

Maciej Witek

MYŚLENIE WOLNE, MYŚLENIE SZYBKIE I IMPLIKATURY KONWERSACYJNE. O ZALETACH PODEJŚCIA INTERDYSCYPLINARNEGO*

Wprowadzenie

Kognitywistyka (ang. *cognitive science*) jest interdyscyplinarną nauką o procesach poznawczych. Jej celem jest budowa empirycznie adekwatnych teorii wybranych zdolności mentalnych, na przykład percepcji, wyobraźni, pamięci, empatii, podejmowania decyzji oraz rozumienia mowy. Konstruując takie teorie, kognitywistyka korzysta z idei oraz dorobku badawczego wielu nauk: filozofii, logiki, informatyki, językoznawstwa, psychologii, neurobiologii, antropologii i innych dyscyplin badających mechanizmy i uwarunkowania aktywności poznawczej. Możemy nawet mówić o swoście kognitywistycznym modelu integracji wiedzy pochodzącej z wielu dziedzin (Klawiter, 2004). Model ten charakteryzuje się sporą płodnością heurystyczną: sprzyja powstawaniu idei, pytań i projektów badawczych, które byłyby niedostępne dla uczonych zamkniętych w wąskich, specjalistycznych aparaturach pojęciowych.

Celem niniejszego rozdziału jest przedstawienie przykładu zastosowania kognitywistycznego modelu integracji wiedzy. Chcemy mianowicie zwrócić uwagę na dwie nauki wchodzące w skład interdyscyplinarnego zaplecza kognitywistyki – psychologię oraz pragmatykę – między którymi dochodzi do swoistej wymiany świadczeń. Psychologia proponuje m.in. rozróżnienie szybkich i wolnych procesów poznawczych, które pozwala na oryginalne ujęcie mechanizmów podejmowania decyzji oraz formułowania osądów (Kahneman, 2012). Okazuje się, że posługując się tym rozróżnieniem można również konstruować i testować pragmatyczne teorie rozumienia wypowiedzi i innych działań komunikacyjnych. Tymczasem pragmatyczne koncepcje pośrednich aktów mowy – na

* Wcześniejsza wersja niniejszego tekstu została przedstawiona na konferencji *Między empirią a teorią. Abp Józef Życiński (1948-2011) in memoriam* (Lublin, 26-27.01.2016) oraz na *Zielonogórskim Konwersatorium Filozoficznym* (Zielona Góra, 28.04.2016). Dziękuję uczestnikom tych spotkań za cenne uwagi i sugestie zgłaszane w trakcie dyskusji.

przykład implikatur konwersacyjnych (Gazdar, 1979, Grice, 1980, Levinson, 2010, Witek, 2011) – umożliwiają pełniejsze zrozumienie mechanizmów decyzji, osądów i ocen formułowanych na podstawie informacji komunikowanych przez innych. Wydaje się, że wynikiem tej międzydyscyplinarnej współpracy jest obiecujący projekt badawczy, który doprowadzi do konstrukcji nowych teorii rozumienia mowy, formułowania osądów i podejmowania decyzji, czyli zdolności poznawczych leżących u podstaw zachowań badanych m.in. w naukach ekonomicznych (zob. Wawrzyniak i Wąsikowska, 2016).

Niniejszy rozdział składa się z trzech części. Część 1 przedstawia w dużym skrócie Amosa Tversky’ego i Daniela Kahnemana koncepcję dwóch typów procesów poznawczych: procesów autonomicznych, mimowolnych i szybkich oraz procesów refleksyjnych, kontrolowanych i wolnych. Omawiamy też ideę heurystyk – tj. uproszczonych schematów wnioskowania – rozumianych jako narzędzia myślenia szybkiego. Jednym z takich narzędzi jest heurystyka reprezentatywności, która – zdaniem Tversky’ego i Kahnemana – dochodzi do głosu podczas formułowania intuicji statystycznych. W części 2 przechodzimy do omówienia implikatur konwersacyjnych, czyli tych aspektów komunikowanego znaczenia, które wykraczają poza to, co nadawca komunikatu dosłownie mówi. W szczególności rozważamy pytanie o procesy poznawcze umożliwiające planowanie i odczytywanie implikatur. Zauważamy, że pytanie to można rozpatrywać w ramach Tversky’ego i Kahnemana modelu dwóch typów procesów poznawczych. Omawiamy też przypadki tzw. implikatur skalarnych, które zdają się funkcjonować głównie dzięki zaangażowaniu automatycznych i mimowolnych procesów szybkich. W części 3 wracamy do Tversky’ego i Kahnemana badania nad intuicjami statystycznymi i twierdzimy, że w komunikacji z uczestnikiem takiego eksperymentu pojawiają się nieprzewidziane przez Tversky’ego i Kahnemana implikatury skalarne, które mogą wpłynąć na wynik badania.

1. Myślenie szybkie i myślenie wolne

Według Tversky’ego i Kahnemana, w ludzkim umyśle zachodzą procesy dwojakiego rodzaju: procesy myślenia szybkiego, które są automatyczne, mimowolne i w większości wypadków asocjacyjne, oraz refleksyjne i kontrolowane procesy myślenia wolnego. Nawiązując do prac Keitha Stanovicha i Richarda Westa, Kahneman pisze o dwóch systemach – Systemie 1 oraz Systemie 2 – w których zachodzą, odpowiednio, procesy

szybkie i wolne (Kahneman, 2012, s. 31). Należy jednak pamiętać, że jest to tylko wygodna konwencja terminologiczna: mówiąc o Systemie 1 oraz Systemie 2 nie zakładamy, że ludzki umysł składa się z dwóch podzespołów czy modułów; chodzi raczej o symboliczne i przystępne wyrażenie idei, w myśl której różnorodne procesy poznawcze zachodzące w naszym umyśle można z grubsza podzielić na dwa typy.

Myślenie szybkie umożliwia nam wykonanie takich zadań, jak (a) uzupełnienie w myślach schematu „ $2 \times 2 = \dots$ ”, (b) przeczytanie ze zrozumieniem słowa „STOP” napisanego wyraźną, ciemną czcionką na jasnym tle, (c) rozpoznanie, że osoba, z którą rozmawiamy, jest zdenerwowana, oraz (d) odwrócenie głowy w kierunku, z którego dochodzi nagły i głośny pisk. Wspólną cechą ww. zadań jest to, że wykonujemy je szybko, bez większego wysiłku, mimowolnie i bez udziału uwagi. Tymczasem procesy myślenia wolnego dochodzą do głosu wtedy, gdy np. chcemy (e) obliczyć w pamięci iloczyn liczb 24 oraz 17, (f) ustalić, jak należy rozumieć instrukcję wypełniania zeznania podatkowego, (g) rozstrzygnąć, jaki prezent kupić bliskiej osobie oraz (h) zaplanować na mapie optymalną trasę wakacyjnej wycieczki, której celem jest odwiedzenie przyjaciół ze studiów. Ww. zadania wymagają czasu i energii, towarzyszą im poczucia świadomej kontroli oraz wysiłku związanego z utrzymywaniem uwagi. Ryzykując pewne uproszczenie możemy przyjąć, że zadania od (a) do (d) rozwiązuje intuicja, a zadania od (e) do (h) są domeną rozumu.

Kluczowa teza Tversky’ego i Kahnemana głosi, że podczas myślenia szybkiego nasz umysł poszukuje prostych rozwiązań, które w sposób spójny łączą aktualnie dostępne dane. Spójność, o której mowa, jest zazwyczaj skojarzeniowa (Kahneman, 2012, s. 89). Na przykład słysząc krótką opowieść „Rodzice Freda przyjechali spóźnieni. Niedługo mieli przyjechać ludzie z firmy cateringowej. Fred był wściekły”, bez zastanowienia przyjmujemy, że przyczyną zdenerwowania Freda było spóźnienie jego rodziców, a nie spodziewany przyjazd pracowników firmy cateringowej; rzecz w tym, że w naszej sieci skojarzeń istnieje silny związek między złością i brakiem punktualności, a nie między złością i oczekiwaniem na firmę cateringową (Kahneman, 2012, s. 103). Taka interpretacja nie jest osiągnięciem Systemu 2, do którego nawet nie dociera potencjalna dwuznaczność opowieści o Fredzie. Rzecz w tym, że drugie z możliwych rozwiązań – czyli idea „przyczyną złości Freda jest spodziewany przyjazd pracowników firmy cateringowej” – jest skutecznie stłumiona przez System 1. Ten ostatni podsuwa ideę „przyczyną złości Freda jest spóźnienie jego rodziców” jako jedyną i oczywistą interpretację rozważanej opowieści.

Związki asocjacyjne wykorzystywane przez System 1 umożliwiają szybkie i w większości wypadków trafne rozwiązania wielu bieżących problemów, z którymi umysł ludzki musi się ciągle mierzyć. System 2 ma inne zadanie: kontroluje i krytycznie ocenia rozwiązania oferowane przez System 1 lub samodzielnie podejmuje zadania, które – tak jak przypadki od (e) do (h) – wykraczają poza kompetencje Systemu 1. Możemy więc powiedzieć, że między rozważanymi typami procesów poznawczych zachodzi dość czytelny podział obowiązków. Zadaniem Systemu 1 jest tworzenie prostych i spójnych rozwiązań. Tymczasem krytyczny System 2 albo ocenia i modyfikuje wytwory Systemu 1, albo przejmuje inicjatywę tam, gdzie System 1 jest bezsilny. Warto jednak pamiętać, że procesy wolne, refleksyjne i kontrolowane wymagają sporych nakładów energii i uwagi, a zasoby tych ostatnich są ograniczone. Dlatego System 2 może zazwyczaj skupić się na jednym zadaniu należącym do zakresu jego kompetencji. Na przykład prowadząc samochód trudno byłoby jednocześnie szukać miejsca do zaparkowania przy zatłoczonej ulicy i mnożyć dwie liczby dwucyfrowe. Tego typu ograniczenia nazywa się obciążeniem poznawczym (Kahneman, 2012, s. 58). Drugim rodzajem ograniczenia charakterystycznego dla Systemu 2 jest wyczerpywanie ego. Ujawnia się ono w trakcie rozwiązywania serii złożonych zadań, z których każde wymaga przechowywania wyników cząstkowych w pamięci roboczej, jak to ma miejsce w wypadku mnożenia liczb dwucyfrowych; w takich sytuacjach coraz trudniej przychodzi nam utrzymywanie uwagi na kolejnych zadaniach, aż w końcu poddajemy się i odmawiamy udziału w grze.

Podział obowiązków i kompetencji między Systemem 1 i Systemem 2 jest dość czytelny. Czasem jednak dochodzi do jego zakłócenia. Mowa o sytuacjach, w których System 2 – a w zasadzie nasze świadome ego, które odczuwa siebie jako wykonawcę i kontrolera naszych rozumowań i decyzji – błędnie traktuje rozwiązania podsuwane przez System 1 jako własne. Dochodzi wtedy do błędów systemowych, które szczególnie wyraźnie pojawiają się w dwóch obszarach: podejmowania decyzji oraz formułowania intuicji statystycznych. Okazuje się, że zarówno podejmując decyzje (konsumenckie, inwestycyjne itp.), jak i oceniając prawdopodobieństwo możliwych, przyszłych zdarzeń, dość często nieświadomie i bezkrytycznie akceptujemy wskazania Systemu 1. Dochodzi wtedy do szczególnego złudzenia: wydaje nam się, że postępujemy zgodnie z zasadami rozumu, gdy tymczasem myślimy na skrótu za pomocą prostych heurystyk, czyli uproszczonych schematów wnioskowania. Powstałe w takiej sytuacji błędy są zaś nie tyle przypadkowe, ile systemowe:

nie chodzi o to, że przez nieuwagę mylimy się w racjonalnych kalkulacjach, ale o to, że takich kalkulacji w ogóle nie przeprowadzamy, a wskazania systemu Systemu 1 mylnie bierzemy za osiągnięcia Systemu 2.

Sytuacją, w której może pojawić się błąd systemowy, jest egzamin ustny zorganizowany dla licznej grupy studentów. Celem egzaminatora jest rzetelna ocena stanu wiedzy studenta: musi z uwagą śledzić jego wypowiedź, a cząstkowe oceny przechowywać w pamięci roboczej do momentu wystawienia oceny ogólnej. Zadanie to angażuje System 2, czyli wolne, refleksyjne i świadomie kontrolowane procesy myślowe, które wymagają sporych nakładów energii i uwagi. Zazwyczaj egzaminator jest w stanie sprostać temu wyzwaniu słuchając kilku pierwszych studentów. Z czasem jednak dochodzi do wyczerpania ego, gdyż zasoby uwagi niezbędne do mobilizacji Systemu 2 kurczą się. Mimo to egzaminator nadal wysłuchuje studentów i wystawia oceny. Dochodzi wtedy u niego do głosu *heurystyka afektu*: choć oficjalnie szuka odpowiedzi na trudne pytanie o stan wiedzy studenta, podświadomie zastępuje je łatwiejszym pytaniem o emocje, które budzi w nim osoba, z którą rozmawia. Jeśli wygląd, zachowanie, tembr głosu itp. studenta podobają się egzaminatorowi, ten pierwszy ma większe szanse na otrzymanie wysokiej noty. Krótko mówiąc, System 1 odpowiedzialny za rejestrowanie emocjonalnych reakcji na rozmówcę zaczyna odgrywać dominującą rolę w trakcie wykonywania zadania leżącego w zakresie kompetencji Systemu 2. Oczywiście System 1 działa od chwili rozpoczęcia egzaminu. Początkowo jednak jego wskazania są korygowane przez System 2. Dopiero po wyczerpaniu ego ten pierwszy bierze górę nad drugim, z czego znużony egzaminator niekoniecznie zdaje sobie sprawę.

Przedstawiony przykład ilustruje pewną wspólną własność heurystyk analizowanych przez Tversky'ego i Kahnemana. Polegają one mianowicie na nieświadomej zamianie pytania trudnego na łatwe: choć osoba wykonująca pewne trudne zadanie jest przekonana, że mierzy się z nim za pomocą rozumu, faktycznie myśli na skróty i rozwiązuje problem łatwiejszy. Dokładniej rzecz ujmując, osoba ta nie zauważa, że w miejsce przewidziane dla odpowiedzi na pytanie trudne wstawia odpowiedź na pytanie łatwe. Nasz umysł działa w zgodzie z zasadą minimalizacji wysiłku poznawczego, w myśl której należy rozwiązywać jak najwięcej problemów jak najmniejszym kosztem, a zasada ta promuje stosowanie heurystyk. Z taką sytuacją mamy do czynienia w wyżej przedstawionym wypadku znużonego egzaminatora, który ulega heurystyce afektu. Podobnie jest w wypadku sytuacji, w której mamy ocenić prawdopodobieństwo pewnych alternatywnych zdarzeń czy sytuacji, co ilustruje następujący

eksperyment przeprowadzony przez Tversky'ego i Kahnemana.

Osoby badane otrzymały informację o młodej kobiecie o imieniu Linda:

Linda ma trzydzieści jeden lat, jest niezamężna, wygadana i bardzo inteligentna. Skończyła filozofię. Na studiach bardzo się angażowała w zwalczanie dyskryminacji i promowanie sprawiedliwości społecznej, a także brała udział w demonstracjach przeciw energii jądrowej. (Kahneman, 2012, s. 211)

Przedstawiono im też sześć alternatywnych wersji kariery zawodowej Lindy:

- (1) Linda działa w ruchu feministycznym.
- (2) Linda jest pracownicą opieki społecznej i pomaga osobom z zaburzeniami psychicznymi.
- (3) Linda jest członkinią Wyborczej Ligi Kobiet.
- (4) Linda jest kasjerką bankową.
- (5) Linda jest agentką ubezpieczeniową.
- (6) Linda jest kasjerką bankową i działa w ruchu feministycznym.

Następnie poproszono uczestników eksperymentu o uporządkowanie ww. wariantów od najbardziej do najmniej prawdopodobnego. Głównym celem badania było sprawdzenie, jakie miejsce w rankingach sporządzonych przez uczestników zajmują wersje (4) i (6). Rzecz w tym, że zgodnie z zasadami rachunku prawdopodobieństwa wariant (6) nie może być bardziej prawdopodobny od wariantu (4). Wariant (6) ma bowiem postać koniunkcji „A & B”, a wariant (4) jest jednym z argumentów tej koniunkcji. Tymczasem jedna z zasad rachunku prawdopodobieństwa głosi, że

$$(7) \quad P(A \& B) \leq P(A)$$

czyli prawdopodobieństwo spełnienia koniunkcyjnego warunku postaci „A & B” jest mniejsze lub równe prawdopodobieństwu spełnienia warunku stanowiącego jeden z argumentów tej koniunkcji. Tymczasem ponad 85% badanych – którymi byli studenci kilku dobrych uniwersytetów – wskazało wariant (6) jako bardziej prawdopodobny od wariantu (4)

(Kahneman, 2012, s. 214). Krótko mówiąc, popełnili błąd koniunkcji polegający na pogwałceniu zasady (7).

Zdaniem Kahnemana i Tversky'ego, za obserwowany efekt odpowiedzialna jest *heurystyka reprezentatywności*: badani zamienili trudne pytanie o *prawdopodobieństwo* wariantów od (1) do (6) na łatwiejsze pytanie o *podobieństwo* podanej na wstępie charakterystyki Lindy do stereotypów, odpowiednio, działaczki ruchu feministycznego, pracownicy opieki społecznej, członkini Wyborczej Ligi Kobiet, kasjerki bankowej i agentki ubezpieczeniowej. Nic więc dziwnego, że wskazali wariant (6) jako bardziej prawdopodobny niż wariant (4), skoro podana charakterystyka Lindy lepiej pasuje do stereotypu feministki pracującej jako kasjerka w banku niż do ogólniejszego stereotypu kasjerki bankowej; innymi słowy, przedstawiona na początku eksperymentu charakterystyka Lindy i wariant (6) stanowią prostą i skojarzeniowo spójną całość, która przemawia do badanych jako bardziej wiarygodna od całości obejmującej konkurencyjny wariant (4). Co ciekawe, wiele z badanych osób dobrze znało rachunek prawdopodobieństwa. Jednak w wypadku wykonywanego zadania doszło do konfliktu między intuicją reprezentatywności oferowaną przez System 1 a logiką prawdopodobieństwa, którą może posłużyć się System 2; jak się okazuje, konflikt ten wygrał System 1, choć System 2 nie zdaje sobie sprawy ze swojej przegranej.

2. Implikatury konwersacyjne jako znaczenia wywnioskowane

Rozważmy następującą wymianę zdań:

(8) A: Choć ze mną do kina!

B: Muszę napisać zaległy artykuł.

Przebieg tego krótkiego dialogu można opisać następująco: A *proponuje* B wspólne wyjście do kina, a B tę *propozycję odrzuca*. Frazy czasownikowe wyróżnione pochyłą czcionką oznaczają dwa dopełniające się akty mowy – złożenie propozycji i jej odrzucenie – które składają się na spójny fragment pewnej rozmowy. Możemy powiedzieć, że A składa swoją propozycję bezpośrednio, tj. za pomocą zdania w trybie rozkazującym, które jest konwencjonalnym środkiem formułowania poleceń, prośb, rad, sugestii, propozycji i innych dyrektywnych aktów mowy (zob. Searle 1979). Tymczasem odmowa B jest pośrednia:

mówiąc, że musi napisać zaległy artykuł, B komunikuje treść (9), tym samym odrzucając propozycję A.

(9) B nie pójdzie z A do kina.

Możemy przyjąć, że całkowite znaczenie aktu mowy zawiera dwa aspekty: *znaczenie pierwotne*, które można utożsamić z tym, co nadawca mówi, oraz *znaczenie wtórne*, rozumiane jako to, co nadawca sugeruje. Znaczenie wtórne słów rozmówcy B z dialogu (8) obejmuje, po pierwsze, treść (9), która jest *implikaturą konwersacyjną* rozważanej wypowiedzi, oraz, po drugie, odrzucenie propozycji złożonej przez A, czyli *pośredni akt mowy* wykonany przez B. Rozważmy, za Paulem H. Gricem – od którego pochodzi zarówno sama idea, jak i pierwsza systematyczna teoria implikatur konwersacyjnych (zob. Grice, 1980, 1989; por Witek, 2011) – dystynktywne własności rozważanych aspektów znaczenia. Po pierwsze, w przeciwieństwie do tego, co powiedziane, implikatury konwersacyjne są zależne od kontekstu: gdyby B wypowiedział zdanie „Muszę napisać zaległy artykuł” w odpowiedzi na pytanie „Odpoczniesz w weekend?”, powiedziałby z grubsza to samo, co w kontekście dialogu (8), ale implikowałby inną treść stanowiącą pośrednią odpowiedź na wyjściowe pytanie (co więcej, byłaby to odpowiedź przecząca lub twierdząca, w zależności od obowiązujących w danym kontekście standardów weekendowego odpoczynku). Po drugie, implikatury konwersacyjne są odwoływalne lub podatne na uchylenie: rozmówca B z dialogu (8) może dodać „Zrobię to po powrocie z kina”, uchylając tym samym implikaturę pierwszej części swojej wypowiedzi. Po trzecie, implikatury konwersacyjne nie wchodzą w skład warunków prawdziwości, dzięki czemu możemy wprowadzać słuchaczy w błąd mówiąc im prawdę (dlatego mową pośrednią chętnie posługują się politycy, akwizytorzy i kaznodzieje). Aby tę możliwość zilustrować, rozważmy następujący dialog:

(10) A: Mam ochotę na pizzę.

B: Tuż za rogiem jest niezła restauracja.

Słyszac wypowiedź rozmówcy B, rozmówca A ma prawo przyjąć do wiadomości treść (11), czyli komunikowaną przez B implikaturę konwersacyjną:

(11) Tuż za rogiem jest niezła restauracja, w której można zjeść pizzę.

Gdyby się jednak okazało, że w restauracji za rogiem nie podaje się pizzy – jest to np. *sushi bar* – to byłby to powód do oskarżenia B nie tyle o kłamstwo, ile o wprowadzenie w błąd. Kłamiemy, gdy komunikujemy fałszywe treści na poziomie znaczenia pierwotnego; tymczasem komunikowanie takich treści na poziomie znaczenia wtórnego jest jedynie wprowadzaniem w błąd.

Najciekawsza z naszego punktu widzenia jest czwarta własność implikatur: o ile znaczenie pierwotne jest *odkodowane*, o tyle znaczenie wtórne jest *wywnioskowane*. Okazuje się jednak, że ideę tę można rozwinąć przynajmniej na dwa sposoby. Pierwszy z nich znajdujemy w pismach Grice’a oraz jego kontynuatorów: Johna R. Searle’a (1979), Kenta Bacha i Roberta M. Harnisha (Bach i Harnish, 1979; por Witek 2011) i innych. Badacze ci zakładają, że proces ustalający każdą implikaturę konwersacyjną ma charakter rozbudowanego, zasadniczo kontrolowanego rozumowania, korzystającego z wielu informacji i założeń pomocniczych, a przede wszystkim z założenia o racjonalności rozmówcy; innymi słowy, rozumienie wszystkich implikatur przypomina proces, który Kahneman umieściłby w Systemie 2. Takie rozwinięcie idei implikatur jako znaczeń wywnioskowanych nazwijmy *racjonalistycznym*. Istnieje jednak alternatywne, *koherencjonistyczne* ujęcie rozważanych procesów interpretacyjnych. Na jego gruncie nie wyklucza się, że rozumienie *wielu* implikatur wymaga przeprowadzenia złożonych i kontrolowanych wnioskowań; dodaje się jednak, że przynajmniej *niektóre* znaczenia wtórne ustala się za pomocą szybkich i automatycznych procesów, którymi kieruje zasada szukania prostych i spójnych rozwiązań, a więc za pomocą mechanizmów podobnych do operacji wykonywanych przez Kahnemanowski System 1. Z ujęciem, o którym mowa, mamy do czynienia w wypadku Nicolasa Ashera i Alex Lascarides teorii reprezentacji dyskursu segmentowanego (ang. *Segmented Discourse Representation Theory*, zob. Asher i Lascarides, 2001, 2003) oraz Geralda Gazdara teorii implikatur skalarnych (zob. Gazdar 1979).

Okazuje się, że Kahnemanowska dystynkcja na automatyczne procesy szybkie i refleksyjne procesy wolne pozwala na interesujące rozwinięcie dyskusji o naturze procesów odpowiedzialnych za rozumienie implikatur. Według ujęcia racjonalistycznego, interpretacja wszystkich implikatur konwersacyjnych wymaga przeprowadzenia złożonego wnioskowania, którym kieruje założenie o racjonalności i kooperatywności nadawcy; możemy więc przyjąć,

że wnioskowanie to, choć zachodzi zazwyczaj bardzo szybko i nie wymaga większych nakładów uwagi, stanowi przykład procesu zbliżonego pod względem swojej struktury do operacji zachodzących w Kahnemanowskim Systemie 2. Tymczasem zwolennicy ujęcia koherencjonistycznego przyjmują, że przynajmniej niektóre implikatury konwersacyjne funkcjonują dzięki temu, że rozmówcy posługują się prostymi schematami oraz heurystykami do budowy spójnej reprezentacji dialogu, w którym uczestniczą.

Zacznijmy od ujęcia racjonalistycznego. Zdaniem Grice'a, wnioskowaniem prowadzącym do ustalenia implikatury konwersacyjnej danej wypowiedzi kieruje założenie, że nadawca tej wypowiedzi jest nastawiony na współpracę i działa racjonalnie. Dokładniej rzecz ujmując, uczestnicy konwersacji oczekują wzajemnie od siebie, że będą działać w zgodzie z zasadą współpracy (ZW) oraz podpadającymi pod nią czterema maksymami konwersacyjnymi (Grice, 1980, 1989).

(ZW) Mów tak, aby twój wkład do wymiany zdań, w której uczestniczysz, był taki, jak tego się aktualnie oczekuje.

Maksyma ilości: Podawaj tyle informacji, ile trzeba.

Maksyma jakości: Mów prawdę i wygłaszaj poglądy, dla których masz uzasadnienie.

Maksyma stosunku: Mów na temat.

Maksyma sposobu: Mów w sposób zwięzły, jasny i uporządkowany.

W myśl ujęcia Grice'owskiego, rozmówca A z dialogu (8) zauważa, że rozmówca B zdaje się łamać maksymę stosunku: nie mówi nic o zaproponowanym temacie konwersacji, czyli wyjściu do kina; A zakłada jednak, że B jest racjonalny, tj. narusza ww. maksymę celowo, dla uzyskania pewnego efektu komunikacyjnego; przyjmuje też, że B działa w zgodzie z zasadą (ZW) oraz podpadającymi pod nią maksymami konwersacyjnymi; A przyjmuje więc, że maksyma stosunku, choć złamana na poziomie tego, co powiedziane, jest zachowana na poziomie tego, co implikowane konwersacyjnie; zakłada też, że do kontekstu dialogu, w którym uczestniczy, należy idea „konieczność pracy nad zaległym artykułem ogranicza możliwość wyjścia do kina”; ostatecznie dochodzi do wniosku, że implikaturą konwersacyjną spełniającą (i) założenie o racjonalności i kooperatywności B, a także (ii) zidentyfikowane ograniczenia natury kontekstowej, jest treść (9). Podobną postać przyjmie Grice'owska

rekonstrukcja rozumowania umożliwiającego rozpoznanie treści (11) jako implikatury konwersacyjnej wypowiedzi rozmówcy B z dialogu (10): motorem tego wnioskowania będzie założenie, że formułując tę wypowiedź B działa zgodnie z maksymą stosunku.

Przejdźmy do omówienia Ashera i Lascarides (2001, 2003) teorii reprezentacji dyskursu segmentowanego, którą można potraktować jako wariant ujęcia koherencjonistycznego. Według Ashera i Lascarides, interpretacją niemal każdej wypowiedzi kieruje oczekiwanie, że pojawia się ona jako element spójnego dyskursu. O dyskursie powiemy, że jest spójny, jeśli każda jego wypowiedź odnosi się retorycznie do innej wypowiedzi tego dyskursu, gdzie u podstaw takiego odniesienia leżą relacje retoryczne typu *pytanie-odpowiedź*, *propozycja-przyjęcie*, *propozycja-odrzućenie*, *prośba-pomoc* itp. (W rzeczywistości Asher i Lascarides proponują bogatszy repertuar bardziej złożonych relacji retorycznych. Dodajmy, że relacje te można przedstawić jako konwencjonalne wzorce organizacji dyskursu; zob. Witek 2014, 2015. Dla potrzeb niniejszego wywodu poprzestaniemy jednak na tym niewielkim zestawie dość uproszczonych relacji.) Innymi słowy, szukamy takiego ujęcia wypowiedzi, które przedstawiałaby ją jako kolejny krok w budowie spójnego dyskursu lub dialogu, przy czym chodzi w tym wypadku o spójność nie tyle skojarzeniową, ile retoryczną. Przyjmując ten punkt widzenia możemy podać alternatywne względem racjonalistycznego wyjaśnienia dialogów (8) i (10): wypowiedź B z dialogu (8) rozumiemy jako odrzucenia propozycji, ponieważ taka interpretacja stanowi element reprezentacji rozważanego dialogu jako spójnego dyskursu zbudowanego na podstawie relacji *prośba-odrzućenie*; podobnie interpretacja wypowiedzi rozmówcy B z dialogu (10) prowadzi do ustalenia treści (11), ponieważ takie jej przedstawienie stanowi element reprezentacji rozważanej wymiany zdań jako spójnego dyskursu opartego na schemacie *prośba-pomoc*.

Kolejną teorią powstającą w ramach ujęcia koherencjonistycznego jest Geralda Gazdara koncepcja implikatur skalarnych. Rozważmy następujące zdanie:

(12) Piotr ma dwie córki.

Słyszac wypowiedź zdania (12) mamy prawo uzupełnić nasz zasób przekonań o treść (13).

(13) Piotr ma dokładnie dwie córki.

Ta ostatnia jest implikaturą konwersacyjną, czyli wykracza poza to, co powiedziane w rozważanej wypowiedzi. Świadczy o tym choćby możliwość odwołania treści (13) przez wypowiedzenie zdania (14):

- (14) A nawet trzy: Annę i Karolinę, które chodzą do gimnazjum,
i Monikę, która skończyła studia i pracuje.

Zauważmy też, że w sytuacji, w której Piotr ma trzy córki, wypowiedź zdania (12) nie jest kłamstwem: jeśli prawdą jest, że Piotr ma trzy córki, to prawdą jest, że ma on dwie córki. W takiej sytuacji wypowiedź zdania (12) wprowadzałaby jedynie w błąd.

Rozważana treść (13) jest *implikaturą skalarną* wypowiedzi zdania (12). Charakterystyczną własnością implikatur tego typu jest to, że ich Grice'owska rekonstrukcja – a w zasadzie Grice'owska rekonstrukcja wnioskowania umożliwiającego ich odczytanie – zawiera założenie o przestrzeganiu przez nadawcę maksymy ilości. Możemy mianowicie oczekiwać, że wypowiadając zdanie (12), racjonalny i kooperatywny rozmówca podaje tyle informacji, ile trzeba; innymi słowy, przyjmujemy, że gdyby wiedział on, że Piotr ma więcej niż dwie córki, to by to powiedział; skoro jednak tego nie mówi, to znaczy, że Piotr ma dokładnie dwie córki.

Istnieje alternatywne ujęcie procesów poznawczych umożliwiających odczytywanie implikatur skalarnych. Główną rolę odgrywa w nim prosta idea Geralda Gazdara (1979), którą przedstawimy tu w uproszczonej postaci. Zacznijmy od spostrzeżenia, że zdanie (12) zawiera jako swój element liczebnik główny „dwie”. Przyjmijmy, że słowo to aktywuje w naszym umyśle nie tylko pojęcie odpowiedniej liczby, ale również – być może dzięki odpowiedniej dyspozycji asocjacyjnej – skalę (15):

- (15) <jedna, dwie, trzy, cztery, ... >

Skala (15) jest *leksykalna*. Składa się mianowicie ze słów, które można wstawiać w puste miejsce następującego schematu:

- (16) Piotr ma ... córki.

W wyniku tych podstawień otrzymujemy skalę sądów:

$$(17) \quad \langle p_1, p_2, p_3, p_4, \dots \rangle$$

gdzie sąd p_1 wyraża się zdaniem „Piotr ma jedną córkę”, sąd p_2 – zdaniem „Piotr ma dwie córki”, sąd p_3 – zdaniem „Piotr ma trzy córki” itd. Zauważmy, że każdy kolejny element tej skali jest logicznie silniejszy od poprzednich w tym sensie, że jeśli $k > j$, to z p_k wynika logicznie p_j , ale z p_j nie wynika logicznie p_k . Na przykład w każdym możliwym świecie, w którym prawdą jest, że Piotr ma trzy córki, prawdą jest, że Piotr ma dwie córki; nie jest jednak tak, że w każdym możliwym świecie, którym prawdą jest, że Piotr ma dwie córki, jest również prawdą, że Piotr ma trzy córki (stosujemy tu standardową eksplikację pojęcia wynikania logicznego za pomocą kategorii prawdziwości w możliwym świecie).

Kluczowa idea teorii Gazdara głosi, że implikaturę skalarną danej wypowiedzi ustala się dzięki zastosowaniu prostej zasady:

$$(18) \quad \text{To, co powiedziane, jest najsilniejszym logicznie z prawdziwych elementów skali.}$$

Nazwijmy zasadę (18) *heurystyką skalarną*. Zauważmy, że jej zastosowanie do interpretacji wypowiedzi (12) prowadzi automatycznie do wniosku (13): w rozważanym wypadku tym, co powiedziane, jest element p_2 skali (17); skoro jest to najsilniejszy logicznie z prawdziwych elementów tej skali, to w myśl zasady (18) każdy sąd p_i , gdzie $i > 2$, należy uznać za fałszywy.

Podobnie można przedstawić przebieg procesu odpowiedzialnego za ustalenie treści (20) jako implikatury konwersacyjnej wypowiedzi zdania (19):

$$(19) \quad \text{Niektórzy studenci zdali egzamin.}$$

$$(20) \quad \text{Nie jest prawdą, że wszyscy studenci zdali egzamin.}$$

Treść (20) faktycznie jest implikaturą konwersacyjną, gdyż można ją odwołać. Rozważmy wypowiedź dwuzdaniowego tekstu (21):

$$(21) \quad \text{a. Niektórzy studenci zdali egzamin. b. W rzeczywistości zdali wszyscy.}$$

Powiemy, że wypowiedź fragmentu (21a) generuje implikaturę (20), która następnie jest odwołana za pomocą wypowiedzi fragmentu (21b).

Rekonstrukcję mechanizmu odpowiedzialnego za ustalenie implikatury (20) zaczniemy od spostrzeżenia, że zdanie (19) zawiera kwantyfikator „niektórzy”. Przyjmijmy, że aktywuje ono w umyśle słuchacza – na mocy odpowiedniej dyspozycji skojarzeniowej – skalę leksykalną (22), która następnie pozwala zbudować skalę sądów (23), gdzie sąd p_1 wyraża się zdaniem „Niektórzy studenci zdali egzamin”, a sąd p_2 – zdaniem „Wszyscy studenci zdali egzamin”.

(22) < *niektórzy, wszyscy* >

(23) < p_1, p_2 >

Stosując heurystykę (18) do skali (23) dochodzimy automatycznie do wniosku, że wypowiadając zdanie (19) nadawca komunikuje treść (20).

Wspólną własnością rozważanych przez Gazdara implikatur skalarnych (13) i (20) jest to, że komunikuje się je za pomocą zdań zawierających słowa – liczebniki główne i kwantyfikatory – aktywujące odpowiednie skale leksykalne; te ostatnie pozwalają szybko skonstruować odpowiednie skale sądów, do których można zastosować heurystykę (18). Dlatego implikatury (13) i (20) są nie tylko skalarne, ale również *uogólnione* w sensie Grice’a (zob. Grice 1980; por Levinson 2010): pojawiają się niezależnie od specyfiki kontekstu. Zauważmy jednak, że istnieją jednak implikatury skalarne, które nie posiadają tej własności. Mowa o skalarnych implikaturach *uszczegółowionych*, w wypadku których skala sądów jest skonstruowana za pomocą informacji pochodzącej z kontekstu, a nie z aktywowanej skali leksykalnej. Rozważmy następującą wymianę zdań między dyrektorem szkoły A i uczniem B:

(24) A: Jasiu, paliłeś w szkolnej ubikacji?

B: Zenek palił!

W sposób naturalny możemy przyjąć, że wypowiadając swoją kwestię rozmówca B – czyli Jaś – komunikuje treść (25):

(25) Jaś nie palił w szkolnej ubikacji.

Treść (25) jest implikaturą: można ją odwołać, o czym świadczy finał następującego dialogu:

(26) A: Jasiu, paliłeś w szkolnej ubikacji?

B: Zenek palił! I ja też.

Implikatura (25) jest skalarna. Uzasadnienie tej tezy zacznijmy od przypomnienia skali (23)

(23) $\langle p_1, p_2 \rangle$

Tym razem przyjmijmy, że sąd p_1 wyraża się zdaniem „Zenek palił w szkolnej toalecie”, a sąd p_2 – koniunkcją postaci „Zenek palił w szkolnej toalecie i Jaś palił w szkolnej toalecie”. Skala ta jest wyznaczona nie tyle leksykalnie, ile pragmatycznie lub kontekstowo: jej drugi element jest koniunkcją tego, co rozmówca B z dialogu (24) mówi, oraz tego, co jest tematem tej rozmowy. W wyniku zastosowania do omawianej skali heurystyki (18) dochodzimy do wniosku, że koniunkcyjny warunek p_2 jest fałszywy; jeśli tak, to skoro lewy argument koniunkcji p_2 uznajemy za prawdziwy, za fałszywy należy uznać argument lewy, czyli sąd wyrażony zdaniem „Jaś palił w szkolnej toalecie”. W ten sposób dochodzimy do implikatury (25).

Rozważmy, dla porównania, wypowiedź następującego dwuzdaniowego tekstu:

(27) Słyszeliście nowiny? Karol był wczoraj w kinie z pewną kobietą!

Przyjmijmy, że słowa (27) padają z ust znanego plotkarza. Co więcej, uczestnicy rozmowy dobrze znają Karola jako mężczyznę wiodącego uporządkowane życie monogamisty. Słyszając wypowiedź (27) mogą więc dojść do następujących wniosków:

(28) Kobietą, z którą Karol był wczoraj w kinie, nie jest jego żona Mariola.

(29) Karol ma romans.

Obie ww. treści są implikaturami konwersacyjnymi, co łatwo wykazać przedstawiając sobie sytuacje ich możliwego odwołania. Na przykład z odwołaniem implikatury (28) mamy do czynienia w wypowiedzi następującego tekstu, którego część b. uchyla implikaturę generowaną przez część a.:

- (30) a. Słyszeliście nowiny? Karol był wczoraj w kinie z pewną kobietą!
b. To była jego żona Mariola.

Co więcej, implikatura (28) jest skalarna. Aby to wykazać, raz jeszcze przywołajmy skalę (23) i przyjmijmy, że tym razem sąd p_1 wyraża się zdaniem „Karol był wczoraj w kinie z pewną kobietą”, a sąd p_2 – zdaniem „Karol był wczoraj w kinie ze swoją żoną Mariolą”. Faktycznie, p_2 jest sądem logicznie silniejszym od p_1 : z tego, że Karol był w kinie z Mariolą wynika logicznie, że był w kinie z pewną kobietą, ale z tego, że był w kinie z pewną kobietą nie wynika logicznie, że był tam z Mariolą (dla prostoty przyjmijmy, że świat, w którym Mariola nie jest kobietą, nie jest możliwy). Z chwilą, gdy taka skala aktywuje się w umysłach uczestników plotkarskiego spotkania, heurystyka (18) prowadzi do wniosku (28). Dlaczego jednak przyjąć, że aktywowana skala ma właśnie taką postać? Istnieje pewne uzasadnienie tego faktu. Zauważmy mianowicie, że drugie zdanie tekstu (27), które wyraża sąd stanowiący pierwszy element rozważanej skali, zawiera tzw. deskrypcję nieokreśloną „pewna kobieta”; wystąpienie tej frazy sugeruje, że drugi element skali powinien wyrażać się podobnym zdaniem, w którym za „pewna kobieta” wstawiamy imię lub jednoznacznie identyfikujący opis pewnej kontekstowo istotnej osoby. W kontekście spotkania, na którym plotkuje się o Karolu, taką osobą jest Mariola.

Podsumujmy dotychczasowe rozważania. Okazuje się, że wnioskowania odpowiedzialne za ustalanie implikatur konwersacyjnych można modelować na dwa sposoby. Można mianowicie przedstawiać je jako złożone rozumowania o strukturze zbliżonej do procesów zachodzących w Kahnemanowskim Systemie 2. Można też przyjąć, że przynajmniej niektóre implikatury – np. treści pośrednio komunikowane w trakcie dialogu oraz implikatury skalarne – funkcjonują dzięki mechanizmom wykorzystującym schematy, dyspozycje asocjacyjne i heurystyki; są to mechanizmy, które z powodzeniem mogłyby wchodzić w skład Kahnemanowskiego Systemu 1.

3. Implikatury skalarne w eksperymencie Tversky'ego i Kahnemana

W podrozdziale 1 przedstawiliśmy główne idee Tversky'ego i Kahnemana koncepcji dwóch typów procesów mentalnych. W podrozdziale 2 wykazaliśmy, że idee te znajdują zastosowanie w dyskusji o mechanizmach poznawczych odpowiedzialnych za ustalanie implikatur konwersacyjnych. Wydaje się jednak, że relacja między psychologicznym modelem dwóch procesów poznawczych oraz pragmatyczną koncepcją implikatur może być symbiotyczna: z jednej strony, w rozważaniach z zakresu pragmatyki korzystamy z narzędzi pojęciowych psychologii myślenia; z drugiej strony, pragmatyczne teorie mowy pośredniej mogą okazać się cennym narzędziem w dyskusji na temat wyników eksperymentów psychologicznych.

Aby uzasadnić powyższą opinię, rozważmy raz jeszcze dwa warianty kariery zawodowej Lindy wykorzystane w Tversky'ego i Kahnemana badaniu nad intuicjami statystycznymi:

- (4) Linda jest kasjerką bankową.
- (6) Linda jest kasjerką bankową i działa w ruchu feministycznym.

Przypomnijmy, że ponad 85% badanych wskazuje koniunkcyjny wariant (6) jako bardziej prawdopodobny, niż wariant (4). Tym samym zdają się popełniać błąd koniunkcji, polegający na sformułowaniu oceny sprzecznej z zasadą (7):

$$(7) \quad P(A \& B) \leq P(A)$$

Według Tversky'ego i Kahnemana, za efekt ten odpowiedzialna jest heurystyka reprezentatywności: w miejsce przewidziane dla odpowiedzi na trudne pytanie o prawdopodobieństwo badani wstawili nieświadomie odpowiedź na łatwe pytanie o podobieństwo znanej sobie charakterystyki Lindy do stereotypów kasjerki bankowej oraz aktywistki ruchu feministycznego.

Istnieje jednak alternatywne wyjaśnienie efektu uzyskanego przez Tversky'ego i Kahnemana. Nie twierdzimy, że jest ono trafniejsze. Jest jednak na tyle wiarygodne –

tj. umotywowane teoretycznie – że warto je przedstawić. Jego empiryczną adekwatność należałoby zaś sprawdzić za pomocą dodatkowych, odpowiednio zaprojektowanych badań.

Wydaje się, że w eksperymencie Tversky’ego i Kahnemana dochodzą do głosu nieprzewidziane przez nich implikatury skalarne. Co więcej, implikatury te mogą odgrywać zasadniczą rolę w wywoływaniu efektu zdiagnozowanego jako błąd koniunkcji. Zauważmy mianowicie, że warianty (4) i (6) układają się w następującą skalę sądów:

(31) < (4), (6) >

Jej konstrukcja nie jest co prawda motywowana względami leksykalnymi. W kontekście omawianego eksperymentu łatwo jednak ją przeprowadzić: wystarczy połączyć dwa zdania wykorzystane w materiale badawczym. Rozważmy teraz sytuację, w której uczestnik eksperymentu rozważa wariant (4). W jego umyśle aktywizuje się skala (31). Stosując do niej heurystykę (18) dochodzi do wniosku (32):

(32) Linda nie działa w ruchu feministycznym.

który jest skalarną implikaturą wypowiedzi zdania (4). Rzecz w tym, że w świetle heurystyki (18) koniunkcyjny element skali (31) uznajemy za fałszywy. Wiemy też, że prawdziwy jest lewy argument rozważanej koniunkcji. Zatem za fałszywy należy uznać jej prawy argument. Krótko mówiąc, rozważając wariant przedstawiony za pomocą zdania (4), uczestnicy eksperymentu biorą pod uwagę jego pragmatycznie wzmocnioną postać (4’):

(4’) Linda jest kasjerką bankową i *nie działa w ruchu feministycznym*.

Treść (4’) jest mianowicie koniunkcją tego, co powiedziane za pomocą zdania (4), oraz tego, co wypowiedź zdania (4) w rozważanym kontekście implikuje dzięki heurystyce skalarnej.

Można więc sformułować następującą hipotezę: umieszczając warianty (4) i (6) na skali od wariantu najbardziej prawdopodobnego do najmniej prawdopodobnego, uczestnicy eksperymentu porównują wariant (6) nie tyle z tym, co dosłownie komunikuje się za pomocą zdania (4), ile z pełnym, pragmatycznie wzmocnionym znaczeniem (4’) rozważanego komunikatu. Dlatego ocena, że prawdopodobieństwo wariantu *zakomunikowanego* za pomocą

zdania (6) jest większe od wariantu *zakomunikowanego* za pomocą zdania (4), nie stanowi błędu koniunkcji. Nie chcemy powiedzieć, że ocena ta jest trafna. Nie znamy przecież odpowiednich wartości bazowych, czyli proporcji feministek do nie-feministek w grupie zawodowej kasjerek bankowych. Wniosek, który płynie z przedstawionej wyżej hipotezy, jest znacznie słabszy. Głosi on, że rozważana ocena nie zawiera błędu koniunkcji.

Przedstawiona hipoteza oferuje więc alternatywne wyjaśnienie efektu zaobserwowanego przez Tversky'ego i Kahnemana. Czy to znaczy, że można się nią posłużyć do obrony zakwestionowanej przez psychologię podejmowania i ekonomię behawioralną decyzji wizji człowieka racjonalnego i *homo economicus*? Wydaje się, że nie. Chodzi raczej o to, by zwrócić uwagę na ukrytą, choć istotną rolę mechanizmów komunikacyjnych w eksperymentach psychologicznych, a także o to, by do repertuaru heurystyk odpowiedzialnych za nasze osądy i decyzje dodać jeden element: heurystykę skalarną.

Podsumowanie

W niniejszym rozdziale przedstawiliśmy przykład wymiany idei między dwoma dziedzinami naukowymi: psychologią oraz pragmatyką. Naszym zdaniem przykład ten dobrze ilustruje zalety interdyscyplinarnej postawy badawczej. Przyjmując ją, zyskujemy szansę na wzbogacenie poszczególnych dyscyplin o nowe perspektywy teoretyczne, które pozwolą na pełniejsze ujęcie badanych zjawisk. Co więcej, integracja wiedzy pochodzącej z wielu dyscyplin prowadzi do nowych idei i projektów badawczych, których podjęcie byłoby niemożliwe w wąsko zdefiniowanych dyscyplinach szczegółowych.

Literatura

- Asher N., Lascarides, A. (2001), *Indirect speech acts*, Synthese, 128, s. 182–228.
- Asher N., Lascarides, A. (2003), *Logics of Conversation*, Cambridge University Press, Cambridge.
- Bach, K., Harnish, R.M. (1979), *Linguistic Communication and Speech Acts*, MIT Press, Cambridge, Mass.
- Gazdar G. (1979), *Pragmatics, Implicature, Presupposition, and Logical Form*, Academic Press, New York, San Francisco, London.

- Grice P.H. (1980), *Logika a konwersacja*, w: red. B. Stanosz, *Język w świetle nauki* (91–114), Czytelnik, Warszawa.
- Grice P.H. (1989), *Studies in the Ways of Words*, Oxford University Press, Oxford.
- Kahneman D. (2012), *Pułapki myślenia. O myśleniu szybkim i wolnym*, Media Rodzina, Poznań.
- Klawiter A. (2004), *Powab i moc wyjaśniająca kognitywistyki*, *Nauka*, 3, 101-120.
- Levinson, C.S. (2010), *Pragmatyka*, Wydawnictwo Naukowe PWN, Warszawa.
- Searle, J.R. (1979), *Expression and Meaning: Studies in the Theory of Speech Acts*, Cambridge University Press, Cambridge.
- Wawrzyniak A., Wąsikowska B. (2016), *The Study of Advertising Content with Application of EEG*, w: red. K. Nermend, M. Łatuszyńska, *Selected Issues of Experimental Economics* (333-353), Springer.
- Witek, M (2011), *Spór o podstawy teorii czynności mowy*, Wydawnictwo Naukowe Uniwersytetu Szczecińskiego, Szczecin.
- Witek, M. (2014), *Neoaustinowskie ujęcie interakcji illokucyjnej*, w: red. P. Stalmaszczyk, P. Cap, *Pragmatyka, retoryka argumentacja. Obrazy języka i dyskursu w naukach humanistycznych* (139-159), Universitas, Kraków 2014.
- Witek, M. (2015), *An Interactional Account of Illocutionary Practice*, *Language Sciences*, 47, s. 43-55.