

Paweł Grabarczyk

Od wrażenia do wyrażenia, czyli o niezdeterminowaniu przekładu i względności odniesienia

Słowa kluczowe: Quine, niezdeterminowanie przekładu, niezdeterminowanie odniesienia, behawioralna teoria znaczenia

W artykule tym chciałbym przyjrzeć się dwóm słynnym tezom – tezie o niezdeterminowaniu przekładu i tezie o względności odniesienia. Wokół obu narosło wiele zaskakujących niekiedy nieporozumień. Nie ulega wątpliwości, że sam Quine nie ułatwił zadania swym czytelnikom – w jednym z miejsc sam przyznał, że rozróżnienie pomiędzy obiema tezami nigdy nie było dlań jasne¹. Co więcej, jak zobaczymy, przykłady ilustrujące jedną z tez, są zarazem przykładami ilustrującymi drugą. Prowokuje to pytanie: jaki jest właściwie ich związek? Nie ma też wśród filozofów zgody co do oceny ich konsekwencji. Jedni – choćby sam Quine – nie są skłonni odkryciu owych zjawisk (o ile można tak powiedzieć) przypisywać szczególnie druzgocących konsekwencji. Inni – jak Rorty – widzą w nich upadek mitu odniesienia przedmiotowego, jeszcze inni – jak Davidson – przyjmują upadek mitu odniesienia, ale nie widzą w tym niczego druzgocącego. Jak zauważył Searle: „Uderzającą cechą całej dyskusji w tej sprawie jest pośpiech, z jakim szokujące konkluzje wyprowadzane są na podstawie kilku lakonicznych uwag i ledwie naszkicowanych przykładów”². Niektórzy byliby skłonni widzieć w obu tych tezach kolejne zwycięstwo sceptycyzmu – nie możemy już nigdy być pewni, o czym mówią nasi rozmówcy. Być może powinniśmy nawet pójść o krok dalej i powiedzieć, że nie wiadomo już, o czym my sami mówiliśmy, gdy zapisaaliśmy sobie coś na kartce dzień wcześniej?

¹ W.O. Quine, *Na tropach prawdy*, tłum. B. Stanosz, Warszawa 1997, s. 86.

² J.R. Searle, *Niezdeterminowanie, empiryzm i pierwsza osoba*, w: Barbara Stanosz (red. i tłum.) *Filozofia języka. Fragmenty filozofii analitycznej*, Warszawa, 1993, s. 158.

Będę starał się pokazać, że obie te tezy są, by sparafrazować Quine'a, dwiema stronami tego samego medalu³, oraz że nie prowadzą one do tak daleko posuniętego sceptycyzmu.

Zacznijmy od przykładu, który ułatwi nam dostrzeżenie różnicy pomiędzy niezdeterminowaniem przekładu a czymś, co moglibyśmy określić różnicą czysto werbalną. Wyobraźmy sobie następującą sytuację. Załóżmy, że my i nasz towarzysz znajdujemy się w jakimś pomieszczeniu – naszym zadaniem jest opisanie tego, co w tym pomieszczeniu widzimy. Przyjmijmy, że nasz opis składa się z następujących zdań:

- 1 Wazon jest na obrusie
- 2 Obrus jest na stole
- 3 Stół jest na dywanie
- 4 Dywan jest pod stołem
- 5 Stół jest pod obrusem
- 6 Obrus jest pod wazonem

Zadaniem naszego towarzysza jest dokładnie to samo – ma opisać, co widzi w pokoju. Okazuje się jednak, że, ku naszemu zdumieniu, tworzy on następującą listę zdań:

- 1' Wazon jest pod obrusem
- 2' Obrus jest pod stołem
- 3' Stół jest pod dywanem
- 4' Dywan jest na stole
- 5' Stół jest na obrusie
- 6' Obrus jest na wazonie

Wyjaśnieniem, które natychmiast się tu nasuwa, jest przyjęcie, że nasz towarzysz używa słów „na” i „pod” dokładnie odwrotnie niż my. Jest to różnica czysto werbalna, istnieją po prostu dwa równobrzmiące idiolekty, w których dwa słowa mają zamienione znaczenia. Zdanie 1 z naszego idiolektu ma to samo znaczenie, co zdanie 1' w idiolekcie naszego towarzysza, zdanie 2 ma to samo znaczenie, co zdanie 2' w idiolekcie naszego towarzysza itd. Nie będziemy, jak sądzę, przypisywać temu przypadkowi żadnej filozoficznej głębi. Quine rozpatruje podobny przypadek, podając przykład kogoś, kto używa terminu *alternatywa*, tak jak my używamy terminu *koniunkcja* – i odwrotnie, również nie przypisując temu przypadkowi większego znaczenia⁴.

³ Ale wcale nie podejrzanego.

⁴ W.O. Quine, *Filozofia logiki*, tłum. B. Stanosz, Warszawa 2002, s. 152.

Zauważmy jednak, że w grę wchodzi jeszcze jedna możliwość. Może się okazać, że nasz towarzysz używa wyrażen „na” i „pod” w taki sam sposób jak my. Okazuje się bowiem, że różnice pojawiają się w innych partiach jego idiolektu – wystarczy, że używa słowa „wazon” tak, jak my używamy słowa „dywan”, a słowa „obrus” tak, jak my używamy słowa „stół”. Przekładem naszego zdania 1 będzie wtedy zdanie 4’, zdania 2 zdanie 5’, 3–6’, 4–1’, 5–2’, a 6–3’. Czy ten drugi przypadek należy uznać za jakościowo różny względem poprzedniego? Nie sądzę – mam wrażenie, że to znów jedynie różnica czysto werbalna, nie pociągająca za sobą żadnych interesujących filozoficznie konsekwencji. A może samo istnienie dwóch konkurencyjnych przekładów moglibyśmy nazwać niezdeterminowaniem przekładu? Sądzę, że i tę możliwość należy odrzucić – o konkurencyjnych przekładach mówimy tu tylko dlatego, że rozpatrujemy pojedynczą, dość sztuczną sytuację. Wystarczyłoby dokonać kilku dodatkowych obserwacji, by ustalić, z którą sytuacją mamy do czynienia – czyli ustalić, który z przekładów jest tym właściwym. Gdyby niezdeterminowanie przekładu miało polegać właśnie na tym, to nikt nie zaprzętałby nim sobie głowy. Przyjrzymy się więc przykładom Quine’a i sprawdzmy, czym różnią się od podanej, nieinteresującej sytuacji.

W artykule „Ontological Relativity” Quine przywołuje ciekawy, bo nie wymyślony na potrzeby teorii, przykład. Oddala to podejrzenie, że niezdeterminowanie przekładu i względność odniesienia są jedynie wymysłem filozofów o bujnej wyobraźni. Okazuje się, że w języku japońskim występuje kategoria wyrażen klasyfikujących, które dołączają się do liczebników, tworząc w ten sposób złożone wyrażenie funkcjonujące jako liczebnik nadający się do liczenia przedmiotów danego rodzaju⁵. Przekładamy więc liczebnik japoński na liczebnik polski (w oryginale oczywiście występuje język angielski), rzeczownik japoński na rzeczownik polski, a klasyfikator pomijamy, jako kategorię w języku polskim niewystępującą. Moglibyśmy jednak postąpić inaczej – moglibyśmy uznać, że większość rzeczowników w języku japońskim to rzeczowniki masowe, a klasyfikator, którego poprzednio nie przekładaliśmy, odpowiada naszym słówkom indywidualizującym wyrażenia masowe, takim jak: *sztuka*, *kawałek* czy *kropla*. Spróbujmy przeanalizować ten przykład dokładniej. Jeśli dobrze go rozumiem, to spór dotyczyłby na przykład tego, czy dane wyrażenie japońskie przekładać należy zawsze na polskie „pięć bydląt” czy też na „pięć sztuk bydląt”. Istotne jest dla nas to, że różnica nie tkwi jedynie w stylistyce, wybór jednego z przekładów pociąga za sobą przypisanie Japończykom pewnej ontologii. Niezdeterminowanie przekładu owocuje więc wątpliwościami, co do odniesienia przedmiotowego – od decyzji twórcy podręcznika przekładu zależy nasz pogląd na ontologię języka japońskiego. W pierwszym przypadku przyjmujemy, że Japończycy, mówiąc o rogatych przeżuwaczach,

⁵ Czymś trochę podobnym byłyby wyrażenia „kopa” i „mendel” w języku polskim. Są to specjalne liczebniki, używane tylko do wybranych przedmiotów.

posługują się wyrażeniem o podzielonym odniesieniu, czyli przyjmują wielość przedmiotów należących do pewnego zbioru. W drugim przypadku zakładamy, że Japończycy traktują je jak pojedynczy przedmiot masowy, o którego częściach, zwanych „sztukami”, chcą niekiedy mówić. Względność odniesienia miałyby więc polegać na tym, że nie istnieje sposób na ustalenie właściwego przekładu, co pociąga za sobą niemożność odkrycia, do czego w istocie odnoszą się Japończycy. Żadna, najbardziej nawet skrupulatna obserwacja nie jest w stanie nas o tym pouczyć. Jest tak, ponieważ zachowania ludzi mówiących o pięciu sztukach bydła nie różnią się od zachowań ludzi mówiących o pięciu bydłach.

Wygląda więc na to, że względność odniesienia jest po prostu efektem niezeterminowania przekładu. Ale co jest powodem niezeterminowania przekładu? Przyjrzyjmy się klasycznemu przykładowi tego zjawiska – opowieści o przekładzie radykalnym⁶. Podobieństwo do naszych trudności z przekładem z języka japońskiego jest uderzające. Lingwista dokonujący przekładu radykalnego zbiera wszystkie dostępne mu dane empiryczne (część z nich to opis zachowań werbalnych tubylców, część to opis rzeczywistości, którą wraz z nimi obserwuje), a następnie dobiera w swoim języku wyrażenie, które uznaje za najwłaściwszy korelat tłumaczonego zdania (w tym przypadku jednowyrazowego zdania „Królik!”). Rozmaite inne kandydatury, takie jak „Zbiór faz czasowych królika!” albo „Fragment czasoprzestrzeni wypełniony królikiem!” odrzuca, jako zbyt dziwaczne. Okazuje się jednak, że tym niewinnym posunięciem determinuje wiele swych przyszłych przekładów. Korelując zdania, koreluje też ich składniki. Korelując składniki zdań, snuje hipotezy na temat reguł, które pozwalają na konstruowanie nowych zdań języka tubylców. Korelacja składników wraz z postulowanymi regułami determinuje ontologię, którą przypisuje tubylcom, czyli odniesienie składników zdań, których używają, dokładnie tak, jak to się stało w przykładzie z językiem japońskim. Rozjaśnia to, jak sądzę, uwagę Quine’a, iż niezeterminowanie przekładu dotyczy zdań, a względność odniesienia terminów⁷. Istotą niezeterminowania przekładu jest to, że nawet gdyby uczulić naszego lingwistę na arbitralność jego pierwszych translatorskich kroków, to nie byłby on skłonny skorygować swych wyborów. Korekta byłaby niezwykle kosztowna – musiałby w zasadzie przepisać cały swój podręcznik, a zysk z tego posunięcia byłby niezwykle mizerny. Nowy podręcznik sprawdzałby się w praktyce dokładnie tak samo, a wprowadzona korekta nie byłaby ani o jotę mniej arbitralna niż wybór pierwotny. Najprzyjaźniej nawet nastawieni tubylcy nie mogliby pomóc mu w wyborze, ponieważ nie mogliby w nim wywołać pobudzeń, które nakierowałyby go na właściwą ontologię. Lingwista nie może też poprosić tubylców o przejrzenie jego notatek albo zweryfikowanie niektórych ryzykownych hipotez.

⁶ Ponieważ jest to przykład doskonale wszystkim znany, nie będę go w tym miejscu szczegółowo opisywał.

⁷ W.O. Quine, *Na tropach prawdy*, tłum. B. Stanosz, Warszawa 1997, s. 84

Wróćmy więc na moment do przykładu różnicy czysto werbalnej, od której rozpocząłem. Czym podane przez Quine'a przykłady różnią się od tego nieinteresującego filozoficznie przypadku? Sądzę, że zachodzą tu dwie istotne różnice – po pierwsze, niezdeterminowanie przekładu dotyczy wyrażen, których z jakichś powodów nie jesteśmy w stanie powiązać z żadnym bodźcem mogącym naprowadzić nas na poprawny przekład. Po drugie, przykład różnicy czysto werbalnej był przekładem słowo w słowo. Przykłady Quine'a bazują zaś na tym, że konkurencyjne przekłady różnią się syntaktycznie – pewne wyrażenia języka tubylców (gavagai) zostały przełożone na wyrażenia proste naszego języka (królik), choć mogłyby zostać przełożone na wyrażenia złożone o tym samym znaczeniu bodźcowym (fragment czasoprzestrzeni wypełniony królikami).

Aby uniknąć nieporozumień, wspomnę tylko, że Quine mówi o znaczeniu bodźcowym zdań, a nie znaczeniu bodźcowym dowolnych wyrażen. Nie widzę jednak powodu, dla którego nie moglibyśmy mówić o znaczeniu bodźcowym części zdań. Wystarczy, że język, na który przekładamy, zawiera wyrażenie, które dołączone do dowolnego wyrażenia tworzy zdanie – nazwijmy je wyrażeniem zdaniotwórczym. W języku polskim czymś takim będzie choćby wyrażenie „oto...” albo „to jest...” (w wielu przypadkach funkcjonuje tak również wykrzyknik). Wystarczy wtedy przyjąć następujące zasady:

1. Dwa wyrażenia mają to samo znaczenie bodźcowe, gdy zdania utworzone z tych wyrażen (i z żadnych innych, poza wyrażeniem zdaniotwórczym) mają to samo znaczenie bodźcowe.
2. Dane wyrażenie posiada znaczenie bodźcowe, gdy zdanie utworzone z tego wyrażenia (i z żadnych innych poza wyrażeniem zdaniotwórczym) posiada znaczenie bodźcowe.

Dzięki tym modyfikacjom teorii Quine'a możemy powiedzieć, że wyrażenie *królik* ma znaczenie bodźcowe, ponieważ zdanie „*Oto królik*” je ma, oraz że wyrażenie „fragment czasoprzestrzeni wypełniony królikami” ma to samo znaczenie bodźcowe, co wyrażenie „królik”, ponieważ zdanie „*Oto królik*” ma to samo znaczenie bodźcowe, co zdanie „*Oto fragment czasoprzestrzeni wypełniony królikami*”.

Fakt, iż pewne wyrażenia proste mogą być bodźcowo synonimiczne z pewnymi wyrażeniami złożonymi, nasuwa podejrzenie, że części tych wyrażen złożonych wzięte w izolacji nie mają żadnego znaczenia bodźcowego. Sądzę, że tak właśnie jest w rozpatrywanych przypadkach – wyrażeniami pozbawionymi znaczenia bodźcowego są słowa takie jak „zbiór”, „fragment” czy japońskie klasyfikatory. Znaczenie tych wyrażen zależy od powiązań, jakie mają z innymi wyrażeniami języka, a nie z bodźcami. Wyobraźmy sobie następujący przykład – założmy, że prezentujemy komuś kawałek zbitej filiżanki, a następnie całą filiżankę. Założmy, że w pierwszym przypadku potwierdza on jednowyrazowe zdanie „*Oto fragment*”, a w drugim je odrzuca. Założmy teraz, że drugi z testowanych przez nas ludzi

robi odwrotnie. Czy uznamy, że pierwszy z nich posługuje się słowem „fragment” poprawnie, a drugi nie? Gdybyśmy to zrobili, byłby to wniosek przedwczesny. Pierwszy mógł przeciwstawić fragment filizanki filizance. Drugi mógł przeciwstawić filizankę kompletowi filizanek. Jeszcze inny mógłby przeciwstawić komplet filizanek asortymentowi jakiegoś magazynu. Aby ustalić, czy ludzie ci posługują się terminami takimi jak „fragment”, „komplet”, „asortyment” w sposób taki jak my, musimy z nimi wejść w dialog, którego nasz lingwista z dżungli nie może nawiązać z tubylcami. Gdy stykamy się z obcym językiem, rozpoczynamy przekład od tych zdań, które łatwo nam zweryfikować. Ale nasz język zawsze oferuje nam kilka alternatyw, które są bodźcowo synonimiczne, ale syntaktycznie różne. Wyboru dokonujemy na podstawie zdroworozsądkowych spekulacji na temat potrzeb tubylców i jednej ekstrawaganckiej hipotezy – że ontologia naszego języka jest czymś naturalnym. W pewnym momencie możemy jednak niepostrzeżenie skorelować wyrażenia proste ze złożonymi i nie będzie nikogo, kto mógłby nas skorygować.

Nie odpowiedzieliśmy sobie jednak na pytanie – po co nam wyrażenia, które nie mają żadnego znaczenia bodźcowego? Dlaczego rozmaity ich przekład miałby owocować różnicami w przypisywanym odniesieniu przedmiotowym?

Aby to zobaczyć, należy przyrzeć się behawioryzmowi Quine’a. Przyda się to nam w dwójnasób. Sądzę bowiem, że większość nieporozumień związanych z niezdeteminowaniem przekładu bierze się z trudności, jakie czytelnicy Quine’a mają z zaakceptowaniem jego behawioryzmu. Warto więc uświadomić sobie powód, dla którego Quine obstawał przy budowaniu teorii behawioralnej. Zupełnie naturalne wydają się przecież co najmniej dwie alternatywne opcje – teorię języka można zacząć budować, wychodząc od stanów mózgów użytkowników języka albo od przedmiotów, z którymi wchodzi w interakcje.

Szczególnie druga z możliwości wydaje się niezwykle kusząca – jaki właściwie zysk mamy z mówienia o bodźcach? Czy chodzi tu o opory przed wypowiadaniem się na temat tajemniczych „rzeczy samych w sobie”, o których wiemy tylko tyle, że drażnią naszą powierzchnię i niektórych naszych filozofów? Wydaje się, że nie to jest powodem – mówienie o bodźcach z góry przesądza istnienie przedmiotów, przecież jesteśmy nimi my i nasze receptory.

Sądzę, że bodźce są jednak w pewnym sensie „ontologicznie niewinne” i to usprawiedliwia obstawanie przy nich. Pamiętajmy, że rozpoczynamy od analizy bodźców całościowych – całkowitych pobudzeń w danym momencie. Odpowiadałyby im zdania złożone z predykatów zeroargumentowych. Trudno byłoby podać przykład takiego predykatu wyrażającego pobudzenie całościowe (te predykaty zeroargumentowe, które posiadamy, takie choćby jak „grzmi”, nie są całościowe).

Język istoty, która zatrzymała się na tym etapie, byłby nieodróżnialny od zwykłego rejestrowania zachodzących w otoczeniu zmian. Można to sobie wyobrazić mniej więcej tak – wyobraźmy sobie fikcyjny gatunek zwierząt – *kamereona*.

Przyjmijmy, że zwierzęta te zmieniają wygląd swej skóry w taki sposób, że co sekundę fotografują otoczenie i przyjmują jego wygląd. Czy bylibyśmy skłonni powiedzieć, że opisują one rzeczywistość, czy też jedynie na nią reagują (rejestrując ją)? Poszczególne obrazki na skórze naszych kamereonów będą dość podobne do siebie, ale prawie nigdy takie same (wszystko zależy od tego, jak duży stopień szczegółowości, rozumiany jako rozdzielczość ich receptorów, przyjmujemy). Ale zachodzenie podobieństwa stwierdzić może jedynie ten obserwator, który dysponuje kategoriami części i całości. Kamereony ich nie mają, ponieważ nie posiadają żadnych narzędzi, za pomocą których mogłyby dokonywać zróżnicowań w ramach pojedynczego bodźca całościowego (czyli jednej fotografii rzeczywistości). Wypowiadanie zdań obserwacyjnych skorelowanych z nowymi pobudzeniami całościowymi nie różni się od tego przykładu w żaden istotny sposób. Jest to jedynie lingwistyczny półprodukt, język w fazie powstawania. Aby mówić o gotowym języku, musimy mieć coś więcej niż tylko takie zdania – nazywanie ich zdaniami jest z naszej strony hojnym komplementem. Aby stały się językiem, musimy w nich wyróżnić części, które służyć będą do budowy nowych zdań. Język, w którym dowolny ciąg dźwięków byłby dopuszczalny, jest niewyuczalny. Musimy więc objąć pewne konstrukcje zakazem. Twierdzimy na przykład, że pewne wyrażenia nie mogą występować w zdaniach w roli orzecznika (co w zależności od konkretnego języka sprowadza się do rozmaitych konkretnych zakazów – mogą one dotyczyć kolejności słów w zdaniu, dopuszczalnych końcówek albo czegoś jeszcze innego). Zakazy te dzielą nasz słownik na kategorie wyrażen, a my, hipostazując te kategorie, wyróżniamy w naszych bodźcach całościowych takie aspekty, które uważamy za zasadniczo różne z punktu widzenia ontologii. W tym sensie bodźce poprzedzają konkretne rozstrzygnięcia ontologiczne i jest sens zaczynać analizę właśnie od nich.

Przewaga bodźców nad stanami mózgu miałyby zaś polegać na tym, że łatwiej o zbieżności pomiędzy pobudzeniami moimi i moich rozmówców niż o zbieżności pomiędzy stanami naszych neuronów. Jak pisze Quine – chciał trzymać się z daleka od idiosynkratycznych ścieżek połączeń nerwowych⁸. Jest to jednak wyjaśnienie mylące. Nasze pobudzenia mogą być równie idiosynkratyczne jak stany naszych mózgów. Warunków tożsamości bodźców nie możemy przecież sprowadzić do samych tylko zbiorów receptorów. Każdy użytkownik języka ma swoje własne receptory i jedyne, o czym możemy mówić, to korelacje pomiędzy pobudzeniami moimi i moich rozmówców. Zgodność ta polega na tym, że zawsze, gdy ja potwierdzam zdanie *p*, towarzyszy mi pewien bodziec *x*, a mojemu towarzyszowi bodziec *y*. Bodźce te mogą być do siebie zupełnie niepodobne. Ważne jest tylko to, że zawsze, gdy mnie przydarza się *x*, mojemu rozmówcy przydarza

⁸ W.O. Quine, *Słowo i przedmiot*, tłum. C. Cieśliński, Warszawa 1999, s. 44.

się y. Zgodność tę Quine nazywa przedustawną harmonią podobieństw⁹. Pozornie może wydawać się to niewystarczającą podstawą do porozumienia. Czy język nie okazuje się przez to być po prostu konglomeratem idiolektów powiązanych w zupełnie tajemniczy sposób, nazwany tu „harmonią przedustawną”? Sądzę, że obiekcja taka powstaje, ponieważ żyjemy wciąż „mitem galerii”, który Quine krytykował¹⁰. Zamiast wspólnych wszystkim rozmówcom idei w umysłach oczekujemy wspólnych wszystkim rozmówcom pobudzeń i gdy ich nie znajdujemy, z przerażeniem konstatujemy, że komunikacja jest Wielką Tajemnicą. Wydaje się, że nawet Quine ulegał czasem temu niepokojowi.

Wydaje mi się, że jeśli odrzucimy ten mit galerii raz na zawsze, to odkryjemy, że komunikację da się jednak w ramach teorii behawioralnej objaśnić. Aby wprowadzić wspólne odniesienia, wystarczy dołączyć do języka takie wyrażenia, które są niezależne od jakichkolwiek konkretnych pobudzeń poszczególnych członków społeczności. Jest tak, ponieważ chcemy zsynchronizować pobudzenia nasze i naszych współplemieńców, niezależnie od tego, jakie one w danym momencie są. Na tym właśnie polega publiczność języka. Potrzebujemy więc terminów, które nie będą miały żadnego określonego znaczenia bodźcowego. Są to wyrażenia takie jak „to coś” albo „to samo, co wczoraj”. W jaki sposób możemy je do języka wprowadzić? Nie jesteśmy przecież w stanie dostrzec, jak w danym momencie pobudzone są receptory naszych rozmówców (ludzie nie są kamereonami) i czy nie powiążą oni tych wyrażen z jakimś przypadkowym bodźcem. Możemy to zrobić, wprowadzając te wyrażenia do języka za pomocą takich zdań, które należy zawsze potwierdzać, i takich, które należy niezależnie od okoliczności odrzucać – nie ryzykujemy w ten sposób wyróżnienia jakiegoś partykularnego pobudzenia, które któryś z członków społeczności mógłby w danym momencie mieć. Oto rola, jaką odgrywają wyrażenia pozbawione znaczenia bodźcowego – pozwalają na synchronizację naszych idiolektów. Ich obecność w języku jest przyczyną niezdeteminowania przekładu – mamy bowiem sporą dowolność w korelowaniu naszych wyrażen pozbawionych znaczenia bodźcowego z tym, co wygląda nam na ich odpowiedniki w językach innych społeczności. Zła wiadomość jest więc taka, że nigdy nie dowiemy się, czy Japończycy liczą bydło, czy bydłeta.

Jest jednak jeszcze dobra wiadomość – gdy nie pojawia się problem przekładu, nie pojawiają się wątpliwości co do odniesienia przedmiotowego. Wydaje się więc, że w pewnej szczególnej sytuacji mogę być pewny co do tego, do czego odnosi się mój rozmówca – jest tak, gdy zarówno ja, jak i mój rozmówca nauczyliśmy się języka polskiego jako swego pierwszego języka. Uczące się dziecko zastaje kompletny system, który musi zinternalizować. Ucząc się języka, przejmuje pewien

⁹ W.O. Quine, „Progress on Two Fronts”, w: „The Journal of Philosophy”, vol. 93, nr 4, s. 161.

¹⁰ W.O. Quine, *Filozofia logiki*, tłum. B. Stanosz, Warszawa 2002, s. 21.

aparatu służący do dzielenia bodźców na części. Nie konstruuje hipotetycznych korelacji, ponieważ nie ma jeszcze z czym korelować. Gdy w dalszych etapach zaczyna konstruować pewne hipotezy, to, w odróżnieniu od naszego badacza z dżungli, stałej kontroli nie podlegają tylko same ich efekty w postaci testowanych zdań, ale i same te hipotezy. Są one bowiem formułowane w języku, który jest zrozumiały dla wszystkich, a nie jedynie dla ich twórcy. W języku, który zinternalizowałem, nie ma miejsca na względność ontologiczną. Nie mam wątpliwości co do tego, że słowo „królik” odnosi się do królików. Gdyby ktoś powiedział mi, że „królik” odnosi się do królika albo, że „królik” odnosi się do zbioru faz czasowych królika albo jakiegoś inaczej skonstruowanego zbioru, to powiedziałbym po prostu, że osoba ta nie posługuje się językiem polskim. W języku polskim, do zbioru faz czasowych królika odnosi się bowiem wyrażenie „zbiór faz czasowych królika”.

Jeżeli wydaje się nam, że człowiek ten może używać terminu „królik”, a mimo to nie odnosić się do królików, to tylko dlatego, że po cichu wierzymy, że oprócz dyspozycji do zachowań ma on jeszcze „coś na myśli”. Jesteśmy po prostu zwolennikami języka myśli, którego odniesienie może różnić się od języka publicznego. Jeżeli zaś uważamy, że wolno nam twierdzić, że „królik” to na przykład nazwa szczególnego zbioru faz czasowych – i twierząc tak, nie zmieniamy automatycznie języka na inny niż polski, to znaczy, że jesteśmy przekonani o możliwości zmiany ontologii bez zmiany języka. Jest to jednak metafizyczna tęsknota za prawdziwym, niezależnym od języka umeblowaniem świata, nie różniącą się niczym od sporu o „rzeczywisty” charakter liczb, który Carnap zdemaskował swego czasu jako spór o kategorię syntaktyczną wyrażań liczbowych.

Podsumujmy więc – niezdeterminowanie przekładu powstaje, gdy zabieramy się do przekładu takich wyrażań, które nie posiadają żadnego znaczenia bodźcowego. Alternatywne wybory, które oferuje nam nasz język, różnią się często pod względem syntaktycznym. Nie ma żadnego sposobu na ustalenie, który z przekładów lepiej odpowiada składni języka tubylców, ponieważ składnia staje się dla nas zauważalna dopiero wtedy, gdy zrekonstruowaliśmy już jakiś jej kawałek. Sceptycy mówili o znaczeniu wyrażań, że jest to ta rzecz, którą postrzega znający język, a której nie postrzega nieznający¹¹. Ale ze składnią jest dokładnie tak samo! Gdy usłyszymy zupełnie obcą nam mowę, to nie wiemy, gdzie kończą i zaczynają się poszczególne słowa i jak zbudowane są zdania. Drobne różnice składniowe pomiędzy naszymi synonimicznymi bodźcowo wyrażeniami, takie choćby jak opozycja wyrażenia prostego i złożonego albo zasada powiązania z tym czy innym wyrażeniem w zdaniu (klasyfikatory), wydawać się mogą niewinne, ale to one są podstawą naszych hipostaz. Ścisły związek względ-

¹¹ Uwagę tę przytacza Barbara Stanosz w: B. Stanosz, *Logika języka naturalnego*, Warszawa 1999, s. 11.

ności ontologicznej z niezdeteminowaniem przekładu pozwala nam jednak spać spokojnie. Gdy ludzie, którzy zinternalizowali tę samą sieć wyrażen ustalających odniesienie, używają tych samych słów co ja, w tych samych sytuacjach co ja, to odnoszą się do tych samych przedmiotów. Mamy to zagwarantowane, ponieważ jest to jedyny sens wyrażenia „odnosić się do tych samych przedmiotów”.

From Impression to Expression, or Indeterminacy of Translation and Reference

Key Words: Quine, indeterminacy of translation, indeterminacy of reference, behavioral theory of meaning.

The paper analyzes notions of indeterminacy of translation and indeterminacy of reference. The author discusses the differences between the two and contrasts them both with the cases of merely verbal differences that can arise when several translators attempt to translate one text. He then sets apart the different mechanisms that generate the two kinds of incongruence. He takes indeterminacy of translation to be the effect of the existence of two concurrent translations of a given expression that are stimulus-synonymous but syntactically different. Because of the close connection between the indeterminacy of translation and the indeterminacy of reference the latter does not occur between people speaking their native language. It is so, because they have internalized their grammatical apparatus without possessing a prior grammatical apparatus to correlate it with.