

Przekłady

Poniżej drukujemy fragment książki Willarda Van Ormana Quine'a *Różności*, która w tym roku ukaze się nakładem wydawnictwa *Aletheia*.

Różności to rodzaj autorskiego słownika, zainspirowanego przez Wolterowski *Słownik Filozoficzny*. Niniejsza praca – pisze we wstępie autor – *jest po części filozoficzna, ale znaczna część książki poświęcona jest nie tak wzniosłym tematom, co sprawiło mi więcej niż połowę przyjemności, gdyż filozofia nie jest na ogół dyscypliną do śmiechu*.

Willard Van Orman Quine

Różności (fragment)

I

Idee

Przypuszcza się, że idee są czymś w naszym umyśle. Kiedy staramy się dokładniej określić ich naturę, na początku przychodzą nam do głowy wyobrażenia mentalne. Jak na pojęcia mentalne są one względnie dobrze scharakteryzowane. Kiedy mamy wyobrażenie mentalne czerwonego jabłka, w naszym mózgu dokonuje się prawdopodobnie proces, który przypomina trochę to, co dzieje się, kiedy w rzeczywistości widzimy czerwone jabłko.

Jednakże słowo „idea” nie odnosi się do wyobrażeń mentalnych. Na ogół jego referencja (jeśli w ogóle jakąś ma) jest bardziej ulotna. Kiedy Biały Skoczek myślał

... o sposobie,
Jak wąsy przefarbować
Na piękną, jasną zieleń sobie
I pod wachlarzem schować

to istotnie miał pewne wyobrażenia mentalne, ale nie były one jego ideą; ideę stanowił raczej cały jego plan, czymkolwiek by on był. Kiedy ktoś mówi „Pomyślałem, że mógłbym wykupić udziały w Tronics”, a my odpowiadamy „Ten sam pomysł przyszedł mi do głowy”, to nasze wyobrażenia nie muszą przypominać wyobrażeń naszego rozmówcy; mówiąc „idea” znów mamy na

myśli plan. Przypuszczam, że wszystkie plany, poza planami na papierze, są ideami, chociaż nie na odwrót. Jednakże wyjaśnienie jednych za pomocą drugich byłoby wyjaśnianiem niejasnego i ulotnego pojęcia przez pojęcie, które jest przynajmniej równie niejasne i ulotne.

Jakiś pajac mówi „Sądzę, że alergię mają charakter psychosomatyczny”, a jego równie niezorientowany rozmówca odpowiada „Mnie to samo przyszło do głowy”. Tutaj wspólną ideą nie jest plan, lecz PRZEKONANIE, które, jak zobaczymy, samo jest dosyć mglistą kategorią.

Popularność zwrotu „ta sama idea” wydaje się sugerować, że idee są czymś znanym – i to do tego stopnia, że potrafimy rozpoznać ich tożsamość i różnicę. W rzeczywistości jednak nasza gotowość do mówienia o identyczności idei nie opiera się na żadnym jasnym pojęciu idei. W wielu przypadkach tego rodzaju zwroty sprowadzają się jedynie do powtórzenia wypowiedzianego przez kogoś zdania. Kiedy człowiek w poprzednim przykładzie powiedział „Mnie też to przyszło do głowy”, to miał po prostu na myśli „Ja również myślę, że alergię mają charakter psychosomatyczny”. Kiedy we wcześniejszym przykładzie ktoś inny powiedział „Ten sam pomysł przyszedł mi do głowy”, to miał na myśli „Ja również pomyślałem, że mógłbym wykupić udziały w Tronics”.

Jeżeli urzędnik biura patentowego oświadcza, że dwa wynalazki opierają się na tej samej idei, to mówi, że jeden i ten sam opis, pomijający nieistotne szczegóły, byłby ich równie trafną charakterystyką.

Widzimy dzięki tym przykładom, że jeśli chcemy posługiwać się bardziej precyzyjnym językiem i mówić w nim o ideach, to powinniśmy nie tyle wyjaśnić, czym są idee, co pokazać, w jaki sposób parafrazować zdania o ideach na zdania o języku. Na ironię zakrawa fakt, że pojęcie idei ma według niektórych ludzi wyjaśniać kwestie językowe, podczas gdy kierunek wyjaśniania jest akurat odwrotny.

Słyszymy czasem, że zdania są równoznaczne albo synonimiczne gdy wyrażają tę samą ideę. Trudno to kwestionować, ale byłoby lepiej, gdybyśmy poszli w przeciwnym kierunku i wyjaśnili w ten sposób identyczność idei. Jak możemy więc wyjaśnić równoznaczność zdań? Wskazując, że wypowiedzenie jednego zdania służyłoby tym samym celom, co wypowiedzenie drugiego, jeśli naszymi słuchaczami są użytkownicy danego języka czy języków. To wciąż nie jest zbyt precyzyjne, ale przynajmniej ma mocne oparcie. Zob. ZNACZENIE.

Słyszymy, że główną funkcją języka jest KOMUNIKOWANIE idei. Twierdzi się, że jeśli chcemy przekazać komuś ideę, to musimy wytworzyć w umyśle odbiorcy tę właśnie ideę. Skąd wiemy, że są to te same idee? Możemy jedynie stwierdzić, że nasza informacja spowodowała właściwą reakcję lub odpowiedź. Poważna analiza języka i komunikacji skoncentruje się na percepcji, mowie i działaniu i nie będzie korzystać z pojęcia idei, gdyż jest to jedynie mglisty opar, który prowadzi tylko do złudnych wyjaśnień.

Mówienie o ideach utrwaliło się w języku potocznym i jest całkiem wygodne.

Trudno nam uniknąć go na co dzień i nie warto próbować. Jest to jednak pułapka dla filozofa i naukowca, którzy pozwalają sobie na używanie takiego języka w swoich teoriach. Proszę zauważyć, że słabym punktem idei nie jest ich abstrakcyjność. Niektóre przedmioty abstrakcyjne albo UNIWERSALIA zajmują niezaprzeczalne miejsce w naukach przyrodniczych. Jednakże nie ma w nauce miejsca dla idei¹.

Identyczność

Terminu tego używamy czasem w luźny sposób. Mówimy o identycznych bliźniakach. Mówimy, że ty i ja mamy identyczne samochody. Pomimo tej potocznej swobody termin ten w swym właściwym sensie jest tak ścisły, jak to tylko możliwe. Przedmiot jest identyczny z samym sobą i z niczym innym, nawet ze swoim identycznym bliźniakiem.

David Hume wyczuwał w tym zagadkę. Identyczność wygląda na relację, nie łączy jednak ze sobą przedmiotów w pary, tak jak to robią inne relacje: przedmioty są identyczne tylko z nimi samymi. Skoro tak, to jaka jest różnica między identycznością a zwykłą własnością? Co więcej, identyczność dotyczy wszystkiego. Jeśli tak, to czym różni się ona od zwykłej własności istnienia, którą ma każdy przedmiot?

Trudno jest zrozumieć niepokoje nawet tak utalentowanych myślicieli jak Hume, z chwilą gdy zostały one wyeliminowane przez postęp nauki. Uznaje się obecnie, że relacja składa się z par przedmiotów; relacja bycia wujem obejmuje wszystkie pary wuj–siostrzeniec i wuj–siostrzenica. Do relacji identyczności należą wszystkie i tylko pary typu $\langle x, x \rangle$ czego nie należy mylić z samym tylko x -em.

W sprawie kłopotów z identycznością, zob. także UŻYCIE A WYMIENIANIE. Załączki kolejnej trudności kryją się w następnej refleksji: powiedzieć, że coś jest identyczne z samym sobą, to wygłosić banał, a powiedzieć, że coś jest identyczne z czymś innym – to wygłosić absurd. Jeśli tak, to po co nam identyczność? Wittgenstein zadał to pytanie.

Poważne kwestie dotyczące identyczności powstają, gdyż możemy odnosić się do czegoś na dwa sposoby i zastanawiać się, czy odnosimy się do tego samego przedmiotu. Wspominam na przykład o Simonie, ktoś wspomina o Peterze i w końcu dochodzimy do wniosku, że Simon jest Peterem, są oni identyczni. Nie jest to banalna konstatacja i można w to wątpić nie wygłaszając absurdów.

Nie potrzebujemy na ogół nadawać ludziom dwóch imion i nie warto wprowadzać pojęcia identyczności tylko po to, by radzić sobie z takimi

¹ Zachowałem się powściągliwie i nie wydrukowałem tego ostatniego zdania kursywą — mam nadzieję, że zostanie to zauważone. Nie mam jednak wątpliwości, że niektórzy ludzie urządzią sobie *sparring* cytując moją wypowiedź poza kontekstem. No cóż, albo *eh bien, vive le sport!*

przypadkami. Ważniejsza jest sytuacja, w której odnosimy się do czegoś nie za pomocą dwóch nazw, lecz dwóch deskrypcji albo nazwy i deskrypcji. Musimy umieć utożsamić Ralpha z człowiekiem, który strzyże trawnik, a jego dom z budynkiem stojącym najbliżej stacji. Takich identyczności pełno jest w naszych codziennych rozmowach.

W starożytności wymyślono filozoficzną zagadkę, związaną z pojęciem identyczności. Statek Tezeusza został całkowicie przebudowany od deski do deski. Czy teraz mamy do czynienia z tym samym statkiem co przed przebudową? Heraklit sformułował podobny problem: twierdził, że nie można wstąpić dwa razy do tej samej rzeki, gdyż jej zawartość wciąż się zmienia. Skoro już o tym mowa, to czy Ralph jest tym samym człowiekiem, który strzygł trawnik osiem lat temu? Mówi się przecież, że ciało człowieka ulega całkowitej wymianie w ciągu siedmiu lat. Czy naprawdę wciąż jesteśmy sobą po tym czasie?

Uznaje się, że te trzy zagadki – a tak naprawdę jedna – godzą w pojęcie identyczności. Jest to błąd – nie chodzi w nich o naturę identyczności, lecz o to, co będziemy uznawać za statek, rzekę, osobę. Słowa są narzędziami i toleruje się ich nieostrość, jeśli nie naraża to na szwank ich przydatności.

Orzekamy identyczność jednostki w czasie nie dlatego, że zachowała się jej substancja cielesna, lecz ze względu na ciągłość zmiany kształtu, masy i przyzwyczajęń. Oczekujemy również ciągłości pamięci, ale od czasu do czasu przechodzimy do porządku dziennego nad przypadkami nieciągłości. Trudno powiedzieć, w którym dokładnie momencie zaczyna się osoba – przy porodzie, przy poczęciu czy gdzieś pośrodku – gdyż do tej pory nic istotnego nie zależało od tej kwestii.

To dziwne, ale nie wszyscy zdają sobie sprawę z faktu, że jest to po prostu pytanie o rzeczywiste lub proponowane użycie słowa „osoba”. Nie chodzi tu o wyróżnienie niezauważonego przedtem znaczenia tego wyrażenia. W słowach, jak zauważył Humpty Dumpty, nie kryje się nic więcej poza tym, co z nimi robimy.

Zastanawiałem się tu nad osobami, ale to samo dotyczy rzeki Heraklita i statku Tezeusza. Prawdziwość sądu stwierdzającego identyczność zależy od tego, jakie terminy ogólne są w nim użyte lub zasugerowane – „osoba”, „rzeka”, „statek”. Ralph jest teraz tą samą osobą co osiem lat temu, tylko czasowe stadia są różne. Z drugiej strony, kiedy ornitolog mówi „To jest tym samym, co to” wskazując w dwóch różnych kierunkach, to nie ma sensu oskarżać go o to, że ma na myśli dokładnie to, co głosi wypowiedziane przez niego zdanie. Chodzi mu raczej o to, że gatunek, do którego należy ten ptak, jest identyczny z gatunkiem tamtego ptaka.

Kiedy używamy takich wyrażęń, jak „tylko”, „nikt inny”, „nic oprócz tego”, to w ukryty sposób posługujemy się identycznością. Kiedy mówię, że tylko Ralph i nikt inny zna tę kryjówkę, to chodzi mi o dwie rzeczy: że Ralph zna kryjówkę i że każda osoba znająca kryjówkę jest identyczna z Ralphem. Jeśli

mówimy, że nie ma Boga prócz Allacha, to stwierdzamy, że ktokolwiek jest Bogiem, jest identyczny z Allachem.

Idiotyzmy

Nasze słowo *idiom* odnosi się do szczególnych wyrażen językowych, których użycie nie daje się odczytać z szerszych regularności występujących w danym języku, ani ze znaczenia składowych słów w innych kontekstach. Używa się również tego słowa, podobnie jak słów *idiome* i *idioma* w językach romańskich, w odniesieniu do całego języka albo dialektu. W języku francuskim słowu *idiom* w pierwszym sensie odpowiada termin *idiotisme*. Jeśli ktoś zaczyna się uczyć francuskiego i boryka się z takimi idiomami jak *j'ai beau faire* albo dowiaduje się, że *le bel âge* oznacza młodość, a *un bel âge* starość, to z kwaśnym rozbawieniem dojdzie do wniosku, że słowo *idiotisme* jest bardzo trafne.

Przyznajemy, że skojarzenie z idiotyzmem jest tu nie na miejscu; idiosynkrazja jest lepszym modelem. W grece samo słowo *idiota* nie miało żadnej przykłej konotacji; *idiotēs* to po prostu zwykły człowiek z ulicy, członek populacji, którego wyróżniającą cechą był brak wyróżniających cech. Słowo to stawało się coraz ostrzejsze w łacinie i późniejszych językach dzięki smutnemu procesowi eufemizacji (zob. EUFEMIZM).

Idiomy w pierwszym sensie wykraczają poza poważną definicję, którą zaproponowałem na początku tego rozdziału. Czasami nazwiemy dany zwrot idiomem, gdy odzwierciedla on punkt widzenia różny od naszego. W operze słyszymy „*fin che non giunga il Re*” i zgadujemy na chybcika, że może to znaczyć „żeby król nie przyszedł”. Jesteśmy zdumieni, kiedy słyszymy prosty przekład: „aż przyjdzie król”. Po co tu słowo *nie*? Wszystko zaczyna do siebie pasować, kiedy zdajemy sobie sprawę z faktu, że czekanie, aż przyjdzie król, dokonuje się dokładnie w okresie, w którym król nie przychodzi. Użytkownicy języka angielskiego myślą w terminach pozytywnych o tym, co się dzieje do chwili przybycia króla („*Until the King comes*”). Donizetti i jego rodacy myślą o tym w terminach negatywnych. Wychodzi na jedno.

Inny podobny przykład również pochodzi z piosenki, tym razem meksykańskiej: „*ya no puede caminar*”. Dosłownie: „już nie może wędrować”. Dla użytkownika języka hiszpańskiego on już jest w stanie, w którym nie może wędrować. Anglik powiedziałby raczej *he can't get around any more* – on już nie jest w stanie, w którym mógłby wędrować. Oba zwroty są równie proste i znaczą to samo. Jeżeli jednak ktoś jest użytkownikiem jednego z tych języków, to mógłby dojść do wniosku, że ci obcokrajowcy mówią w niezwykle idiotyczny sposób.

Informacja

Ile informacji, prawdziwej lub fałszywej, interesującej lub banalnej, uzyskaliśmy od momentu spożycia śniadania, albo od momentu otwarcia tej książki? To

pytanie jest pozbawione sensu – nie scharakteryzowaliśmy żadnej miary informacji. Claud Shannon i Warren Weaver zdefiniowali jednak miarę czegoś, co nazwali informacją i co można zastosować w pewnych dokładnie opisanych sytuacjach – jest to właśnie ich teoria informacji. Otrzymała ona nagrodę Bell Labs za użyteczność w projektach komunikacyjnych.

Przypuśćmy, że mamy daną dziedzinę ATOMÓW. Może być ich skończenie lub nieskończenie wiele, muszą się jednak dzielić na skończoną ilość rodzajów. Rozważmy prosty przykład półtonów, o których już wspominałem w rozdziale ATOMY. Atomami są tam miejsca na siatce półtonowej, w których może, lecz nie musi występować czarna kropka. Mamy dwa rodzaje atomów: czarne miejsca i niezapełnione białe miejsca. Kiedy określamy kolor danego punktu – biały lub czarny – to przekazujemy jeden bit informacji, dotyczącej obrazu na siatce półtonowej. Każdy kolejny punkt niesie ze sobą następny bit. Ilość bitów to ilość przekazanej informacji. Każdy bit informacji wiąże się więc z pewnym binarnym wyborem. W mniej banalnych systemach występują więcej niż dwa rodzaje atomów, a sposób pomiaru informacji odpowiednio się zmienia. Istnieje ogólna metoda radzenia sobie z dowolną skończoną liczbą rodzajów atomów. Przypuśćmy na przykład, że mamy osiem rodzajów i każdemu z nich nadajemy nazwę w stylu alfabetu Morse’a, składającą się z kropek i kresek. Pierwszy rodzaj nazywa się „...”, drugi „-.”, trzeci „.-.” i tak dalej, aż do „- - -”. Jeśli mówimy o jednym z ośmiu rodzajów, to przekazana przez nas informacja sprowadza się do takich oto trzech binarnych możliwości: nazwa może się zaczynać od kropki albo kreski, na jej drugim miejscu może stać kropka albo kreska i może się ona kończyć kropką albo kreską. Skoro wybór jednej z dwóch możliwości jest jednym bitem informacji, przekazana przez nas informacja mieści w sobie trzy bity.

Pewna nieskomplikowana formuła obejmuje wszystkie przypadki, w których informacja jest względna wobec atomizmu z określoną i skończoną liczbą rodzajów. Czytelnicy przyzwyczajeni do logarytmów domyślili się pewnie z ostatniego przykładu, że dla n rodzajów liczba bitów informacji, którą przekazujemy kwalifikując dany przedmiot do jednego z tych rodzajów, wynosi $\log_2 n$.

Gdyby można było zbudować atomizm percepcyjny, o którym marzyliśmy w rozdziale ATOMY, to byłibyśmy w stanie mówić obiektywnie o ilości empirycznej informacji odbieranej przez dowolną osobę przy danej okazji lub w danym okresie. Obecnie jednak teoria informacji przydaje się przede wszystkim w komunikacji i komputerach.

J

Jednostki

Możemy sobie wyobrazić jard kwadratowy i milę kwadratową, ale co z kwadratową godziną? No cóż, wiem, co możemy z tym zrobić. Pom-

nóży ją przez przyspieszenie, a otrzymamy długość, jak za chwilę zobaczymy.

Zacznijmy od prędkości. Sześćdziesiąt mil na godzinę. Możemy to zwięźle zapisać w takiej oto postaci: 60 m/h. Jest to kreska dzielenia, takiego samego jak w przypadku wyrażenia „8 procent”, albo „per centum”, to znaczy osiem podzielone przez sto albo 0,08. Tyle jeśli chodzi o lekcję nr 1.

Przyspieszenie to prędkość podzielona przez czas. Wielkość ta określa, w jakiej mierze prędkość zwiększa się co godzinę albo co sekundę – w zależności od tego, jaką jednostkę wybierzemy. Prędkość spadającego ciała mierzona w stopach na sekundę zwiększa się o 32 w każdej sekundzie; przyspieszenie wynosi zatem mniej więcej 32 st/sek². Przyspieszenie jest więc długością podzieloną przez czas do kwadratu. Pomnożmy przyspieszenie spadającego ciała przez sekundę do kwadratu, a otrzymamy długość: 32 stopy. Pomnożmy przyspieszenie spadającego ciała przez godzinę do kwadratu, a również otrzymamy długość: 414 720 000 stóp, czyli mniej więcej 78 545 mil.

To *nie* znaczy, że ciało spadnie aż tak daleko przez godzinę. Przebędzie jedynie połowę tego dystansu, jak wskazuje na to cyfra 32, która zresztą nie ma zastosowania przy takiej odległości od ziemi. Nie stosuje się ona również przy bliższych odległościach ze względu na opór powietrza. Podane przeze mnie cyfry, choć fascynujące, są cudownie bezużyteczne.

Liczba sekund do kwadratu w minucie do kwadratu to kwadrat liczby sekund w minucie, czyli 3 600. Jest więc tyle sekund kwadratowych w minucie kwadratowej, co sekund w godzinie. Z drugiej strony, w godzinie kwadratowej występuje 12 960 000 sekund kwadratowych.

Mnożenie rozmaitych jednostek to znana sprawa. Pracę mierzy się w stopach na funt. Pieniądże podzielone przez pracę – dolary na stopy-na-funty – określają stopień łatwości zarabiania pieniędzy, a odwrotność tej wielkości – stopy-na-funty na dolary – określa stopień trudności. Czysta liczba podzielona przez pieniądź – jeden na dolara – mierzy progresję podatkową: wzrost opodatkowania wraz ze wzrostem dochodów. Progresja podatkowa jest odwrotnością pieniądza.

Stopa na akr to mało znana jednostka, ale jej sens jest oczywisty: jest ona miarą objętości, podobnie jak litr. Chodzi tu o wystarczającą ilość wody, żeby pokryć obszar jednego akra na głębokość jednej stopy. Nie słyszy się zwykle o litrze na godzinę, ale jest to godna uwagi jednostka, gdyż służy do pomiaru czasoprzestrzennej objętości jednostki ludzkiej lub innej istoty w ciągu jej całego życia. Jeśli ktoś w trakcie swojego życia ma przeciętnie objętość sześćdziesięciu litrów, to przez siedemdziesiąt lat uzbiera mu się olbrzymia objętość prawie 37 milionów litrów na godzinę.

Jednostka długości podniesiona do kwadratu staje się jednostką powierzchni: kwadratowe jardy, hektary, akry. Pierwiastek kwadratowy z akra to mniej więcej siedemdziesiąt jardów. Jaki jest pierwiastek kwadratowy z jarda? Przekracza to

naszą wyobraźnię, albo lepiej: jest to jednostka urojona, tak samo jak pierwiastek kwadratowy $z - 1$ jest liczbą urojoną.

Kolejne pytanie: czym jest akr kwadratowy? Czy jest to objętość sześciangu, którego każdy bok ma powierzchnię jednego akra? Takie wyobrażenie może się nam narzucać, ale nie jest to poprawna odpowiedź. Akr kwadratowy musi mieć cztery wymiary. Stawia go to w jednym rzędzie z litrem na godzinę, tyle tylko, że ten czwarty wymiar nie musi być związany z czasem. Jeśli przyjmimy, że jest to czas, to akr kwadratowy będzie czasoprzestrzenną objętością sześciangu, którego boki mają powierzchnię jednego akra, w czasie, będącym temporalnym odpowiednikiem siedemdziesięciu jardów (zob. CZASOPRZESTRZEŃ).

Możemy odnaleźć akry i godziny kwadratowe w pojęciu energii, jeśli przyjrzymy się sloganowi teorii względności, „ $E = mc^2$ ”; c jest tu prędkością, mianowicie prędkością światła, a więc odległością podzieloną przez czas; c^2 jest zatem odległością do kwadratu podzieloną przez czas do kwadratu, mamy tu zatem akry na godziny kwadratowe. Wreszcie m jest masą, nasz slogan utożsamia więc energię z akrami i tonami na godziny kwadratowe. Pozwolę sobie dodać, że jasność koncepcji nie jest moim obecnym celem.

Skale temperatury nie pozwalają już na tak frywolne myśli. Stopnie Fahrenheita i stopnie Celsjusza są do siebie proporcjonalne, ale różne punkty pełnią tam rolę zera. W obu skalach ujemne wartości poniżej -273°C lub $-459,4^{\circ}\text{F}$ nie mają sensu, gdyż ciepło w ogóle wówczas nie występuje: nie ma ruchu molekuł. W skali Kelvina ten punkt zostaje uznany za zero, a następnie liczy się w górę w stopniach Celsjusza, dlatego 0°C 273°K .

Interesujące jest to, że możemy mimo wszystko wynaleźć skale temperatury, w których nie ma wielkości najmniejszej, a które jednak są sensowne we wszystkich swoich punktach. Załóżmy, że naszym zerem jest stare 0°C , a naszą jedynką stary 1°C . Teraz jednak przyjmimy, że nasze dodatnie stopnie stają się stopniowo coraz większe niż stopnie Celsjusza, a nasze ujemne stopnie stają się od nich coraz mniejsze. Naszym pierwszym stopniem przekraczającym zero jest 1°C . Możemy zwiększyć kolejny stopień o drobną wielkość $1/273$, jeszcze następny o dwukrotność tej wielkości i tak dalej; jednocześnie modyfikujemy odpowiednio nasze wartości ujemne. Ogólna formuła jest następująca: n stopni, dodatnich czy ujemnych, na nowej skali jest równe $274^n/273^{n-1}$ w skali Kelvina.

Najciekawsze w tym wszystkim jest to, że istnienie lub nieistnienie dolnej temperatury nie jest kwestią fizyki, lecz konwencjonalnego pomiaru, nawet jeśli obecność lub nieobecność ciepła jest faktem fizycznym. Nasza nowa skala zawiera nieskończenie wiele stopni pomiędzy zerem Celsjusza a zerem Kelvina. Jest to skala logarytmiczna. Równie dobrze można by zresztą powiedzieć, że to zwykłe skale są logarytmiczne ze względu na nią.

Moglibyśmy wykonać podobną sztuczkę w związku z pomiarem czasu, żeby uspokoić ludzi (oczywiście ani ciebie, ani mnie), którzy zastanawiają się, co mogło się dzieć przed Wielkim Wybuchem (zob. STWORZENIE). Jeśli

uciekamy się do skali logarytmicznej, możemy odsunąć Wielki Wybuch na nieskończoną odległość i utrzymywać, że świat istniał zawsze. Pierwsze trzy minuty Stevena Weinberga rozciągnęłyby się na połowę wieczności. Teorie naukowe pozostają nietknięte, zmieniają się tylko jednostki. Taki przekład wymaga jednak odpowiedniego przekształcenia jednostek przestrzennych, co ma pewne niepożądane konsekwencje: przeszłe wymiary się zwiększają, a przyszłe stają się mniejsze. Atomy z odległej przeszłości przybierają kosmiczne proporcje.

Języki sztuczne

W średniowieczu na zachodzie Europy łacina była międzynarodowym językiem dyplomatów i uczonych, jacy by oni nie byli. W trzynastym stuleciu naszej ery wieki ciemnoty miały się ku końcowi; rozwijała się myśl spekulatywna, nauka wisiała w powietrzu. Łacina odgrywała istotną rolę aż do początków ery nowożytnej, kiedy to w pełni rozwinęła się nauka.

Tymczasem francuski, angielski i inne języki wkroczyły do ambasad, a w osiemnastym wieku opanowały również akademie i przyszłe laboratoria. Tylko przy okazji rozpraw doktorskich i inauguracyjnych posługiwano się łaciną.

Na skutek tego doszło do ograniczenia ogólnokulturowego dialogu. Zauważono ten fakt i ubolewano nad nim, a w drugiej połowie dziewiętnastego wieku próbowano podjąć kroki zaradcze. Zamiast odrodzenia łaciny, wysunięto jednak bardziej skrajne i wizjonerskie projekty: proponowano ten czy inny sztuczny język ze względu na prostotę jego gramatyki i łatwość jego opanowania.

Już w siedemnastym wieku kierunek ten zwiastowały dwie próby. Jedna z nich została podjęta przez szkockiego nauczyciela, George'a Dalgarno, druga zaś przez angielskiego biskupa Johna Wilkinsa, przełożonego Wadham. Za ich propozycjami kryły się raczej utopijne wizje racjonalnej semantyki i składni, niż chęć zaradzenia kryzysowi w komunikacji międzynarodowej, gdyż wśród uczonych i w kołach dyplomatycznych wciąż można się było uciec do łaciny. Dopiero w 1880 roku niemiecki ksiądz J.M. Schleyer doszedł do wniosku, że językowa wieża Babel powinna należeć do przeszłości i stworzył swój międzynarodowy język: volapük. Słownictwo wywodziło się na ogół z języka angielskiego, chociaż było nie do rozpoznania dla Anglika. Obowiązywała w tym języku uproszczona wersja składni niemieckiej. W ciągu dziewięciu lat volapük stał się popularny w całym cywilizowanym świecie, dając początek dwustu osiemdziesięciu pięciu towarzystwom volapük. Uczyło się go podobno około miliona ludzi. Potem zaczął się rozpadać na rywalizujące ze sobą dialekty, tworzone przez zapaleńców, którzy mieli własne pomysły i potrzeby twórcze. Ludzie, którzy entuzjazzmują się reformą języka, sami mają na ogół ochotę wprowadzić jakieś innowacje. Na ironię zakrawa fakt, że międzynarodowy język zbliżył się do uniwersalności tylko po to, aby za chwilę stworzyć swoją własną wieżę Babel.

Tymczasem rosyjski* lekarz L.L. Zamenhof zajmował się konstruowaniem esperanto. Język ten powstał w 1887 roku, a moda na międzynarodowe języki okazała się dla niego bardzo sprzyjająca. W 1965 roku miał podobno osiem milionów użytkowników, chociaż również i on rozpadł się na wiele wariantów.

Najbardziej znanym z tych wariantów jest Ido, skonstruowany w 1907 roku. Było to w dużej mierze dzieło Louisa Couturata, który jest również znany jako jeden z pomniejszych prekursorów nowoczesnej logiki matematycznej. Zajmował się tą dyscypliną przed 1910 rokiem – datą ukazania się pierwszego tomu *Principia Mathematica* Russella i Whiteheada.

Ido było językiem agresywnie międzynarodowym – starano się zmaksymalizować międzynarodowy charakter każdego słowa doń należącego. Czy rdzeń *man* powinien sugerować, że mowa o rękach, jak w łacinie i językach romańskich, czy też o ludziach, jak w językach germańskich? Rozstrzygnijmy tę kwestię licząc użytkowników języków romańskich i germańskich. Czy w ten sposób udaje nam się zmaksymalizować liczbę ludzi, którzy odgadną znaczenie takiego słowa? Raczej nie; wielu z nas będzie się natomiast daremnie zastanawiać nad liczebnością tych populacji.

Sztuczne języki apelują do podobnych gustów i impulsów co logika matematyczna. Inny i bardziej znaczący prekursor logiki matematycznej, Giuseppe Peano, stworzył w 1903 roku najmniej sztuczny i moim zdaniem najatrakcyjniejszy ze wszystkich konkurencyjnych języków: *latino sine flexione*, nazwany później *Interlingua*. Darujmy sobie maksymalnie międzynarodowy charakter; każde słowo tego języka odnajdziemy w klasycznej łacinie, wyjątki stanowią jedynie nazwy leków i tym podobne przypadki. Gramatyka jest za to prostsza: nie ma deklinacji z wyjątkiem *-s* dla liczby mnogiej i nie ma koniugacji. Strukturę gramatyczną tworzą przyimki i czasy, jak w językach romańskich. Tekst napisany w tym języku jest zupełnie zrozumiały dla kogoś, kto włada jakimś współczesnym językiem romańskim lub trochę pamięta łacinę. *Interlingua* stanowi propozycję przywrócenia dawnej naukowej pozycji łacinie, która została beztrząsco odrzucona dwa wieki temu, i jest poważnym konkurentem dla swoich bardziej sztucznych, popularnych rywali. Ruch *Interlingua* przetrwał aż po dzień dzisiejszy, ale, jak to można było przewidzieć, język ten nie mógł zadowolić swoich hałaśliwych entuzjastów. Splamili oni klasyczną czystość słownika i zepsuli przejrzystość języka wprowadzając „ulepszenia”, które, choć łatwe, wymagają dodatkowej nauki.

Ta smutna opowieść nie jest wezwaniem do czynu. Angielski staje się w coraz większym stopniu językiem międzynarodowym, a dotyczy to zwłaszcza nauki. Nie ma więc potrzeby wprowadzania dodatkowego pomocniczego języka. Tego rodzaju projekty zamieniają się w nieszkodliwe hobby, jak w czasach Dalgarno i Wilkinsa, kiedy łacina wciąż cieszyła się popularnością. Ogarnia nas jednak melan-

* L. Zamenhof był lekarzem polskim (przyp. red.).

cholia: sentymentalna tęsknota za utraconą łaciną sprzed trzech wieków i przejrzystą prawie-łaciną, która w obecnym stuleciu mogłaby nam przynosić pociechę.

Sztuczna notacja (choć nie język w pełnym tego słowa znaczeniu) już od wielu wieków kwitnie w matematyce. Nie rozwinęto jej po to, żeby przekraczać granice między językami, lecz po to, żeby ułatwić myślenie o pewnych szczególnych zagadnieniach. Postępy matematyki byłyby psychologicznie niemożliwe bez obrazowej i giętkiej symboliki stworzonej specjalnie do tych celów. Podziały, które za sprawą wścibskich entuzjastów stały się zmorą takich języków jak volapuk, oszczędziły matematykę, bo to nie notacja jest oczkiem w głowie matematyków. Notacja jest jedynie środkiem i instrumentem, ale prawdziwy cel leży gdzie indziej. Kiedy pojawiają się rozmaite notacje, nikt się tym za bardzo nie przejmuje: matematyk załatwia się z nimi krótko i obserwuje, jakie informacje przekazuje się za ich pomocą.

W ostatnich latach notacja matematyczna, a zwłaszcza notacja logiki matematycznej, bardzo się rozwinęła i jest obecnie trochę bliższa duchowi starych języków sztucznych. Chodzi mi tu o Fortran, Loglan i inne języki sztuczne, które stworzono dla potrzeb programowania komputerów. Mnożą się one teraz podobnie jak dawne sztuczne języki, a sprzyja temu, tak jak kiedyś, czysta radość entuzjasty-majsterkowicza. Kryteria selekcji są jednak obecnie ostre i sztywne, gdyż przydatność tych języków może być ściśle mierzona czasem komputerowym. Jest to nowa generacja języków sztucznych, języków sztucznej inteligencji w coraz bardziej sztucznym świecie.

K

Klasy a własności

Jeżeli mówimy coś o przedmiocie, to przypisujemy mu cechę albo atrybut. Dawnymi czasy atrybut rzeczy lub gatunku nazywano cechą tylko wtedy, gdy był on charakterystyczną właściwością tej rzeczy lub gatunku. Później zapomniano o tej subtelności i zaczęto używać w ten sam sposób obu terminów. Od tej pory będę mówił tylko o cechach i zrezygnuję z terminu „atrybut”.

Jeśli mówienie o cechach ma sens, to mówienie o identyczności i różności cech również powinno mieć określony sens; niestety te ostatnie pojęcia są bezsensowne. Jeśli przedmiot posiada tę cechę, a nie tamtą, to z pewnością chodzi tu o różne cechy. Ale co będzie, jeśli wszystko, co ma jedną z tych cech, ma również i tę drugą? Czy powiemy wówczas, że są to te same cechy? Jeśli tak, to wszystko w porządku; nie ma problemu. Ludzie jednak nie udzielają takiej odpowiedzi. Słyszałem, że każda istota posiadająca serce ma również nerki i na odwrót, ale czy ktoś odważy się powiedzieć, że cecha posiadania serca jest tym samym, co cecha posiadania nerek?

Krótko mówiąc, twierdzi się, że koekstensywność cech nie jest wystarczającym warunkiem ich identyczności. W takim razie co nim jest? Często się mówi, że

cechy są identyczne nie wtedy, gdy są koekstensywne, ale wtedy, gdy są koniecznie koekstensywne. Jednakże KONIECZNOŚĆ jest zbyt mglistym pojęciem, aby miało nas to zadowolić.

W przeszłości mogliśmy sobie pozwolić na bez troskę i nie nadawać sensu pojęciu identyczności własności, gdyż użyteczność tego pojęcia nie zależy od utożsamiania lub rozróżniania cech. Jeśli tak, to czemu nie oczyścić atmosfery i nie stwierdzić po prostu, że cechy koekstensywne są identyczne? Jedyna trudność polega na tym, że sprzeciwiałoby się to potocznym intuicjom, jak widać na przykładzie serca i nerek. Pragnąc złagodzić szok, zmieniamy słowo: nie mówimy już o cechach, lecz o klasach.

To rozwiązanie stale spotyka się z niezrozumieniem. Sądzi się, że klasy różnią się od własności nie tylko ze względu na swą ekstensjonalność, tj. identyczność tego, co koekstensywne. Niektórzy charakteryzują tę ostatnią własność mówiąc, poprawnie lecz nieostrożnie, że klasa jest wyznaczona przez swoje elementy. Takie ujęcie jest nieostrożne, gdyż sugeruje, że elementy są w pewnym sensie przyczyną klasy, a przecież rzeczy nie są przyczyną swoich własności. Mówi się, że w idealnym przypadku można scharakteryzować klasę wyliczając jej elementy – metoda niewykonalna w przypadku prawie każdej interesującej klasy. Opisuje się klasy jako „zbiorowiska” lub „agregaty”, przyjmując nieuprawnioną metaforę wybierania i gromadzenia przedmiotów poprzez zmianę ich położenia przestrzennego. Są to szkodliwe metafory. W bardziej neutralnych terminach można by określić klasę jako *zbiorowość* przedmiotów – pomija się tu jednak klasy jednoelementowe i puste. Klasa w przydatnym znaczeniu tego słowa jest po prostu cechą w zwykłym sensie, jeśli odrzucimy rozróżnienia pomiędzy cechami koekstensywnymi.

Słowo „klasa” w tych logicznych kontekstach, ze szkodliwą metaforą lub bez niej, nie jest stare. Zdaniem Meilleta jego najdawniejszym znanym nam zastosowaniem jest łacińskie słowo *classis*, używane w związku z klasyfikacją poborowych. Sugeruje on, że to słowo zostało zapożyczone z etruskiego. Jeszcze w czasach rzymskich zaczęto za jego pomocą nazywać klasy społeczne, obecne marksistowskie rozumienie jest więc bardzo konserwatywne. W osiemnastym wieku pojawiło się w taksonomii, ale używano go tylko na pewnym poziomie klasyfikacji, i tak mamy klasę ssaków, klasę skorupiaków.

W późniejszych czasach sens słowa „klasa” stawał się coraz bogatszy, aż zaczął przypominać znaczenie takich wyrażen jak „cecha” albo „atrybut”, pomijając wymóg ekstensjonalności. W dziewiętnastym wieku było dla wszystkich oczywiste, że każdy warunek przynależności wyznacza klasę, tak jak wszystko, co się mówi o rzeczy, charakteryzuje jakąś własność. Właśnie w tym momencie klasy uzyskały status cech w ekstensjonalnym sensie, który, jak twierdzą, wciąż nie jest należycie doceniony.

A potem wybuchła bomba Rusella, przekreślając komuną, zgodnie z którym każdy warunek przynależności wyznacza klasę. Zob. PARADOKSY, a także

NIEPREDYKATYWNOŚĆ. To jednak nie przeczy temu, że klasy są własnościami w ekstensjonalnym sensie. Rozumowanie kryjące się za paradoksem Russella stosuje się równie dobrze do własności, jak i do klas, i przekreśla komunał, zgodnie z którym wszystko, co mówi się o rzeczy, charakteryzuje jakąś własność. Jeśli nakładamy na klasy pewne teoriomnogościowe ograniczenia, żeby uniknąć sprzeczności, to musielibyśmy nałożyć podobne ograniczenia na własności, gdybyśmy byli na tyle przewrotni, by uznawać ich istnienie obok albo zamiast istnienia klas.

Dla zwykłych celów zadowolamy się językiem potocznym, zaś słowo „cecha” stanowi jego integralną część. Pojęcie cechy albo jakiś jego odpowiednik jest nam jednak również potrzebne dla celów technicznych w teorii naukowej, a zwłaszcza w matematyce. W tych przypadkach rolę cech mogą przejąć klasy, gdyż nie zależy nam tu na odróżnianiu od siebie własności koekstensywnych. Pod hasłem **DEFINICJA**, gdzie mowa o definicji liczby, znajdziecie jeden z wielu przykładów wykorzystania klas w matematyce. W nauce klasom mówimy *si*, cechom *no*.

Klasy a zbiory

Człowiek jest zwierzęciem praktycznym, a nawet skąpym, i dlatego nie lubi, kiedy na ten sam przedmiot nalepia się dwie różne etykiety. Niektórzy mówią „ziemniak”, a inni „pyra”, nikt jednak nie będzie używał obu tych wyrażений – taka jest naturalna kolej rzeczy. Jeśli zauważamy, że dwa terminy oznaczają ten sam przedmiot, to mamy ochotę doszukać się mimo wszystko jakiejś różnicy. W języku angielskim dwa wyrażenia: *ape* i *monkey* nazywają zwierzęta, które po niemiecku zwą się *Affen* a po polsku „małpy”. Użytkownicy języka angielskiego najwyraźniej znaleźli rozwiązanie: podzielili małpy według wielkości. Wyczuwam w tym próbę trzymania się maksymy „dwa słowa, dwa sensy”.

Tak się właśnie stało z terminami „zbiór” i „klasa”. Były one używane zamiennie, co skłaniało do wprowadzenia jakiejś różnicy znaczeniowej. Niektórzy logicy i matematycy mówią więc zarówno o klasach, jak i o zbiorach. Ich zdaniem klasy przypominają cechy (biję im za to brawo, zob. powyżej), a zbiory są w pewnym sensie bardziej konkretne, chociaż wciąż abstrakcyjne. Ten kontrast jest na razie tylko mglistą metaforą, jednakże matematycy nadają jej uchwytny kształt tworząc silną lub niepredykatywną teorię zbiorów (zob. **NIEPREDYKATYWNOŚĆ**). Następnie dodają do tego klasy, których elementami według tej koncepcji mogą być zbiory, ale nie inne klasy. Jak wiemy, niektóre warunki przynależności nie wyznaczają żadnego zbioru (w przeciwnym razie grozi nam **PARADOKS**), mogą jednak wyznaczać klasy. Na przykład nie może istnieć zbiór wszystkich zbiorów, które nie są własnymi elementami, ale może istnieć klasa wszystkich takich zbiorów.

John von Neumann był pierwszym matematykiem, który w 1925 roku zaproponował taki schemat. Upraszcza to dowodzenie w teorii mnogości,

a także wzmacnia system bez narażania go na ryzyko paradoksu. Realizacja tego pomysłu wiąże się jednak ze zbędną duplikacją i wymyślnymi rozróżnieniami. Oprócz dodatkowych klas, takich jak ta, którą przed chwilą opisałem, w teorii tej przyjmuje się, że dla każdego zbioru istnieje klasa z nim koekstensywna.

Możemy się pozbyć niepotrzebnej duplikacji i wciąż cieszyć się zaletami tego rozwiązania. Możemy po prostu utożsamiać zbiory z koekstensywnymi klasami, tak jak to proponowałem od 1940 roku. Wciąż zachowujemy podwójną terminologię „zbiór” i „klasa”, gdyż zbiory stają się klasami specjalnego rodzaju. Klasa jest zbiorem, gdy jest elementem pewnej klasy. Klasa jest klasą ostateczną, gdy nie jest elementem żadnej klasy.

Niektórzy ludzie sądzili, że w świetle rozróżnienia „zbiór – klasa” paradoks Russella i podobne paradoksy są zwykłymi błędami, a nie antynomiami. Komunał, zgodnie z którym każdy warunek przynależności wyznacza klasę, nie prowadzi, jak mówili, do żadnego paradoksu, gdyż klasy nie są przedmiotami, które mogłyby być elementem czegoś. Takimi przedmiotami są zbiory, ale przecież nikt nigdy nie twierdził, że wyznacza je każdy warunek przynależności. Ludzie ci twierdzą, że zbiory pojawiły się dopiero sto lat temu w pracach Cantora i od początku było jasne, że mają być one wyznaczone zgodnie z zasadami wyraźnie sformułowanymi w teorii mnogości Zermelo.

Jest to niebezpieczne rozumowanie. Klasy i zbiory miały to samo zastosowanie, tę samą rację bytu i tę samą genezę. Każda niejasność związana z jednym pojęciem wiązała się również i z drugim. Przypnijmy, że Cantorowskie „zbiory” (*Mengen*) były zbiorami punktów, ale to nie jest żadną wskazówką. W dojrzałej Cantorowskiej teorii mnogości można odnaleźć pewne ograniczenia nałożone na dziedzinę zbiorów, ale chodziło najprawdopodobniej o to, że Cantor bystro wyczuł groźbę antynomii. Istnieją fragmenty prac Cantora, które ktoś mógłby uznać za zapowiedź rozróżnienia tak owocnie wprowadzonego później przez von Neumanna, ale to również można złożyć na karb tych przeczuć. Mit, zgodnie z którym zbiory zostały wymyślone niezależnie od klas, a później Russell i inni pomieszały ze sobą te dwa pojęcia, jest jeszcze jednym przejawem skłonności, by za różnicą werbalną dostrzegać różnicę rzeczową.

Korzyści rozwiązania von Neumanna są niezależne od tego mitu. Jak widzieliśmy, możemy akceptować to rozwiązanie i uznawać jedynie klasy; innymi słowy – tylko zbiory i ostateczne klasy. Wykorzystujemy podwójną terminologię, ale nie tłumaczymy jej żadną wymyślną opowieścią.

Komunikacja

Przekazujemy innym ludziom nie tylko choroby, ale również idee. Idea mieszcząca się w jednym umyśle zostaje skopiowana w drugim. Kiedy staramy się „zglebić mroki umysłu drugiego człowieka”, jak to ujmie Santayana, nie zawsze możemy stwierdzić, jak wierna jest to kopia. Sama idea IDEI jest tak

niejasna, że trudno jest nawet powiedzieć, jaką formę, treść i granice mają nasze własne idee. Niech zgadnie kto potrafi.

Możemy rzucić pewne światło na naturę i granice komunikacji, jeśli odrzucimy mgliste pojęcie idei i zajmiemy się dotykalną, widzialną i słyszalną rzeczywistością. Proste zdania na takie konkretne tematy są świetnym narzędziem komunikacji, a zwłaszcza wtedy, gdy zarówno my, jak i nasi rozmówcy, od czasu do czasu stykamy się z przedmiotami, o których mowa. Wszyscy nauczyliśmy się tego rodzaju słów i zwrotów w obecności przedmiotów, do których się one odnoszą, albo w obecności rzeczy, które są do nich podobne. Znaczenie nadawane przez nas tym słowom i zwrotom było wciąż uaktualniane i sprawdzane w procesie komunikacji. Nic więc dziwnego, że skuteczna komunikacja na tym poziomie jest możliwa, i to bez względu na to, czym miałyby być idee.

Podstawa komunikacji jest mniej prosta i oczywista, kiedy staram się zakomunikować mojemu rozmówcy, że ktoś ukradł stary miecz, do którego byłem przywiązany, gdyż ojczym mojej matki posługiwał się nim w bitwie pod Gettysburgiem. Zwiedzanie muzeów nudzi mojego rozmówcę, toteż nigdy nie widział on miecza. Nic mu o tym nie wiadomo, żeby kiedykolwiek się natknął na czyjegós ojczyma. Żaden z nas nie był nigdy świadkiem kradzieży ani bitwy. Żaden z nas nie odwiedził Gettysburga. Jeśli zaś chodzi o uczucie przywiązania, to nie bardzo nawet wiadomo, od czego zacząć. A jednak fakt komunikacji jest tu niewątpliwy.

Potrafimy to wyjaśnić. Mój rozmówca słyszał i widział słowo „miecz” w różnych kontekstach, znane mu są również rozmaite wyjaśnienia tego terminu. Jeśli o mnie chodzi, to słyszałem i widziałem to słowo w *innych* kontekstach i natknąłem się na inne wyjaśnienia, które niekiedy wiązały się ze wskazaniem odpowiedniego przedmiotu. Te różnorodne sposoby poznania znaczenia danego terminu tworzą w skali społecznej spójną sieć. Ta spójność nie jest niczym przypadkowym, gdyż wspomniana sieć ma właściwości samonaprawcze. Kiedy komunikacja się nie udaje ze względu na to, że ktoś w niewłaściwy sposób posługuje się jakimś terminem, przywołuje się go do porządku i ustawia z powrotem w szeregu. Podobnie można wyjaśnić jednolite rozumienie wszystkich słów w naszym przykładzie, nawet słowa „przywiązanie”.

Można by podawać jeszcze trudniejsze przykłady, np. wypowiedź Hegla „Prawda jest sprzymierzeńcem rzeczywistości przeciwko świadomości” albo moją uwagę „Logika ściga prawdę na drzewie gramatyki”. Jestem przekonany, że rozumiem, co chciał powiedzieć Hegel i że istnieją filozofowie logiki, którzy rozumieją, co ja miałem na myśli. Jednakże samo tylko potwierdzenie, choćby nawet szczere – „Złapałem, o co ci chodzi” albo „Słyszę cię, bez odbioru.” – nie jest konkluzywną oznaką udanej komunikacji. Uczeń na lekcji łaciny dostaje niski stopień, kiedy mówi „Och, wiem, co to znaczy, ale nie potrafię tego wyrazić”. Nieudana komunikacja jest wdzięcznym tematem komedii – postaci sceniczne wdają się w długie, skomplikowane dialogi, co wywołuje zabawny

efekt, gdyż publiczność, w przeciwieństwie do nich, świetnie zdaje sobie sprawę z nieporozumienia.

Istnieją tu pewne obiektywne wskaźniki. Uznajemy, że udało się nam zakomunikować to, o co nam chodziło, gdy nasz rozmówca reaguje we właściwy sposób, np. naciska szybko na hamulec, patrzy na właściwy obszar nocnego nieba albo kontynuuje rozmowę w tak dogłębny sposób, że nie może być mowy o nieporozumieniu. Jesteśmy również w stanie zauważyć nieudaną komunikację. Porażka umożliwia nam lepsze opanowanie języka; po odrzuceniu błędnych interpretacji zwierają się szeregi naszej społeczności.

Im bardziej oddalamy się od zwykłej rozmowy o znanych, konkretnych sprawach, tym rzadziej możemy korzystać z obiektywnych wskaźników, a nawet kiedy z nich korzystamy, często nie przynosi to konkluzyjnych rezultatów. Przemawiamy długo do cierpliwego słuchacza i tylko od czasu do czasu uzyskujemy nie wiążące potwierdzenie tego, że przekazaliśmy nasze idee (wybaczcie to wyrażenie) albo wywołaliśmy takie, których nie chcieliśmy wywołać. Brak wiadomości to dobra wiadomość. Interpretujemy umysł słuchacza zgodnie z zasadą życzliwości, jak to określił Neil Wilson. Wyobrażamy sobie, że bardzo dobrze nas zrozumiano tylko dlatego, że nic nie świadczy przeciwko temu. Na zewnętrznych rubieżach cud komunikacji przypomina trochę cud przeistoczenia: jakiego przeistoczenia?

Konieczność

Mówi się, że „musi” implikuje „jest”. Cokolwiek jest faktem koniecznym, jest faktem. Kiedy jednak mówię „On musi być koło domu”, to zgaduję i dopuszczam możliwość pomyłki. W przeciwnym razie powiedziałbym po prostu „On jest koło domu”. Najwyraźniej coś tu musi ustąpić, a tym czymś jest właśnie słowo „musi”. Niewykluczone, że w połowie przypadków słowo to wyraża konieczność, ale w pozostałych kontekstach konotuje ono brak konieczności, albo przynajmniej brak pewności. Słowo „musi” jest prawem dla siebie samego, jeśli rzeczywiście chcemy tu mówić o prawie. Zob. FLEKSJA i NEGACJA.

Pozostawmy więc słowu „musi” jego gierki i zabawy i zajmijmy się samą koniecznością. Nie jest to łatwe zadanie. Nowocześni filozofowie biorą przykład z Leibniza i wyjaśniają konieczność jako prawdziwość we wszystkich możliwych światach. Jeśli jednak wyjaśnianie konieczności w terminach możliwości w ogóle coś nam daje, to ten sam zysk można osiągnąć w bardziej bezpośredni sposób: zdanie jest koniecznie prawdziwe, gdy nie jest możliwe, że jest fałszywe. „Konieczne” znaczy „niemożliwe, że nie”. Równie dobrze możemy wyjaśnić możliwość w terminach konieczności: „możliwe” znaczy „niekonieczne nie”. Albo rozumiemy oba te terminy, albo nie rozumiemy żadnego. Są one tym bardziej użyteczne, że jeden z nich można wyjaśnić dzięki drugiemu, musimy jednak poszukać jakiejś pomocy z zewnątrz.

Dwieście lat temu David Hume stracił nadzieję, że uda nam się odróżnić fakty konieczne od tych, które po prostu zachodzą. Twierdzi się zwykle, że prawdy matematyki oraz prawa przyrody wraz z ich logicznymi konsekwencjami są konieczne. W ten sposób przesuwamy tylko problem. Które prawdy o przyrodzie są konieczne, a które nie? Cóż, mówi się, że powinny one być ogólne. To jednak nie pomaga; najbanalniej w świecie typowe zdania szczegółowe są równoważne ze zdaniem ogólnym. Zdanie „Garfield urodził się w Orange” jest równoważne ogólnemu zdaniu „Każdy urodził się w Orange lub jest różny od Garfielda”. Proponuje się więc kolejny wymóg, zgodnie z którym prawo przyrody nie powinno wyszczególniać żadnego konkretnego przedmiotu, takiego jak Orange albo Garfield. Kłopot polega jednak na tym, że taki wymóg dyskwalifikuje prawa geologii, w których mowa o naszej planecie, a także prawa dotyczące słońca i systemu słonecznego. Pozostałyby nam tylko najogólniejsze prawa fizyki i niewiele mielibyśmy okazji do używania przysłowka „koniecznie”.

Sądzę, że Hume miał słuszość dyskredytując metafizyczną konieczność. Prawa przyrody różnią się od innych prawd o przyrodzie jedynie ze względu na to, jak do nich dochodzimy. Myślę, że prawdziwe zdanie ogólne jest prawem, gdy dochodzimy do niego dzięki metodzie indukcyjnej lub hipotetyczno-dedukcyjnej (zob. PREDYKCJA), a nie dzięki takim sztuczkom, jak w przykładzie z Garfieldem. *Sub specie aeternitatis* nie ma konieczności ani przypadku; wszystkie prawdy mają taki sam charakter.

Przysłówek „koniecznie” jest używany w rozmowach potocznych o wiele częściej niż wymagałyby tego prawa przyrody albo metafizyczna konieczność. W tym potocznym użyciu uwidacznia się ludzki element, który przypisałem prawu; tak zwana konieczność pojawia się i znika z okazji na okazję. W trakcie rozmowy będziemy skłonni zastosować ten przysłówek do zdania, które wynika z czegoś, na co zgodziliśmy się wraz z naszym rozmówcą, w przeciwieństwie do kwestii, które wciąż pozostają dyskusyjne. Jeśli wyjaśniamy na piśmie jakąś kwestię, to zastosujemy go do zdania, wynikającego z czegoś, co już powiedzieliśmy, w przeciwieństwie do hipotezy, która dopiero musi zostać dowiedziona.

Warto tu wspomnieć o jeszcze jednej sprawie. Mówiłem wcześniej o rzekomej konieczności prawd matematyki i praw przyrody, a następnie wykazywałem, że powoływanie się na prawa przyrody jest w tym kontekście chybione. Matematyczną konieczność trzeba jednak wyjaśnić w zupełnie inny sposób. Chciałbym zatrzymać się na chwilę i omówić stanowisko w teorii świadectw zwane holizmem. Zgodnie z tezą holizmu, sformułowaną przez Pierre Duhema osiemdziesiąt lat temu, obserwacyjne konsekwencje, za pomocą których sprawdzamy hipotezę naukową, nie wynikają na ogół z samej tylko hipotezy; są one konsekwencjami zbioru zdań, do którego należy ta hipoteza. Ujmując sprawę w podobnych terminach jak w rozdziale PREDYKCJE: hipoteza nie implikuje sama z siebie obserwacyjnego zdania kategorycznego.

Ma to następujący związek z matematyczną koniecznością. Do zbioru zdań,

potrzebnych do zakotwiczenia naszej implikacji, będą należały nie tylko zdania tej czy innej nauki, np. fizyki, ale również formuły matematyczne oraz rozmaite zdroworozsądkowe prawdy, które są dla nas oczywiste. Jeśli przewidziana obserwacja nie zajdzie, możemy w zasadzie wyjaśnić to niepowodzenie i odrzucić którekolwiek z tych zdań. Będziemy się starali zrobić to tak, by optymalizować przyszłe predykcje i właśnie dlatego w punkcie wyjścia wskazujemy na tę a nie inną hipotezę. Niezbyt dobrze wiemy, jakie względy rządzą tu naszym wyborem, ale jedną z oczywistych maksym jest maksyma minimalnej szkody: z dwóch konkurencyjnych rozwiązań lepsze jest to, które wymaga mniejszych zmian ogólnosystemowych. Ta maksyma stoi na straży prawd matematyki, gdyż każda zmiana w tej dziedzinie będzie miała rozległe reperkusje w całej nauce.

Jest to całkiem niezłe wyjaśnienie aury konieczności otaczającej prawdy matematyki. Matematyka ma tę samą treść empiryczną co reszta nauki, gdyż bez niej nie moglibyśmy wyprowadzać obserwacyjnych konsekwencji, a swą aurę konieczności zawdzięcza temu, że jesteśmy ostrożni i nie chcemy zbyt mocno kołysać łodzią.

Konstruktywizm

Sens tego terminu w matematyce nie jest do końca określony, można go jednak zdefiniować jako praktykę, projekt lub plan uprawiania matematyki ze związanymi rękami. Istnieli heterodoksyjni matematycy o nieco nominalistycznych skłonnościach (zob. UNIWERSALIA), którzy chętnie rezygnowali ze swobody. Byli również inni, którzy z radością podążaliby po każdej ścieżce prowadzącej do prawdy, uznawali jednak, że rozróżnienie między twierdzeniami posiadającymi konstruktywny dowód, a całą resztą twierdzeń jest interesujące z metodologicznego i filozoficznego punktu widzenia i lepiej jest szukać najpierw dowodu konstruktywnego, zanim zwrócimy się ku innym metodom.

Jeśli ktoś chciałby usłyszeć przykład teorii konstruktywistycznej, to może nim być predykatywna teoria zbiorów; zob. NIEPREDYKATYWNOŚĆ. Predykatywna teoria zbiorów jest zbyt słaba, żeby dowieść w niej, że istnieją niespecyfikowalne klasy i liczby rzeczywiste. W samej rzeczy, według jednej z interpretacji tezę konstruktywizmu jest to, że każdy przedmiot abstrakcyjny jest specyfikowalny.

Istnieje pewien rodzaj ZMIENNYCH i kwantyfikatorów, z których można korzystać, kiedy dziedzina, którą przebiegają zmienne, obejmuje tylko specyfikowalne przedmioty. (LICZBY NATURALNE stanowią przykład takiej dziedziny; każdej z nich jest przyporządkowana cyfra arabska). Przy interpretacji klasycznej lub *przedmiotowej*, kwantyfikator ogólny „ $\forall x$ ” RACHUNKU PREDYKATÓW znaczy „każde x jest takie, że” i tworzy zdanie prawdziwe zawsze i tylko wtedy, gdy formuła, przed którą stoi, jest spełniona przez każdy przedmiot x w danej dziedzinie. W innej wersji, zwanej substytucyjną, wymagane jest raczej to, żeby

formuła, przed którą stoi kwantyfikator, okazywała się prawdziwa przy każdym gramatycznie dopuszczalnym podstawieniu za „ x ”.

Kiedy naszą dziedziną jest zbiór liczb naturalnych, obie wersje są zbieżne, jeśli jednak niektóre przedmioty nie dają się wyszczególnić za pomocą żadnego terminu jednostkowego w naszym języku, to pojawiają się rozbieżności między obiema wersjami. Formuła przy kwantyfikatorze może być spełniona przez wszystkie specyfikowalne przedmioty ale nie przez niektóre inne. W tym wypadku zdanie jest prawdziwe przy interpretacji substytucyjnej, a fałszywe przy interpretacji przedmiotowej.

Podobnie ma się sprawa z kwantyfikacją egzystencjalną. Przy interpretacji przedmiotowej „ $\exists x$ ” znaczy „pewne x jest takie, że” i tworzy zdanie prawdziwe z formułą, która jest spełniona przez pewien przedmiot należący do dziedziny, przebieganej przez zmienne. Przy odczytaniu substytucyjnym wymagane jest natomiast to, by formuła była prawdziwa przy pewnym podstawieniu w miejsce „ x ”. Obie interpretacje pozostają w konflikcie, kiedy formuła jest spełniona przez pewien niespecyfikowalny przedmiot i nie jest spełniona przez żadne specyfikowalne przedmioty.

Kwantyfikacja substytucyjna jest nierealistyczna dla konkretnych przedmiotów. Czy można wyszczególnić każdy konkretny przedmiot z osobna – każdego przysłego i przeszłego ptaka i pszczołę, każdy atom i elektron? W zasadzie tak; moglibyśmy nałożyć na całą czasoprzestrzeń system liczbowych koordynat, wykorzystując po prostu liczby rzeczywiste. Jednakże zwykła kwantyfikacja przedmiotowa jest tu o wiele bardziej naturalna.

Dla predykatywnej teorii zbiorów kwantyfikacja substytucyjna jest natomiast wykonalna i atrakcyjna. Jest ona atrakcyjna, gdyż odnosimy wrażenie, że przedmioty abstrakcyjne, w przeciwieństwie do konkretów, zostały wyhodowane na naszym języku. Taką właśnie tezę wygłosił Charles Parsons. Nie chodzi tu o to, że substytucyjna kwantyfikacja po przedmiotach abstrakcyjnych usuwa je z naszego ontologicznego inwentarza. Wyrażenia językowe są przecież przedmiotami abstrakcyjnymi (zob. TYP A EGZEMPLARZ), ale nie w takim stopniu jak przedmioty, o których mówi wyższa teoria zbiorów. Można uznać, że kwantyfikacja substytucyjna przyznaje wartościom swoich zmiennych związanych pewien rodzaj rozrzedzonego istnienia, innego niż proste istnienie przypisywane konkretnym przez kwantyfikację przedmiotową. Jest to kompromis zawarty z wojującym nominalizmem.

To atrakcyjne stanowisko łączy kwantyfikację przedmiotową po konkretach z kwantyfikacją substytucyjną po abstraktach, które są klasami, jak w predykatywnej teorii mnogości. Parsons wykazał, że oba typy kwantyfikacji mogą istnieć obok siebie bez żadnych konfliktów, przy wykorzystaniu różnych rodzajów zmiennych.

Kłopot z tą propozycją, jak zauważa sam Parsons, polega na tym, że predykatywna teoria mnogości nie może spełniać roli klasycznej matematyki

liczb rzeczywistych; zob. znów NIEPREDYKATYWNOŚĆ. Russell odrzucił predykatywną teorię mnogości i wybrał pełnowartościową koncepcję, gdyż w ramach tej pierwszej nie jesteśmy w stanie zbudować matematyki potrzebnej nam w naukach przyrodniczych. W koncepcji pełnowartościowej przyjmuje się jednak istnienie niespecyfikowalnych liczb rzeczywistych oraz innych niespecyfikowalnych klas i nie można jej sformułować przy użyciu kwantyfikacji substytucyjnej.

Od czasu do czasu pojawiają się jednak entuzjaści, wyobrażający sobie konstruktywistyczną matematykę wystarczającą dla wszelkich potrzeb nauki. Dwa pokolenia temu jednym z nich był L.E.J. Brouwer, ale jego podejście wiązało się z nieatrakcyjną zmianą standardowej logiki. Hermann Weyl pracował nad tymi zagadnieniami w ramach zwykłej logiki, a później robili to Paul Lorenzen, Erret Bishop, Hao Wang i Sol Feferman. Autorzy ci starają się wykazać, że dzięki pewnym pomysłowym określonym konstrukcjom wystarczy nam sama predykatywna teoria mnogości. Oczywiście nie przedstawimy niezbitego dowodu adekwatności takiej teorii, dopóki nie odpowiemy na pytanie, jakiej właściwie matematyki potrzebujemy w naukach przyrodniczych. Możemy jednak mieć nadzieję nie na nominalizm w prawdziwym tego słowa znaczeniu, ale na atrakcyjne przybliżenie.

Przekład Cezary Cieśliński