszej deskrypcji. Jeśli przyjmujemy wiązkową teorię nazw, to możemy dopuścić tego typu scenariusze i twierdzić, iż nawet jeśli Sienkiewicz nie napisał

Pana Wołodyjowskiego, to mówiąc „Sienkiewicz” nadal odnosimy się do tego samego autora. Dysponujemy bowiem innymi deskrypcjami, które pozwalają nam odnieść się do Sienkiewicza, takimi jak: „autor *Quo Vadis*”, „pierwszy Polak, który otrzymał nagrodę Nobla w dziedzinie literatury” itp.

Jednak to rozwiązanie nie pozwala poradzić sobie z problemem w sposób zadowalający. Są bowiem sytuacje, w których nie potrafimy sformułować dostatecznie bogatej wiązki deskrypcji. Kripke (1972/1988: 84–85) podaje przykład nazwy „Peano”. Jeśli już jesteśmy w stanie związać ją z pewną deskrypcją, to będzie ona brzmiała zapewne: „człowiek, który sformułował aksjomaty arytmetyki”. Czy jednak przy takim założeniu da się w ogóle stworzyć hipotezę, zgodnie z którą to nie Peano sformułował aksjomaty arytmetyki? Wydaje się to niemożliwe, co jest wnioskiem dość paradoksalnym — bowiem zdanie takie byłoby zwykłą hipotezą historyczną i nie widać powodu, dla którego miałyby być ona z góry uznana za fałszywą.

Kolejnym zarzutem Kripkego przeciw teorii deskrypcyjnej jest argument modalny. Jego sedno sprowadza się do stwierdzenia, iż osoba, o której mówimy, mogłaby nie mieć własności, które jej przypisujemy — nawet jeśli faktycznie je posiada. Jest możliwe, że Henryk Sienkiewicz w młodości zdecydowałby się na kontynuowanie studiów medycznych i resztę swych dni spędziłby praktykując medycynę. Jednak taka możliwość — jakkolwiek możemy o niej swobodnie rozprawiać — jest trudna do opisanego z punktu widzenia deskryptywnej teorii nazw. Jeśli „Henryk Sienkiewicz” znaczy to samo co „autor *Pana Wołodyjowskiego*”, sytuacja, w której „autor *Pana Wołodyjowskiego* nie zostałby autorem *Pana Wołodyjowskiego*”, zdaje się wykluczona, mimo że mamy silne przeświadczenie, że była to realna możliwość.

Wobec niedomagań teorii deskrypcyjnej Kripke proponuje całkowicie inną wizję mechanizmu odnoszenia się do przedmiotów za pomocą nazw własnych. Według niego (Kripke 1972/1988: 92–94), sytuacja z nazwami własnymi przedstawia się

następująco: na początku mamy do czynienia z „chrztem” — to jest pierwszym nadaniem nazwy przedmiotowi przez użytkownika języka. Akt chrztu ustala relację odniesienia pomiędzy nazwą a przedmiotem, do którego ta nazwa się odnosi. Następnie nazwa jest przekazywana innym użytkownikom języka, przy czym nie muszą oni być świadomi tego, od kogo przejęli swój sposób użycia. Ważne jest to, iż poprzez chrzest i następne akty komunikacji ustanawiają pomiędzy użytkownikiem języka a przedmiotem, do którego się odnosi pewien „łańcuch”. To właśnie istnienie tego łańcucha umożliwia posługiwanie się nazwami, także w sytuacji, w której nasze deskrypcje okazałyby się fałszywe lub nie potrafilibyśmy sformułować deskrypcji jednoznacznie wskazującej na obiekt, o który nam chodzi.

Hilary Putnam w podobny sposób ujmuje kwestię wyrazów oznaczających tak zwane rodzaje naturalne (Putnam 1975/1998: 124–129). Przykładami takich wyrazów mogą być nazwy pierwiastków i związków chemicznych („woda”, „złoto” itp.) oraz rodzajów i gatunków biologicznych („cytryna”, „tygrys” itp.). Putnam sprzeciwia się koncepcji znaczenia tych terminów opartej na „definicjach operacyjnych”, według której o tym, co możemy poprawnie nazwać np. „złotem” decydują testy, którymi poddajemy próbki metalu.

Problem z tą operacyjną wizją znaczenia polega na tym, że proponowane przez nią kryteria rozpoznawania, czy dany przedmiot zalicza się do określonego rodzaju naturalnego, są historycznie zmienne. Jeszcze Kant uważał za zdanie analityczne tezę „złoto jest żółtym metalem” (Kant 1783/1993: 17). Nasze dzisiejsze pojęcie złota jest już znacząco różne — w mniejszym stopniu opiera się ono na zewnętrznych, fenomenologicznych cechach złota, w większym zaś bierze pod uwagę wewnętrzną strukturę tego pierwiastka. Jeśli jednak zgodzilibyśmy się z operacyjną teorią znaczenia, to musielibyśmy przyznać, iż ekstensje takich terminów, jak „złoto” różnią się historycznie w zależności od rozwoju naszych technik operacyjnych. Jeśli dysponujemy dziś lepszymi metodami rozpoznawania złota niż ludzie w czasach Archimedesza, to (jakkolwiek paradoksalnie by to nie

brzmiało) ekstensja wyrazu „złoto” zgodna ze współczesnym jego stosowaniem jest węższa, niż odpowiednika tego słowa w starożytnej grece. Konsekwencją takiego podejścia byłby również relatywizm w odniesieniu do prawdy — jeśli bowiem zakresy poszczególnych terminów różnią się historycznie, to i zbiór zdań prawdziwych w różnych okresach czy językach będzie różny.

Kripke (1972/1988: 116 - 124) atakuje z kolei operacjonistyczne podejście do nazw rodzajów naturalnych używając podobnego argumentu jak w przypadku krytyki deskrypcyjnej teorii nazw. Zwraca więc uwagę, że mogłoby się okazać, iż rzeczy kwalifikujące się do zakresu danej nazwy nie mają którejs z cech, które im powszechnie przypisujemy, np. tygrysy nie są zwierzętami (tak jak kiedyś okazało się, że wieloryby nie są rybami).

Według Putnama i Kripkego, w stosunku do terminów odnoszących się do rodzajów naturalnych musimy przyjąć zatem inną koncepcję znaczenia. Gdy mówimy, że coś jest złotem (wodą, cytryną itp.), chodzi nam o to, że ta rzecz ma taką samą strukturę wewnętrzną jak inne kawałki złota (inne cytryny, próbki wody itp.). Historia ustalania i przekazywania terminów odnoszących się do rodzajów naturalnych jest podobna do tej, z którą mieliśmy do czynienia w przypadku nazw własnych. Dzięki łańcuchowi użyć mogą się odnosić np. do wiązków, nie potrafiąc wskazać na operacyjne kryteria odróżniania ich od innych drzew. Odniesienie zachodzi dzięki temu, iż jestem jednym z ogniw łańcucha, na którego końcu znajduje się egzemplarz należący do tego rodzaju naturalnego.

W jakim sensie zarysowana tu teoria może stanowić odpowiedź na problem postawiony przez Kripkensteina? Otóż uznanie, iż to, do czego odnoszą się nazwy jednostkowe i terminy naturalnorodajowe jest, przynajmniej po części, zdeterminowane przez coś zewnętrznego w stosunku do podmiotu, daje asumpt do wyjaśnienia kwestii normatywności. Rozważmy przypadek słowa „wiąz”. Jak zostało to już ustalone, z mojej dyspozycji do używania tego słowa nie da się odczytać, czy używam go poprawnie czy nie — ta sama dyspozycja może powodować zachowania

poprawne i niepoprawne. Jednak to, czy odnoszę to słowo do przedmiotu należącego do tego samego rodzaju naturalnego czy też nie, określa, czy można to użycie nazwać poprawnym (por. Maddy 1984: 475).

Przyczynowa teoria odniesienia dostarcza nam zatem podstaw do mówienia o normatywności w przypadku dwóch typów wyrażeń językowych — nazw własnych i nazw rodzajów naturalnych. Zarówno przedmioty jednostkowe, jak i rodzaje naturalne mają warunki tożsamości niezależne od naszych przekonań na ich temat — to, czy dany obiekt jest sobą oraz to, czy należy do danego rodzaju naturalnego, jest obiektywnym faktem.

Założeniem rozwiązania tego typu jest niewątpliwie, iż podstawową formą normatywności języka jest normatywność wynikająca z reguł poprawnego stosowania. Kluczową kwestią w tym ujęciu jest więc to, czy używając omawianych wyrażeń „trafiamy” w odpowiedni desygnat. Takie rozumienie normatywności może wydawać się nietrafne w świetle rozważań zawartych w drugim rozdziale niniejszej pracy. Zwolennik przyczynowej teorii odniesienia może jednak odpowiedzieć, iż skoro nie udało nam się wyjaśnić, co miałyby konstituować standardy poprawnego użycia, to musimy zadowolić się określeniem standardów poprawnego stosowania.

Kolejny możliwy do postawienia zarzut mógłby wskazywać, iż przyczynowa teoria odniesienia stosuje się do niezwykle ograniczonego fragmentu języka — obejmuje jedynie nazwy własne i terminy odnoszące się do rodzajów naturalnych. Skoro jednak niemożliwe okazuje się wyjaśnienie fenomenu znaczenia na poziomie ogólnym, to być może należy skoncentrować się na wyjaśnieniu prostych zjawisk językowych — takich właśnie, jak odniesienie się do przedmiotów jednostkowych i rodzajów naturalnych. Taka wizja zdaje się optymistycznie zakładać, że da się wyjaśniać skomplikowane zjawiska językowe poprzez odniesienie ich do zjawisk prostych. Nie widać jednak powodu, dla którego należałoby taki projekt *a priori* odrzucić.

Problemy, które rodzi próba rozwiązania problemu Kripkensteina za pomocą przyczynowej teorii odniesienia, mają bardziej zasadniczy charakter. Dotyczą dwóch kluczowych elementów tej koncepcji — chrztu, czyli pierwotnego ustalenia odniesienia oraz dalszego przekazywania tegoż odniesienia, czyli łańcucha przyczynowego.

Słabości swojej koncepcji był świadom sam Kripke., który *explicite* twierdził, że nie podaje teorii odniesienia, lecz jedynie proponuje lepszy obraz tego, jak funkcjonuje odniesienie (Kripke 1972/1988: 94–95). O stworzeniu teorii odniesienia można by mówić wtedy, gdyby podano zbiór koniecznych i wystarczających warunków, jakie musi spełniać zjawisko, aby zasłużyć na miano „odniesienia” — przy czym w sformułowaniu tych warunków nie może występować słowo „odniesienie”. Tymczasem, jak zauważa omawiany filozof:

Zauważmy, że poprzednie warunki nie usuwają raczej pojęcia odniesienia; przeciwnie, uznają za dane pojęcie zamiaru, by posłużyć się tym samym odniesieniem. Występuje tu również odwołanie do pierwszego chrztu, który wyjaśnia się w terminach bądź ustalania odniesienia przez deskrypcję, bądź wskazywania (Kripke 1988: 97–98).

Rozważmy najpierw kwestię chrztu. Na pierwszy rzut oka sprawa jest jasna — aby nadać nazwę jakiemuś przedmiotowi czy też rodzajowi naturalnemu, potrzebna jest intencja nadania właśnie takiej nazwy. Bez towarzyszącego mojemu wypowiedzeniu aktu mentalnego, wydanie pewnego dźwięku nie będzie aktem nazwania — równie dobrze może to przecież być nieartykułowany wyraz moich emocji związanych z danym przedmiotem.

Po drugie, wprowadzenie nazwy do języka musi odbywać się zgodnie z pewnymi konwencjami, które mają językowy charakter. Jak zauważył Wittgenstein, „wiele musiało być już przygotowane w języku, by proste nazywanie miało swój sens” (Wittgenstein 1953/2000: § 257). Samo wypowiedzenie słowa w obecności przedmiotu nie stanowi jeszcze nazwania.

Kolejną kwestią, z którą musi poradzić sobie propozycja Maddy, jest tak zwany *qua-problem* czyli „problem *jako*” (por. Devitt i Sterelny 1999: 79–81; Hattiangadi 2007: 143–144, Kusch 2006: 134–135, Odrowąż-Sypniewska 2006: 115 – 116). Załóżmy, że widzę pewnego swoiście wyglądającego ptaka, i chcąc dokonać „pierwszego chrztu” wypowiadam słowo „puchacz”. Rodzi się pytanie, co sprawia, iż wypowiadając to słowo, nazywam ten przedmiot jako przedstawiciela pewnego gatunku biologicznego, a nie jako jednostkowego osobnika. Akt chrztu może się przecież również odnosić do swoistego ubarwienia, jakie cechuje tego ptaka, kształtu jego uszu itp. Z bardziej prawdopodobnych opcji, mogę również chcieć nadać nazwę całej rodzinie ptaków, do której należy ów gatunek. Wydaje się, iż bez odwołania do pewnej szczególnej intencji, do tego, iż chcę nazwać w ten sposób gatunek biologiczny właśnie, nie da się rozstrzygnąć, która z tych interpretacji mojego zachowania jest właściwa. Jednak intencja, do której jestem zmuszony się odwołać, jeśli chcę wyjaśnić akt chrztu, jest stanem mentalnym posiadającym określoną treść — musi się bowiem w nim zawierać odniesienie do pojęcia „gatunku biologicznego”. W tym momencie powstaje problem „regresu interpretacji”: co miałyby bowiem konstituować znaczenie owego pojęcia „rodzaju biologicznego”? Okazuje się, iż w ramach teorii, która miała wyjaśniać pojęcie „odniesienia” zakładamy, że pewne rzeczy (czyli moje intencje) mają ustalone znaczenie.

Maddy zdaje sobie sprawę z problemów z tym, do czego odnosi się chrzest. Jej odpowiedź odwołuje się „po trosze do teorii neurologicznej, a po trosze do spekulacji” (Maddy 1984: 465). Powołuje się ona na teorię zespółów komórkowych, które mają odpowiadać za percepcję przedmiotów określonego rodzaju, takich jak kąty proste (por. np. Maruszewski 2001: 44). Zdaniem Maddy, o tym, że nadając czemuś nazwę „puchacz”, wprowadzam do języka termin odnoszący się do rodzaju naturalnego, a nie nazwę przedmiotu indywidualnego czy barwy, decyduje fakt, iż w tym momencie wytwarza bądź uaktywnia się w moim mózgu określony

zespół komórkowy. Zespół ten będzie się następnie aktywował, kiedy będę miał do czynienia z przedstawicielem tego gatunku biologicznego. Tego typu reakcja nie nastąpi tylko wtedy, gdy będę miał do czynienia z tym samym konkretnym osobnikiem lub gdy będę widział przedmiot tej samej barwy. To właśnie wzorce aktywności neuronalnej mają, w myśl tej teorii, decydować, do jakiego rodzaju rzeczy się odnoszę.

Koncepcja Maddy zdaje się opierać w większej mierze na spekulacji niż na neurobiologii — nie wiadomo bowiem nic o tym, aby w mózgu istniały zespoły komórkowe odpowiadające za percepcję gatunków czy rodzajów biologicznych. Teoria zespołów komórkowych dotyczy raczej percepcji takich rzeczy jak kształty (Maruszewski 2001: 44). Jednak można przyjąć życzliwe dla badaczki założenie, iż neurobiologom uda się ostatecznie wyodrębnić stany mózgowe odpowiadające za percepcję gatunków biologicznych czy pierwiastków chemicznych. Mielibyśmy wówczas do czynienia z następującym obrazem działania języka: gdy po raz pierwszy nazywam jakąś rzecz „złotem” to uaktywnia się w moim mózgu pewien zespół komórkowy. Kiedy stosuję ową nazwę do próbki metalu o tej samej strukturze wewnętrznej, co pierwotna próbka, uaktywnia się w moim mózgu ten sam zespół komórkowy, co w przypadku pierwszego użycia.

W tym momencie pojawia się jednak pytanie: czy teoria Maddy nie jest jedynie próbą restytucji koncepcji dyspozycjonalnej, ze wszystkimi aporiami, w jakie się ona wklęła? Jeśli bowiem o tym, do czego się odnoszę, ma decydować fakt, iż w moim mózgu uaktywnia się pewien zespół komórkowy, to jak wyjaśnić możliwość popełniania przeze mnie błędów? Jeśli o tym, do jakiego rodzaju przedmiotów się odnoszę, decyduje aktywacja pewnego obszaru w moim mózgu, to czy nie wyklucza to automatycznie sytuacji, w której mógłbym się pomylić co do tego, jak stosować dane wyrażenie?

Założmy, iż słowem „puchacz” ochrzciłem przedstawiciela pewnego gatunku biologicznego. Jednak później omyłkowo nazwałem tym słowem również inne

zwierzę, dość podobne z wyglądu do puchacza i należące do tej samej rodziny puszczykowatych. Co więcej, w moim mózgu uaktywniła się ta sama grupa neuronów, która uaktywniła się w sytuacji pierwotnego chrztu. Czy można stwierdzić, że popełniłem błąd? Jeśli o tym, co poprawne, decyduje aktywacja pewnej grupy neuronów, to nie, mimo że wydawałoby się to naturalne w tym kontekście.

Rozwiązaniem tego problemu wydaje się stwierdzenie, iż z poprawnym użyciem mamy do czynienia wtedy, gdy mam przed sobą przedstawiciela lub próbkę danego rodzaju naturalnego oraz gdy w moim mózgu aktywuje się odpowiednia grupa neuronów. Jednak i to rozwiązanie budzi wątpliwości — nietrudno wyobrazić sobie sytuację, w której wypowiedziałbym słowo „puchacz” w obecności stworzenia należącego do odpowiedniego gatunku biologicznego, lecz nie nastąpiłaby aktywacja odpowiedniej grupy komórkowej — mogę być przecież pod wpływem narkotyków, mieć uszkodzony układ nerwowy i tak dalej. Założenie, iż za każdym razem, gdy moje zastosowanie danego słowa będzie poprawne, mój stan neuronalny będzie tego samego rodzaju, wydaje się niczym nieuzasadnioną wiarą w *harmonia praestabilita*. Dość oczywiste wydaje się przypuszczenie, iż te same skutki behawioralne mogą być powodowane przez różne konfiguracje neuronalne. Zatem podwójny wymóg — zarówno aktywacji stanu neuronalnego jak i obecności przedmiotu odniesienia — jest wymogiem zbyt mocnym.

Drugim problemem dla przyczynowej teorii odniesienia jest kwestia przekazywania odniesienia innym użytkownikom języka. Wydaje się bowiem, iż do tego, aby przejąć odniesienie od drugiego mówcy, potrzebny jest akt intencji — co znowu rodzi problem regresu interpretacji. Można by argumentować, że przekazanie odniesienia jest sprawą czysto przyczynową — skutek zetknięcia się ze sposobem użycia pewnego terminu przez nadawcę w odbiorcy wykształca się stosowna dyspozycja do używania właśnie tego słowa. W takim obrazie nie mielibyśmy do czynienia z koniecznością odwoływania się do jakichkolwiek intencji.

Jednak ten obraz byłby nieprzekonujący ze względu na swą mechaniczność. Nie jest prawdą, iż ilekroć zetknę się z osobą używającą pewnego słowa, to automatycznie, jedynie na skutek tego kontaktu, przejmę od niej odniesienie. Możliwe jest przecież popełnienie błędu przy przekazywaniu odniesienia. Kripke (1972/1988: 162) podaje przykład słowa „Madagaskar”, które było używane pierwotnie w odniesieniu do części kontynentu afrykańskiego, a dopiero później, na skutek nieporozumienia, zaczęło funkcjonować jako nazwa wyspy. Jest również możliwe, iż zupełnie świadomie postanowię zignorować fakt, iż wszyscy wokół mnie używają słowa „Cezar” na oznaczenie rzymskiego dyktatora i będę go używał jako nazwy psa. Aby wyjaśnić te fenomeny, potrzebne jest wprowadzenie pojęcia ponownego ugruntowania nazwy (por. Devitt i Sterelny 1999: 76, Odrowąż-Sypniewska 2006: 84). Jednak pojęcie ponownego ugruntowania ma charakter jednoznacznie intencjonalny: cóż bowiem innego, jeśli nie intencja użytkownika języka decyduje o tym, że pewna wypowiedź nie jest niepoprawnym zastosowaniem danego terminu, lecz stanowi jego ponowne ugruntowanie? Rozważania nad tymi kwestiami doprowadziły Odrowąż-Sypniewską (2006: 86) do wniosku, iż nie powinniśmy mówić o „łańcuchu przyczynowym”, lecz o „łańcuchu komunikacyjnym”. Nie jest to zmiana jedynie stylistyczna — użycie terminu „łańcuch komunikacyjny” wskazuje, że omawiana teoria nie dostarcza narzędzi do redukcji pojęć „komunikacji” i „znaczenia” do kategorii przyczynowych.

Te wnioski nie prowadzą do odrzucenia teorii przyczynowej jako takiej. Możliwe, iż Putnam (1981/1998: 315 - 316) ma rację, gdy twierdzi, iż warunkiem koniecznym odnoszenia się do czegoś jest pozostawanie z tym czymś w kontakcie przyczynowym. Rozważenie tej kwestii wykraczałoby jednak poza ramy niniejszej pracy. Można wszakże stwierdzić, że samo istnienie relacji przyczynowej — zapośredniczonej lub bezpośredniej — nie stanowi wystarczającego warunku odniesienia. Z tego też względu teoria przyczynowa zawodzi jako próba rozwiązania problemu Kripkensteina.

IV.3 Koncepcja teleologiczna

Przedstawione dotychczas koncepcje nie potrafiły wyjaśnić kwestii normatywności znaczenia. Być może w celu rozwikłania tej zagadki konieczne jest zatem odwołanie się do koncepcji, w której wprowadza się wprost rozróżnienie na poprawne i niepoprawne działania. Taki charakter mają wszelkie koncepcje teleologiczne, to jest takie, które wskazują na pewien cel działania. Podejście teleologiczne zostało we współczesnej nauce zrehabilitowane za sprawą teorii ewolucji, która jawnie posługuje się takimi kategoriami jak sukces ewolucyjny, cel i temu podobne.

Na teorii ewolucji i zawarte w niej treści teleologiczne jako możliwe rozwiązanie problemu Kripkensteina wskazała Ruth Millikan w swoim tekście *Truth rules, hoverflies and Kripke-Wittgenstein paradox*. Millikan wskazuje na fakt, iż w zachowaniu zwierząt można wyróżnić dwa rodzaje reguł, którymi one się kierują — reguły bliższe i dalsze. Reguły bliższe to standardowe schematy postępowania zwierząt określonego rodzaju (Millikan 1990/2002: 215 - 221). Reguły dalsze określają natomiast, które schematy postępowania przyczyniają się do osiągnięcia sukcesu ewolucyjnego. Autorka przedstawia to rozróżnienie na przykładzie owadów z rodziny bzygowatych. Samiec ma dyspozycję do kierowania się regułą bliższą polegającą na gonieniu za wszelkimi obiektami poruszającymi się w określony sposób. Jednak dalszą regułą kierującą jego zachowaniem jest chęć zdobycia partnerki. Nie każde zachowanie, które jest zgodne z dyspozycją do gonienia za przedmiotami wyglądającymi i zachowującymi się jak samice, będzie sukcesem z punktu widzenia reguły dalszej — samiec może się bowiem pomylić i uznać za samicę coś, co nią nie jest. Wówczas jego zachowanie, jakkolwiek zgodne z jego dyspozycją, będzie w pewnym nietrywialnym sensie błędne.

Podział na reguły bliższe i dalsze odnosi się również do zachowań wyuczonych. Millikan (1990/2002: 222–223) rozważa przykład szczura, który nauczył się, iż nie powinien jeść pewnych rzeczy smakujących w określony sposób, gdyż w jego

przypadku wywołało to chorobę. Jest to dla niego reguła bliższa. Jednak dalszą regułą, którą kieruje się szczur, jest unikanie zatrutego pożywienia. Dyspozycja, którą wyrobił w sobie rzeczony szczur może, ale nie musi być zgodna z tak wyrażoną regułą dalszą. Unikanie pożywienia o pewnym smaku nie musi uchronić zwierzęcia przed zatruciem.

Wprowadzenie rozróżnienia na reguły bliższe i dalsze pozwala Millikan na proponowanie rozwiązania kluczowego problemu: co konstytuuje standard, z którym zgodność decyduje o poprawności moich zachowań językowych? Autorka twierdzi, że rolę tę odgrywa sukces ewolucyjny. Wskazanie na to rozwiązanie pozwala wyjaśnić klęskę dotychczasowych prób rozwiązania sceptycznego paradoksu: wszystkie one koncentrowały się na faktach, które już się wydarzyły. Tymczasem sukces ewolucyjny jest sprawą przyszłości. Nie możemy jeszcze stwierdzić, które z naszych zachowań są poprawne a które nie. Ewolucja jest grą otwartą.

Millikan nie musi formułować żadnej nowej koncepcji nabywania i używania języka — zwykła teoria dyspozycjonalna jest wszystkim, czego potrzebuje, by wyjaśnić te fenomeny. Odwołanie się do teorii ewolucji pozwala natomiast na wyjaśnienie tego, iż pewne użycia — te, które nie doprowadzą mnie do sukcesu ewolucyjnego, rozumianego czysto biologicznie, jako przeżycie i przekazanie genów — mimo że zgodne z moją dyspozycją, są niepoprawne.

Rozwiązanie to jest narażone na szereg zarzutów. Martin Kusch (2006: 73) interpretuje stanowisko Millikan jako tezę, iż „normatywność związana z wyrażonymi intencjami może być ostatecznie zredukowana do normatywności niewyrażanych biologicznych celów” (Kusch 2006: 73). Wsuwa on jednak wątpliwość, czy mówienie o normatywności wynikającej z celów biologicznych jest czymś więcej niż tylko metaforą. Czy ma sens twierdzenie, że mamy jakkolwiek pojęty obowiązek postępowania tak, by zwiększyć nasze szanse ewolucyjne?

Zarzut ten byłby trafny, gdybyśmy uznali, iż w przypadku zdań należących do dyskursu semantycznego mamy do czynienia z normatywnością w sensie moc-

nym — to znaczy takim, w jakim normatywne są zdania z dyskursu etycznego. Jednak rozważania zawarte w rozdziale drugim doprowadziły nas do wniosku, iż w przypadku zdań o znaczeniu możemy mówić jedynie o normatywności w sensie słabym. W tej sytuacji, można argumentować, że jedyne czego nam potrzeba, aby dokonać analizy pojęcia znaczenia, to kryterium, które umożliwiłoby odróżnienie wypowiedzi poprawnych od niepoprawnych. Sukces ewolucyjny jest właśnie takim kryterium. To prawda, że nie musimy w żadnym sensie czuć się zmotywowani do postępowania tak, aby prowadziło to do sukcesu ewolucyjnego. Nawet gdybyśmy potrafili stwierdzić, które zachowanie jest korzystne z punktu widzenia przetrwania i przekazania genów, możemy zdecydować się na zgoła odmienne postępowanie. Jednak jeśli uznajemy, iż teza internalizmu motywacyjnego nie stosuje się do zdań o znaczeniu, to nie widać, dlaczego brak konieczności starania się o osiągnięcie sukcesu ewolucyjnego miałby stanowić problem. Możemy uznawać dane zdanie za poprawne — czyli, w myśl omawianej koncepcji, za prowadzące do sukcesu ewolucyjnego — i, świadomie bądź nie, sformułować je zupełnie inaczej.

Hattiangadi (2007: 130) z kolei atakuje rozwiązanie Millikan twierdząc, iż odwołanie się do celów ewolucyjnych nie wystarcza, aby wyjaśnić kwestię indywiduacji pojęć. Rozważa ona następujący przykład: ewolucyjnym sensem jedzenia cytryn jest dostarczenie organizmowi witaminy C, niezbędnej do prawidłowego funkcjonowania organizmu. Jednak witaminę C mogą również sobie dostarczyć poprzez spożycie pomarańczy. Czy oznacza to, że znaczenia słów „cytryna” i „pomarańcza” są takie same? Jeśli powiem komuś „podaj mi cytrynę”, a ten poda mi pomarańczę, z punktu widzenia moich biologiczno-ewolucyjnych potrzeb nie będzie to stanowiło różnicy. Czy zachowanie mojego interlokutora jest zatem poprawne?

Argument ten opiera się na założeniu, iż cele biologiczne nie różnicują pewnych zachowań. Jest to przesłanka o tyle kontrowersyjna, że nie wiadomo, jakie zachowania okażą się korzystne ewolucyjnie w dłuższej perspektywie czasowej. Być może okaże się kiedyś, iż umiejętność rozróżniania między cytrynami a pomarań-

czami jest istotna z punktu widzenia sukcesu ewolucyjnego; jedne z tych owoców mogą się na przykład okazać zgubne dla mojej zdolności rozmnażania. Dopiero w przyszłości okaże się, które z moich zachowań były korzystne z ewolucyjnego punktu widzenia.

Przeciwko koncepcji ewolucyjnej można jednak sformułować inny zarzut. Otóż teoria ta jest *de facto* jedynie pewną wersją koncepcji ujmującej kwestię normatywności znaczenia w kategoriach imperatywów hipotetycznych. Zewnętrznym celem, którego realizację ma umożliwiać stosowanie się do reguł semantycznych, jest, w myśl tego poglądu, sukces ewolucyjny. Podstawowa forma wyjaśniania normatywności znaczenia przez odwołanie do zewnętrznych celów jest tu jednak zachowana. W myśl koncepcji Millikan „nakaz” stosowania się do reguł semantycznych ma w rzeczywistości następującą formę: „Jeśli chcesz osiągnąć sukces ewolucyjny, to powinieneś stosować wyrażenia językowe w określony sposób”. Przyjmowane *implicitie* w tej teorii założenie, iż każdy organizm biologiczny w jakimś sensie „dąży” do sukcesu ewolucyjnego, uwiarygodnia ten schemat.

Do koncepcji Millikan stosuje się zatem dokładnie ten sam zarzut, co do wszystkich innych prób analizy znaczenia w kategoriach imperatywów hipotetycznych. Otóż jeśli okaże się, że stosowanie się do reguł semantycznych jest nieskuteczną metodą osiągania celu, którym jest sukces ewolucyjny, to całą analizę należy odrzucić. I, jakkolwiek nie można oczywiście stwierdzić, które dokładnie zachowania będą powodować sukces ewolucyjny, to można podać przykłady zachowań, które, jakkolwiek były zgodne ze standardami poprawnego użycia języka, to z pewnością nie doprowadziły do sukcesu ewolucyjnego.

Drastycznym przykładem może być sytuacja, jaka miała miejsce w Kambodży w czasie rządów Czerwonych Khmerów. W tym czasie „często «zbrodniarzy» zabijano za to, że byli piśmienni” (Polit 2001: 206). Osoba, która znalazła się w takiej dramatycznej sytuacji, a posługiwała się poprawnie językiem khmerskim, znacząco *zmniejszała* zatem swoje szanse na przeżycie. Nie sposób w takim

wypadku wysnuć wniosku, iż za każdym razem, gdy przejawiam dobre opanowanie reguł jakiegoś języka, zwiększam swoje szanse na sukces ewolucyjny. Uznanie, iż zachodzi ścisła zależność pomiędzy umiejętnością poprawnego używania języka a przetrwaniem i przekazaniem genów jest więc nieuzasadnione.

Można by tu odpowiedzieć, iż miarą sukcesu ewolucyjnego nie jest przekazanie genów należących do jednostki, ani pewnej grupy, lecz sukces gatunku jako takiego. Można jednak mieć zasadne wątpliwości, czy kierowanie się regułami poprawności językowej musi z konieczności doprowadzić do sukcesu ludzkości jako gatunku. Gdyby Chruszczow w czasie kryzysu kubańskiego wydał rozkaz do ataku nuklearnego, to im lepiej posługiwałby się językiem rosyjskim, tym większe byłoby prawdopodobieństwo wykonania jego rozkazu — zmniejszyłoby to jednak szanse ludzkości na przetrwanie.

Jedynym wyjściem dla zwolenników semantyki opartej na pojęciu sukcesu ewolucyjnego byłoby upieranie się, iż w takich sytuacjach naprawdę poprawnymi zachowaniami byłyby te, które uznalibyśmy potocznie za niepoprawne. Jeśli zatem, w danej sytuacji, odpowiedzenie „5” na pytanie o to, ile jest „ $67 + 58$ ” zwiększy moje szanse na przetrwanie, to należy uznać, iż jest to poprawna odpowiedź. Nie da się wykazać niespójności takiego stanowiska, lecz można argumentować, że jest to pozorne rozwiązanie problemu. Wyjściowe pytanie zadane przez Kripkensteina brzmiało bowiem: co konstytuuje fakt, iż pewne użycia języka można uznać za poprawne a inne nie? Zwolennicy teorii ewolucyjnej twierdzą, że odwołanie się do kryterium sukcesu ewolucyjnego pozwala nam wyróżnić pewne zachowania językowe i nazwać je „poprawnymi”, jakkolwiek z punktu widzenia zwykłego użytkownika języka te same zachowania mogłyby zostać zakwalifikowane odmiennie.

Ostatecznie koncepcja ta sprowadza się do trywialnego w gruncie rzeczy stwierdzenia, że pewne sposoby używania języka — nie wiadomo dokładnie które — przyczyniają się do sukcesu ewolucyjnego. Ta konstatacja nie stanowi odpowiedzi

na sceptyczne wyzwanie — nie ma bowiem gwarancji, że zachowania językowe ewolucyjnie korzystne będą tożsame z tymi, które intuicyjnie określilibyśmy mianem poprawnych. Oczywiście, można argumentować na rzecz tego, iż teoria ewolucji może dostarczyć nam odpowiedzi na pytania o genezę i funkcję języka. Jednak trudno spodziewać się, że teoria ewolucji zastąpi semantykę — nie da się określić znaczenia poszczególnych wyrażeń odwołując się do roli, jaką nasze użycia języka mogą odegrać w przetrwaniu i reprodukcji.

IV.4 *Qualia*

Wszystkie rozważane dotąd koncepcje mające stanowić odpowiedź na pytanie Kripkensteina przyjmowały ontologię w szerokim sensie fizykalistyczną. Innymi słowy, fakty, które miałyby według nich odpowiadać zdaniom o znaczeniu, mają w myśl tych koncepcji mieścić się w ramach wyznaczonych przez współczesną naukę. Być może to właśnie to założenie uniemożliwia znalezienie rozwiązania problemu Kripkensteina i wyjście poza dziedzinę faktów naturalnych umożliwiłoby nam rozwikłanie zagadki.

We współczesnej filozofii analitycznej pojawiło się wiele argumentów na rzecz tezy, iż to właśnie *qualia*, czyli jakościowe przeżycia psychiczne, są tym, czego nie da się ująć w ramach ontologii fizykalistycznej. Do typowych przykładów *qualiów* można zaliczyć przeżycia percepcyjne, takie jak widok czegoś zielonego, doznania cielesne, jak ból czy swędzenie, powidoki, wyobrażenia czy halucynacje. Wszystkie te stany psychiczne charakteryzują się pewną subiektywną jakością, tym, że jeśli znajdujemy się w danym stanie, możliwe jest pytanie: jak to jest się w nim znajdować?

David Chalmers (2002/2008: 446–449) rekonstruuje trzy główne argumenty przeciwko tezie o redukowalności fenomenalnych stanów psychicznych do faktów fizycznych. Pierwszy z nich to argument z wyjaśnienia, wskazujący na fakt, iż ujęcia fizyczne wyjaśniają jedynie struktury i funkcje, a to nie wystarcza do wytłu-

maczenia fenomenu świadomości. Drugi to argument z pojmovalności. Jego struktura jest następująca: przyjmuje się najpierw, że jest możliwe do pomyślenia istnienie istoty identycznej z człowiekiem pod względem budowy fizycznej oraz zachowania, lecz nieposiadającej żadnych przeżyć świadomych. Istotę taką określa się w literaturze mianem *zombi*. Z faktu pojmovalności *zombi* wnosi się następnie o metafizycznej możliwości ich istnienia. Dowodzi to fałszywości fizykalizmu, ponieważ okazuje się, iż istoty tak samo zbudowane pod względem fizycznym mogą się różnić co do stanów psychicznych. Trzeci argument odnosi się do wiedzy. Podnosi się w nim, że z informacji o stanach fizycznych nie można wyprowadzić wniosków na temat świadomości. Kanoniczny przykład na poparcie tego rozumowania pochodzi od Franka Jacksona (1982 : 130). Opisuje on przypadek Marii, genialnej neurobiolożki, która wie wszystko na temat fizycznych i biologicznych procesów będących podstawą widzenia. Jednak Maria nigdy nie widziała koloru czerwonego — całe życie spędziła w czarno-białym pokoju (w alternatywnych wersjach historii cierpi na achromatyzm lub jest niewidoma). Dopiero gdy opuściła swój pokój (resp. ozdrowiała) i zobaczyła coś czerwonego, dowiedziała się, jak to jest widzieć na czerwono. Wiedzy tej nie mogła była uzyskać na podstawie swojej znajomości odpowiednich faktów fizycznych. Możemy zatem wysnuć wniosek o istnieniu luki epistemologicznej pomiędzy faktami fizycznymi a fenomenalnymi, z czego z kolei wyprowadza się wniosek o istnieniu luki ontologicznej (Chalmers 2002/2008: 450).

Dysponujemy dobrymi argumentami na rzecz tezy, iż fakty fenomenalne nie mogą zostać zredukowane do faktów fizycznych. Co więcej, część zwolenników nieredukcyjnego ujęcia faktów fenomenalnych twierdzi, iż mają one charakter epifenomenów — nie wpływają zatem przyczynowo na zachowania (Jackson 1982: 130). Przyjęcie tezy o przyczynowej nierelewanencji fenomenalnych własności mentalnych kłóci się jednak z przywoływanym w rozdziale drugim założeniem przyjmowanym przez Kima i Shoemakera, że realne własności muszą mieć moc przyczynową. Mimo

to, istnienie qualiów wydaje się czymś niepodważalnym — jesteśmy bezpośrednio świadomi tego, że doznajemy bólu, że widzimy coś zielonego itp.

Teza o istnieniu epifenomenalnych qualiów stanowi wyłom w fizykalistycznym obrazie świata. Czy jednak rozszerzenie naszej ontologii o tego typu fakty pozwala uporać się z problemem postawionym przez Kripkensteina?

Kłopoty z ujmowaniem znaczenia jako pochodnego względem introspekcyjnie dostępnych jakości były już omawiane w *Rozdziale I* niniejszej pracy — Kripke analizował to podejście jako jedno z rozwiązań problemu sceptycznego. Wskazywał on na dwa problemy związane z tym ujęciem: pierwszy dotyczył kwestii metody projekcji łączącej domniemany mentalny obraz z przedmiotem, do którego miałbym się odnosić. Drugim, bardziej ogólnym argumentem było wskazanie na to, iż trudno zrozumieć, w jaki sposób posiadanie pewnego szczególnego *quale* może wyjaśniać normatywność swoistą dla zdań o znaczeniu.

Warto jednak rozważyć tę kwestię głębiej, tym bardziej, że we współczesnej filozofii na powrót zyskała uznanie idea, iż posiadanie danego nastawienia sądzeniowego wiąże się z pewnym jakościowym przeżyciem (por. np. Pitt 2004: 1). Zdaniem autorów głoszących ten pogląd, jeśli chcemy zjawisko bezpośredniej świadomości treści naszych intencjonalnych stanów mentalnych wyjaśnić, to musimy przyjąć, że stany te cechują się pewną introspekcyjnie dostępną jakością. Rozumowanie to może budzić wątpliwości, gdyż zakłada się w nim, iż bycie świadomym jest możliwe tylko wtedy, jeśli rzecz, której jesteśmy świadomi, posiada pewną fenomenalną własność (por. np. Lormand 1996: 245–246).

Pitt (2004: 27–29) na poparcie swojej tezy o fenomenalnym charakterze nastawień sądzeniowych przywołuje przykład czytania. Według niego doświadczenia jakie mamy gdy czytamy zdanie, które jest dla nas niezrozumiałe ze względu na jego eliptyczność, w oczywisty sposób różnią się od analogicznych doświadczeń w przypadku zdania, które jest dla nas składniowo przejrzyste. Podobnie, inne odczucie mamy mieć w sytuacji, gdy czytamy zdanie, którego nie rozumiemy

ze względu na użycie w nim niezrozumiałych słów, niż w przypadku zdania przekazującego podobną myśl, lecz wyrażającego ją prostszymi, bardziej zrozumiałymi słowami.

Kwestię fenomenalnego poczucia obecnego przy czytaniu poruszał wcześniej Wittgenstein (2000 § 165–171), do którego Pitt, co dość znamienne, się nie odnosi. Autor *Dociekań filozoficznych* zdaje się przeczyć przeświadczeniu Pitta:

A teraz przeczytaj kilka zdań druku, tak jak zwykle to robisz, gdy nie myślisz o pojęciu czytania; i zastanów się, czy miałeś podczas czytania takie przeżycia jedności, wpływu itd. - Nie mów, jakobyś miał je nieświadomie! Nie dajmy się też zwieść idei, że zjawiska te pojawiają się przy 'bliższym wejrzeniu'! (Wittgenstein 1953/2000 § 171).

Spór zdaje się niemożliwy do rozstrzygnięcia. Mamy bowiem do czynienia z wzajemnie się wykluczającymi sprawozdaniami dotyczącymi wewnętrznych stanów podmiotów. Trudno polemizować ze szczerym wyznaniem Pitta, iż w momencie, gdy ma jakieś nastawienie sądzeniowe lub gdy czyta coś ze zrozumieniem, doświadcza pewnej szczególnej jakości. Niemal z definicji to właśnie podmiot ma uprzywilejowany dostęp do własnych stanów psychicznych o charakterze jakościowym. Z tego samego względu nie ma żadnych podstaw, by przeczyć szczerym wyznaniom Lormanda czy Wittgensteina, gdy mówią, iż oni żadnego takie doświadczenia nie mają. Nie wiadomo, jakie argumenty można by przytaczać w tego typu sporze, co może być (kolejnym) przyczynkiem do stwierdzenia, iż prowadzenie obserwacji własnych stanów psychicznych nie jest właściwą metodą uprawiania filozofii.

Przyjmijmy jednak, że zwolennicy tezy o fenomenalnym charakterze nastawień sądzeniowych mają rację i rzeczywiście, zawsze gdy mamy jakieś przekonanie, doświadczamy pewnej swoistej jakości. Czy takie założenie pomaga jednak w znalezieniu rozwiązania paradoksu sceptycznego?

Odpowiedź na to pytanie jest negatywna. Problem sceptycyzmu semantycznego jest — jak to było w niniejszej książce wielokrotnie podkreślanie — problemem ontologicznym a nie epistemologicznym. Gdyby nasze pytanie brzmiało: skąd wiem, że aktualnie mam takie a nie inne przekonanie, można by odpowiedzieć, iż wiem to, ponieważ odczuwam pewną fenomenalną jakość. Pytanie jednak dotyczy tego, co konstytuuje fakt posiadania przeze mnie takiego a nie innego nastawienia sądzeniowego.

Teorię, według której znaczenie wyjaśniałoby się przez odwołanie do qualiów, można by próbować rekonstruować w następujący sposób: posiadanie pewnego odczucia wyznacza standard poprawności użycia danego wyrażenia. Jeśli mojemu użyciu pewnego słowa towarzyszy pewne swoiste *quale*, to jest ono poprawne, jeśli zaś mojemu użyciu nie towarzyszy rzeczzone doznanie, to moja wypowiedź jest błędna. Epifenomenalny charakter qualiów pozwala na wyjaśnienie faktu, iż znaczenie nie ma charakteru wewnętrznie motywującego — wystąpienie pewnej jakości nie wpływa przyczynowo na żadne zachowania, w tym i językowe. Może się zdarzyć, że nawet jeśli dane *quale* nie pojawi się w mojej świadomości, to i tak będę używał słowa z nim „związanego” i *vice versa*.

Taki pomysł na wyjaśnienie fenomenu znaczenia jest jednak całkowicie chybiony. Jak to zostało pokazane w części poświęconej przyczynowym teoriom odniesienia, często posługujemy się językiem po to, aby mówić o rzeczach znajdujących się w rzeczywistości zewnętrznej wobec naszej świadomości. Chcemy mówić o tygrysach, rzekach i innych ludziach, a nie tylko o subiektywnych odczuciach i doznaniach. Warto zauważyć, iż żaden z teoretyków głoszących tezę o istnieniu fenomenalnych aspektów towarzyszących nastawieniom sądzeniowym nie postuluje przyjęcia idealizmu subiektywnego, czyli przeświadczenia (przypisywanego zazwyczaj Berkeley’owi), iż doświadczamy jedynie subiektywnych wrażeń.

W *Rozdziale II*, przy okazji omawiania rozróżnienia na poprawne użycie i poprawne zastosowanie, został wysnuty wniosek, że istnieją konteksty, w których

poprawnym sposobem użycia danego wyrażenia będzie jego poprawne zastosowanie. W niektórych sytuacjach poprawność językowa sprowadza się do tego, by nazwać złoto „złotem”. Czy tak rozumianą poprawność można wyjaśnić odwołując się do mojego wewnętrznego przeżycia? Nic przecież nie stoi na przeszkodzie, abym wypowiadając słowo „złoto”, gdy przede mną znajduje się próbka tombaku, miał takie samo wewnętrzne doznanie, jak wówczas, gdy wypowiadam to słowo w obecności prawdziwego złota.

Nic nie uzasadnia przekonania, że moje jakościowe przeżycia będą przejawiały regularności zgodne ze zmianami w otoczeniu fizycznym lub kontekście społecznym — a naturalne jest przypuszczenie, iż to właśnie od tych zmian zależeć będzie poprawność wypowiedzi. Nie da się zatem uznać, że subiektywne przeżycia mogą być standardem dla poprawnego użycia wyrażań. Aby obronić tę tezę należałoby całkowicie zmienić nasze pojęcie znaczenia, tak by nie obejmowało ono zależności od zewnętrznego kontekstu.

Kompromisową pozycję w sporze o rolę qualiów w eksplikacji znaczenia zajmuje Galen Strawson. Jest on zwolennikiem tezy, iż posiadanie stanu intencjonalnego wiąże się z pewnym swoistym odczuciem. Z drugiej jednak strony twierdzi, że to odczucie nie może być utożsamiane z treścią pojęcia, którym się posługujemy (Strawson 2005: 51–52). Zdaniem tego autora, założenie, że każdemu nastawieniu sądzeniowemu odpowiada swoista jakość nie wiąże się w sposób konieczny z odrzuceniem eksternalizmu semantycznego, to jest przekonania, iż treść pojęć jest zależna od okoliczności zewnętrznych. Strawson uważa, że odrzucenie internalizmu semantycznego oraz obrazkowej teorii znaczenia było słuszne — nie można utożsamiać znaczenia z wewnętrznym odczuciem. Przy okazji doszło jednak do „wylania dziecka z kąpielą” — niesłusznie uznano, iż należy zaprzestać podejmowania jakichkolwiek badań o charakterze fenomenologicznym, a dotyczących nastawień sądzeniowych właśnie.

Cokolwiek by nie sądzić o sensowności prób powrotu do fenomenologii nastawień sądzeniowych, wydaje się, że te badania i ich ewentualne rezultaty nie mają

najmniejszego wpływu na pytanie o to, czy zdania należące do dyskursu semantycznego mają wartość logiczną. Takie ujęcie relacji między treścią pojęć czy znaczeniem słów a subiektywnymi odczuciami, w którym te pierwsze miałyby redukować się do tych drugich, jest niemożliwe do utrzymania.

IV.5 Platonizm i nonredukcjonizm semantyczny

Najbardziej radykalną formą odrzucenia ontologicznego fizykalizmu jest przekonanie iż istnieją byty lub własności, jakich istnienie postuluje się w ramach stanowiska zwanego potocznie platonizmem. Zwięzłej charakterystyki tego stanowiska dokonał Chrudzimski, wedle którego:

Platonik twierdzi (...), iż istnieje odrębny świat bytów ogólnych, zaś posiadanie przez indywiduum danej cechy (lub stanie indywiduów w określonej relacji) polega na stosunku egzemplifikacji, który zachodzi między tymi dwoma światami. (...) Jeśli zatem umysł, odnosząc się intencjonalnie, ujmuje zawsze pewną abstrakcyjną cechę, wówczas w ramach teorii platońskiej bardzo kuszące jest przypisanie mu zdolności bezpośredniego docierania do platońskiego nieba. Nieekstensjonalna pseudo-relacja intencjonalna, w jakiej podmiot stoi do indywiduum, sprowadza się w myśl tej teorii do autentycznej, ekstensjonalnej relacji uchwytywania przez ten podmiot cechy, która z kolei egzemplifikowana jest przez owo indywiduum (Chrudzimski 2003: 106 - 107).

Nie jest tu istotne, czy historyczny Platon był zwolennikiem tak scharakteryzowanego platonizmu. Nie ma bowiem jasności, które z tez zawartych w dialogach Platona możemy traktować dosłownie, a które jako metafory. Ważne jest to, czy odwołanie się do takiego stanowiska, niezależnie od jego historycznych korzeni, pozwoli na rozstrzygnięcie problemu istnienia własności i faktów semantycznych.

Platonik na pytanie o standard poprawności zachowań językowych odpowiedziałby, iż jeśli używam jakiegoś słowa będącego nazwą własności, to używam tego słowa poprawnie tylko wówczas, jeśli rzecz, do której się odnoszę, rzeczywiście partycypuje w idei, która stanowi pierwotne odniesienie słowa.

Tak rozumiany platonizm pozwala na utrzymanie tezy, która stanowiła o atrakcyjności eksternalizmu semantycznego, mówiącej, że fakty determinujące znaczenie muszą być w jakiś sposób zewnętrzne wobec świadomości użytkowników języka. Z drugiej strony, niefizykalny charakter bytów platońskich pozwala wyjaśnić, dlaczego znaczenie nie jest przyczynowo relewantne — idee nie wchodzą bowiem w relacje przyczynowe z elementami świata fizycznego. Nie można przy wyjaśnianiu zachowania odwoływać się do czegoś, co nie mieści się w granicach świata zmysłów. Taki wniosek wprost wypływa z założenia, przyjętego już w *Rozdziale III*, iż sfera fizyczna jest przyczynowo domknięta — to znaczy, że do wyjaśnienia dowolnego zdarzenia fizycznego wystarcza odwołanie się do innych zdarzeń fizycznych. A jako że zachowania werbalne użytkowników języka są niewątpliwie faktami fizycznymi, nie mogą one być determinowane przyczynowo przez platońską ideę.

Jednak wskazanie na idee platońskie jako na rozwiązanie problemu sceptycznego samo generuje kolejne pytania. Najważniejsze z nich dotyczy tego, co łączy mnie, moje stany psychiczne oraz dyspozycje do zachowań z tak rozumianym powszechnikiem. Podobny problem pojawił się w przypadku teorii terminów naturalnorodzajowych — trudno było wyjaśnić, co konstytuuje relację zachodzącą pomiędzy użytkownikami języka a rodzajami naturalnymi. W tamtym przypadku możliwa jednak była odpowiedź polegająca na wskazaniu na relacje przyczynową zachodzącą pomiędzy danym egzemplarzem rodzaju naturalnego a dyspozycją, jaka miałaby się wytwarzać w momencie chrztu. Jakkolwiek rozwiązanie to ostatecznie okazało się niesatysfakcjonujące, zwolennik owej teorii dysponował przynajmniej pewnym obrazem tego, jak zachowania użytkownika są powiązane z faktami określającymi ich poprawność.

W przypadku realizmu w stylu platońskim takiej odpowiedzi nie widać. Nie wiadomo, jakiego typu relacja miałyby zachodzić pomiędzy ideami a użytkownikami języka. Zazwyczaj platonizmowi zarzuca się, że nie wyjaśnia on relacji partycypacji, zachodzącej pomiędzy ideą a podpadającym pod nią przedmiotem. W tym miejscu okazuje się ponadto, iż platonizm postuluje istnienie jeszcze jednej, trudnej do wyjaśnienia relacji — między podmiotem a uchwytywaną przez niego ideą.

Platonik jest zmuszony stwierdzić, iż relacja uchwytywania zachodząca pomiędzy użytkownikiem języka a ideą jest czymś pierwotnym i nieredukowalnym do niczego innego. W ten sposób jednak platonizm, rozumiany jako stanowisko w sporze o uniwersalia, staje się pewną mutacją tezy o pierwotności znaczenia czy też antyredukcjonizmu semantycznego. Główną cechą takiego podejścia do problemu znaczenia jest przekonanie, iż własności i relacje semantyczne są czymś swoistym i nieredukowalnym. W myśl tego stanowiska, zdania należące do dyskursu semantycznego odnoszą się do faktów i mają określoną wartość logiczną, lecz nie można dokonać redukcji tych faktów do żadnych innych.

Antyredukcjonizm semantyczny w najprostszej formie polega na przyjęciu, iż istnieją pewne szczególne byty — „znaczenia”, których „uchwycenie” konstytuowałoby rozumienie poszczególnych wyrażeń. Owe znaczenia miałyby status ontologiczny zbliżony do platońskich idei. David Finkelstein tak charakteryzuje to stanowisko:

Można myśleć o platonizmie jako o desperackiej próbie zablokowania regresu interpretacji, który odpowiada za powstanie paradoksu (...). Platonik postuluje istnienie szczególnych przedmiotów, które — w odróżnieniu od dźwięków, napisów czy gestów — są, w pewnym sensie, wewnętrznie znaczące. Według platonika, nasze słowa przed pustką ratuje fakt, iż za nimi stoją takie przedmioty. Problem regresu nie powstaje, gdyż te wewnętrznie znaczące przedmioty ani nie muszą, ani nie mogą być interpretowane (Finkelstein 2003: 55).

W podobnym duchu tę propozycję opisał Wittgenstein, który tak scharakteryzował stanowisko swojego adwersarza:

Każdy znak można poddać interpretacji; ale *znaczenia* nie powinno się już dać interpretować. To jest ostateczna interpretacja (Wittgenstein 1958/1998: 67).

Postulowanie znaczeń jako osobnych bytów prowadzi do dokładnie tych samych problemów, co „zwykły” platonizm. Niezbędne jest bowiem wyjaśnienie dwóch relacji. Po pierwsze — tej, w jakiej pozostają owe specjalne byty do poszczególnych przypadków użycia wyrażenia, które ma owo znaczenie, czy też do przedmiotów podpadających pod dane pojęcie. Po drugie, wyjaśnienia domaga się relacja pomiędzy użytkownikami języka a znaczeniami. Frege (1882/1977) posługuje się w tym kontekście metaforą „uchwytywania“, lecz nie wyjaśnia dokładnie, na czym relacja „uchwytywania“ miałaby polegać.

Antyredukcjonizm w kwestii natury faktów semantycznych nie musi wszakże pociągać za sobą uznania znaczeń za byty platońskie — wystarczy przyjąć, że mamy do czynienia z dalej niewyjaśnianymi własnościami czy relacjami. Takie rozwiązanie paradoksu Kripkensteina proponuje m.in. Boghossian (1989: 547 - 549), który twierdzi, iż sądy na temat znaczenia dotyczą faktów, są nieredukowalne i niezależne od opinii użytkowników języka. Według niego taki nieredukcyjny realizm co do znaczenia jest jedyną opcją, jaka pozostaje wobec niemożności podania redukcyjnych warunków prawdziwości dla zdań o znaczeniu.

Teza o pierwotności znaczenia nie jest, co warto podkreślić, w sposób istotny związana ze skrajnym realizmem w kwestii uniwersaliów (choć skrajny realizm w sporze o uniwersalia wymaga, jak się wydaje, przyjęcia antyredukcjonizmu semantycznego). Zwolennik przekonania o nieredukowalnym charakterze relacji semantycznych może przyjmować, że pierwotny charakter ma na przykład relacja oznaczania, zachodząca pomiędzy stanem mentalnym użytkownika a przedmiotem. Taką koncepcję głosił chociażby, uważany za wzorcowego nominalistę, Ockham,

wedle którego „pojęcie, czyli cecha umysłu w sposób naturalny oznacza wszystko to, co oznacza, natomiast termin wypowiedziany lub napisany oznacza coś tylko na mocy dowolnego ustanowienia” (Ockham 1971: 13).

Wydaje się, iż tak scharakteryzowany platonizm unika zarówno problemu regresu interpretacji, jak i problemu normatywności. Jest bowiem jasne, iż jeśli relacje semantyczne po prostu zachodzą, to wyznaczają one, niejako automatycznie, standardy poprawnego użycia. Jeśli rozpatrzymy prosty przypadek, w którym poprawność językową można utożsamić z poprawnym zastosowaniem, to istnienie pierwotnej relacji oznaczania, łączącej symbol z przedmiotami odniesienia, to można z łatwością stwierdzić, iż poprawnymi będą tylko te zastosowania symbolu, w których stosujemy go do tych rzeczy, z którymi jest połączony ową pierwotną relacją.

Platonizm płaci za poradzenie sobie z tymi problemami wysoką cenę. Finkelstein (2003: 55) zauważa, że w tej wizji komunikacja jest zjawiskiem całkowicie tajemniczym. Wydaje się bowiem, iż zmuszeni jesteśmy zgadywać, jaka idea czy też pierwotna relacja semantyczna konstytuuje znaczenie słowa, którego używa nasz interlokutor. Wiara w to, iż możliwy jest w takiej sytuacji sukces komunikacyjny, ma dość kruche podstawy.

Kripke rozważał rozwiązanie podobne do platonizmu — mianowicie przekonanie, iż używanie symbolu w danym znaczeniu jest stanem *sui generis*, jedynym w swoim rodzaju, niepodobnym do żadnego innego (Kripke 1982/2007: 86). Odnosi się do tego pomysłu w następujący sposób:

W pewnym sensie argument taki jest nie do odparcia, a jeśli odpowiednio by się go przedstawiło, to Wittgenstein mógłby go chyba przyjąć. Wydaje się on jednak desperacki: pozostawia naturę owego postulowanego pierwotnego stanu (...) całkowicie owianą tajemnicą. Ma to nie być stan introspekcyjny, a jednak rzekomo jesteśmy go z pewnością dozą pewności świadomi za każdym razem, gdy zachodzi. (Kripke 1982/2007: 86).

Mamy tutaj do czynienia z dwoma problemami. Pierwszy dotyczy konieczności wyjaśnienia faktu wiedzy na temat tego, co znaczą słowa w używanym przeze mnie języku. Jeśli znaczenie jest determinowane przez pewnego rodzaju szczególne byty lub relacje, to niezbędne staje się postulowanie specjalnej władzy poznawczej, która odpowiadałaby za wiedzę na temat znaczenia.

Po drugie — co podkreśla m.in. Finkelstein (2003: 56) — tego typu odpowiedź wydaje się raczej aktem teoretycznej bezradności, niż rzetelnym rozwiązaniem. Postuluje się tu bowiem istnienie pewnych bytów, których jedyną znaną nam cechą, jest to, że rozwiązują problem, z jakim się borykamy.

Warto w tym miejscu przyrzeć się szczególnej formie antyredukcjonizmu, czy też platonizmu w sensie szerszym, jakim była teoria semiozy Peirce'a. Jest ona warta rozważania, gdyż w jej przypadku widać wyraźnie, że bardziej istotne dla antyredukcjonizmu w kwestii znaczenia jest postulowanie pewnych swoistych relacji niż bytów.

Wedle amerykańskiego filozofa znak jest relacją triadyczną. Pierwszy jej człon określa on jako podstawę znaku, co można w uproszczeniu rozumieć jako fizyczny byt, dźwięk lub napis. Kolejnym elementem jest interpretant, który nie powinien być utożsamiany ze zjawiskiem psychicznym (por. Nowak 2002: 158 - 159); ostatnim zaś przedmiot. Swoistą cechą semantyki Peirce'a jest przekonanie, iż również poszczególne elementy tej relacji są znakami (Peirce 1931: 339). Przekonanie, iż interpretant znaku jest czymś, co samo podlega dalszej interpretacji, wydaje się dość zrozumiałe; taka teza stanowi wyraz dość intuicyjnego sceptycyzmu wobec twierdzenia, iż możliwa jest ostateczna interpretacja.

Poważniejsze zastrzeżenia budzi twierdzenia, iż znakowy charakter mają — a co za tym idzie podlegają dalszej interpretacji — zarówno podstawa znaku, jak i jego przedmiot. Ta ostatnia, być może najbardziej trudna do zaakceptowania teza, wynika z Peirce'owskiego przeświadczenia, iż przedmiot również musi być środkiem przekazu myśli (Peirce 1931: 538). Jednak nawet jeśli zgodzilibyśmy się

z Peirce'm jedynie co do tego, iż to interpretant znaku podlega dalszej interpretacji, to i tak jego koncepcja prowadzi do regresu w nieskończoność, albowiem znak jest tu definiowany przez coś, co samo w sobie ma znakowy charakter (por. np. Nowak 2002: 45–47, 159).

Peirce odrzuca zatem jedno z założeń charakterystycznych dla platonizmu w rozumieniu Finkelsteina: nie istnieją według niego żadne samointerpretujące się byty czy relacje, wszystko co ma charakter znakowy podlega dalszej interpretacji. Można powiedzieć, iż Peirce neutralizuje argument z regresu, zabójczy dla wszystkich prób rozwiązania paradoksu Kripkensteina wskazujących na fakty, które same w sobie mają charakter znaczący. Peirce twierdzi po prostu, iż fakty semantyczne ze swej natury generują regres w nieskończoność i nie można tego traktować jako zarzutu.

Jednak z innej nieco perspektywy zakwalifikowanie Peirce'a jako platonika — w sensie szerszym — wydaje się zasadne. Przyjmuje on bowiem, iż istnienie relacji semantycznych jest czymś absolutnie swoistym, czymś, czego nie da się wyjaśnić przez odwołanie się do faktów innego rodzaju. To, iż między podstawą znaku, interpretantem a przedmiotem zachodzą takie a nie inne relacje trzeba uznać za fakt pierwotny. Podobnie w ontologii platońskiej za pierwotny i niewyjaśniany fakt należało uznać zachodzenie relacji łączącej ideę z partycypującymi w niej przedmiotami oraz relację uchwytowania, łączącą podmiot z ideą.

Platonizm w sensie szerszym, czyli antyredukcjonizm — czy to w wersji zakładającej samointerpretujące się znaczenia, czy to regres w nieskończoność — prezentuje się jako skuteczne, acz paradoksalne rozwiązanie problemu. Sceptyk stawia przed nami pytanie o to, co konstytuuje fakty semantyczne i wskazuje na różnego rodzaju trudności związane z odpowiedzią na nie — choćby na to, iż znaczenie generuje problem nieskończonego regresu interpretacji oraz jest normatywne. Platonik odpowiada na to, iż fakty semantyczne zachodzą po prostu. To zaś, że mają one właśnie taki charakter, jaki dla sceptyka był problematyczny, jest czę-

ścią ich natury. Rozstrzygnięcie sporu między tymi stanowiskami wymagać będzie zatem pewnej refleksji nad metodologicznymi i metafizycznymi założeniami pytania o faktualizm, jak również przyjrzenia się temu, czy oba te stanowiska — antyredukcjonizm i nonfaktualizm — da się w spójny sposób obronić. Warto również będzie zadać pytanie, czy rzeczywiście wyczerpują one zakres możliwych odpowiedzi.

Rozdział V

Nonfaktualizm i minimalizm

V.1 Co przemawia za nonfaktualizmem?

Jakie wnioski można wyciągnąć z rozważań zawartych w powyższych rozdziałach? Czy dysponujemy konkluzywną metodą wykazania, że fakty semantyczne nie istnieją, czy też musimy zgodzić się z przywoływaną w *Rozdziale II* opinią Hattiangadi (2007: 210), że zwolennik nonfaktualizmu w tej kwestii ma do swojej dyspozycji jedynie argument z niewiedzy? Czyż cała argumentacja sceptyczna nie dowodzi jedynie, iż nie potrafimy odnaleźć faktów, które konstytuują znaczenie?

Argumentacja przedstawiona w poprzednich rozdziałach jest jednak mocniejsza. Jej celem było pokazanie, iż przedstawione propozycje wyczerpują zakres możliwych odpowiedzi na pytanie o ontologiczny status faktów znaczeniowych. Rozumowanie to można przedstawić za pomocą serii następujących dylematów:

1. Czy domniemane fakty semantyczne miałyby być relewantne przyczynowo?

Odpowiedź „Tak” na to pytanie prowadzi nas do dyspozycjonalnej teorii znaczenia, która została odrzucona. Odpowiedź „Nie” prowadzi do dalszych pytań.

2. Czy fakty semantyczne miałyby mieć charakter fizyczny?

Odpowiedź „Tak” wymusza przyjęcie koncepcji faktów, które miałyby mieć charakter fizyczny, a zarazem nie być relewantne przyczynowo. Wydaje się, że możliwe są tu tylko dwie opcje: albo fakty te mają charakter nie-lokalny, albo teleologiczny. Obie te propozycje zostały poddane krytyce — zarówno eksternalizm, jak i odwołanie się do teorii ewolucji nie pozwalają na skuteczną obronę pojęcia znaczenia.

3. Jeśli fakty odpowiadające zdaniom o znaczeniu miałyby mieć charakter niefizyczny, to czy mają one charakter semantyczny lub intencjonalny?

Jeśli odpowiedź brzmi: „Nie”, to wydaje się, że jedynymi faktami, które mają zarazem charakter niefizyczny, jak i niesemantyczny — czy też ogólnie, nieintencjonalny — są jakościowe przeżycia psychiczne, czyli *qualia*. Jak jednak zostało już pokazane, nie da się zredukować semantyki do *qualiów*.

Mamy tu więc do czynienia nie tyle z argumentem z niewiedzy, ile z wyczerpaniem opcji — nasza ontologia nie zawiera innych faktów, niż te wyżej wymienione, a na które można by się powołać celem odpowiedzi na pytanie sceptyka. Jedynym wyjściem dla zwolennika faktualizmu pozostaje zatem postulowanie swoistych faktów i własności/relacji semantycznych.

Propozycja nonredukcjonisty na pierwszy rzut oka wydaje się trudna do odparcia. Nie wiadomo bowiem, na jakiej podstawie można by zaprzeczyć istnieniu pewnych swoistych faktów, które są dostępne za pomocą specjalnej władzy poznawczej.

Można tu odpowiedzieć, że problem nonfaktualisty z refutacją platonizmu jest tylko szczególną postacią bardziej ogólnego zagadnienia: jak można wykazać, iż jakiegoś rodzaju fakty nie istnieją? Na pozór nie ma na to pytanie dobrej odpowiedzi. Rorty w przywoływanym już w *Rozdziale I* eseju o eliminatywizmie zwraca uwagę na problem, jaki ma zwolennik ubogiej ontologii w konfrontacji ze zwolennikiem ontologii bogatej. Jeśli mielibyśmy na przykład do czynienia z szamanem, upierającym się przy tezie o istnieniu demonów, to nie byłoby

nam łatwo przekonać go, iż tego typu bytów nie ma (Rorty 1965/1995: 220–221). Nawet gdybyśmy wykazali, iż wszystko, co chcemy wyjaśnić, odwołując się do istnienia demonów, można wyjaśnić prościej, to i tak nie przeszkadzałoby mu to w twierdzeniu, iż demony istnieją. Gdybyśmy np. wskazali, iż za występowanie chorób odpowiadają nie demony, lecz bakterie i wirusy, to mógłby on twierdzić, że bakterie i wirusy po prostu współtowarzyszą demonom.

Riposta Rorty’ego jest prosta — jego zdaniem, fakt, że dysponujemy lepszym wyjaśnieniem niż to, które oferuje nam zwolennik rozszerzonej ontologii, pozwala nam uznać, że bytów postulowanych w tej ontologii po prostu nie ma. Wniosek taki wynika wprost z zasady oszczędności. W omawianym przez Rorty’ego przykładzie, na wszelkie argumenty szamana mamy prawo odpowiedzieć, że jeśli demony są zbędne, to ich nie ma.

Omawiany w *Rozdziale II* argument Kima z przyczynowego wykluczenia, jest, jak to już było wspomniane, przeniesieniem zasady oszczędności z problemu istnienia przedmiotów na problem istnienia własności. Jeśli bez przyjmowania pewnego typu własności możemy przyczynowo wyjaśnić zachodzenie pewnych zjawisk, to mamy prawo stwierdzić, iż tych własności po prostu nie ma.

W omawianym tu przypadku, zwolennik nonfaktualizmu może dowodzić, że wszystkie fenomeny związane z użyciem języka da się wyjaśnić bez postulowania faktów odpowiadających zdaniom o znaczeniu oraz własności semantycznych. Do tego, by przyczynowo wyjaśnić werbalne zachowania użytkowników języka, wystarczy odwołać się do dyspozycji i warunkującej owe dyspozycje architektury neurobiologicznej.

Zastosowanie argumentu z wykluczenia zatem przenosi ciężar dowodu na faktualistę. Musi on wskazać na jakiś powód do tego, aby w ogóle fakty semantyczne w naszej ontologii przyjmować. Faktualista nie może po prostu powołać się na fakt, iż używamy w języku zdań oznajmujących mówiących o znaczeniach. Wstępnym założeniem sporu faktualisty z nonfaktualizmem (w jakiegokolwiek kwestii) jest

bowiem to, że nie każdy dyskurs zawierający poprawne językowo zdania oznajmujące będzie dyskursem faktualnym. Innymi słowy, zakłada tu się, że możliwe jest odróżnienie zdań, które mówią o faktach, od tych, których rola w języku polega na czymś innym.

Najważniejszą racją, by włączać fakty semantyczne do naszej ontologii jest według wielu zwolenników faktualizmu semantycznego to, iż przyjęcie stanowiska przeciwnego w omawianej kwestii prowadzi do sprzeczności. Za tezę o niespójności postawy nonfaktualistycznej w odniesieniu do znaczenia przemawiają dwa argumenty. Pierwszy z nich to przedstawiony przez Boghossiana argument z dwuznaczności pojęcia prawdy; drugi to sformułowany przez Wrighta i doprecyzowany przez Kuschę argument z uogólnienia.

V.2 Niespójność nonfaktualizmu I: Argument z dwuznaczności

Boghossian dowodzi, że nonfaktualizmu w sprawie znaczenia nie da się wyrazić w sposób spójny, gdyż nonfaktualista posługuje się jednocześnie dwoma pojęciami prawdy: substancjalnym, wedle którego prawda jest rozmaicie rozumianą „realną własnością”, oraz deflacyjnym, zakładającym, iż nie ma takiej własności jak prawdziwość.

Sednem deflacionizmu (por. np. Soames 1997: 2) jest przekonanie, iż sens pojęcia „prawdziwości” sprowadza się do tak zwanych T-równoważności. Pojęcie to zostało wprowadzone przez Tarskiego i oznacza zdania następującej formie:

(T) *s* jest zdaniem prawdziwym wtedy i tylko wtedy gdy *p*.

Gdzie *p* jest dowolnym zdaniem zaś *s* nazwą tego zdania (Tarski 1933/1995: 18). Najpopularniejszą (choć nie jedyną) metodą generowania T-równoważności jest stosowanie nazw cudzysłowowych zdań. Przykładowa T-równoważność ma

wówczas postać: „Zdanie ‘Śnieg jest biały’ jest prawdziwe wtedy i tylko wtedy gdy śnieg jest biały”.

Sam Tarski, jak wiadomo, nie był deflacionistą — sens pojęcia prawdy nie wyczerpywał się według niego w T-równoważnościach. Służyły one raczej do sformułowania kryterium trafności dla definicji prawdy. Z poprawnej definicji prawdy muszą bowiem wynikać wszystkie T-równoważności. (Tarski 1933/1995: 9–10). Dopiero późniejsi filozofowie (por. np. Horwich 1996/2004: 38) uznali, iż sens pojęcia prawdy wyczerpuje się w T-równoważnościach; innymi słowy — nie da się o prawdzie powiedzieć nic, poza wymienieniem zdań, będących podstawieniami podanego przez Tarskiego schematu.

Przyjęte tutaj rozumienie deflacionizmu jest dość szerokie — mieszczą się w nim również koncepcje określane mianem „minimalistycznych”. Przyznają one, że „prawdziwość” jest w pewnym sensie cechą, a predykat prawdy niesie ze sobą pewną treść, jednak jest to treść minimalna. Zdaniem minimalistów, teza, że predykat prawdy nie jest pozorny, „nie pociąga za sobą całej lawiny twierdzeń na temat tego, jaka jest głęboka struktura tej własności, przez co jest ona konstytuowana, od czego jej przypisywanie jest ostatecznie uzależnione itp.” (Szubka 2000: 378). Takie koncepcje pozostają deflacyjne tak długo, jak długo zgadzają się na to, iż sens pojęcia prawdy wyczerpuje się w T-równoważnościach. Mówienie o jednoznacznym rozróżnieniu między substancjalnymi teoriami prawdy a deflacionizmem w szerokim sensie jest, rzecz jasna, pewnym uproszczeniem. Szubka (1999: 306) wskazuje na teorię Alstona jako na przykład koncepcji, która jest „zarazem realistyczna [czyli substancjalna — przyp. K.P] i minimalistyczna”. Na potrzeby dalszych rozważań zostanie jednak przyjęte założenie, że możliwe jest przeprowadzenie ścisłego rozróżnienia między teoriami deflacyjnymi a nie-deflacyjnymi.

Motywacja zwolennika nonfaktualizmu do przyjęcia teorii deflacyjnej wydaje się oczywista. Wedle powszechnego mniemania filozofów, podać znaczenie zda-

nia to podać warunki jego prawdziwości. Skoro żadne zdanie, w którym mowa o znaczeniu, nie może być prawdziwe (ani fałszywe), to w szczególności żadne zdanie przypisujące warunki prawdziwości innemu zdaniu nie może być prawdziwe (ani fałszywe). Aby zaś przyjąć jakąkolwiek substancjalną teorię prawdy, musimy zgodzić się, że zdania mają warunki prawdziwości (Boghossian 1989: 524). Jedynym sposobem na mówienie o prawdzie w sytuacji, gdy odrzucamy znaczenie i warunki prawdziwości jest więc przyjęcie koncepcji, wedle której prawda nie jest realną własnością, a gdy mówimy, że zdanie jest prawdziwe, nie przypisujemy mu żadnej realnej cechy (lub, co najwyżej, cechę rozumianą minimalistycznie). Ujmując rzecz nieco inaczej — jeśli uznamy, że zagadnienia prawdy i znaczenia są ze sobą ściśle związane i jeśli przyjmujemy antyrealistyczny pogląd w tej drugiej kwestii, to w naturalny sposób będziemy mieli skłonność do takiej postawy również w pierwszej.

Może się jednak wydawać, że, wbrew intuicji, nonfaktualista w kwestii znaczenia może być zmuszony do przyjęcia substancjalnej teorii prawdy. Zdaniem Boghossiana (1990: 163–164) tylko przyjęcie teorii niedeflacyjnej pozwala na twierdzenie, że zdania pewnego typu nie mogą być ani prawdziwe, ani fałszywe. Deflacionista bowiem — w myśl tej argumentacji — zmuszony jest przyznać, iż każde poprawnie skonstruowane zdanie oznajmujące musi być prawdziwe lub fałszywe. Jeśli prawdziwość nie jest „realną własnością”, to nie wiadomo, na jakiej podstawie mielibyśmy odmawiać możliwości przysługiwania jej poprawnym zdaniom języka naturalnego. Argumentacja ta opiera się na założeniu, iż „każde proponowane wymaganie, jakie spełniać ma kandydat na zdanie prawdziwe, musi być ugruntowane w uznawanej koncepcji prawdy” (Boghossian 1990: 165). Innymi słowy, tylko w ramach jakiejś wersji substancjalnej koncepcji prawdy możliwa jest teza, że pewnych zdań nie należy oceniać w kategoriach prawdziwości lub fałszywości.

Wyznawcy nonfaktualistycznych poglądów w innych dziedzinach dyskursu mogą bez sprzeczności przyjmować substancjalne teorie prawdy i na tej podstawie

określać, które zdania nie kwalifikują się do oceny w kategoriach prawda/fałsz. Szczególna sytuacja nonfaktualizmu w sprawie znaczenia polega na tym, iż jednocześnie, ze względu na argumenty przytoczone powyżej, winien on założyć deflacyjną teorię prawdy. Mamy tu więc do czynienia, jeśli nawet nie ze sprzecznością sensu stricto, to z pewnego rodzaju napięciem teoretycznym, które sprawia, iż nonfaktualizm w sprawie znaczenia jest — zdaniem Boghossiana — stanowiskiem niemożliwym nie do utrzymania.

Jak jednak wskazywało wielu autorów (m.in. Devitt 1990, Devitt i Rey 1991 oraz Holton 1993), argumentacja Boghossiana nie jest poprawna. Nieskuteczna okazuje się bowiem próba „wymuszenia” na zwolenniku nonfaktualizmu semantycznego przyjęcia substancjalnej teorii prawdy. Aby to pokazać, należy wprowadzić pewne rozróżnienia natury pojęciowej.

Pojęcie zdolności do prawdy będzie rozumiane następująco: zdania, które są zdolne do prawdy, mogą być prawdziwe lub fałszywe. Można wyróżnić dwa podejścia do tego zagadnienia. Pierwsze z nich to minimalizm, pogląd, którego najbardziej znanym orędownikiem jest Crispin Wright. Jest to stanowisko, wedle którego zdolność do prawdy jest „stosunkowo powszechną własnością” (Wright 1999/2000: 187). Zdaniem Wrighta, aby zdanie mogło cechować się zdolnością do prawdy, wystarczy spełnienie dwóch warunków: składni i dyscypliny. Warunek składni wymaga, abyśmy mieli do czynienia z poprawnie skonstruowanym zdaniem oznajmującym, które następnie może być umieszczane w kontekstach propozycyjalnych i okresach warunkowych. Do spełnienia warunku dyscypliny konieczne jest natomiast, by dla interesujących nas zdań istniały standardy poprawnego użycia. Minimalizm uznaje za zdolne do prawdy zdania należące do dyskursów tradycyjnie uważanych za „metafizycznie podejrzane”: etycznego, estetycznego czy tzw. psychologii potocznej.

Podejście przeciwne do minimalizmu Devitt i Rey (1990: 94) ochrzczili mianem selektywnego realizmu. Nie sposób przypisać tego stanowiska jednemu

autorowi — jest to raczej zbiór koncepcji, dla których wspólna jest myśl, iż nie wszystkie poprawnie skonstruowane zdania oznajmujące są zdolne do prawdy. Oznacza to, że aby zdania danej dziedziny dyskursu mogły zostać uznane za prawdziwe lub fałszywe, muszą spełniać dodatkowe kryteria. W koncepcjach tego typu przyjmuje się, że istnieje kryterium selekcji, pozwalające odróżnić dyskursy faktualne od nefaktualnych.

Otrzymujemy zatem cztery możliwe stanowiska:

	deflacionizm	substancjalna teoria prawdy
minimalizm	I	II
selektywny realizm	III	IV

Zgodnie z podejściem Boghossiana, tak naprawdę mamy do wyboru jedynie stanowiska I oraz IV, co znaczy, że możemy przyjąć jednocześnie albo deflacionizm i minimalizm, albo substancjalną teorię prawdy i selektywny realizm. Kluczowe znaczenie dla oceny argumentacji Boghossiana ma pytanie o to, czy sensowne jest stanowisko III — czy można równocześnie twierdzić, iż prawda nie jest „realną” własnością oraz że tylko zdaniom z wybranych dyskursów przysługuje prawo do miana „prawdziwych”, ewentualnie „fałszywych”.

Zdaniem Boghossiana (1990: 164) oraz Szymury (1998: 187–188), trudno sobie wyobrazić, na jakiej podstawie możliwy byłby wybór kryterium selekcji zdań prawdziwych lub fałszywych przy założeniu, że prawda nie jest własnością. Wedle Szymury, „nonfaktualizm przybrałszy postać deflacionizmu uniemożliwił realizację podstawowego zadania każdego nonfaktualizmu — zadania polegającego na odróżnieniu sytuacji, w której mowa o faktach, od sytuacji, które co najwyżej sprawiają takie wrażenie” (Szymura 1998: 188).

Obiekcje te nie wydają się jednak w pełni zasadne. Na poziomie ogólnym na możliwość pogodzenia deflacionizmu z selektywnym realizmem w najbardziej przekonujący sposób wskazuje analogia podana przez Richarda Holtona (1993: 52–

53). Wyobraźmy sobie automat do gry hazardowej na żetony. Akceptuje on tylko żetony o określonych własnościach — kształcie, masie etc. Jednak to, czy dany żeton wygra czy nie, jest dziełem przypadku. Żadna cecha wrzucanego żetonu nie pozwala przewidzieć, czy automat pokaże wygraną, czy przegraną.

Przykład Holtona pokazuje, że warunki uczestnictwa w grze są logicznie niezależne od warunków zwycięstwa w grze — surowym kryteriom selekcji potencjalnych uczestników gry może towarzyszyć brak jakiegokolwiek „realnej cechy”, od której zależałoby zwycięstwo. W szczególnym przypadku dyskusji o prawdzie oznacza to, iż warunki, które musi spełnić zdanie, aby wziąć udział w grze w przypisywanie prawdy i fałszu, mogą być niezależne od tego, co decyduje o „zwycięstwie” w tej grze. Warunki, które musi spełniać żeton, aby zostać zaakceptowanym przez maszynę, są analogonem warunków, jakie musi spełniać zdanie, aby mogło zostać uznane za zdolne do prawdy, zaś fakt, że wybór zwycięskiego żetonu ma charakter czysto losowy — deflacyjnego charakteru prawdy.

Istnieje wiele powodów, dla których można chcieć odróżniać dyskursy faktualne od non-faktualnych, a które nie są wprost związane z teorią prawdy i mogą, *eo ipso*, być akceptowane przez zwolennika deflacionizmu. Deflacionista wyznający scjentyzm będzie przykładowo uważał, iż tylko zdania w obrębie dyskursów spełniających kryteria „naukowości” mogą być uznane za prawdziwe lub fałszywe — na taką ewentualność wskazywał Devitt (1990). Jackson, Oppy i Smith (1994: 298) twierdzą, iż możliwym kryterium faktualności dyskursów jest to, czy zdania w nich występujące są wyrazem przekonań. Wypowiedzi należące do dyskursu etycznego, z racji zawierania terminów normatywnych, nie spełniałyby tego kryterium, albowiem — według teoretyków przyjmujących tzw. Hume’owską teorię motywacji (w sprawie tego pojęcia por. np. Miller 2003: 6–7) — zdania takie nie są wyrazem przekonań, lecz pragnień. Mówiąc zaś bardziej ogólnie, wyrażają one nie stany poznawcze podmiotu, lecz jego stany motywacyjne. Kryterium faktualności zdań, w myśl tego poglądu, jest zatem to, czy wyrażają one stany poznawcze.

Wright (1992: 235) oraz Szymura (1998: 187) odrzucają wariant III jako niesatysfakcjonujący z uwagi na to, że formułowane w ramach tego stanowiska twierdzenia o faktualnym bądź nefaktualnym charakterze danego dyskursu są jedynie deflacyjnie prawdziwe. Trudno uznać zarzut ten za przekonujący. Niewątpliwie, konsekwentny deflacionista, niezależnie od tego, czy zdecyduje się sformułować sąd na temat pogody, czy też tego, które zdania mogą podlegać ocenie prawdziwościowej, winien przyznać, iż jego twierdzenia mają walor jedynie deflacyjnej prawdziwości. Być może ogólnym problemem deflacionizmu jest to, że nie pozwala na wyrażenie intuicji dotyczących tego, co robią użytkownicy języka, gdy mówią „na serio” i wygłaszają rozmaite twierdzenia „z pełnym przekonaniem”. Jeśli jednak uznamy, że deflacionizm można z sensem przyjąć i że deflacionista może wygłaszać rozmaite twierdzenia, to może również stawiać tezy na temat tego, które zdania są zdolne do prawdy.

Źródło słabości argumentacji Boghossiana można doszukiwać się po pierwsze w tym, że nie wyklucza ona tego, by nonfaktualista posługiwał się deflacyjnym pojęciem prawdy. Po drugie zaś, w tym, iż traktuje nonfaktualizm w kwestii znaczenia jako pogląd lokalny, to znaczy taki, który dotyczy jedynie wybranej dziedziny dyskursu.

V.3 Niespójność nonfaktualizmu II: Argument z uogólnienia

Odrzucenie powyższych założeń stanowi sedno argumentu z uogólnienia. Argument ten pierwotnie został (dość skrótowo) przedstawiony przez Crispina Wrighta, zaś jego precyzacji dokonał Martin Kusch. Przeprowadzone przez nich rozumowanie przebiega w dwóch etapach. W pierwszym pokazuje się, że ze stanowiska lokalnego nonfaktualizmu w sprawie znaczenia wynika konieczność przyjęcia nonfaktualizmu globalnego — tezy głoszącej, iż żadne zdanie nie jest zdolne do prawdy. Drugim krokiem jest wykazanie, że globalny nonfaktualizm jest poglądem niespójnym.

Przejście od nonfaktualizmu w sprawie znaczenia do nonfaktualizmu globalnego opiera się na założeniu, że „wartość logiczna S [dowolnego zdania zdolnego do prawdy] jest determinowana przez jego znaczenie oraz to, jaki jest świat” (Wright 1984: 769). Na przykład, wartość logiczna zdania „Cukier jest słodki”, jest określana zarówno przez właściwości sacharozy, jak i przez to, że słowo „cukier” w języku polskim odnosi się do tej właśnie substancji.

Kolejny krok rozumowania polega na pokazaniu, iż niefaktualny charakter tego, co determinuje, przenosi się na to, co determinowane. Przykładowo, jeśli to, jakie prawo należy przyjąć, zależy między innymi od tego, co jest dobre, zaś pytanie o to, co jest dobre, nie jest pytaniem o fakt, to także rozstrzygnięcie pierwszej kwestii nie będzie miało charakteru faktualnego. Podobnie, jeśli zdania o znaczeniu (w tym zdania typu „S ma takie a takie warunki prawdziwości”), nie traktują o faktach, to zdania typu „S jest prawdziwe” — również nie. Przenosząc to wnioskowanie na omawiany wcześniej przykład, jeśli zdanie „słowo 'cukier' odnosi się do sacharozy”, nie jest prawdziwe ani fałszywe, to również zdanie „cukier jest słodki”, nie będzie podlegać ocenie w kategoriach prawdy i fałszu.

W ostatnim kroku rozumowania konieczne jest odwołanie się do T-równoważności (czyli zdań typu „S jest prawdziwe wtedy i tylko wtedy, gdy S”). Według Wrighta, uznawszy pewną równoważność oraz to, że jedna z jej stron odnosi się do faktów, musimy przyznać, iż to samo dotyczy zdania znajdującego się po jej drugiej stronie. I odwrotnie — jeśli jakieś zdanie nie podlega ocenie prawdziwościowej, to zdanie mu równoważne ma taki sam, nonfaktualny, charakter. W omawianym przypadku mamy zatem do czynienia ze swoistym przeniesieniem niefaktualnego charakteru zdania „S jest prawdziwe” na zdanie S.

Powyższy argument wygląda na przekonujący, aczkolwiek jego skrótowy charakter uniemożliwia rzetelną ocenę. Z tego względu warto poddać analizie precyzację tego argumentu, przedstawioną przez Martina Kusch.

Wychodzi on (Kusch 2006: 155–156) od następującej formalizacji stanowiska nonfaktualisty w sprawie znaczenia:

- (1) Dla dowolnych zdań s oraz p i dowolnego sądu Φ : nieprawda, że zdanie ' s ma warunki prawdziwości p ' ma warunki prawdziwości Φ .

Teza (1) jest wyrazem przekonania, iż podać znaczenie jakiegoś zdania, to podać warunki jego prawdziwości. W związku z tym, konsekwentny nonfaktualista w sprawie znaczenia będzie przeczył temu, że zdania, w których przypisujemy warunki prawdziwości innym zdaniom, same mają warunki prawdziwości.

- (2) Jeśli zdanie z jest prawdziwe, to istnieje sąd Φ będący warunkami prawdziwości zdania z (gdzie z jest nazwą dowolnego zdania).

Z (2) przez kontrapozycję i negację kwantyfikatorów otrzymujemy:

- (3) Jeśli dla dowolnego sądu Φ nieprawda, że z ma warunki prawdziwości Φ , to z nie jest prawdziwe.

Ze zdań (1) oraz (3), po podstawieniu pod z zdania ' s ma warunki prawdziwości p ', wynika:

- (4) Dla dowolnych zdań s oraz p : nieprawda, że zdanie ' s ma warunki prawdziwości p ' jest prawdziwe.

Z czego na mocy T-równoważności wynika (przez obcięcie „jest prawdziwe”):

- (5) Dla dowolnych zdań s oraz p : nieprawda, że s ma warunki prawdziwości p .

Zdanie (5) jest wyrazem globalnego nonfaktualizmu. Mówi ono, że żadnemu zdaniu nie można przyporządkować warunków prawdziwości. Przy założeniu, że tylko zdania mające warunki prawdziwości mogą być prawdziwe, oznacza to, że żadne zdanie nie posiada wartości logicznej prawdy (ani fałszu). Przejście od (1) do (5) stanowi sedno „argumentu z uogólnienia” — pokazuje, iż lokalny nonfaktualizm w kwestii znaczenia prowadzi wprost do nonfaktualizmu obejmującego wszystkie dziedziny dyskursu.

Przejście to można by zablokować jedynie poprzez odrzucenie którejś z przesłanek dodatkowych: tezy (2) lub T-równoważności. Żadna z tych opcji nie wydaje się sensowna. Odrzucenie tezy (2) nie ma teoretycznego uzasadnienia, jeśli warunki prawdziwości rozumiane są możliwie szeroko — czyli niezależnie od przyjętej teorii prawdy, nie wyłączając deflacyjnej. Rezygnacja z T-równoważności byłaby z kolei rozwiązaniem czysto *ad hoc* — powszechnie uważa się bowiem, że chwytają one nasze intuicje co do użycia słowa „prawdziwy”.

To, iż w powyższym dowodzie jako drugi argument relacji „...ma warunki prawdziwości...” podstawiano czasem zdania, a czasem sądy, wydaje się zabiegiem czysto stylistycznym. Nic nie wskazuje na obecność jakiegokolwiek poważnej przeszkody, która uniemożliwiałaby przeprowadzenie całego tego dowodu jedynie przy użyciu zmiennych zastępujących zdania.

Potencjalnie kontrowersyjnym założeniem powyższego rozumowania jest natomiast to, iż jeśli zdanie nie ma znaczenia, to nie może mieć żadnych warunków prawdziwości — nawet deflacyjnych. Jak uzasadnić przejście od stwierdzenia, że danemu zdaniu nie można przypisać znaczenia do tezy, iż T-równoważność zbudowana na podstawie tego zdania nie będzie mogła być uznana za prawdziwą?

Aby uznać, że pewne zdanie typu T jest prawdziwe, należy przyjąć, że zachodzi równoważność pomiędzy nim a zdaniem przypisującym mu prawdziwość. Aby jednak „przypisać zdaniu prawdziwość”, musimy użyć nazwy tego zdania. Zatem w T-równoważności zdanie występuje dwa razy: po lewej stronie jest wymienione a po drugiej użyte. Trzeba w tym miejscu przyjąć, iż po obu stronach równoważności występuje to samo zdanie. Jak jednak możemy być pewni, że mamy do czynienia z „tym samym” zdaniem? Czy wystarczy tu równość napisów?

Rozważmy T-równoważność: „Zdanie 'Anna kupiła zamek' jest prawdziwe wtedy i tylko wtedy, gdy Anna kupiła zamek”. Jeśli zdanie w języku przedmiotowym odnosi się do krawcowej robiącej zakupy w pasmanterii, a w metajęzyku do pośredniczki nieruchomości dokonującej zakupu zabytkowego obiektu, to ja-

sne jest, że z żadną równoważnością nie mamy tu do czynienia, choć zdanie użyte i wymienione „wyglądają” tak samo. Aby stwierdzić, że mamy do czynienia z tym samym zdaniem, musimy przyjąć, że dwa — nawet identyczne co do kształtu — napisy mają to samo znaczenie.

Problem ten nie pojawia się jedynie w językach sformalizowanych, bowiem w nich „sens każdego wyrażania jest jednoznacznie wyznaczony przez jego kształt” (Tarski 1995: 31). Co ciekawe, to właśnie tę tezę Tarski uznawał za definicyjny warunek bycia językiem sformalizowanym. Tego założenia nie da się jednak zastosować do języka potocznego (co może było jednym z powodów, dla których Tarski był sceptyczny wobec możliwości zastosowania jego teorii do języków innych niż sformalizowane).

Jeśli uznajemy, iż wszelkie zdania na temat znaczenia są pozbawione wartości logicznej, to nie będziemy mogli prawdziwie orzec, że jakieś dwa zdania mają to samo znaczenie. Uniemożliwia to uznanie jakichkolwiek T-równoważności za prawdziwe — co uniemożliwia z kolei przypisywanie jakimkolwiek zdaniom warunków prawdziwości, nawet rozumianych deflacyjnie.

Drugi etap argumentacji Kuschy polega na wykazaniu, iż globalny nonfaktualizm jest doktryną niespójną. Przedstawiony przez niego dowód jest dość prosty (Kusch 2006: 156–157). Oznaczmy tezę nonfaktualizmu w sprawie znaczenia, czyli tezę (1), literą *a*. Jeśli w tezie (5) w miejsce zmiennej *s* podstawimy stałą *a*, to otrzymujemy:

(6) Dla dowolnego zdania *p*, nieprawda, że *a* ma warunki prawdziwości *p*.

Następnym krokiem jest teza (3) wyrażona w postaci, w której kwantyfikujemy po zdaniach:

(3') Jeśli dla dowolnego zdania *p* nieprawda, że *z* ma warunki prawdziwości *p*,
to *z* nie jest prawdziwe.

Z przesłanek (3') oraz (6) przez podstawienie *a* za *z* otrzymujemy:

(7) Zdanie *a* nie jest prawdziwe.

Ze zdania (7) nie można wyciągać wniosku, że zdanie *a* jest fałszywe — skoro mamy do czynienia z nonfaktualizmem globalnym, to zapewne trzeba będzie uznać, że zdanie *a* jest pozbawione wartości logicznej. Nie mamy tu zatem do czynienia ze sprzecznością *sensu stricto*. Jednakowoż to, iż nonfaktualista w kwestii znaczenia, przyjąwszy przesłankę (1), musi w konsekwencji zaakceptować tezę (7) o nieprawdziwości wyjściowej przesłanki, wskazuje jednoznacznie na niespójność nonfaktualizmu w kwestii znaczenia. Jeśli bowiem głosi się jakąś tezę, to *implicite* przyjmuje się jej prawdziwość. Tymczasem, nonfaktualista w sprawie znaczenia, nie może twierdzić, iż jego własna teoria jest prawdziwa.

Być może mniej radykalna wersja postawy antyrealistycznej, czyli przedstawiona w *Rozdziale I* teoria błędu, jest mniej narażona na tego typu zarzuty? Wnikliwy czytelnik sam z pewnością zauważył, iż argument z uogólnienia można łatwo zastosować do tezy o systematycznym fałszu dyskursu o znaczeniu, głoszącej, iż wszystkie (elementarne) zdania, w których przypisujemy warunki prawdziwości, są fałszywe. Z tezy tej niemal automatycznie wynika natomiast teza (4) — bo jeśli wszystkie zdania mówiące o znaczeniu są fałszywe, to żadne zdanie mówiące o tym, że jakieś zdanie ma dane warunki prawdziwości, nie może być prawdziwe. Dalsza część dowodu przebiega tak samo, jak w przypadku argumentu przeciwko nonfaktualizmowi. Również konkluzja jest ta sama — zwolennik teorii błędu w odniesieniu do znaczenia zmuszony jest przyznać, iż jego własna teoria nie jest prawdziwa.

Zarówno nonfaktualizm, jak i teoria błędu w odniesieniu do dyskursu o znaczeniu są więc poglądami niespójnymi. Nasuwają się w tej sytuacji dwa pytania: po pierwsze, na czym polega szczególny charakter dyskursu o znaczeniu — dlaczego właśnie ten dyskurs (w odróżnieniu od innych) nie może być ujmowany w kategoriach antyrealistycznych? Po drugie zaś: do przyjęcia jakiego stanowiska prowadzi nas refutacja podejścia nonfaktualistycznego i teorii błędu?

Źródłem tego, że dyskurs dotyczący znaczenia nie może być analizowany w kategoriach antyrealistycznych, należy upatrywać przede wszystkim w związku, jaki zachodzi między pojęciami znaczenia i prawdy. Jeśli przyjmiemy, iż odmówienie zdaniu znaczenia jest równoznaczne z zaprzeczeniem, jakoby zdanie to miało warunki prawdziwości — niezależnie od przyjętej teorii prawdy — to okaże się, że nie sposób antyrealizmu w sprawie znaczenia utrzymać jako stanowiska lokalnego. Musi się on w pewnym sensie rozlać na wszystkie pozostałe dyskursy.

V.4 Konsekwencje argumentu z uogólnienia

Jakie wnioski należałoby jednak wyciągnąć z takiej konkluzji? Jaki pogląd przyjąć, żeby uniknąć niespójności? Badając tę kwestię, trzeba najpierw zastanowić się, jakie są możliwe sformułowania antyrealizmu w sprawie znaczenia. Szymura (1998: 186) przedstawił cztery sposoby:

- (i) Zdanie „*s* ma warunki substancjalnej prawdziwości *p*” nie ma warunków deflacyjnej prawdziwości.
- (ii) Zdanie „*s* ma warunki deflacyjnej prawdziwości *p*” nie ma warunków deflacyjnej prawdziwości.
- (iii) Zdanie „*s* ma warunki substancjalnej prawdziwości *p*” nie ma warunków substancjalnej prawdziwości.
- (iv) Zdanie „*s* ma warunki deflacyjnej prawdziwości *p*” nie ma warunków substancjalnej prawdziwości.

Z każdej z powyższych tez można wyprowadzić konsekwencje następującej treści:

- (i') Zdanie „*s* ma warunki substancjalnej prawdziwości *p*” nie jest deflacyjnie prawdziwe.

(ii') Zdanie „s ma warunki deflacyjnej prawdziwości p” nie jest deflacyjnie prawdziwe.

(iii') Zdanie „s ma warunki substancjalnej prawdziwości p” nie jest substancjalnie prawdziwe.

(iv') Zdanie „s ma warunki deflacyjnej prawdziwości p” nie jest substancjalnie prawdziwe.

(Rzecz jasna, wszystkie wyżej wymienione tezy należy traktować jako odnoszące się do wszystkich zdań).

Tezy (i') — (iv') wynikają zarówno z nonfaktualizmu, jak i z teorii błędu. Są to uszczegółowienia przesłanki (4) — kluczowej w zaprezentowanej powyżej argumentacji Kuscha, dowodzącej niespójności teorii antyrealistycznych w odniesieniu do znaczenia.

Warto zadać sobie pytanie, które z tych sformułowań prowadzi do sprzeczności. Na pewno nie dotyczy to tezy (iv'). Jest ona bowiem do zaakceptowania przez zwolennika stanowiska, które można by nazwać konsekwentnym deflacionizmem, tj. przekonania, iż wszystkie zdania mają jedynie deflacyjne warunki prawdziwości i żadne zdanie nie jest prawdziwe w innym sensie. Refutacja argumentu Boghossiana prowadzi do wniosku, iż nie ma powodów, by uważać podejście takie za niespójne. Zwolennik konsekwentnego deflacionizmu nie będzie się też wzbraniał przed akceptacją tezy (i'); jest to bowiem logiczna konsekwencja, którą musi przyjąć, jeśli nie dopuszcza, by jakiegokolwiek zdanie mogło mieć inne warunki prawdziwości niż deflacyjne.

Teza (ii') wydaje się natomiast jednoznacznie niespójna. Jest ona wyrazem postawy, którą Devitt (1990: 259) określił mianem surowego eliminatywizmu — stanowiska, zgodnie z którym żadne zdanie nie jest w żadnym sensie prawdziwe i nie ma warunków prawdziwości. To właśnie stanowisko, jak pokazuje argument z uogólnienia, jest niemożliwe do utrzymania.

Najbardziej intrygująco przedstawia się kwestia spójności tezy (iii'). Z jednej strony, może być ona uznana za element konsekwentnego deflacionizmu. Wynika z niego bowiem, iż skoro brak podstaw, by przypisywać czemuś substancjalne warunki prawdziwości, to zdania mówiące o takich warunkach nie są w żadnym — a więc i substancjalnym — sensie prawdziwe.

Można sobie jednak wyobrazić inne stanowisko, na gruncie którego teza (iii') byłaby prawdziwa — na potrzeby tych rozważań nazwane teorią substancjalnej prawdy ograniczonego zasięgu. Zdaniem hipotetycznego zwolennika tego stanowiska, tylko niektórym zdaniom przysługiwałaby prawda rozumiana substancjalnie, zaś w przypadku pozostałych zdań zdolnych do prawdy, moglibyśmy o niej mówić jedynie w deflacyjnym sensie. W szczególności, wszystkie tezy dotyczące warunków prawdziwości, niezależnie od tego, czy mowa o warunkach substancjalnych, czy deflacyjnych, są jedynie deflacyjnie prawdziwe. Jest to zatem pogląd, wedle którego jakaś wersja substancjalnej teorii prawdy może być lokalnie prawdziwa, lecz zdaniom z dyskursu semantycznego mogą co najwyżej przysługiwać deflacyjne warunki prawdziwości.

Zdaniem Szymury (1998: 188) oraz Wrighta (1992: 236), podejście takie byłoby nieadekwatne. Wright (1992: 215–220) rozważa możliwość zastosowania argumentu z uogólnienia do deflacionizmu w sprawie prawdziwości zdań o znaczeniu. Celem tej argumentacji byłoby pokazanie, iż deflacionizm w tej kwestii prowadzi do deflacionizmu globalnego — stanowiska utrzymującego, iż żadne zdanie nie może mieć substancjalnych warunków prawdziwości. Rozumowanie to jest (pozornie) proste do przeprowadzenia: wystarczy pod „warunki prawdziwości” podstawić „substancjalne warunki prawdziwości”, a pod „prawdziwy” — „substancjalnie prawdziwy”.

Przejście od (1) do (4) oraz dodatkowe przesłanki (2) i (3) zdają się poprawne przy takim podstawieniu. Ale już przejście od (4) do (5) jest wysoce problematyczne. Jeśli dopuszczamy, iż „prawda” jest terminem dwuznacznym — w sensie

takim, iż do niektórych dyskursów stosuje się substancjalne, a do niektórych deflacyjne pojęcie prawdziwości — to z tezy:

- (4') Dla dowolnych zdań s oraz p : nieprawda, że zdanie ' s ma substancjalne warunki prawdziwości p ' jest substancjalnie prawdziwe.

nie wynika wcale:

- (5) Dla dowolnych zdań s oraz p : nieprawda, że s ma substancjalne warunki prawdziwości p .

Zdanie w postaci „ s ma substancjalne warunki prawdziwości p ” może bowiem nie mieć substancjalnych warunków prawdziwości, lecz deflacyjne, a co za tym idzie, może być deflacyjnie prawdziwe. Wykluczenie takiej opcji byłoby klasycznym błędem *petitio principii* — założylibyśmy bowiem z góry, że zdania przypisujące substancjalne warunki prawdziwości nie mogą być „tylko” deflacyjnie prawdziwe, a właśnie to miało być przecież rezultatem rozumowania.

Teorii substancjalnej prawdy ograniczonego zasięgu nie da się zatem zarzucić sprzeczności. Wright (1992: 236) i Szymura (1998: 188–189) twierdzą jednak, że jest to koncepcja dziwaczna. Głosiłaby ona, ich zdaniem, że można przeprowadzić rozróżnienie na „twarde” fakty i fakty „deflacyjne”, tyle że samo to rozróżnienie nie będzie miało charakteru „twardego” faktu. Stanowisko to być może rzeczywiście wydaje się nieintuicyjne, lecz stwierdzenie małej mocy jego przekonywania nie jest tym samym co wykazanie jego niespójności.

Zwolennicy teorii substancjalnej prawdy ograniczonego zasięgu mogliby argumentować, podobnie jak filozofowie optujący za połączeniem deflacyjnej teorii prawdy i selektywnego realizmu, iż istnieją kryteria podziału nie związane z teorią prawdy. Taki pogląd mógłby wyznawać na przykład scjentyista, głoszący iż tylko zdania należące do zakresu nauk szczegółowych mają substancjalne warunki prawdziwości. Mógłby być on, co więcej, scjentyistą świadomym, tzn. zdającym sobie sprawę z faktu, iż sam scjentyzm, jako stanowisko, nie jest przez naukę

uzasadniony, w związku z czym przywiązanie do niego jest czymś na kształt estetycznej preferencji. Dopóki dopuszczamy możliwość, iż używanie prawdy w sensie deflacyjnym może być wyrazem szczerzej asercji, to nie ma powodu, aby taką deklarację uważać za niespójną.

Podsumowując, mimo że da się dowieść niespójności antyrealistycznych podejść do znaczenia, wnioski teoretyczne, jakie płyną z tego faktu, nie są zbyt daleko idące. Nie do utrzymania okazują się obie postacie skrajnego antyrealizmu — zarówno nonfaktualizm, jak i teoria błędu — odmawiające jakiejkolwiek formy prawdziwości zdaniom o znaczeniu. Aby uniknąć sprzeczności, musimy przyznać, iż pewne zdania o znaczeniu — a w szczególności T-równoważności — są w jakimś (deflacyjnym lub substancjalnym) sensie prawdziwe. Najsłabsze stanowisko, które umożliwia ominięcie sprzeczności, to takie, w którym mówi się, że pewne zdania przypisujące innym zdaniom deflacyjne warunki prawdziwości same mogą być deflacyjnie prawdziwe. Jedną z koncepcji, która umożliwia wypowiedzenie tego typu tezy, jest minimalizm w kwestii zdolności do prawdy.

V.5 Minimalizm jako interpretacja rozwiązania sceptycznego

Przyjęcie minimalizmu w sprawie zdolności do prawdy zdań o znaczeniu wystarczy, aby uwolnić się od zarzutu niespójności. Co ciekawe, wedle niektórych interpretatorów, sceptyczne rozwiązanie Kripkensteina (omówione w *Rozdziale I*) należy odczytywać jako wyraz minimalizmu właśnie — zwolennikami takiego ujęcia byli, między innymi Byrne (1996: 342–343), Kusch (2006: 158–176), Soames (1998a: 323) oraz Wilson (1998: 108–109).

Trzeba przyznać, iż wobec niejasnej natury tekstu Kripkego, który nie używa takich technicznych terminów, jak „nonfaktualizm” czy „zdolność do prawdy”, trudno jest podać niekontrowersyjną interpretację jego poglądów. Co więcej,

dość łatwo byłoby znaleźć w dziele Kripkego fragmenty, które można uznać za wspierające wzajemnie znoszące się poglądy.

Na pierwszy rzut oka wydaje się, iż Kripkenstein jednoznacznie opowiadał się za nonfaktualizmem w kwestii znaczenia. Świadczyć mogą o tym następujące cytaty:

Odważę się jednak powiedzieć: Wittgenstein utrzymuje, w zgodzie ze sceptykiem, że nie istnieje żaden fakt określający to, czy mam na myśli [*I mean*] plus czy kwus (Kripke 2007: 116 -117).

I dalej:

Przypomnijmy sobie sceptyczną konkluzję Wittgensteina: żadne warunki prawdziwości nie odpowiadają [*correspond to*] stwierdzeniom takim jak: „Jones ma na myśli [*means*] dodawanie, kiedy używa '+'” (Kripke 2007: 126).

Wreszcie:

Sceptyczne rozwiązanie Wittgensteina zgadza się ze sceptykiem co do tego, iż nie ma w świecie żadnych „warunków prawdziwości” czy też „odpowiadających faktów”, które czyniłyby prawdziwymi stwierdzenia takie, jak: „Jones, jak wielu z nas, gdy używa '+' ma na myśli dodawanie” (Kripke 2007: 139)

Z drugiej jednak strony, na tej samej stronie, co ostatni z przytoczonych fragmentów, znajduje się następująca wypowiedź:

Czyż nie nazywamy tego typu stwierdzeń „prawdziwymi” bądź „fałszywymi”? Czy nie możemy w sposób poprawny poprzedzać takich stwierdzeń frazami: „Jest faktem, że...” lub „Nie jest faktem, że...”? Wittgenstein krótko rozprawia się z takimi zarzutami. Podobnie jak

wielu innych akceptuje on „redundancyjną” teorię prawdy: stwierdzić, że zdanie jest prawdziwe (a być może także poprzedzić je frazą „Jest faktem, że...”), to po prostu stwierdzić samo to zdanie, zaś stwierdzić, że jest ono nieprawdziwe, to po prostu mu zaprzeczyć ($'p'$ jest prawdziwe $= p$). (...) *Nazývámy* coś sądem, a w związku z tym określamy to coś mianem „prawdziwego” lub „fałszywego” wówczas, gdy w naszym języku stosujemy do tego czegoś rachunek funkcji prawdziwościowych. To, iż funkcje prawdziwościowe stosują się do pewnych zdań, jest po prostu pierwotną częścią naszej gry językowej, niemożliwą do wyjaśnienia w bardziej podstawowy sposób (Kripke 2007: 139–140).

W świetle powyższych cytatów można stworzyć dwie interpretacje stanowiska Kripkensteina. W pierwszej z nich jest on zwolennikiem nonfaktualizmu w kwestii znaczenia, konsekwentnie odmawiającym zdaniom z dyskursu semantycznego zdolności do posiadania wartości logicznej. W drugiej przyjmuje minimalistyczną wersję realizmu i przyznaje, że zdania z dyskursu o znaczeniu mogą być deflacyjnie prawdziwe lub fałszywe. Uznanie, że zdania o znaczeniu mają deflacyjne warunki prawdziwości wiąże się z odrzuceniem nonfaktualizmu, którego zasadnicza teza mówi, że zdania z danej dziedziny dyskursu nie posiadają wartości logicznej.

Skoro sam tekst nie umożliwia jednoznacznej interpretacji stanowiska Kripkensteina, to należy odwołać się do argumentacji pośredniej. Przy założeniu interpretacji nonfaktulistycznej, jego poglądy stają się w sposób jednoznaczny niezgodne z tymi, które w jakiś sposób żywimy na poziomie przedfilozoficznym. Gdy używamy zdań z dyskursu semantycznego, to wydaje się nam, że dokonujemy rzetelnych asercji, a wypowiedane zdania mają określoną wartość logiczną. Nonfaktualista odrzuca to przekonanie. Gdyby Wittgensteina czytać w ten sposób, należałoby go uznać za głosiciela reformistycznego podejścia do języka naturalnego. Reformista, w sytuacji, w której argumentacja filozoficzna wykazuje, iż dany fragment potocznego dyskursu jest z jakiegoś względu nieuprawniony,

stwierdza, iż należy dążyć do eliminacji owego dyskursu lub reinterpretacji jego roli w języku. Przykładami tego typu podejścia mogą być stanowisko Churchlanda w odniesieniu do dyskursu mentalnego lub opis zdań etycznych autorstwa Ayera.

Można z dużą dozą pewności stwierdzić, iż Wittgensteinowi, w późnym okresie jego twórczości, tego typu reformistyczna postawa w odniesieniu do języka naturalnego była całkowicie obca. To przecież właśnie Wittgenstein stwierdził, że „Gdyby w filozofii chciano wysuwać *tezy*, to nie mogłyby one stać się nigdy przedmiotem dyskusji, gdyż wszyscy by się na nie godzili” (Wittgenstein 1953/2000 § 128). Twierdził on zatem, że filozofowie nie mogą przeczyć potocznym przeświadczeniom — a reformiści językowi właśnie to chcą czynić.

Aby jednak móc uznać, że stanowisko Kripkensteina nie jest stanowiskiem skrajnego antyrealizmu semantycznego, należałoby wyjaśnić sens twierdzeń Kripkego, o tym, iż nie istnieją „fakty” ani „warunki prawdziwości”, które miałyby odpowiadać zdaniom o znaczeniu. Wilson (1998: 109) twierdzi, iż powinniśmy odróżniać od siebie stanowiska Kripkensteina oraz sceptyka w sprawie znaczenia. „Radykalny wniosek sceptyczny”, według Wilsona, brzmi: nie ma czegoś takiego jak znaczenie. Stąd właśnie w tekście tezy o „nieistnieniu faktów” etc. Jednak sam Kripkenstein tak radykalnych wniosków nie akceptuje, uznając stanowisko bardziej umiarkowane, czyli minimalistyczne właśnie.

Zdaniem omawianego autora, argument Kripkensteina ma strukturę dowodu nie-wprost: pokazuje się w nim, iż przyjęcie pewnych założeń na temat znaczenia (nazywanych przez Wilsona „klasycznym realizmem”) prowadzi do niemożliwej do zaakceptowania konkluzji, czyli radykalnego wniosku sceptycznego. Za kluczowy element „klasycznego realizmu” Wilson uznaje przekonanie, iż aby dany termin coś znaczył w ustach danej osoby, musi istnieć zbiór cech, które konstytuują standardy poprawnego użycia tego wyrażenia przez tę osobę. Sceptyk i Kripkenstein są zgodni co do fałszywości klasycznego realizmu, jednak różnią się co do tego, jakie konsekwencje należałoby z niej wyciągnąć.

Podobną interpretację, uznającą Kripkego za zwolennika faktualizmu w kwestii znaczenia, lecz w większym stopniu posługującą się kategoriami przyjmowanymi w niniejszej pracy, przedstawił Byrne. Zdaniem tego filozofa, Kripkenstein uznaje minimalizm w kwestii zdolności do prawdy, oraz tezę, iż zdania o znaczeniu mają wyłącznie deflacyjne warunki prawdziwości. Domniemany nonfaktualizm Kripkensteina ogranicza się w tej interpretacji do poglądu, iż niemożliwe jest podanie substancjalnych warunków prawdziwości dla zdań dyskursu semantycznego. Byrne rozwija tę myśl następująco:

Posiadanie czegoś na myśli (*meaning something*) może wydawać się bardzo dziwnym zjawiskiem. (...) *Co więc czyni prawdziwym* to, iż obecnie mam na myśli (*I mean*) dodawanie używając „+”? Nie sprawia tego moje przeszłe użycie tego znaku, ani też nic, co znajduje się przed moim umysłem gdy używam tego znaku. Nie są to również moje przeszłe dyspozycje. W rzeczywistości, *nic interesującego* nie czyni prawdziwym, tego, że przez 'plus' mam na myśli dodawanie (Byrne 1996: 342-343).

Sceptyczny charakter stanowiska Kripkensteina polegać miałby więc na tym, iż nie da się sformułować warunków prawdziwości dla zdań o znaczeniu, które miałyby inny charakter niż deflacyjny. Nie istnieje zatem żaden nietrywialny i niekolisty opis, który można by podstawić po prawej stronie T-równoważności konstruowanych dla zdań semantycznych. Prawdziwości zdań o znaczeniu nie da się wyjaśnić w żaden sposób. Wszystkie przedstawione dotychczas teorie znaczenia próbowały pokazać, co może pełnić funkcję faktów odpowiadających prawdziwym zdaniom o znaczeniu i żadna z nich nie spełniła tego zadania.

Byrne odwołuje się tutaj do stworzonego przez Dummetta pojęcia zdań, które są prawdziwe tak po prostu (*barely true*). Dummett wprowadza to pojęcie opierając się na przykładzie nierzeczywistych okresów warunkowych. Koncepcja Dummeta polega na tym, że jeśli uznajemy jakiś nierzeczywisty okres warunkowy

za prawdziwy, przyjmujemy, że musi istnieć inne zdanie, z którego można ten okres wyprowadzić. Na przykład, jeśli mówimy, że gdybyśmy przyłożyli kartkę papieru do rozgrzanego piecyka, to zapaliłaby się, zakładamy, że prawdziwość tego okresu warunkowego opiera się na jakimś sądzie, dotyczącym cech fizycznych kartki papieru. „Innymi słowy, musi istnieć, dla każdego prawdziwego i nierzeczywistego okresu warunkowego, nietrywialna odpowiedź na pytanie 'Co sprawia, że jest on prawdziwy?'" (Dummett 1976/1993: 53). Dummett twierdzi zatem, że nierzeczywiste okresy warunkowe nie mogą być prawdziwe „tak po prostu”. Byrne twierdzi natomiast, iż sceptyczne rozwiązanie Kripkensteina sprowadza się do powiedzenia, iż zdania o znaczeniu są tylko prawdziwe w takim właśnie sensie. Na pytanie typu: „Co sprawia, że zdanie *ś*łowo 'plus' oznacza funkcję dodawania, jest prawdziwe?”, mogę odpowiedzieć jedynie wzruszeniem ramion lub stwierdzeniem: „słowo 'plus' oznacza funkcję dodawania”.

V.6 Zdania o znaczeniu a warunek dyscypliny

Minimalizm w sprawie zdolności do prawdy jest stanowiskiem, które pozwala uniknąć niespójności, w jaką wikła nas nonfaktualizm, przy jednoczesnej cesji na rzecz sceptyka. Minimalizm jest do pogodzenia ze sceptyczną argumentacją — da się na jego gruncie utrzymać twierdzenie o nieistnieniu faktów semantycznych, o ile „fakty” będziemy rozumieć mocniej niż deflacyjnie.

Jak wspomnieliśmy wcześniej, minimalizm narzuca zdaniom, które miałyby być zdolne do prawdy dwa warunki. I o ile można uznać, że zdania o znaczeniu bez problemu spełniają pierwszy z nich — warunek składni, o tyle problematyczny jest warunek dyscypliny. Edwards (1994: 64) oraz Miller (2007: 199) poddali w wątpliwość przekonanie, iż zdania o znaczeniu go spełniają.

Według tych autorów warunek dyscypliny ma charakter w dużej mierze normatywny. Oznacza on, że istnieją pewne standardy, których spełnianie warunkuje fakt, czy dane zdanie podlega ocenie w kategoriach prawdziwości lub fałszywo-

ści. Lecz to właśnie istnienie standardów poprawnego użycia było celem ataku sceptyka Kripkensteina. Miller wyraża tę wątpliwość następująco:

(...) warunki stwierdzalności — standardy właściwego użycia — które leżą u podstaw minimalnej prawdy są normatywne w tym sensie, który został zakwestionowany przez negatywną argumentację sceptyka Kripkego-Wittgensteina. Ponieważ są to standardy, które dzielą użycia X'a na kategorie *poprawne* i *niepoprawne*, to przyjmując minimalistyczną koncepcję zdolności do prawdy, tym samym uznajemy za dane pojęcie poprawności, kwestionowane przez sceptyka (Miller 2007: 199).

Przyjrzyjmy się dokładniej temu argumentowi. Po pierwsze, należy zauważyć, iż koncepcja Wrighta nie jest bynajmniej jednoznacznie minimalistyczna. Nie jest tak, iż każde zdanie jest w myśl tej koncepcji prawdziwe lub fałszywe. Możliwe jest bardziej radykalne stanowisko w tej kwestii, które Jackson, Oppy i Smith (1994: 291) ochrzcili mianem składnizmu (*syntacticism*). W tym ujęciu podleganie ocenie prawdziwościowej jest wyłącznie kwestią składni — każde zdanie oznajmujące, zbudowane według zasad danego języka naturalnego będzie posiadać określoną wartość logiczną. Oczywiście, stanowiska tego nikt nigdy nie wyznawał, albowiem w takim ujęciu musielibyśmy przypisywać prawdę lub fałsz zdaniom, które są jawnie absurdalne, lecz poprawne składniowo, jak choćby „Zielone idee wściekle śpią” — by przywołać klasyczny przykład Chomsky'ego. Wprowadzenie warunku dyscypliny pozwala wyeliminować takie zdania z grona kandydatów do wartości logicznej.

Teoria Wrighta dostarcza pewnego kryterium uznawania zdań za zdolne do prawdy. Pewne zdania nie będą ani prawdziwe ani fałszywe — do oceny w kategoriach wartości logicznej nadają się tylko te, które funkcjonują w zdyscyplinowanym dyskursie. Składniowo poprawne absurdy tego warunku ewidentnie nie spełniają.

Miller i Edwards twierdzą, iż jeśli przyjmujemy „zdyscyplinowany składnizm” to tym samym zobowiązujemy się do uznania, iż istnieją pewne standardy poprawnego użycia, których istnienie zostało zakwestionowane przez sceptyczną argumentację Kripkego. Wydaje się więc, że nie można używać minimalizmu jako sposobu na poradzenie sobie z problemem sceptycznym bez uprzedniego rozwiązania tego problemu. Rozwiązanie minimalistyczne jawi się tu jako klasyczne błędne koło.

Argumentacja ta jest jednak błędna. Zakłada się w niej bowiem niezwykle silne rozumienie pojęcia „dyscypliny”, w myśl którego, aby było można mówić o spełnianiu tego warunku przez daną dziedzinę dyskursu, muszą istnieć pewne obiektywne kryteria podziału użyć wyrażen należących do tego dyskursu na poprawne i niepoprawne.

Podstawową motywacją dla Wrighta, przy wprowadzaniu minimalizmu było tymczasem uzasadnienie tezy, iż „zdolność do prawdy jest stosunkowo powszechną własnością” (Wright 1999/2000: 187). Miała ona w zamierzeniu obejmować np. dyskurs estetyczny czy dyskurs o tym, co jest zabawne (Wright 1992: 7 -12). Czy ten cel można osiągnąć, jeśli będziemy rozumieli warunek dyscypliny tak, jak proponują to Edwards i Miller?

Odpowiedź na pytanie jest o tyle trudna, że sam Wright w dość niejasny sposób wyraża się na temat charakteru warunku dyscypliny. Na pewno muszą istnieć pewne standardy poprawności, jednak nie wiadomo dokładnie, jaki jest ich status.

Edwards i Miller zdają się sądzić, iż owe „standardy” mają status podobny do reguł semantycznych, tak jak były one pojmowane przez klasycznych realistów w sprawie znaczenia. Uważają oni, że standardy poprawnego użycia powinny, po pierwsze, wyznaczać, który sposób użycia będzie poprawny, a który nie, dla każdego możliwego kontekstu. Kolejną cechą obiektywistycznej wersji minimalizmu jest przeświadczenie, iż o tym, które użycia są poprawne, nie decyduje wyłącznie skłonność mówiących do wypowiedzania się w określony sposób. Normy poprawności

językowej powinny charakteryzować się pewną dozą obiektywności i niezależności od przekonań użytkowników. Nawet jeśli kompetentny użytkownik języka jest przekonany, iż używa danego wyrażenia poprawnie, to nie oznacza to, że jest tak w istocie.

Taka charakterystyka obiektywistycznego minimalizmu może wydać się nierzetelnym podsumowaniem stanowiska Millera i Edwardsa — przypisuje im bowiem przekonania niewiele różniące się od tych, które wyznawał oponent Kripkensteinowskiego sceptyka. Wydaje się jednak, że aby zarzut kolistości miał sens, przyjęcie, iż standardy poprawności muszą mieć te same cechy co reguły semantyczne, byłoby konieczne. W innym wypadku nie miałyby sensu twierdzenie, że brak rozwiązania problemu Kripkensteina uniemożliwia stosowanie minimalizmu w sprawie zdolności do prawdy.

Aby wykazać, że obiektywistyczny minimalista musi przyjmować taką właśnie koncepcję, przedstawmy alternatywną, „radykalną”, interpretację minimalizmu. Po pierwsze, nie musimy uznawać, iż w każdej możliwej sytuacji istnieje jednoznaczna odpowiedź na pytanie, czy użycie danego wyrażenia należącego do problematycznego dyskursu jest poprawne. Zamiast tego, w tym ujęciu zakłada się, iż wszystko, czego potrzeba do spełnienia warunku dyscypliny to istnienie społecznego konsensusu co do poprawności pewnych wzorcowych przypadków użycia omawianych wyrażen. Rozważmy dyskurs na temat tego, co jest zabawne. W myśl radykalnego minimalizmu, tym, co musimy stwierdzić, aby móc określić ten dyskurs mianem zdyscyplinowanego jest fakt, iż w pewnych szczególnych przypadkach niemal wszyscy z nas będą mieli tendencję do podobnych zachowań werbalnych; przynajmniej większość kompetentnych użytkowników języka uzna np. film z Chaplinem za zabawny. Natomiast ci, których zachowanie będzie w rażący sposób odstawać od większości — na przykład ktoś, kto uzna za „prześmieszoną” instrukcję obsługi edytora Microsoft Word — będą przez nas uznawani za dewiantów nieposiadających kompetencji językowych. Kripke opisał

taki mechanizm pisząc, że gdy uznajemy, że ktoś posługuje się pewnym symbolem w pewnym znaczeniu, to „zobowiązujemy się raczej do tego, że jeśli w przyszłości (...) będzie się [on] zachowywał (odpowiednio często) w odpowiednio dziwny sposób, to nie będziemy się dalej upierać przy twierdzeniu, że kieruje się on zwykłą regułą (...)” (Kripke 1982/2007: 152). Oprócz zachowań jednoznacznie poprawnych i zdradzających ewidentną nieznajomość reguł istnieje jednak dość spora grupa zachowań, o których nie można orzec ani poprawności, ani jej braku.

Radykalny minimalista twierdzi również, iż nie ma potrzeby przyjmować — tak jak to czynił zwolennik obiektywizmu — że większość użytkowników może się mylić jeśli idzie o ocenę poprawności jakiegoś zachowania językowego. Pytanie o to, czy w języku stosujemy się do jakichś niezależnych reguł, jest nieistotne, jeśli szukamy kryteriów spełniania przez dany dyskurs warunek dyscypliny. Jedyna rzecz, która liczy się w tym kontekście, to obserwacja, iż zachowujemy się podobnie. Kripke pisze w tym kontekście o „nagim fakcie, że się zazwyczaj ze sobą zgadzamy” (Kripke 1982/2007: 155). Przy czym zgoda nie może być rozumiana jako uchwytowanie tych samych pojęć czy kierowanie się tymi samymi regułami. Prowadziłoby to wprost to błędnego koła — chcąc uzasadnić zdolność do prawdy zdań z dyskursu o znaczeniu, odwoływalibyśmy się do faktów, które same w sobie mają semantyczny charakter. Aby uniknąć tego typu zarzutów, „zgoda” należy rozumieć w kategoriach czysto behawioralnych.

Jeśli warunek dyscypliny będziemy interpretować w sposób postulowany przez radykalny minimalizm, to trudno zrozumieć, w jaki sposób to, do czego przekonuje nas sceptyk, miałyby zagrozić twierdzeniu, iż dyskurs semantyczny jest zdyscyplinowany. Sceptyczne rozważania nie kwestionują tego, iż w szczególnych wzorcowych przypadkach większość użytkowników języka ma tendencję do podobnych zachowań werbalnych. Większość z nas będzie skłonna akceptować (w czysto behawioralnym sensie) proste zdania wyjaśniające znaczenie typu: „'kawaler' znaczy tyle, co 'nieżonaty mężczyzna'”, zaś odrzucać zdania typu „'lodówka' znaczy to samo co 'lokówka'”.

Przyjęcie radykalnego minimalizmu pozwala wyjaśnić rolę, jaką odgrywa społeczność w sceptycznym rozwiązaniu problemu Kripkensteina. W *Rozdziale IV* zostały przedstawione powody, dla których nie można utrzymywać, iż odwołanie się do wspólnoty i zgody społecznej pozwoli na podanie warunków prawdziwości dla zdań o znaczeniu. Możemy jednak stwierdzić, że istnienie zgody wśród wspólnoty językowej umożliwia zdolność do prawdy tychże zdań. W tym ujęciu, nagi fakt, że mamy tendencję do podobnych zachowań językowych, pozwala wysnuć wniosek, że zdania o znaczeniu mają wartość logiczną prawdy lub fałszu. Nie prowadzi to jednak do akceptacji przekonania, że warunki prawdziwości tych zdań muszą zawierać odniesienie do zgody społecznej. Warunki prawdziwości i warunki zdolności do prawdy są od siebie niezależne.

Przeciw temu ostatniemu stwierdzeniu zdecydowanie protestowali Jackson, Oppy i Smith (1994: 291 - 293). Ich zdaniem nie można rozdzielać kwestii warunków zdolności do prawdy od warunków prawdziwości. W swojej polemice ze składnizmem, czyli przekonaniem, iż składnia stanowi wystarczający warunek podlegania ocenie prawdziwościowej, przedstawiają następujący argument: zdanie nie może być prawdziwe lub fałszywe, jeśli nie ma określonych warunków prawdziwości. Jeśli przyjmujemy zatem, że dane zdanie jest zdolne do prawdy, to twierdzimy zarazem, iż posiada one pewne, w miarę określone, warunki prawdziwości. Wymienieni autorzy wyciągają stąd wniosek, iż zwolennik składnizmu musiałby przyjąć następującą tezę: „Same fakty dotyczące składni są wystarczające do tego, by dane zdanie miało te warunki prawdziwości a nie inne” (Jackson, Oppy, Smith 1994: 293).

Taki wniosek stanowi podstawę do odrzucenia składnizmu. Nie sposób bowiem wyczytać z samej składni danego zdania tego, jakie ma ono warunki prawdziwości. Jeśli zarówno argumentacja, jak i założenia Jacksona i innych są poprawne, to trzeba wprowadzić dodatkowe kryteria zdolności do prawdy, które pozwolą na określenie warunków prawdziwości zdań podlegających ocenie logicznej.

Widać jednak, iż wątpliwości wysunięte przez omawianych autorów w stosunku do składnizmu, stawiają pod znakiem zapytania również minimalizm uznający konieczność wprowadzenia dodatkowego warunku dyscypliny, zwłaszcza jeśli warunek ten ujmować w taki sposób, w jaki proponuje to zwolennik minimalizmu w wersji radykalnej. Musimy bowiem zadać sobie pytanie, czy ustanawiane przezeń warunki zdolności do prawdy dostarczają nam sposobu przypisywania warunków prawdziwości omawianym zdaniom?

Gdyby tak było, to przyjęcie radykalnego minimalizmu musiałoby nas prowadzić do wniosku, iż jeśli chcemy podać warunki prawdziwości dla zdań o znaczeniu, to musimy się przy ich formułowaniu odwołać do faktów dotyczących zgody społecznej. To właśnie zgoda społeczna gwarantuje zdolność do prawdy owych zdań. W ten sposób jednak musielibyśmy powrócić do zdezwuowanej wcześniej społecznej koncepcji warunków prawdziwości dla zdań z dyskursu semantycznego.

Aby obronić się przed tego typu zarzutami, zwolennik radykalnego minimalizmu musi zakwestionować istnienie związku między warunkami zdolności do prawdy a warunkami prawdziwości. Nie jest to zresztą to rozwiązanie tak dziwne, jak mogłoby się początkowo wydawać. Rozważmy raz jeszcze przywoływaną już analogię Holtona i płynące z niej konsekwencje. Miała ona służyć pokazaniu, iż warunki uczestnictwa w grze nie muszą być tożsame z warunkami odniesienia zwycięstwa. Przekładając to na kategorie warunków zdolności do prawdy i warunków prawdziwości można powiedzieć, że te pierwsze nie muszą wyznaczać tych drugich. Co więcej, może być tak, iż nawet jeśli uda nam się powiedzieć coś nietrywialnego o tym, co determinuje fakt, iż jakieś zdania podlegają ocenie prawdziwościowej bądź nie, to nie będziemy nic umieli powiedzieć o tym, dlaczego jakieś zdania są prawdziwe. Tak właśnie mają się sprawy ze zdaniami o znaczeniu.

Nawet radykalny minimalizm jest w pewnym sensie bogatą teorią zdolności do prawdy — nie jest bowiem tak, iż absolutnie każde zdanie jest w myśl tej koncepcji zdolne do prawdy. Warunek dyscypliny pozwala wyeliminować pewne

zdania z grona kandydatów do prawdy. Nie ma jednak istotnego powodu, by przyjąć, iż musimy warunek dyscypliny czynić także warunkiem prawdziwości zdań.

V.7 Minimalizm a nonredukcjonizm semantyczny

Dwie podstawowe tezy radykalnego minimalizmu w sprawie znaczenia można ująć w sposób następujący. Po pierwsze, zdaniom o znaczeniu można przypisać jedynie deflacyjne warunki prawdziwości. Po drugie, to, że w ogóle można im te warunki prawdziwości przypisywać jest wynikiem jedynie tego, że zdania te spełniają minimalistyczne warunki zdolności do prawdy. Innymi słowy, nie da się powiedzieć nic interesującego ani o tym, co sprawia, że zdania te mogą w ogóle być oceniane jako prawdziwe, ani o tym, dlaczego przynajmniej niektórym z nich własność prawdziwości przysługuje. Są to zdania „po prostu prawdziwe”.

Alex Miller (2010: 179 - 180) zarzuca takiej koncepcji, iż jest ona *de facto* równoznaczna z przyjęciem „desperackiego” i „tajemniczego” nonredukcjonizmu w sprawie znaczenia. Argument ten wydaje się poważny. Trudno bowiem zaprzeczyć, iż stanowiska te mają wiele wspólnego — łączy je choćby odrzucenie wszelkich form redukcjonizmu znaczeniowego. Zarówno minimalistą, jak i zwolennik platonizmu twierdzą, iż nie można wskazać na żadne fakty, które nie miałyby same w sobie charakteru semantycznego a które odpowiadałyby zdaniom z dyskursu o znaczeniu. Kluczową sprawą staje się zatem zidentyfikowanie różnicy pomiędzy tymi stanowiskami.

Dla anty-redukcjonisty, to, iż nie możemy podać substancjalnych warunków prawdziwości dla zdań o znaczeniu, jest świadectwem tego, że istnieją specjalne fakty semantyczne, które mają zupełnie inny charakter od wszelkich innych faktów. Można odnieść wrażenie, jakoby dla zwolenników tej koncepcji fakty semantyczne tworzyły coś w rodzaju osobnej substancji, podobnej do substancji myślącej u Kartezjusza. Zadaniem teorii znaczenia byłby wówczas opis ontologiczny tejże substancji.

Podejście minimalistyczne jest niechętnie tego typu rozstrzygnięciom natury metafizycznej. Teza o nieredukowalności zdań o znaczeniu ma dla minimalisty charakter „gramatyczny”, w Wittgensteinowskim sensie tego słowa. To, że nie można podać substancjalnych warunków prawdziwości dla zdań o znaczeniu wynika bardziej z natury „gry językowej” w przypisywanie znaczeń, niż rozważań natury ontologicznej. Ta gra ma po prostu taki charakter, iż nie jest możliwe zastąpienie zdań o znaczeniu jakimikolwiek innymi — czy to o dyspozycjach, czy o jakościowych stanach mentalnych itp.

Jeżeli jednak pominąć tego typu deklaracje, to nie jest łatwo dokładnie wyrazić różnicę między nonredukcjonizmem a minimalizmem. Wright (1992: 24–29) zwraca uwagę na podobny fenomen w dyskusjach pomiędzy deflacionistami a zwolennikami substancjalnej teorii prawdy. Wydaje się, iż nie ma takiego potocznie akceptowanego zdania na temat prawdy — jak to określa Wright banału (*platitude*) — którego deflacionista nie mógłby ująć w swoim schemacie pojęciowym i uznać za poprawne. Rozważmy np. twierdzenie „zdanie jest prawdziwe wtedy i tylko wtedy gdy koresponduje z faktami”. Deflacionista może uznać to zdanie za prawdziwe; wystarczy, że zinterpretuje pojęcia „faktu” oraz „korespondencji” w sposób deflacyjny — np. przez „fakt” może rozumieć „to, co jest opisywane przez zdanie prawdziwe” a przez „korespondencję” — „relację zachodzącą pomiędzy zdaniem prawdziwym a tym, co ono opisuje”. Wright twierdzi, iż jest to problem dla zwolenników substancjalnej teorii prawdy — nie potrafią oni wskazać takiego potocznie używanego zdania na temat prawdy, które odróżniałoby ich stanowisko od deflacionizmu. Sytuacja ta jest jednak kłopotliwa również dla deflacionistów. Rodzi się bowiem pytanie: czym deflacionizm miałby się różnić od bogatych teorii prawdy, jeśli akceptuje każde potoczne sformułowanie wyrażające intuicje stojące za ujęciem prawdy jako korespondencji? Zwolennik bogatej teorii prawdy będzie przedstawiał inną filozoficzną interpretację tych samych banałów — np. inaczej będzie rozumiał pojęcie faktu, czyniącego zdanie prawdziwym (por. Szubka 1999:

310). Szubka (2000: 384) zwraca jednak uwagę na to, iż teoria korespondencyjna, posługująca się słabym pojęciem korespondencji oraz teoria minimalistyczna, dopuszczająca mówienie o prawdzie jako cesze i o faktach, do których odnoszą nas zdania prawdziwe, stają się w pewnym momencie od siebie nieodróżnialne. Być może różnica między tymi stanowiskami polega na tym, że deflacionista, w przeciwieństwie do zwolennika teorii substancjalnej, będzie twierdził, że nie można powiedzieć nic więcej na temat tego, co sprawia, że dane zdanie jest prawdziwe, niż tylko przywołać T-równoważność. Być może zwolennik teorii korespondencyjnej, jeśli chce się odróżnić od deflacionisty, musi przyjąć mocne rozumienie korespondencji. Kwestie te są niełatwe do rozstrzygnięcia, co skłania do wniosku, że odróżnienie deflacionizmu od korespondencyjnej teorii prawdy nie jest sprawą trywialną.

Podobne kłopoty zdaje się sprawiać wysłowienie różnicy pomiędzy tezami nonredukcyjisty w sprawie znaczenia a przedstawionym powyżej radykalnym minimalizmem. Problem zdaje się polegać na tym, iż każde wyrażenie stanowiska nonredukcyjistycznego można zinterpretować w sposób, który będzie zgodny z minimalistycznym sposobem spojrzenia na status zdań o znaczeniu (i *vice versa*).

Kripke miał świadomość problemu, jaki pojawia się, gdy chcemy wysłowić sceptyczną hipotezę:

Kiedykolwiek nasz oponent upiera się, iż nasze zwyczajne sposoby wyrażania się (np. „przejścia są już wyznaczone przez wzór”, „przyszłe użycie jest już obecne”) są całkowicie poprawne, możemy upierać się przy tym, że się z tym zgadzamy, o ile wyrażenia te właściwie rozumieć. Niebezpieczeństwo pojawia się, gdy próbujemy w sposób ścisły przedstawić, czym dokładnie jest to, czemu *zaprzeczamy* — *jaka* „błędna interpretację” zwykłego sposobu wyrażania się głosi nasz oponent. Możemy mieć trudności ze zrobieniem tego bez wypowiedzenia kolejnego zdania, które, trzeba przyznać, jest *nadal* „doskonale poprawne, jeśli je właściwie rozumieć” (Kripke 1982/2007: 116).

Rozważmy pewien przykładowy banał na temat znaczenia: „poprawne użycie danego symbolu, to użycie zgodne z jego znaczeniem”. Można go rozumieć po platońsku, jako przekonanie, że istnieje pewien abstrakcyjny byt zwany „znaczeniem”, którego „treść” wyznacza poprawne zachowania. Można go również ujmować minimalistycznie, jako wyraz przekonania, iż pewne zdania przypisujące poprawność użyciom mogą być deflacyjnie prawdziwe oraz że logika naszego języka pozwala na opisanie użyć, którym ta deflacyjnie prawdziwa poprawność przysługuje, mianem „zgodnych ze znaczeniem”.

W ostateczności można stwierdzić, iż różnica między anty-redukcjonizmem a stanowiskiem minimalizmu sprowadza się do różnicy tonu. Obie strony tego sporu głoszą podobne tezy — że zdania o znaczeniu służą do opisywania faktów oraz że nie jest możliwe zredukowanie ich do zdań innego typu. Nonredukcyjniści przedstawiają owe tezy tonem obwieszczającym dokonanie metafizycznego odkrycie na temat głębokiej natury świata; radykalny minimalista z kolei przedstawia je trzeźwym tonem komunikującym pewien zwykły fakt dotyczący natury naszych gier językowych.

Jednak nawet i ta różnica nie jest oczywista. Peter Strawson (1983: 77–91) przyjmuje postawę, którą określa mianem „radosnego platonizmu”. Według niego, opis użycia języka musi odwoływać się do przedmiotów i własności, które nie mają charakteru naturalnego. Zdaniem tego filozofa, nie ma w tym nic niepokojącego ani filozoficznie podejrzanego, o ile nie przyjmujemy bezzasadnych (według Strawsona) naturalistycznych przesądów. Mówienie o nie-naturalnych przedmiotach i własnościach jest w kontekście znaczenia (podobnie jak w wielu innych) czymś zupełnie zwyczajnym. Nie jest potrzebna żadna szczególna filozoficzna teoria uzasadniająca mówienie o tego typu cechach i przedmiotach. Na pierwszy rzut oka trudno wskazać rzeczowe różnice pomiędzy „radosnym platonizmem” a minimalizmem — oba stanowiska dopuszczają możliwość mówienia o nie-naturalistycznie ujmowanych własnościach semantycznych, deprecjonując zarazem istotność zobowiązań ontologicznych tego typu wypowiedzi.

Wedle Boghossiana (1989: 547–549), który sam był nonredukcjonistą, zwolennik tego stanowiska musi odpowiedzieć na kilka pytań. Pierwsze z nich dotyczy istnienia relacji przyczynowej łączącej stany mentalne użytkownika z jego działaniem. Jeśli bowiem przyjmiemy, iż fakty, własności i relacje semantyczne tworzą pewnego rodzaju odrębną dziedzinę bytową, to trzeba w jakiś sposób wyjaśnić co, jeśli cokolwiek, łączy tę dziedzinę z dziedziną faktów fizycznych, do których zaliczają się fakty behawioralne.

Nonredukcjonista ma dwa wyjścia. Może postulować istnienie specjalnej relacji przyczynowej łączącej fakty semantyczne z ludzkimi zachowaniami lub też przeczyć, jakoby fakty semantyczne w ogóle wchodziły w relacje przyczynowe z zachowaniami i faktami fizycznymi. Pierwsze z tych rozwiązań wiąże się z odrzuceniem, wprowadzonej w *Rozdziale II*, zasady przyczynowego domknięcia, głoszącej, iż każde zdarzenie fizyczne ma dostateczną przyczynę fizyczną. Rozwiązanie to wydaje się całkowicie nieumotywowane — zmusza nas bowiem do rewizji naszego obrazu świata fizycznego w imię rozważań dotyczących zupełnie innych kwestii. To, czy dziedzina fizykalna jest domknięta przyczynowo czy nie, jest pytaniem dotyczącym filozofii fizyki i nie może być rozstrzygnięte na gruncie teorii znaczenia.

Pozostaje zatem odpowiedź epifenomenalna, polegająca na odmawianiu faktom semantycznym roli przyczynowej. Wydaje się ona jednak trudna do zaakceptowania. Całkowicie zasadny wydaje się zarzut, że postulowanie specjalnej dziedziny bytowej, która jest całkowicie nieefektywna przyczynowo, jest łamaniem zasady oszczędności ontologicznej.

Opcja epifenomenalna jest łatwiejsza do zaakceptowania, jeśli przyjmiemy minimalizm. Fakty semantyczne są tu rozumiane tylko i wyłącznie jako coś, dzięki czemu zdania o znaczeniu mogą być prawdziwe; zaś ich zdolność do prawdy uzasadniamy spełnianiem pewnych minimalnych warunków. Tak słabo rozumiane fakty nie muszą wchodzić w żadne relacje z niczym. Mówienie o tym, że istnieją,

jest uwarunkowane tylko tym, że tak działa logika naszego języka — niejako wypycha nas ona w dyskurs zobowiązań ontologicznych, ilekroć chcemy przypisać pewnym zdaniom wartość logiczną. Nie widać powodu, dla którego tego typu fakty miałyby mieć jakąkolwiek wartość eksplanacyjną. Byłaby to wręcz niepożądana konsekwencja, gdybyśmy musieli uznać, iż fakty, które wprowadzamy do ontologii tylko ze względu na chęć uniknięcia niespójności, przyczynowo wpływają na fakty fizyczne. Można zatem wskazać na pierwszą istotną przewagę minimalistycznej koncepcji faktów semantycznych nad nonredukcjonizmem: pozwala ona zaakceptować epifenomenalizm dotyczący treści stanów mentalnych jako pożądaną konsekwencję i pozwala zachować czysto fizyczny opis zachowania werbalnego, uwalniając nas od poszukiwania dziwnej quasi-przyczynowej relacji łączącej fakty fizyczne z нефizycznymi.

Kolejną różnicą między tymi dwoma stanowiskami jest zagadnienie wiedzy o treści własnych stanów mentalnych. Boghossian odnosi się do tej kwestii w następujący sposób:

(...) nie widzę żadnego powodu, dla którego zwolennik antyredukcjonizmu miałby mieć *szczególny* problem z samowiedzą. Jeśli o mnie chodzi, to nikt nie dysponuje zadowalającą teorią tego, w jaki sposób poznajemy własne myśli (Boghossian 1989: 548).

Taka deklaracja jest jednak nieuzasadniona. Wydaje się, iż zwolennik nonredukcjonizmu istotnie ma *szczególny* problem z wyjaśnieniem tego, jak możliwa jest wiedza na temat znaczeń, a co za tym idzie, również na temat treści własnych stanów mentalnych. Skoro tworzą one *szczególną*, niesprowadzalną do żadnej innej, dziedzinę faktów, to naturalnym wydaje się domniemanie, że do ich uchwycenia potrzebna jest pewna *szczególna* władza poznawcza. Nawet jeśli nie istnieje powszechnie akceptowane wyjaśnienie naszej zdolności do zdawania sprawy z treści naszych nastawień sądzeniowych, to wydaje się, iż twierdzenie, że mamy do czynienia ze *szczególnym* rodzajem faktami semantycznymi, raczej komplikuje obraz, niż wskazuje na drogę rozwiązania zagadki samowiedzy.

Zwolennik minimalizmu może odwołać się do kilku już istniejących teorii samowiedzy. Martin Kusch (2006: 211–219) sugeruje, iż teorię minimalistyczną można w łatwy sposób pogodzić z teorią samowiedzy zaproponowaną przez Crispina Wrighta. Zdaniem tego ostatniego autora, sądy na temat zgodności jakiegoś zachowania z regułą nie mają „epistemologii śledzącej” (Wright 2002: 129–130). Nie możemy zatem twierdzić, iż zdają one sprawę z jakiejś, istniejącej niezależnie od naszych przekonań, rzeczywistości. Zamiast tego powinniśmy przyjąć, iż prawdziwość zdań traktujących o znaczeniu i regułach jest zależna od sądu (*judgment-dependent*).

Pojęcie zależności od sądu Wright wprowadza na przykładzie zdań dotyczących barw, które, jego zdaniem, nie służą do opisu istniejącej niezależnie rzeczywistości. W tym przypadku, to nasze sądy wyznaczają ekstensję predykatów, których używamy. Według Wrighta, testem na to, czy dana dziedzina zawiera zdania zależne od sądu czy też nie, jest to, czy w jej ramach za prawdziwą można uznać następującą formułę:

$$C(x) \rightarrow ([x \text{ sądzi, że } P] \leftrightarrow P)$$

Gdzie x oznacza użytkownika języka, P sąd, zaś C jest zbiorem warunków. Teza ta głosi, iż jeśli podmiot spełnia odpowiednie warunki, to prawdziwość jego twierdzenia, że jakieś zdanie jest prawdą, jest równoważne prawdziwości tego zdania. Innymi słowy, podmiot spełniający określone warunki jest nieomylny co do sądów z danej dziedziny.

Rozważmy rzecz na przykładzie sądu na temat barw. Jeśli mówię, że mój zeszyt jest niebieski, to, w myśl tej koncepcji, o ile znajduję się w odpowiednich warunkach (mój zmysł wzroku oraz układ nerwowy musi sprawnie funkcjonować, oświetlenie musi być korzystne, powietrze przejrzyste etc.), to moje stwierdzenie będzie w nieunikniony sposób prawdziwe. Jest tak dlatego, że barwy są własnościami wtórnymi. Wright kontrastuje barwy z kształtami — jego zdaniem, w przypadku sądów o kształtach, fakt, iż podmiot znajduje się w optymalnych warunkach,

nie gwarantuje jeszcze, iż jego sąd będzie prawdziwy. Iluzje dotyczące kształtu czy długości mogą zachodzić również, gdy warunki do obserwacji są idealne (Wright 2002: 134).

Cokolwiek by nie sądzić o koncepcji Wrighta w kontekście sporu o jakości pierwotne i wtórne, można ją łatwo zastosować do problemu sądów o własnych stanach psychicznych. Czasami Wright pisze tak, jakby jego rozważania na temat barw miały jedynie przybliżyć czytelnikowi jego teorię poznawania stanów mentalnych (Wright 1989/2001: 191). Jak nietrudno się domyślić, Wrightowi chodzi o to, że prawdziwość zdań, w których zdajemy sprawę z własnych stanów mentalnych, jest zależna od sądu. Jeśli spełniam odpowiednie warunki, to moja wypowiedź na temat treści moich stanów mentalnych jest automatycznie prawdziwa. Dokonując takiej asercji, nie staram się zdać sprawy z istniejących uprzednio faktów, lecz raczej je konstytuuję.

Kluczowe znaczenie dla powodzenia tej koncepcji ma wyszczególnienie warunków *C*, których spełnienie ma gwarantować podmiotowi nieomylność w sprawach własnych stanów intencjonalnych. Jest bowiem oczywiste, iż człowiek może się czasem w takich kwestiach mylić.

Wright kładzie nacisk na to, że warunki *C* nie mogą być formułowane w sposób trywialny — to znaczy nie mogą mieć postaci: *x* ma takie cechy, jakie są potrzebne do tego, by zdobywać wiedzę na temat określonego zdania (Wright 2002: 131). Rozważmy dyskurs o przyszłych wynikach mistrzostw świata w piłkę nożną: jeśli sformulowalibyśmy warunki *C* w następujący sposób: „podmiot znajduje się w odpowiednich warunkach, jeśli doskonale przewiduje wyniki przyszłych mistrzostw świata w piłkę”, to moglibyśmy stwierdzić, że zdania z tego dyskursu też są zależne od sądu. Konsekwencje tego typu ujęcia są jawnie absurdalne — w ten sposób możliwe byłoby podanie warunków *C* dla każdego dyskursu. Dlatego muszą one być sformułowane w taki sposób, by nie zawierały odniesienia do wiedzy o faktach opisywanych przez dany dyskurs. Takie niekolidujące sformułowanie warunków *C* jest

możliwe w przypadku np. zdań o barwach — podmiot znajduje się w odpowiednich warunkach, jeśli posiada odpowiednio zbudowany zmysł wzroku i układ nerwowy. Opis działania tych układów można podać w terminach czysto neurobiologicznych, bez odniesienia do zdolności widzenia.

Czy w przypadku zdań o znaczeniu możemy sformułować warunki *C* tak, by nie zawierały odniesienia do wiedzy na temat znaczeń? Wright (2002: 136) ma świadomość, iż jest to niemożliwe — jeśli chcielibyśmy sformułować warunki *C* dla zdań o własnych stanach mentalnych, to musielibyśmy zastrzec w nich, że podmiot nie oszukuje sam siebie. Jednak wymóg ten jest jawnie kolisty — zawiera bowiem odniesienie do samowiedzy. Mimo to, wedle Wrighta, takie sformułowanie warunków *C* pozwala na poczynienie istotnej obserwacji na temat zdań o znaczeniu: gdy mamy do czynienia z osobą dokonującą szczerego wyznania, to istnieje domniemanie, iż wyznanie to będzie prawdziwe. Aby móc zanegować takie wyznanie, musimy dysponować jakąś racją na rzecz przypuszczenia, że nasz interlokutor sam siebie oszukuje, jest niespełna rozumu etc. To domniemanie prawdziwości, w myśl omawianego autora, daje się wyjaśnić jedynie, jeśli przyjmiemy, że zdania o znaczeniu są zależne od sądu.

Teoria Wrighta rodzi wiele pytań i została tu przedstawiona z konieczności skrótowo. Istotne jest jednak to, że wskazuje ona na możliwość stworzenia takiej teorii samowiedzy, która będzie do pogodzenia z nieredukcyjnym charakterem faktów semantycznych i, w konsekwencji, nie będzie postulować specjalnej władzy poznawczej. Kluczowa dla tej teorii idea, iż wiedza na temat treści własnych stanów mentalnych nie jest wynikiem odkrywania niezależnej rzeczywistości, lecz raczej konstituowania odpowiednich stanów rzeczy przez sam fakt ich stwierdzania zdaje się współgrać z minimalistycznym spojrzeniem na naturę tego, co odpowiada zdaniom o znaczeniu.

Inną teorią, która stara się wyjaśnić naszą zdolność do opisywania własnych stanów mentalnych, a przy tym jest do pogodzenia z minimalistycznym ujęciem

znaczenia, jest koncepcja Davidsona. Zdaniem tego autora, podstawowym faktem, z jakiego powinna zdać sprawę teoria samowiedzy jest to, iż zazwyczaj przypisujemy podmiotowi uprzywilejowany dostęp poznawczy do własnych stanów psychicznych. Powinna ona jednak wyjaśniać ten fakt bez odwoływania się do szczególnej władzy introspekcji i tej podobnych (Davidson 1984/2001: 3–5). Zadanie to staje się trudniejsze, gdy, jak Davidson akceptuje się eksternalistyczne tezy wypowiedziane przez Putnama i Burge’a o zależności naszych treści mentalnych od czynników zewnętrznych — czy to mających charakter społeczny, czy to naturalny, takich jak istnienie rodzajów naturalnych (Davidson 1987/2001: 29).

Na czym, w takiej sytuacji, może polegać uprzywilejowany dostęp do stanów mentalnych? Przypuśćmy, że spotykamy osobę — choćby samego Davidsona — która twierdzi (przykładowo), iż „Wagner umarł szczęśliwy” (Davidson 1984/2001: 12). Załóżmy teraz, iż ktoś inny chce przytoczyć jego słowa i mówi „Davidson powiedział, że Wagner umarł szczęśliwy”. Pojawia się w tym momencie pytanie, jak należy rozumieć słowo „Wagner”. Interlokutor Davidsona mógł, w narzucający się sposób, zrozumieć to słowo jako odnoszące się do znanego niemieckiego kompozytora. Z kolei dla samego Davidsona słowo to mogło odnosić się przede wszystkim do noszącego to nazwisko sąsiada.

Jeśli przyjmie się eksternalizm, to nie można twierdzić, że ktokolwiek, nawet mówiący, jest nieomylny w kwestii tego, do czego odnoszą się jego słowa. Zapytany o to, do kogo odnosi się słowo „Wagner”, Davidson mógłby odpowiedzieć, że chodzi mu o jakiegoś austriackiego kompozytora z XVIII wieku, jednak nie zmieniłoby to wcale faktu, że *tak naprawdę* odnosiłby się do rzeczywistego Ryszarda Wagnera.

Na czym zatem polega przewaga podmiotu nad jego interlokutorem? Sprawadza się ona do tego, że jeśli na pytanie: „Co masz na myśli mówiąc ‘Wagner umarł szczęśliwy’?” sam Davidson odpowie: „To, że Wagner umarł szczęśliwy”, to możemy uznać, że jego druga wypowiedź trafnie zdaje sprawę z jego intencji. Jeśli natomiast zadalibyśmy jego rozmówcy pytanie: „Co Davidson ma na myśli

mówiąc 'Wagner umarł szczęśliwy'?", to odpowiedź „To, że Wagner umarł szczęśliwy”, nie byłaby już satysfakcjonująca. Nie mamy bowiem pewności, iż Davidson i jego rozmówca rozumieją tak samo choćby słowo „Wagner”. I choć nie możemy Davidsonowi przypisać nieomyślności co do sensu tego słowa, to możemy założyć, iż w obu tych kontekstach użyje on tego słowa w tym samym znaczeniu.

Uprzywilejowany dostęp sprowadza się w myśl tej koncepcji do obserwacji, iż jeśli podmiot używa dwa razy tego samego zdania, to oba te użycia znaczą to samo — w związku z tym może podać on deflacyjne warunki prawdziwości dla własnych zdań bez ryzyka pomyłki. Inny podmiot zaś może się pomylić nawet w takiej sytuacji. Taka koncepcja samowiedzy, podobnie jak ta zaproponowana przez Wrighta, współgra z minimalizmem. Mój autorytet co do treści własnych stanów mentalnych ogranicza się tylko do możliwości bezbłędnego podawania deflacyjnych warunków prawdziwości dla zdań wyrażających moje przekonania.

Rzecz jasna, samo przyjęcie radykalnego minimalizmu w sprawie znaczenia nie zobowiązuje jeszcze do akceptacji teorii Wrighta czy Davidsona. Przywołane zostały one tutaj by pokazać, iż możliwe jest takie wyjaśnienie fenomenu samowiedzy nie odwołujące się do szczególnej władzy introspekcji, lecz do mechanizmów językowych. To, że podmiotowi przysługuje autorytet w kwestii jego stanów psychicznych, jest, w myśl tych teorii, po prostu kwestią gramatyki wypowiedzi w pierwszej osobie liczbie pojedynczej. Takie podejście współgra z minimalistycznym przekonaniem, że fakty, do których odnosi nas dyskurs semantyczny to „fakty prawdziwe tak po prostu” w sensie Dummetta, których nie można w żadnym interesującym sensie „odkrywać”.

Ostatnią cechą różnicującą minimalizm i nonredukcjonizm jest kwestia społecznego charakteru języka. Zwolennik tezy, iż dyskurs o znaczeniu odnosi się do pewnych pierwotnych i nieredukowalnych bytów nie będzie miał powodu, by odrzucać tezę o możliwości języka prywatnego, to jest takiego, który jest zrozumiały wyłącznie dla jego jedyne go użytkownika. Skoro np. fakt, iż moje słowa odnoszą się

do jakiejś rzeczy, będziemy wyjaśniać istnieniem pewnej pierwotnej, tajemniczej relacji, w której pozostaję z owym przedmiotem, to nie jest potrzebne odwoływanie się do faktu, iż jestem członkiem wspólnoty językowej. Nonredukcjonizm rehabilituje ideę, że można mówić o znaczeniu w całkowicie indywidualistycznym schemacie pojęciowym.

Z kolei w schemacie pojęciowym radykalnego minimalizmu znajduje się miejsce dla mówienia o odniesieniu do wspólnoty jako warunku sensowności dyskursu semantycznego. Nie polega ono jednak na twierdzeniu, iż fakty semantyczne redukują się do faktów społecznych. Odwołanie się do zgody społecznej jest natomiast konieczne przy podawaniu warunków zdolności do prawdy dla zdań o znaczeniu. Tylko fakt, iż jesteśmy członkami wspólnoty, która jest w zasadzie zgodna, na poziomie czysto behawioralnym, w swoich zachowaniach językowych pozwala stwierdzić, iż dyskurs semantyczny jest dyskursem zdyscyplinowanym. Zatem istnienie wspólnoty językowej jest koniecznym warunkiem do uznania dyskursu o znaczeniu za zdolny do prawdy. Widać tu wyraźnie, jak istotne znaczenie teoretyczne ma odróżnienie warunków zdolności do prawdy od warunków prawdziwości. Pozwala nam ono uniknąć opisanych w poprzednim rozdziale groźnych konsekwencji utożsamienia znaczenia z dyspozycjami społecznymi. Z kolei uznanie, że istnienie wspólnoty językowej jest koniecznym warunkiem tego, by zdania semantyczne w ogóle miały wartość logiczną wyklucza możliwość istnienia języka prywatnego.

Można zatem wskazać na trzy różnice pomiędzy minimalizmem a tradycyjnymi formami nonredukcjonizmu semantycznego. Po pierwsze, minimalista uzna to, że fakty semantyczne nie są relewantne w wyjaśnieniach kauzalnych — nie musi postulować żadnej quasi-przyczynowej relacji łączącej treściowe stany psychiczne z zachowaniami. Po drugie, minimalista przy wyjaśnianiu fenomenu samowiedzy nie będzie odwoływać się do specjalnej władzy poznawczej, lecz raczej do tego, jak funkcjonuje dyskurs, w którym mówimy o własnych stanach mentalnych. Po trzecie

wreszcie, zwolennik minimalizmu będzie uwzględniał konsensus społeczności jako konieczny warunek zdolności do prawdy zdań o znaczeniu.

Rozdział VI

Zakończenie: Konsekwencje minimalizmu

Minimalizm pozwala na częściową akceptację sceptycznego rozumowania i na uznanie, że nie ma żadnej interesującej z filozoficznego punktu widzenia odpowiedzi na pytanie o to, co czyni prawdziwymi zdania o znaczeniu. Unika on zarazem paradoksalnych konsekwencji nonfaktualizmu przyznając, iż przynajmniej niektóre zdania o znaczeniu mogą być deflacyjnie fałszywe (lub prawdziwe), co uniemożliwia zastosowanie zabójczego dla nonfaktualizmu argumentu z uogólnienia. Można pokusić się o stwierdzenie, że minimalizm, stanowi zniesienie — w quasi-heglowskim sensie — paradoksu sceptycznego. Pozwala na zachowanie tego, co w nonfaktualizmie cenne — przekonania, że nie ma żadnych szczególnych czynników, które czynią zdania o znaczeniu prawdziwymi — przy zanegowaniu niepożądanych elementów tej doktryny.

Jakie konsekwencje dla innych obszarów filozofii niesie za sobą przyjęcie minimalizmu w kwestii znaczenia? Rozważmy to pytanie w odniesieniu do trzech problemów filozoficznych. Na początek zajmijmy się zagadnieniem prawdy. Może się wydawać, iż przyjęte rozwiązanie będzie zobowiązywać nas do akceptacji

deflacyjnej teorii prawdy. Minimalizm semantyczny opiera się wszak na założeniu, że zdania o znaczeniu mają jedynie deflacyjne warunki prawdziwości.

Rozważania zawarte w ostatnim rozdziale pracy pokazują jednak jasno, iż argument z uogólnienia nie działa w przypadku teorii deflacyjnej. Inaczej mówiąc, deflacionizm nie „rozlewa się” na inne dziedziny dyskursu. O ile minimalista musi przyjąć, iż zdania z dyskursu semantycznego mogą być jedynie deflacyjnie prawdziwe, o tyle poza dziedziną znaczenia, może on pozostać zwolennikiem korespondencyjnej, koherencyjnej albo pragmatystycznej teorii prawdy. Może też twierdzić, jak to czynił Wright (1992: 38), że różne dyskursy wymagają przyjęcia różnych teorii prawdy: w przypadku dyskursu etycznego możemy mieć do czynienia z inną „prawdziwością” niż w przypadku zdań o makroskopowych obiektach fizycznych.

Radykalny minimalizm zmusza nas jednak do przyznania, iż deflacionizm może być poglądem przynajmniej lokalnie prawdziwym. Jeśli chcemy zatem uznawać jakąś wersję substancjalnej teorii prawdy w innych dziedzinach, to musimy odrzucić tezę Boghossiana (1990: 165), iż konieczne jest dokonanie wyboru między substancjalnymi a deflacyjnymi teoriami. Taki wybór możemy uznać za konieczny tylko lokalnie, w odniesieniu do konkretnej dziedziny dyskursu. Chęć utrzymania którejś z substancjalnych teorii prawdy w odniesieniu do innych dyskursów niż semantyczny zmusza nas do przyjęcia pluralizmu prawdziwościowego — to znaczy przekonania, iż słowo „prawdziwy” może różnie funkcjonować w różnych kontekstach. Prezentowana tu koncepcja różni się więc znacząco od globalnego minimalizmu, który promował w swoich dziełach chociażby Horwich (por. np. 2010: v-vi). Nie twierdzi się tutaj jednak, iż jakiegokolwiek formy substancjalnej teorii prawdy są skazane na niepowodzenie.

Kolejne pytanie, jakie rodzi zaproponowane rozwiązanie, jest następujące: zdaniem Michaela Dummetta, rozstrzygnięcia na gruncie teorii znaczenia pozwalają odpowiedzieć na pytanie o realizm — w tym sensie, jakie temu pojęciu nadał

ów autor. Według Dummetta (1991/1995: 21–23), podstawową różnicą między realistą a antyrealistą jest stosunek do zasady wyłączonego środka. Realista będzie ją uznawał za obowiązującą zawsze, zaś antyrealista będzie chciał ją odrzucić, motywując to twierdzeniem, że nie każde zdanie posiada określoną wartość logiczną. Kluczem do rozwikłania tej zagadki ma być rozstrzygnięcie czy znaczenie zdań określa warunki ich prawdziwości, czy też warunki słusznej stwierdzalności.

Minimalistyczne podejście do kwestii znaczenia pozostawia to pytanie bez odpowiedzi. Rozważmy choćby zdania o odległej przeszłości — np. „Mieszko I w dniu swoich 30 urodzin zachorował na gripę”. Nie mamy żadnej metody weryfikacji takiego zdania. Antyrealista będzie twierdził, iż jest ono pozbawione wartości logicznej, realista uzna zaś, iż brak metody weryfikacji nie zmienia faktu, że zdanie to jest prawdziwe lub fałszywe. Nie widać powodu, dla którego minimalista musiałby się odpowiadać po którejś ze stron tego sporu. Do rozwiązania problemu sceptycznego potrzebne jest jedynie założenie, iż pewnym zdaniom przysługują deflacyjne warunki prawdziwości. Nie rozstrzyga się tu natomiast o tym, czy takie warunki przysługują wszystkim zdaniom. Nawet jeśli uznamy, że dyskurs historyczny globalnie zawiera zdania zdolne do prawdy, to nie wyklucza to tego, że niektóre zdania mogą mieć status „luk prawdziwościowych”. Taka konkluzja wskazuje jednak na to, iż wnioski na temat znaczenia, jakie możemy wyprowadzić z rozważań nad sceptycznym paradoksem, są względnie ubogie. Nie pozwalają one na udzielenie odpowiedzi na pytania, które, przynajmniej wedle niektórych, miała rozstrzygać teoria znaczenia.

Jak to już jednak było zaznaczane w *Rozdziale I*, przedmiotem zainteresowania niniejszej pracy był raczej problem realizmu w znaczeniu egzystencjalnym, aniżeli kwestia podnoszona przez Dummetta. Sednem egzystencjalnie rozumianego problemu antyrealizmu jest natomiast pytanie o to, w jakich sytuacjach zdania, których używamy, odnoszą nas do obiektywnie istniejącej rzeczywistości, a kiedy tylko sprawiają takie wrażenie. Idea, iż można dokonać rozróżnienia dyskursów

„rzetelnych” od „pozornych” wydaje się towarzyszyć filozofii analitycznej niemal od samego jej początku.

Już w neopozytywizmie widoczna była chęć wprowadzenia takiego podziału. Podnoszono wówczas problem kryterium demarkacji, czyli odróżnienia zdań sensownych od nonsensownych. Podstawą takiego rozróżnienia miało być to, czy zdania danego typu są weryfikowalne (por. np. Ayer 1936/1952: 35). To właśnie nieweryfikowalność zdań z zakresu etyki, estetyki czy metafizyki sprawiać miała według neopozytywistów, iż są one pozbawione sensu — to znaczy, że nie można ich traktować jako zdań wyrażających sądy, które mogą być prawdziwe albo fałszywe.

Odrzucenie kryterium weryfikowalności w późniejszej filozofii analitycznej nie doprowadziło jednak do całkowitego wycofania się z pomysłu odróżniania od siebie dyskursów, które są zdolne do opisywania rzeczywistości od tych, które robią to tylko pozornie. Debaty na temat tego, czy zdania z zakresu etyki albo psychologii potocznej odnoszą się do faktów, trwają niemal od początku dwudziestego wieku po dziś dzień. Próba zidentyfikowania kryterium takiego podziału na poziomie całkowicie ogólnym była przywoływana w *Rozdziale III* propozycja Shoemakera, wedle którego kryterium bycia realną własnością stanowić miało posiadanie mocy przyczynowej. Konsekwencją tego podejścia jest wniosek, iż dyskursy, w których chcielibyśmy odnosić się do faktów, a które nie są potrzebne w wyjaśnieniach przyczynowych, nie są dyskursami faktualnymi.

Nie zawsze jednak tego typu projekty opierają się na wyrażonych wprost założeniach. Częściej propozycje uznania pewnego rodzaju dyskursów za nieodnoszące się do rzeczywistości opierają się *implicite* na scjentyzmie, czyli na przekonaniu, iż jedynymi bytami, którymi powinniśmy się zajmować w naszej ontologii, są te, których istnienie zakłada współczesna nauka; jeśli jakichś bytów czy własności nie ma w tym obrazie świata, to ich nie ma wcale. Oczywiście scjentyzm nie jest jedyną możliwą podstawą do podejmowania prób selekcji dyskursów — wydaje się jednak najczęściej przyjmowaną przesłanką.

Przyjęcie radykalnego minimalizmu w sprawie statusu zdań o znaczeniu zmusza do ponownej refleksji nad sensownością przeprowadzania tego typu podziałów w ogóle. Okazuje się, że w szczególnym przypadku dyskursu semantycznego zmuszeni jesteśmy do przyznania, iż, z jednej strony, na fakty rzekomo odpowiadające zdaniom z tego dyskursu nie ma miejsca w naszej ontologii, z drugiej zaś musimy — pod groźbą popadnięcia w sprzeczność — uznawać, że zdania te do czegoś się odnoszą.

Stajemy więc przed pytaniem o to, co jest dostateczną racją do uznawania jakiegoś dyskursu za faktualny. Wydaje się, iż w większości przypadków próbowano rozstrzygnąć ten problem przyjmując — nie zawsze *explicite* — pewne ogólne założenia natury metafizycznej, a następnie sprawdzając, czy zgodnie z tymi założeniami dany dyskurs można uznać za odnoszący się do faktów. Rozważania dotyczące dyskursu semantycznego stawiają jednak takie podejście pod znakiem zapytania. Wychodząc z niemal dowolnych przesłanek natury metafizycznej, możemy bowiem dojść do wniosku, iż żadnych faktów semantycznych po prostu nie ma.

Ta konkluzja jest niemożliwa do zaakceptowania, co wskazuje na to, że musimy odrzucić prowadzącą do niej metodę rozstrzygania o faktulaności dyskursów. Minimalizm tym się różni od takiego podejścia, że wychodzi się w nim od strukturalnych własności dyskursów, a nie od apriorycznych przesłanek natury metafizycznej. W schemacie pojęciowym minimalizmu nie ma przesłanek, które uzasadniałyby podważanie zdolności jakiegokolwiek dyskursu do odnoszenia się do faktów — o ile tylko rzeczywiście się nim posługujemy. Jakkolwiek dziwne nie byłyby fakty, do których odnosimy się za pomocą zdań należących do jakiegoś dyskursu, musimy je włączyć do naszej ontologii.

Jak to już było wspomинane, nie można dokonać uogólnienia deflacyjnej teorii prawdy na inne dyskursy. Minimalista nie ma jednak powodu, by ograniczać swoją teorię zdolności do prawdy do zdań z dyskursu semantycznego. Można nawet

argumentować, iż jeśli przyjmujemy minimalizm w kwestii jednego dyskursu, to musimy tak postąpić w odniesieniu do wszystkich pozostałych. Jeśli uznajemy, że zdania z dyskursu semantycznego są zdolne do prawdy tylko dzięki temu, iż spełniają warunki składni i dyscypliny, to nie mamy prawa odmawiać tej zdolności spełniającym te warunki zdaniom z innych dyskursów. Gdybyśmy tak postąpili, minimalizm okazałby się rozwiązaniem o charakterze *ad hoc* — wprowadzilibyśmy minimalizm tylko po to, aby uniknąć niespójności nonfaktualizmu w sprawie znaczenia.

Czy jednak można uznać globalny minimalizm w kwestii zdolności do prawdy za przekonujące stanowisko? Czy jest do przyjęcia teza, że do podjęcia sensownych rozważań na temat wartości logicznej dowolnego zdania oznajmującego wystarczyć spełnienie warunków poprawności składni i przynależności do zdyscyplinowanego dyskursu?

Akceptacja takiej wizji zależy od przyjętych założeń natury metafizycznej. Jeśli uznajemy, iż zadaniem filozofii jest krytyczna refleksja nad obrazem świata sugerowanym przez zdrowy rozsądek i refutacja niektórych założeń przyjmowanych na poziomie pre-filozoficznym, to minimalizm wyda się rozwiązaniem niesatysfakcjonującym. Uniemożliwia on bowiem uznanie przekonania kompetentnego użytkownika języka, że jego wypowiedź opisuje obiektywną rzeczywistość, za mylne.

Minimalizm współgra za to z koncepcją filozofii, którą można by nazwać „kwietystyczną”. W tej wizji rolą filozofii nie jest dzielenie dyskursów na „porządne” i „nieporządne”, lecz zdawanie sprawy z tego, jak używany jest język. W tym właśnie duchu wypowiadał się Wittgenstein:

Filozofia nie może w żaden sposób naruszać faktycznego użycia języka, a więc może je w końcu tylko opisywać.

Albowiem nie może ona go też uzasadniać.

Zostawia ona wszystko tak, jak jest (Wittgenstein 1953/2000 §

Prezentowane tu podejście minimalistyczne nie jest jednak całkowicie kwietystyczne. Umożliwia ono do pewnego stopnia różnicowanie dyskursów — podstawowym narzędziem takiego podziału jest przypisywanie substancjalnych i deflacyjnych warunków prawdziwości. Zdania należące do dyskursów, które chcielibyśmy określić jako „rzetelne”, to zdania, którym przysługują warunki substancjalne. Za kryterium rozróżniania dyskursów można również przyjąć zdolność do prawdy w sensie korespondencyjnym; klasa zdań, które spełniają ten warunek, może być jeszcze węższa niż tych, które mają substancjalne warunki prawdziwości. Zdania z zakresu etyki mogą być uznane za np. za prawdziwe substancjalnie, ale nie korespondencyjnie. Mielibyśmy zatem do czynienia z następującą gradacją: zdania zdolne do prawdy, zdania zdolne do prawdy substancjalnej oraz zdania zdolne do prawdy korespondencyjnej. W przypadku każdego z dyskursów można zapytać o to, do której kategorii należy, można też zastanawiać się nad ogólnymi kryteriami tego podziału. Podobny projekt został zarysowany przez Wrighta w *Truth and objectivity*. Brytyjski filozof akceptuje tezę głoszącą, iż wszystkie dyskursy są zdolne do prawdy, jednak (chcąc uniknąć kwietyzmu) wprowadza pewne rozróżnienia dotyczące ich „realności”. Rozważania zawarte w mojej pracy są w dużej mierze komplementarne w stosunku do zawartych w *Truth and objectivity* — można wręcz powiedzieć, że celem niniejszej pracy było uzasadnienie tezy, którą Wright przyjmuje za wyjściową dla swoich rozważań.

Podstawowy wniosek tej książki jest zatem następujący: minimalizm w sprawie zdolności do prawdy jest stanowiskiem, którego akceptacja jest konieczna, jeśli chcemy — z jednej strony — wziąć pod uwagę sceptyczną krytykę faktualizmu semantycznego, z drugiej zaś, uniknąć paradoksalnych konsekwencji nonfaktualizmu. Przyjęcie minimalizmu prowadzi nas jednak do wniosku, że zdolność opisywania rzeczywistości jest powszechną cechą poprawnych zdań oznajmujących.

Bibliografia

- [–] AJDUKIEWICZ K.
(1934/1985) Sprache und Sinn. „Erkenntnis” IV: 100–138 [Język i znaczenie
tłum. Zeidler F. [w:] Język i poznanie. Tom. I. Warszawa: PWN]
- [–] ALEXANDER S.
(1920) Space, Time and Deity. London: Macmillan
- [–] ARMSTRONG D. M.
(1997) A world of states of affairs. Cambridge: Cambridge UP
- [–] AUSTIN J.L.
(1962/1993) How to do things with words. Oxford: Clarendon Press. [Jak
działać słowami? [w:] Mówienie i poznawanie. Tłum. B. Chwedeńczuk.
Warszawa: PWN]
- [–] AYER A. J.
(1936/1952) Language, truth and logic. London: Victor Gollancz Ltd [New
York: Dover]
- [–] BLACKBURN S.
(1984) The individual strikes back. „Synthese” 58: 281–301

[–] BOGHOSSIAN P.

(1989) The rule-following considerations. „Mind” 98: 507–49

(1990) The status of content. „The Philosophical Review” 99: 157–184

[–] BURGE T.

(1979/1996) Individualism and the mental. „Midwest studiem in Philosophy” 4: 73–121. [[w:] The Twin Earth Chronicles. A. Pessin i S. Goldberg (red.) New York: Sharpe]

[–] BYRNE A.

(1996) On misinterpreting Kripke’s Wittgenstein. „Philosophy and Phenomenological Research” 56: 339–343

[–] CHALMERS D.

(2002/2008) Consciousness and its place in nature [w:] Philosophy of Mind. Classical and Contemporary Readings. D. Chalmers (red.) Oxford: Oxford OUP. [Świadomość i jej miejsce w naturze tłum. T. Ciecierski i R. Poczobut [w:] Analityczna metafizyka umysłu. M. Miłkowski i R. Poczobut (red.) Warszawa: IFiS PAN]

[–] CHRUDZIMSKI A.

(2003). Metafizyczny realizm. „Studia Philosophiae Christiane” 39. No 2: 97–120

[–] CHURCHLAND P.

(1981) Eliminative materialism and the propositional attitudes. “The Journal of Philosophy” 78: 67–90

[-] DAVIDSON D.

(1970/1992) Mental events. [w:] Experience and Theory. L. Foster i J. W. Swanson (red.) London: Duckworth. [Zdarzenia mentalne. tłum. T. Baszniak [w:] Eseje o prawdzie, języku i umyśle. Warszawa: PWN.]

(1984/2001) First person authority „Dialectica” 38: 101–111 [[w:] Subjective, intersubjective, objective. Oxford: Clarendon Press]

(1986/2005) Nice dearangment of epitaphs [w:] Philosophical Grounds of Language. R. Grandy i R. Werner (red.) Oxford: Oxford UP. [[w:] Truth, language and history. Oxford: Clarendon Press]

(1987/2001) Knowing One's Own Mind. „Proceedings and Addresses of the American Philosophical Association” 1987: 441–458 [Subjective, intersubjective, objective. Oxford: Clarendon Press]

(1994/2005) The social aspect of language [w:] The philosophy of Michael Dummett. B. McGuinness i G. Olivieri. (red.) Dodrecht: Kluwer [[w:] Truth, language and history. Oxford: Clarendon Press]

(1989/2005) James Joyce and Humpty Dumpty. „Proceedings of the Norwegian Academy of Science and Letters”: 54–66 [[w:] Truth, language and history. Oxford: Clarendon Press]

[-] DEVITT M.

(1990) Transcendentalism about content. „Pacific Philosophical Quaterly” 71: 247–263

[-] DEVITT M., REY G.

(1991) Transcending transcendentalism. „Pacific Philosophical Quaterly” 72: 87–100

[-] DEVITT M., STERELNY K.

(1999) Language and reality. Cambridge, Mass: The MIT Press

[–] DUMMETT M.

(1976/1992) Realism [w:] Truth and Other enigmas. Cambridge, Mass.: Harvard University Press [Realizm. tłum. T. Placek i P. Turnau „Principia” VI: 5–31]

(1976/1993) What is a theory of meaning? (II) [w:] Truth and Meaning, G. Evans and J. McDowell (red.) Oxford: Clarendon Press [[w:] The seas of language. Oxford: Oxford University Press]

(1991/1995) Logical basis of metaphysics. Cambridge, Mass.: Harvard University Press [Logiczna podstawa metafizyki. tłum. W. Sady. Warszawa: PWN]

[–] FINKELSTEIN D.

(2003). Wittgenstein on rules and platonism [w:] The new Wittgenstein A. Carry i R. Read (red.) London: Routledge

[–] FODOR J.

(1974/2008) Special sciences. „Synthese” 28: 97–115 [Nauki szczegółowe. tłum. M. Gokieli [w:] Analityczna metafizyka umysłu. M. Mikłowski i R. Poczobut. Warszawa: IFiS PAN]

[–] FREGE G.

(1882/1977) Über Sinn und Bedeutung. „Zeitschrift für Philosophie und philosophische Kritik” 1882: 22–50. [Sens i znaczenie [w:] Pisma semantyczne. tłum. B. Wolniewicz Warszawa: PWN]

[–] GOODMAN N.

(1949) On Likeness of Meaning. „Analysis” 10:1–7

(1955) Fact, Fiction and Forecast. Cambridge, Mass.: Harvard UP

- [–] GLOCK H.-J.
(2005) The normativity of meaning made simple [w:] *Philosophy — Science — Scientific Philosophy*. A. Beckermann, i C. Nimtz (red.) Paderborn: Mentis.
- [–] GLÜER K.
(2001) *Dreams and nightmares* [w:] *Interpreting Davidson*. P. Kotatko et al. (red.) Stanford: CLSI Publications
- [–] GLÜER K., PAGIN P.
(1999) Rules of meaning and practical reasoning. „*Synthese*“ 117: 207–227
- [–] GŁÓD B.
(2003) Problem istnienia i poznawalności faktów semantycznych w świetle sceptycznej argumentacji Kripkego. „*Kwartalnik Filozoficzny*” XXXI, Z. 2: 85–106
- [–] GREEN K.
(2001) Davidson’s derangement: Of the conceptual priority of language. „*Dialectica*” 55: 239–258
- [–] GRICE H. P.
(1957). Meaning. „*The Philosophical Review*” 66: 377–388
(1975/2000) *Logic and Conversation* [w:] *The Logic of Grammar*. D. Davidson and G. Harman (red.) Encino, CA: Dickenson. [Perspectives in the philosophy of language. R. Stainton (red.) Ontario: Broadview]
- [–] HARMAN G.
(1967) Quine on Meaning and Existence I. „*Review of Metaphysics*” 21:1: 124–157

- [–] HATTIANGADI A.
(2006) Is meaning normative? „Mind and Language” 21: 220–40
(2007) Oughts and Thoughts. Oxford: Clarendon Press
- [–] HOLTON R.
(1993) Minimalisms about truth [w:] Themes from Wittgenstein. B. Garrett i K. Mulligan (red.) Canberra: ANU: 45–61
- [–] HORWICH P.
(2004) Realism and truth. “Philosophical Perspectives” 10: 187–97 [[w:] From a deflationary point of view. Oxford: Clarendon Press.]
(2010) Truth - meaning - reality. Oxford: Clarendon Press
- [–] HUME D.
(1739/1963) A Treatise of Human Nature. [Traktat o naturze ludzkiej. tłum. Cz. Znamierowski Warszawa: PWN]
- [–] JACKSON F.
(1982) Epiphenomenal Qualia. „The Philosophical Quarterly” 32: 127–136
- [–] JACKSON F., OPPY G., SMITH M.
(1994) Minimalism and truth aptness. „Mind” 103: 287–302
- [–] KALLESTRUP J.
(2006) The causal exclusion argument. „Philosophical Studies” 131:459–485
- [–] KANT I.
(1783/1993) Prolegomena zu einer jeden künftigen Metaphysik [Prolegomena do wszelkiej przyszłej metafizyki. tłum. B. Bornstein Warszawa: PWN]
(1787/2001) Kritik der reinen Vernunft [Krytyka czystego rozumu tłum. R. Ingarden Kęty: Antyk]

- (1785/2001) *Grundlegung zur Metaphysik der Sitten* [Uzasadnienie metafizyki moralności. tłum. M. Wartenberg Kęty: Antyk.]
- [–] KIM J.
- (1989/2008) *The Myth of Non-Reductive Materialism*. “Proceedings and Addresses of the American Philosophical Association” 63 (3): 31–47. [Mit nieredukcyjnego materializmu. tłum. P. Gutowski i T. Szubka [w:] *Analityczna metafizyka umysłu*. M. Miłkowski i R. Poczubut (red.) Warszawa: IFiS PAN]
- [–] KRIPKE S.
- (1972/1998) *Naming and Necessity* [w:] *Semantics of Natural Language*. D. Davidson i G. Harman (red.) New York: Humanities [Nazywanie a konieczność. tłum. B. Chwedeńczuk Warszawa: PwsAX]
- (1982/2007) *Wittgenstein on Rules and Private Language*. Harvard: Harvard UP. [Wittgenstein o regułach i języku prywatnym. tłum. K. Posłajko i L. Wroński, Warszawa: Aletheia]
- [–] KUSCH M.
- (2006) *A sceptical guide to meaning and rules*. Chesham: Acumen
- [–] LORMAND E.
- (1996) Nonphenomenal consciousness. „Noûs”. 30: 242–261
- [–] MACKIE J.
- (1977) *Ethics. Inventing right and wrong*. New York: Penguin
- [–] MADDY P.
- (1984) How the causal theorist follows a rule. “Midwest Studies in Philosophy” IX: 457–477.

- [–] MARTIN C. B., HEIL J.
(1998) Rules and powers. „Philosophical Perspectives” 12: 283–312
- [–] MARUSZEWSKI T.
(2001) Psychologia poznania. Gdańsk: GWP
- [–] MCDOWELL J.
(1984) Wittgenstein on following a rule. „Synthese” 58: 326–363.
- [–] MCGINN C.
(1987) Wittgenstein on meaning. Oxford: Basil Blackwell
- [–] MCLAUGHLIN B., BENNETT K.
(2008) Supervenience. [w:] The Stanford Encyclopedia of Philosophy (Fall 2008 Edition). Edward N. Zalta [red.] URL = <
<http://plato.stanford.edu/archives/fall2008/entries/supervenience/> >.
- [–] MILLAR A.
(2002) The normativity of meaning [w:] Logic, thought and language. A. O’Hear (red.) Cambridge: Cambridge University Press
- [–] MILLER A.
(2003) An introduction to contemporary metaethics. Cambridge: Polity
(2006) Meaning scepticism. [w:] The Blackwell guide to the philosophy of language. M. Devitt i R. Hanley (red.) Oxford: Blackwell
(2008) Realism. [w:] The Stanford Encyclopedia of Philosophy (Fall 2008 Edition) E. Zalta (red.) URL = <
<http://plato.stanford.edu/archives/fall2008/entries/realism/> >
(2009) Semantic realism and the argument from motivational internalism.
[w:] What is meaning? R. Schantz (red.) Berlin: De Gruyter

- (2010) Kripke's Wittgenstein, factualism, and meaning. [w:] The later Wittgenstein on language. D. Whiting (red.) Basingstoke: Palgrave
- [–] MILLIKAN R.
- (1990/2002) Truth rules, hoverflies and Kripke-Wittgenstein paradox. „The Philosophical Review” 99: 323–353. [[w:] Rule-following and meaning. A. Miller i C. Wright (red.) Chesham: Acumen]
- [–] NOWAK A.
- (2002) Świat człowieka. Kraków: Collegium Collumbinum.
- [–] OCKHAM W.
- (1971) Suma logiczna. tłum. T. Włodarczyk Warszawa: PWN
- [–] ODROWĄŻ-SYPNIEWSKA J.
- (2006) Rodzaje naturalne. Warszawa: Semper
- [–] PAPRZYCKA K.
- (2005) O możliwości antyredukjonizmu. Warszawa: Semper
- [–] PEIRCE C.S.
- (1931). Collected papers. T 1. C. Hartshorne i P. Weiss (red.). Cambridge Mass.: Harvard University Press
- [–] PITT D.
- (2004) The phenomenology of cognition or ”What is it like to think that P?” „Philosophy and Phenomenological Research” 69: 1–36
- [–] POLIT J.
- (2001) Daleki Wschód w latach siedemdziesiątych [w:] Najnowsza historia świata. Tom II. A. Patek et. al. (red.) Kraków: Wydawnictwo Literackie

[–] POPPER K.

(1972/1992) *Objective knowledge : An Evolutionary Approach*. Oxford: Clarendon Press [Wiedza obiektywna. tłum. A. Chmielewski. Warszawa: PWN]

[–] PUTNAM H.

(1975/1998) The meaning of 'Meaning'. „Minnesota Studies in the Philosophy of Science” 7: 131–193. [Znaczenie wyrazu „znaczenie”. tłum. A. Grobler [w:] *Wiele twarzy realizmu*. Warszawa: PWN]

(1982/1998) Why reason can't be naturalized. „Synthese” 52: 3–23 [Dlaczego rozumu nie można znaturalizować. tłum. A. Grobler [w:] *Wiele twarzy realizmu*. Warszawa: PWN]

(1981/1998) *Brains in a vat* [w:] *Reason, Truth and History*. Cambridge: Cambridge UP. [Mózgi w naczyniu. tłum. A. Grobler [w:] *Wiele twarzy realizmu*. Warszawa: PWN]

[–] QUINE WVO.

(1953/1969a) *On what there is* [w:] *From a logical point of view*. Cambridge, Mass.: Harvard University Press [O tym co istnieje. tłum. B. Stanosz [w:] *Z punktu widzenia logiki*. Warszawa: PWN]

(1953/1969b) *Two dogmas of empiricism* [w:] *From a logical point of view*. Cambridge, Mass.: Harvard University Press [Dwa dogmaty empiryzmu tłum. B. Stanosz [w:] *Z punktu widzenia logiki*. Warszawa: PWN]

(1960/1999) *Word and Object*. Cambridge, Mass.: The MIT Press, [Słowo i przedmiot. tłum. C. Cieśliński Warszawa: Aletheia]

(1970) On the reasons for indeterminacy of translation. „The Journal of Philosophy” 67: 178–183

[-] RORTY R.

(1965/1995) Mind-Body Identity, Privacy, and Categories. „Review of Metaphysics” 19: 24–54 [Identyczność ciała i umysłu. tłum. M. Szczubiałka [w:] Fragmenty filozofii analitycznej. Filozofia umysłu. B. Chwedeńczuk (red.) Warszawa: Aletheia]

[-] RUSSELL B.

(1905/1967) On Denoting. „Mind” 14: 479–493 [Denotowanie. tłum. J. Pelc [w:] Logika i język. J. Pelc (red.) Warszawa: PWN]

[-] RYLE G.

(1949/1970) The concept of mind. Chicago: Chicago UP [Czym jest umysł? tłum. W. Marciszewski Warszawa: PWN]

[-] SEARLE J.

(1965/1996) What is a speech act [w:] Philosophy in America. M. Black (red.) London: Allen & Unwin. [[w:] The Communication theory Reader. P. Cobley (red.) Oxford: Routledge]

(1971) Introduction [w:] The philosophy of language. J. Searle (red.). Oxford: Oxford UP.

(1987/1993) Indeterminacy, empiricism and first person. „Journal of Philosophy” LXXXIV No.3: 123–146 [Niezdeterminowanie, empiryzm i pierwsza osoba tłum. B. Stanosz [w:] Fragmenty filozofii analitycznej. Filozofia języka. B. Stanosz (red.) Warszawa: Aletheia]

[-] SHOEMAKER S.

(1980/1999) Causality and properties. [w:] Time and Cause. P. van Inwagen (red.) Dordrecht: D.Reidel [[w:] Metaphysics. An Anthology. J. Kim i E. Sosa (red.) Blackwell, Oxford]

[–] SMITH M.

(1994) *The Moral Problem* Oxford: Blackwell

[–] SOAMES S.

(1997) The truth about deflationism. „*Philosophical Issues*” 8: 1–44

(1998a) Facts, truth-conditions and the sceptical solution to the rule-following paradox. „*Philosophical perspectives*” 12: 314–348

(1998b) Skepticism about meaning: Indeterminacy, normativity and the rule-following paradox. „*Canadian Journal of Philosophy*” Supp. Vol. 23: 211–49

[–] STRAWSON G.

(2005) Intentionality and experience. [w:] *Phenomenology and Philosophy of Mind*. D. W. Smith i A. Thomasson (red.) Oxford: Oxford University Press

[–] STRAWSON P.

(1970/1993) *Meaning and truth*. Oxford: Oxford UP [Znaczenie i prawda. tłum. T. Szubka [w:] *Fragmenty filozofii analitycznej* Filozofia języka. B. Stanosz (red.) Warszawa: Aletheia]

(1983) *Skepticism and naturalism*. New York: Columbia UP.

[–] SZYMURA J.

(1998) Czy można pozbyć się pojęcia prawdy? „*Studia Semiotyczne*” XXI–XXII: 165–198

[–] SZUBKA T.

(1999) Prawda minimalna a realizm metafizyczny. „*Roczniki Filozoficzne*”. T. XLVII. Z. 2: 299–313

(2000) Minimalistyczne a korespondencyjne teorie prawdy [w:] *Filozofia i logika. W stronę Jana Woleńskiego*. J. Hartman (red.) Kraków: Aureus

[–] TARSKI A.

(1933/1995) Pojęcie prawdy w językach nauk dedukcyjnych. Warszawa: Towarzystwo Naukowe Warszawskie [[w:] Pisma logiczno-filozoficzne. Tom 1. Warszawa: PWN]

[–] VAN ROOJEN M.

(2008) Moral Cognitivism vs. Non-Cognitivism [w:] The Stanford Encyclopedia of Philosophy (Fall 2008 Edition), E. Zalta (red.) URL = <
<http://plato.stanford.edu/archives/fall2008/entries/moral-cognitivism/> >

[–] WHITING D.

(2007) The normativity of meaning defended. „Analysis” 67: 133–140

[–] WILSON G.

(1998) Semantic realism and Kripke’s Wittgenstein. „Philosophy and Phenomenological Research” 58: 99–122

[–] WIKFORSS Å.

(2001) Semantic normativity. „Philosophical Studies” 102: 203–226

[–] WITTGENSTEIN L.

(1953/2000) Philosophische Untersuchungen. Philosophical Investigations. Oxford: Basil Blackwell [Dociekania filozoficzne. tłum. B. Wolniewicz Warszawa: PWN]

(1958/1998) The Blue and Brown Books. Oxford: Blackwell [Niebieski i brązowy zeszyt. Tłum. A. Lipszyc i Ł. Sommer. Warszawa: Spacja]

[–] WRIGHT C.

(1989/2001) Wittgenstein’s rule-following considerations and the central project of theoretical linguistics [w:] A. George (red.) Reflections on Chomsky.

Oxford: Blackwell. [[w:] Rails to infinity. Cambridge Massachusetts: Harvard University Press]

(1989/2002) Critical notice of Colin McGinn's Wittgenstein on meaning. „Mind” 98: 289–305. [[w:] Rule-following and meaning A. Miller i C. Wright (red.) Chesham: Acumen]

(1992) Truth and Objectivity. Cambridge, Massachusetts: Harvard University Press

(1999/2000) Truth. A traditional debate revisited. [w:] C. Misak (red.) Pragmatism. Alberta: University of Calgary Press. [Prawda. Przegląd tradycyjnego sporu. tł. T. Szubka, „Kwartalnik Filozoficzny” XXXVIII: 153–198]

(2002) Meaning and intention as judgment dependent. [w:] Rule-following and meaning. A. Miller i C. Wright (red.) Chesham: Acumen

Indeks nazwisk

Ajdukiewicz K. 64
Alexander S. 92
Armstrong D. M. 99
Austin J.L. 28
Ayer A. J. 4–5, 42, 159, 184
Bennett K. 97
Blackburn S. 40, 42, 103–104
Boghossian P. 7–9, 38–40, 46–57, 59, 79, 132, 140, 142–144, 146, 153, 172–173, 182
Burge T. 52, 100, 177
Byrne A. 156, 160–161
Chalmers D. 123–124
Chrudzimski A. 129
Churchland P. 34–35, 159
Davidson D. 70–78, 90, 177–178
Devitt M. 114, 117, 143, 145, 153
Dummett M. 5, 160–161, 178, 182–183
Finkelstein D. 131, 134–135
Fodor J. 90
Frege G. 6, 132
Goodman N. 23, 29–32
Glock H.-J. 52–55, 57, 69
Glüer K. 61–64, 66, 68, 83

- Glód, B. 37
- Green K. 74
- Grice H. P. 37, 74
- Harman G. 32
- Hattiangadi A. 36, 42, 49, 57–59, 64–66, 80–81, 83, 87–88, 114, 120, 137
- Heil J. 86–87
- Holton R. 143–145, 167
- Horwich P. 182
- Hume D. 41, 145
- Jackson F. 124, 145, 162, 166
- Joyce J. 72, 75–76
- Kallestrup J. 90–92
- Kant I. 31–32, 58, 110
- Kim J. 89–90, 92, 94, 98, 124, 139
- Kusch M. 16, 107, 114, 119, 140, 146–147, 150, 153, 156, 174
- Leśmian B. 72
- Lormand E. 125–126
- Mackie J. 8, 58
- Maddy P. 107, 112, 114–115
- Martin C. B. 86–87
- Maruszewski T. 114–115
- Matejko J. 4
- McDowell J. 44, 46–47
- McGinn C. 15, 37
- McLaughlin B. 97
- Millar A. 52–55, 57
- Miller A. 7, 13, 42, 55, 59, 67–69, 145, 161–164, 168
- Millikan R. 118–121
- Nowak A. 13, 134–135
- Ockham W. 75, 93, 132–133
- Oppy G. 145, 162, 166

- Odrowąż-Sypniewska J. 114, 117
- Pagin P. 61–63, 66, 68, 83
- Paprzycka K. 7, 97–98
- Peirce C.S. 134–135
- Pitt D. 125–126
- Platon 92, 129
- Polit J. 121
- Popper K. 22
- Putnam H. 101–102, 104, 106–107, 110–111, 117, 177
- Quine WVO. 6, 31–35, 74, 108
- Russell B. 107–108
- Rey G. 143
- Ryle G. 24–25
- Searle J. 11, 34, 61, 108
- Shoemaker S. 92, 124, 184
- Sienkiewicz H. 108–109
- Smith M. 82, 145, 162, 166
- Soames S. 28–29, 97–98, 140, 156
- Sokrates 92
- Sterelny K. 114, 117
- Strawson G. 128
- Strawson P. 37, 171
- Szubka T. 13, 141, 169–170
- Szymura J. 13, 144, 146, 152, 154–155
- Tarski A. 140–141, 150
- van Roojen M. 42, 67
- Wagner R. 177–178
- Whiting D. 48–49, 65–66, 76
- Wikforss Å. 49, 61, 73
- Wilson G. 156, 159
- Woleński J. 13
- Wright C. 44, 140, 143, 146–147, 154–155, 162–163, 169, 174–176, 178, 182, 187



„PRIMA FACIE, PROBLEM ZNACZENIA BLISKI JEST TEMU, CO DOŚWIADCZAMY NIEMAL BEZ PRZERWY. CZYTAMY KSIĄŻKI, ROZMAWIAMY, A W KAŻDYM Z TYCH DZIAŁAŃ UŻYWAMY JĘZYKA. NA OGÓŁ ROZUMIEMY TO, CO JEST PRZEDMIOTEM NASZEJ LEKTURY, TO, CO SŁYSZYMY OD INNYCH, I TO, CO SAMI MÓWIMY (...). OKAZUJE SIĘ JEDNAK, ŻE TE PROSTE I OCZYWISTE KONTEKSTY NIEWIELE POMAGAJĄ W ODPOWIEDZI NA PYTANIE: CO TO JEST ZNACZENIE I O CZYM WŁAŚCIWIE MÓWIMY, GDY WYPOWIADAMY SIĘ O ZNACZENIU?”

JAN WOLEŃSKI



KRZYSZTOF POŚAJKO JEST PRACOWNIKIEM INSTYTUTU FILOZOFII UNIWERSYTETU JAGIELLOŃSKIEGO. ZAJMUJE SIĘ ANALITYCZNĄ FILOZOFIĄ JĘZYKA I ONTOLOGIĄ.

FOT. ALICJA DYBOWSKA